

Lineare Algebra I, WS 2020/21

Ulrich Görtz

Version vom 18. April 2021.

Ulrich Görtz

Universität Duisburg-Essen

Fakultät für Mathematik

45117 Essen

`ulrich.goertz@uni-due.de`

Ich freue mich über Kommentare und Berichtigungen.

Ich bedanke mich für Bemerkungen/Korrekturen bei Fereshteh Fattahi, Lukas Fußangel, Michelle Ganske, Florian Leptien, Enes Türk.

© Ulrich Görtz, 2020–21.

Lizenz: [CC BY-NC-ND 4.0](#)¹. [Lesbare Kurzform](#)². Das bedeutet insbesondere: Sie dürfen die PDF-Datei (unverändert) ausdrucken und als Datei oder ausgedruckt weitergeben, wenn es nicht kommerziellen Zwecken dient.

Gesetzt in der Schrift [Vollkorn](#)³ von F. Althausen mit LuaLaTeX, TikZ und anderen T_EX-Paketen. Einige Abbildungen wurden mit [IPE](#)⁴ erstellt. Die HTML-Version wird mit [plasTeX](#)⁵ erzeugt.

¹<https://creativecommons.org/licenses/by-nc-nd/4.0/legalcode.de>

²<https://creativecommons.org/licenses/by-nc-nd/4.0/deed.de>

³<http://vollkorn-typeface.com/>

⁴<http://ipe.otfried.org/>

⁵<https://github.com/plastex/plastex>

Inhaltsverzeichnis

Kapitel 1. Einführung, Motivation	7
1.1. Zum Mathematikstudium und zum kommenden Semester	7
1.2. Dieses Skript	8
1.3. Wie denkt man über ein mathematisches Problem nach? *	10
1.4. Gute wissenschaftliche Praxis	12
1.5. Zu erwartender Arbeitsaufwand	12
Kapitel 2. Die Vorlesung »Lineare Algebra I«	15
2.1. Lernziele	15
2.2. Wozu Lineare Algebra?	16
2.3. Konkrete Fragen, die mit LA beantwortet werden können *	21
2.4. Alternativer Aufbau der Vorlesung *	28
2.5. Lineare Gleichungssysteme	28
Kapitel 3. Grundlagen	33
3.1. Worum geht es eigentlich?	33
3.2. Was ist ein Beweis? *	33
3.3. Beweise – einige Klassiker *	35
3.4. Falsche Beweise erkennen *	38
3.5. Offene Fragen *	41
3.6. Wie schreibe ich einen Beweis (oder: die Lösung einer Aufgabe) auf?	41
3.7. Grundbegriffe, Notation	46
3.8. Mengen	48
3.9. Teilmengen, Konstruktionen von Mengen	51
3.10. Kartesisches Produkt, Abbildungen	53
3.11. Injektive, surjektive und bijektive Abbildungen	57
3.12. Vollständige Induktion	59
3.13. Endliche Mengen	67
3.14. Relationen *	70
3.15. Mächtigkeit von Mengen *	75
Kapitel 4. Körper	79
4.1. Körper	79
4.2. Endliche Körper	87
4.3. Polynomfunktionen	93
Kapitel 5. Lineare Gleichungssysteme und Matrizen	97
5.1. Lineare Gleichungssysteme	97
5.2. Der Gauß-Algorithmus	102
5.3. Das Matrizenprodukt	113
Kapitel 6. Vektorräume	135
6.1. Vektorräume	135
6.2. Erzeugendensysteme	140
6.3. Lineare Unabhängigkeit	144

6.4. Basen, Basissätze	146
6.5. Wie berechne ich ...?	152
6.6. Produkt und direkte Summe von Vektorräumen	155
6.7. Ergänzungen *	157
Kapitel 7. Lineare Abbildungen	163
7.1. Lineare Abbildungen	163
7.2. Kern und Bild einer linearen Abbildung	170
7.3. Der Zusammenhang zwischen linearen Abbildungen und Matrizen	172
7.4. Der Rang einer Matrix	179
7.5. Der Dualraum eines Vektorraums	182
7.6. Wie berechne ich ...?	188
7.7. Ergänzungen *	190
Kapitel 8. Gruppen	201
8.1. Grundlagen	201
8.2. Die spezielle lineare Gruppe	212
8.3. Permutationen	214
8.4. Die Bruhat-Zerlegung *	219
8.5. Ergänzungen *	222
Kapitel 9. Die Determinante	237
9.1. Determinantenfunktionen	237
9.2. Die Determinante eines Endomorphismus	246
9.3. Die Cramersche Regel	249
9.4. Die Spur einer Matrix	252
9.5. Ergänzungen *	253
Kapitel 10. Eigenwerte	257
10.1. Eigenwerte, Eigenvektoren	257
10.2. Eigenräume	262
10.3. Ergänzungen *	265
Kapitel 11. Analytische Geometrie *	271
11.1. Affine Räume	274
11.2. Das Standard-Skalarprodukt auf $\mathbb{R}^n$	276
11.3. Weglängen und die trigonometrischen Funktionen	285
11.4. Die trigonometrischen Funktionen	289
11.5. Der Begriff des Winkels	294
11.6. Volumina und die Determinante	297
Kapitel 12. Kodierungstheorie *	307
12.1. Einführung und Definitionen	307
12.2. Lineare Codes	310
12.3. Die Gilbert-Varshamov-Schranke	312
12.4. Hamming-Codes	312
12.5. Das McEliece-Verfahren	312
12.6. Quellen/Weiterführende Literatur	313
Kapitel 13. Graphentheorie *	315
13.1. Definition	315
13.2. Ramsey-Zahlen	316
13.3. Die Adjazenzmatrix eines Graphen	317
13.4. Teildreiecke suchen	319
13.5. Eigenwerte der Adjazenzmatrix	320

13.6. Ausblick: Planare Graphen	321
13.7. Weitere Literatur	324
Anhang A. Zusammenfassung *	325
A.1. Gruppen und Körper	325
A.2. Vektorräume	326
A.3. Matrizen	330
A.4. Determinante	332
A.5. Eigenwerte	333
Anhang B. Mathematische Ergänzungen *	335
B.1. Das Lemma von Zorn	335
Anhang C. Verschiedenes *	339
C.1. Mathematische Sprechweisen, Code words	339
C.2. Mathematische Texte am Computer schreiben.	342
Anhang D. Bemerkungen zur Literatur *	345
D.1. Lehrbücher	345
D.2. Vorlesungsskripte zur Linearen Algebra	347
D.3. Bücher übers Problemlösen und Beweisen	347
D.4. Englische Lehrbücher zur linearen Algebra	348
D.5. Verschiedene Buchempfehlungen	349
D.6. Webseiten, Videos	351
Anhang. Literaturverzeichnis	353
Anhang. Index	355

Einführung, Motivation

1.1. Bevor es losgeht – Allgemeines zum Mathematikstudium und speziell zum kommenden Wintersemester 2020/21

Herzlich Willkommen zur Vorlesung *Lineare Algebra I* im Wintersemester 2020/21 an der Universität Duisburg-Essen – ich freue mich, dass Sie dabei sind! Gerne hätte ich Sie persönlich im Hörsaal begrüßt, aber dieses Semester findet wegen der Corona-Epidemie unter besonderen Umständen statt, so dass es für den Moment bei einer Begrüßung auf dem Papier und per Video bleiben muss.

Die üblichen Präsenzvorlesungen können nicht stattfinden, sie werden ersetzt durch online bereitgestelltes Material wie dieses Skript, Videos und Aufgaben, die Sie am Rechner bearbeiten können. Andere Bestandteile des Moduls Lineare Algebra bleiben: Insbesondere die Hausaufgaben, die wöchentlich abgegeben werden müssen, und die für Sie gleichermaßen ein wichtiges mathematisches »Training« und die Möglichkeit sind, eine detaillierte Rückmeldung über Ihre Fortschritte zu erhalten. Hoffentlich wird es auch möglich sein, die Übungsgruppen in Präsenz durchzuführen, so dass Sie die Möglichkeit haben, in diesen kleinen Gruppen Ihre Fragen loszuwerden, mathematisch zu arbeiten, und andere Studierende kennenzulernen, mit denen Sie auch außerhalb der Uni-Veranstaltungen studieren können.

Gleich bleiben auch die Lernziele (siehe auch Abschnitt 2.1) des Moduls.

Und auch die wichtigste Tätigkeit, um diese Lernziele zu erreichen, bleibt: Dass Sie selbst über mathematische Probleme nachdenken! Auch wenn er offensichtliche Schwachstellen hat, mag ich den folgenden Vergleich:

Mathematik lernen ist wie schwimmen lernen

(oder ersetzen Sie, wenn Sie möchten, das Schwimmen durch eine andere Sportart oder dadurch, ein Musikinstrument zu lernen ...). Um richtig schnell schwimmen zu lernen, ist es hilfreich, darüber Bücher zu lesen und sich Videos von Schwimmstars anzuschauen. Noch wichtiger ist es, Ratschläge von einer Trainer*in zu bekommen, was man falsch macht und worauf man achten muss. Vor allem zählt aber, wie viel Zeit man im Wasser verbringt, und dass man diese Zeit sinnvoll nutzt. Das können Sie auch aufs Klavierspielen übertragen: Natürlich kann es helfen, sich ein Klavierkonzert auf CD anzuhören. Wichtiger ist eine Klavierlehrer*in, die konkrete Hinweise gibt, und das Lernen strukturiert. Vor allem zählt aber die Zeit, die man selbst am Klavier sitzt und übt. Und so ist es eben auch im Mathematikstudium: Bücher, Skripte, Videos (und Vorlesungen) sind wichtig. Fragen zu diskutieren und Tipps zu bekommen – in der Vorlesung oder vor allem in den Übungsgruppen – ist ebenso wichtig, oder vielleicht noch wichtiger. In erster Linie zählt aber die Zeit, in der Sie sich selbst den Kopf zerbrechen und selbst mathematisch arbeiten. Im Studium wird das auch konkret eingefordert in Form der schon erwähnten Hausaufgaben. Die wöchentlich verteilten Übungszettel werden Ihnen einiges abverlangen. Sie sind das essenzielle Training, und auch das beste Messinstrument für die erzielten Erfolge. Es zeigt sich immer wieder, dass die auf den Übungszetteln erreichte Punktzahl sehr direkt mit dem Erfolg in der Abschlussklausur zusammenhängt.

Zusätzliches Material (Bücher, Skripte, Videos, Aufgaben, ...) gibt es im Überfluss. Schauen Sie sich Anhang D an oder geben Sie *Lineare Algebra* in einer Suchmaschine ein. Es kann durchaus nützlich sein, mit einigen dieser Quellen zu arbeiten, *es ersetzt aber nicht, dass Sie sich selbst mit dem Stoff auseinandersetzen und selber Probleme knacken*. Selber Probleme zu lösen ist eine der wichtigsten Fertigkeiten, die in dieser (und allen folgenden) Mathematikvorlesungen erworben werden sollen. Damit Sie darin immer besser werden, müssen Sie das trainieren, und je mehr Sie passiv konsumieren, desto mehr Gelegenheiten verschenken Sie, denselben Stoff »aktiv« zu bearbeiten. Es ist nicht so einfach, hier eine gute Balance zu finden – im Zweifelsfall sollten Sie eher mehr Zeit darauf verbringen »sich selbst den Kopf zu zerbrechen«, als weiteres Material zu suchen.

”

But after having finished the classes, I came to an unfortunate conclusion—the videos don't matter too much. Having a video explanation of a concept is nice, but it's rarely superior to the same explanation in text. Text even has the advantages of searchability and nonlinearity, features missing in video. What mattered was having practice problems and projects. [...]

I feel the fuss over video has two causes:

- (1) Students associate university subjects with lectures, and don't know how to learn without them.
- (2) Watching videos is fun, doing practice problems is hard.

[...] **One reason I'm pessimistic about video is that it gives a false sense of accomplishment.**

Scott H. Young^a

^a<https://www.scotthyoung.com/blog/2012/11/13/why-lectures/>

1.2. Dieses Skript

Das Ziel dieses Skripts ist in erster Linie, den Stoff der Vorlesung zu dokumentieren und so ausführlich darzustellen, dass Sie ihn lernen können. Wie bei mathematischen Texten üblich, stelle ich mir eine »aktive« Leser*in vor: Nehmen Sie sich Papier und einen Stift zur Hand und betrachten Sie alle Schritte, die Sie nicht verstehen, als kleine Übungsaufgaben, und Rechenschritte, die Ihnen zu schnell gehen, als Rechenaufgaben, die Sie lösen sollten, bevor Sie weiterlesen. Notieren Sie sich die Fragen, die bleiben, damit Sie am nächsten Tag noch einmal darauf zurückkommen können, oder sie im Forum auf der Moodle-Seite oder in Ihrer Übungsgruppe stellen können.

Ich versuche an vielen Stellen implizit vorausgesetztes Wissen explizit machen. Siehe zum Beispiel Anhang C.I. Dies ist ohnehin zu Beginn des Studiums besonders wichtig, und umso mehr in diesem Semester, wo die Möglichkeiten, Fragen zu stellen, eingeschränkt sind. Dadurch sind vielleicht manche Passagen etwas länglich; ich hoffe, dass es Ihnen leicht fallen wird, das zu überspringen, was Ihnen schon bekannt ist.

Darüberhinaus habe ich versucht, Ihnen einige zusätzliche Informationen bereitzustellen, mit denen Sie sich beschäftigen können, wenn Sie Lust haben – zum Beispiel, wenn es Ihrer Motivation dafür dient, sich mit dem eigentlichen Stoff zu beschäftigen, oder einfach zur Abwechslung. Zum Beispiel:

- In Abschnitt 3.3 einige berühmte Beweise, die »jede*r mal gesehen haben sollte«.

- In Abschnitt 2.3 und an vielen Stellen im Text Hinweise zu Anwendungen der Linearen Algebra.
- In Anhang D einige Literaturhinweise.

Ich hoffe, dass das Skript damit gleichzeitig zu einer »Fundgrube« für Sie wird, um die vielen Verbindungen, die von der Linearen Algebra zu anderen Fragen innerhalb und außerhalb der Mathematik bestehen, zu erkunden. Und ich hoffe, dass diese Ergänzungen Ihrer Motivation dienen, sich mit dem eigentlichen Stoff der Vorlesung zu beschäftigen (wenn Ihre Motivation an dieser Stelle der Unterstützung bedarf), und Sie nicht zu sehr ablenken.

Einige Abschnitte sind mit einem Sternchen (*) markiert. Dies sind Ergänzungen, die nicht zum klausurrelevanten Vorlesungsstoff gehören. Die als *Ergänzung* bezeichneten Einschübe sind ebenfalls optional und nicht klausurrelevant. Gleichwohl können sie natürlich manchmal das Verständnis des klausurrelevanten Stoffs befördern.

Den Kern der Vorlesung bilden die Kapitel 4 bis 10.

”

The best teacher, the most successful teacher I had at Breslau, was typical of the point of view of education, was a man in algebra, his name was Rosanes. His name is not known any more. His great success as a teacher was really due to the fact that he didn't teach things very well. He came to the platform. There was a blackboard. In his right hand he had some chalk, in the left hand he had a wet sponge. He turns his back to the audience, and he mumbled something towards the blackboard, and scribbled something in small letters on the blackboard, covering it up with his body. And as he moved along he erased what he had written. And then the student always had to try to snatch a few words. Then there was an enormous task after class, one sat there for another half hour to try to put together the pieces. If one succeeded, one really had learned enormously much.

Interview mit Richard Courant^a

^a<https://www.aip.org/history-programs/niels-bohr-library/oral-histories/4562>

Zur Auflockerung habe ich an einigen Stellen Zitate eingestreut. Ich denke, dass sie alle (mindestens) ein Körnchen Wahrheit enthalten. Aber sie sollten sie nicht zu ernst nehmen.

Neben dem Skript umfasst die Vorlesung die folgenden Teile:

- Auf der Moodle-Seite zur Vorlesung finden Sie alle wichtigen Informationen zur Vorlesung sowie Links zu allen Materialien, die im Rahmen der Vorlesung bereitgestellt werden.
- Lernvideos – ich werde Ihnen wöchentlich Lernvideos zur Verfügung stellen, in denen ein Überblick über die wichtigen Themen der Vorlesungswoche gegeben wird; einige Ergebnisse und Beweise im Detail diskutiert werden; und eine oder mehrere Übungsaufgaben der Vorwoche erklärt werden,
- Fragestunde – einmal wöchentlich findet eine Videokonferenz statt, in der Ihre Fragen behandelt werden können und eventuell weitere Aufgaben besprochen werden können,
- Übungsgruppe – einmal wöchentlich findet (nach jetzigem Stand in Präsenz am Campus Essen) eine Übungsgruppe statt, in der Sie in einer Gruppe von ca. 20 Studierenden Fragen stellen und diskutieren können, und zusammen einfache Aufgaben zum Vorlesungsstoff bearbeiten können,

- Hausaufgaben – jede Woche wird ein Übungsblatt mit Aufgaben herausgegeben, deren Lösungen Sie in der Folgewoche abgeben müssen (gescannt, per Email). Die Übungsleiter*innen korrigieren die Lösungen und schicken Sie Ihnen zurück. Für die Zulassung zur Klausur benötigen Sie mindestens 50% der in den Hausaufgaben erreichbaren Punkte.
Sie können und sollten die Aufgaben in einer Gruppe von zwei oder drei Personen abgeben. Die mit dem abgegebenen Übungsblatt erzielten Punkte werden dann jedem Mitglied dieser Gruppe zugeschrieben.
- Online-Aufgaben und Online-Tests – regelmäßig werden Ihnen auf der Moodle-Seite (und per Link im Skript) Online-Aufgaben zur Verfügung gestellt, an denen Sie vor allem Ihre Rechenfertigkeiten üben können. Diese Aufgaben sind freiwillig und fließen nicht in Ihre Punktzahlen ein. Sie können sie auch mehrfach bearbeiten. Meistens werden die Zahlenwerte jeweils neu erzeugt, so dass es nicht so schnell langweilig wird.
Dreimal im Semester werde ich einen Online-Test auf die Moodle-Seite stellen. Für die Zulassung zur Klausur benötigen Sie mindestens 50% der in den Online-Tests erreichbaren Punkte.
- Eine Online-Version des Skripts finden Sie unter math.ug/1a1-ws2021/¹. Dort können Sie einige Teile »einklappen« (zum Beispiel die Ergänzungen) und dadurch eventuell besser das ausblenden, womit Sie sich nicht beschäftigen möchten. Ganz zufrieden bin ich noch nicht mit der Formatierung; ich hoffe, im Laufe der Zeit noch einige Sachen verbessern zu können. Rückmeldungen sind aber auch jetzt schon willkommen.

Einige dieser Sachen sind (auch) für mich ein Experiment – teilweise, aber nicht nur, der speziellen Situation in diesem Semester geschuldet. Insofern bin ich dankbar für Kommentare, was aus Ihrer Sicht gut, und was weniger gut funktioniert. Ich werde versuchen darauf einzugehen und gegebenenfalls, auch im Licht meiner eigenen Erfahrungen, einige Dinge im Lauf des Semesters ändern. Die Moodle-Seite wird jedenfalls als Fixpunkt bleiben, wo Sie immer die aktuellen Informationen finden können.

1.3. Wie denkt man über ein mathematisches Problem nach? *

... oder umformuliert: Wie löst man ein mathematisches Problem? (In Ihrem Fall heißt das erstmal: Wie löse ich eine Übungsaufgabe? Und in einigen Monaten dann: Wie löse ich eine Klausuraufgabe?)

Dies ist offenbar eine entscheidende Frage, um erfolgreich das Mathematik-Studium zu absolvieren (und auch, um darüberhinaus Mathematik zu erforschen).

Leider kennt niemand eine einfache Antwort darauf. Allerdings sind sich alle einig, dass man diese Fertigkeit trainieren und weiterentwickeln kann. Deshalb werden wir das Problemlösen ausführlich üben.

Einige allgemeingültige Hinweise kann man aber trotzdem geben, hier gemünzt auf die Bearbeitung einer Übungsaufgabe:

- Stellen Sie sicher, dass Sie die Aufgabenstellung verstehen. Vergewissern Sie sich die Definitionen aller Fachbegriffe, die in dieser Aufgabe vorkommen, oder schlagen Sie sie gegebenenfalls nach. Machen Sie sich klar, was zu tun ist: eine Rechnung, der Beweis einer Gleichheit oder einer Folgerung, das Finden eines Beispiels, ...
- Erinnern Sie sich an Ergebnisse, die zu den in der Aufgabe vorkommenden Begriffen in der Vorlesung behandelt wurden. Gibt es eine äquivalente Charakterisierung? Gab es Beispiele?

¹<https://math.ug/1a1-ws2021/>

- Überlegen Sie sich, ob die Aussage, die Sie zeigen sollen, für Sie plausibel ist. Können Sie Beispiele angeben, in denen sie richtig ist? Oder versuchen Sie, ein Gegenbeispiel zu finden. (Wenn Sie zeigen können, dass das nicht geht, sind Sie auch fertig ...)
- Welche Standardverfahren könnte man anwenden, um die Aufgabe zu lösen? Das kann ein Rechenverfahren (Gauß-Algorithmus) oder ein Beweisverfahren (Beweis durch vollständige Induktion, Widerspruchsbeweis, ...) sein, das Sie schon kennengelernt haben.

Natürlich wird es viele Aufgaben geben, wo Sie diese Schritte nicht alle durchgehen müssen, weil Sie direkt sehen, was zu tun ist. Und manchmal wird es auch eine Aufgabe geben, die Sie auch mit dieser »Anleitung« noch nicht lösen können.

In diesem Fall: **Stellen Sie Fragen!** Sie sollten versuchen, viele Fragen zu stellen. Nicht nur zu den Übungsaufgaben, die Sie im Moment nicht lösen können, sondern auch zu Schritten im Skript oder den Videos, die Sie nicht verstanden haben, und ganz allgemein zu den Begriffen, die Sie in der Vorlesung kennenlernen.

Wieso, weshalb, warum?

Manchmal ist es nicht so einfach, gute Fragen zu stellen. Arbeiten Sie daran, Ihre Fragen zu konkretisieren (vielleicht, indem Sie sie aufschreiben) und stellen Sie sie – sich selbst, anderen Studierenden, Ihrer Übungsleiter*in, in der Fragestunde. Nehmen Sie sich Zeit, gute Fragen zu finden. Je mehr Sie sich in Übungs und Fragestunde aktiv beteiligen statt nur zuzuhören, desto mehr werden Sie profitieren. Es lohnt sich, Zeit zur Vorbereitung zu investieren.

Sie werden sehen, dass es einem selbst hilft, etwas zu erklären, das man schon verstanden hat (oder glaubt, verstanden zu haben). Deshalb können Sie grundsätzlich davon ausgehen, dass jede*r gerne bereit sein wird, sich Ihre Fragen anzuhören und Ihnen zu helfen.

Gute Fragen zu finden ist nicht leicht! Sie zu stellen, kostet vielleicht Überwindung; in der Corona-Zeit kann man nicht direkt in der Vorlesung, und nicht so leicht in der Vorlesungspause/auf dem Flur/im LuDi anderen Studierenden eine Frage stellen. Machen Sie sich bewusst, dass es trotzdem den Aufwand wert ist, Ihre Fragen zu stellen und anderen mit deren Fragen zu helfen.

Versuchen Sie, Ihre Arbeitszeit sinnvoll zu nutzen: Reservieren Sie sich genügend Qualitätszeit, in der Sie ausgeruht und ungestört sind, für das Studium. Dass man die richtige Idee hat, muss man vorbereiten – es bringt nichts, stundenlang auf ein leeres Blatt zu starren, wenn man gar nicht weiß, wo man anfangen soll. Versuchen Sie dann, die oben genannten Schritte durchzuführen (Definitionen und verwandte Sätze nachschauen, Beispiele finden, Fragen formulieren). Wenn nichts hilft, dann gehen Sie zur nächsten Aufgabe.

Wie bearbeite ich ein Übungsblatt^a von Manfred Lehn (Univ. Mainz).

^a<https://www.agtz.mathematik.uni-mainz.de/wie-bearbeitet-man-ein-uebungsblatt-von-prof-dr-manfred-lehn/>

Es gibt viele Bücher zu dem Thema »Wie löse ich ein mathematisches Problem« (oder eine Aufgabe), siehe Anhang D.3 für eine kleine Auswahl.

1.4. Gute wissenschaftliche Praxis

Ein Thema, das mir am Herzen liegt und das ich daher hier ganz an den Anfang stelle, sind die Regeln Guter Wissenschaftlicher Praxis. »Gute wissenschaftliche Praxis« bedeutet, Wissenschaft ehrlich und redlich zu betreiben. Für die Veröffentlichung von wissenschaftlichen Texten bedeutet das insbesondere:

- Keine Ergebnisse als korrekt/vollständig darzustellen, von denen man weiß, dass sie das nicht sind,
- keine Daten (oder Ergebnisse) zugrundelegen, von denen man weiß, dass sie nicht korrekt/stichhaltig sind (oder die man gar selbst verfälscht oder den eigenen Wünschen angepasst hat),
- die Urheberschaft der Ergebnisse vollständig und zutreffend anzugeben.

Gute wissenschaftliche Praxis^a an der Universität Duisburg-Essen.

^a<https://www.uni-due.de/de/gute-wissenschaftliche-praxis/>

Zu Recht erhält dieses Thema in Bezug auf Plagiate in Doktorarbeiten in den letzten Jahren auch in der Öffentlichkeit Aufmerksamkeit. Es ist eine wesentliche Bedingung für das Funktionieren von Wissenschaft, dass diese Grundregeln eingehalten werden. Das beginnt »im Kleinen« – auch an Master-Arbeiten und Bachelor-Arbeiten wird der Anspruch gestellt, dass Ergebnisse (und erst recht wörtliche Zitate, die allerdings in mathematischen Texten selten sind), die aus anderen Quellen übernommen werden, entsprechend gekennzeichnet sind.

Gute wissenschaftliche Praxis^a bei der Deutschen Forschungsgemeinschaft.

^ahttps://www.dfg.de/foerderung/grundlagen_rahmenbedingungen/gwp/

Dies sollten Sie sich von vornherein zur Gewohnheit machen. Für Ihre Vorlesungen bedeutet das: Lösungen von Übungsaufgaben werden nicht von Kommiliton*innen abgeschrieben. Andere Quellen (Bücher, Skripte, Internetseiten), die Sie benutzen, müssen genannt werden.

Verschenken Sie nicht die Gelegenheit, eine Rückmeldung zu Ihren eigenen Bemühungen zu erhalten. Verschenden Sie nicht die Zeit Ihrer Übungsleiter*in dadurch, dass diese eine abgeschriebene Lösung korrigieren muss.

1.5. Zu erwartender Arbeitsaufwand

Das Mathematikstudium ist nicht unbedingt ein einfaches Studium – der Arbeitsaufwand ist hoch. Dafür ist es auch ein besonders interessantes und besonders nützliches Studium. Sie werden viel lernen, und zwar interessante Mathematik und gleichzeitig Ihre Kompetenzen erweitern, komplexe Probleme zu analysieren und zu lösen, und diese Lösungen präzise und verständlich anderen Personen darzustellen. Dass Mathematiker*innen das können, ist aus meiner Sicht der Hauptgrund dafür, dass sie sich um ihre Berufsaussichten keine Sorgen machen müssen: Mathematiker*innen sind gefragt in vielen verschiedenen Branchen, und das war in den letzten Jahrzehnten durchgängig genauso.

Die Einrichtung von Bachelor- und Master-Studiengängen hat es mit sich gebracht, dass das Studium kleinteiliger organisiert wurde als vorher und dass man versucht, viele Schritte genau zu quantifizieren. Der durchschnittliche Arbeitsaufwand für einen Credit-Punkt wird auf 30 Stunden festgelegt. Die 9 Credit-Punkte, die für die Vorlesung *Lineare Algebra I* erreicht

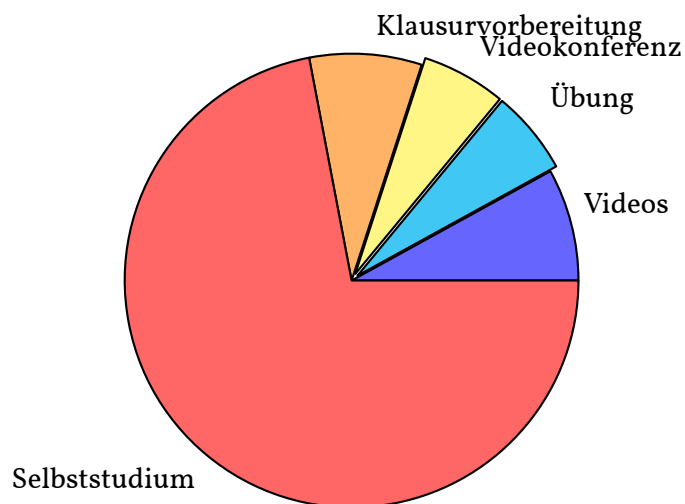
werden, entsprechen also ungefähr einem Arbeitsaufwand von 270 Stunden. Das ist eine ganze Menge!

Das soll nun nicht bedeuten, dass ich dieser Rechnung zustimme. Mindestens ist sie mit Vorsicht zu genießen, schon allein weil der Aufwand individuell sehr verschieden ist (einerseits, wie viel nötig ist; andererseits, zu wie viel die/der Einzelne Lust hat).

Etwas genauer ist es an der Uni Duisburg-Essen so, dass die Credit-Punkte für das *Modul Grundlagen der Linearen Algebra* vergeben werden, das die beiden Vorlesungen Lineare Algebra 1 und 2 umfasst, und das dann 18 CP wert ist. Insofern ist, wenn man überhaupt Arbeitsstunden ausrechnen möchte, ein Teil der oben genannten 270 Stunden in die vorlesungsfreie Zeit zur Nachbereitung zu verlagern, und ein Teil für die Vorbereitung auf die abschließende mündliche Prüfung (die auch die Note bestimmt) einzuplanen.

Jedenfalls: **Das Studium ist ein Vollzeitjob.**

Wie das Diagramm illustriert, das eine grobe Aufteilung der veranschlagten Arbeitszeit zeigt, entfällt nur ein kleiner Teil der veranschlagten Arbeitszeit auf »Veranstaltungen« (in diesem Semester: Video-Fragestunde, Übungsgruppe). Das ist auch unter normalen Umständen nicht viel anders, mit dem Unterschied, dass Sie jetzt keine festen Termine für die Vorlesung selber haben, sondern selbst entscheiden können, wann Sie das Skript/die Videos anschauen (und sich eventuell mehr disziplinieren müssen, das tatsächlich zu tun, als wenn Sie zu einer festen Zeit zur Vorlesung kommen würden).



Was hier unter »Selbststudium« gefasst wird, kann (und sollte) natürlich viele verschiedene Formen annehmen, und soll insbesondere nicht heißen, dass Sie in dieser Zeit durchgehend allein arbeiten müssten. Zum Beispiel fällt mir ein:

- Bearbeitung der Hausaufgaben,
- Bearbeitung von Online-Aufgaben,
- sich mit anderen Studierenden treffen und zusammen arbeiten (Fragen stellen und beantworten),
- auf der Moodle-Seite Fragen stellen und beantworten,
- Vorbereitung auf die Übungsgruppe und die Fragestunde
- das Nacharbeiten des Skripts,
- »selber denken«,
- gegebenenfalls andere Quellen lesen/durcharbeiten, ...

Bei einem Problem »festzustecken« ist für die Mathematiker*in nicht ungewöhnlich. Auch in der Forschung steckt man sozusagen einen großen Teil der Zeit fest. Das bedeutet, dass man sich davon allein nicht zu sehr irritieren lassen sollte. Andererseits bedeutet es, dass man Strategien finden muss, um in diesem Zustand weiterzukommen. Denn das Feststecken einfach hinzunehmen, bringt einen der Lösung natürlich auch nicht näher. Oft lohnt es sich, noch einmal die Definitionen der relevanten Begriffe darüber nachzuschauen und sich an Beispielen zu illustrieren; die Ergebnisse zu wiederholen, die man über diese Begriffe kennt (bzw. kennen sollte), und nachzuvollziehen, wie sie bewiesen wurden; zu überlegen, ob man das Problem lösen könnte, wenn man es ein bisschen verändert...

Ein Link zum Thema zum Blog »Math with bad drawings« von B. Orlin: [The state of being stuck](https://mathwithbaddrawings.com/2017/09/20/the-state-of-being-stuck/)^a

^a<https://mathwithbaddrawings.com/2017/09/20/the-state-of-being-stuck/>

Mein Ziel ist es, Ihnen das Lernen so leicht wie möglich zu machen. Ich bin für Anregungen, was man verbessern könnte, offen und dankbar. Einen Weg, auf dem dies möglich wäre, ohne dass Sie sich anstrengen müssten, kenne ich leider nicht.

Meine ehrliche Hoffnung ist, dass Sie den hohen Arbeitsaufwand nicht als etwas Negatives sehen, sondern sich darauf einlassen. Die Zeit des Studiums kann und soll eine schöne Zeit des Lebens sein, in der Sie die Möglichkeit haben, viel Zeit mit etwas zu verbringen, das Sie wirklich interessiert. Mathematik zu verstehen und seine eigenen Fortschritte zu sehen, ist sehr befriedigend, und Durchhaltevermögen wird zusätzlich belohnt mit einem Studienabschluss, der Ihnen exzellente Berufsaussichten eröffnet.

Die Vorlesung »Lineare Algebra I«

2.1. Lernziele

In diesem Abschnitt möchte ich Ihnen kurz die Lernziele dieser Vorlesung vorstellen. Die mathematischen Begriffe, die in 2.1.2 und 2.1.3 genannt werden, kennen Sie jetzt natürlich größtenteils noch nicht; die Listen dort dienen dazu, dass Sie gegebenenfalls später darauf zurückkommen können.

2.1.1. Allgemeines. (Dies ist gewissermaßen das Wichtigste, und vielleicht auch das Schwierigste ...)

- (1) Mathematische Sprache verstehen,
- (2) Definitionen und Aussagen verstehen und präzise wiedergeben,
- (3) Beweise nachvollziehen, auf Korrektheit/Vollständigkeit prüfen und reproduzieren,
- (4) Beweise selbst finden und verständlich dokumentieren.

Diese Kompetenzen gehen ineinander über. (Zum Verstehen muss man »Details einfüllen«, d.h. Mini-Beweise selbst entwickeln. Je weiter man kommt, desto mehr ...)

Auch wenn man nicht forschende Mathematiker*in werden möchte, muss man im Mathematikstudium beweisen lernen. Denn um den Stoff des Studiums zu meistern, muss sich die Lerngeschwindigkeit im Laufe der Zeit erhöhen, das bedeutet, dass in Vorlesungen späterer Semester Details, die in den Anfängervorlesungen ausführlich begründet werden, nicht weiter erklärt werden: Die Student*in muss sich selbst überlegen, warum etwas richtig ist, d.h., einen (kleinen oder manchmal auch größeren) Beweis »einfüllen«. Umso mehr gilt das, wenn man seine Bachelor-Arbeit schreibt oder Forschungsartikel lesen möchte, wie bei der Vorbereitung auf ein (Master-)Seminar oder bei der eigenen Masterarbeit.

Diese Fertigkeit (Problemlösen, Strukturen erkennen, ...) ist das Wichtigste, was diejenigen aus dem Mathematikstudium mitnehmen, die später nicht an der Uni bleiben (oder in der Schule unterrichten). Und sie ist auch der Hauptgrund, warum Mathematiker*innen keine Sorgen zu haben brauchen, ob sie einen Arbeitsplatz finden.

2.1.2. Begriffe der Linearen Algebra. Sie sollen am Ende der Vorlesung vertraut sein mit den folgenden Themen/Objekten/Begriffen:

- (1) Mengen, Abbildungen
- (2) Körper
- (3) Lineare Gleichungssysteme
- (4) Vektorraum, Untervektorraum
- (5) linear (un-)abhängig, Erzeugendensystem, Basis, Dimension
- (6) lineare Abbildungen, Kern, Bild, Rang, Dimensionsformel für lineare Abbildungen
- (7) Zusammenhang lineare Abbildungen und Matrizen, Zeilenrang/Spaltenrang
- (8) Gruppen, S_n , $GL_n(K)$, $SL_n(K)$

- (9) Determinanten (einer Matrix, eines Endomorphismus), Cramersche Regel
- (10) Eigenwerte und Eigenvektoren (Grundbegriffe)

Schauen Sie auch einmal in das Modulhandbuch zu Ihrem Studiengang:
Bachelor Mathematik^a (Seite 6), **Lehramt GyGeBK^b** (Seiten 9ff.)

^ahttps://www.uni-due.de/imperia/md/content/mathematik/2016_04_12_modulhandbuch-ba-aktiv.pdf

^bhttps://www.uni-due.de/imperia/md/content/didmath/lehre/mhb_ba_gyge_190523.pdf

2.1.3. Rechenmethoden.

- (1) Gauß-Algorithmus (zum Beispiel zur Lösung von linearen Gleichungssystemen; Bestimmung von Kern und Bild einer Matrix; Bestimmung von Basen; Berechnung von Determinanten)
- (2) Rechnen mit Matrizen (Matrizenprodukt, Determinanten)
- (3) Bestimmung von Eigenwerten und Eigenräumen eines Endomorphismus/einer Matrix

Die Rechenmethoden stehen hier bewusst an letzter Stelle (auch wenn wir in der Vorlesung recht bald zu dem zentralen Gauß-Algorithmus kommen werden), weil die anderen Punkte wichtiger sind. Dennoch empfehle ich, auch viel zu rechnen, weil das erstens ausreichend geübt werden muss, damit es fehlerfrei und zügig funktioniert, und zweitens das Verständnis der abstrakten Begriffe profitiert, wenn man sich genügend viele Beispiele rechnerisch erarbeitet hat.

Manchmal wird unterschieden zwischen *Rechenaufgaben* und *Beweisaufgaben*. Der Übergang ist natürlich fließend und auch für eine »Rechenaufgabe« werden Sie oft erstmal ein bisschen nachdenken müssen, um herauszufinden, welche Rechnung überhaupt gefragt ist. Aber in der Tat gibt es verschiedene Kompetenzen, die beide wichtig sind: Einerseits, eine Rechnung korrekt und in angemessener Zeit durchführen zu können; andererseits, eine Aussage begründen (oder gegebenenfalls widerlegen) zu können. In beiden Fällen ist es wichtig, dass Sie Ihre Lösung angemessen dokumentieren können.

Die Rechenkompetenz allein lässt sich gut in den Online-Aufgaben abdecken, außerdem fallen die Rechenaufgaben in der Regel den Studienanfänger*innen leichter. Deshalb liegt in den Hausaufgaben, die Sie abgeben müssen und die dann per Hand korrigiert werden, und in der Klausur der Schwerpunkt auf den Beweisaufgaben. Dabei geht es nicht darum, geniale Ideen zu entwickeln, sondern einfache Zusammenhänge (wieder) zu erkennen und verstandene Argumente klar dokumentieren zu können.

Online-Aufgaben mit NUMBAS

Zur Vorlesung werden Online-Aufgaben angeboten, die auf dem System **NUMBAS^a** basieren, das an der Newcastle University entwickelt wird. Die Online-Aufgaben werden auf der Moodle-Seite der Vorlesung verlinkt.

^a<https://www.numbas.org.uk/>

2.2. Wozu Lineare Algebra?

2.2.1. Wozu man Lineare Algebra braucht (oder gebrauchen kann), lässt sich natürlich leichter jemandem erklären, der schon etwas über Lineare Algebra weiß (und am besten

auch etwas über die Problemstellungen weiß, wo man eine Anwendung finden kann). Trotzdem dazu einige Bemerkungen. Die Lineare Algebra hat sich aus der Theorie der linearen Gleichungssysteme entwickelt, und so ist auch diese Vorlesung aufgebaut. Wir suchen also beispielsweise alle reellen Zahlen x, y , die die beiden Gleichungen

$$3x + y = 7$$

$$x - 2y = \frac{2}{3}$$

erfüllen. Diese Art von Fragestellung kennen Sie aus der Schule und haben wahrscheinlich keine Schwierigkeiten, die Lösungsmenge des Gleichungssystems zu bestimmen. Wir werden aber in Kürze Fragen stellen und Antworten kennenlernen, die über den Schulstoff weit hinausgehen.

”

Ich habe bey mir von Jugend auff eine unersättliche Begierde die Wahrheit gewiß zu erkennen und anderen zu dienen gefunden. Daher als ich bey Zeiten vernahm, daß man der Mathematick eine ungezweiffelte Gewißheit zuschreibe, und absonderlich die Algebra als eine richtige Kunst verborgene Wahrheiten zuentdecken rühme; Hingegen aus den so vielfältigen und wiedrigen Meinungen der Gelehrten in anderen Sachen, die zur Mathematick nicht gehören, und aus den steten Aenderung, die darinnen vorgenommen werden, mir auch dazumahl genung begreiflich war, daß es ausser der Mathematick an einer völligen Gewißheit meistens fehle; Erweckte bey mir die Begierde zur Warheit eine Liebe zur Mathematick und sonderlich eine Lust zur Algebra, um zusehen, was doch die Ursache sey, warum man in der Mathematick so große Gewissheit habe, und nach was vor Regeln man daselbst dencke, wann man verborgene Wahrheiten zum Vorscheine bringen will, damit ich mich desto sicherer bemühen möchte auch ausser der Mathematick dergleichen Gewißheit zu suchen ...

Aus der Vorrede zu *Mathematisches Lexicon*, C. Wolffens, Leipzig 1716;
gefunden in [Jä].

Wenn Sie fragen »**Wozu brauche ich das?**«, dann hoffe ich, dass die erste Antwort ist: Weil es Sie interessiert, eine mathematische Theorie kennenzulernen. Oder genauer: Das Problemgebiet der Linearen Gleichungssysteme strukturell und systematisch kennenzulernen, daran zu sehen, wie Mathematik auf Hochschulebene betrieben wird, offene Fragen kennenzulernen und eine Grundlage für das gesamte weitere Mathematikstudium zu legen.

”

In the broad light of day mathematicians check their equations and their proofs, leaving no stone unturned in their search for rigour. But, at night, under the full moon, they dream, they float among the stars and wonder at the miracle of the heavens. They are inspired. Without dreams there is no art, no mathematics, no life.

Michael Atiyah (NAMS Jan 2010 p.8)

Fundort: <https://www.jmilne.org/math/>

Ich werde versuchen, dieses Interesse zu erhalten und weiter zu nähren, indem ich im Skript an vielen Stellen optionale Ergänzungen einfüge, die Verbindungen zu anderen Bereichen

der Mathematik und zu konkreten Fragen aufzeigen, die man mit Linearer Algebra lösen kann. In der »Vorlesung« selbst werden wir uns allerdings auf den Kernstoff konzentrieren. Wenn Ihnen das ausreicht, ist es völlig in Ordnung, wenn Sie diese Ergänzungen nicht weiter beachten. Wenn es Ihnen leichter fällt, sich durch den eigentlichen Stoff durchzubeißen, wenn Sie wissen, dass er auch zu etwas gut ist (und sich nicht darauf verlassen möchten, dass ich Ihnen das versichere), dann können Sie auf die Ergänzungen zurückgreifen, die Sie ansprechen.

2.2.2. Berufsbezug des Mathematikstudiums. Regelmäßig wird in den Umfragen zu den (Anfänger-)Vorlesungen ein größerer Berufsbezug gewünscht. In der Vorlesungsumfrage zu dieser Linearen Algebra I (WS 2020/21) wurde die Frage *Ein Bezug zwischen Veranstaltungsinhalten und Berufsfeld bzw. Anwendung wird hergestellt.* mit der Durchschnittsnote 3,3 bewertet. Ich kann gut verstehen, dass man schon (und vielleicht besonders) in den ersten Semestern darauf schaut, wohin die Reise beruflich nach dem Abschluss des Studiums gehen kann und ob das Mathematikstudium überhaupt eine gute Vorbereitung auf eine Berufstätigkeit ist. Zum Teil wird hier sicher auch eine zusätzliche Unsicherheit »von außen« erzeugt, weil es so schwierig ist, die Inhalte des Studiums Nicht-Mathematiker*innen zu erklären.

Den konkreten Berufsbezug herzustellen, ist aus mehreren Gründen schwierig (Und aus meiner Sicht ist das auch nicht die Aufgabe eines Universitätsstudiums.)

- Natürlich gibt es nur wenige Berufe, in denen man die Determinante oder die Eigenwerte einer Matrix ausrechnen muss (oder überhaupt wissen muss, was das ist). Das macht es auch schwierig, von den mathematischen Inhalten einen ganz *konkreten* Berufsbezug herzustellen.
- Es gibt *sehr viele* und *sehr verschiedene* Berufe, die man als Mathematiker*in ergreifen kann. Zuerst einmal gibt es natürlich den Unterschied zwischen Lehramtsstudium und Fachstudiengang. Aber auch wenn man sich auf Absolventen mit Bachelor-/Master-Abschluss einschränkt: Mit Mathematik kann man wirklich alles mögliche machen.

Darum kann ich den Berufsbezug auch im Vorlesungsskript nicht mal eben im Vorübergehen herstellen. Ich möchte aber versuchen zu begründen, dass das Mathematikstudium in hervorragender Weise auf das Berufsleben vorbereitet und ausgezeichnete und sehr vielseitige berufliche Perspektiven eröffnet.

Denn im Mathematikstudium werden viele Fertigkeiten eingeübt, die in jedem Beruf nützlich und wichtig sind – in den Vorlesungen, den Seminaren, dem Verfassen der Abschlussarbeiten ... Zum Beispiel:

- Schwierige Probleme systematisch zu analysieren, kreative Lösungen (unter genau definierten Rahmenbedingungen) zu finden und Lösungen auf ihre Korrektheit zu überprüfen.
- Komplexe Sachverhalte korrekt und sehr präzise darzustellen, sowohl schriftlich als auch im Einzelgespräch und in Vorträgen vor einer Gruppe
- Fragen zu stellen
- Sich in neue Theorien selbständig einzuarbeiten
- Durchhaltevermögen
- (meistens) Englischkenntnisse – jedenfalls ist es häufig erforderlich, mathematische Texte auf Englisch zu lesen; je nach weiterer Spezialisierung hat man im Masterstudium vielleicht auch die Gelegenheit, selbst Vorträge auf Englisch zu halten,

Je nach Spezialisierung werden auch Kompetenzen erworben, die eine »konkretere« fachliche Natur haben (aber: diese Sachen kann man sich sonst auch nachträglich noch aneignen), beispielsweise

- Programmieren
- das Arbeiten mit Softwarepaketen aus Bereichen wie Numerik, Optimierung, Statistik
- Methoden der Finanz-/Versicherungsmathematik

Gerade bei der ersten Liste kann man sicherlich die Frage stellen, ob es nicht möglich ist, sich diese Fähigkeiten auf direkterem Wege als durch ein Mathematikstudium anzueignen – vermutlich schon, auch wenn ich keine direkte Empfehlung auf Lager habe. Sicher gilt aber: Ein Mathematikstudium empfehle ich nur der/demjenigen, die sich für die mathematischen Inhalte interessiert. Dieses Interesse muss sich nicht auf die gesamte Mathematik beziehen, und natürlich sind auch einmal »Durststrecken« erlaubt, in denen sich das Studium etwas zäher anfühlt. Wenn die Motivation »für die Mathematik« aber grundsätzlich vorhanden ist, dann ist das Mathematikstudium eine wunderbare Gelegenheit, gleichzeitig diesem Interesse nachzugehen und sich im Vorübergehen in Hinsicht auf die genannten Punkte weiterzuentwickeln.

Ein paar konkrete Beispiele für die Berufstätigkeit als Mathematiker finden Sie unter <https://math.ug/posts/berufsbezug-des-mathematikstudiums/>

2.2.3. Bemerkungen zum Lehramtsstudium. Die Vorlesungen Lineare Algebra 1 und 2 sind auch Bestandteil des Lehramtsstudienganges Gymnasium/Gesamtschule/Berufskolleg. Ich möchte einige Bemerkungen dazu machen, warum ich das für sinnvoll halte. Natürlich ist der Vorlesungsstoff überwiegend nicht Unterrichtsstoff. Und auch wenn Sie mit Sicherheit vom fachmathematischen Anteil des Studium für die Tätigkeit an der Schule profitieren, will ich nicht verheimlichen, dass es eines nicht geringen Aufwands bedarf, wenn Sie aus dem Vorlesungsstoff das maximal Mögliche für den Schulunterricht herausholen wollen. Warum sollten Sie trotzdem die regulären Anfängervorlesungen des Bachelor-Studiengangs Mathematik besuchen?

- Am wichtigsten aus meiner Sicht: Es ist für eine Mathematiklehrer*in essenziell, präzise formulieren zu können (und das ihren Schüler*innen mitzugeben); und entscheiden zu können, ob eine mathematische Aussage wahr oder falsch ist, und das begründen zu können.
Das lernt man in der Schule oft nicht ausreichend. Um hier eine hinreichend hohe Sicherheit zu erreichen, muss man diese Fertigkeiten trainieren. Und dann bietet es sich ja an, dieses Training anhand des grundlegenden Stoffs der Mathematik durchzuführen, wie er in den Vorlesungen Lineare Algebra und Analysis behandelt wird.
Bedenken Sie, dass es ein großer Luxus des Faches Mathematik ist, dass man sich sicher sein kann, dass alles, was man erklärt, auch richtig ist, und man das auch überprüfen kann.
- Lineare Algebra ist Schulstoff (natürlich nicht alles, was in der Vorlesung vorkommt; aber es ist wichtig, dass die Lehrer*in einen Überblick und einen deutlichen Wissensvorsprung vor den Schüler*innen hat.)
- Eine der Aufgaben der gymnasialen Oberstufe ist die Vorbereitung auf ein mögliches Studium und eine Mathematiklehrer*in sollte vermitteln können, wie Mathematik an der Universität betrieben wird.

Es ist nun so, dass der Stoff in der Vorlesung nicht mundgerecht für den Unterricht aufbereitet wird – nicht einmal näherungsweise: Teilweise ist eine erhebliche »Übersetzungsarbeit« zu leisten. Das kann nicht innerhalb dieser Vorlesung geschehen. Schon alleine aus Zeitgründen, aber auch, weil das besser jemand macht, der mehr Kontakt zur Lehrerbildung hat als eine Professor*in in der Fachmathematik. Es gibt fachdidaktische Veranstaltungen in Ihrem Studium, die sich dem widmen und Materialien, die Sie dabei unterstützen können.

Übrigens ist diese Diskrepanz ein klassisches Thema, siehe die Bücher [Elementarmathematik vom höheren Standpunkt](#)¹ von Felix Klein (erste Auflage um 1908), die historisch und mathematisch interessant sind, die man mit insgesamt über 850 Seiten allerdings nicht mal eben an einem Nachmittag durchliest.

Weitere Quellen:

B. Schwarz, P. Herrmann, [Bezüge zwischen Schulmathematik und Linearer Algebra in der hochschulischen Ausbildung angehender Mathematiklehrkräfte – Ergebnisse einer Dokumentenanalyse](#)^a, Mathematische Semesterberichte **62** (2015), 195–217.

B. Schwarz, P. Herrmann, G. Kaiser, B. Richter, J. Struckmeier, [Lineare Algebra in der Lehramtsausbildung – Wenig Bezug zum Mathematikunterricht?](#)^b

^a<https://link.springer.com/article/10.1007/s00591-015-0147-7>

^b<https://www.math.uni-hamburg.de/home/richter/la-lehramtsbezug.pdf>

2.2.4. Anwendungen der Linearen Algebra. Eine etwas andere als die oben genannte ist die Frage »**Wozu braucht man das/kann man das gebrauchen?**«. Mit Antworten darauf könnte man speziell in Bezug auf die Lineare Algebra ganze Bücher füllen, und diese Anwendungen können natürlich auch in Ihre persönliche Antwort auf die Frage einfließen, warum Sie lineare Algebra lernen möchten.

Wichtige Gründe, warum die Lineare Algebra an so vielen Stellen nützlich ist, sind, dass man lineare Gleichungssysteme einerseits theoretisch viel besser versteht als Systeme von allgemeineren Gleichungen, und sie andererseits tatsächlich lösen kann, selbst wenn sie einigermaßen groß sind, anders als viele andere Arten/Systeme von Gleichungen (irgendwo gibt es natürlich auch hierbei Grenzen). Außerdem kommen lineare Strukturen praktisch »überall« vor, zumindest näherungsweise.

Ein gutes Beispiel ist die **Analysis**. Es ist die grundlegende Strategie der Differentialrechnung, zu einer (differenzierbaren) Funktion die bestmögliche Annäherung durch eine *lineare* Funktion zu suchen. Die Steigung dieser Geraden ist die Ableitung an der betrachteten Stelle. Während die Theorie der linearen Funktionen $\mathbb{R} \rightarrow \mathbb{R}$, also von den reellen Zahlen in sich, so einfach ist, dass man alles Benötigte ad hoc bereitstellen kann, ist es in höheren Dimensionen sinnvoll, diesen Teil der Theorie auf eine systematische Grundlage zu stellen. In der Analysis 2 und Analysis 3 werden die Begriffe des Vektorraums (Definition 6.1) und der linearen Abbildung (Definition 7.1) ein wichtiges Hilfsmittel sein.

Ähnlich ist es in Numerik, Optimierung, Wahrscheinlichkeitstheorie, ... In fast allen Bereichen der Mathematik wird Lineare Algebra benutzt. Die Begriffe des Vektorraums und der Gruppe (Definition 8.1, sozusagen die mathematische Art und Weise, über Symmetrie zu sprechen, vergleiche Abschnitt 8.1.6) spielen praktisch überall in der Mathematik eine Rolle. Besonders natürlich in den fortgeschrittenen Vorlesungen der Algebra-Schiene: Algebra, Kommutative Algebra, Algebraische Geometrie, Algebraische Zahlentheorie. Da

¹<https://gdz.sub.uni-goettingen.de/id/PPN376497785>

diese sequentiell aufeinander aufbauen, ist es aber – im Vergleich etwa zur Graphentheorie (siehe Kapitel 13), die nicht viele Vorkenntnisse benötigt – schwieriger, das an dieser Stelle zu illustrieren – Sie müssen sich also noch etwas gedulden, um die Lineare Algebra dort im Einsatz zu sehen.

2.3. Konkrete Fragen, die wir im Laufe der Vorlesungen Lineare Algebra I/2 werden beantworten können *

Um den vorherigen Abschnitt nicht vollständig im Vagen zu lassen, nenne ich hier einige Fragen/Problemstellungen, die wir mit den im Laufe der Vorlesung entwickelten Methoden beantworten oder mindestens besser verstehen können.

Zum Teil handelt es sich um »Anwendungen« von Methoden der linearen Algebra auf mathematische Probleme, wo diese nicht offensichtlich sind. Zum Teil handelt es sich um Anwendungen, die außerhalb der Universität wichtig sind. Viele (mindestens ebenso interessante) Anwendungen innerhalb der Mathematik müssen leider außen vor bleiben, weil mehr Mathematik benötigt wird, als uns im Moment zur Verfügung steht, um über diese Probleme überhaupt zu sprechen.

(Für viele der folgenden Fragen/Probleme gibt es mehrere verschiedene Lösungsansätze. Man muss nicht immer Lineare Algebra benutzen – aber oft ist es möglich und nützlich.)

Wenn Ihnen die Formelsprache zu kompliziert ist, dann überspringen Sie die Frage erstmal. Es wird nicht lange dauern, bis Ihnen das keine Probleme mehr bereitet, und dann können Sie noch einmal hierher zurückkommen.

FRAGE 2.1. Die **Fibonacci-Folge**² (siehe auch [Wikipedia \(Englisch\)](#)³) ist die Folge natürlicher Zahlen, die gegeben ist durch

$$F_0 = 0, \quad F_1 = 1, \quad F_n = F_{n-1} + F_{n-2} \text{ für alle } n \geq 2.$$

Die ersten Terme der Folge lauten also 0, 1, 1, 2, 3, 5, 8, 13, 21.

Auf der Webseite mathe-vital.de^a finden Sie (neben vielen anderen sehenswürdigen Visualisierungen, auch zur linearen Algebra) mehrere Seiten zu den Fibonacci-Zahlen.

Die Fibonacci-Zahlen und ... [der Goldene Schnitt I](#)^b, [der Goldene Schnitt II](#)^c, [die Ananas](#)^d.

^a <https://mathe-vital.de/>

^b <https://mathe-vital.de/Botanik/4-1.html>

^c <https://mathe-vital.de/Botanik/4-3.html>

^d <https://mathe-vital.de/Botanik/4-6.html>

Wir werden später die folgenden Fragen leicht mit Linearer Algebra beantworten können:

- Wie kann man möglichst schnell eine einzelne Fibonacci-Zahl F_n für großes n berechnen (zum Beispiel $F_{10\,000\,000\,000}$), ohne alle Fibonacci-Zahlen dazwischen berechnen zu müssen? Siehe Beispiel 5.60.
- Was ist eine geschlossene Formel für die n -te Fibonacci-Zahl F_n , in der die kleineren Fibonacci-Zahlen nicht auftreten? Siehe Ergänzung 6.60, Beispiel 10.19.

² <https://de.wikipedia.org/wiki/Fibonacci-Folge>

³ https://en.wikipedia.org/wiki/Fibonacci_number

Die dabei verwendeten Methoden sind oft auch nützlich, um andere Folgen, die in ähnlicher Weise (durch eine »lineare Rekursionsgleichung«) definiert sind, zu analysieren. (Zum Teil lassen sich diese Fragen natürlich auch auf anderem Wege beantworten. Versuchen Sie es ruhig einmal!)

□ Frage 2.1

FRAGE 2.2. Der folgende Satz ist wichtig für das Verfahren des **Quadratischen Siebes**⁴, das ist eines der besten bekannten Verfahren, um sehr große ganze Zahlen in ihre Primfaktoren zu zerlegen. Die Frage, ob/wie man das »schnell« machen kann, ist von hoher Bedeutung für Verschlüsselungsverfahren (beziehungsweise die Frage, ob man sie knacken kann), die an allen möglichen Stellen eingesetzt werden (Online-Banking, ...), zum Beispiel das **RSA-Verfahren**⁵.

SATZ 2.3. Gegeben seien eine natürliche Zahl $n \geq 1$ und n verschiedene Primzahlen $p_1, \dots, p_n$. Wenn $a_1, \dots, a_{n+1}$ natürliche Zahlen > 1 sind, in deren Primfaktorzerlegungen nur die Primzahlen $p_1, \dots, p_n$ vorkommen, dann gibt es eine Möglichkeit, einige der Zahlen a_i so auszuwählen, dass ihr Produkt eine Quadratzahl ist.

BEISPIEL 2.4. Wir betrachten die Primzahlen 2, 3 und 7 und die vier Zahlen 7, 12, 18, 21. Dann ist $7 \cdot 12 \cdot 21 = 4 \cdot 9 \cdot 49 = (42)^2$ eine Quadratzahl. ◇

Wir kommen in Ergänzung 6.63 auf diese Frage zurück.

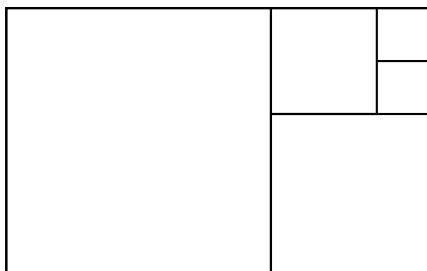
Siehe auch:

R.-H. Schulz, H. Witten, **Faktorisieren mit dem quadratischen Sieb**^a, LOG IN **172/173** (2011/12), 70–78.

^ahttps://bscw.schule.de/pub/bscw.cgi/d1024028/Schulz_Witten_Quadratisches_Sieb.pdf

□ Frage 2.2

FRAGE 2.5. Wann kann man ein Rechteck wie in der Abbildung lückenlos durch Quadrate überdecken?



Es gilt der folgende Satz:

SATZ 2.6. Sei R ein Rechteck mit Seitenlängen $a, b \in \mathbb{R}_{>0}$. Dabei sei $a \in \mathbb{Q}$ und $b \in \mathbb{R} \setminus \mathbb{Q}$. Dann lässt sich R nicht vollständig durch (endlich viele) Quadrate überdecken, die sich nicht überlappen.

Wir werden das in Kapitel 7 beweisen können, siehe Ergänzung 7.64.

□ Frage 2.5

Die nächste »Frage« ist ziemlich lang, aber ein gutes (und inzwischen ziemlich prominentes) Beispiel für eine Anwendung von Methoden der Linearen Algebra, mit denen (fast) jede*r täglich in Kontakt ist.

⁴https://de.wikipedia.org/wiki/Quadratisches_Sieb

⁵<https://de.wikipedia.org/wiki/RSA-Kryptosystem>

FRAGE 2.7 (Googles Page-rank-Algorithmus). Wir wollen das Problem betrachten, eine gute Internet-Suchmaschine zu bauen. Wir stellen uns vor, dass bereits eine Datenbank aller Webseiten, die eine Benutzer*in finden können soll, existiert. Die Frage, die wir betrachten wollen, ist, wie die Suchmaschine entscheidet, in welcher Reihenfolge die Treffer zu einem Suchbegriff angezeigt werden. Dass Google ziemlich schnell alle damaligen Mitbewerber fast vollständig vom Markt verdrängen konnte, lag wesentlich mit daran, dass bei Google viel verlässlicher die relevanten Treffer weit oben in der Liste der Suchergebnisse angezeigt wurden. Der Ansatz, den wir hier beschreiben, ist die Basis des Google-Algorithmus. (Genauer des Page-rank-Algorithmus, der die Grundlage für die Google-Suche in den ersten Jahren nach der Gründung war. Sicher wurden nicht alle Feinheiten verraten, und inzwischen wurde der Algorithmus weiterentwickelt und/oder durch andere Methoden ersetzt. Das Prinzip, das wir hier kennenlernen, ist aber natürlich nach wie vor von Bedeutung; die Frage, wie man »Wichtigkeit/Relevanz« in einem »Netzwerk« messen kann, stellt sich ja an vielen Stellen.)

Wir stellen uns vor, dass wir die Webseiten in unserer Datenbank durchnummerieren. Für jede Webseite wollen wir eine Zahl berechnen, die misst, wie »wichtig« diese Seite ist. Für die erste Seite in unserer Datenbank bezeichnen wir die zu findende Zahl mit x_1 , für die zweite mit x_2 usw., also für die i -te Seite mit x_i . Je höher der Wert x_i ist, desto höher würde die Seite in der Liste der Suchergebnisse angezeigt, wenn sie bei den Treffern dabei ist. (Wenn Sie die Formelsprache hier oder in den nächsten Absätzen stört, dann springen Sie erstmal zu Beispiel 2.8.)

Dabei messen die Zahlen x_i die »Relevanz« einer Seite unabhängig von dem jeweiligen Suchbegriff. Wenn dann eine Suche ausgeführt wird, werden die entsprechenden Webseiten als »Treffer« aus der Datenbank ausgewählt und dann in der Reihenfolge ihrer Relevanz angezeigt.

Die erste wesentliche Überlegung ist, die Relevanz einer Webseite nicht durch eine komplizierte Analyse ihres Inhalts zu messen, sondern mit einer Methode, die sich leicht mit den erhobenen Daten umsetzen lässt. Im ersten Schritt stellen wir fest, dass eine Seite umso relevanter sein dürfte, je mehr andere Seiten auf sie verlinken. Wir könnten also einfach die Anzahl dieser Links für jeden Suchtreffer zählen und die Treffer dementsprechend anordnen. Wenn wir mit L_i die Menge aller Webseiten bezeichnen, die auf die i -te Seite verlinken, so könnten wir diesen Ansatz schreiben als

$$x_i = \#L_i, \quad \text{die Anzahl der Elemente in } L_i.$$

Wir wollen dabei nicht berücksichtigen, wenn eine Seite auf sich selbst verlinkt; es soll also i kein Element von L_i sein.

Es ist aber vernünftig, das Verfahren noch etwas zu verfeinern. Es sollte eine Rolle spielen, ob ein Link auf die betrachtete Seite von einer »wichtigen« oder »unwichtigen« Seite kommt. Wir sollten in der vorherigen Formel nicht für jedes Element j von L_i (wir schreiben dann $j \in L_i$) eine 1 zählen, sondern den Einfluss entsprechend gewichten – und zwar gerade mit der Zahl x_j , die die Relevanz der Seite j angibt. Ein besserer Ansatz wäre also

$$x_i = \sum_{j \in L_i} x_j,$$

das heißt x_i ist die Summe aller Werte x_j , wo die j -te Seite einen Link auf die i -te Seite hat. Wir benutzen hier das Summensymbol $\sum$ (der große griechische Buchstabe Sigma, siehe Beispiel 3.43) als Abkürzung für die Summe aller derjenigen x_j mit $j \in L_i$. Das macht einen komplizierten Eindruck. Denn wir können x_i nicht mehr direkt berechnen, weil wir für die Berechnung der x_j ja möglicherweise x_i schon kennen müssten (wenn $i \in L_j$ ist). Trotzdem ist es nicht so schlimm, wie es aussieht: Wir können die Gesamtheit dieser Gleichungen (für alle i) als ein sehr großes lineares Gleichungssystem in den Unbestimmten x_i auffassen.

Diese Änderung eröffnet noch die folgende Interpretationsmöglichkeit für die Zahlen x_i . Es ist klar, dass das obige Gleichungssystem (wenn es überhaupt eine Lösung hat, für die nicht alle x_i gleich Null sind) keine eindeutige Lösung haben kann; denn wenn wir eine Lösung haben, erhalten wir eine neue Lösung, indem wir jedes x_i mit derselben Zahl multiplizieren. Für die resultierende Sortierung nach Relevanz tut das aber nichts zu Sache (jedenfalls, wenn wir nur mit positiven Zahlen multiplizieren). Wenn wir optimistisch sind, könnten wir versuchen, die Zahlen x_i so zu suchen, dass sie alle zwischen 0 und 1 liegen und dass die Summe aller x_i gleich 1 ist. Damit stellt man sicher, dass die Skala nicht ausufert und man die Werte möglichst konkret festnagelt. So kann man zum Beispiel besser Vergleiche zwischen solchen Berechnungen für verschiedene »Netzwerke« anstellen.

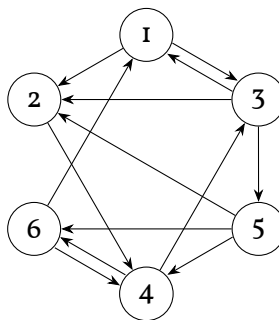
Mit dieser Konvention kann man die Zahl x_i auch als *Wahrscheinlichkeit* interpretieren. Dazu stellen wir uns eine Internet-Surfer*in vor, die auf jeder Seite zufällig irgendeinen der Links auf andere Seiten aufruft. Dann ist x_i die Wahrscheinlichkeit, dass sie sich gerade auf der i -ten Seite aufhält. Zum Beispiel würde $x_2 = 0,013 = 1,3\%$ bedeuten, dass die Wahrscheinlichkeit, sich dabei gerade auf der zweiten Seite zu befinden, $1,3\%$ ist.

Eine weitere Verbesserung des Verfahrens ist die folgende. Mit der obigen Formel könnte selbst eine relativ unwichtige Seite die Bewertung vieler anderer Seiten dadurch beeinflussen, dass sie sehr viele Links auf andere Seiten einbaut. Es ist daher vernünftig, ein Stimmgewicht einzuführen, das bewirkt, dass eine Seite ihre Relevanz sozusagen auf alle Seiten aufteilt, auf die sie verlinkt. Wir ersetzen daher die Gleichung oben durch

$$x_i = \sum_{j \in L_i} \frac{1}{n_j} x_j,$$

wobei n_j die Anzahl der ausgehenden Links von Seite j ist. (Weil $j \in L_i$ für alle j , die einen Beitrag zu der Summe liefern, verlinkt Seite j mindestens auf Seite i , also ist $n_j > 0$ und wir können durch diese Zahl teilen.)

BEISPIEL 2.8. Betrachten wir als Beispiel das in der Abbildung dargestellte Mini-Internet; Links sind dort als Pfeile dargestellt.



Die Daten, die wir brauchen, um das gesuchte Gleichungssystem aufzustellen, sind in der folgenden Tabelle gesammelt:

Seite	Verlinkt von	Anzahl ausgehende Links
1	3, 6	2
2	1, 3, 5	1
3	1, 4	3
4	2, 5, 6	2
5	3	3
6	4, 5	2

Aus diesen Daten stellen wir nach dem obigen Rezept das folgende lineare Gleichungssystem auf. Zum Beispiel erhalten wir folgendermaßen die erste Gleichung: Auf die erste Seite

verlinken die Seiten 3 und 6, also sind nur die Koeffizienten von x_3 und x_6 auf der rechten Seite der Gleichung $\neq 0$. Für den Koeffizienten von x_3 erhalten wir $\frac{1}{3}$, weil die dritte Seite 3 ausgehende Links hat. Die sechste Seite hat 2 ausgehende Links, also ist der Koeffizient von x_6 gleich $\frac{1}{2}$.

$$\begin{array}{rcll}
 x_1 & = & & \frac{1}{3}x_3 & + & \frac{1}{2}x_6 \\
 x_2 & = & \frac{1}{2}x_1 & + & \frac{1}{3}x_3 & + & \frac{1}{3}x_5 \\
 x_3 & = & \frac{1}{2}x_1 & & + & \frac{1}{2}x_4 & \\
 x_4 & = & & x_2 & & + & \frac{1}{3}x_5 & + & \frac{1}{2}x_6 \\
 x_5 & = & & \frac{1}{3}x_3 & & & & \\
 x_6 & = & & & \frac{1}{2}x_4 & + & \frac{1}{3}x_5 &
 \end{array}$$

An dieser Stelle können wir schon eine Besonderheit feststellen, die bei der weiteren Betrachtung eine große Rolle spielen wird. Wenn man auf der rechten Seite alle Koeffizienten einer Spalte aufsummiert, dann ist die Summe immer gleich 1. Überzeugen Sie sich, dass das in der Tat eine direkte Konsequenz davon ist, wie wir das Gleichungssystem aufbauen. (Es gibt allerdings die Möglichkeit, dass als »Sonderfall« Spalten auftreten, in denen alle Koeffizienten = 0 sind. Wann wäre das der Fall?)

Eine Lösung dieses Gleichungssystems ist gegeben durch (32, 36, 45, 58, 15, 34). Jedenfalls lässt sich leicht nachprüfen, dass diese Zahlen wirklich alle 6 Gleichungen erfüllen. Es ist auch nicht »schwierig«, diese Lösung (zum Beispiel durch Einsetzen und/oder Elimination von Variablen) zu finden. Allerdings macht es sich bei einem Gleichungssystem dieser Größe bezahlt, einen systematischen Ansatz zu wählen, wie wir ihn demnächst kennenlernen werden (den Gauß-Algorithmus, Abschnitt 5.2).

Wir hatten auch gesagt, dass wir als Summe der x_i gerne 1 erhalten möchten. Das ist natürlich mit der obigen Lösung nicht erfüllt. Wir können die Lösung aber skalieren: Wenn wir alle x_i mit derselben Zahl multiplizieren, oder durch dieselbe Zahl teilen, entsteht aus einer Lösung des Gleichungssystems eine neue Lösung. In diesem Fall teilen wir alle Einträge durch die Gesamtsumme 220. Wir erhalten die Lösung $(\frac{32}{220}, \frac{36}{220}, \frac{45}{220}, \frac{58}{220}, \frac{15}{220}, \frac{34}{220})$, wobei wir ausnahmsweise nicht gekürzt haben, weil es so ein bisschen übersichtlicher erscheint. Für diese Lösung ist die Summe aller x_i gleich $\frac{220}{220} = 1$. Außerdem liegen alle x_i zwischen 0 und 1. Man kann auch zeigen: Dies ist die *einzige* Lösung, so dass sich die Einträge zu 1 summieren. Das ist die Lösung des Page-rank-Problems in diesem Fall. Die wichtigste Seite ist Seite 4, mit Abstand am unwichtigsten (aus Sicht des Algorithmus) ist Seite 5. $\diamond$

Auch im allgemeinen Fall handelt es sich hier nach wie vor um ein lineares Gleichungssystem (das wir über den rationalen oder über den reellen Zahlen betrachten können). Dies ist noch nicht ganz das endgültige Gleichungssystem des Page-rank-Algorithmus, aber wir verschieben die Diskussion der weiteren Verbesserungen auf Ergänzung 5.61, wo wir die zur Verfügung stehende Sprache schon etwas erweitert haben.

Quellen zum Page-rank-Algorithmus:

K. Bryan, T. Leise, [The 25 000 000 000 eigenvector. The Linear Algebra behind Google^a](#), SIAM Review, Vol. **48**, No. 3 (2006), 569–581

D. Austin, [How Google Finds Your Needle in the Web's Haystack^b](#) (dort finden Sie auch Verweise auf weitere Literatur zum Page-rank-Algorithmus und allgemeiner zu Suchmaschinen-Algorithmen)

Wikipedia: [Deutsch^c](#), [Englisch^d](#) (dort finden Sie auch eine Diskussion der Defizite des Page-rank-Algorithmus und (wenige) Informationen, wie Google ihn weiterentwickelt und inzwischen größtenteils ersetzt hat)

^a<https://www.rose-hulman.edu/%7Ebryan/googleFinalVersionFixed.pdf>

^b<http://www.ams.org/publicoutreach/feature-column/fcarc-pagerank>

^c<https://de.wikipedia.org/wiki/PageRank>

^d<https://en.wikipedia.org/wiki/PageRank>

Wir wollen es hier dabei belassen, zwei Fragen zu formulieren:

- (1) Gibt es eine Lösung für dieses Gleichungssystem, d.h. gibt es Zahlen x_i , die die Gleichungen erfüllen? Ist die Lösung (bis auf Skalieren) eindeutig bestimmt? (Denn wenn es mehrere Lösungen, also mehrere mögliche Rankings der Suchergebnisse gibt, ständen wir vor dem neuen Problem, wie wir uns zwischen diesen entscheiden könnten.)
- (2) Wenn es eine eindeutige Lösung gibt, wie kann man sie (angesichts der in der Praxis riesigen Zahl an Gleichungen und Variablen) berechnen?

Wir werden in den Ergänzungen 5.61, 7.66 und 10.24 auf diese Fragen zurückkommen und sehen, dass Frage (1) nach Durchführung der angekündigten Verbesserungen eine positive Antwort hat (Korollar 7.69), und es für Teil (2) Methoden gibt, die die spezielle Form dieses Gleichungssystems ausnutzen, siehe Ergänzung 10.24. □ Frage 2.7

FRAGE 2.9. Wie kann man große Datenmengen analysieren, effizient abspeichern und komprimieren? Diese Frage ist recht allgemein gehalten. Wir werden an mehreren Stellen darauf zurückkommen; teilweise erst in der Linearen Algebra 2.

Konkrete Anwendungen sind

- die (verlustfreie oder verlustbehaftete) Kompression von Bilddaten: Wie speichern Sie eine Bilddatei möglichst effizient ab, ohne Information zu verlieren – statt den Farbwert jedes einzelnen Bildpunktes abzuspeichern, möchte man ausnutzen, dass in typischen Bildern benachbarte Bildpunkte oft ähnliche Farbwerte haben. Was hat man für Möglichkeiten, den benötigten Speicherplatz noch deutlich stärker zu reduzieren, wenn man (kleine) Qualitätseinbußen in Kauf nimmt? Siehe Ergänzung 5.63.
- Wie kann man die Qualität einer Datensammlung (beispielsweise wieder einer Bilddatei) verbessern, indem man Fehler (»noise«) sozusagen ausbügelt?

□ Frage 2.9

Wir werden nach und nach mit den Begrifflichkeiten, die zum jeweiligen Zeitpunkt zur Verfügung stehen, weitere Fragen entdecken/sehen, die wir dann erst später beantworten können. Zum Beispiel:

- Es gibt keine Divisionsalgebra, die die reellen Zahlen als Teilkörper enthält und über $\mathbb{R}$ die Vektorraumdimension 3 hat. Siehe Ergänzungen 4.9, 6.65 für die Erläuterung der Frage und 10.20 für die Lösung.

Mit dieser Frage hat sich R. Hamilton um 1840 beschäftigt, es handelte sich damals um ein aktuelles Forschungsproblem, das ihn 1843 zur »(Er-)findung« der **Quaternionen**⁶ führte. Diese spielen auch in der heutigen Mathematik (und Physik, und zum Beispiel auch in der Computergeometrie) eine Rolle. Es handelt sich dabei um einen Zahlbereich, der die reellen und die komplexen Zahlen enthält, aber auch noch zusätzliche Elemente. Es gelten dort die üblichen Rechenregeln bis auf das Kommutativgesetz der Multiplikation. Siehe Ergänzungen 4.11, 5.64, 9.44.

”

Every morning, on my coming down to breakfast, you used to ask me: “Well, Papa, can you multiply triplets?” Where to I was always obliged to reply, with a sad shake of the head: “No, I can only add and subtract them.”

R. Hamilton in einem **Brief an seinen Sohn**^a aus dem Jahr 1865, in dem er auch von seiner Erfindung der Quaternionen berichtet.

^a<https://www.maths.tcd.ie/pub/HistMath/People/Hamilton/Letters/BroomeBridge.html>

Als zwei weitere Beispiele für Gebiete, aus denen die Lineare Algebra nicht wegzudenken ist, seien hier genannt:

- die *Kodierungstheorie*, die sich damit befasst, wie man Informationen so über einen Kommunikationskanal (wie eine Funkverbindung oder ein »Internetkabel«) übertragen kann, dass sich Übertragungsfehler durch den Empfänger feststellen und bestenfalls automatisch korrigieren lassen. Dies ist eine sehr anwendungsnahe Theorie mit großen Überschneidungen mit der Informatik und den Ingenieurwissenschaften. Siehe Kapitel 12,
- die *Graphentheorie*, in der Konfigurationen von Punkten (»Knoten«), die durch Strecken (»Kanten«) miteinander verbunden sind, untersucht werden. Beispiele für *gerichtete* Graphen, in denen jede Kante mit einer Richtung versehen ist, sind die Darstellungen eines Netzwerks im Abschnitt über den Page-rank-Algorithmus (Frage 2.7). Es ist vielleicht ein bisschen überraschend, dass man Methoden der Linearen Algebra (wie zum Beispiel die Theorie der Eigenwerte, deren Anfänge wir zum Ende der Linearen Algebra 1 hin kennenlernen werden und die uns dann auch noch im kommenden Semester beschäftigen wird) auf das Studium von solchen Graphen anwenden kann, sie erweisen sich aber oft als sehr nützlich. Siehe Kapitel 13.
- Die lineare Algebra ist auch eng mit der *analytischen Geometrie* verwoben. Siehe Kapitel 11 für einige Bemerkungen dazu. In der Folgevorlesung Lineare Algebra 2 werden wir dieses Thema noch weiter vertiefen können.
- Der Begriff der *Symmetrie* spielt in vielen Bereichen der Mathematik eine Rolle und Symmetrien werden meist mithilfe des Begriffs der *Gruppe* beschrieben und untersucht, den wir in Kapitel 8 kennenlernen werden.

⁶<https://de.wikipedia.org/wiki/Quaternion>

MathOverflow: Why Linear Algebra is fun! (or?)^a

^a<https://mathoverflow.net/questions/33911/why-linear-algebra-is-funor>

Weitere Beispiele finden Sie in den Büchern [LM], [Ba], in der Einleitung von [Fi] (und den dort gegebenen Literaturverweisen), [Ma] (Anwendungen von linearer Algebra auf mathematische Probleme), [RS]. Siehe auch Anhang D.

2.4. Alternativer Aufbau der Vorlesung *

Der Aufbau der Vorlesung, den ich gewählt habe, führt theoretische Begriffe erst dann ein, wenn sie benötigt werden und (hoffentlich) durch vorher aufgetretene Fragen motiviert werden können. Es wird trotzdem viele Definitionen geben, die nicht auf der Hand liegen und deren Nutzen nicht von vorneherein klar ist. Ich möchte aber versuchen, soweit es geht, den falschen Eindruck zu vermeiden, dass hier eine Theorie um ihrer selbst willen entwickelt wird, die losgelöst ist von den ursprünglichen Fragestellungen. Eines der Bücher, das einen ähnlichen Zugang wählt, ist das Buch [Lo] von Lorenz.

Alternativ kann man eine Darstellung wählen, die für diejenigen, die die Theorie schon kennen, möglichst systematisch ist. Damit kann man an (wenigen) Stellen etwas effizienter arbeiten, muss aber noch öfter Begriffe präsentieren, deren Nutzen sich erst wesentlich später herausstellt. Es gibt auch für diesen Aufbau gute Argumente, und er ist ebenfalls gebräuchlich; siehe zum Beispiel das Buch [Bo] von Bosch.

Dann würde man die Themen der Vorlesung zum Beispiel in der folgenden Reihenfolge abarbeiten: Gruppen, Körper, Vektorräume, Basen, Lineare Abbildungen, Matrizen und lineare Gleichungssysteme, Determinanten, Eigenwerte und Eigenvektoren.

(Und wenn man wollte, könnte man den Begriff der Matrix und die »Anwendung« auf die Lösung von linearen Gleichungssystemen vollständig vermeiden, was den Aufbau der Theorie angeht. Das würde dann aber, wie ich denke, doch ein falsches Bild vermitteln.)

2.5. Lineare Gleichungssysteme

Zum Schluss dieses einführenden Kapitels möchte ich wenigstens ein kleines bisschen lineare Algebra besprechen, nämlich die einfachsten Fälle linearer Gleichungssysteme. Dieses Material ist für Sie wahrscheinlich nicht neu (höchstens die systematische Form der Betrachtung). Es soll ein kleiner Ausblick auf das erste große Ziel der Vorlesung sein – die Lösung linearer Gleichungssysteme mit dem Gauß-Algorithmus (5.2). Auch wenn dann im späteren Verlauf die linearen Gleichungssysteme weniger sichtbar sind, bleibt der Gauß-Algorithmus ein wichtiges Mittel für die allermeisten Aufgaben, in denen etwas berechnet werden soll.

Eine *Gleichung* ist ein Ausdruck, in dem zwei Terme gleichgesetzt werden, in denen eine oder mehrere *Unbestimmte* (die wir unten $X, Y, X_1, X_2, \dots$ nennen) vorkommen, sowie Zahlen (»Koeffizienten«, unten mit kleinen Buchstaben bezeichnet) aus einem Zahlbereich K (zum Beispiel: die rationale Zahlen oder die reelle Zahlen). Ein *Gleichungssystem* besteht einfach aus mehreren Gleichungen.

Wir suchen die *Lösungsmenge* der Gleichung oder der Gleichungssystems, also alle Möglichkeiten, für alle Unbestimmten Elemente aus K so einzusetzen, dass alle Gleichungen erfüllt sind. Zum Beispiel ist für die Gleichung $x^2 = 2$ die Lösungsmenge die leere Menge, wenn wir die Gleichung über den rationalen Zahlen betrachten, und die Menge $\{\sqrt{2}, -\sqrt{2}\}$, wenn wir die Gleichung über den reellen Zahlen betrachten.

2.5.1. Lineares Gleichungssystem mit 1 Gleichung und 1 Unbestimmten. Wir beginnen ganz einfach und betrachten die Gleichung

$$aX = b$$

für Elemente $a, b \in K$. Die Lösungsmenge $\mathbb{L}$ ist dann eine Teilmenge von K , genauer die Teilmenge $\{x \in K; ax = b\}$.

Fall 1: $a \neq 0$. Dann ist die Lösungsmenge $\mathbb{L} = \{\frac{b}{a}\}$.

Fall 2: $a = 0, b \neq 0$. Dann ist die Lösungsmenge leer: $\mathbb{L} = \emptyset$.

Fall 3: $a = b = 0$. Dann ist jedes x eine Lösung, $\mathbb{L} = K$.

2.5.2. Lineares Gleichungssystem mit 1 Gleichung und n Unbestimmten. Nun betrachten wir eine Gleichung der Form

$$a_1X_1 + a_2X_2 + \dots + a_nX_n = b,$$

in der mehrere Unbestimmt auftreten (aber keine höheren Potenzen der X_i – es handelt sich um eine *lineare* Gleichung). Die Lösungsmenge ist dann eine Menge von » n -Tupeln«, also Listen der Länge n , von Elementen in K . Wir bezeichnen die Menge aller n -Tupel $(x_1, \dots, x_n)$, $x_i \in K$, mit K^n (vergleiche 3.10.1).

Fall 1: Es gibt ein i mit $a_i \neq 0$. Dann erhalten wir alle Lösungen der Gleichung folgendermaßen. Wir wählen beliebige Elemente $x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n \in K$ und definieren

$$x_i = \frac{1}{a_i}(b - a_1x_1 - \dots - a_{i-1}x_{i-1} - a_{i+1}x_{i+1} - \dots - a_nx_n).$$

Wir schreiben die Lösungsmenge in der Form

$$\mathbb{L} = \left\{ \left(x_1, \dots, x_{i-1}, \frac{1}{a_i}(b - a_1x_1 - \dots - a_{i-1}x_{i-1} - a_{i+1}x_{i+1} - \dots - a_nx_n), x_{i+1}, \dots, x_n \right); \right. \\ \left. x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n \in K \right\}.$$

Fall 2: Alle Koeffizienten a_i sind gleich 0.

Fall 2a: $b = 0$. In diesem Fall lösen alle $x_1, \dots, x_n \in K$ die Gleichung. Wir schreiben die Lösungsmenge als

$$\mathbb{L} = K^n = \{(x_1, \dots, x_n); x_1, \dots, x_n \in K\}.$$

Fall 2b: $b \neq 0$. In diesem Fall ist die Gleichung $0 = b$ nicht erfüllt, also ist die Lösungsmenge die leere Menge.

2.5.3. Lineares Gleichungssystem mit 2 Gleichungen und 2 Unbestimmten. Nun betrachten wir ein System von zwei Gleichungen und suchen alle Paare (x, y) von Elementen aus K , die beide Gleichungen

$$aX + bY = e$$

$$cX + dY = f$$

gleichzeitig erfüllen. Die Lösungsmenge ist dann eine Teilmenge von K^2 , der Menge aller Paare (x, y) von Elementen in K .

Fall 1: $a \neq 0$. Wir ziehen das $\frac{c}{a}$ -fache der ersten Gleichung von der zweiten Gleichung ab, multiplizieren danach die zweite Gleichung mit a und erhalten das lineare Gleichungssystem

$$\begin{aligned} aX + bY &= e \\ (ad - bc)Y &= af - ec. \end{aligned}$$

(Machen Sie sich klar, dass sich die Lösungsmenge nicht ändert: Das neue System hat die gleiche Lösungsmenge, wie das ursprüngliche System. Vergleiche Lemma 5.12.)

Fall 1a: $a \neq 0, ad - bc \neq 0$. Wir können dann die zweite Gleichung durch $ad - bc$ teilen, dies in die erste Gleichung einsetzen, und die erste Gleichung durch a dividieren. Wir erhalten:

$$\begin{aligned} X &= \frac{e}{a} - \frac{b}{a} \cdot \frac{af - ec}{ad - bc} = \frac{e(ad - bc) - b(af - ec)}{a(ad - bc)} = \frac{ed - bf}{ad - bc} \\ Y &= \frac{af - ec}{ad - bc}. \end{aligned}$$

Fall 1b: $a \neq 0, ad - bc = 0, af - ec = 0$. Dann bekommen wir

$$\begin{aligned} aX + bY &= e \\ 0 &= 0, \end{aligned}$$

d.h. die zweite Gleichung fällt weg. Für beliebiges y erhalten wir die Lösung $(\frac{e-by}{a}, y)$. (Vergleiche den Fall einer Gleichung mit 2 Unbestimmten.)

Fall 1c: $a \neq 0, ad - bc = 0, af - ec \neq 0$. Dann ist die zweite Gleichung niemals erfüllt, egal, was wir für X und Y einsetzen – die Lösungsmenge ist leer.

Fall 2: $a = 0, c \neq 0$. Diesen Fall können wir analog zu Fall 1 behandeln. Oder wir können ihn durch Vertauschen der beiden Gleichungen, also durch Umbenennen der Koeffizienten) auf Fall 1 zurückführen. Weil $ad - bc = -(cb - ad)$ gilt, sind die Bedingungen $ad - bc \neq 0$ und $cb - ad = 0$ äquivalent. Wir erhalten also wieder genau dasselbe Kriterium für die eindeutige Lösbarkeit des gegebenen linearen Gleichungssystems. Weil sich bei den Ausdrücken $ed - bf$ und $af - ec$ bei Vertauschen der beiden Zeilen auch gerade das Vorzeichen ändert, gilt in diesem Fall auch wieder die gleiche Formel für die eindeutige Lösung.

Fall 3: $a = c = 0$. In diesem Fall spielt die Wahl von X überhaupt keine Rolle. Die Lösungsmenge des linearen Gleichungssystems ist die Schnittmenge der Lösungsmengen von $bY = e$ und $dY = f$, die jeweils wie oben besprochen (1 Unbestimmte, 1 Gleichung) bestimmt werden können. Im Ergebnis kann die Lösungsmenge leer sein, oder Y ist eindeutig bestimmt (und X ist frei wählbar), oder Y (und X) sind beide frei wählbar.

Wir können unser Ergebnis in der folgenden Form zusammenfassen:

- (1) Das lineare Gleichungssystem ist eindeutig lösbar genau dann, wenn $ad - bc \neq 0$. Dann ist

$$\mathbb{L} = \left\{ \left(\frac{ed - bf}{ad - bc}, \frac{af - ec}{ad - bc} \right) \right\}.$$

- (2) Wenn $ad - bc = 0$, aber nicht alle Koeffizienten a, b, c, d verschwinden:
 (a) Gilt $af - ec = ed - bf = 0$, so gibt es so viele Lösungen wie die Anzahl der Elemente des Zahlbereichs K (also unendlich viele, wenn K unendlich viele Elemente hat).

- (b) Gilt $af - ec \neq 0$ oder $ed - bf \neq 0$, so gibt es keine Lösungen: die Lösungsmenge ist leer.
- (3) Gilt $a = b = c = d = 0$, dann können die folgenden Fälle auftreten:
 - (a) Ist $e = f = 0$, so sind alle Paare (x, y) Lösungen des linearen Gleichungssystems.
 - (b) Ist $e \neq 0$ oder $f \neq 0$, so ist die Lösungsmenge leer.

Wir sehen auch: Wir wollen zwar auch größere Gleichungssysteme (mehr Unbestimmte, mehr Gleichungen) verstehen, aber nicht in dieser Art und Weise weitermachen. Die Fallunterscheidungen dafür würden viel zu unübersichtlich werden. Wir benötigen einen systematischeren Ansatz, der besser die zugrundeliegende Struktur des Problems ausnutzt und sichtbar macht.

Grundlagen

3.1. Worum geht es eigentlich?

In der Mathematik möchten wir (mathematische) Strukturen verstehen und Erkenntnisse darüber als Aussagen formulieren, an deren Wahrheit kein Zweifel bestehen kann. Etwas konkreter sind die drei wesentlichen Bestandteile des Aufbaus der Mathematik

- (1) Definition,
- (2) Satz und
- (3) Beweis.

In einer *Definition* werden zu verwendende Begriffe erklärt. Der *Satz* formuliert eine Aussage über die definierten Begriffe, und im *Beweis* wird eine vollständige und korrekte Argumentation angegeben, die belegt, dass die Aussage des Satzes wahr sein muss.

Das klingt einerseits ziemlich abstrakt, andererseits ist es nicht sehr präzise. Für den Moment ist mir das wichtigste, dass Sie mitnehmen, dass es uns vorrangig darum gehen wird, »strukturelle Aussagen« zu machen (und weniger darum, konkrete Rechnungen durchzuführen) und die Korrektheit dieser Aussagen zu beweisen. Etwas ausführlicher wird die Frage *Was ist ein Beweis?* im nächsten Abschnitt behandelt, den Sie aber auch erstmal überspringen können.

In der Praxis ist meist ziemlich klar, ob eine Argumentation einen Beweis darstellt oder nicht, und einen Beweis für eine mathematische Aussage zu *finden* ist ein ganz anderes Problem als die formale Frage, was ein Beweis eigentlich ist. Uns interessiert hier nicht so sehr die formale, sondern mehr die »praktische« Seite dieser Medaille (die ja immer noch abstrakt genug ist ...).

Man muss kein Genie sein, um Beweise zu verstehen und selbst zu finden. Von selbst kann das aber auch so gut wie niemand. Das wichtigste ist die Übung: Beschäftigen Sie sich viel mit Mathematik, und nehmen Sie sich genügend Zeit, um über Aufgaben nachzudenken und sich mit anderen darüber zu unterhalten.

Siehe auch Anhang D.3.

3.2. Was ist ein Beweis? *

Das heute allgemein anerkannte Grundprinzip der Mathematik ist die sogenannte *axiomatische Methode*:

- ausgehend von einer (kleinen) Liste von Annahmen (»Axiomen«)
- unter Ausnutzung »offensichtlich« korrekter Schlussweisen
- wahre Aussagen zu beweisen.

Die Axiome setzen also sozusagen den Rahmen für die Theorie. Bei der Auswahl der Axiome muss man eine Balance finden, weil man genügend viele/starke Axiome benötigt, um damit

dann auch etwas beweisen zu können, aber andererseits nur solche Axiome verwenden möchte, die »offensichtlich wahr« sind, d.h., die sich in der Anschauung widerspiegeln, die man von der mathematischen Theorie hat, die man aufbauen möchte.

Je nachdem, worüber man mathematisch sprechen möchte, kann man unterschiedliche Axiomensysteme benutzen. Drei der bekanntesten Systeme sind

- (1) die [Peano-Axiome](#)¹, die die Eigenschaften der natürlichen Zahlen formalisieren,
- (2) [Euklids Axiome der Geometrie](#)², die die geometrischen Eigenschaften der Ebene (und des dreidimensionalen Raums) formalisieren. Auch wenn Euklids Ausführungen nicht den heutigen formalen Ansprüchen genügen (eine modernere Variante ist [Hilberts Axiomensystem](#)³), wird das Prinzip der axiomatischen Methode bei Euklid sehr gut sichtbar.
- (3) Die [Zermelo-Fraenkelsche Mengenlehre](#)⁴ (kurz ZFC), die nach E. Zermelo und A. Fraenkel benannt ist, ist ein Axiomensystem für den Begriff der Menge. Damit lässt sich praktisch die gesamte heutige Mathematik beschreiben. Das C in ZFC steht für die Hinzunahme des *Auswahlaxioms* (englisch *axiom of choice*) zu den anderen Axiomen, siehe Anhang B.1.

ERGÄNZUNG 3.1 (Das Parallelenaxiom). Eines der Axiome der euklidischen Geometrie ist das sogenannte Parallelenaxiom, das besagt, dass zu jeder Gerade und jedem Punkt, der nicht auf dieser Gerade liegt, genau eine Gerade existiert, die durch den Punkt geht und zu der ersten Gerade parallel ist. Das ist sicherlich eine Eigenschaft, die unserer Anschauung entspricht.

Es war lange Zeit nicht bekannt, ob das Parallelenaxiom schon aus den anderen Axiomen folgt (so dass man es nicht als Axiom hätte hinzufügen müssen). Die Entwicklung der [nicht-euklidischen Geometrie](#)⁵ hat aber gezeigt, dass dies nicht der Fall ist: Das Parallelenaxiom ist unabhängig von den anderen Axiomen.

Eine Empfehlung für ein Buch (in Romanform mit mathematischen Einschüben) zu diesem Thema: G. Suri, H. S. Bal, *Eine gewisse Ungewissheit oder Der Zauber der Mathematik*, DuMont 2008. (oder das englische Original: *A certain ambiguity. A Mathematical Novel*, Princeton Univ. Press, 2nd ed., 2010).

□ Ergänzung 3.1

Für den formal korrekten Aufbau der Theorie ist es sinnvoll, an dieser Stelle auch zu formalisieren, was man unter einem Beweis, also unter den oben angesprochenen »offensichtlich« korrekten Schlussweisen versteht. Der üblicherweise gewählte Ansatz ist die sogenannte [Prädikatenlogik erster Stufe](#)⁶.

Literaturverweise (wenn Sie unbedingt mehr wissen müssen):

Ein bisschen mehr sagt Wolfram Pohlers (dessen Forschungsgebiet die mathematische Logik ist) in seinem [Skript zur Linearen Algebra](#)^a (Kapitel 0).

Auch in dem (englischen) Buch [Hu] von Hungerford finden Sie etwas mehr Informationen zu diesem Thema.

^ahttps://www.uni-muenster.de/imperia/md/content/logik/Skripte/pohlers.lineare_algebra.pdf

¹<https://de.wikipedia.org/wiki/Peano-Axiome>

²https://de.wikipedia.org/wiki/Euklidische_Geometrie

³https://de.wikipedia.org/wiki/Hilberts_Axiomensystem_der_euklidischen_Geometrie

⁴<https://de.wikipedia.org/wiki/Zermelo-Fraenkel-Mengenlehre>

⁵https://de.wikipedia.org/wiki/Nichteuklidische_Geometrie

⁶https://de.wikipedia.org/wiki/Prädikatenlogik_erster_Stufe

In dieser Vorlesung werden wir aber darauf verzichten, das Axiomensystem auszuformulieren (das oben erwähnte ZFC ist eine Möglichkeit, die für die Vorlesungen Lineare Algebra 1+2 und die meisten (wahrscheinlich: alle) anderen Vorlesungen, die Sie in Ihrem Studium hören werden, ausreicht). Stattdessen werden wir mit einem »naiven« Mengenbegriff arbeiten. Sie müssen sich sozusagen darauf verlassen, dass alle Operationen, die wir mit Mengen ausführen, auch innerhalb von ZFC erlaubt sind (oder sich eigenständig die Mühe machen, das nachzuprüfen; es ist aber nicht unbedingt eine gute Idee, das zu Beginn des Studiums zu tun). Wir werden auch darauf verzichten, die erwähnte Prädikatenlogik ganz formal einzuführen, aber die nötigen Begriffe in Abschnitt 3.7 in der Form bereitstellen, wie wir sie verwenden werden. (Und fast alle anderen Lehrbücher machen es genauso.)

”

If logic is the hygiene of the mathematician, it is not his source of food; the great problems furnish the daily bread on which he thrives.

André Weil^a, *The future of mathematics*,
American Mathematical Monthly, May 1950

^ahttps://de.wikipedia.org/wiki/Andr%C3%A9_Weil

Auch wenn man sich an dieser Stelle mehr Zeit nehmen würde, wären die Beweise im Hauptteil des Textes genauso geschrieben, wie hier. Dieser Formalismus (hat zwar im Gesamtbild seine Berechtigung, aber) hilft nicht beim Verständnis der Linearen Algebra.

Außerdem: Irgendwo muss man anfangen – selbst bei maximaler Formalisierung muss man *auch* auf einer Meta-Ebene arbeiten, die man nicht in gleichem Maße formalisieren kann. Man kann eben vor der allerersten Definition nicht schon definiert haben, was man unter einer *Definition* versteht. Darüberhinaus gibt es auch tiefliegendere Einschränkungen an die vollständige Formalisierung der mathematischen Theorie (Stichwort: [Gödelscher Unvollständigkeitssatz](#)⁷). Diese Phänomene werden in der mathematischen Logik genauer untersucht. Dort analysiert man auch die »Beweiskraft« verschiedener formaler Systeme.

Das Grundprinzip, dass alle Aussagen auf die zu Beginn vorausgesetzten Axiome zurückgeführt werden können und dass jede Behauptung mit einer vollständigen und fehlerfreien Argumentation begründet werden muss, ist aber auch in diesem Text fundamental.

3.3. Beweise – einige Klassiker *

In diesem Abschnitt gehen wir einige ”klassische” Beweise durch. Einige der Ergebnisse sind Ihnen sicher bekannt, vielleicht sogar die meisten.

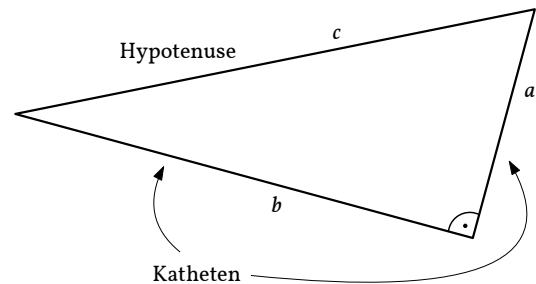
Die Beweise sind alle einfach in dem Sinne, dass sie nur Schulmathematik benötigen. Das heißt aber nicht, dass man leicht darauf kommen kann. Trotzdem ist es eine gute Idee, über die Behauptungen, die unten aufgestellt werden, ein bisschen nachzudenken, bevor man den Beweis liest. Vielleicht finden Sie einen eigenen Beweis? Und wenn nicht, dann ist es oft leichter, einen Beweis zu verstehen (und zu würdigen), wenn man sich vorher schon einmal selbst daran versucht hat. Für Ihre Vorlesung(en) empfehle ich dasselbe!

⁷https://de.wikipedia.org/wiki/Gödelscher_Unvollständigkeitssatz

3.3.1. Der Satz des Pythagoras. Sei Δ ein rechtwinkliges Dreieck. Wie üblich nennen wir die beiden Seiten von Δ , die an dem rechten Winkel anliegen, die *Katheten*, und die Seite, die dem rechten Winkel gegenüberliegt, die *Hypotenuse* des Dreiecks.

SATZ 3.2 (Pythagoras). Sei Δ ein rechtwinkliges Dreieck. Seien a und b die Längen der beiden Katheten, und sei c die Länge der Hypotenuse von Δ . Dann gilt

$$a^2 + b^2 = c^2.$$



BEMERKUNG 3.3. Auch wenn der Satz traditionell *Pythagoras*⁸ (ca. 550 v. Chr.) zugeschrieben wird, war er mit ziemlicher Sicherheit schon vorher bekannt. $\diamond$

BEWEIS. Es gibt viele Beweise für den Satz des Pythagoras. Der folgende ist wohl einer der einfachsten.

Wir betrachten die nebenstehende Figur, in der das gegebene rechtwinklige Dreieck mit Kathetenlängen a und b viermal vorkommt.

Wir können die Fläche des großen Quadrats auf zweierlei Arten berechnen: Die Seitenlänge des Quadrats ist $a + b$, also ist die Fläche

$$(a + b)^2 = a^2 + 2ab + b^2.$$

Andererseits ist die Fläche die Summe der Teilflächen: viermal die Fläche des Dreiecks, also zweimal die Fläche des Rechtecks mit Kantenlängen a und b ; und die Fläche des Quadrats in der Mitte, das Seitenlänge c hat. Damit ergibt sich für die Fläche

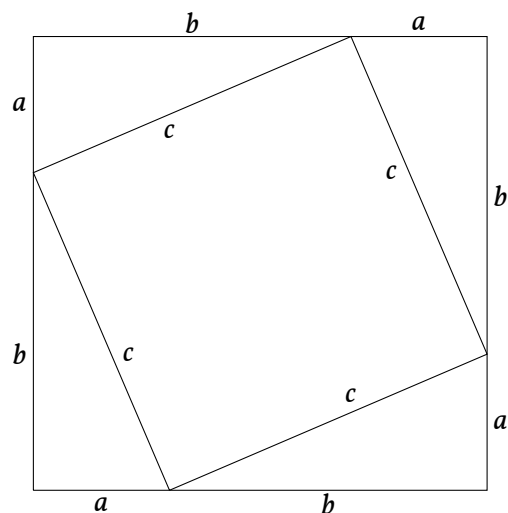
$$2ab + c^2.$$

Wir erhalten also

$$a^2 + 2ab + b^2 = 2ab + c^2,$$

und indem wir auf beiden Seiten $2ab$ abziehen, das gewünschte Ergebnis

$$a^2 + b^2 = c^2.$$



Eine Bemerkung zum Schluss: Prüfen Sie noch einmal genau, dass die geometrische Figur wirklich alle Eigenschaften hat, die wir im weiteren Beweis ausgenutzt haben. Ist der Bereich in der Mitte wirklich ein Quadrat, oder könnte es ein Viereck mit vier gleichlangen Seiten (eine *Raute*) sein, das aber kein Quadrat ist?

In [Ho] Kapitel 19 wird der Beweis noch ausführlicher aufgedröselt. $\square$

⁸<https://de.wikipedia.org/wiki/Pythagoras>

3.3.2. $\sqrt{2}$ ist keine rationale Zahl. Mit $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, 3, \dots\}$ bezeichnen wir wie üblich die Menge der ganzen Zahlen und mit $\mathbb{Q}$ die Menge der rationalen Zahlen (also aller Bruchzahlen von ganzen Zahlen).

SATZ 3.4. *Die Menge $\mathbb{Q}$ der rationalen Zahlen enthält kein Element, dessen Quadrat gleich 2 ist.*

Zum Beweis benutzen wir das folgende Lemma. Wir nennen dabei wie üblich eine ganze Zahl $a \in \mathbb{Z}$ gerade, wenn sie durch 2 teilbar ist (also wenn $b \in \mathbb{Z}$ existiert mit $a = 2b$) und andernfalls ungerade.

LEMMA 3.5. (1) *Eine ganze Zahl $a \in \mathbb{Z}$ ist genau dann gerade, wenn ihr Quadrat a^2 gerade ist.*

(2) *Ist a eine ungerade ganze Zahl, so ist $2a$ nicht durch 4 teilbar.*

BEWEIS. Versuchen Sie, bevor Sie den Beweis lesen, erst einmal, einen eigenen Beweis zu finden und aufzuschreiben!

zu (1). Ist a gerade, etwa $a = 2b$, dann ist $a^2 = 4b^2$ ebenfalls ein Vielfaches von 2, also gerade. Wenn a ungerade ist, dann ist $a - 1$ gerade, also existiert eine ganze Zahl b mit $a = 2b + 1$.

$$a^2 = (2b + 1)^2 = 4b^2 + 4b + 1 = 2(2b^2 + 2b) + 1$$

und dies ist eine ungerade Zahl, da sie sich um 1 von der geraden Zahl $2(2b^2 + 2b)$ unterscheidet.

zu (2). Da a nach Voraussetzung ungerade ist, können wir wie im vorherigen Teil $a = 2b + 1$ für eine geeignete ganze Zahl b schreiben. Dann ist $2a = 4b + 2$. Wäre diese Zahl durch 4 teilbar, so könnten wir auch $2a = 4c$ für eine ganze Zahl c schreiben, und würden dann erhalten, dass $4c = 4b + 2$, also $2 = 4(c - b)$. Dies ist offenbar nicht möglich. $\square$

BEWEIS DES SATZES. Angenommen, es gäbe eine Bruchzahl, deren Quadrat gleich 2 ist. Wir schreiben diese Zahl als gekürzten Bruch $\frac{a}{b}$, d.h. a und b sind ganze Zahlen, $b \neq 0$, und a und b haben keinen gemeinsamen Teiler > 1 . Und es gilt $(\frac{a}{b})^2 = 2$, also

$$a^2 = 2b^2.$$

Damit sehen wir zunächst, dass a^2 gerade ist, denn offenbar ist $2b^2$ gerade. Nach dem Lemma (Teil (1)) folgt, dass a eine gerade Zahl ist, wir können also $a = 2c$ für eine ganze Zahl c schreiben. Deshalb ist $a^2 = (2c)^2 = 4c^2$ sogar durch 4 teilbar.

Weil a gerade und der Bruch $\frac{a}{b}$ gekürzt ist, ist andererseits b ungerade. Teil (2) des Lemmas zeigt dann, dass $2b^2$ nicht durch 4 teilbar ist – ein Widerspruch zu der Gleichheit $a^2 = 2b^2$, denn wir haben gezeigt, dass die linke Seite durch 4 teilbar ist, die rechte jedoch nicht. $\square$

Dieser Beweis wird meist, wie wir es auch getan haben, als »Widerspruchsbeweis« aufgeschrieben: Man macht eine Annahme (dass eben doch eine rationale Zahl mit Quadrat gleich 2 existiert) und leitet daraus durch gültige logische Schlüsse eine Aussage ab, die offensichtlich falsch ist. Der einzige Teil dieser Ableitung, der falsch sein könnte, ist die ursprüngliche Annahme. Man hat also bewiesen, dass die Annahme falsch ist. Oft lassen sich Widerspruchsbeweise auch ersetzen durch »direkte« Beweise. Versuchen Sie einmal direkt zu zeigen, dass für jede rationale Zahl x gilt, dass $x^2 \neq 2$.

Häufig wird dieser Satz formuliert als: $\sqrt{2}$ ist nicht rational. Um das zu tun, muss man aber erst einmal wissen, dass in einem größeren Zahlbereich – in diesem Fall in den reellen Zahlen $\mathbb{R}$ – eine (positive) Zahl mit Quadrat 2 existiert; diese bezeichnet man dann mit $\sqrt{2}$. Das ist, wie Sie wissen, richtig. Es zu beweisen, ist aber nicht so einfach, schon deshalb, weil man zuerst präzisieren muss, was genau man eigentlich unter den reellen Zahlen versteht. Diese Aufgabe überlassen wir der Analysis-Vorlesung.

Schließlich nutze ich noch die Gelegenheit, auf das Büchlein [Go] hinzuweisen, wo (in Kapitel 3) noch genauer hinterfragt wird, ob unser Beweis wirklich vollständig und schlüssig ist. Die anderen Teile des Buchs sind auch empfehlenswert!

3.3.3. Es gibt unendlich viele Primzahlen. Für den nächsten Satz brauchen wir den Begriff der Primzahl: Eine *Primzahl* ist eine ganze Zahl > 1 , die keine Teiler außer 1 und sich selbst hat. Zum Beispiel sind 2, 3, 61 Primzahlen, 57 jedoch nicht, da 57 durch 3 teilbar ist. (Etwas ausführlicher befassen wir uns in Ergänzung 3.44 mit dem Begriff des Teilers.)

SATZ 3.6. *Es gibt unendlich viele Primzahlen.*

BEWEIS (NACH EUKLID). Wir zeigen, dass zu jeder endlichen Liste $p_1, \dots, p_n$ von Primzahlen noch mindestens eine weitere existiert.

Als Vorüberlegung zeigen wir, dass jede ganze Zahl $a > 1$ durch irgendeine Primzahl teilbar ist. In der Tat, ist a selbst eine Primzahl, so ist diese Behauptung richtig, dann a ist durch sich selbst teilbar. Ist a keine Primzahl, so besitzt a einen Teiler b mit $1 < b < a$. Ist b eine Primzahl, so sind wir fertig, denn b ist ein Teiler von a . Sonst finden wir einen Teiler von b , der noch kleiner ist; und jeder Teiler von b ist auch ein Teiler von a . Dieser Prozess muss irgendwann enden, weil wir in jedem Schritt zu einer kleineren Zahl übergehen. (Siehe auch Beispiel 3.42.)

Nun definieren wir

$$a = p_1 p_2 \cdots p_n + 1,$$

d.h. wir bilden das Produkt aller gegebenen Primzahlen und addieren dann noch 1. Nach der Vorüberlegung ist a teilbar durch eine Primzahl p . Andererseits ist a nicht durch irgendeine der Primzahlen $p_1, \dots, p_n$ teilbar, denn der Rest bei Division von a durch eine von diesen Zahlen ist 1.

Also ist p eine Primzahl, die von $p_1, p_2, \dots, p_n$ verschieden ist. □

Probieren Sie das Verfahren im Satz einmal aus, indem Sie mit der Liste 2, 3 von Primzahlen starten und nach und nach eine weitere konstruieren. Sie sehen schon im ersten Schritt, dass nicht alle Primzahlen der Größe nach auftreten, und nach wenigen Schritten, dass die Zahl a , die im Beweis konstruiert wird, nicht unbedingt selbst eine Primzahl ist.

Die Tatsache, dass jede natürliche Zahl > 1 als Produkt von Primzahlen geschrieben werden kann, und dass diese Zerlegung eindeutig ist bis auf die Reihenfolge der Faktoren (*eindeutige Primfaktorzerlegung in $\mathbb{Z}$*) ist übrigens etwas schwieriger zu zeigen. Siehe Satz 3.56.

3.4. Falsche Beweise erkennen *

Fehler passieren immer wieder und es ist wichtig diese zu erkennen: Eigene Fehler (bevor Sie Ihre Lösungen zu den Übungsaufgaben abgeben) und Fehler von anderen – Studierenden, Dozent*innen, in Büchern, Skripten usw. Siehe auch Abschnitt 1.6 in Pinks Skript [Pi].

Gleichzeitig kann man aus seinen eigenen Fehlern besonders viel lernen. Das soll zwar kein Aufruf sein, absichtlich Fehler zu machen; aber zu viel Angst sollte man davor auch nicht haben. Und: Wenn Sie in den Hausaufgaben Fehler gemacht haben, sollten Sie sicherstellen, dass Sie genau verstehen, was falsch war und wie es richtig gewesen wäre. Fragen Sie gegebenenfalls bei Ihrer Übungsgruppenleiter*in nach, falls nötig auch mehrfach!

”

He was not a very careful person as a mathematician. He made a lot of mistakes. But he made mistakes in a good direction. I tried to imitate him. But I've realized that it's very difficult to make good mistakes.

G. Shimura über Y. Taniyama^a

^ahttps://en.wikipedia.org/wiki/Yutaka_Taniyama

BEISPIEL 3.7. Analysieren Sie den folgenden Beweis:

Behauptung: $1 \leq 0$

Beweis: Sei $a = 1$, $b = 0$. Dann

$$1 \leq 0$$

$$a \leq b \quad | \cdot b$$

$$ab \leq b^2$$

$$0 \leq 0 \quad \checkmark$$

◇

”

An expert is someone who knows some of the worst mistakes that can be made in his subject, and how to avoid them.

W. Heisenberg, in: *Physics and Beyond*

Gefunden auf <http://math.furman.edu/~mwoodard/mqs/data.html>

BEISPIEL 3.8. Analysieren Sie die folgende Rechnung.

Aufgabe. Bestimmen Sie die Lösungsmenge des folgenden Gleichungssystems:

$$x + y = 0$$

$$x - y = 3$$

Wie viele Elemente hat die Lösungsmenge?

Lösung. Wir ersetzen die zweite Gleichung durch das Produkt der beiden Gleichungen und wenden die dritte binomische Formel an:

$$x + y = 0$$

$$x^2 - y^2 = 0$$

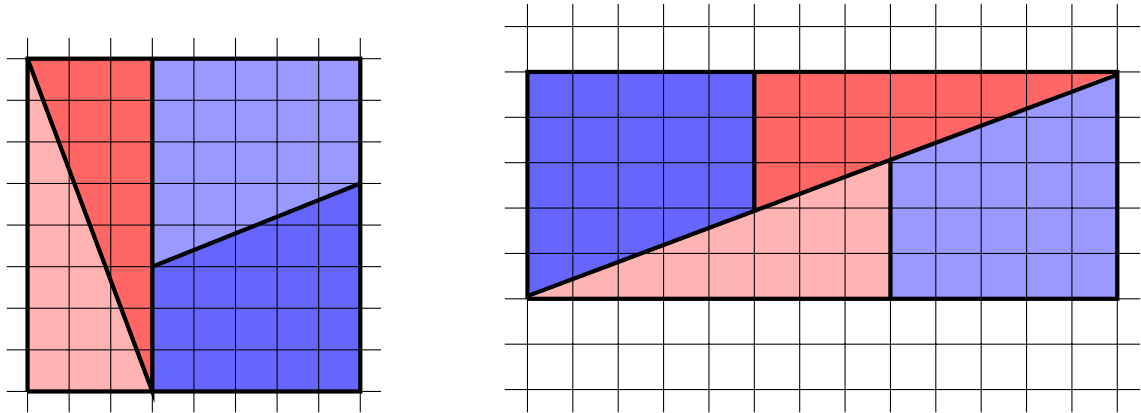
Die erste Gleichung können wir umschreiben als $x = -y$. Wenn das gilt, dann ist die zweite Gleichung automatisch erfüllt. Die Lösungsmenge ist also $\{(x, -x); x \in \mathbb{R}\}$. Sie hat unendlich viele Elemente. ◇

BEISPIEL 3.9. Es kommt nicht nur auf das Endergebnis an, sondern auch darauf, ein korrektes Argument verständlich darzulegen. Für die folgende Rechnung gibt es nicht die volle Punktzahl:

$$\frac{\cancel{16}}{\cancel{64}} = \frac{1}{4} \quad (\text{mit } 6 \text{ kürzen})$$



BEISPIEL 3.10. Betrachten Sie die folgenden Figuren ([Schachbrett-Paradoxon](#)⁹):



Wir sehen:

$$64 = 8 \cdot 8 = \text{hellrot} + \text{dunkelrot} + \text{hellblau} + \text{dunkelblau} = 13 \cdot 5 = 65$$

Was geht hier schief? Können Sie sich ähnliche Beispiele ausdenken?

Übrigens: Es ist kein Zufall, dass die Zahlen 3, 5, 8, 13, die als Kantenlängen auftreten, ein Teil der Fibonacci-Folge (siehe Frage 2.1) sind.



BEISPIEL 3.11. Dieses Beispiel benötigt die komplexen Zahlen, siehe 4.5. Analysieren Sie den folgenden Beweis, dass in den komplexen Zahlen $i = -i$ gilt.

$$i = \sqrt{i} = \sqrt{(-i) \cdot (-i)} = \sqrt{-i} \sqrt{-i} = (\sqrt{-i})^2 = -i.$$



Ein subtiler falscher Beweis, dass jedes Dreieck gleichschenkelig ist, wurde von Charles Dodgson erfunden (der unter dem Pseudonym Lewis Carroll die Geschichte von Alice im Wunderland geschrieben hat).

[Wikipedia \(Englisch\)](#)^a

^ahttps://en.wikipedia.org/wiki/Mathematical_fallacy#Fallacy_of_the_isosceles_triangle

Auch die »Profis« täuschen sich manchmal selbst und gehen fälschlich davon aus, ein Ergebnis bewiesen zu haben ...

E. Burger, F. Morgan, [Fermat's Last Theorem, the Four Color Theorem, and Bill Clinton for April Fools' Day](#)^a, Amer. Math. Monthly **104** no. 3 (1997), 246–255.

^a<https://doi.org/10.1080/00029890.1997.11990629>

⁹<https://de.wikipedia.org/wiki/Schachbrett-Paradoxon>

3.5. Offene Fragen *

Manchmal hört man, dass in der Mathematik ja wohl seit langem alles erforscht sei. Dabei ist das Gegenteil der Fall. In der Mathematik und auch in der linearen Algebra gibt es viele noch offene Fragen. Der Stoff der Vorlesungen über Lineare Algebra ist allerdings so grundlegend und so stark »optimiert«, dass uns in der Vorlesung solche offene Fragen kaum begegnen werden.

Besonders in der Zahlentheorie gibt es aber Vermutungen, die ganz leicht zu formulieren sind, die aber seit Jahrzehnten (oder teilweise seit Jahrhunderten) offen sind. Hier zwei Beispiele:

VERMUTUNG 3.12 ([Goldbachsche Vermutung](https://de.wikipedia.org/wiki/Goldbachsche_Vermutung)¹⁰, 1742). *Jede gerade Zahl, die größer als 2 ist, lässt sich als die Summe von zwei Primzahlen schreiben.*

Es ist bekannt, dass die Vermutung für alle geraden Zahlen bis 4 000 000 000 000 000 000 richtig ist. Man weiß auch, dass jede ungerade Zahl sich als Summe von höchstens 3 Primzahlen schreiben lässt (dies ist die sogenannte *schwache Goldbachsche Vermutung*, die 2013 von Helfgott bewiesen wurde).

Um die nächste Vermutung zu formulieren, brauchen wir eine kleine Vorbereitung. Wir nennen natürliche Zahlen teilerfremd, wenn sie keinen gemeinsamen Teiler > 1 besitzen – zum Beispiel sind 12 und 55 teilerfremd, 15 und 55 jedoch nicht. Das *Radikal* einer natürlichen Zahl n ist das Produkt aller Primzahlen, die n teilen. (Der Unterschied zwischen der Zahl n und ihrem Radikal ist also, dass in n einige Primfaktoren mit einer höheren Potenz auftreten können.) Das Radikal von $1152 = 2^7 \cdot 3^2$ ist $2 \cdot 3 = 6$.

VERMUTUNG 3.13 ([abc-Vermutung](https://de.wikipedia.org/wiki/abc-Vermutung)¹¹, Oesterlé, Masser, ca. 1985). *Sei $\varepsilon > 0$ eine positive reelle Zahl. Dann gibt es nur endlich viele teilerfremde natürliche Zahlen a, b und c , so dass $a + b = c$ und*

$$c > \text{rad}(abc)^{1+\varepsilon}.$$

Zum Beispiel sind $a = 3, b = 125, c = 128$ teilerfremd und es gilt $a + b = c$. Die Primzahlen, die eine der drei Zahlen teilen, sind 2, 3 und 5, also gilt $\text{rad}(abc) = 2 \cdot 3 \cdot 5 = 30$. Es gilt also $c > \text{rad}(abc)$ (und es ist nicht so schwer zu sehen, dass die obige Aussage falsch würde, wenn man einfach ε durch 0 ersetzt). Man kennt aber kein einziges Beispiel für teilerfremde Zahlen a, b, c mit $a + b = c$ und $c > \text{rad}(abc)^2$.

Die [Millennium-Probleme](https://de.wikipedia.org/wiki/Millennium-Probleme)¹² sind sieben Probleme, für deren Lösung das Clay Institute im Jahr 2000 ein Preisgeld von jeweils 1 Million Dollar ausgelobt hat. Von diesen Problemen konnte bisher nur die Poincaré-Vermutung bewiesen werden (G. Perelman, 2002; der Preis wurde ihm 2010 zugesprochen, aber Perelman hat den Preis abgelehnt).

3.6. Wie schreibe ich einen Beweis (oder: die Lösung einer Aufgabe) auf?

Oder allgemeiner: Wie schreibt man ein »Stück Mathematik« auf, das man verstanden hat und nun dokumentieren möchte? Sei es ein Beweis, den man sich überlegt hat, sei es die Lösung einer Hausaufgabe oder die Lösung einer Klausuraufgabe. Natürlich gibt es zwischen diesen Kategorien gewisse Unterschiede, aber viele Grundsätze sind doch in allen Fällen gleich.

Auch das ist etwas, was man lernen und üben muss (und kann!). Machen Sie sich bewusst, dass gute Aufschreiben eines der Lernziele der Vorlesung ist. Am wichtigsten dafür ist das »learning by doing« durch das Abgeben der Hausaufgaben. Aber es ist vielleicht auch hilfreich, sich bewusst zu machen, was am anstrebt und worauf man achten sollte.

¹⁰ https://de.wikipedia.org/wiki/Goldbachsche_Vermutung

¹¹ <https://de.wikipedia.org/wiki/Abc-Vermutung>

¹² <https://de.wikipedia.org/wiki/Millennium-Probleme>

3.6.1. Grundprinzipien. Drei Punkte, die ich ganz an den Anfang stellen möchte:

- Als erstes muss man den mathematischen Inhalt dessen verstehen, was man aufschreiben möchte. Natürlich sollte man auch schon für das Knacken eines mathematischen Problems Papier und Stift zur Hand haben, so dass man sich Notizen machen, Rechnungen ausführen, Informationen zusammentragen und Zusammenhänge festhalten und »Sachen ausprobieren« kann.

Der aufgeschriebene Beweis ist aber in aller Regel nicht die Version, die vorher auf einem Schmierzettel entstanden ist: Nachdem man die Lösung verstanden hat, ist es ein wichtiger Schritt, die optimale Struktur herauszuarbeiten, wie man sie aufschreibt (und sich zu überlegen, welche Teile der vorherigen Überlegungen sich als überflüssig herausgestellt haben und nun weggelassen werden sollten).

- Was sollte das Ziel beim Aufschreiben sein? Dass die gedachte Leser*in den Beweis/die Lösung der Aufgabe gut verstehen kann. Eine gute Daumenregel ist dabei, sich eine Leser*in vorzustellen, die ungefähr dieselben Kenntnisse hat/Voraussetzungen mitbringt wie man selbst *bevor man sich mit dieser konkreten Aufgabe beschäftigt hat*. (Es ist nicht so sinnvoll, an eine Leser*in zu denken, der die Lösung schon vorher bekannt ist; dann könnte man in vielen Fällen ja einfach schreiben »Die Lösung ist offensichtlich.«)

Die Kompetenz, auch schwierige Sachverhalte für andere verständlich zu dokumentieren ist ein wichtiges Lernziel der Mathematikvorlesungen der ersten Semester (und die »Kompetenz«, eine Lösung für jemanden aufzuschreiben, der die Lösung ohnehin schon kennt, ist im Gegensatz dazu ziemlich nutzlos).

Schreiben Sie die Lösung so auf, dass Sie damit wirklich zufrieden sind, und so, wie Sie es sich als Leser*in wünschen würden! Insbesondere sollte alles, was Ihnen Schwierigkeiten bei der Lösungsfindung bereitet hat, angesprochen werden. Versuchen Sie, es der Leser*in so leicht wie möglich zu machen.

- Lassen Sie weg, was nicht erforderlich ist. Besonders bei Lösungen zu Klausuraufgaben: Dies ist nicht der Ort um darzulegen, was Sie sonst noch alles wissen! Es ist auch kein Pluspunkt, wenn sich die Korrektor*in aus vielen richtigen, aber meist irrelevanten Aussagen diejenigen herausuchen muss, die überhaupt für die Lösung eine Rolle spielen – im Gegenteil! Versuchen Sie, die Lösung auf den Punkt zu bringen, indem Sie die erforderlichen Argumente alle nennen und andere Sachen ausblenden.

Bei anderen Texten kann es sinnvoll sein, auch einmal »über den Tellerrand zu schauen« und zusätzliche Bemerkungen oder Hinweise zu geben. Diese sollten dann aber getrennt werden vom Beweis des entsprechenden Satzes und als Zusatzstoff erkennbar sein.

3.6.2. Einen Beweis/eine Lösung strukturieren. Als nächstes ist es wichtig, der Leser*in einen gut strukturierten Text zu präsentieren.

Es muss klar sein, was die zu beweisende/erklärende Aussage überhaupt ist. Wenn man einen Beweis aufschreibt, muss also erstmal der Satz formuliert werden, der überhaupt bewiesen werden soll. Schreibt man die Lösung einer Haus-/Klausuraufgabe auf, dann ist natürlich die Aufgabe vorgegeben; es ist dann nicht erforderlich, die Aufgabenstellung abzuschreiben.

Der Beweis selbst beziehungsweise die Aufgabenlösung muss eine klare Struktur haben. Beispielsweise muss sichtbar sein, ob man sich in einem direkten Beweis oder einem Widerspruchsbeweis befindet (und dann, was die zu widerlegende Annahme ist). Dazu muss man nicht unbedingt schreiben *Wir führen einen Widerspruchsbeweis.*, man kann das auch anders signalisieren (zum Beispiel: *Wir wollen die folgende Annahme zum Widerspruch führen, oder Angenommen, das wäre nicht der Fall.*). In längeren Beweisen sollten die Beweisschritte klar getrennt sein. Unter Umständen lohnt es sich, zu Beginn die Beweisstrategie zu erklären.

Ähnlich ist es mit Fallunterscheidungen, »Vorüberlegungen«, die man dem eigentlichen Beweis voranstellt, und so weiter. Gegebenenfalls kann man einzelne Teile auslagern in Hilfssätze oder Lemmata, die man eigenständig formuliert und beweist.

Insbesondere, bevor gerechnet wird, muss erklärt werden, *was die Rechnung ist*, gegebenenfalls *nach welchem Verfahren sie durchgeführt wird*, und idealerweise *welchem Zweck sie dient*. Je nach Situation sind manche dieser Punkte selbsterklärend (zum Beispiel: Soll die Lösungsmenge eines linearen Gleichungssystems bestimmt werden, kann es genügen zu sagen *Wir bringen die Matrix A mit dem Gauß-Algorithmus auf reduzierte Zeilenstufenform.*; der Zweck dieser Rechnung sollte im Kontext der Aufgabe klar sein).

Das Ziel einer Rechnung anzugeben hat bei der Lösung von Aufgaben den zusätzlichen Nutzen, dass Rechenfehler nicht so stark ins Gewicht fallen, wenn klar ist, dass Sie wussten, was zu tun ist und sich wirklich nur verrechnet haben. Wenn andererseits der Eindruck entsteht, dass Sie einfach irgendetwas rechnen, ohne einen Zusammenhang mit der Aufgabe herstellen zu können, gibt es auch für eine korrekte Rechnung nicht die volle Punktzahl.

Dazu, den Beweis gut zu strukturieren, gehört es auch, die Implikationspfeile $\Rightarrow$, $\Leftarrow$ und den Äquivalenzpfeil $\Leftrightarrow$ richtig zu benutzen (wenn man sie überhaupt benutzt und nicht durch Text ersetzt; siehe unten). Als erstes heißt das natürlich, nur korrekte Aussagen hinzuschreiben und nicht eine Implikation/Äquivalenz zu behaupten, die gar nicht gilt. Zweitens ist aber darauf zu achten, dass alle Implikationen, die für den Beweis benötigt werden, auch angegeben werden. Je nach Aufbau eines Beweises können das Äquivalenzen sein, oder Folgerungen der Form $\Rightarrow$. Reduktionsschritte (*es genügt dafür zu zeigen, dass ...*) entsprechen aber Folgerungen der Richtung $\Leftarrow$.

Es sollte im Beweis/in der Lösung sichtbar werden, wo die Voraussetzungen eingehen (besonders dann, wenn diese Stellen etwas »versteckt« sind). Ein Satz sollte so formuliert werden, dass nur solche Voraussetzungen gemacht werden, die auch erforderlich sind (sei es, damit die Aussage überhaupt richtig ist, sei es, damit der dann angegebene Beweis gültig ist; je nachdem könnte man darauf hinweisen, dass eine Voraussetzung mit entsprechender Mehrarbeit weggelassen werden kann). Bei Aufgaben ist es in der Regel auch so, dass die gemachten Voraussetzungen ihre Bedeutung haben; und wenn man eine der Voraussetzungen in seiner Lösung gar nicht benötigt, ist das ein Warnsignal, das darauf hinweisen kann, dass vielleicht etwas noch nicht stimmt.

3.6.3. Einige Regeln, um mathematische Aussagen aufzuschreiben. Allgemein sind beim Aufschreiben mathematischer Aussagen einige Sachen zu beachten.

- Symbole, die Sie verwenden, müssen definiert werden! Natürlich können Sie allgemeine Bezeichnungen, die »Standard« sind bzw. in der Vorlesung eingeführt wurden, verwenden. Aber Symbole, die »auf einmal auftauchen«, müssen beim ersten Auftreten eingeführt werden. Verlassen Sie sich nicht darauf, dass der Leser »errät«, dass A die Matrix bezeichnet, die ein paar Zeilen vorher betrachtet wurde, oder dass v ein Element des Vektorraums ist, den Sie betrachten.

- Es ist wichtig, Klarheit hinsichtlich der Quantoren herzustellen (egal, ob die Formelsymbole $\forall$, $\exists$ benutzt werden, oder ob sie durch Text ersetzt werden): Ist gemeint, dass eine Aussage *für alle* Objekte gilt, oder dass es überhaupt *ein Objekt gibt*, für das die Aussage gilt. Wenn mehrere Quantoren im Spiel sind, ist es wichtig, die Reihenfolge klarzumachen – es ist ein großer Unterschied, ob es *für alle x ein y gibt, so dass ...*, oder ob es *ein y gibt, so dass für alle x ...*

- Um zu zeigen, dass ein Objekt mit einer gewissen Eigenschaft existiert, ist es am einfachsten, ein möglichst konkretes Beispiel zu geben. Natürlich ist das nicht immer möglich. Aber wenn man kein konkretes Beispiel geben kann, muss man besonders aufpassen, dass man

nicht nur darüber redet, dass die Existenz des gewünschten Objektes nicht im Widerspruch zu anderen Tatsachen steht. Das reicht nicht aus, die Existenz zu belegen! Denn es könnte ja immer noch andere Gründe geben, die einem gerade nicht eingefallen sind, und die nach sich ziehen, dass es ein solches Objekt eben doch nicht geben kann. Ganz besonders gilt diese Bemerkung für Aufgaben, in denen gezeigt werden soll, dass eine Aussage *nicht allgemein richtig* ist. Man sollte dann immer versuchen ein (möglichst) konkretes Gegenbeispiel anzugeben.

- Schreiben Sie nur dann $\Rightarrow$, wenn beide Seiten gleich sind! Zum Beispiel nicht, wenn es sich um Matrizen handelt, die durch eine Zeilenumformung auseinander hervorgehen. Oder wenn auf der einen Seite eine Matrix, auf der anderen ihre Determinante steht. Manchmal ist es wichtig zu wissen, als was die beiden Seiten verstanden werden, um über die Gleichheit zu entscheiden (es ist zum Beispiel zulässig, $1 = 3$ zu schreiben, wenn klar ist, dass es sich hier um Elemente des Körpers

$\mathbb{F}_2$ handelt).

- Gleichungsketten $a = b = c = d$ sollten so geschrieben werden, dass die Gültigkeit jedes einzelnen $=$ klar ist beziehungsweise direkt begründet wird. Im obigen schematischen Beispiel sollten also die Gleichheiten $a = b$, $b = c$, $c = d$ die Gleichheiten sein, die einzeln erklärt werden können (oder sogar offensichtlich sind). Das Ziel des Arguments ist dann, die Gleichheit $a = d$ zu folgern.

3.6.4. Unterschiede zwischen verschiedenen Formaten. Es ist ein Unterschied, ob man – womöglich unter Zeitdruck wie in einer Klausur – die Lösung einer Aufgabe aufschreibt, ob man eine Abschlussarbeit (Bachelorarbeit, Masterarbeit, ...), einen Artikel für eine Fachzeitschrift oder ein Lehrbuch schreibt. Es ist wieder etwas anderes, ob man an der Tafel einen Vortrag hält, oder ein Lernvideo macht.

Speziell in der Klausur und an der Tafel ist es erlaubt und sinnvoll, sich kurz zu fassen. Das bedeutet auch, dass Abkürzungen benutzt werden können (VR statt Vektorraum, l.u. statt linear unabhängig) und Formelsprache eingesetzt werden kann ($\forall$, $\exists$, $\Leftrightarrow$), und dass nicht unbedingt voll ausformulierte Sätze verwendet werden müssen.

Aber: Das darf nicht auf Kosten der Klarheit gehen und bedeutet, dass man sich *umso mehr Mühe geben muss, dass die wichtige Information vollständig vorhanden ist!* Besonders in diesen Formaten ist es wichtig, Beweise/Lösungen klar zu strukturieren. Erklärungen sind mindestens so wichtig wie (an der Tafel: wichtiger als) einzelne Rechenschritte!

Bei einem getippten Text sieht es natürlich etwas anders aus. Dieser sollte flüssig zu lesen sein, und das bedeutet auch, dass von »Formelsprache« nur sehr zurückhaltend Gebrauch gemacht werden sollte. (Auf den ca. 350 Seiten dieses Skripts kommt das Symbol $\forall$ weniger als zehnmal vor; davon fünfmal in dem Abschnitt, in dem das Symbol definiert wird.) Ein bisschen unterscheidet sich dieser Punkt je nach Geschmack der Autor*in und auch nach mathematischen Gebiet. Jedenfalls sollte aber die Lesbarkeit an erster Stelle stehen, und nicht der Wunsch, mit möglichst wenigen Schriftzeichen auszukommen.

Bei getippten Texten sollten einige Grundregeln des mathematischen Satzsatzes beachtet werden. Siehe C.2 für eine kurze Liste und weitere Verweise (siehe auch die Box unten).

In einem Buch (und auch bei einer Abschlussarbeit, vielleicht mit einem etwas anderen Maßstab) kann es auch sinnvoll sein, bewusst einzelne Schritte der Leser*in zu überlassen. Anders als bei einer Aufgabenlösung steht hier der Punkt, dass die Leser*in bei der Lektüre etwas Neues lernt im Vordergrund, und der Punkt, dass die Autor*in dokumentiert, den Stoff selbst verstanden zu haben, im Hintergrund. Und neue Mathematik zu lernen, geht eben nur mit aktivem Mitdenken, dass der Text deshalb auch explizit einfordern darf, indem Lücken

gelassen werden, die die Leser*in selbst füllen muss. Gerade bei Rechnungen oder Ausführungen, bei denen es genauso lange dauern würde, sie sich mit Verständnis durchzulesen, wie sie selbst durchzuführen, ist das sinnvoll. Allerdings: Diese »Lücken« sollten explizit als solche erkennbar sein, und nicht dadurch entstehen, dass die Autor*in selbst vergessen hat, über wichtige Punkte nachzudenken, oder dadurch begründet sein, dass die Autor*in sich selbst Arbeit sparen möchte! (Es ist eine häufige Fehlerquelle, dass man Beweisschritte, die man sich nur ungefähr überlegt hat, als *offensichtlich* oder *trivial* bezeichnet. Wenn immer man diese Wörter schreibt, sollte man sich zu besonderer Aufmerksamkeit zwingen.)

3.6.5. Praxistraining. Mathematik aufzuschreiben, muss man üben (genau so, wie mathematische Probleme zu lösen). Im Studium haben Sie dazu Gelegenheit, indem Sie regelmäßig die Hausaufgaben abgeben, die Korrekturen sorgfältig durchgehen und gegebenenfalls Nachfragen dazu stellen.

Darüber hinaus ist es sinnvoll, sich dann und wann Zeit zu nehmen und bewusst auf den Aspekt der Darstellung/des Aufschreibens zu achten. Man kann sich zum Beispiel die »Musterlösungen«/Lösungsvorschläge in Hinsicht darauf anschauen und mit den eigenen Lösungen vergleichen. Oder die Vorlesung, das Skript, Lehrbücher anschauen und gegebenenfalls auch mal vergleichen und bewerten. (Was gefällt Ihnen gut, was nicht so gut? Was ist gut oder nicht so gut zu lesen und woran liegt das? Mit welchem Stil kommen Sie gut klar?)

Das gewissenhafte Aufschreiben ist eine sehr gute Qualitätskontrolle, ob man die Lösung der Aufgabe wirklich verstanden hat (oder ob die Aussage, die man glaubt bewiesen zu haben, überhaupt korrekt ist). Oft genug stellt man beim Aufschreiben (bzw. wenn man versucht, seinen *Beweis* jemand anderem zu erklären) fest, dass man doch noch etwas übersehen hat. Auch deshalb lohnt es sich, das klare Aufschreiben zu trainieren.

Hier noch einige weitere Links (überwiegend auch/mehr zum Thema *Verfassen einer Abschlussarbeit*, aber Vieles lässt sich auch auf das Aufschreiben von Aufgabenlösungen übertragen).

Hier im Skript sind die Anhängen C.2 und D.5.3 relevant.

Clara Löh, Univ. Regensburg, [Exercises in academic writing](http://www.mathematik.uni-regensburg.de/loeh/seminars/eaw.pdf)^a

Mirjam Dür, Univ. Trier, [Schreiben mathematischer Texte](https://www.math.uni-trier.de/duer/MathSchreiben.pdf)^b

Peter Pfaffelhuber, Univ. Freiburg, [Schreiben einer mathematischen Arbeit](https://www.stochastik.uni-freiburg.de/professoren/pfaffelhuber/inhalte/writing)^c

^a<http://www.mathematik.uni-regensburg.de/loeh/seminars/eaw.pdf>

^b<https://www.math.uni-trier.de/duer/MathSchreiben.pdf>

^c<https://www.stochastik.uni-freiburg.de/professoren/pfaffelhuber/inhalte/writing>

3.6.6. Lebenslanges Lernen. Klar und fehlerfrei über schwierige und komplizierte Dinge zu kommunizieren ist eine der Fertigkeiten, die man sich im Mathematikstudium aneignet (denn ohne sie kann man das Studium kaum erfolgreich abschließen) und die auch im Berufsleben – egal wo – von sehr hohem Nutzen sind. Auch wenn es nicht um Mathematik, sondern um anderes geht. Dies ist aus meiner Sicht einer der wichtigen Gründe, warum Mathematiker*innen auf dem Arbeitsmarkt so gefragt sind. (Und bezogen auf die Lehramtsstudierenden: Etwas, wo Ihre Schüler*innen von Ihrem guten Beispiel mindestens genauso profitieren können, wie von dem konkreten mathematischen Stoff, den Sie vermitteln.)

Wie so vieles ist das gute Aufschreiben etwas, wo man sich immer weiter verbessern kann, und was je nach Tagesform mal besser und mal schlechter klappt. Und das gilt auch für dieses Skript! Ich hoffe zwar, dass es in weiten Teilen als gutes Beispiel taugt, aber das ist sicher nicht überall so. »Beschweren« Sie sich gerne, zum Beispiel im Forum auf der Moodle-Seite zur Vorlesung, wenn etwas für Sie unverständlich ist!

3.7. Grundbegriffe, Notation

In diesem Abschnitt klären wir einige Grundbegriffe, die wir benutzen werden, um über Mathematik zu sprechen.

3.7.1. Aussagen. Eine *Aussage* formuliert Eigenschaften mathematischer Objekte und kann wahr oder falsch sein.

BEISPIEL 3.14. (1) Beispiele für Aussagen sind:

- Die Zahl 5 ist eine Primzahl.
 - Die Zahl 5 ist durch 3 teilbar.
- (2) In den folgenden Sätzen kommen zwar auch mathematische Symbole vor, es handelt sich aber nicht um Aussagen:
- Die Zahl 5 ist grün. (Denn *grün* hat keine mathematische Bedeutung.)
 - Die Zahlen 5, 7 größer kleiner. (Kein sprachlich verständlicher Satz.)

◇

Wir können Aussagen miteinander verknüpfen: Seien $A, B, C, \dots$ mathematische Aussagen.

- (1) $A \wedge B$ (A und B) ist genau dann wahr, wenn A und B wahr sind.
- (2) $A \vee B$ (A oder B) ist genau dann wahr, wenn A oder B wahr sind. (Es ist auch erlaubt, dass A und B beide wahr sind.)
- (3) $\neg A$ (nicht A) ist genau dann wahr, wenn A falsch ist. Man nennt $\neg A$ auch die *Negation* der Aussage A .
- (4) $A \Rightarrow B$ bedeutet, dass B aus A folgt, das heißt: wenn A wahr ist, dann ist B wahr. (Für den Fall, dass A falsch ist, wird keine Aussage über B gemacht; in diesem Fall ist die Aussage $A \Rightarrow B$ unabhängig von B wahr!)

Mit anderen Worten: $A \Rightarrow B$ ist genau dann wahr, wenn B wahr oder A falsch ist. (Wir können also das Symbol $\Rightarrow$ als Abkürzung sehen und könnten $A \Rightarrow B$ immer umschreiben als $B \vee \neg A$. Diese »formalistische« Sichtweise des Folgepfeils ist sicher etwas gewöhnungsbedürftig.)

- (5) $A \Leftarrow B$ ist gleichbedeutend mit $B \Rightarrow A$.
- (6) $A \Leftrightarrow B$ (genau dann A , wenn B) ist wahr wenn A beide wahr oder beide falsch sind. Damit ist $A \Leftrightarrow B$ gleichbedeutend mit $(A \Rightarrow B) \wedge (B \Rightarrow A)$.

Von diesen Zeichen benutzen wir im weiteren Text eigentlich nur $\Rightarrow$ und $\Leftrightarrow$ häufiger, und selbst diese Symbole schreiben wir meist aus.

Um die Äquivalenz von zwei Arten, Aussagen zu verknüpfen, zu überprüfen – zum Beispiel die oben genannte Äquivalenz von $A \Rightarrow B$ und $B \vee \neg A$ – muss man prüfen, dass für alle möglichen Werte von A und B , also für alle Kombinationen von wahr/falsch, beide Verknüpfungen dasselbe Ergebnis liefern. Diese Werte kann man in einer Tabelle (»Wahrheitstafel«) auflisten. Die Zeilen entsprechen den Werten von A (also: in der ersten Zeile ist A in allen Fällen wahr; in der zweiten Zeile ist A in allen Fällen falsch), die Spalten den Werten von B (erste Spalte: B wahr, zweite Spalte: B falsch) und in den Einträgen geben wir das Ergebnis des betrachteten Ausdrucks an. Wir schreiben w für wahr, f für falsch.

$A \setminus B$	w	f
w	w	f
f	w	w

Wahrheitstafel für $A \Rightarrow B$.

$A \setminus B$	w	f
w	w	f
f	w	w

Wahrheitstafel für $B \vee \neg A$.

Diese Überprüfungen können also ohne weiteres Nachdenken durchgeführt werden (sind aber dafür ziemlich langweilig).

Mit derselben Methode kann man folgendes überprüfen: Wenn $A \Rightarrow B$ und $B \Rightarrow C$ gilt, dann gilt auch $A \Rightarrow C$.

Oft wird es vorkommen, dass wir für eine Liste von mehreren Aussagen zeigen wollen, dass je zwei dieser Aussagen äquivalent sind. Sind zum Beispiel Aussagen (i), (ii), (iii), (iv) gegeben, deren Äquivalenz gezeigt werden soll, dann genügt es wegen der obigen Bemerkung, statt alle 12 Implikationen zu beweisen, die Implikationen (i) $\Rightarrow$ (ii), (ii) $\Rightarrow$ (iii), (iii) $\Rightarrow$ (iv) und (iv) $\Rightarrow$ (i) zu zeigen. (Denn aus (ii) $\Rightarrow$ (iii) und (iii) $\Rightarrow$ (iv) folgt (ii) $\Rightarrow$ (iv), und mit (iv) $\Rightarrow$ (i) erhalten wir (ii) $\Rightarrow$ (i), also insgesamt (i) $\Leftrightarrow$ (ii), usw.) Man nennt diese Vorgehensweise manchmal einen *Ringschluss*.

3.7.2. Quantoren. Der *Allquantor* $\forall$ bedeutet, dass eine Aussage für *alle* Elemente einer Menge gelten soll. Nach dem $\forall$ wird angegeben, auf welche »Variable« er sich bezieht, nach einem Doppelpunkt kommt dann die eigentliche Aussage. Zum Beispiel:

$$\forall n \in \mathbb{Z} : n^2 \geq 0,$$

in Worten: Für jede ganze Zahl $n \in \mathbb{Z}$ ist $n^2 \geq 0$.

Der *Existenzquantor* $\exists$ drückt aus, dass (mindestens) ein Element existiert, so dass die Aussage wahr ist. Zum Beispiel:

$$\exists n \in \mathbb{Z} : (n > 5 \wedge n < 8),$$

in Worten: Es gibt eine ganze Zahl n , die größer als 5 und kleiner als 8 ist.

Es ist wichtig, die *Reihenfolge der Quantoren* zu beachten: Genauso wie $\forall m \in \mathbb{Z} \exists n \in \mathbb{Z} : m = n + 1$ richtig ist, ist $\exists n \in \mathbb{Z} \forall m \in \mathbb{Z} : m = n + 1$ offensichtlich falsch – machen Sie sich das klar.

3.7.3. Beweismethoden. Um Beweisen zu lernen (was wie gesagt ein wichtiges Ziel von Ihnen für diese Vorlesung sein sollte) muss man es an konkreten Aufgaben üben; dazu werden Sie ausgiebig Gelegenheit haben. In diesem Abschnitt soll es aber erst einmal nur darum gehen, einige Begrifflichkeiten zu erklären.

Unter einem *direkten Beweis* für eine zu zeigende Aussage B (die *Behauptung*) versteht man eine Kette von Folgerungen $A_1 \Rightarrow A_2 \Rightarrow \dots \Rightarrow A_n \Rightarrow B$, für die bekannt ist, dass A_1 wahr ist. Damit folgt dann auch, dass B wahr ist. Wir haben zum Beispiel für den Satz des Pythagoras (Satz 3.2) einen direkten Beweis gegeben.

Ein *indirekter Beweis* oder *Widerspruchsbeweis*, *Beweis durch Widerspruch* einer Behauptung B besteht aus einer Kette von Folgerungen, die *ausgehend von der Negation von B* eine *falsche* Aussage ableitet. Wenn aber $\neg B \Rightarrow A$ für eine falsche Aussage A gilt, dann muss $\neg B$ falsch, also B wahr sein. Wir haben den Beweis, dass keine rationale Zahl existiert, deren Quadrat 2 ist, als indirekten Beweis geführt (Satz 3.4).

Um die Wahrheit einer Aussage der Form $A \Rightarrow B$ zu beweisen (ohne sich auf die Wahrheit von A und B festzulegen), kann man genausogut die *äquivalente* Aussage $\neg B \Rightarrow \neg A$ beweisen, die sogenannte *Kontraposition* der ursprünglichen Aussage.

Name	Groß	Klein	Name	Groß	Klein
alpha	A	α	ny	N	ν
beta	B	β	xi	Ξ	ξ
gamma	Γ	γ	omikron	O	$\omicron$
delta	Δ	δ	pi	Π	π
epsilon	E	ε	rho	P	ρ (auch: ϱ)
zeta	Z	ζ	sigma	Σ	σ (auch: ς)
eta	H	η	tau	T	τ
theta	Θ	θ (auch: ϑ)	ypsilon	Y	υ
iota	I	ι	phi	Φ	φ (auch: ϕ)
kappa	K	κ	chi	X	χ
lambda	Λ	λ	psi	Ψ	ψ
my	M	μ	omega	Ω	ω

TABELLE 1. Die griechischen Buchstaben

Manchmal ist es für einen Beweis hilfreich, eine *Fallunterscheidung* vorzunehmen, also verschiedene Fälle getrennt zu behandeln. Das kann man machen, wenn man dabei alle Möglichkeiten abdeckt. Manchmal bietet es sich dann an, die Ergebnisse aus bereits abgeschlossenen Fällen in den weiteren Fällen zu benutzen.

Eine weitere wichtige Beweismethode (für Aussagen über natürliche Zahlen) ist die *vollständige Induktion*, der wir einen eigenen Abschnitt (Abschnitt 3.12) widmen werden.

3.7.4. Notationen. Manchmal benutzen wir das Symbol $:=$ um anzuzeigen, dass das Symbol auf der linken Seite (der Seite mit dem Doppelpunkt) neu definiert wird, und gleich dem Ausdruck auf der rechten Seite sein soll.

Für zwei (natürliche, rationale oder allgemein zwei reelle) Zahlen a, b bezeichnen wir mit $\min(a, b)$ die kleinere der beiden, das *Minimum*, und mit $\max(a, b)$ die größere, das *Maximum*. (Sind a und b gleich, so ist $\min(a, b) = \max(a, b) = a = b$.) Analog kann man diese Bezeichnungen für nicht-leere endliche Mengen von Zahlen verwenden. Für unendliche Mengen müssen das Minimum und Maximum nicht unbedingt existieren; dann ist also besondere Vorsicht geboten.

3.7.5. Griechische Buchstaben. Gelegentlich benutzen wir *griechische Buchstaben*. Dass das in der Mathematik nicht unüblich ist, wissen Sie schon – denken Sie an die Kreiszahl π . In Ihrem Mathematikstudium wird Ihnen wahrscheinlich jeder griechische Buchstabe irgendwann einmal begegnen, insofern ist es gut investierte Zeit, sie gleich zu Beginn lesen und schreiben zu lernen (Tabelle 1).

3.8. Mengen

Heutzutage wählen die meisten Mathematiker als die formale Grundlage der Mathematik den Begriff der *Menge*. Kompliziertere Objekte drückt man dann in geeigneter Weise als/durch Mengen aus. Zum Beispiel kann man die natürlichen Zahlen als Mengen »kodieren«, indem man die natürliche Zahl n durch eine Menge mit n Elementen »ersetzt«. Unten werden wir sehen, wie man den Begriff der *Funktion* $X \rightarrow Y$ in Termen von Mengen ausdrückt.

Es bleibt allerdings zu sagen, was man überhaupt unter einer *Menge* versteht. Formal wäre der richtige Ansatz hier ein geeignetes Axiomensystem zugrundezulegen, das die essenziellen Eigenschaften von Mengen festschreibt und auf dessen Grundlage wir dann arbeiten könnten – zum Beispiel das Axiomensystem ZFC, benannt nach den Mathematikern Ernst

Zermelo und Abraham Fraenkel. (Das C steht für das »Auswahlaxiom«, englisch: axiom of choice, siehe Anhang B.1.)

Wir wollen uns diese Arbeit aber hier ersparen und wählen stattdessen einen »naiven Zugang«.

Eine *Menge* ist eine Zusammenfassung mathematischer Objekte. Eine Menge M ist gegeben durch ihre Elemente: Für alle Objekte x gilt entweder $x \in M$ (d.h. x ist ein Element von M ; wir sagen, x liege in M) oder $x \notin M$ (d.h. x ist kein Element von M). Zwei Mengen M, M' sind genau dann gleich, wenn sie dieselben Elemente haben, also wenn für alle x gilt:

$$x \in M \quad \text{genau dann, wenn} \quad x \in M'.$$

Literaturverweise: (Meine Empfehlung bleibt aber, sich diese für einen späteren Zeitpunkt aufzuheben.)

Ein bisschen mehr sagt Wolfram Pohlers (dessen Forschungsgebiet die mathematische Logik ist) in seinem [Skript zur Linearen Algebra](#)^a (Kapitel 0).

Auch in dem (englischen) Buch [Hu] von Hungerford finden Sie etwas mehr Informationen zu diesem Thema.

Wesentlich ausführlicher ist das Skript [Grundkurs Mengenlehre](#)^b von W. Rautenberg.

^a https://www.uni-muenster.de/imperia/md/content/logik/Skripte/pohlers.lineare_algebra.pdf

^b <http://page.mi.fu-berlin.de/raut/Mengenlehre/m.pdf>

BEISPIEL 3.15. • Die *leere Menge* bezeichnen wir mit $\emptyset$ oder mit $\{\}$. Dies ist die einzige Menge, die kein Element enthält, d.h. für alle x gilt $x \notin \emptyset$. Achtung: »Doppelt gemoppelt« hält in diesem Fall nicht besser: $\{\emptyset\}$ ist *nicht* die leere Menge, sondern eine Menge mit einem Element: $\emptyset \in \{\emptyset\}$.

- Wir können Mengen angeben, indem wir alle ihre Elemente hinschreiben: $\{1, 2, 3, 4, 5\}$. Beachten Sie, dass

$$\{1, 2, 3, 4, 5\} = \{1, 4, 2, 5, 3\} = \{1, 1, 3, 5, 4, 4, 2\},$$

die Reihenfolge der Elemente spielt ebensowenig eine Rolle wie die Frage, ob Elemente mehrfach genannt sind.

- Für die meisten Mengen, die wir betrachten, ist es nicht möglich, alle Elemente direkt hinzuschreiben. Oft werden wir Mengen dadurch angeben, dass wir aus einer Grundmenge alle Elemente auswählen, die eine gewisse Eigenschaft haben. Ist zum Beispiel $\mathbb{N}$ die Menge der natürlichen Zahlen, dann ist

$$\{n \in \mathbb{N}; \exists m \in \mathbb{N} : n = m^2\}$$

die Menge der Quadratzahlen. Statt des Semikolons findet man auch häufig die Notation mit einem senkrechten Strich |.

- Manchmal kann man alle Elemente einer Menge »aufzählen«, indem man beschreibt, wie sie durch eine Konstruktion aus anderen Elementen entstehen, zum Beispiel können wir die Menge aller Bruchzahlen schreiben als

$$\mathbb{Q} = \left\{ \frac{a}{b}; a, b \in \mathbb{Z}, b \neq 0 \right\}.$$

◇

BEISPIEL 3.16. Wichtige Beispiele von Mengen sind natürlich die üblichen Zahlbereiche:

- Die Menge der *natürlichen Zahlen* bezeichnen wir mit $\mathbb{N}$:

$$\mathbb{N} = \{0, 1, 2, 3, \dots\}.$$

Die Frage, ob die Zahl 0 eine natürliche Zahl ist, und dementsprechend ob die mit dem Symbol $\mathbb{N}$ bezeichnete Menge die Zahl 0 enthält, wird in der Literatur nicht einheitlich gehandhabt. In diesem Text schließen wir die 0 mit ein. (Gegebenenfalls schreiben wir $\mathbb{N}_{>0}$ oder $\mathbb{Z}_{>0}$ für die Menge der natürlichen Zahlen ohne 0.)

- Mit $\mathbb{Z}$ bezeichnen wir die Menge der *ganzen Zahlen*:

$$\mathbb{Z} = \{\dots, -3, -2, -1, 0, 1, 2, 3, \dots\}.$$

Wir erhalten die Menge $\mathbb{Z}$ aus der Menge $\mathbb{N}$ der natürlichen Zahlen, indem wir zu jeder natürlichen Zahl $\neq 0$ ihr Negatives hinzufügen. Ganze Zahlen kann man addieren, subtrahieren und multiplizieren, und es gelten die »üblichen Rechenregeln«.

- Mit $\mathbb{Q}$ bezeichnen wir die Menge der *rationalen Zahlen*, mit anderen Worten die Menge der Bruchzahlen von zwei ganzen Zahlen:

$$\mathbb{Q} = \left\{ \frac{a}{b}; a, b \in \mathbb{Z}, b \neq 0 \right\}.$$

Es gelten die üblichen Bruchrechenregeln: Zunächst einmal gilt $\frac{a}{c} = \frac{b}{d}$ genau dann, wenn $ad = bc$ (das beschreibt die Möglichkeit, einen Bruch zu kürzen bzw. zu erweitern). Die Addition und Multiplikation von Brüchen sind definiert als

$$\frac{a}{c} + \frac{b}{d} = \frac{ad + bc}{cd}, \quad \frac{a}{c} \cdot \frac{b}{d} = \frac{ab}{cd}.$$

Wenn man mit konkreten Zahlen rechnet, sollten die Ergebnisse natürlich in gekürzter Form angegeben werden. Wir betrachten $\mathbb{Z}$ als Teilmenge (siehe Definition 3.18) von $\mathbb{Q}$ und schreiben a statt $\frac{a}{1}$, auch wenn wir a als Element von $\mathbb{Q}$ auffassen. Beachte, dass die Addition und Multiplikation auf $\mathbb{Z}$ und $\mathbb{Q}$ kompatibel sind.

- Die Menge der *reellen Zahlen* hat das Symbol $\mathbb{R}$. Auch wenn das der Rechenbereich ist, der Ihnen (vermutlich) aus der Schule am geläufigsten ist, ist er aus mathematischer Sicht deutlich komplizierter als der Bereich der rationalen Zahlen.

Man kann jede reelle Zahl durch ihre Dezimalbruchentwicklung beschreiben (die unendlich viele Stellen hinter dem Komma haben kann), aber auch hier gibt es das Problem, dass die Darstellung nicht eindeutig ist (es ist $0,999\ldots = 1$), und es ist nicht ganz trivial, die Rechenoperationen für solche Zahlen »sauber« hinzuschreiben. Weitere Ausführungen zur Konstruktion (bzw. zur Charakterisierung) der reellen Zahlen überlassen wir der Analysis-Vorlesung.

◇

ERGÄNZUNG 3.17 (Russells Paradoxon). Dass man die Definition einer Menge als Zusammenfassung mathematischer Objekte nicht ohne jegliche Vorsicht verwenden kann, zeigt das Paradoxon von [Bertrand Russell](https://de.wikipedia.org/wiki/Bertrand_Russell)¹³. Wenn wir für jede Menge M und jedes mathematische Objekt x entweder $x \in M$ oder $x \notin M$ haben, können wir auch fragen, ob $M \in M$ oder $M \notin M$ gilt. (Für die Mengen, die wir kennen, gilt sicher $M \notin M$, aber es könnte ja auch Mengen geben, die sich selbst als Element enthalten.)

¹³https://de.wikipedia.org/wiki/Bertrand_Russell



Einem wissenschaftlichen Schriftsteller kann kaum etwas Unerwünschteres begegnen, als daß ihm nach Vollendung einer Arbeit eine der Grundlagen seines Baues erschüttert wird. In diese Lage wurde ich durch einen Brief des Herrn Bertrand Russell versetzt, als der Druck dieses Bandes sich seinem Ende näherte.

G. Frege im Nachwort des zweiten Bands seiner Grundgesetze der Arithmetik von 1903, in dem er ein Axiomensystem der Arithmetik aufzubauen versuchte, das auf dem naiven Mengenbegriff basierte

Sei nun X die Menge aller Mengen M , für die gilt $M \notin M$.

Wenn dann $X \in X$ wäre, so folgt $X \notin X$. Wäre $X \notin X$, so folgt $X \in X$. So oder so erhalten wir einen Widerspruch!

Solche Paradoxien werden durch die Einführung eines geeigneten Axiomensystems wie ZFC vermieden, da sich die »Menge«, über die Russells Paradox spricht, innerhalb dieses Systems nicht konstruieren lässt.

Buchtip (in Comicform): A. Doxiadis, C. Papadimitriou, *Logicomix: Eine epische Suche nach Wahrheit*, (der Vollständigkeit halber: [...the reader should provide his/her grain of salt](#)¹⁴)

□ Ergänzung 3.17

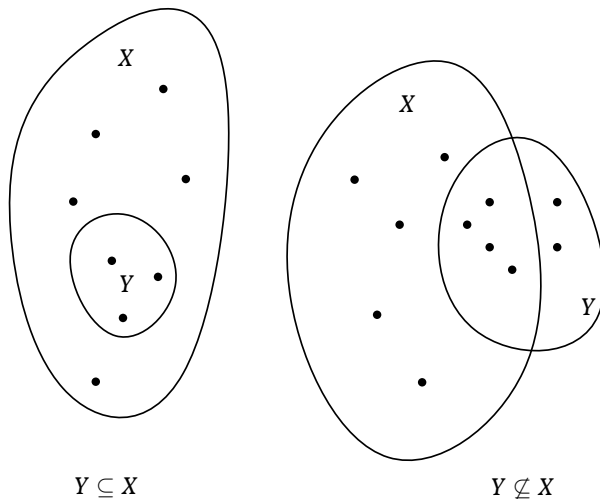
3.9. Teilmengen, Konstruktionen von Mengen

DEFINITION 3.18. Sei X eine Menge. Wir sagen, eine Menge Y sei eine *Teilmenge* von X und schreiben $Y \subseteq X$, wenn für alle $y \in Y$ gilt: $y \in X$.

Statt des Begriffs *Teilmenge* verwendet man auch den Begriff *Untermenge*, oder sagt einfach, die Menge Y sei *in X enthalten* (Wenn die Elemente von X selbst Mengen sind, muss man unter Umständen etwas aufpassen: $Y \subseteq X$ und $Y \in X$ sind zwei sehr unterschiedliche Aussagen.) Wir schreiben $Y \subsetneq X$, wenn $Y \subseteq X$ und $Y \neq X$; in dieser Situation sagt man auch, Y sei eine *echte* Teilmenge von X . $\dashv$

Wie üblich wird das entsprechende Symbol mit einem Schrägstrich durchgestrichen, um die entsprechende Aussage zu verneinen (wie bei $=$ gleich, $\neq$ ungleich). Also bedeutet $Y \not\subseteq X$, dass Y keine Teilmenge von X ist.

¹⁴https://en.wikipedia.org/wiki/Logicomix#Historical_accuracy



Statt $Y \subseteq X$ wird oft auch $Y \subset X$ geschrieben, d.h. auch das Symbol $\subset$ schließt in den meisten mathematischen Texten die Gleichheit $Y = X$ als eine Möglichkeit ein. (Die Situation ist also anders als beim Kleiner-/Größer-Zeichen.) In manchen Quellen wird das Zeichen $\subset$ allerdings auch im Sinne von $\subsetneq$ verwendet. Ich bemühe mich, das Symbol $\subset$ dann zu verwenden, wenn die Gleichheit zwar nicht formal ausgeschlossen ist, aber dieser Fall nicht relevant ist bzw. aus dem Kontext klar ist, dass er nicht auftreten kann.

Statt $Y \subseteq X$ kann man auch $X \supseteq Y$ schreiben, und man kann dies auch lesen als » X ist eine Obermenge von Y «. Entsprechendes gilt für $\supset$ und $\supsetneq$.

Wir können damit die Gleichheit zweier Mengen M, M' formulieren als

$$M = M' \quad \text{genau dann, wenn} \quad M \subseteq M', M' \subseteq M.$$

Oft zeigt man eine Gleichheit $M = M'$ von Mengen, indem man die beiden Inklusionen $M \subseteq M'$ und $M' \subseteq M$ beweist.

DEFINITION 3.19. Seien X und Y Teilmengen einer Menge M .

(1) Der *Durchschnitt* (oder die *Schnittmenge*) von X und Y ist

$$X \cap Y = \{m \in M; m \in X \text{ und } m \in Y\}.$$

(2) Die *Vereinigung* von X und Y ist

$$X \cup Y = \{m \in M; m \in X \text{ oder } m \in Y\}.$$

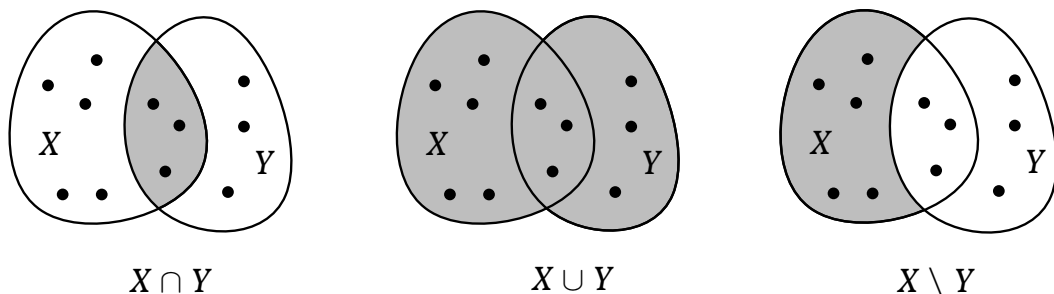
(3) Die *Differenz* von X und Y ist

$$X \setminus Y = \{m \in M; m \in X \text{ und } m \notin Y\}.$$

(4) Das *Komplement* von X in M ist

$$X^c = M \setminus X = \{m \in M; m \notin X\}.$$

⊢



Allgemeiner können wir Durchschnitte und Vereinigungen von mehr als zwei Mengen bilden:

DEFINITION 3.20. Seien M und I Mengen, und sei für jedes $i \in I$ eine Teilmenge $M_i \subseteq M$ gegeben.

(1) Der *Durchschnitt* der Teilmengen M_i ist

$$\bigcap_{i \in I} M_i = \{x \in M; x \in M_i \text{ für alle } i \in I\}.$$

(2) Die *Vereinigung* der Teilmengen M_i ist

$$\bigcup_{i \in I} M_i = \{x \in M; \text{ es gibt ein } i \in I \text{ mit } x \in M_i\}.$$

⊢

Wir nennen zwei Teilmengen M_1, M_2 einer Menge M *disjunkt*, wenn $M_1 \cap M_2 = \emptyset$.

3.10. Kartesisches Produkt, Abbildungen

3.10.1. Produkte.

DEFINITION 3.21. (1) Sind X und Y Mengen, so bezeichnen wir mit $X \times Y$ die Menge aller Paare (x, y) mit $x \in X$ und $y \in Y$. Wir nennen $X \times Y$ das (*kartesische*) *Produkt* der Mengen X und Y .

(2) Analog können wir das Produkt von mehr als zwei Mengen bilden: Sind $X_1, X_2, \dots, X_n$ Mengen, so ist

$$\prod_{i=1}^n X_i = X_1 \times \dots \times X_n = \{(x_1, \dots, x_n); \forall i = 1, \dots, n : x_i \in X_i\}$$

die Menge aller »Listen« $(x_1, \dots, x_n)$ von Elementen der Mengen X_i . Die Elemente des Produkts nennen wir *n-Tupel*.

Im Fall, dass $X_1 = \dots = X_n = X$ für eine Menge X ist, so schreiben wir X^n statt $X \times X \times \dots \times X$ für das *n-fache Produkt*.

(3) Noch allgemeiner ist für eine Menge I und Mengen X_i für alle $i \in I$ das Produkt

$$\prod_{i \in I} X_i = \{(x_i)_{i \in I}; \forall i \in I : x_i \in X_i\}$$

definiert. Man nennt I die *Indexmenge* des Produkts.

Für den Spezialfall, dass $I = \emptyset$ die leere Menge ist (und daher gar kein X_i gegeben ist) hat sich die Konvention bewährt, dass das Produkt genau ein Element hat. Man spricht dann vom *leeren Produkt* (aber leer ist die Indexmenge, nicht die Produktmenge selbst).

Ist $X_i = X$ für eine Menge X und alle i , so schreiben wir manchmal $X^I = \prod_{i \in I} X$.

⊢

Bei den Elementen des Produkts $X \times Y$ kommt es auf die Reihenfolge an, auch wenn $X = Y$ ist. Zum Beispiel hat $\{1, 2\} \times \{1, 2\}$ die vier verschiedenen Elemente $(1, 1)$, $(1, 2)$, $(2, 1)$, $(2, 2)$. Ein 3-Tupel nennt man meist *Tripel*, ein 4-Tupel manchmal *Quadrupel*, ein 5-Tupel manchmal *Quintupel*.

¹⁵https://commons.wikimedia.org/w/index.php?title=File%3ADescartes_La_G%C3%A9om%C3%A9trie.djvu

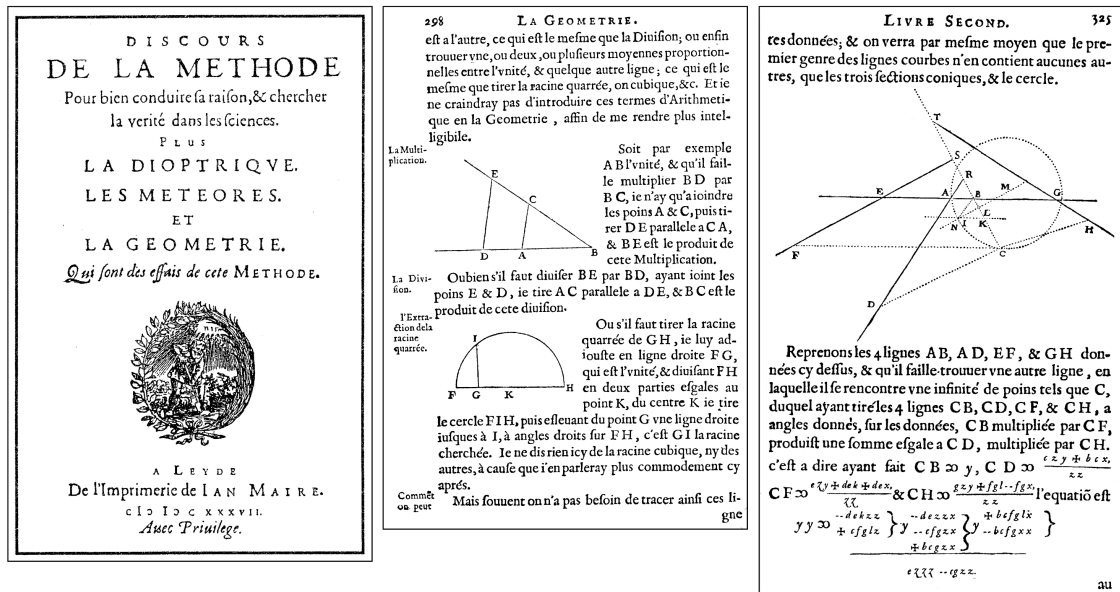
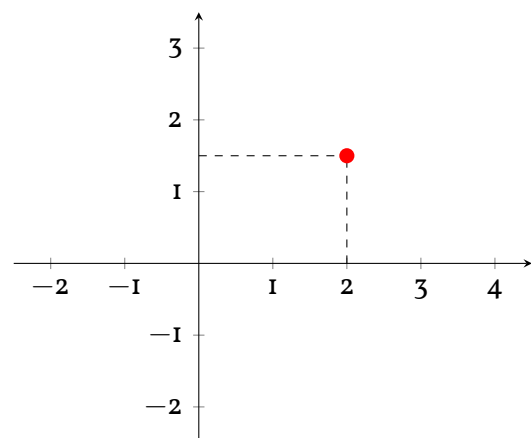


ABBILDUNG 1. Drei Seiten aus dem Buch über Geometrie von René Descartes (erschienen 1637). Descartes hat die Verwendung eines Koordinatensystems eingeführt, um die Lage von Punkten anzugeben. Koordinaten eines Punktes in der Ebene bilden ein Paar (x, y) von reellen Zahlen, also ein Element des Produkts $\mathbb{R} \times \mathbb{R}$. Deshalb nennt man das Produkt auch das *kartesische Produkt*.
Quelle: [Wikimedia](#)¹⁵

Wie Descartes veranschaulichen wir das Produkt $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$ als die Ebene, in der wir die Lage eines Punkts durch seine beiden Koordinaten angeben können. Zum Beispiel hat der rot markierte Punkt in der Abbildung die Koordinaten $(2, 1, 5)$. Die waagerechte Koordinatenachse, deren Koordinate als erste angegeben wird, heißt üblicherweise die x -Achse, die senkrechte Achse ist die y -Achse. Je nachdem kann man natürlich auch andere Namen verwenden; wird zum Beispiel ein zeitlicher Verlauf dargestellt, verwendet man oft den Parameter t (für lateinisch *tempus*) und stellt diesen auf der waagerechten Achse dar.



3.10.2. Abbildungen. Wir kommen nun zum Begriff der *Abbildung*; neben dem Begriff der Menge der grundlegendste Begriff aus diesem Kapitel.

DEFINITION 3.22. (I) Seien X, Y Mengen. Eine *Abbildung* f von X nach Y , geschrieben $f: X \rightarrow Y$, ist eine Zuordnung, die jedem Element $x \in X$ genau ein Element $f(x) \in Y$ zuordnet. Man schreibt $x \mapsto f(x)$ und nennt $f(x)$ das *Bild* von x unter f oder den *Wert* der Abbildung f bei x .

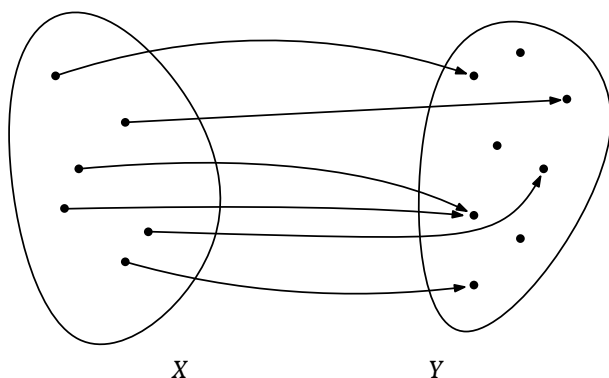
Formaler können wir den Begriff folgendermaßen definieren: Die *Zuordnung* ist gegeben als Teilmenge $F \subset X \times Y$ mit der Eigenschaft, dass zu jedem $x \in X$ genau ein $y \in Y$ existiert mit $(x, y) \in F$. Dieses Element y wird dann als $f(x)$ bezeichnet.

- (2) Ist $f: X \rightarrow Y$ eine Abbildung, so nennen wir X den *Definitionsbereich* und Y den *Wertebereich* (oder manchmal das *Ziel*) der Abbildung f .

—

Eine *Funktion* ist nichts anderes als eine Abbildung (und der Funktionswert an einer Stelle ist dann der Wert an dem entsprechenden Element des Definitionsbereichs), allerdings wird der Begriff Funktion vor allem in spezielleren Situationen verwendet, beispielsweise spricht man in der Analysis oft von Funktionen $\mathbb{R} \rightarrow \mathbb{R}$. In der Linearen Algebra wird er kaum benutzt.

Auch wenn es von der Wortherkunft etwas miteinander zu tun hat, ist der hier definierte mathematische Begriff der Abbildung natürlich zu trennen vom Begriff einer Abbildung (im Sinne von Zeichnung, Foto, ...) der Alltagssprache.



Schematische Darstellung einer Abbildung. Die Elemente von X und Y sind durch Punkte dargestellt, die Zuordnung durch Pfeile. Die *Abbildungseigenschaft* besteht darin, dass bei jedem Punkt in X genau ein Pfeil startet, und alle Pfeile in Y enden.

BEISPIEL 3.23. (I) Beispiele für Funktionen sind

(i) $f: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$,

(ii) $g: \mathbb{Z} \rightarrow \mathbb{Z}, x \mapsto 2x - 3$,

(iii) $g: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto \begin{cases} x & x \geq 0 \\ -x & x < 0 \end{cases}$.

Die folgenden »Vorschriften« definieren keine Funktionen:

(i) $f: \mathbb{N} \rightarrow \mathbb{N}, x \mapsto 2x - 3$ (denn der Wert liegt nicht immer in dem angegebenen Wertebereich),

(ii) $g: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto \begin{cases} x & x \geq 0 \\ -x & x \leq 1 \end{cases}$
(denn für x zwischen 0 und 1 ist kein eindeutiger Funktionswert definiert),

(iii) $h: \mathbb{Q} \rightarrow \mathbb{Q}, x \mapsto \frac{1}{x}$
(denn für $x = 0$ ist kein Funktionswert definiert, weil $\frac{1}{0}$ keine rationale Zahl ist).

- (2) Eine Abbildung muss nicht durch eine »Formel« gegeben sein. Zum Beispiel ist $f: \{0\} \rightarrow \{0, 1\}$ mit

$$f(0) = \begin{cases} 0 & \text{wenn die Goldbachsche Vermutung wahr ist,} \\ 1 & \text{wenn die Goldbachsche Vermutung falsch ist,} \end{cases}$$

eine Abbildung. Allerdings kennt niemand den Wert $f(0)$.

- (3) Es gibt ziemlich »erstaunliche« Abbildungen. Zum Beispiel kann man zeigen, dass eine Abbildung $f: \mathbb{R} \rightarrow \mathbb{R}^2$ existiert, so dass jedes Element von $\mathbb{R}^2$ als Wert $f(x)$ auftritt.

- (4) Zu einer Abbildung $X \rightarrow Y$ (oder Funktion) können wir den *Funktionsgraph* betrachten, das ist die Teilmenge

$$\{(x, f(x)); x \in X\}$$

von $X \times Y$ (also gerade die Teilmenge F in Definition 3.22 (1); eine Funktion ist aus dieser Sichtweise dasselbe wie ihr Funktionsgraph). Da wir in der linearen Algebra meist Abbildungen zwischen »höherdimensionalen« Räumen betrachten, ist es meistens nicht möglich, den Funktionsgraph zu zeichnen; daher spielt er für uns eigentlich, anders als in der Analysis, keine Rolle.

◇

Statt vom Wertebereich einer Abbildung f spricht man manchmal auch vom *Bildbereich* von f . Diese Sprechweise ist aber weniger günstig, weil man unter dem *Bild* von f etwas anderes versteht (siehe Definition 3.25).

Die Menge aller Abbildungen von X nach Y bezeichnen wir mit $\text{Abb}(X, Y)$.

BEMERKUNG 3.24. Überlegen Sie sich, dass wir für Mengen X und Y die Menge $\text{Abb}(X, Y)$ mit dem Produkt $Y^X (= \prod_{x \in X} Y)$ identifizieren können. Wir können also eine Abbildung von X nach Y auch als ein Element des Produkts Y^X auffassen.

Hier zeigt sich, dass die Konvention, dass $Y^\emptyset$ genau ein Element haben soll, sinnvoll ist, denn die Menge $\text{Abb}(\emptyset, Y)$ hat nach unserer Definition für jedes Y genau ein Element. ◇

Ist X eine Menge, so bezeichnen wir mit $\text{id}_X: X \rightarrow X$ die Abbildung, die jedes $x \in X$ auf sich selbst abbildet: $x \mapsto x$ für alle $x \in X$. Diese Abbildung heißt die *identische Abbildung* oder *Identitätsabbildung* von X (oder manchmal einfach die *Identität* von/auf X).

DEFINITION 3.25. Sei $f: X \rightarrow Y$ eine Abbildung.

- (1) Das *Bild* von f ist

$$\text{Im}(f) = \{y \in Y; \text{es existiert } x \in X \text{ mit } f(x) = y\}.$$

- (2) Ist $Z \subseteq X$ eine Teilmenge, so nennt man

$$f(Z) = \{f(z); z \in Z\}$$

das *Bild von Z unter f* . Es gilt also $\text{Im}(f) = f(X)$.

- (3) Ist $Z \subseteq Y$ eine Teilmenge, so heißt

$$f^{-1}(Z) = \{x \in X; f(x) \in Z\}$$

das *Urbild von Z unter f* .

⊥

Für jede Abbildung $f: X \rightarrow Y$ gilt, dass $f^{-1}(Y) = X$, deshalb gibt es für das Urbild des gesamten Wertebereichs kein eigenes Symbol.

Wir nennen eine Abbildung f *konstant*, wenn alle Funktionswerte unter f gleich sind, mit anderen Worten, wenn $\text{Im}(f)$ nur ein einziges Element enthält (oder, in dem Fall, dass der Definitionsbereich von f die leere Menge ist, leer ist).

DEFINITION 3.26. Seien X, Y, Z Mengen und seien $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ Abbildungen. Wir definieren eine Abbildung $g \circ f: X \rightarrow Z$ durch

$$(g \circ f)(x) = g(f(x)).$$

(Diese Definition ist sinnvoll, da $f(x) \in Y$, so dass wir die Abbildung g auf dieses Element anwenden können.)

Die Abbildung $g \circ f$ heißt die *Verkettung*, *Verknüpfung* oder manchmal die *Komposition* der Abbildungen f und g . Das Symbol $g \circ f$ liest man auch als *g nach f* . $\dashv$

Achtung: Wenn man die Abbildungen f und g als Pfeile schreibt:

$$X \xrightarrow{f} Y \xrightarrow{g} Z$$

so steht f links und g rechts, aber die Verkettung der beiden Abbildungen ist $g \circ f$.

Ist $f: X \rightarrow Y$ eine Abbildung und $U \subseteq X$ eine Teilmenge, so bezeichnen wir mit $f|_U$ die Abbildung $U \rightarrow Y, x \mapsto f(x)$. Wir behalten also die Zuordnungsvorschrift unverändert bei und verkleinern lediglich den Definitionsbereich. Wir nennen $f|_U$ die *Einschränkung* der Abbildung f auf U .

Oftmals möchten wir Objekte »auflisten«, zum Beispiel, indem wir sie durchnummerieren (»Seien x_1, x_2, x_3 reelle Zahlen.«), je nach Situation kann es dabei um endlich viele oder unendlich viele Objekte gehen. Dabei möchten wir erlauben, dass Elemente mehrfach vorkommen (es wäre in dem gerade genannten Beispiel erlaubt, dass $x_1 = x_2$ ist), und ihre Reihenfolge festhalten. Daher ist der Mengenbegriff für diese Zwecke nicht ausreichend. Mithilfe des Abbildungsbegriffs können wir aber eine geeignete Definition machen:

DEFINITION 3.27. Sei X eine Menge. Eine *Familie* (oder: ein *System*) von Elementen aus X (mit *Indexmenge* I) ist gegeben durch eine Menge I und für jedes Element $i \in I$ ein Element $x_i \in X$. Formal kann man eine solche Familie als Abbildung $I \rightarrow X, i \mapsto x_i$, oder als ein Element des Produkts X^I betrachten bzw. definieren. Wir schreiben eine durch I indizierte Familie oft in der Form $(x_i)_{i \in I}$. $\dashv$

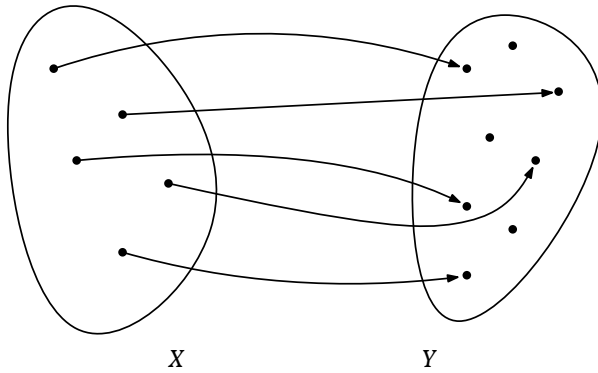
Der häufigste Fall wird sein, dass $I = \{1, \dots, n\}$ für eine natürliche Zahl n , oder dass $I = \mathbb{N}$ ist.

3.11. Injektive, surjektive und bijektive Abbildungen

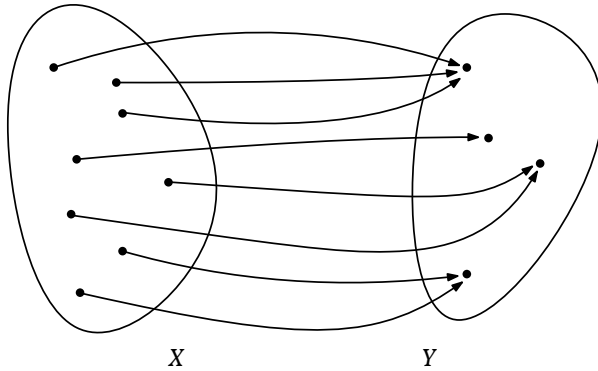
DEFINITION 3.28. Sei $f: X \rightarrow Y$ eine Abbildung.

- (1) Die Abbildung f heißt *injektiv*, wenn für alle $x, x' \in X$ mit $x \neq x'$ gilt, dass $f(x) \neq f(x')$. Man nennt f in diesem Fall auch eine *Injektion*.
- (2) Die Abbildung f heißt *surjektiv*, wenn für alle $y \in Y$ ein $x \in X$ existiert mit $f(x) = y$, mit anderen Worten: wenn $\text{Im}(f) = Y$. Man nennt f in diesem Fall auch eine *Surjektion*.
- (3) Die Abbildung f heißt *bijektiv*, wenn sie injektiv und surjektiv ist. Man nennt f in diesem Fall auch eine *Bijektion*.

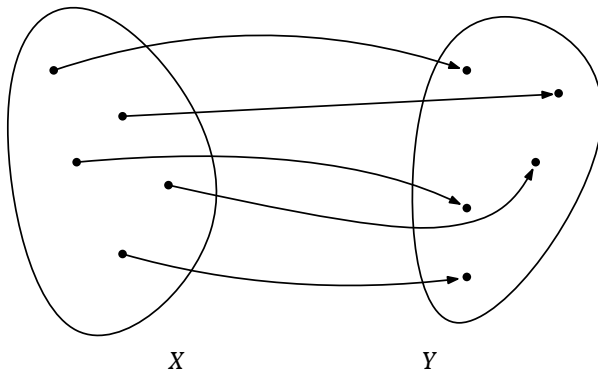
$\dashv$



Injektiv: Jedes Element von Y wird von höchstens einem Pfeil erreicht. Mit anderen Worten: Zwei verschiedene Pfeile dürfen nicht denselben Endpunkt haben.



Surjektiv: Jedes Element von Y wird von mindestens einem Pfeil erreicht.



Bijektiv: Jedes Element von Y wird von genau einem Pfeil erreicht. (Indem man die Richtungen aller Pfeile umdreht, erhält man deswegen eine Abbildung von Y nach X , die *Umkehrabbildung* der ursprünglichen Abbildung, siehe Definition 3.30.)

Eine injektive Abbildung nennt man manchmal auch eine *Einbettung*. Ist $X \subseteq Y$ eine Teilmenge, so ist die *Inklusionsabbildung* (oder kurz: *Inklusion*) $X \rightarrow Y, x \mapsto x$, eine injektive Abbildung.

BEISPIEL 3.29. Betrachte die folgenden Abbildungen:

- (i) $f: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$,
- (ii) $f: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^3$,
- (iii) $f: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}, x \mapsto x^2$,
- (iv) $f: \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}, x \mapsto x^2$,
- (v) $f: \mathbb{Z}_{\geq 0} \rightarrow \mathbb{Z}_{\geq 0}, x \mapsto x^2$.

- (1) Die folgenden Abbildungen sind injektiv: (ii), (iii), (v). Die anderen Abbildungen aus der Liste sind nicht injektiv.
- (2) Die folgenden Abbildungen sind surjektiv: (ii), (iv). Die anderen Abbildungen aus der Liste sind nicht surjektiv.

◇

DEFINITION 3.30. Sei $f: X \rightarrow Y$ eine Abbildung. Eine Abbildung $g: Y \rightarrow X$ heißt *Umkehrabbildung* von f , wenn $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$ gilt. $\dashv$

In der Situation ist dann also auch f eine Umkehrabbildung von g . Außerdem sind dann f und g automatisch bijektiv, wie der folgende Satz zeigt:

SATZ 3.31. Sei $f: X \rightarrow Y$ eine Abbildung. Es existiert genau dann eine Umkehrabbildung g von f , wenn die Abbildung f bijektiv ist. In diesem Fall ist die Umkehrabbildung von f eindeutig bestimmt.

BEWEIS. Wenn f eine Umkehrabbildung hat, dann ist f surjektiv (denn für $y \in Y$ gilt $f(g(y)) = y$) und injektiv (denn für $x, x' \in X$ mit $f(x) = f(x')$ gilt $x = g(f(x)) = g(f(x')) = x'$).

Sei nun f bijektiv. Gegeben $y \in Y$, so existiert $x \in X$ mit $f(x) = y$, weil f surjektiv ist. Zudem ist x eindeutig bestimmt, denn f ist injektiv. Wir setzen $g(y) := x$. Damit ist eine Abbildung $g: Y \rightarrow X$ definiert, und nach Konstruktion gilt $g(f(x)) = x$ für alle $x \in X$. Es ist noch zu zeigen, dass $f \circ g = \text{id}_Y$ ist. Sei dazu $y \in Y$, und sei x das eindeutig bestimmte Element von X mit $f(x) = y$. Dann gilt $f(g(y)) = f(x) = y$.

Alternativ kann man die Umkehrfunktion g von f über ihren Funktionsgraphen angeben; dieser ist

$$\{(f(x), x); x \in X\} \subseteq Y \times X.$$

Wir begründen noch, dass die Umkehrabbildung von f eindeutig bestimmt ist. Wegen der Bedingung $g \circ f = \text{id}_X$ ist g jedenfalls auf allen Elementen der Form $f(x)$ eindeutig bestimmt: Es muss $g(f(x)) = x$ gelten. Weil f surjektiv ist, hat aber jedes Element von Y diese Form. $\square$

Wir bezeichnen die (eindeutig bestimmte) Umkehrabbildung einer bijektiven Abbildung $f: X \rightarrow Y$ oft mit f^{-1} . In diesem Fall ist für $y \in Y$ also $f^{-1}(y)$ ein Element von X . Man muss hier etwas aufpassen, da man die Bezeichnung $f^{-1}(y)$ (für nicht notwendig bijektive Abbildungen) auch manchmal als Abkürzung für $f^{-1}(\{y\})$ benutzt, und dies ist nach Definition die *Teilmenge*

$$\{x \in X; f(x) \in \{y\}\} = \{x \in X; f(x) = y\}$$

von X . (Ist f bijektiv, so hat diese Teilmenge aber nur ein einziges Element.)

Man zeigt leicht die folgenden Aussagen:

LEMMA 3.32. (1) Die Verkettung von zwei injektiven Abbildungen ist eine injektive Abbildung.

(2) Die Verkettung von zwei surjektiven Abbildungen ist eine surjektive Abbildung.

(3) Die Verkettung von zwei bijektiven Abbildungen ist eine bijektive Abbildung.

LEMMA 3.33. Seien $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ Abbildungen. Dann gilt:

(1) Ist die Verkettung $g \circ f$ injektiv, dann ist f injektiv.

(2) Ist die Verkettung $g \circ f$ surjektiv, dann ist g surjektiv.

3.12. Vollständige Induktion

3.12.1. Das Grundprinzip der Induktion. In diesem Abschnitt behandeln wir das Prinzip der *vollständigen Induktion*, einer wichtigen Methode, um zu beweisen, dass eine Aussage für alle natürlichen Zahlen gilt. Betrachten Sie als Beispiel die Aussage

$$P(n) : \quad n^3 - n \text{ ist teilbar durch } 3,$$

die für jede natürliche Zahl n sinnvoll ist. Man kann die Aussage in konkreten Fällen überprüfen – zum Beispiel ist $7^3 - 7 = 343 - 7 = 336$ tatsächlich durch 3 teilbar. Aber wie würde man so etwas für *alle* n beweisen?

Das Induktionsprinzip gibt uns eine Methode dafür an die Hand. Formal ausgedrückt wird das in dem folgenden Satz:

SATZ 3.34 (Prinzip der vollständigen Induktion). Sei $P(n)$ eine Eigenschaft, die die natürliche Zahl n haben oder nicht haben kann (oder: sei für jede natürliche Zahl n eine Aussage $P(n)$ gegeben), so dass die folgenden beiden Bedingungen erfüllt sind:

- (a) $P(0)$ ist wahr, und
- (b) für jede natürliche Zahl $n \geq 1$ gilt: Wenn $P(n-1)$ wahr ist, dann ist auch $P(n)$ wahr.

Dann ist $P(n)$ wahr für alle $n \in \mathbb{N}$.

Oft wird die Aussage dieses Satzes als eines der Axiome eingesetzt, die für die natürlichen Zahlen angenommen werden. Jedenfalls ist klar, dass für seinen Beweis ein Axiom über die natürlichen Zahlen erforderlich ist, das darüber hinaus geht, dass jede natürliche Zahl $\neq 0$ einen eindeutig bestimmten Nachfolger hat. Wir wollen hier vom *Prinzip des kleinsten Elements* ausgehen, das vielleicht eingängiger ist als das Induktionsprinzip:

AXIOM 3.35 (Prinzip des kleinsten Elements). Jede nicht-leere Teilmenge von $\mathbb{N}$ besitzt ein kleinstes Element.

BEWEIS VON SATZ 3.34. Sei $M = \{n \in \mathbb{N}; P(n) \text{ ist falsch}\}$. Wir wollen zeigen, dass M die leere Menge ist, denn das bedeutet gerade, dass $P(n)$ für alle n wahr ist. Wenn M nicht leer ist, dann besitzt M nach dem Prinzip vom kleinsten Element ein kleinstes Element n . Weil nach Voraussetzung $P(0)$ wahr ist, kann nicht $n = 0$ sein. Dann ist aber $n-1 \notin M$, weil n das kleinste Element von M ist. Also ist $P(n-1)$ wahr, aber $P(n)$ falsch. Das ist ein Widerspruch zu Teil (b) der Voraussetzung. $\square$

Um mit dem Prinzip der vollständigen Induktion eine Aussage über natürliche Zahlen zu beweisen, muss man also die Eigenschaften (a) und (b) in Satz 3.34 nachweisen. Dabei nennt man Teil (a) den *Induktionsanfang* und Teil (b) den *Induktionsschritt*. In den meisten Beweisen hier im Skript und in der Literatur werden die Wörter Induktionsanfang und Induktionsschritt aber gar nicht mehr hingeschrieben, sondern es wird normalerweise nur darauf hingewiesen, dass ein Induktionsbeweis folgt und dann werden die Fälle $n = 0$ und $n > 0$ betrachtet.

BEISPIEL 3.36. *Behauptung.* Für alle natürlichen Zahlen n gilt, dass $n^3 - n$ durch 3 teilbar ist. *Beweis* durch vollständige Induktion nach n .

Induktionsanfang $n = 0$: In diesem Fall gilt $n^3 - n = 0^3 - 0 = 0$, und $0 = 3 \cdot 0$ ist ein Vielfaches von 3, mit anderen Worten: Durch 3 teilbar.

Induktionsschritt $n > 0$. Wir dürfen nun als Induktionsvoraussetzung annehmen, dass für $n-1$ die Aussage gilt, das bedeutet, dass $(n-1)^3 - (n-1)$ durch 3 teilbar ist. Wir müssen zeigen, dass $n^3 - n$ ebenso ein Vielfaches von 3 ist.

Wir rechnen¹⁶

$$(n-1)^3 - (n-1) = (n^3 - 3n^2 + 3n - 1) - (n-1) = n^3 - 3n^2 + 2n.$$

Um zu sehen, dass das hilfreich ist, schreiben wir den letzten Ausdruck noch ein bisschen um:

$$(n-1)^3 - (n-1) = n^3 - 3n^2 + 2n = n^3 - n + 3(n - n^2).$$

Nun ziehen wir auf beiden Seiten $3(n - n^2)$ ab und erhalten

$$n^3 - n = (n-1)^3 - (n-1) + 3(n - n^2).$$

Nach Induktionsvoraussetzung ist $(n-1)^3 - (n-1)$ ein Vielfaches von 3. Offensichtlich ist $3(n - n^2)$ ein Vielfaches von 3. Deswegen ist auch die Summe, und das ist gerade $n^3 - n$ ein Vielfaches von 3. $\square$

¹⁶Allgemein gilt $(a+b)^3 = a^3 + 3a^2b + 3ab^2 + b^3$. Das benutzen wir im ersten Schritt.

In diesem Fall ist es nicht allzu schwierig, statt des Induktionsbeweises einen direkten Beweis zu finden, der nur die Eigenschaften natürlicher Zahlen benutzt, die wir als bekannt voraussetzen. Haben Sie eine Idee?

Wie ist es, wenn wir den Exponenten durch eine andere Zahl als 3 ersetzen? Für $k \in \{2, 5, 7\}$ gilt: $n^k - n$ ist ein Vielfaches von k . Können Sie das beweisen? Für $k = 4$ ist diese Aussage nicht richtig. Finden Sie ein Gegenbeispiel? (Siehe Abschnitt 4.2.3.) $\diamond$

Statt im Induktionsschritt zu beweisen, dass für $n > 0$ die Aussage $P(n)$ aus $P(n-1)$ folgt, kann man natürlich genau so gut zeigen, dass für $n \geq 0$ die Aussage $P(n+1)$ aus $P(n)$ folgt. (Schreiben Sie den Beweis des vorherigen Beispiels in diesem Stil um. Es ist dann vielleicht besser, zuerst den Ausdruck $(n+1)^3 - (n-1)$ umzuformen.)

Wir haben das Induktionsprinzip schon benutzt (ohne es wirklich auszusprechen) im Beweis des Satzes, dass es unendlich viele Primzahlen gibt (Satz 3.6, vergleiche Beispiel 3.42 unten).

3.12.2. Varianten des Induktionsprinzips. Es gibt einige Abwandlungen des Prinzips der vollständigen Induktion:

SATZ 3.37. Sei $M \subseteq \mathbb{N}$ eine Teilmenge mit den Eigenschaften

- (a) $0 \in M$, und
- (b) für alle natürlichen Zahlen $n \geq 1$ mit $n-1 \in M$ gilt $n \in M$.

Dann gilt $M = \mathbb{N}$.

BEWEIS. Definiere $P(n)$ als die Aussage $P(n) : \Leftrightarrow n \in M$. Dann gilt $P(0)$, und für alle $n \in \mathbb{N}$ folgt aus $P(n)$, dass $P(n+1)$ gilt. Nach dem Prinzip der vollständigen Induktion ist $P(n)$ für alle $n \in \mathbb{N}$ eine wahre Aussage. Das besagt genau, dass $M = \mathbb{N}$. $\square$

BEMERKUNG 3.38. Statt des Axioms 3.35 könnte man auch einen der beiden Sätze als Axiom hernehmen, und dann die beiden anderen Aussagen daraus ableiten. $\diamond$

Manchmal sind die folgenden Varianten des Induktionsprinzips nützlich:

SATZ 3.39 (»Induktionsanfang bei n_0 «). Sei $n_0 \in \mathbb{N}$ eine natürliche Zahl. Sei für alle $n \geq n_0$ die Aussage $P(n)$ gegeben. Es gelte:

- (a) $P(n_0)$ ist wahr, und
- (b) für jede natürliche Zahl $n > n_0$ gilt: Wenn $P(n-1)$ wahr ist, dann ist auch $P(n)$ wahr.

Dann ist $P(n)$ wahr für alle natürlichen Zahlen $n \geq n_0$.

BEWEIS. Wir definieren die Aussage $P'(n)$ (für $n \in \mathbb{N}$) als $P(n+n_0)$. Aus dem üblichen Induktionsprinzip folgt dann, dass P' für alle $n \in \mathbb{N}$ wahr ist. Das bedeutet, dass $P(n)$ wahr ist für alle $n \geq n_0$. $\square$

BEISPIEL 3.40. Behauptung. Für alle natürlichen Zahlen $n \geq 5$ gilt $2^n > n^2$.

Beweis per Induktion nach n . Induktionsanfang: $n = 5$. Es gilt

$$2^5 = 32 > 25 = 5^2.$$

Induktionsschritt: $n > 5$. Wir müssen zeigen, dass aus der »Induktionsvoraussetzung« $2^{n-1} > (n-1)^2$ folgt, dass $2^n > n^2$. Wir rechnen dazu:

$$2^n = 2 \cdot 2^{n-1} > 2(n-1)^2 = 2n^2 - 4n + 2 = n^2 + (n^2 - 4n + 4) - 2 = n^2 + (n-2)^2 - 2 > n^2 + 3^2 - 2 > n^2,$$

wobei wir die Induktionsvoraussetzung (für das erste $>$) und die Abschätzung $n > 5$ (für das zweite $>$) benutzt haben.

Man beachte, dass die Aussage $2^n > n^2$ für $n \leq 4$ falsch ist. $\diamond$

SATZ 3.41 (»Induktionsvoraussetzung für $0, \dots, n$ «). Sei $P(n)$ eine Eigenschaft, die die natürliche Zahl n haben oder nicht haben kann (oder: sei $P(n)$ eine Aussage über alle natürlichen Zahlen n) mit den folgenden beiden Eigenschaften:

- (a) $P(0)$ ist wahr, und
- (b) für jede natürliche Zahl $n > 1$ gilt: Wenn $P(0), P(1), \dots, P(n-1)$ alle wahr sind, dann ist auch $P(n)$ wahr.

Dann ist $P(n)$ wahr für alle $n \in \mathbb{N}$.

Machen Sie sich klar, dass die Voraussetzung (b) in diesem Satz schwächer ist als in Satz 3.34, weil man auf mehr Informationen zurückgreifen kann, um $P(n)$ zu zeigen. Der Satz ist also a priori stärker als Satz 3.34. (Man spricht manchmal von *starker* Induktion.) Das heißt: Es ist klar, dass die Aussage von Satz 3.34 aus Satz 3.41 folgen würde, aber wir müssen erst beweisen, dass es auch umgekehrt der Fall ist. Zum Glück ist das nicht schwierig.

BEWEIS. Wir definieren für $n \in \mathbb{N}$ die Aussage $P'(n)$ durch

$$P'(n) : \Longleftrightarrow \quad \forall m \in \{0, \dots, n\} : P(m).$$

Dann erfüllt P' die Voraussetzungen des üblichen Induktionsprinzips (Satz 3.34) und daher gilt $P'(n)$ für alle n . Damit folgt auch $P(n)$ für alle n , wie gewünscht. $\square$

Manchmal kombiniert man auch die beiden vorherigen Varianten:

BEISPIEL 3.42. *Behauptung.* Jede natürliche Zahl $n \geq 2$ wird von einer Primzahl geteilt. (Vergleiche den Beweis von Satz 3.6)

Begründung. Induktionsanfang: $n = 2$. In diesem Fall ist die Aussage richtig, da 2 eine Primzahl und ein Teiler von sich selbst ist.

Induktionsschritt: $n > 2$. Wir verwenden das Prinzip von Satz 3.41 und dürfen daher annehmen, dass die Aussage für alle natürlichen Zahlen m mit $2 \leq m < n$ richtig ist. Wenn nun n eine Primzahl ist, dann ist die Sache klar. Ansonsten besitzt n eine Zerlegung als Produkt $n = ab$ mit $a, b > 1$. Dann gilt, dass $2 \leq a < n$, also wird a nach Induktionsvoraussetzung von einer Primzahl geteilt. Als Teiler von a ist diese auch ein Teiler von n , und wir sind fertig. $\diamond$

BEISPIEL 3.43. Auch einige Definitionen werden nach dem Induktionsprinzip vorgenommen, zum Beispiel definiert man (für Zahlen $a_0, \dots, a_n$ oder andere »Objekte«, für die eine Addition definiert ist) das Summensymbol $\sum_{i=0}^n a_i$, das die Summe der a_i bezeichnen soll. (Σ ist der griechische Großbuchstabe Sigma.) Anschaulich, aber informell möchte man definieren

$$\sum_{i=0}^n a_i = a_0 + a_1 + \dots + a_n,$$

aber die Pünktchen auf der rechten Seite sind kein mathematisches Objekt (wir haben dieses Symbol jedenfalls nicht definiert). Als formale Definition für das Summensymbol setzen wir

$$\sum_{i=0}^0 a_i = a_0, \quad \sum_{i=0}^n a_i = \sum_{i=0}^{n-1} a_i + a_n \quad (n > 0).$$

Das Induktionsprinzip zeigt, dass dann $\sum_{i=0}^n a_i$ für alle natürlichen Zahlen n definiert ist. Analog definiert man $\sum_{i=m}^n a_i$, $\sum_{i=-m}^n a_i$, usw.

Wenn in dem (Zahl-)Bereich, dessen Addition benutzt wird, das Kommutativgesetz und Assoziativgesetz der Addition gelten (also $a + b = b + a$ und $(a + b) + c = a + (b + c)$ für alle a, b, c), so kommt es auf die Reihenfolge der Summanden nicht an und man kann für jede endliche Menge I und Familie $(a_i)_{i \in I}$ die Summe $\sum_{i \in I} a_i$ definieren. Ist $I = \emptyset$ so spricht man auch von der *leeren* Summe und definiert ihren Wert als 0. Das ist eine sinnvolle Konvention; zum Beispiel gilt dann für jede Indexmenge I , Familie $(a_i)_{i \in I}$ und Teilmenge $J \subseteq I$:

$$\sum_{i \in I} a_i = \sum_{i \in J} a_i + \sum_{i \in I \setminus J} a_i.$$

Analog zum Summenzeichen definiert man das Produktzeichen (Das Symbol $\prod$ ist das griechische große Pi.):

$$\prod_{i=0}^0 a_i = a_0, \quad \prod_{i=0}^n a_i = \left(\prod_{i=0}^{n-1} a_i \right) \cdot a_n \quad (n > 0).$$

Auch hier kann man natürlich weitere Varianten definieren. Das *leere Produkt* $\prod_{i \in \emptyset} a_i$ hat per Konvention den Wert 1. $\diamond$

ERGÄNZUNG 3.44 (Teilbarkeit). Je nachdem, wie viel Lust/Interesse Sie daran haben, zu diesem Zeitpunkt auch »offensichtliche« Eigenschaften der ganzen Zahlen formal zu beweisen, können Sie diese ergänzende Bemerkung lesen oder überspringen.

DEFINITION 3.45. (1) Sei $n \in \mathbb{Z}$ eine ganze Zahl. Wir sagen, eine ganze Zahl $d \in \mathbb{Z}$ sei ein *Teiler* von n (oder d *teile* n), wenn eine ganze Zahl $k \in \mathbb{Z}$ existiert mit $n = dk$. Wir schreiben dann $d \mid n$, gesprochen » d teilt n «.

(2) Eine *Primzahl* ist eine ganze Zahl $p > 1$, deren einzige positive Teiler 1 und p selbst sind.

—

Wenn d kein Teiler von n ist, drücken wir das in Symbolen aus als $d \nmid n$. Zum Beispiel gilt für alle $n \in \mathbb{Z}$: $1 \mid n$, $-1 \mid n$, $n \mid n$, $n \mid 0$. Wenn $n \neq 0$ ist, dann gilt $0 \nmid n$.

Einige weitere Eigenschaften der Teilbarkeit sind in dem folgenden Satz gesammelt. Wir schreiben dabei $|a|$ für den *Absolutbetrag* (manchmal sagt man einfach *Betrag*) einer (reellen) Zahl a . Ist $a \geq 0$, so ist $|a| := a$, ist $a < 0$, so setzt man $|a| := -a$. In jedem Fall gilt also $|a| \geq 0$.

SATZ 3.46. Seien a, b, c ganze Zahlen.

- (1) Gilt $a \mid b$ und $b \neq 0$, so gilt $|a| \leq |b|$. Gilt $a \mid b$ und $a, b > 0$, so gilt $a \leq b$.
- (2) Aus $a \mid b$ und $b \mid c$ folgt $a \mid c$.
- (3) Gilt $a \mid b$ und $b \mid a$, so ist $a = b$ oder $a = -b$.
- (4) Gilt $a \mid b$ und $a \mid c$, so gilt $a \mid (b + c)$ und $a \mid (b - c)$.

BEWEIS. Alle diese Aussagen sind leicht zu beweisen. Versuchen Sie es als erstes selbst einmal!

zu (1). Wenn $b = ka$ und $b \neq 0$, so folgt $|k| |a| = |ka| = |b|$ und damit $|a| \leq |b|$. Wenn sogar a und b beide positiv sind, dann muss auch k positiv sein, d.h. $k \geq 1$. Daraus folgt $b = ak \geq a$.

zu (2). Gilt $b = ka$ und $c = lb$, so folgt $c = (kl)a$, also $a \mid c$.

zu (3). Gilt $b = ka$ und $a = lb$, so folgt $a = (kl)a$, also $a(kl - 1) = 0$. Weil das Produkt $a(kl - 1)$ Null ist, muss einer der Faktoren Null sein. Ist $a = 0$, so folgt $b = ka = 0$, also $a = b$. Ist

$kl = 1$, so muss $k = l = 1$ oder $k = l = -1$ sein, denn k und l sind ganze Zahlen. Daraus folgt die Behauptung.

zu (4). Gilt $b = ka$ und $c = la$, so folgt $b + c = (k + l)a$ und $b - c = (k - l)a$. Das zeigt die Behauptung. $\square$

SATZ 3.47 (Division mit Rest). Seien x und n ganze Zahlen, $n > 0$. Dann existieren eindeutig bestimmte ganze Zahlen q und r mit

$$x = qn + r, \quad 0 \leq r < n.$$

Man sagt, die Division mit Rest von x durch n ergebe q , Rest r .

BEWEIS. Zuerst zeigen wir die Eindeutigkeit: Gilt $x = qn + r = q'n + r'$ mit $0 \leq r, r' < n$, so folgt $(q - q')n + (r - r') = x - x = 0$. Weil $|r - r'| < n$ ist, folgt $|q - q'| \cdot |n| = |r - r'| < n$ und daraus $q - q' = 0$, also $q = q'$ und damit auch $r = r'$.

Es bleibt noch die Existenz der Zahlen q und r zu zeigen. Wir betrachten zunächst den Fall $x \geq 0$ und führen Induktion nach x . Für $x = 0$ setzen wir $q = r = 0$.

Ist $x > 0$, so können wir per Induktionsvoraussetzung annehmen, dass wir $x - 1$ in der Form $q'n + r'$ schreiben können, mit $0 \leq r' < n$. Ist r' sogar kleiner als $n - 1$, so setzen wir $q = q'$, $r = r' + 1$. Dann gilt $n = (n - 1) + 1 = q'n + r' + 1 = qn + r$ und $0 \leq r < n$, wie gewünscht. Sonst ist $r' = n - 1$, und dann können wir $q = q' + 1$ und $r = 0$ setzen.

Ist $x < 0$, so können wir das schon Bewiesene auf $-x$ anwenden und erhalten $-x = q'n + r'$ für Zahlen q', r' mit $0 \leq r' < n$. Ist $r' = 0$, so folgt $x = -q'n$ und wir setzen $q = -q'$, $r = 0$. Ist $r' > 0$, so gilt $x = -q'n - r' = (-q' - 1)n + n - r'$ und wir können $q = -q' - 1$ und $r = n - r' < n$ setzen. $\square$

Es gibt natürlich die Möglichkeit, den Satz auf den Fall $q < 0$ zu verallgemeinern (mit der Bedingung $0 \leq r < |q|$).

DEFINITION 3.48. Seien a, b ganze Zahlen. Eine Zahl $d \in \mathbb{Z}$ heißt ein *gemeinsamer Teiler* von a und b wenn $d \mid a$ und $d \mid b$.

Das größte Element der Menge aller gemeinsamen Teiler von a und b heißt der *größte gemeinsame Teiler* von a und b , geschrieben $\text{ggT}(a, b)$ (sofern nicht $a = b = 0$ gilt).

Sind $a = b = 0$, so ist jede ganze Zahl ein gemeinsamer Teiler von a und b , weswegen wir diesen Fall oben ausschließen. Wir definieren $\text{ggT}(0, 0) = 0$. $\dashv$

Da für jeden Teiler d einer Zahl $a \neq 0$ gilt, dass $|d| \leq |a|$, haben a und b nur endlich viele gemeinsame Teiler, wenn nicht $a = b = 0$ gilt, und daher hat die Menge der gemeinsamen Teiler tatsächlich ein größtes Element.

BEISPIEL 3.49. Die gemeinsamen Teiler von 24 und 45 sind $-3, -1, 1, 3$, also $\text{ggT}(24, 45) = 3$. $\diamond$

Wir halten noch die folgenden Eigenschaften des größten gemeinsamen Teilers fest:

LEMMA 3.50. Seien $a, b \in \mathbb{Z}$.

- (1) Es gilt $\text{ggT}(a, b) = \text{ggT}(-a, b) = \text{ggT}(a, -b) = \text{ggT}(-a, -b)$.
- (2) Es gilt $\text{ggT}(a, b) = \text{ggT}(a - b, b)$.

BEWEIS. zu (1). Dies ist klar, da die Teiler von a mit den Teilern von $-a$ übereinstimmen (und ebenso für b).

zu (2). Wir zeigen, dass die Menge der gemeinsamen Teiler von a und b übereinstimmt mit der Menge der gemeinsamen Teiler von $a - b$ und b . Daraus folgt die Behauptung.

Ist d eine Zahl mit $d \mid a, d \mid b$, so folgt mit Lemma 3.46 (4), dass $d \mid a - b$. Also ist d ein gemeinsamer Teiler von $a - b$ und b .

Ist umgekehrt d eine ganze Zahl mit $d \mid a - b$ und $d \mid b$, so folgt mit demselben Lemma, dass $d \mid (a - b) + b = a$. $\square$

Zum Berechnen des größten gemeinsamen Teilers verwendet man den [Euklidischen Algorithmus](#)¹⁷. $\square$ Ergänzung 3.44

ERGÄNZUNG 3.51 (Die eindeutige Primfaktorzerlegung in $\mathbb{Z}$). In dieser Ergänzung beweisen wir den Satz über die eindeutige Primfaktorzerlegung in den ganzen Zahlen. Der entscheidende Punkt ist die folgende Eigenschaft von Primzahlen:

SATZ 3.52 (Primeigenschaft). *Sei p eine Primzahl. Seien a, b ganze Zahlen, so dass p ein Teiler des Produkts ab ist. Dann ist p ein Teiler von a oder von b (oder von beiden).*

Um den Satz zu beweisen, benutzen wir das folgende Lemma.

LEMMA 3.53. *Sind $a, b \in \mathbb{Z}$, so existieren $x, y \in \mathbb{Z}$ mit*

$$\text{ggT}(a, b) = xa + yb.$$

BEWEIS. Wegen Lemma 3.50 (i) können wir gegebenenfalls a und/oder b durch ihr Negatives ersetzen und daher annehmen, dass $a, b \geq 0$. Ist eine der Zahlen a, b gleich Null, so ist die Aussage von vorneherein klar. Wir können daher sogar voraussetzen, dass $a, b > 0$.

Wir führen nun Induktion nach der Zahl $\max(a, b)$, dem Maximum von a und b . Ist dieses $= 1$, so gilt $a = b = 1$ und $\text{ggT}(a, b) = 1$, und wir nehmen $x = 1, y = 0$.

Sei nun $\max(a, b) > 1$. Ohne Einschränkung sei $a \geq b$ (sonst vertauschen wir einfach a und b .) Gilt $a = b$, so ist $\text{ggT}(a, b) = a$ und wir können wieder $x = 1, y = 0$ setzen. Sonst gilt $a - b, b \geq 1$ und $\max(a - b, b) < \max(a, b)$. Dann erhalten wir

$$\text{ggT}(a, b) = \text{ggT}(a - b, b) = x'(a - b) + y'b = x'a + (y' - x')b$$

für geeignete ganze Zahlen x', y' , wobei wir für die erste Gleichheit Lemma 3.50 und für die zweite die Induktionsvoraussetzung verwenden. Wir setzen also $x := x', y := y' - x'$ und erhalten das gewünschte Ergebnis. $\square$

BEWEIS VON SATZ 3.52. Sei p eine Primzahl und seien $a, b \in \mathbb{Z}$ mit $p \mid ab$. Wenn p nicht a teilt, dann gilt $\text{ggT}(p, a) = 1$ (denn die einzigen positiven Teiler von p sind ja nach Definition eine Primzahl 1 und p).

Nach dem Lemma können wir also $x, y \in \mathbb{Z}$ finden mit

$$1 = xp + ya,$$

also

$$b = xpb + yab.$$

Nun werden xpb (offensichtlich) und yab (nach Voraussetzung) von p geteilt, also auch ihre Summe: $p \mid b$. $\square$

¹⁷ https://de.wikipedia.org/wiki/Euklidischer_Algorithmus

BEMERKUNG 3.54. Es ist nicht schwer zu sehen, dass jede ganze Zahl $p > 1$ mit der Eigenschaft, dass $p \mid ab \Rightarrow p \mid a$ oder $p \mid b$ eine Primzahl ist.

In der Tat, wenn p diese Eigenschaft hat und $p = ab$ gilt, dann gilt ja erst recht $p \mid ab$, also $p \mid a$ oder $p \mid b$, und wenn zum Beispiel $p \mid a$ gilt, so folgt $ab = p \leq |a|$, also $b = 1$ oder $b = -1$ und damit $a = p$ oder $a = -p$. $\diamond$

KOROLLAR 3.55. Sei p eine Primzahl, und seien $a_1, \dots, a_n$ ganze Zahlen. Wenn p das Produkt $a_1 \cdots a_n$ teilt, dann teilt p (mindestens) einen der Faktoren a_i .

BEWEIS. Dies folgt aus dem Satz über die Primeigenschaft und einer einfachen Induktion. $\square$

SATZ 3.56. Sei $a \neq 0$ eine ganze Zahl. Dann gibt es $n \geq 0$ und Primzahlen $p_1, \dots, p_n \in \mathbb{Z}$, so dass

$$a = \varepsilon \prod_{i=1}^n p_i,$$

wobei $\varepsilon = 1$, wenn $a > 0$, und $\varepsilon = -1$, wenn $a < 0$ ist. Dabei sind die Primzahlen p_i bis auf ihre Reihenfolge eindeutig bestimmt.

Für den Fall, dass $a = 1$ oder $a = -1$ ist, verstehen wir die Aussage so, dass a als Produkt von $\varepsilon = 1$ (bzw. $\varepsilon = -1$) und dem »leeren Produkt«, dessen Wert 1 ist, geschrieben wird. Man beachte, dass die p_i in der Aussage des Satzes in der Regel nicht paarweise verschieden sein werden. Insbesondere geht die Eindeutigkeitsaussage darüber hinaus zu behaupten, dass die Menge $\{p_1, \dots, p_n\}$ der Primzahlen, die im Produkt überhaupt auftreten, eindeutig bestimmt sei; auch die Anzahl der Faktoren, die gleich einer gegebenen Primzahl sind, ist eindeutig bestimmt.

(Aus dem Satz 3.52 über die Primeigenschaft folgt, dass $\{p_1, \dots, p_n\}$ genau die Menge der Primzahlen ist, die a teilen.)

BEWEIS. Es ist klar, dass es ausreicht, den Fall $a > 0$ zu behandeln, weil sich der andere Fall leicht daraus ableiten lässt.

Existenz der Zerlegung. Wir zeigen die Existenz durch vollständige Induktion nach a . Für $a = 1$ ist nach der obigen Bemerkung über die Interpretation der Aussage in diesem Fall nichts mehr zu zeigen. Sei nun also $a > 1$. Wir haben in Beispiel 3.42 gesehen, dass es eine Primzahl p gibt, die a teilt, etwa $a = pk$. Nach Induktionsvoraussetzung lässt sich k als ein Produkt von Primzahlen schreiben. Indem wir den Faktor p hinzufügen, erhalten wir eine entsprechende Darstellung von a .

Eindeutigkeit der Zerlegung. Wir wenden wiederum Induktion nach a an, wobei der Fall $a = 1$ klar ist. Betrachten wir eine Gleichheit der Form

$$p_1 \cdots p_m = a = q_1 \cdots q_n$$

für Primzahlen p_i, q_j . Wir müssen zeigen, dass $m = n$ und dass jede Primzahl auf der linken Seite genauso oft vorkommt, wie auf der rechten Seite.

Da $a > 1$ ist, muss $m \geq 1$ (und $n \geq 1$) gelten. Nun gilt $p_m \mid p_1 \cdots p_m = q_1 \cdots q_n$, die Primzahl p_m teilt also das Produkt $q_1 \cdots q_n$. Wegen Korollar 3.55 gibt es ein $i \in \{1, \dots, n\}$ mit $p_m \mid q_i$. Weil p_m und q_i Primzahlen sind, muss $p_m = q_i$ gelten. Weil wir uns nur für die Eindeutigkeit der Faktoren *bis auf ihre Reihenfolge* interessieren, können wir die q 's so umnummerieren, dass $i = n$ ist; das vereinfacht die Notation ein bisschen.

Es gilt dann also $p_m = q_n$ und daher auch

$$p_1 \cdots p_{m-1} = q_1 \cdots q_{n-1}.$$

Auf dieses Produkt können wir die Induktionsvoraussetzung anwenden, das heißt: $m - 1 = n - 1$, und die Familien $p_1, \dots, p_{m-1}$ und $q_1, \dots, q_{n-1}$ unterscheiden sich höchstens durch ihre Reihenfolge. Damit sind wir fertig. $\square$

Vermutlich kennen Sie die Aussage dieses Satzes schon seit langem und halten sie für klar. Jedenfalls wird sie meistens in der Schule irgendwann angegeben. Oftmals wird aber nicht darauf hingewiesen, dass sie keineswegs selbstverständlich ist. Die Existenz einer solchen Zerlegung ist dabei noch recht eingängig, denn man kann ja, wie wir es auch im Beweis tun, jede Zahl immer weiter aufspalten, bis man es nur noch mit Primzahlen als Faktoren zu tun hat. Warum man aber eine Zahl wie 244 609 nur in der einen Weise $244\,609 = 331 \cdot 739$ (oder eben $= 739 \cdot 331$) als Produkt von Primzahlen geschrieben werden kann und es nicht noch andere Möglichkeiten geben könnte, ist nicht offensichtlich.

$\square$ Ergänzung 3.51

3.13. Endliche Mengen

In diesem Abschnitt definieren wir, wann eine Menge *endlich* ist, und wie viele Elemente sie dann hat. Wenn Sie (nicht ganz zu unrecht) denken, dass das ohnehin klar ist, können Sie ihn auch erstmal überspringen, weil die Beweise ein bisschen »technisch« sind. (Und wenn Sie sich irgendwann fragen, wie der Begriff *endlich* formal definiert wird, darauf zurückkommen.)

Oder Sie betrachten diese technischen Beweise als Fingerübungen, um den Umgang mit Injektivität, Surjektivität und Bijektivität von Abbildungen und der Beweismethode der Induktion zu trainieren.

Für jede natürliche Zahl n betrachten wir die Menge $[n] := \{1, \dots, n\}$. (Wenn $n = 0$, dann soll das bedeuten, dass $[0] = \emptyset$ die leere Menge bezeichnet.) Das ist für uns der Prototyp für eine endliche Menge mit n Elementen.

LEMMA 3.57. *Seien m, n endliche Zahlen.*

- (1) *Wenn es eine injektive Abbildung $[m] \rightarrow [n]$ gibt, dann gilt $m \leq n$.*
- (2) *Wenn es eine bijektive Abbildung $[m] \rightarrow [n]$ gibt, dann gilt $m = n$.*

BEWEIS. zu (1). Wir führen Induktion nach n . Ist $n = 0$, so ist $[n] = \emptyset$. Ist $M \rightarrow \emptyset$ irgendeine Abbildung, so muss auch $M = \emptyset$ gelten. Es folgt $[m] = \emptyset$ und damit $m = 0 = n$ (denn sonst wäre $1 \in [m]$).

Sei nun $n > 0$. Ist $m = 0$, so ist nichts zu zeigen, wir nehmen also auch an, dass $m > 0$. Sei $f: [m] \rightarrow [n]$ eine injektive Abbildung. Gilt $f([m-1]) \subseteq [n-1]$, so folgt $m-1 \leq n-1$ nach Induktionsvoraussetzung, also $m \leq n$. Sonst ist $n \in f([m-1])$ und wegen der Injektivität gibt es eine eindeutig bestimmte Zahl i mit $1 \leq i < m$ und $f(i) = n$. Andererseits muss (wieder wegen der Injektivität) $f(m) \neq n$ gelten, also $f(m) \in [n-1]$. Wir definieren die Abbildung $g: [m-1] \rightarrow [n-1]$ wie folgt: $g(i) = f(m)$, $g(j) = f(j)$ für $j \in [m-1] \setminus \{i\}$. Dann ist g eine Injektion $[m-1] \rightarrow [n-1]$ und nach Induktionsvoraussetzung folgt $m-1 \leq n-1$, also $m \leq n$.

zu (2). Dies folgt direkt aus Teil (1): Ist f eine Bijektion zwischen $[m]$ und $[n]$, so ist f injektiv, also $m \leq n$ nach Teil (1), und die Umkehrabbildung von f ist eine Injektion $[n] \rightarrow [m]$, also gilt auch $n \leq m$. $\square$

DEFINITION 3.58. Eine Menge M heißt *endlich*, wenn eine natürliche Zahl $n \geq 0$ und eine Bijektion $[n] \rightarrow M$ existiert. Wir sagen dann, M habe n Elemente und schreiben $\#M = n$. (Oft schreibt man auch $|M|$ statt $\#M$. Diese Zahl heißt auch die *Mächtigkeit* oder *Kardinalität* von M .) $\dashv$

Statt einer Bijektion $[n] \rightarrow M$ könnte man natürlich ebenso gut eine Bijektion $M \rightarrow [n]$ betrachten. Durch Übergang zur Umkehrabbildung kann man ja zwischen diesen beiden Standpunkten hin und her gehen.

Wegen des Lemmas kann es für gegebenes M höchstens für eine einzige Zahl n eine Bijektion wie in der Definition geben. Die Zahl n ist also durch M eindeutig bestimmt, so dass die Definition der Mächtigkeit überhaupt sinnvoll ist. Wir sagen, der Begriff der Mächtigkeit sei *wohldefiniert* (siehe Abschnitte 3.14.2, C.1.1).

Ist M eine Menge und M nicht endlich, dann sagen wir, M sei unendlich (und schreiben manchmal $\#M = \infty$). Wichtig ist aber zu beachten, dass es zwischen zwei unendlichen Mengen nicht unbedingt eine Bijektion gibt. Zum Beispiel gibt es keine Bijektion zwischen $\mathbb{Q}$ und $\mathbb{R}$. Siehe Abschnitt 3.15.

Der Mächtigkeitsbegriff hat die Eigenschaften, die man erwartet; allerdings müssen diese, formal betrachtet, natürlich erst einmal bewiesen werden, bevor man sie dann benutzen kann. Zum Beispiel:

LEMMA 3.59. Sei X eine Menge und $x \in X$ ein Element. Dann gilt

$$\#X = \#(X \setminus \{x\}) + 1.$$

BEWEIS. Sei etwa n die Mächtigkeit von $X \setminus \{x\}$, es gibt also eine Bijektion $f: [n] \rightarrow X \setminus \{x\}$. Wir definieren die Abbildung $g: [n+1] \rightarrow X$ durch $g(i) = f(i)$ für alle $i = 1, \dots, n$, und setzen $g(n+1) = x$. Dann ist g ebenfalls bijektiv, also $\#X = n+1 = \#(X \setminus \{x\}) + 1$. $\square$

Den Beweis des folgenden, ähnlichen Lemmas lassen wir als Übungsaufgabe.

LEMMA 3.60. Seien $X, Y \subseteq M$ endliche Teilmengen einer Menge M . Wenn $X \cap Y = \emptyset$, dann gilt $\#(X \cup Y) = \#X + \#Y$.

Eine andere Aussage, die Sie nicht überraschen wird, aber die eben auch eines Beweises bedarf:

LEMMA 3.61. Seien n eine natürliche Zahl, und $X_1, \dots, X_n$ nicht-leere Mengen. Dann ist das Produkt $\prod_{i=1}^n X_i$ nicht leer.

BEWEIS. Für $n = 1$ ist nichts zu zeigen, denn $\prod_{i=1}^1 X_i = X_1$ ist nach Voraussetzung nicht leer. (Und für $n = 0$ ist die Aussage auch richtig angesichts unserer Konvention, dass das Produkt mit leerer Indexmenge genau ein Element hat.)

Ist $n > 1$, so haben wir nach Definition

$$\left(\prod_{i=1}^{n-1} X_i \right) \times X_n.$$

Nach Induktionsvoraussetzung ist $\left(\prod_{i=1}^{n-1} X_i \right)$ nicht leer; sei x ein Element dieser Menge.

Nach Voraussetzung ist X_n nicht leer. Sei $x' \in X_n$. Dann ist $(x, x') \in \left(\prod_{i=1}^{n-1} X_i \right) \times X_n$, also ist auch diese Menge nicht leer. $\square$

KOROLLAR 3.62. Sei $f: X \rightarrow Y$ eine surjektive Abbildung von einer Menge X in eine endliche Menge Y . Dann existiert eine Abbildung $g: Y \rightarrow X$ mit $f \circ g = \text{id}_Y$.

BEWEIS. Nach Lemma 3.61 ist $\prod_{y \in Y} f^{-1}(\{y\}) \neq \emptyset$. Ein Element dieses Produkt gibt uns für jedes $y \in Y$ ein Element in X mit $f(x) = y$. Wir definieren $g(y) := x$.

(Wenn wir $\prod_{y \in Y} f^{-1}(\{y\}) \neq \emptyset$ als Teilmenge von $\prod_{y \in Y} X = \text{Abb}(Y, X)$ betrachten, dann ist jedes Element dieser Teilmenge eine Abbildung $Y \rightarrow X$ mit der gewünschten Eigenschaft.) $\square$

So einleuchtend die Aussagen von Lemma 3.61 und Korollar 3.62, und so kurz (wenn auch »technisch«) die Beweise sind: Die Beweise, die wir hier gegeben haben, benötigen die Voraussetzung, dass es sich um ein Produkt mit endlicher Indexmenge handelt beziehungsweise dass Y eine endliche Menge ist. Für beliebige Indexmengen und Mengen Y sind diese (zueinander äquivalenten) Aussagen genau der Inhalt des Auswahlaxioms, eines der Axiome der Mengenlehre. Siehe Anhang B.1.

LEMMA 3.63. Sei $f: X \rightarrow Y$ eine Abbildung zwischen endlichen Mengen.

- (1) Wenn f injektiv ist, dann gilt $\#X \leq \#Y$.
- (2) Wenn f surjektiv ist, dann gilt $\#X \geq \#Y$.

BEWEIS. Seien m und n die Mächtigkeiten von X und Y . Dann existieren Bijektionen $g: [m] \rightarrow X, h: [n] \rightarrow Y$. Die Verkettung $h^{-1} \circ f \circ g$ ist dann eine Injektion $[m] \rightarrow [n]$, und es folgt $m \leq n$ mit Lemma 3.57.

Für den Beweis von Teil (2) wenden wir Korollar 3.62 an und finden eine Abbildung $g: Y \rightarrow X$ mit $f \circ g = \text{id}_Y$. Dann ist g notwendigerweise injektiv (Lemma 3.33) und die Behauptung folgt aus Teil (1), angewandt auf g . $\square$

SATZ 3.64. Seien X, Y endliche Mengen mit $\#X = \#Y$. Sei $f: X \rightarrow Y$ eine Abbildung. Dann sind äquivalent:

- (i) Die Abbildung f ist injektiv.
- (ii) Die Abbildung f ist surjektiv.
- (iii) Die Abbildung f ist bijektiv.

BEWEIS. Es ist offenbar ausreichend, die Äquivalenz von (i) und (ii) zu zeigen, denn nach Definition gilt (iii) $\Leftrightarrow$ (i) und (ii).

Sei zunächst f injektiv. Ist f nicht surjektiv, dann existiert $y \in Y \setminus \text{Im}(f)$, also können wir f als Abbildung $X \rightarrow Y \setminus \{y\}$ betrachten (die natürlich ebenfalls injektiv ist). Das würde wegen Lemma 3.59 und Lemma 3.63 bedeuten, dass $\#X \leq \#Y - 1$, ein Widerspruch.

Ist andererseits f surjektiv, so gibt es nach Korollar 3.62 eine Abbildung $g: Y \rightarrow X$ mit $f \circ g = \text{id}_Y$. Lemma 3.33 zeigt, dass g injektiv ist. Aus Teil (1) folgt nun, dass g sogar bijektiv ist. Da $f \circ g$ und g bijektiv sind, ist auch f bijektiv. $\square$

ERGÄNZUNG 3.65 (Hilberts Hotel). Für eine unendliche Menge ist das Lemma nicht richtig. Überlegen Sie sich eine injektive Abbildung $\mathbb{N} \rightarrow \mathbb{N}$, die nicht surjektiv ist, und eine surjektive Abbildung $\mathbb{N} \rightarrow \mathbb{N}$, die nicht injektiv ist.

Die Existenz von injektiven, aber nicht surjektiven Abbildungen $\mathbb{N} \rightarrow \mathbb{N}$ wird in dem Gedankenexperiment des Hotels mit unendlich vielen Zimmern (zu jeder natürlichen Zahl n gibt es das Zimmer Nummer n) illustriert, das D. Hilbert in seiner Vorlesung »Über das Unendliche« 1924 beschrieben hat: Sind alle Zimmer mit Gästen belegt und kommt ein weiterer Gast an, so bittet der Hotelchef einfach jeden Gast, ein Zimmer weiter zu ziehen (von Zimmer n nach Zimmer $n + 1$). Dann ist Zimmer 0 frei.

Hilberts Hotel auf [Wikipedia](#)¹⁸ und auf Youtube: [Video von C. Spannagel](#)¹⁹, [Steven Strogatz and Hilbert's Infinite Hotel/WorldScienceFestival](#)²⁰ (englisch), [Video von J. Dekofsky/Ted-ED](#)²¹ (englisch).

Siehe Abschnitt 3.15.

□ Ergänzung 3.65

3.14. Relationen *

3.14.1. Definition. Die Definitionen aus diesem Abschnitt kann man zunächst überspringen. Thematisch gehören sie aber dennoch ins Grundlagenkapitel und sind daher hier einsortiert. Zum Teil kommen die Begriffe in den Ergänzungen vor, jedenfalls implizit. In der Linearen Algebra 2 werden wir dann noch einmal darauf zurückkommen.

DEFINITION 3.66. Eine *Relation* zwischen zwei Mengen X und Y ist eine Teilmenge $R \subseteq X \times Y$. ⊢

Sofern man von einer Relation zwischen X und Y keine weiteren Eigenschaften kennt, ist der Begriff eher uninteressant (und es gäbe keinen Grund, dafür eine eigene Bezeichnung einzuführen). Der Sinn der Sache ist, Relationen zu betrachten, die durch zusätzliche Eigenschaften besonders ausgezeichnet sind. Dabei gibt es mehrere Arten von Eigenschaften, die es zu betrachten lohnt.

Ein wichtiges Beispiel haben wir bereits gesehen: Eine Abbildung $X \rightarrow Y$ ist eine Relation zwischen X und Y mit der Eigenschaft, dass für jedes $x \in X$ genau ein $y \in Y$ existiert, so dass $(x, y) \in R$.

In den folgenden beiden Abschnitten betrachten wir bestimmte Relationen zwischen einer Menge X und sich selbst, also Teilmengen von $X \times X$. Einerseits die *Äquivalenzrelationen*, die dazu dienen, Objekte zusammenzufassen, die zwar nicht unbedingt gleich, aber doch gleichartig sind, was gewisse Eigenschaften angeht. Andererseits (*partielle*) *Ordnungen*, die beschreiben, wie man Elemente einer Menge vergleichen und anordnen kann.

3.14.2. Äquivalenzrelationen. Sei $R \subseteq X \times X$ eine Relation. Oft wählt man ein Symbol, zum Beispiel $\sim$ und definiert $x \sim y$ als $(x, y) \in R$. Wir sagen dann auch, dass $\sim$ eine Relation auf X sei.

DEFINITION 3.67. Sei $\sim$ eine Relation auf einer Menge X .

- (1) Die Relation $\sim$ heißt *reflexiv*, wenn für alle $x \in X$ gilt: $x \sim x$.
- (2) Die Relation $\sim$ heißt *symmetrisch*, wenn für alle $x, y \in X$ genau dann $x \sim y$ gilt, wenn $y \sim x$ gilt.
- (3) Die Relation $\sim$ heißt *transitiv*, wenn für alle $x, y, z \in X$ mit $x \sim y$ und $y \sim z$ gilt, dass $x \sim z$.
- (4) Die Relation $\sim$ heißt eine *Äquivalenzrelation*, wenn sie reflexiv, symmetrisch und transitiv ist.

⊢

BEISPIEL 3.68. Ist X eine Menge, so ist Gleichheit $=$ eine Äquivalenzrelation. ◇

¹⁸https://de.wikipedia.org/wiki/Hilberts_Hotel

¹⁹<https://www.youtube.com/watch?v=XTsaZRKx9UI>

²⁰<https://www.youtube.com/watch?v=wE9f16tUWhc>

²¹https://www.youtube.com/watch?v=Uj3_KqkI9Zo

BEISPIEL 3.69. Sei $X \subset \mathbb{Z} \times \mathbb{Z}$ die Menge aller Paare von ganzen Zahlen (a, b) mit $b \neq 0$. Wir definieren für $(a, b), (c, d) \in X$:

$$(a, b) \sim (c, d) \iff ad = bc$$

Dies ist offenbar eine Relation zwischen X und X . Es ist nicht schwierig nachzuprüfen, dass es sich um eine Äquivalenzrelation handelt. (Sie sollten das zur Übung tun.) Siehe Beispiel 3.72 für die Fortsetzung dieses Beispiels. Wenn Ihnen dieses Beispiel ziemlich künstlich vorkommt, dann ist das in Ordnung, aber Sie sollten gerade dann auch Beispiel 3.72 bis zum Ende lesen. $\diamond$

Sei $\sim$ eine Äquivalenzrelation auf X . Wir sagen dann, $y \in X$ sei *äquivalent* zu x , wenn $x \sim y$ gilt. Für $x \in X$ nennen wir

$$\{y \in X; y \sim x\},$$

die Menge aller Elemente, die bezüglich $\sim$ in Relation zu x stehen, die *Äquivalenzklasse* von x . Oft schreibt man $[x]$ für die Äquivalenzklasse von x .

BEISPIEL 3.70. Sei $X = \mathbb{Z}$ die Menge der ganzen Zahlen. Wir definieren für $x, y \in \mathbb{Z}$:

$$x \sim y \iff x - y \text{ ist durch } 3 \text{ teilbar.}$$

Dies ist offenbar eine Relation zwischen $\mathbb{Z}$ und $\mathbb{Z}$. Man prüft leicht nach, dass es sich um eine Äquivalenzrelation handelt. (Sie sollten das zur Übung tun.) Oft schreibt man $x \equiv y \pmod{3}$ statt $x \sim y$.

In diesem Fall gibt es drei Äquivalenzklassen: Erstens die Teilmenge von $\mathbb{Z}$, die aus allen Zahlen besteht, die durch 3 teilbar sind; zweitens die Teilmenge aller Zahlen, die bei Division durch 3 Rest 1 haben. Und drittens die Teilmenge derjenigen Zahlen, die bei Division durch 3 Rest 2 haben. Die Äquivalenzklassen für diese spezielle Äquivalenzrelation nennt man auch *Restklassen modulo 3*.

Siehe Beispiel 3.73 für die Fortsetzung dieses Beispiels. $\diamond$

Aus Reflexivität, Symmetrie und Transitivität der Äquivalenzrelation folgt, dass je zwei Elemente einer Äquivalenzklasse zueinander äquivalent sind. Außerdem gilt: Ist y in der Äquivalenzklasse von x enthalten und gilt $y \sim z$, so liegt auch z in der Äquivalenzklasse von x ; das ist einfach eine Umformulierung der Transitivität.

LEMMA 3.71. Sei $\sim$ eine Äquivalenzrelation auf X . Seien $A, B \subseteq X$ Äquivalenzklassen. Dann gilt entweder $A = B$ oder $A \cap B = \emptyset$.

Mit anderen Worten: Zwei Äquivalenzklassen sind entweder gleich (die gleiche Teilmenge von X) oder disjunkt.

BEWEIS. Wir zeigen, dass aus $A \cap B \neq \emptyset$ folgt, dass $A = B$. Sei dazu $x \in A \cap B$. Ist dann $y \in A$, so folgt $x \sim y$ (denn $x, y \in A$) und damit $y \in B$ (denn $x \in B$ und B ist eine Äquivalenzklasse). $\square$

Ist $A \subseteq X$ eine Äquivalenzklasse (bezüglich $\sim$) und ist $x \in A$, so nennt man x einen *Repräsentanten* der Äquivalenzklasse. Dann gilt $A = \{y \in X; y \sim x\}$.

Wir bezeichnen mit $X/\sim$ die Menge aller Äquivalenzklassen. Dies ist also eine Menge, deren Elemente Teilmengen von X sind. Die Abbildung $X \rightarrow X/\sim, x \mapsto [x]$, die jedes Element von X auf seine Äquivalenzklasse abbildet, bezeichnet man auch als die *kanonische Projektion*. Per Definition ist diese Abbildung surjektiv, aber in aller Regel nicht injektiv: Denn für $x \sim y$ gilt ja $[x] = [y]$, also haben äquivalente Elemente x und y dasselbe Bild unter dieser Abbildung.

Um eine Abbildung f von der Menge $X/\sim$ in eine Menge Y zu definieren, gibt man $f([x])$ oft an, indem man x verwendet. Weil für $x \sim y$ aber $[x] = [y]$ gilt und deswegen $f([x]) = f([y])$ gelten muss, ist das problematisch. Wenn die »Formel« für $f([x])$ wirklich von x abhängt und für y mit $y \sim x$ ein anderes Ergebnis liefern würde, dann hätten wir gar keine Zuordnung definiert. Siehe Beispiel 3.72 für Beispiele. Wenn die gegebene Vorschrift für alle Elemente der Äquivalenzklasse dasselbe Ergebnis liefert, also unabhängig ist von der Wahl des Repräsentanten der Äquivalenzklasse, dann sagt man, die Vorschrift sei *wohldefiniert*. Oft sagt man auch, die Abbildung sei wohldefiniert.

BEISPIEL 3.72. Wir nehmen wieder die Notation von Beispiel 3.69 auf. Wir hatten eine Äquivalenzrelation $\sim$ auf der Menge aller Paare (a, b) von ganzen Zahlen mit $b \neq 0$ definiert. Wir bezeichnen mit Q die Menge der Äquivalenzklassen, also $Q = (\mathbb{Z} \times (\mathbb{Z} \setminus \{0\})) / \sim$. Wir wollen als erstes die im vorherigen Absatz angesprochene Problematik (Stichwort »wohldefiniert«) aufgreifen.

Betrachten wir die Vorschrift $[(a, b)] \mapsto a + b$. Definiert diese eine Abbildung $f: Q \rightarrow \mathbb{Z}$? Nein, denn die Vorschrift ist nicht wohldefiniert! Es gilt nämlich zum Beispiel $(1, 1) \sim (2, 2)$, also $[(1, 1)] = [(2, 2)]$, aber nicht $1 + 1 = 2 + 2$.

Die Vorschrift $[(a, b)] \mapsto \frac{a}{b}$ ist hingegen wohldefiniert, denn wenn $(a, b) \sim (c, d)$, dann bedeutet das $ad = cb$, also tatsächlich $\frac{a}{b} = \frac{c}{d}$. Wir erhalten so eine Abbildung $i: Q \rightarrow \mathbb{Q}$.

Die folgende Vorschrift ist ein anderes Beispiel für eine wohldefinierte Zuordnung, und zwar ordnen wir jedem Paar von Elementen in Q ein neues Element in Q zu:

$$M: Q \times Q \longrightarrow Q, \quad ([a, b], [c, d]) \mapsto [(ac, bd)].$$

Wir haben hier die Wohldefiniertheit schon vorweggenommen und so getan, als hätten wir schon eine Abbildung $Q \times Q \rightarrow Q$ in der Hand. Wir müssen sie aber natürlich überprüfen. Sei also $[(a, b)] = [(a', b')]$, das bedeutet $ab' = ba'$, und $[(c, d)] = [(c', d')]$, das heißt $cd' = dc'$. Dann gilt tatsächlich $[(ac, bd)] = [(a'c', b'd')]$, denn das heißt ja genau, dass $(ac, bd) \sim (a'c', b'd')$, und wir haben

$$acb'd' = (ab')(cd') = (ba')(dc') = bda'c'.$$

In ähnlicher Weise definieren wir eine Abbildung $A: Q \times Q \longrightarrow Q$:

$$A: Q \times Q \longrightarrow Q, \quad ([a, b], [c, d]) \mapsto [(ad + bc, bd)].$$

Wieder muss man überprüfen, dass diese Vorschrift wohldefiniert ist, also dass im Fall $[(a, b)] = [(a', b')]$, das bedeutet $ab' = ba'$, und $[(c, d)] = [(c', d')]$ auch $[(ad + bc, bd)] = [(a'd' + b'c', b'd')]$ gilt. Führen Sie diese Rechnung durch.

Um das Beispiel abzuschließen, führen wir noch die Notationen

$$[(a, b)] \cdot [(c, d)] := M([(a, b)], [(c, d)]) \quad \text{und} \quad [(a, b)] + [(c, d)] := A([(a, b)], [(c, d)])$$

ein, wir betrachten also M und A als Multiplikation und Addition auf der Menge Q .

Wir kommen nun noch einmal auf die Abbildung $i: Q \rightarrow \mathbb{Q}$ zurück. Sie hat die folgenden Eigenschaften:

- (1) i ist bijektiv,
- (2) $i([(a, b)] + [(c, d)]) = i([(ad + bc, bd)]) = \frac{ad+bc}{bd} = \frac{a}{b} + \frac{c}{d} = i([(a, b)]) + i([(c, d)])$
- (3) $i([(a, b)] \cdot [(c, d)]) = i([(ac, bd)]) = \frac{ac}{bd} = \frac{a}{b} \cdot \frac{c}{d} = i([(a, b)]) \cdot i([(c, d)])$

Zur ersten Aussage: Die Surjektivität ist klar, denn jedes Element von $\mathbb{Q}$ hat die Form $\frac{a}{b}$ für geeignete ganze Zahlen a und $b \neq 0$, und $\frac{a}{b} = i([(a, b)])$. Zur Injektivität: Wenn $i([(a, b)]) = i([(c, d)])$, also $\frac{a}{b} = \frac{c}{d}$, dann gilt $ad = bc$. Das bedeutet aber $(a, b) \sim (c, d)$, also $[(a, b)] = [(c, d)]$.

Das bedeutet, dass die Abbildung i eine Identifikation von Q und $\mathbb{Q}$ erlaubt, die mit Addition und Multiplikation verträglich ist. Eine andere Sichtweise ist, dass wir eine Konstruktion der

rationalen Zahlen ausgehend von den ganzen Zahlen kennengelernt haben, denn wenn wir $\mathbb{Q}$ noch nicht kennen würden, ist die Menge $\mathbb{Q}$ mit den Rechenoperationen, die wir definiert haben, ein vollwertiger Ersatz. $\diamond$

BEISPIEL 3.73. Wir nehmen wieder die Notation von Beispiel 3.70 auf.

Die Ausführungen hier sind etwas skizzenhaft. Betrachten Sie das Beispiel als erweiterte Übungsaufgabe und/oder melden Sie sich, wenn Sie gerne weitere Details hätten.

Wie üblich bezeichnen wir die Äquivalenzklasse von x mit $[x]$. Wir hatten schon festgestellt, dass $\mathbb{Z}/\sim = \{[0], [1], [2]\}$ gilt. (Die eckigen Klammern haben hier eine andere Bedeutung als in Abschnitt 3.13.)

Eine interessante Beobachtung ist, dass man mit Restklassen modulo 3 (das war unser Name für die Äquivalenzklassen in diesem Beispiel) ähnlich rechnen kann wie ganzen Zahlen: Für alle $x, x', y, y' \in \mathbb{Z}$ mit $x \sim x'$ und $y \sim y'$ gilt

$$x + y \sim x' + y', \quad xy \sim x'y',$$

also

$$[x + y] = [x' + y'], \quad [xy] = [x'y'].$$

Ähnlich wie in Beispiel 3.72 (aber sogar noch einfacher) haben wir also eine wohldefinierte Addition und Multiplikation auf der Menge der Äquivalenzklassen, die gegeben ist durch

$$[x] + [y] := [x + y], \quad [x] \cdot [y] := [xy].$$

Wir dürfen also, salopp gesagt, die Rechenzeichen $+$ und $\cdot$ beliebig in die oder aus den eckigen Klammern ziehen.

Ein paar Beispielrechnungen:

$$\begin{aligned} [10] &= [1] && \text{denn 10 hat bei Division durch 3 Rest 1,} \\ [100] &= [1] && \text{denn 100 hat bei Division durch 3 Rest 1,} \\ &&& \text{oder wir rechnen } [100] = [10 \cdot 10] = [10] \cdot [10] = [1] \cdot [1] = [1 \cdot 1] = [1], \\ [10^i] &= [1] && \text{für alle } i \in \mathbb{N}, \text{ mit einem ähnlichen Argument,} \end{aligned}$$

und

$$[752] = [7 + 5 + 2] = [14] = [2]$$

denn

$$[752] = [7 \cdot 100 + 5 \cdot 10 + 2] = [7] \cdot [100] + [5] \cdot [10] + [2] = [7] + [5] + [2] = [7 + 5 + 2].$$

Die letzte Rechnung lässt sich offensichtlich auf beliebige natürliche Zahlen verallgemeinern und zeigt: Eine natürliche Zahl hat denselben Rest bei Division durch 3 wie ihre Quersumme (die Summe aller ihrer Ziffern).

Insbesondere haben wir das bekannte Kriterium für Teilbarkeit durch 3 bewiesen: Eine natürliche Zahl ist genau dann durch 3 teilbar, wenn ihre Quersumme durch 3 teilbar ist.

Können Sie diese Betrachtungen auf den Fall $n = 9$ übertragen? Was ist zum Beispiel im Fall $n = 7$ anders?

Vergleiche Abschnitt 4.2.1. $\diamond$

Zum Schluss wollen wir noch zwei etwas andere Sichtweisen auf den Begriff der Äquivalenzrelation angeben (was hoffentlich unterstreicht, dass es sich vom Prinzip her um etwas sehr einfaches handelt).

BEMERKUNG 3.74. Ist $f: X \rightarrow Y$ eine (surjektive) Abbildung, so wird durch

$$x \sim x' \iff f(x) = f(x')$$

eine Äquivalenzrelation auf X definiert. Umgekehrt hat jede Äquivalenzrelation diese Form, denn man kann für f die kanonische Projektion auf die Menge aller Äquivalenzklassen verwenden. $\diamond$

BEMERKUNG 3.75. Ist $\sim$ eine Äquivalenzrelation, so bilden die Äquivalenzklassen eine Familie von paarweise disjunkten Teilmengen von X , deren Vereinigung ganz X ist.

Ist umgekehrt eine solche Darstellung von $X = \bigcup_{i \in I} X_i$ als Vereinigung von paarweise disjunkten nicht-leeren Teilmengen X_i gegeben, so können wir eine Äquivalenzrelation auf X definieren durch

$$x \sim x' \iff \text{es gibt } i \text{ mit } x, x' \in X_i.$$

$\diamond$

3.14.3. Partielle und totale Ordnungen.

DEFINITION 3.76. Sei X eine Menge und $\preceq$ eine Relation zwischen X und sich selbst.

- (1) Die Relation $\preceq$ heißt *antisymmetrisch*, wenn für alle $x, y \in X$ mit $x \preceq y$ und $y \preceq x$ gilt, dass $x = y$.
- (2) Die Relation $\preceq$ heißt eine *partielle Ordnung* (oder *Halbordnung* oder manchmal einfach *Ordnung*), wenn sie reflexiv, transitiv und antisymmetrisch ist.

$\dashv$

BEISPIEL 3.77. Sei M eine Menge, und sei $P(M)$ die Potenzmenge von M , also die Menge aller Teilmengen von M . Die Relation $\subseteq$ der Inklusion von Teilmengen ist dann eine partielle Ordnung auf $P(M)$ (und natürlich auch auf allen Teilmengen von $P(M)$). $\diamond$

Ein wichtiger Punkt (und das soll durch das Wort *partiell* betont werden) ist, dass es in der Situation der Definition Elemente x, y geben kann, für die *weder* $x \preceq y$ *noch* $y \preceq x$ gilt. In der Situation des Beispiels ist das in der Tat klar: Sind $A, B \subseteq M$ Teilmengen, dann kann es passieren, dass weder $A \subseteq B$ noch $B \subseteq A$ gilt.

DEFINITION 3.78. Sei X eine Menge und $\preceq$ eine Relation zwischen X und sich selbst.

- (1) Die Relation $\preceq$ heißt *total*, wenn für alle $x, y \in X$ gilt, dass $x \preceq y$ oder $y \preceq x$.
- (2) Eine Relation $\preceq$ heißt eine *totale Ordnung* (oder *lineare Ordnung*), wenn sie reflexiv, transitiv, antisymmetrisch und total ist.

$\dashv$

Mit anderen Worten: Eine totale Ordnung ist eine partielle Ordnung, bezüglich derer je zwei Elemente stets »vergleichbar« sind, d.h. in Relation stehen (in der einen oder anderen Reihenfolge).

BEISPIEL 3.79. Die übliche $\leq$ -Relation ist eine totale Ordnung auf der Menge der reellen Zahlen (und ebenso auf $\mathbb{Q}, \mathbb{Z}, \mathbb{N}$). $\diamond$

DEFINITION 3.80. Sei $\preceq$ eine partielle Ordnung auf X .

- (1) Ein Element $x \in X$ heißt *minimales Element* (bezüglich $\preceq$), wenn für alle $y \in X$ mit $y \preceq x$ gilt: $y = x$.
- (2) Ein Element $x \in X$ heißt *kleinstes Element* (bezüglich $\preceq$), wenn für alle $y \in X$ gilt: $x \preceq y$.
- (3) Ein Element $x \in X$ heißt *maximales Element* (bezüglich $\preceq$), wenn für alle $y \in X$ mit $x \preceq y$ gilt: $y = x$.
- (4) Ein Element $x \in X$ heißt *größtes Element* (bezüglich $\preceq$), wenn für alle $y \in X$ gilt: $y \preceq x$.

⊢

Im allgemeinen muss es weder minimale noch maximale Elemente (und erst recht kein kleinstes oder größtes Element) geben; betrachten Sie zum Beispiel die $\leq$ -Ordnung auf $\mathbb{Z}$.

Wenn es ein kleinstes Element (bezüglich einer partiellen Ordnung) gibt, dann ist dieses eindeutig bestimmt (und ist ein minimales Element). Wenn es ein eindeutig bestimmtes minimales Element gibt, dann ist dieses das kleinste Element. Entsprechendes gilt für maximale Elemente und das größte Element.

Wenn $\preceq$ eine totale Ordnung ist, dann fallen die Begriffe des minimalen Elements und des kleinsten Elements zusammen; ebenso die Begriffe des maximalen und des größten Elements.

BEISPIEL 3.81. Wir betrachten auf $\mathbb{N}$ die Relation $d \mid n$ der Teilbarkeit. Dies ist eine partielle Ordnung. Für alle $n \in \mathbb{N}$ gilt $1 \mid n$, also ist 1 das kleinste Element in $\mathbb{N}$ bezüglich der Teilbarkeitsordnung. Weil $n \mid 0$ für alle n gilt, ist 0 das größte Element in $\mathbb{N}$ für diese partielle Ordnung!

Für natürliche Zahlen a, b ist $\text{ggT}(a, b)$ das größte Element (bezüglich Teilbarkeit) der Menge aller positiven gemeinsamen Teiler von a und b . Mit dieser Beschreibung ist es nicht notwendig, den Fall $a = b = 0$ gesondert zu betrachten.

Um das zu beweisen, müssen wir zeigen, dass jeder gemeinsame Teiler d von a und b auch ein Teiler von $\text{ggT}(a, b)$ ist. Das folgt aus Lemma 3.53, das besagt, dass wir $\text{ggT}(a, b) = xa + yb$ schreiben können (mit ganzen Zahlen x und y).

In der Menge $\{2, 3, 5, 7, 11, \dots\}$ der Primzahlen ist jedes Element gleichzeitig minimal und maximal bezüglich Teilbarkeit. Es gibt weder ein kleinstes noch ein größtes Element. $\diamond$

3.15. Mächtigkeit von Mengen *

Für endliche Mengen haben wir die Mächtigkeit in Abschnitt 3.13 definiert. Wir wollen hier die Erweiterung dieses Begriffs auf den Fall unendlicher Mengen skizzieren. (Für den Moment ohne Beweise.)

In Kapitel 5 in Soergels Skript [So-AZT] und in dem (englischen) Buch [Hu] (Introduction, Abschnitt 8) von Hungerford finden Sie mehr Informationen zu diesem Thema.

DEFINITION 3.82. Wir nennen zwei Mengen M, M' *gleichmächtig*, wenn eine bijektive Abbildung $M \rightarrow M'$ existiert. Wir schreiben dann $\#M = \#M'$

⊢

Achtung: Zwei unendliche Mengen sind nicht unbedingt gleichmächtig (siehe unten). Aus der unpräzisen Aussage $\#M = \infty$, $\#M' = \infty$ lässt sich also nicht die Gleichheit $\#M = \#M'$ im Sinne dieser Definition folgern.

Wir drücken uns hier darum, genau zu sagen, was für ein Objekt $\#M$ eigentlich ist. (Man nennt diese Ausdrücke **Kardinalzahlen**²².) Es soll für uns genügen zu wissen, wie man Kardinalzahlen vergleicht.

DEFINITION 3.83. Eine Menge M heißt *abzählbar* (oder genauer *abzählbar unendlich*), wenn M gleichmächtig ist zur Menge $\mathbb{N}$ der natürlichen Zahlen. Man schreibt dann auch $\#M = \aleph_0$.
 $\dashv$

($\aleph$, ausgesprochen Aleph, ist der erste Buchstabe des hebräischen Alphabets.)

Wenn man von einer Menge M sagt, sie sei *höchstens abzählbar*, so meint man, dass M endlich oder abzählbar unendlich sei. Eine unendliche Menge, die nicht abzählbar ist, heißt *überabzählbar*.

SATZ 3.84. (I) $\mathbb{Q}$ ist abzählbar.

(2) $\mathbb{R}$ ist nicht abzählbar, also überabzählbar.

Siehe auch Ergänzung 3.65.

Wir können die Mächtigkeiten von Mengen folgendermaßen anordnen: Sind M, M' Mengen, so schreiben wir $\#M \leq \#M'$, wenn es eine injektive Abbildung $M \rightarrow M'$ gibt. Wir schreiben $\#M < \#M'$, wenn $\#M \leq \#M'$ und nicht $\#M = \#M'$ gilt, d.h. wenn es eine Injektion $M \rightarrow M'$, aber keine Bijektion zwischen M und M' gibt. Diese Definition für $\leq$ erfüllt die Eigenschaften einer totalen Ordnung: Offenbar folgt aus $\#M \leq \#M'$ und $\#M' \leq \#M''$, dass $\#M \leq \#M''$, weil die Verkettung injektiver Abbildungen wieder injektiv ist. Es ist auch klar, dass $\#M \leq \#M$ für alle M gilt, da die Identität eine injektive Abbildung ist.

Etwas schwieriger sind die folgenden beiden Ergebnisse, die die Antisymmetrie und Totalität zeigen:

THEOREM 3.85 (Satz von Schröder-Bernstein). Seien M, M' Mengen. Wenn es injektive Abbildungen $M \rightarrow M'$ und $M' \rightarrow M$ gibt, dann gibt es eine Bijektion $M \rightarrow M'$, d.h. M und M' sind gleichmächtig.

THEOREM 3.86. Seien M, M' Mengen. Dann gilt genau eine der folgenden drei Aussagen:

$$\#M < \#M', \quad \#M = \#M', \quad \#M > \#M'.$$

SATZ 3.87. Sei M eine unendliche Menge. Dann gilt $\#M \geq \#\mathbb{N}$.

BEISPIEL 3.88. Sei M eine Menge und $P(M)$ ihre Potenzmenge, d.h. die Menge aller Teilmengen von M . Dann gilt $\#M < \#P(M)$.

Es ist klar, dass es eine Injektion $M \rightarrow P(M)$ gibt, zum Beispiel die Abbildung $m \mapsto \{m\}$. Wir müssen daher zeigen, dass es keine Surjektion $M \rightarrow P(M)$ gibt. Sei $\varphi: M \rightarrow P(M)$ eine Abbildung. Wir behaupten, dass $X := \{m \in M; m \notin \varphi(m)\}$ nicht im Bild von φ liegt (insbesondere ist φ nicht surjektiv). In der Tat, nehmen wir an, dass $X = \varphi(m)$ für ein $m \in M$. Wenn $m \in X$, dann folgt $m \notin \varphi(m) = X$, ein Widerspruch. Wenn $m \notin X$, dann folgt $m \in \varphi(m)$, also $m \in X$, auch ein Widerspruch. Weil weder $m \in X$ noch $m \notin X$ richtig sein kann, kann die Teilmenge X von M nicht im Bild von φ liegen. $\diamond$

BEISPIEL 3.89. Es gilt $\#P(\mathbb{N}) = \#\mathbb{R}$. $\diamond$

²²[https://de.wikipedia.org/wiki/Kardinalzahl_\(Mathematik\)](https://de.wikipedia.org/wiki/Kardinalzahl_(Mathematik))

ERGÄNZUNG 3.90 (Die Kontinuumshypothese). Unter der *Kontinuumshypothese* versteht man die Aussage, dass jede Menge M mit $\#\mathbb{N} \leq \#M \leq \#P(\mathbb{N})$ entweder abzählbar ist (also $\#\mathbb{N} = \#M$ gilt), oder die Mächtigkeit von $P(\mathbb{N})$ hat, also $\#M = \#P(\mathbb{N}) (= \#\mathbb{R})$.

Es wurde von K. Gödel und P. Cohen bewiesen, dass die Kontinuumshypothese unabhängig von dem üblichen Axiomensystem ZFC ist – sie lässt sich weder widerlegen (Gödel, 1938), noch beweisen (Cohen, 1960). Man könnte also entweder die Kontinuumshypothese zu den anderen Axiomen hinzunehmen, oder ihre Negation.

Die *verallgemeinerte Kontinuumshypothese* ist die Aussage, dass für jede unendliche Menge M keine Mächtigkeit existiert, die strikt zwischen $\#M$ und $\#P(M)$ liegt. Sie ist ebenfalls unabhängig von ZFC.

□ Ergänzung 3.90

KAPITEL 4

Körper

Um Lösungsmengen von Gleichungssystemen sinnvoll zu betrachten, müssen wir sagen, in welchem Bereich wir nach Lösungen suchen. Zum Beispiel ist es bei der Gleichung $x^2 = 2$ wichtig zu sagen, ob Lösungen in $\mathbb{Q}$ oder in $\mathbb{R}$ gesucht sind – im ersten Fall gibt es keine, im zweiten Fall existieren zwei Lösungen.

Um die Theorie nicht für jeden Zahlbereich neu entwickeln zu müssen, legen wir daher zunächst fest, welche Bedingungen wir an die Grundmenge stellen wollen, in der wir nach Lösungen der zu betrachtenden Gleichungssysteme schauen. Diese Bedingungen werden im Begriff des *Körpers* zusammengefasst.

4.1. Körper

4.1.1. Anschaulich gesprochen formalisiert der Begriff des *Körpers*, der in der folgenden Definition erklärt wird, die essenziellen Eigenschaften der Grundrechenarten (Addition $+$, Subtraktion $-$, Multiplikation $\cdot$ und Division $/$), wie sie zum Beispiel für rationale und reelle Zahlen gelten, aber eben auch in anderen »Zahlbereichen«.

Andere Eigenschaften der rationalen und reellen Zahlen (speziell die Eigenschaft der »Anordnung«, also dass man sinnvoll von positiven und negativen Zahlen sprechen kann) und Eigenschaften, in denen sich die Bereiche der rationalen und der reellen Zahlen unterscheiden, werden in der Definition eines Körpers nicht angesprochen.

Ein Körper in diesem Sinne hat nichts mit einem Körper im geometrischen Sinne zu tun. Das englische Wort für Körper im Sinne von Definition 4.1 ist »field«.

Die Definition wird uns erlauben, im folgenden viele Tatsachen mit einem einzigen Beweis für die rationalen Zahlen, für die reellen Zahlen und für alle anderen Körper auch zu beweisen. Dies ist ein Grundprinzip der Mathematik: Versuche zu abstrahieren, welche Eigenschaften eines Objekts für eine gewisse Aussage wirklich erforderlich sind, und »baue« aus diesen essenziellen Eigenschaften eine Definition.

Für eine Menge K verstehen wir unter einer *Verknüpfung* (auf K) eine Abbildung $K \times K \rightarrow K$. Eine Verknüpfung erlaubt uns also, zwei Elemente von K herzunehmen (also ein Element von $K \times K$) und daraus ein neues Element von K zu »produzieren«. Typische Beispiele sind die Addition und Multiplikation von (ganzen, rationalen, reellen, ...) Zahlen. Diese Beispiele werden in der folgenden Definition abstrahiert.

DEFINITION 4.1. Ein *Körper* ist ein Tripel $(K, +, \cdot)$ bestehend aus einer Menge K und Verknüpfungen

$$+: K \times K \rightarrow K, \quad \cdot: K \times K \rightarrow K,$$

so dass gilt:

- (I) Die Verknüpfung $+$ ist assoziativ, kommutativ, hat ein eindeutig bestimmtes neutrales Element 0 und es existieren (eindeutig bestimmte) inverse Elemente, genauer:

- (a) Es gilt das *Assoziativgesetz*, d.h. für alle $a, b, c \in K$ gilt

$$(a + b) + c = a + (b + c).$$

- (b) Es gilt das *Kommutativgesetz*, d.h. für alle $a, b \in K$ gilt

$$a + b = b + a.$$

- (c) Es gibt ein eindeutig bestimmtes *neutrales Element* bezüglich der Addition, d.h. ein eindeutig bestimmtes Element $0 \in K$, so dass

$$0 + a = a + 0 = a \quad \text{für alle } a \in K.$$

- (d) Jedes Element $a \in K$ besitzt ein *inverses Element* bezüglich der Addition, d.h. es existiert $b \in K$ mit

$$a + b = 0 = b + a,$$

wobei 0 das im vorherigen Punkt geforderte neutrale Element bezeichnet. (Wir sehen unten, dass dann zu jedem $a \in K$ *genau ein* inverses Element existiert, d.h. b ist durch a eindeutig bestimmt. Wir bezeichnen das Inverse von a bezüglich der Addition mit $-a$; es gilt dann also $a + (-a) = 0 = (-a) + a$.)

- (2) Die Verknüpfung $\cdot$ ist assoziativ, kommutativ, hat ein eindeutig bestimmtes neutrales Element 1 , das von 0 verschieden ist, und für alle $x \in K \setminus \{0\}$ existiert ein inverses Element, genauer:

- (a) Es gilt das *Assoziativgesetz*, d.h. für alle $a, b, c \in K$ gilt

$$(a \cdot b) \cdot c = a \cdot (b \cdot c).$$

- (b) Es gilt das *Kommutativgesetz*, d.h. für alle $a, b \in K$ gilt

$$a \cdot b = b \cdot a.$$

- (c) Es gibt ein eindeutig bestimmtes *neutrales Element* bezüglich der Multiplikation, d.h. ein Element $1 \in K$, so dass

$$1 \cdot a = a \cdot 1 = a \quad \text{für alle } a \in K.$$

- (d) Jedes Element $a \in K \setminus \{0\}$ besitzt ein *inverses Element* bezüglich der Multiplikation, d.h. es existiert $b \in K$ mit

$$a \cdot b = 1 = b \cdot a,$$

wobei 1 das im vorherigen Punkt geforderte neutrale Element bezeichnet. (Wir sehen unten, dass dann zu jedem $a \in K \setminus \{0\}$ *genau ein* inverses Element existiert, d.h. b ist durch a eindeutig bestimmt. Wir bezeichnen das Inverse von a bezüglich der Multiplikation mit a^{-1} ; es gilt dann also $a \cdot a^{-1} = 1 = a^{-1} \cdot a$. Beachte, dass dieser Punkt auf das neutrale Element der Addition Bezug nimmt. Das Element 0 besitzt kein multiplikatives Inverses.)

- (e) Es gilt $0 \neq 1$.

- (3) Es gilt das *Distributivgesetz*:

$$(a + b) \cdot c = (a \cdot c) + (b \cdot c) \quad \text{für alle } a, b, c \in K.$$

⊥

Üblicherweise sagt man, K sei ein Körper und erwähnt die Verknüpfungen $+$, $\cdot$ nicht explizit. Wenn es unbedingt nötig ist, die Addition/Multiplikation auf verschiedenen Körpern in der Notation zu unterscheiden, schreibt man $+_K$, $\cdot_K$ (und gegebenenfalls 0_K statt 0 , 1_K statt 1). In Termen mit Elementen aus K wird der Multiplikationspunkt $\cdot$ üblicherweise weggelassen, d.h. ab steht für $a \cdot b$. Wir verwenden die übliche Konvention *Punkt- vor Strichrechnung*, um nicht zu viele Klammern schreiben zu müssen. Das Distributivgesetz schreiben wir also auch in der Form $(a + b)c = ac + bc$.

Für $a, b \in K$ schreiben wir $a - b := a + (-b)$. Damit haben wir für jeden Körper auch eine Subtraktion $-: K \times K \rightarrow K$ definiert. Auch die Bruchschreibweise $\frac{a}{b} := ab^{-1} (= b^{-1}a)$ (für $a, b \in K, b \neq 0$) verwendet man manchmal für allgemeine Körper. Wir schreiben $K^\times = K \setminus \{0\}$, die sogenannte *multiplikative Gruppe* von K . Diese Bezeichnung wird erst später klar werden (Bemerkung 8.5 (4)), wenn wir den Begriff der *Gruppe* definieren; sie hängt damit zusammen, dass das Produkt von Elementen auf $K^\times$ wieder in $K^\times$ liegt (siehe Abschnitt 4.1.3) und dass jedes Element von $K^\times$ ein Inverses bezüglich der Multiplikation hat. Die Elemente von $K^\times$ heißen auch die *Einheiten* des Körpers K .

Für $n \in \mathbb{N}$ und $a \in K$ schreiben wir $a^n = a \cdot \dots \cdot a$ (n Faktoren) und nennen dieses Element von K die n -te Potenz von a . Eine formale(re) Definition wäre zu sagen, dass $a^0 := 1$ und $a^n = a^{n-1} \cdot a$ für alle $n > 1$. Ist $a \neq 0$, so kann man auch $a^{-n} = (a^{-1})^n$ definieren.

4.1.2. Beispiele für Körper.

BEISPIEL 4.2. (1) Der Körper $\mathbb{Q}$ der rationalen Zahlen:

$$\mathbb{Q} = \left\{ \frac{a}{b}; a, b \in \mathbb{Z}, b \neq 0 \right\}.$$

Die Addition und Multiplikation sind durch die üblichen Bruchrechenregeln gegeben:

$$\frac{a}{b} + \frac{c}{d} = \frac{ad + cb}{bd}, \quad \frac{a}{b} \cdot \frac{c}{d} = \frac{ac}{bd}.$$

Das neutrale Element bezüglich der Addition ist $0 = \frac{0}{1}$, das neutrale Element bezüglich der Multiplikation ist $1 = \frac{1}{1}$. Das Inverse von $\frac{a}{b}$ bezüglich der Addition ist $-\frac{a}{b}$, das Inverse bezüglich der Multiplikation ist (für $a \neq 0$) $\frac{b}{a}$.

- (2) Die Menge $\mathbb{R}$ der reellen Zahlen bildet bezüglich der »üblichen« Addition und Multiplikation einen Körper. (Auch wenn Sie die reellen Zahlen aus der Schule kennen und sie dort wohl häufiger verwendet wurden als der Körper der rationalen Zahlen, sollten Sie sich klarmachen, dass es wesentlich schwieriger ist, die reellen Zahlen und die Verknüpfungen $+$ und $\cdot$ darauf präzise zu definieren.) Siehe zum Beispiel [He] I.2. Eine axiomatische Charakterisierung finden Sie in praktisch allen Lehrbüchern der Analysis. Wie man die reellen Zahlen konstruieren kann, wird meistens nicht erklärt. Eine Möglichkeit bilden die [Dedekindschen Schnitte](#)¹. Eine andere ist, die reellen Zahlen als [Äquivalenzklassen von Cauchy-Folgen von rationalen Zahlen](#)² zu betrachten; diese Sichtweise geht auf Cantor zurück. Die Algebra-Vorlesung wird dazu die nötigen Hilfsmittel bereitstellen.

◇

BEISPIEL 4.3. Die Menge $\mathbb{Z}$ der ganzen Zahlen ist (mit der üblichen Addition und Multiplikation) *kein* Körper, denn es gibt zwar ein neutrales Element für die Multiplikation – die ganze Zahl 1 – aber zum Beispiel hat 2 kein multiplikatives Inverses: Es gibt keine ganze Zahl a mit $2a = 1$. (In der Tat sind 1 und -1 die einzigen Elemente, die ein multiplikatives Inverses haben, und zwar jeweils sich selbst.) Was die Bedingungen an die Addition (und die Distributivität) angeht, sind alle Bedingungen erfüllt.

Die Menge der natürlichen Zahlen ist ebenfalls kein Körper. In diesem Fall gibt es für Elemente $\neq 0$ nicht einmal ein additives Inverses.

◇

¹https://de.wikipedia.org/wiki/Dedekindscher_Schnitt

²https://en.wikipedia.org/wiki/Construction_of_the_real_numbers#Construction_from_Cauchy_sequences

Literatur zu den oben genannten (und anderen) Zahlbereichen:

T. Körner, *Where do Numbers Come From?*, Cambridge Univ. Press, 2020.

J. Kramer, A. von Pippich, *Von den natürlichen Zahlen zu den Quaternionen*, Springer 2013, <https://doi.org/10.1007/978-3-658-02055-2>, oder die erweiterte englische Übersetzung *From Natural Numbers to Quaternions*, <https://doi.org/10.1007/978-3-319-69429-0>.

Oder der Klassiker

H. D. Ebbinghaus et al., *Zahlen*, Springer 1992.

DEFINITION 4.4. Seien $K \subseteq L$ Körper, deren zugrundeliegende Mengen ineinander enthalten sind. Wenn $x +_K y = x +_L y$ und $x \cdot_K y = x \cdot_L y$ für alle $x, y \in K$ gilt, dann heißt K ein *Teilkörper* von L , und L ein *Erweiterungskörper* von K . $\dashv$

Zum Beispiel ist $\mathbb{R}$ ein Erweiterungskörper von $\mathbb{Q}$.

BEISPIEL 4.5 (Die komplexen Zahlen). Sei $\mathbb{C}$ die Menge aller Paare (a, b) von reellen Zahlen. Eine komplexe Zahl können wir also auch als einen Punkt in der reellen Ebene $\mathbb{R}^2$ auffassen, wenn wir die beiden Einträge des Pairs (a, b) als Koordinaten verstehen. Man spricht daher auch von der *komplexen Zahlenebene*.

Wir definieren eine Addition

$$(a, b) + (c, d) := (a + c, b + d),$$

und eine Multiplikation

$$(a, b) \cdot (c, d) := (ac - bd, ad + bc).$$

Es ist dann nicht schwer zu überprüfen, dass $\mathbb{C}$ mit diesen Verknüpfungen ein Körper ist. Dass die Kommutativgesetze für $+$ und $\cdot$ gelten, kann man unmittelbar nachprüfen. Die Gültigkeit der Assoziativgesetze und des Distributivgesetzes kann man auch direkt nachrechnen; weil es etwas lästig ist, lassen wir die Rechnung hier aus. Wir werden aber später eine Möglichkeit kennenlernen, die Gültigkeit dieser Gesetze ohne weiteren Rechenaufwand zu beweisen, siehe Bemerkung 5.42.

Es ist direkt ersichtlich, dass $0_{\mathbb{C}} = (0, 0)$ neutrales Element bezüglich der Addition und dass $1_{\mathbb{C}} = (1, 0)$ neutrales Element bezüglich der Multiplikation ist.

Das Negative (d.h. das additive Inverse) von (a, b) ist $(-a, -b)$. Am schwierigsten ist es vielleicht zu sehen, wie das multiplikative Inverse von (a, b) gegeben ist (sofern nicht $a = b = 0$ gilt). Wenn man die Formel

$$(a, b)^{-1} = \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right)$$

einmal hat, ist es aber leicht nachzurechnen, dass das Produkt dieser beiden Element tatsächlich gleich $1_{\mathbb{C}}$ ist. Man beachte, dass $a^2 + b^2 \neq 0$, weil a und b reelle Zahlen sind, die nicht beide verschwinden.

Dieser Körper heißt der *Körper der komplexen Zahlen*. Er enthält die reellen Zahlen als »Teilkörper«, das bedeutet: $\mathbb{R} \rightarrow \mathbb{C}, a \mapsto (a, 0)$, ist eine injektive Abbildung, die mit Addition und Multiplikation in $\mathbb{R}$ bzw. $\mathbb{C}$ verträglich ist, d.h. es gilt

$$(a + b, 0) = (a, 0) + (b, 0), \quad (ab, 0) = (a, 0)(b, 0).$$

Mittels dieser Injektion identifizieren wir üblicherweise $\mathbb{R}$ mit seinem Bild in $\mathbb{C}$ und schreiben einfach a statt $(a, 0)$. Hier benutzen wir den Begriff Teilkörper in einer Situation, die

ein kleines bisschen allgemeiner ist als Definition 4.4, weil $\mathbb{R}$ streng genommen keine Teilmenge von $\mathbb{C}$ ist, sondern wir nur eine injektive Abbildung $\mathbb{R} \rightarrow \mathbb{C}$ fixieren; diese kleine Ungenauigkeit wollen wir uns an dieser Stelle erlauben.

Das Element $(0, 1)$, das oft mit i bezeichnet wird, hat die Eigenschaft

$$(0, 1)^2 = (0, 1) \cdot (0, 1) = (-1, 0) = -1,$$

sein Quadrat ist -1 . In $\mathbb{C}$ besitzt die quadratische Gleichung $x^2 + 1 = 0$ (anders als in $\mathbb{R}$) also eine Lösung, genauer die beiden Lösungen i und $-i$. Allgemeiner gilt für $b \in \mathbb{R}_{\geq 0}$, dass $(i\sqrt{b})^2 = (-i\sqrt{b})^2 = -b$, die quadratische Gleichung $x^2 + b = 0$ ist also in $\mathbb{C}$ lösbar. Es ist nicht sehr schwer zu zeigen, dass alle quadratischen Gleichungen mit Koeffizienten in $\mathbb{C}$ auch eine Lösung in $\mathbb{C}$ besitzen. Es gilt sogar der *Fundamentalsatz der Algebra*: Jedes nicht-konstante Polynom mit Koeffizienten in $\mathbb{C}$ besitzt eine Nullstelle in $\mathbb{C}$.

Es gilt $(a, b) = a + ib$, und meist schreibt man komplexe Zahlen in dieser Form. Man nennt dann a den *Realteil* und b den *Imaginärteil* der komplexen Zahl $a + ib$.

Wir kratzen hier, in der linearen Algebra, nur an der Oberfläche der komplexen Zahlen. Sie spielen sowohl in der Analysis, speziell in der sogenannten *Funktionentheorie*, der Theorie komplex-differenzierbarer Funktionen, eine herausragende Rolle, sind aber beispielsweise auch in der Algebra und der algebraischen Geometrie sehr wichtig.

Ausführliches [Video^a](#) über die »Basics« der komplexen Zahlen von G. Sander-son/3Blue1Brown (auf Englisch).

^a<https://www.youtube.com/watch?v=5PcpBw5Hbw0>

Für die Praxis ist es für Sie erstmal am wichtigsten, dass Sie wissen, wie Sie mit komplexen Zahlen rechnen. Und das ist ganz einfach: Man benutzt das Distributivgesetz, das Rechnen mit reellen Zahlen, und die Tatsache, dass $i^2 = -1$. Dann braucht man sich die Definitionen von Addition und Multiplikation in der obigen Form gar nicht zu merken. Zum Beispiel:

$$(3 + 2i)(-1 - 2i) = -3 - 6i - 2i - 4i^2 = 1 - 8i,$$

und um einen Bruch von komplexen Zahlen zu vereinfachen, erweitert man so, dass man im Nenner die dritte binomische Formel anwenden kann:

$$\frac{1+i}{2-2i} = \frac{(1+i)(2+2i)}{(2-2i)(2+2i)} = \frac{2+2i+2i+2i^2}{4-4i^2} = \frac{4i}{8} = \frac{1}{2}i.$$

Probieren Sie es selbst aus!^a

^a<https://math.ug/applets/komplexe-zahlen.html>



BEISPIEL 4.6. (I) Da jeder Körper ein Nullelement und ein Einselement besitzt und diese nach Definition verschieden sein müssen, hat jeder Körper mindestens zwei Elemente. Es ist nicht schwer zu sehen, dass es genau einen Körper mit zwei Elementen gibt. Dieser wird mit $\mathbb{F}_2$ bezeichnet.

Als Menge ist $\mathbb{F}_2 = \{0, 1\}$ (Achtung! Diese sind das Nullelement und das Einselement des Körpers $\mathbb{F}_2$, nicht (zum Beispiel) das Nullelement und das Einselement des Körpers $\mathbb{Q}$. Während es nicht so wichtig ist, die 0 und 1 von $\mathbb{Q}$ bzw. $\mathbb{R}$ zu unterscheiden – da es für die Addition und Multiplikation rationaler Zahlen unerheblich ist, ob wir sie als rationale oder als reelle Zahlen auffassen, ist diese Unterscheidung hier essenziell. Es

kann ja zum Beispiel im Körper $\mathbb{F}_2$ nicht $1 + 1 = 2$ gelten, da 2 gar kein Element dieses Körpers ist!)

Um zu sehen, dass es tatsächlich die Möglichkeit gibt, auf $\mathbb{F}_2 = \{0, 1\}$ eine Addition und Multiplikation zu definieren, so dass die Körperaxiome erfüllt sind, überlegt man sich zunächst, dass nur die folgende Möglichkeit in Frage kommt:

+	0	1
0	0	1
1	1	0

·	0	1
0	0	0
1	0	1

Die Tabellen sind dabei folgendermaßen zu lesen: Die Summe $a + b$ von Elementen $a, b \in \mathbb{F}_2$ steht in der linken Tabelle in der Zeile zum Element a und der Spalte zum Element b . Zum Beispiel ist $1 + 0 = 1$, $1 + 1 = 0$. Entsprechend sind in der rechten Tabelle alle Produkte von zwei Elementen eingetragen. Es ist dann (etwas langweilig, aber) nicht schwer, die Körperaxiome nachzuprüfen.

Der Körper $\mathbb{F}_2$ ist kein Teilkörper von $\mathbb{Q}$, denn (selbst wenn man die Menge $\mathbb{F}_2 = \{0, 1\}$ als Teilmenge von $\mathbb{Q}$ auffassen könnte): in $\mathbb{F}_2$ gilt $1 + 1 = 0$, aber in $\mathbb{Q}$ ist das nicht richtig. Zwar ist die ganze Zahl 2 kein Element von $\mathbb{F}_2$, es ist aber eine nützliche Sichtweise, in jedem Körper K neben 0 und 1 auch Elemente $2_K, 3_K, \dots$ zu definieren, und zwar als $2_K := 1 + 1$, $3_K = 1 + 1 + 1$, usw. (Und meist schreibt man dann wieder einfach 2, 3, usw. statt $2_K, 3_K$. Wenn man Missverständnisse befürchtet, kann man $2_K, 3_K$ schreiben, oder schreibt speziell im Fall von $\mathbb{F}_2$ auch $\bar{2}, \bar{3}, \dots$, oder manchmal $[2], [3]$, dots für diese Elemente.) In diesem Sinne gilt dann im Körper $\mathbb{F}_2$ die Gleichheit $2 = 0$. Siehe Abschnitt 4.2.2.)

- (2) In ähnlicher Weise kann man einen Körper mit genau 3 Elementen konstruieren, für die wir hier die Bezeichnungen 0, 1, 2 wählen wollen, und dieser wird mit $\mathbb{F}_3$ bezeichnet.

+	0	1	2
0	0	1	2
1	1	2	0
2	2	0	1

·	0	1	2
0	0	0	0
1	0	1	2
2	0	2	1

Im Körper $\mathbb{F}_3$ gilt also $(3 :=) 1 + 1 + 1 = 0$ und $2 = -1$.

Wie im Fall von $\mathbb{F}_2$ kann man auch hier die Körperaxiome direkt nachprüfen, allerdings ist es noch aufwändiger. Siehe Abschnitt 4.2.1 für eine bessere und gleichzeitig allgemeinere Lösung.

- (3) Es gibt keinen Körper mit genau 6 Elementen.

Für den Fall von 6 Elementen kann man das mit ein bisschen Mühe direkt einsehen, indem man nachprüft, dass es keine Möglichkeit gibt, eine Additions- und eine Multiplikationstabelle wie oben für eine sechselementige Menge aufzustellen, die mit den Körperaxiomen verträglich ist.

Allgemeiner können Sie in der Algebra-Vorlesung lernen, dass die Anzahl der Elemente eines Körpers entweder unendlich oder die Potenz einer Primzahl ist (siehe auch Ergänzung 6.57), und dass es zu jeder Primzahlpotenz q genau einen Körper mit q Elementen gibt.

- (4) Die Verknüpfung $(a, b) \mapsto a - b$ auf den rationalen Zahlen erfüllt weder das Assoziativgesetz noch das Kommutativgesetz (geben Sie dafür Beispiele!). Finden Sie weitere Verknüpfungen, die einige der Forderungen aus der Definition eines Körpers nicht erfüllen.

◇

Das nächste Lemma zeigt, dass es ausreichen würde, in den Körperaxiomen die Existenz neutraler Elemente für Addition und Multiplikation zu fordern, da diese automatisch eindeutig bestimmt wären. Da die Axiome über die Existenz von Inversen aber auf die neutralen Elemente Bezug nehmen müssen, haben wir die Eindeutigkeit in unsere Definition aufgenommen. Außerdem sehen wir, dass die inversen Elemente jeweils eindeutig bestimmt sind.

LEMMA 4.7. (1) *In der Situation von Definition 4.1 sind die neutralen Elemente der Addition und der Multiplikation in K notwendigerweise eindeutig bestimmt, auch wenn dies nicht gefordert würde.*

(2) *Für jedes $a \in K$ ist das additive inverse Element eindeutig bestimmt. Für jedes $a \in K^\times$ ist das multiplikative Inverse eindeutig bestimmt.*

BEWEIS. zu (1). Seien e, e' Elemente von K , die beide die Eigenschaft eines neutralen Elements der Addition haben. Wir zeigen, dass dann $e = e'$ gelten muss; das beweist die behauptete Eindeutigkeit. In der Tat gilt

$$e = e + e' = e',$$

wobei wir bei der ersten Gleichheit benutzen, dass e' ein neutrales Element bezüglich $+$ ist, und bei der zweiten Gleichheit die entsprechende Eigenschaft für e ausnutzen.

Analog können wir für die Multiplikation argumentieren. Haben $e, e' \in K$ beide die Eigenschaft eines neutralen Elements bezüglich der Multiplikation, so folgt

$$e = e \cdot e' = e'.$$

zu (2). Sei nun $a \in K$ und seien $b, b' \in K$ inverse Elemente von a bezüglich der Addition, d.h. es gelte $a + b = b + a = 0, a + b' = 0$. Wir erhalten dann

$$b = b + 0 = b + (a + b') = (b + a) + b' = 0 + b' = b'.$$

Die Rechnung zur Eindeutigkeit der multiplikativen Inversen verläuft analog. $\square$

4.1.3. Rechenregeln. Sei K ein Körper, und seien $a, b, c, \dots$ Elemente von K .

Kürzungsregeln. Aus $a + c = b + c$ folgt $a = b$, denn wir können auf beiden Seiten $-c$ addieren und erhalten $a = a + (c - c) = (a + c) - c = (b + c) - c = b$.

Ist $c \neq 0$, so folgt aus $ac = bc$ mit einem analogen Argument, dass $a = b$, da c ein multiplikatives Inverses besitzt.

Multiplikation mit 0. Es gilt $0 \cdot a = 0$. In der Tat gilt $0 \cdot a = (0 + 0) \cdot a = 0 \cdot a + 0 \cdot a$, und die Behauptung folgt, indem wir auf beiden Seiten $0 \cdot a$ abziehen (mit anderen Worten: das Negative des Elements $0 \cdot a$ addieren).

Weitere Rechenregeln Es gilt $-(-a) = a$. Warum?

Es gilt $(-1) \cdot a = -a$. In der Tat, $a + (-1) \cdot a = 1 \cdot a + (-1) \cdot a = (1 + (-1)) \cdot a = 0 \cdot a = 0$, und das bedeutet, dass $(-1) \cdot a$ die charakterisierende Eigenschaft des Negativen von a erfüllt.

Insbesondere folgt $(-1) \cdot (-1) = -(-1) = 1$, und daher auch $(-a) \cdot (-b) = ab$ für alle a, b .

BEMERKUNG 4.8. Beachten Sie, dass wir in einem allgemeinen Körper keinen Begriff von positiven und negativen Elementen haben. Nicht nur haben wir diese Begriffe nicht in unsere Definition aufgenommen – es gibt Körper, in denen sie sich nicht sinnvoll (also mit den erwarteten Eigenschaften) definieren lassen. Im Körper der komplexen Zahlen gibt es ein Element i mit $i \cdot i = -1$. (Auch in $\mathbb{F}_2$ gibt es ein Element mit dieser Eigenschaft.) Aus den üblichen Eigenschaften der Begriffe *positiv* und *negativ* würde aber folgen, dass Quadratzahlen immer positiv (oder Null) sind, und dass -1 negativ ist. Das lässt sich nicht mit der Gleichheit $i^2 = -1$ in Einklang bringen. $\diamond$

Nullteilerfreiheit. Gilt $ab = 0$, so folgt $a = 0$ oder $b = 0$. Denn wenn $a \neq 0$ ist, dann können wir die beiden Seiten der Gleichung $ab = 0$ mit a^{-1} multiplizieren und erhalten $b = 0$. (Oder äquivalent umformuliert: Sind $a, b \in K^\times$, so ist $ab \neq 0$, also $ab \in K^\times$.)

Summen und Produkte mit mehr als zwei Elementen. Die Assoziativgesetze zeigen auch, dass man auch in Summen und Produkten von mehr als drei Elementen die Klammern nach Belieben setzen kann, ohne das Ergebnis zu verändern, zum Beispiel:

$$(a + b) + (c + (d + e)) = a + (b + ((c + d) + e)).$$

Daher können und werden wir es uns üblicherweise erlauben, die Klammern in solchen Ausdrücken wegzulassen. (Wenn mehrere Rechenoperationen in einem Ausdruck vorkommen, soll wie gesagt *Punkt- vor Strichrechnung* gelten, d.h. $\cdot$ und $/$ haben Priorität vor $+$ und $-$. Operatoren aus der gleichen Klasse (also $+$, $-$ einerseits und $\cdot$, $/$ andererseits) verarbeiten wir von links nach rechts.

Wir können in beliebigen Körpern die Summennotation $\sum_{i=1}^n a_i$ und die Produktnotation $\prod_{i=1}^n a_i$ anwenden, und wegen des Kommutativgesetzes sogar Summen (und analog Produkte) der Form $\sum_{i \in I} a_i$ für endliche Mengen I und Familien $(a_i)_{i \in I}$ betrachten.

ERGÄNZUNG 4.9 (Schiefkörper).

DEFINITION 4.10. Ein *Schiefkörper* ist eine Menge D zusammen mit Verknüpfungen $+: D \times D \rightarrow D$, $\cdot: D \times D \rightarrow D$, für die alle Körperaxiome bis auf möglicherweise die Kommutativität der Multiplikation erfüllt sind. $\dashv$

Wie beim Begriff des Erweiterungskörpers kann man auch hier die Situation betrachten, dass ein Körper K und ein Schiefkörper D gegeben sind, so dass $K \subset D$ eine Teilmenge ist und dass die Addition und Multiplikation auf D für Elemente von K dieselben Ergebnisse liefert wie Addition und Multiplikation auf D . Gilt zusätzlich noch $xy = yx$ für alle $x \in K$ und $y \in D$, dann nennt man D eine (assoziative) *Divisionsalgebra* über K . (In der Regel fordert man auch noch, dass D als K -Vektorraum von endlicher Dimension ist. Der Dimensionsbegriff wird uns aber erst später zur Verfügung stehen.)

Es ist ein interessantes Problem zu verstehen, welche Divisionsalgebren es über einem gegebenen Körper gibt. Ein Beispiel werden wir in der folgenden Ergänzung 4.11 kennenlernen.

Siehe auch Ergänzungen 6.65, 10.20.

□ Ergänzung 4.9

ERGÄNZUNG 4.11 (Die Hamiltonschen Quaternionen). Wir definieren auf $\mathbb{H} := \mathbb{R}^4 = \{(a, b, c, d); a, b, c, d \in \mathbb{R}\}$ eine Addition durch

$$(a, b, c, d) + (a', b', c', d') = (a + a', b + b', c + c', d + d').$$

Die Multiplikation definieren wir durch

$$(a, b, c, d) \cdot (a', b', c', d') = (aa' - bb' - cc' - dd', ab' + ba' + cd' - dc', ac' - bd' + ca' + db', ad' + bc' - cb' + da').$$

Man kann zeigen, dass $\mathbb{H}$ mit diesen Operationen einen Schiefkörper bildet. Das neutrale Element der Addition ist $0 = (0, 0, 0, 0)$, das neutrale Element der Multiplikation ist $1 = (1, 0, 0, 0)$. Dass die Addition alle geforderten Eigenschaften hat, ist klar.

Wenn wir $\mathbb{R}$ mittels der Einbettung $\mathbb{R} \rightarrow \mathbb{H}, a \mapsto (a, 0, 0, 0)$, mit der Teilmenge $\{(a, 0, 0, 0); a \in \mathbb{R}\}$ identifizieren, dann ist $\mathbb{H}$ sogar eine Divisionsalgebra über $\mathbb{R}$, genannt die Algebra der Hamiltonschen Quaternionen.

Wir können sogar mit $a + bi \mapsto (a, b, 0, 0)$ die komplexen Zahlen $\mathbb{C}$ als Teilkörper in $\mathbb{H}$ einbetten. Dann wird $i \in \mathbb{C}$ auf das Element $(0, 1, 0, 0) \in \mathbb{H}$ abgebildet, das man wieder mit i

bezeichnet. Außerdem setzt man oft $j := (0, 0, 1, 0)$ und $k := (0, 0, 0, 1)$. Die oben angegebene Multiplikation ist die eindeutig bestimmte Verknüpfung, die $\mathbb{H}$ zu einem Schiefkörper macht, so dass $(a, 0, 0, 0)(a', b', c', d') = (aa', ab', ac', ad')$, und dass

$$i^2 = j^2 = k^2 = ijk = -1$$

gilt. So einfach das auch aussieht – Hamilton hat jahrelang darüber nachgedacht, wie man eine Divisionsalgebra über $\mathbb{R}$ konstruieren könnte, bis ihm die Idee für die hier angegebene Multiplikation kam. Siehe Ergänzung 4.9.

Wir sehen, dass nicht $xy = yx$ für alle $x \in \mathbb{C}$ und $y \in \mathbb{H}$ gilt, zum Beispiel ist

$$ij = i = -ji.$$

Es ist nicht sehr angenehm, die Schiefkörperaxiome anhand der hier gegebenen Formeln zu überprüfen, und es gibt bessere Methoden, wie wir in Kürze sehen werden (Ergänzung 5.64).

Siehe auch Ergänzung 9.44.

□ Ergänzung 4.11



One of the big misapprehensions about mathematics that we perpetrate in our classrooms is that the teacher always seems to know the answer to any problem that is discussed. This gives students the idea that there is a book somewhere with all the right answers to all of the interesting questions, and that teachers know those answers. And if one could get hold of the book, one would have everything settled. That's so unlike the true nature of mathematics.

L. Henkin

in: Steen, Albers (eds.), *Teaching Teachers, Teaching Students*, Birkhäuser, 1981.

Gefunden auf <http://math.furman.edu/~mwoodard/mqs/data.html>

4.2. Endliche Körper

4.2.1. Rechnen mit Restklassen. Sei $n \geq 1$ eine natürliche Zahl. Wir betrachten die Menge

$$\mathbb{Z}/n := \{\bar{0}, \bar{1}, \dots, \overline{n-1}\}$$

mit n Elementen, die wir mit den Symbolen $\bar{a}$ für a in den natürlichen Zahlen von 0 bis $n-1$ bezeichnen. (Die Wahl der Bezeichnung $\mathbb{Z}/n$ – wird im zweiten Semester noch klarer werden. Andere übliche Bezeichnungen sind $\mathbb{Z}/(n)$ und $\mathbb{Z}/n\mathbb{Z}$.)

Wir nennen die Elemente von $\mathbb{Z}/n$ auch *Restklassen (modulo n)*.

Wir definieren die folgenden Verknüpfungen:

$$a + b := \bar{r}, \text{ wobei } r \text{ der Rest von } a + b \text{ bei Division durch } n \text{ ist,}$$

$$a \cdot b := \bar{r}, \text{ wobei } r \text{ der Rest von } a \cdot b \text{ bei Division durch } n \text{ ist.}$$

Wie üblich lassen wir den Multiplikationspunkt $\cdot$ manchmal weg. Es ist klar, dass sowohl für $+$ als auch für $\cdot$ das Kommutativgesetz gilt.

BEISPIEL 4.12. (1) Für $n = 12$ erhalten wir zum Beispiel

$$\bar{9} + \bar{6} = \bar{3}.$$

(Wie beim »Rechnen« auf der Uhr: 6 Stunden nach 9 Uhr ist es 3 Uhr, denn ab 12 beginnt die Zählung wieder bei 1.)

(2) Es gilt für $n = 17$:

$$\bar{3} \cdot \bar{6} = \bar{1} \quad \text{in } \mathbb{Z}/17.$$

(3) Für alle n , und alle $a \in \mathbb{Z}/n$ gilt $\bar{0} + a = a$, also ist $\bar{0}$ ein neutrales Element für $+$. Es gilt $\bar{0} \cdot a = \bar{0}$. Ferner gilt $\bar{1} \cdot a = a$, also ist $\bar{1}$ ein neutrales Element für $\cdot$.

(4) Sei $n = 2$. Dann ist $\bar{1} + \bar{1} = \bar{0}$, und wir sehen, dass die Operationen $+$ und $\cdot$ auf $\mathbb{Z}/2$ gerade die Addition und Multiplikation liefern, die wir auf $\mathbb{F}_2$ definiert haben.

(5) Sei $n = 3$. In diesem Fall kann man $\mathbb{Z}/3$ mit den Operationen $+$ und $\cdot$ mit dem Körper $(\mathbb{F}_3, +, \cdot)$ identifizieren.

◇

Wir bezeichnen mit $\pi: \mathbb{Z} \rightarrow \mathbb{Z}/n$ die Abbildung, die $x \in \mathbb{Z}$ abbildet auf $\bar{r}$, wobei r der Rest von x bei Division durch n ist. Diese Abbildung nennt man manchmal die *kanonische Projektion*. Offenbar handelt es sich um eine surjektive Abbildung. Natürlich ist π nicht injektiv; es gilt genau dann $\pi(a) = \pi(b)$, wenn a und b denselben Rest bei Division durch n haben, oder äquivalent ausgedrückt: wenn $a - b$ durch n teilbar ist. Insbesondere gilt $\pi(a) = \bar{0}$ genau dann, wenn a ein Vielfaches von n ist.

Sehr nützlich ist, dass die Abbildung π mit den Additionen und Multiplikationen auf beiden Seiten verträglich ist, und zwar im folgenden Sinne:

Erstens können wir die Definition der Addition und Multiplikation auf $\mathbb{Z}/n$ auch schreiben als

$$(1) \quad a + b = \pi(a + b), \quad a \cdot b = \pi(ab), \quad a, b \in \mathbb{Z}/n.$$

Zweitens gilt:

LEMMA 4.13. Seien $x, x' \in \mathbb{Z}$. Dann gilt

$$\pi(x + x') = \pi(x) + \pi(x'), \quad \pi(xx') = \pi(x) \cdot \pi(x').$$

BEWEIS. Wir schreiben die Division durch n mit Rest aus als

$$x = qn + r, \quad x' = q'n + r', \quad 0 \leq r, r' < n.$$

Es gilt also $\pi(x) = r, \pi(x') = r'$.

Es ist eine offensichtliche Eigenschaft der Division mit Rest durch n , dass sich der Rest nicht ändert, wenn wir den Dividenten (die Zahl, die durch n geteilt wird) um ein Vielfaches von n abändern.

Also gilt

$$\pi(x) + \pi(x') = \pi(r + r') = \pi(x + x'),$$

wobei die erste Gleichheit aus (1) folgt, und die zweite, weil sich x und r , beziehungsweise x' und r' nur um Vielfache von n unterscheiden.

Weil $xx' = (qq'n + qr' + q'r)n + rr'$ sich ebenfalls um ein Vielfaches von n von rr' unterscheidet, können wir für $\cdot$ ganz analog rechnen:

$$\pi(x) \cdot \pi(x') = \pi(rr') = \pi(xx').$$

□

Dieses Lemma erlaubt es uns, ohne weiteren Aufwand zu zeigen, dass die Verknüpfungen $+$ und $\cdot$ auf $\mathbb{Z}/n$ assoziativ und kommutativ sind und dass das Distributivgesetz gilt. Seien nämlich $x, y, z \in \mathbb{Z}/n$, und seien $\dot{x}, \dot{y}, \dot{z} \in \mathbb{Z}$ mit $\pi(\dot{x}) = x, \pi(\dot{y}) = y, \pi(\dot{z}) = z$ (wir könnten x, y, z als Elemente von $\mathbb{Z}$ auffassen und $\dot{x} = x$ usw. wählen, aber es ist für das Weitere egal; wichtig ist nur, dass so eine Wahl überhaupt möglich ist, d.h. dass π surjektiv ist).

Dann gilt zum Beispiel

$$\begin{aligned}(x + y) + z &= (\pi(\dot{x}) + \pi(\dot{y})) + \pi(\dot{z}) = \pi(\dot{x} + \dot{y}) + \pi(\dot{z}) = \pi((\dot{x} + \dot{y}) + \dot{z}) \\ &= \pi(\dot{x} + (\dot{y} + \dot{z})) = \pi(\dot{x}) + \pi(\dot{y} + \dot{z}) = x + (\pi(\dot{y}) + \pi(\dot{z})) = x + (y + z).\end{aligned}$$

Diese Gleichungskette ist zwar lang, aber die einzelnen Schritte sind ganz formal. Wir benutzen mehrfach das Lemma, und in der Mitte dann das Assoziativgesetz für die Addition in $\mathbb{Z}$. Wenn wir alle $+$ -Zeichen in dieser Rechnung durch $\cdot$ ersetzen, dann ergibt sich ein Beweis des Assoziativgesetzes der Multiplikation. Die beiden Kommutativgesetze und das Distributivgesetz können nach exakt dem gleichen Schema bewiesen werden.

Wir haben oben schon festgestellt, dass $\bar{0} \in \mathbb{Z}/n$ ein neutrales Element bezüglich $+$ und $\bar{1} \in \mathbb{Z}/n$ ein neutrales Element bezüglich $\cdot$ ist. Für $a \in \mathbb{Z}$ ist $\pi(-a)$ ein additives Inverses von $\pi(a)$, denn

$$\pi(a) + \pi(-a) = \pi(a + (-a)) = \bar{0}.$$

Das einzige Axiom, das aus der Liste der Körperaxiome noch fehlt, ist die Existenz von multiplikativen Inversen. Hier zeigt sich eine interessante Situation:

- Anders als in $\mathbb{Z}$ haben manchmal auch Elemente $\neq \bar{1}$ (und $\neq -\bar{1}$) ein multiplikatives Inverses. Wir hatten beispielsweise oben schon festgehalten, dass in $\mathbb{Z}/17$ gilt, dass $\bar{3} \cdot \bar{6} = \bar{1}$. Wir haben auch schon gesehen, dass $\mathbb{Z}/2$ der Körper $\mathbb{F}_2$ und $\mathbb{Z}/3$ der Körper $\mathbb{F}_3$ ist.
- Im allgemeinen ist $\mathbb{Z}/n$ kein Körper; zum Beispiel gilt in $\mathbb{Z}/6$: $\bar{2} \neq \bar{0}, \bar{3} \neq \bar{0}$, aber $\bar{2} \cdot \bar{3} = \bar{0}$. Diese Beobachtung können wir leicht verallgemeinern: Wenn $n = ab$ mit $0 < a, b < n$ gilt, dann gilt $\bar{a}\bar{b} = \bar{0}$ in $\mathbb{Z}/n$, obwohl $\bar{a} \neq \bar{0}$ und $\bar{b} \neq \bar{0}$ ist. Also ist $\mathbb{Z}/n$ kein Körper.

Der zweite Punkt zeigt uns, dass $\mathbb{Z}/n$ höchstens dann ein Körper sein kann, wenn n eine Primzahl (vergleiche Ergänzung 3.44) ist. Wir wollen nun zeigen, dass wir für eine Primzahl p tatsächlich einen Körper $\mathbb{Z}/p$ konstruiert haben.

Dafür benutzen wir die *Primeigenschaft*: Ist p eine Primzahl und ist p ein Teiler des Produkts ab von zwei ganzen Zahlen a, b , so ist p ein Teiler von a oder von b (wie immer ist auch erlaubt, dass p beide Zahlen a und b teilt). Siehe Satz 3.52 in Ergänzung 3.51 für einen Beweis. Wir wollen uns, für den Fall, dass Sie diese Ergänzung ausgelassen haben, auf die Bemerkung beschränken, dass diese Eigenschaft aus der Eindeutigkeit der Primfaktorzerlegung in $\mathbb{Z}$ folgt. Denn die (eindeutige!) Primfaktorzerlegung des Produktes ab erhalten wir, indem wir die Primfaktorzerlegungen von a und von b zusammenfügen (miteinander multiplizieren), und wenn p als Faktor in dem Produkt auftritt, muss es folglich auch in einem der Faktoren dabei sein.

SATZ 4.14. Sei p eine Primzahl. Dann ist $\mathbb{Z}/p$ ein Körper.

BEWEIS. Nach dem bereits Gesagten ist nur noch zu zeigen, dass jedes Element $a \in (\mathbb{Z}/p) \setminus \{\bar{0}\}$ ein multiplikatives Inverses besitzt. Sicher genügt es dafür, zu beweisen, dass die Abbildung

$$m_a: \mathbb{Z}/p \rightarrow \mathbb{Z}/p, \quad x \mapsto a \cdot x,$$

surjektiv ist, denn $m_a(x) = \bar{1}$ besagt ja gerade $a \cdot x = \bar{1}$; dann ist x das gesuchte Inverse.

Weil $\mathbb{Z}/p$ eine endliche Menge ist, ist es äquivalent zu zeigen, dass m_a injektiv ist (Satz 3.64).

Das bedeutet, wir müssen zeigen, dass für $x \neq y \in \mathbb{Z}/p$ stets $ax \neq ay$, oder mit anderen Worten $a(x - y) \neq \bar{0}$ gilt. Wir schreiben $a = \pi(\dot{a})$, $x = \pi(\dot{x})$, $y = \pi(\dot{y})$ für geeignete ganze Zahlen $\dot{a}, \dot{x}, \dot{y}$. Dann können wir umformulieren: Dass $a \neq \bar{0}$ ist, ist gleichbedeutend damit, dass $p \nmid \dot{a}$, ebenso bedeutet $x \neq y$, dass $p \nmid \dot{x} - \dot{y}$. Wegen der oben bemerkten Primeigenschaft erhalten wir daraus, dass p kein Teiler des Produkts $\dot{a}(\dot{x} - \dot{y})$ ist. Folglich ist $a(x - y) = \pi(\dot{a})(\pi(\dot{x}) - \pi(\dot{y})) = \pi(\dot{a}(\dot{x} - \dot{y})) \neq \bar{0}$, wie gewünscht. $\square$

ERGÄNZUNG 4.15. Wir geben noch einen alternativen Beweis dafür, dass $\mathbb{Z}/p$ für Primzahlen p ein Körper ist. Dieser beruht auf Lemma 3.53 aus Ergänzung 3.51. (Wir haben dieses Lemma benutzt, um die Eindeutigkeit der Primfaktorzerlegung zu beweisen; insofern hängt auch der erste Beweis für die Körpereigenschaft von $\mathbb{Z}/p$ bei unserem Aufbau der Dinge von Lemma 3.53 ab. Man kann aber die Primeigenschaft, und damit die Eindeutigkeit der Primfaktorzerlegung auch (etwas) anders beweisen.)

SATZ 4.16. Sei $n > 1$ eine natürliche Zahl, und sei $a \in \mathbb{Z}$. Genau dann besitzt $\pi(a)$ ein multiplikatives Inverses in $\mathbb{Z}/n$, wenn $\text{ggT}(a, n) = 1$.

BEWEIS. Sei $a \in \mathbb{Z}$ mit $\text{ggT}(a, n) = 1$. Nach Lemma 3.53 existieren dann $x, y \in \mathbb{Z}$ mit

$$1 = xa + yn.$$

Das bedeutet aber

$$\pi(x) \cdot a = \pi(xa) = \pi(1 - yp) = \pi(1) + \pi(-yp) = \pi(1) = \bar{1},$$

also ist $\pi(x)$ das gesuchte Inverse von a .

Hat andererseits $\pi(a)$ ein multiplikatives Inverses, so können wir dieses in der Form $\pi(b)$ für $b \in \mathbb{Z}$ schreiben. Dass das Produkt $\pi(a)\pi(b) = \bar{1}$ ist (in $\mathbb{Z}/n$), bedeutet genau, dass $1 - ab$ von n geteilt wird. Ist also d ein gemeinsamer Teiler von a und n , so muss d auch 1 teilen, also $d \in \{-1, 1\}$. $\square$

KOROLLAR 4.17. Sei p eine Primzahl. Dann ist $\mathbb{Z}/p$ ein Körper.

BEWEIS. Sei $a \in \mathbb{Z}/p$, $a \neq \bar{0}$. Wir schreiben $a = \pi(\dot{a})$ für eine ganze Zahl $\dot{a}$. Weil $a \neq \bar{0}$, wird $\dot{a}$ nicht von p geteilt, und weil p eine Primzahl ist, folgt $\text{ggT}(\dot{a}, p) = 1$. Daher können wir den vorherigen Satz anwenden. $\square$

(Wir erhalten so auch einen (etwas) anderen Beweis von Satz 3.52: Der Satz sagt, dass wenn immer die Primzahl p ein Produkt ab teilt, p einen der Faktoren teilt. In der Tat, $p \mid ab$ bedeutet $\pi(a)\pi(b) = \pi(ab) = \bar{0}$, und da $\mathbb{Z}/p$ ein Körper ist, gilt dann $\pi(a) = \bar{0}$ oder $\pi(b) = \bar{0}$. Das bedeutet gerade $p \mid a$ oder $p \mid b$.) $\square$ Ergänzung 4.15

Für eine Primzahl p schreiben wir statt $\mathbb{Z}/p$ oft auch $\mathbb{F}_p$, um zu betonen, dass es sich hierbei um einen Körper handelt.

(Wenn Sie Lust haben, ist jetzt ein guter Zeitpunkt, sich den Sternchen-Abschnitt 3.14.2 über Äquivalenzrelationen anzuschauen, und dort speziell das Beispiel über $\mathbb{Z}/3$.)

Zum Schluss noch einige gebräuchliche Schreibweisen: Oft schreibt man $\bar{a}$ statt $\pi(a)$ für beliebige ganze Zahlen a . Für $0 \leq a < n - 1$ fällt das mit der oben verwendeten Schreibweise zusammen. Manchmal schreibt man auch $[a]$ statt $\bar{a}$.

Wenn keine Missverständnisse zu befürchten sind (oder die Autor*in entscheidet, der Leser*in zuzumuten, das selbst (gedanklich) »aufzuräumen«), dann kann man auch einfach a statt $\bar{a}$ schreiben.

Eine andere übliche Schreibweise ist es, statt $\pi(x) = \pi(y)$ (für $x, y \in \mathbb{Z}$ und ein fixiertes n)

$$x \equiv y \pmod{n}$$

zu schreiben (gesprochen » x ist kongruent zu y modulo n «). Wie oben erläutert ist das äquivalent dazu zu sagen, dass $x - y$ von n geteilt wird.

Die Schreibweise, statt $\bar{a}$ einfach a zu schreiben, kann man im Rahmen der folgenden allgemeinen Konvention sehen: Sei K ein beliebiger Körper. Für eine natürliche Zahl n schreiben wir n_K , oder einfach n , wenn klar ist, dass wir n als Element von K betrachten möchten, für das Element

$$n_K := 1 + \cdots + 1 \quad (n \text{ Summanden})$$

von K . Zum Beispiel gilt (für die übliche Einbettung der Menge der natürlichen Zahlen in den Körper $\mathbb{Q}$) $n_{\mathbb{Q}} = n$ für alle natürlichen Zahlen n . Andererseits gilt $2_{\mathbb{F}_2} = 0_{\mathbb{F}_2}$ und allgemeiner $n_{\mathbb{F}_2} = 0_{\mathbb{F}_2}$ für alle geraden, und $n_{\mathbb{F}_2} = 1_{\mathbb{F}_2}$ für alle ungeraden natürlichen Zahlen n . Man sagt auch, in $\mathbb{F}_2$ gelte $2 = 0$, $3 = 1$, etc.

Man erweitert diese Definition auf alle ganzen Zahlen, indem man $(-n)_K := -n_K$ für natürliche Zahlen n setzt. Dann ist n_K für alle ganzen Zahlen n definiert. Hier ist $-n_K$ das Negative in K des Elements n_K .

4.2.2. Die Charakteristik eines Körpers * In diesem Abschnitt untersuchen wir das Phänomen, dass in manchen Körpern der Ausdruck $1 + \cdots + 1$ gleich 0 ist, etwas genauer. Zum Beispiel gilt in $\mathbb{F}_2$, dass $1 + 1 = 0$, und in $\mathbb{F}_p$, dass $1 + \cdots + 1 = 0$ (mit p Summanden auf der linken Seite).

Wir verwenden weiter die am Ende des vorherigen Abschnitts eingeführte Schreibweise n_K für $n \in \mathbb{Z}$. Wie gesagt, werden wir später dazu übergehen, einfach n statt n_K zu schreiben, aber für den Moment bleiben wir bei dem Index $-_K$ zur besseren Unterscheidung.

Wir können also von der Abbildung $\varphi: \mathbb{Z} \rightarrow K, n \mapsto n_K$ sprechen. Es folgt leicht aus der Definition von n_K , dass $\varphi(m + n) = \varphi(m) + \varphi(n)$ gilt, wobei auf der linken Seite $+$ die Addition in $\mathbb{Z}$, und auf der rechten Seite $+$ die Addition in K bezeichnet. (Überlegen Sie sich, dass das auch hinkommt, wenn m und/oder n negativ sind.) Darüberhinaus gilt auch $\varphi(mn) = \varphi(m)\varphi(n)$, wobei wieder links in $\mathbb{Z}$ und rechts in K multipliziert wird. Das zeigt man zuerst für nicht-negative m und n , indem man m und n als Summen von m bzw. n Summanden 1 schreibt und das Distributivgesetz ausnutzt. Aus $\varphi(-m) = -\varphi(m)$ kann man das Ergebnis dann in der allgemeinen Form herleiten.

DEFINITION 4.18. Sei K ein Körper. Wenn es eine natürliche Zahl $n \geq 1$ gibt, so dass $n_K = 0$, also

$$1 + \cdots + 1 = 0$$

in K gilt (mit n Summanden auf der linken Seite), so nennen wir die kleinste solche Zahl $n \geq 1$ die *Charakteristik* von K .

Gibt es kein solches n , so sagen wir, K habe die Charakteristik 0 . —

Da in jedem Körper per Definition die Elemente 0 und 1 verschieden sind, kann ein Körper nicht die Charakteristik 1 haben. Es gilt sogar die folgende viel stärkere Einschränkung:

SATZ 4.19. Sei K ein Körper der Charakteristik $p \neq 0$. Dann ist p eine Primzahl.

BEWEIS. Sei K ein Körper mit Charakteristik $p \neq 0$. Wir haben bereits bemerkt, dass $p > 1$ gelten muss. Angenommen, p ließe sich als Produkt zweier Zahlen $m, n > 1$ zerlegen. Wir könnten dann p_K (oder mit der oben eingeführten Notation $\varphi(p)$) schreiben als

$$p_K = \varphi(p) = \varphi(mn) = \varphi(m)\varphi(n) = m_K n_K,$$

was bedeutet, dass $m_K = 0$ oder $n_K = 0$ gelten muss. Das steht im Widerspruch zur Minimalität von p in der Definition der Charakteristik. $\square$

BEISPIEL 4.20. Die Charakteristik von $\mathbb{Q}$ ist 0, ebenso die Charakteristik von jedem Erweiterungskörper von $\mathbb{Q}$, insbesondere also von $\mathbb{R}$ und $\mathbb{C}$.

Ist p eine Primzahl, so ist die Charakteristik von $\mathbb{F}_p$ gleich p . In der Tat folgt aus der Definition von $\mathbb{F}_p$ als Menge von Restklassen modulo p , dass $p_{\mathbb{F}_p} = 0$, aber $n_{\mathbb{F}_p} \neq 0$ für alle $n = 1, \dots, p-1$.

Es gibt für jede Primzahl p noch weitere Körper der Charakteristik p , sowohl endliche mit mehr als p Elementen, als auch unendliche. Da in einem Körper K der Charakteristik 0 die Elemente n_K , $n \in \mathbb{Z}$, alle verschieden sind, gibt es aber keine endlichen Körper der Charakteristik 0. $\diamond$

4.2.3. Der Kleine Fermatsche Satz *. Wir kommen noch einmal auf die endlichen Körper zurück und beginnen mit der folgenden simplen Beobachtung:

$$\begin{aligned} 2^7 - 2 &= 128 - 2 = 126 = 7 \cdot 18, \\ 3^7 - 3 &= 2187 - 3 = 2184 = 7 \cdot 312, \\ 4^{13} - 4 &= 67108860 = 13 \cdot 5162220. \end{aligned}$$

In allen Beispielen ist der Ausdruck $a^n - a$ durch n teilbar. (Siehe auch Beispiel 3.36.) Das funktioniert aber nicht immer:

$$2^9 - 2 = 512 - 2 = 510$$

ist nicht durch 9 teilbar.

Der strukturelle Unterschied zwischen diesen Beispielen ist, dass im ersten Fall der Exponent immer eine Primzahl war. In der Tat gilt:

SATZ 4.21 (Kleiner Fermatscher Satz). Sei p eine Primzahl und $a \in \mathbb{Z}$ irgendeine ganze Zahl. Dann ist $a^p - a$ durch p teilbar.

BEWEIS. Wenn wir $K = \mathbb{F}_p$ schreiben, dann können wir die Aussage umformulieren als

$$(a^p - a)_K = 0.$$

Es gilt aber $(a^p - a)_K = (a_K)^p - a_K$, also genügt es zu zeigen, dass für alle Elemente $x \in \mathbb{F}_p$ gilt: $x^p = x$.

Für $x = 0$ ist das klar. Sei also nun $x \neq 0$. Betrachte die Multiplikation mit x als Abbildung $\xi: \mathbb{F}_p^\times \rightarrow \mathbb{F}_p^\times: y \mapsto xy$. Diese Abbildung ist injektiv, weil $x \neq 0$ und $\mathbb{F}_p$ ein Körper ist. Weil $\mathbb{F}_p^\times$ nur endlich viele Elemente hat, ist sie automatisch sogar bijektiv.

Das bedeutet, dass das Produkt aller Bilder $\xi(y)$, $y \in \mathbb{F}_p^\times$, unter der Abbildung ξ mit dem Produkt aller $y \in \mathbb{F}_p^\times$ übereinstimmt, denn die Faktoren sind ja bis auf die Reihenfolge genau dieselben:

$$x^{p-1} \cdot (1 \cdot \dots \cdot (p-1)) = (x \cdot 1) \cdot (x \cdot 2) \cdot \dots \cdot (x(p-1)) = 1 \cdot \dots \cdot (p-1).$$

Die rechte Seite ist wieder in $\mathbb{F}_p^\times$, also $\neq 0$, und wir können durch sie teilen. Wir erhalten damit

$$x^{p-1} = 1,$$

also $x^p = x$, wie gewünscht. $\square$

BEMERKUNG 4.22. Eine kleine Verfeinerung (der Satz von Euler, Korollar 8.56) dieses Satzes ist ein wichtiger Bestandteil des [RSA-Verfahrens](#)³, eines der wichtigsten Public-Key-Verfahren, also Verschlüsselungsverfahren, bei denen Absender und Empfänger verschlüsselt kommunizieren können, ohne vorher einen geheimen Schlüssel auszutauschen. Stattdessen kann der Empfänger den Schlüssel, den der Absender zum Verschlüsseln benutzt, öffentlich machen (public key), ohne dass Außenstehende die Möglichkeit hätten, den öffentlichen Schlüssel zum Entschlüsseln zu verwenden. Siehe Bemerkung 8.58

Ausgangsbasis für alle solchen Public-Key-Verfahren ist eine mathematische Operation, die vergleichsweise schnell berechenbar und eindeutig umkehrbar ist, für die aber die Berechnung der Umkehroperation einen wesentlich höheren Rechenaufwand erfordert. Im Falle des RSA-Verfahrens benutzt man, dass es leicht ist, große Primzahlen zu finden und das Produkt zweier solcher Primzahlen zu berechnen, dass aber kein Verfahren bekannt ist, um in annehmbarer Zeit die Zerlegung eines solchen Produkts in seine Primfaktoren zu bestimmen. Konkret würde man zum Beispiel mit zwei Primzahlen beginnen, die jeweils etwa 1000 Stellen haben. Mittelfristig könnte die Entwicklung leistungsstarker Quantencomputer allerdings eine Möglichkeit darstellen, auch Produkte dieser Größenordnung in ihre Primfaktoren zu zerlegen, so dass das Verfahren dann als geknackt gelten müsste (siehe [Shor-Algorithmus](#)⁴). ◇

BEMERKUNG 4.23. Der Kleine Fermatsche Satz ist eine Möglichkeit zu testen, ob eine (große) Zahl n eine Primzahl ist. Man nimmt eine Zahl a und rechnet aus, ob $a^n - a$ durch n teilbar ist. Diese Rechnung lässt sich *wesentlich* schneller durchführen, als die Zahl n in ihre Primfaktoren zu zerlegen. Wenn $a^n - a$ nicht durch n teilbar ist, dann kann n keine Primzahl sein. Wenn für mehrere Zahlen a die Zahl $a^n - a$ durch n teilbar ist, dann ist die Wahrscheinlichkeit hoch, dass n eine Primzahl ist.

In der Praxis benutzt man etwas ausgefeiltere Tests, die aber oft letztlich auf dem Kleinen Fermatschen Satz beruhen, zum Beispiel den [Miller-Rabin-Test](#)⁵. Diese Tests sind in der Regel probabilistischer Natur, d.h. sie liefern mit einer verschwindend geringen Wahrscheinlichkeit ein falsches Ergebnis. Das nimmt man in der Praxis dann in Kauf.

Um eine große Primzahl (sagen wir mit mehreren hundert Stellen) zu finden, wie sie zum Beispiel für manche Verschlüsselungsverfahren benötigt wird, probiert man dann einfach so lange zufällig gewählte große Zahlen durch, bis man eine findet, die den zugrunde gelegten Primzahltest besteht. ◇

4.3. Polynomfunktionen

4.3.1. Summe und Produkt von Abbildungen in einen Körper. Sei K ein Körper. Wir bezeichnen mit $\text{Abb}(X, K)$ die Menge aller Abbildungen $X \rightarrow K$. Indem wir die Addition und Multiplikation auf K benutzen, können wir die Summe und das Produkt von zwei Abbildungen $f, g \in \text{Abb}(X, K)$ definieren:

$$f + g: X \rightarrow K, \quad x \mapsto f(x) + g(x),$$

und

$$f \cdot g: X \rightarrow K, \quad x \mapsto f(x)g(x).$$

Wir erhalten so eine Addition $+$ und eine Multiplikation $\cdot$ auf $\text{Abb}(X, K)$.

Zwar ist $\text{Abb}(X, K)$ kein Körper (es sei denn, X hat genau ein Element), aber das soll uns an dieser Stelle nicht stören.

³ <https://de.wikipedia.org/wiki/RSA-Kryptosystem>

⁴ <https://de.wikipedia.org/wiki/Shor-Algorithmus>

⁵ <https://de.wikipedia.org/wiki/Miller-Rabin-Test>

4.3.2. Polynomfunktionen auf einem Körper. Sei K weiter ein Körper.

DEFINITION 4.24. Eine Abbildung $f: K \rightarrow K$ heißt *Polynomfunktion*, wenn eine natürliche Zahl $n \geq 0$ und Elemente $a_i \in K, i = 0, \dots, n$, existieren, so dass für alle $x \in K$ gilt:

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = \sum_{i=0}^n a_i x^i.$$

⊢

Die Elemente a_i nennt man die *Koeffizienten* von f . (Aber siehe Bemerkung 4.29: Ist K endlich, so sind die Koeffizienten einer solchen Darstellung nicht eindeutig durch die Abbildung f bestimmt!) Der Koeffizient a_0 heißt der *Absolutterm* oder das *absolute Glied*. Dies ist gleichzeitig der Wert von f an der Stelle 0 : $f(0) = a_0$.

Es ist oft nützlich, den Fall, dass f die konstante Funktion mit Wert 0 ist, separat zu behandeln. Wenn wir diesen Fall ausschließen, brauchen wir nur noch Polynomfunktionen zu betrachten, bei denen mindestens ein Koeffizient $\neq 0$ ist. Dann können wir auch direkt annehmen, dass der Koeffizient mit dem höchsten Index nicht verschwindet, weil wir die Terme mit verschwindenden Koeffizienten weglassen können.

Ist $f: K \rightarrow K$ eine Abbildung, dann nennen wir die Elemente a von K mit $f(a) = 0$ die *Nullstellen* der Funktion f . Eine einfache aber wichtige Beobachtung ist die folgende: Ist f das Produkt von Polynomfunktionen $g, h: K \rightarrow K$ (im Sinne von Abschnitt 4.3.1, das heißt also einfach, dass $f(x) = g(x)h(x)$ für alle $x \in K$ gilt), dann ist f wieder eine Polynomfunktion, und für $a \in K$ gilt $f(a) = 0$ genau dann, wenn $g(a) = 0$ oder $h(a) = 0$ ist.

SATZ 4.25. Sei $f: K \rightarrow K$ eine Polynomfunktion. Gilt

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = \sum_{i=0}^n a_i x^i.$$

(siehe Definition 4.24) und sind nicht alle $a_i = 0$, dann gibt es höchstens n verschiedene Nullstellen von f in K .

Zumindest in Spezialfällen ist Ihnen der Satz bekannt; zum Beispiel kann eine quadratische Gleichung nicht mehr als zwei Lösungen haben. In diesem Semester werden die Polynomfunktionen keine zentrale Rolle spielen, sondern nur gelegentlich als Beispielmaterial dienen. Das wird sich in der Linearen Algebra 2 ändern und wir werden dann noch einmal systematischer auf diesen Begriff zurückkommen. Daher können Sie, wenn Sie möchten, den Beweis des Satzes erstmal überspringen.

BEWEIS. Ähnlich wie wir die Division mit Rest für ganze Zahlen haben, können wir auch Polynomfunktionen (für die wir eine Darstellung wie in der Definition gewählt haben) mit Rest durcheinander dividieren. Wir nennen das Polynomdivision.

LEMMA 4.26 (Polynomdivision). Seien f, g Polynomfunktionen mit $f(x) = \sum_{i=0}^n a_i x^i, g(x) = \sum_{i=0}^m b_i x^i$ und es gelte $a_n \neq 0, b_m \neq 0$ und $0 < m \leq n$. Dann existieren eine Polynomfunktion q mit $q(x) = \sum_{i=0}^{n-m} c_i x^i$ und eine Polynomfunktion r mit $r(x) = \sum_{i=0}^l d_i x^i$, so dass

$$f = qg + r, \quad \text{und } l < m.$$

(Die erste Bedingung ist im Sinne von Abschnitt 4.3.1; sie bedeutet einfach, dass $f(x) = q(x)g(x) + r(x)$ für alle $x \in K$. Es wird nicht behauptet, dass $c_{n-m} \neq 0$ sein muss.)

BEWEIS DES LEMMAS. Wir führen Induktion nach n . Ist $n < m$, so können wir einfach $r(x) := f(x) = \sum_{i=0}^n a_i x^i$ und $q = 0$ setzen.

Ist $n \geq m$, so sei $q_1(x) = \frac{a_n}{b_m} x^{n-m}$ und $f_1 = f - q_1 g$, also gilt für alle $x \in K$:

$$f_1(x) = f(x) - q_1(x)g(x) = \sum_{i=0}^n a_i x^i - \frac{a_n}{b_m} x^{n-m} \sum_{i=0}^m b_i x^i = \sum_{i=0}^{n-1} \left(a_i - \frac{a_n b_{i-n+m}}{b_m} \right) x^i,$$

wobei wir für $j < 0$ den Wert von b_j als 0 ansehen wollen. Wir können daher auf diese Darstellung von f_1 die Induktionsvoraussetzung anwenden und $f_1 = q_2 g + r$ schreiben, wobei die höchste Potenz von x , die in r auftritt, kleiner als m ist. Wir erhalten dann

$$f = f_1 + q_1 g = (q_1 + q_2)g + r,$$

und setzen $q := q_1 + q_2$, um die gewünschte Darstellung von f zu erhalten. $\square$

Nun zum Beweis des Satzes. Wir führen Induktion nach n . Ist $n = 0$, also $f = a_0$, so hat f gar keine Nullstellen, weil nach Voraussetzung nicht alle a_i (und das ist hier nur a_0) Null sein dürfen.

Sei nun $n > 0$. Wir können annehmen, dass $a_n \neq 0$ ist, denn sonst könnten wir den Term $a_n x^n$ einfach weglassen und ein noch besseres Ergebnis beweisen. Sei a eine Nullstelle von f . (Sollte f gar keine Nullstellen haben, ist die Aussage sowieso klar.) Wir wenden das Lemma über die Polynomdivision an auf das gegebene f und $g(x) = x - a$. Das bedeutet $m = 1$, also können wir f schreiben als $f = qg + r$ mit $r(x) = d_0$ (denn die Zahl l im Lemma muss < 1 sein).

Wir haben also $f(x) = q(x)(x - a) + d_0$ für alle $x \in K$. Wenn wir $x = a$ setzen, dann bekommen wir

$$0 = f(a) = q(a)(a - a) + d_0 = d_0,$$

also gilt $d_0 = 0$ und damit $f(x) = q(x)(x - a)$. Nun gibt uns das Lemma auch, dass $q(x) = \sum_{i=0}^{n-1} c_i x^i$. Nach Induktionsvoraussetzung hat q höchstens $n - 1$ Nullstellen in K . Ist b irgendeine Nullstelle von f , so gilt $q(b)(b - a) = f(b) = 0$, also ist $b = a$ oder $q(b) = 0$. Insgesamt folgt die Behauptung. $\square$

Natürlich kann f wie im Satz auch weniger als n Nullstellen haben, betrachten Sie zum Beispiel die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2 + 1$, oder allgemein die Funktion $f: K \rightarrow K, f(x) = x^n$, die 0 als einzige Nullstelle hat.

BEISPIEL 4.27. Sei $K = \mathbb{Q}$ und $f(x) = x^3 - 3x^2 + 3x - 2$. Wir sehen, dass f bei $x = 2$ eine Nullstelle hat. (Man kann zeigen, dass die Nullstellen in $\mathbb{Q}$ eines Polynoms mit Koeffizienten in $\mathbb{Z}$ und höchstem Koeffizient 1 alle in $\mathbb{Z}$ liegen und den Absolutterm teilen; in diesem Fall kommen also nur die Teiler von -2 , d.h. $-2, -1, 1, 2$, als Nullstellen in $\mathbb{Q}$ in Frage. Natürlich kann es weitere Nullstellen in $\mathbb{R} \setminus \mathbb{Q}$, oder auch in $\mathbb{C} \setminus \mathbb{Q}$, geben. Auch in diesem Beispiel gibt es nur eine rationale Nullstelle und zwei reelle Nullstellen, die nicht in $\mathbb{Q}$ liegen.)

Wir führen die Polynomdivision von f durch $x - 2$ durch:

$$\begin{array}{r} (x^3 - 3x^2 + 3x - 2) : (x - 2) = x^2 - x + 1 \\ \underline{-x^3 + 2x^2} \\ -x^2 + 3x - 2 \\ \underline{+x^2 - 2x} \\ x - 2 \\ \underline{-x + 2} \\ 0 \end{array}$$

(Das Vorgehen ist hier nicht so leicht zu illustrieren; versuchen Sie, die obige Rechnung schrittweise wie eine schriftliche Division zu lesen. Mit den unteren Zeilen ergeben sich nach und nach die Summanden des Ergebnisses.)

Wie im Beweis des Satzes bleibt in diesem Fall kein Rest. Die weiteren Nullstellen von f können wir nun ausrechnen, indem wir die quadratische Gleichung $x^2 - x + 1 = 0$ lösen. $\diamond$

KOROLLAR 4.28. *Sei K ein unendlicher Körper. Sei f eine Polynomfunktion auf K , die nicht konstant 0 ist. Dann sind die Zahlen n und a_i in der Darstellung $f(x) = \sum_{i=0}^n a_i x^i$ mit $a_n \neq 0$ eindeutig bestimmt.*

BEWEIS. Seien $f(x) = \sum_{i=0}^n a_i x^i = \sum_{i=0}^{n'} a'_i x^i$ Darstellungen von f , und sei ohne Einschränkung $n \geq n'$. Wir setzen $a'_i = 0$ für $i = n' + 1, \dots, n$ und betrachten die Differenz

$$\sum_{i=0}^n (a_i - a'_i) x^i.$$

Dies ist eine Polynomfunktion, die an jedem Element von K verschwindet, also unendlich viele Nullstellen hat. Ihre Koeffizienten müssen nach Satz 4.25 sämtlich verschwinden. Es folgt $a_i = a'_i$ für alle i , und wenn wir $a_n \neq 0$ voraussetzen, auch $n = n'$, denn im Fall $n > n'$ wäre nach unserer Definition $a'_n = 0 \neq a_n$. $\square$

Ist K ein unendlicher Körper, so nennen wir die Zahl n aus Korollar 4.28 den Grad der Polynomfunktion f . (Der Nullfunktion weisen wir formal den Grad $-\infty$ zu.)

BEMERKUNG 4.29. Ist K ein endlicher Körper, so kann eine Polynomfunktion f konstant $= 0$ sein, also $f(x) = 0$ für alle $x \in K$, und dennoch eine Darstellung $f(x) = \sum_{i=0}^n a_i x^i$ haben, in der nicht alle a_i verschwinden. Die Aussage von Korollar 4.28 ist in diesem Fall falsch. Zum Beispiel ist für $K = \mathbb{F}_p$ die Polynomfunktion $f(x) = x^p - x$ die Nullfunktion; dies ist gerade die Aussage des Kleinen Fermatschen Satzes (Satz 4.21). Wir kommen in der Linearen Algebra 2 auf diesen Punkt noch einmal zurück. $\diamond$

Lineare Gleichungssysteme und Matrizen

5.1. Lineare Gleichungssysteme

5.1.1. Definitionen. Wir halten im gesamten Abschnitt einen Körper K fest. Im Kontext der linearen Gleichungssysteme ist es nützlich, Elemente von K^n , also n -Tupel von Elementen von K , nicht wie üblich als Tupel $(x_1, \dots, x_n)$ zu schreiben, sondern in einer Spalte:

$$\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}.$$

Das von vorneherein so zu tun, erspart uns eine spätere Umstellung der Notation (und damit mögliche Verwirrung beim späteren Zurückblättern). Die Begründung für diese Konvention liegt im Formalismus des Matrizenprodukts, das wir in Abschnitt 5.3 kennenlernen werden.

Da diese *Spaltenvektoren* allerdings im Text sehr viel Platz wegnehmen, wollen wir stattdessen die Notation $(x_1, \dots, x_n)^t$ verwenden, wo das kleine t (» transponiert«) bedeuten soll, dass wir uns dieses Tupel als Spalte denken.

(Eine pragmatische Lösung wäre, einfach das t zu ignorieren, bis Sie zu Bemerkung 5.36 im Abschnitt 5.3 kommen.)

DEFINITION 5.1. Sei K ein Körper, seien m, n natürliche Zahlen. Ein *lineares Gleichungssystem* (LGS) von m Gleichungen in n Unbestimmten $X_1, \dots, X_n$ über K ist gegeben durch m Gleichungen

$$\begin{aligned} a_{11}X_1 + a_{12}X_2 + \cdots + a_{1n}X_n &= b_1 \\ a_{21}X_1 + a_{22}X_2 + \cdots + a_{2n}X_n &= b_2 \\ &\vdots \\ a_{m1}X_1 + a_{m2}X_2 + \cdots + a_{mn}X_n &= b_m \end{aligned}$$

mit $a_{ij} \in K$ für alle $i = 1, \dots, m, j = 1, \dots, n, b_i \in K, i = 1, \dots, m$.

Ein lineares Gleichungssystem heißt *homogenes* lineares Gleichungssystem, falls $b_i = 0$ für alle $i = 1, \dots, m$. Sonst heißt das lineare Gleichungssystem *inhomogen*.

Wir schreiben kurz, das obige lineare Gleichungssystem sei gegeben durch die Elemente (a_{ij}, b_i) von K .

Die Elemente a_{ij} von K heißen auch die *Koeffizienten* des linearen Gleichungssystems. Genauer ist a_{ij} der Koeffizient von X_j in der i -ten Gleichung des Gleichungssystems. Ein Element $(x_1, \dots, x_n)^t$ von K^n heißt *Lösungsvektor* des linearen Gleichungssystems, wenn nach Einsetzen von x_i für X_i (für alle i) alle m Gleichungen erfüllt sind.

Die *Lösungsmenge* $\mathbb{L}$ des linearen Gleichungssystems ist die Menge aller Lösungsvektoren, also die Menge aller n -Tupel $(x_1, \dots, x_n)^t \in K^n$, so dass für alle $i = 1, \dots, m$ gilt:

$$a_{i1}x_1 + a_{i2}x_2 + \cdots + a_{in}x_n = b_i.$$

Sei ein lineares Gleichungssystem M gegeben durch (a_{ij}, b_i) . Das lineare Gleichungssystem M_0 , das gegeben ist durch $(a_{ij}, 0)$, d.h. die Koeffizienten a_{ij} werden beibehalten, aber alle b_i werden durch 0 ersetzt, heißt das zu M gehörige *homogene lineare Gleichungssystem*.

Das Wort »linear« bezieht sich hier darauf, dass die Unbestimmten alle nur in der ersten Potenz auftreten (keine Quadrate oder höheren Potenzen und keine Produkte von Unbestimmten). Geometrisch gesehen ist die Lösungsmenge einer Gleichung der Form $aX_1 + bX_2 = c$ über den reellen Zahlen eine Gerade in $\mathbb{R}^2$, und auch in allgemeineren Fällen sind, wie sich zeigen wird, die Lösungsmengen linearer Gleichungssysteme »lineare Objekte«.

Wir nennen ein lineares Gleichungssystem *lösbar*, wenn die Lösungsmenge nicht leer ist, d.h. wenn es überhaupt einen Lösungsvektor gibt. Sonst heißt das System *unlösbar*. Gibt es genau einen Lösungsvektor, d.h. hat die Lösungsmenge genau ein Element, dann sprechen wir von einem *eindeutig lösbaren* System.

Jedes homogene lineare Gleichungssystem ist lösbar, denn $(0, \dots, 0)^t$ ist ein Lösungsvektor. Diese Lösung wird auch die *triviale Lösung* genannt. Umgekehrt gilt: Wenn ein lineares Gleichungssystem $(0, \dots, 0)^t$ als Lösungsvektor hat, dann muss es ein homogenes Gleichungssystem sein.

BEISPIEL 5.2. Sei $K = \mathbb{R}$. Wir betrachten als Beispiel das folgende Gleichungssystem:

$$\begin{aligned} 2X_1 - 2X_2 &= 3 \\ X_1 + 2X_2 &= 3 \end{aligned}$$

Addieren wir die beiden Gleichungen, so sehen wir, dass $3X_1 = 6$, also $X_1 = 2$ gelten muss. Durch Einsetzen in die zweite Gleichung bekommen wir dann $X_2 = \frac{1}{2}$, und man sieht direkt, dass dies tatsächlich eine Lösung des Gleichungssystems ist: Die Lösungsmenge ist

$$\left\{ \left(2, \frac{1}{2} \right)^t \right\}.$$

Vergleiche auch Abschnitt 2.5.

Wenn wir $\mathbb{R}^2$ als die Zahlenebene sehen, dann können wir jede Teilmenge als eine Menge von Punkten in dieser Ebene betrachten. Zum Beispiel ist die Menge

$$\{(x_1, x_2); 2x_1 - 2x_2 = 3\}$$

eine Gerade in $\mathbb{R}^2$, siehe die Abbildung.

In ähnlicher Weise ist

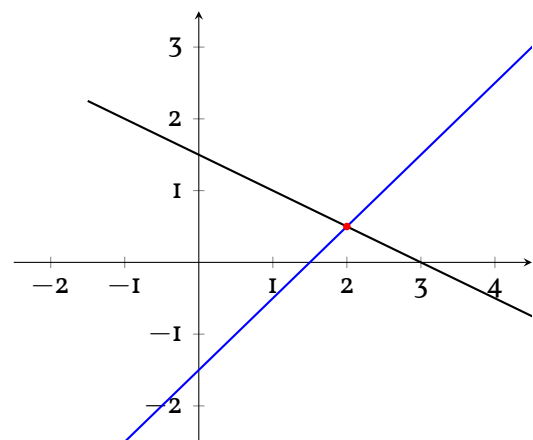
$$\{(x_1, x_2); x_1 + 2x_2 = 3\}$$

eine Gerade. Und die Lösungsmenge des obigen Gleichungssystems ist einfach die Schnittmenge dieser beiden Geraden.

Das ist auch allgemeiner richtig, dann für alle $a_1, a_2, b \in \mathbb{R}$ ist

$$\{(x_1, x_2); a_1x_1 + a_2x_2 = b\}$$

eine Gerade – es sei denn es ist $a_1 = a_2 = 0$; dann ist diese Menge entweder gleich $\mathbb{R}^2$ (wenn $b = 0$) oder leer (wenn $b \neq 0$).



Da zwei Geraden in $\mathbb{R}^2$ entweder gleich sind oder sich in höchstens einem Punkt schneiden, sehen wir, dass es kein lineares Gleichungssystem mit 2 Gleichungen und 2 Unbestimmten

über $\mathbb{R}$ gibt, das mehr als eine, aber nur endlich viele Lösungen hat. (Wie ist es mit mehr Gleichungen/Unbestimmten? Über anderen Körpern? – Später wird es uns leicht fallen, diese Fragen zu beantworten, zum Teil können Sie es vielleicht jetzt schon?) $\diamond$

In ähnlicher Weise wie im vorherigen Beispiel können wir die Lösungsmengen von linearen Gleichungssystemen in $\mathbb{R}^3$ als Durchschnitte von Ebenen betrachten (und, wenn Ihre Anschauung das mitmacht, kann man die Situation auch für mehr als 3 Unbestimmte ähnlich sehen.)

5.1.2. Addition und Skalarmultiplikation auf K^n . Um die Struktur der Lösungsmengen besser zu verstehen, ist es hilfreich, auch auf der Menge K^n Rechenoperationen einzuführen, und zwar eine Addition und eine »Skalarmultiplikation« (also Multiplikation mit Elementen des Körpers K , die auch als *Skalare* bezeichnet werden).

DEFINITION 5.3. Für $(x_1, \dots, x_n)^t, (y_1, \dots, y_n)^t \in K^n$ definieren wir die Summe

$$(x_1, \dots, x_n)^t + (y_1, \dots, y_n)^t := (x_1 + y_1, \dots, x_n + y_n)^t,$$

und für $a \in K$ das Produkt von a mit $(x_1, \dots, x_n)^t$ durch

$$a \cdot (x_1, \dots, x_n)^t := (ax_1, \dots, ax_n)^t.$$

—

BEISPIEL 5.4. (1) Ist $K = \mathbb{C}$, so ist zum Beispiel

$$i \begin{pmatrix} 1 \\ 1+i \end{pmatrix} + \begin{pmatrix} i \\ 1-i \end{pmatrix} = \begin{pmatrix} 2i \\ 0 \end{pmatrix}.$$

(2) Ist $K = \mathbb{F}_7$, dann ist zum Beispiel

$$2 \begin{pmatrix} 5 \\ 2 \end{pmatrix} + \begin{pmatrix} 4 \\ 3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

$\diamond$

Wir nennen die Elemente von K^n auch *Vektoren*. Da die Addition separat in den einzelnen Komponenten erfolgt, ist klar, dass das Assoziativgesetz und das Kommutativgesetz gelten, und dass der sogenannte *Nullvektor* $(0, \dots, 0)^t$ ein neutrales Element für die Addition ist (und zwar das einzige). Wir bezeichnen den Nullvektor meist einfach mit 0 (und damit bleibt der Leser*in die Aufgabe überlassen festzustellen, welches Objekt mit dem Symbol 0 eigentlich gerade gemeint ist). Ist $x = (x_1, \dots, x_n)^t \in K^n$, so ist $-x := (-x_1, \dots, -x_n) = (-1) \cdot x$ das inverse Element von x bezüglich der Addition. Natürlich stellt sich nicht die Frage, ob es sich bei K^n für $n > 1$ um einen Körper handelt, weil wir keine Multiplikation $K^n \times K^n \rightarrow K^n$ definiert haben. Immerhin ist klar, dass eine Art Assoziativgesetz gilt, das heißt $a(bx) = (ab)x$ für $a, b \in K, x \in K^n$.

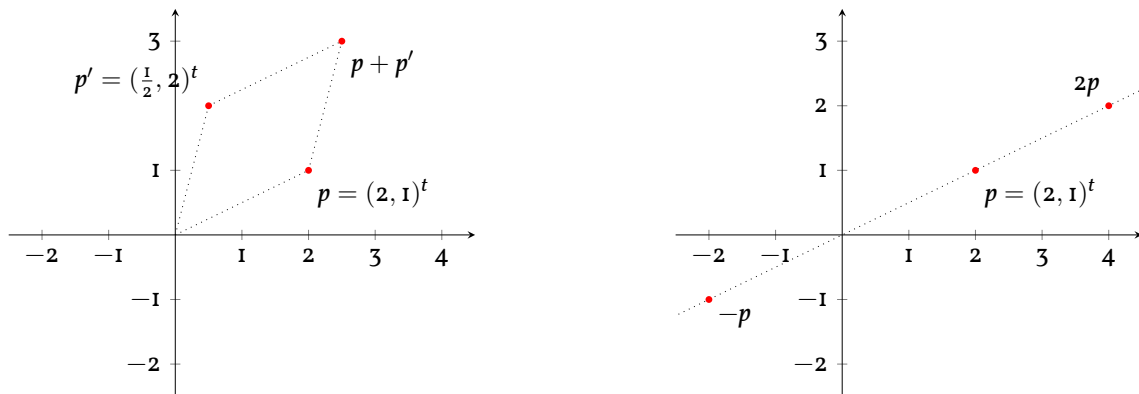
BEMERKUNG 5.5. Die Addition und Skalarmultiplikation in $\mathbb{R}^2$ können wir geometrisch veranschaulichen, vergleiche die Abbildungen: Für $p, p' \in \mathbb{R}^2$ sind $0, p, p'$ und $p + p'$ die vier Eckpunkte eines Parallelogramms.

Für $p \in \mathbb{R}^2 \setminus \{0\}$ liegen alle Punkte $ap, a \in \mathbb{R}$ auf der Geraden durch den Ursprung $(0, 0)^t$ und p .

Selber ausprobieren^a.

^a<https://math.ug/applets/R2.html>

Die Definition der Addition und Skalarmultiplikation sind also der erste Schritt, »mit Punkten zu rechnen«, und damit der Beginn der analytischen Geometrie (siehe Kapitel 11).



Im Prinzip gelten dieselben Beschreibungen in $\mathbb{R}^3$ und allgemein in $\mathbb{R}^n$, allerdings ist es dann weniger leicht, dies an Abbildungen zu illustrieren. $\diamond$

5.1.3. Matrizen. Es ist praktisch, die Koeffizienten eines linearen Gleichungssystems zu »organisieren«, ohne dass wir immer alle Unbestimmten mit ausschreiben müssen (wie wir die Unbestimmten nennen, spielt ja ohnehin keine Rolle). Dafür führen wir den Begriff der Matrix ein. Im Moment ist das eine reine Organisationshilfe für uns, später werden wir allerdings viele, teilweise auch tiefere, Anwendungen sehen, wo Matrizen sehr nützlich sind.

DEFINITION 5.6 (Matrix). Seien m, n natürliche Zahlen. Unter einer *Matrix* der Größe $m \times n$ mit Einträgen in K (man spricht auch von einer $(m \times n)$ -Matrix über K) verstehen wir eine Familie $(a_{ij})_{i=1, \dots, m, j=1, \dots, n}$ von Elementen von K . Die Elemente a_{ij} heißen die *Einträge* oder *Koeffizienten* der Matrix. Eine Matrix stellen wir uns immer in der Form

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

vor, d.h. wir schreiben die Koeffizienten in einem rechteckigen Schema auf, in dem der erste Index die Zeile und der zweite Index die Spalte angibt. In der linken oberen Ecke steht der Eintrag a_{11} , rechts daneben $a_{12}, \dots$, und in der i -ten Zeile und j -ten Spalte steht a_{ij} .

Die Menge aller $(m \times n)$ -Matrizen über K bezeichnen wir mit $M_{m \times n}(K)$. Die Menge aller $(n \times n)$ -Matrizen bezeichnen wir manchmal mit $M_n(K)$ statt mit $M_{n \times n}(K)$; wir sprechen auch von *quadratischen* Matrizen. $\dashv$

Formal betrachtet können wir eine Matrix als eine Abbildung $M: \{1, \dots, m\} \times \{1, \dots, n\} \rightarrow K$ definieren, dann ist a_{ij} in der obigen Definition der Wert $M(i, j)$. Rein formal können wir dann auch die Fälle erlauben, in denen $m = 0$, oder $n = 0$ oder $m = n = 0$ ist. Mit der Konvention $\{1, \dots, m\} = \emptyset$, falls $m = 0$ (und entsprechend für n) ist eine Matrix dann eine Abbildung $\emptyset \rightarrow K$, also gibt es in diesen Fällen genau eine Matrix (die »leere« Matrix der Größe $0 \times n$, $m \times 0$ oder 0×0). Die *Nullmatrix* (der Größe $m \times n$) ist die Matrix $0 \in M_{m \times n}(K)$, deren Einträge alle gleich Null sind. Es gibt also für jede Wahl von m und n eine Nullmatrix, und diese sind natürlich alle verschieden, da sie verschiedene Größen haben. Trotzdem schreibt man meist einfach 0 dafür, oder $0_{m \times n}$, wenn man die Größe mitangeben möchte. Wir können Elemente von K^n auch als Matrizen mit einer einzigen Spalte und n Zeilen betrachten: $K^n = M_{n \times 1}(K)$.

Gegeben ein lineares Gleichungssystem

$$\begin{aligned} a_{11}X_1 + a_{12}X_2 + \cdots + a_{1n}X_n &= b_1 \\ a_{21}X_1 + a_{22}X_2 + \cdots + a_{2n}X_n &= b_2 \\ &\vdots \\ a_{m1}X_1 + a_{m2}X_2 + \cdots + a_{mn}X_n &= b_m \end{aligned}$$

mit $a_{ij} \in K$ für alle $i = 1, \dots, m, j = 1, \dots, n, b_i \in K, i = 1, \dots, m$, so bezeichnen wir die Matrix

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

als die *Koeffizientenmatrix* des linearen Gleichungssystems, und die Matrix

$$\left(\begin{array}{cccc|c} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ a_{21} & a_{22} & \cdots & a_{2n} & b_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{array} \right)$$

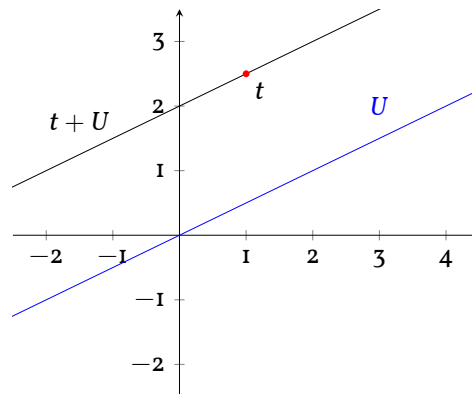
als die *erweiterte Koeffizientenmatrix* des linearen Gleichungssystems. (Der senkrechte Strich vor der letzten Spalte ist nur eine Erinnerung, dass diese Spalte von der rechten Seite der Gleichungen eines linearen Gleichungssystems herkommt, und hat, was die Matrix betrifft, keine mathematische Bedeutung.)

Wir führen für $t \in K^n$ und eine Teilmenge $U \subseteq K^n$ die folgende Schreibweise ein:

$$t + U := \{t + u; u \in U\}.$$

BEISPIEL 5.7. Ein Beispiel in $\mathbb{R}^2$, wo U die Gerade durch den Ursprung mit Steigung $\frac{1}{2}$ und t der Punkt $(1, \frac{5}{2})^t$ ist. $\diamond$

Damit können wir formulieren, wie die Lösungsmenge eines linearen Gleichungssystems und die Lösungsmenge des zugehörigen homogenen linearen Gleichungssystems zusammenhängen.



SATZ 5.8. Seien K ein Körper, $A \in M_{m \times n}(K)$ und $b \in K^m$. Sei $\mathbb{L}$ die Lösungsmenge des linearen Gleichungssystems mit erweiterter Koeffizientenmatrix $(A | b)$, sei $\mathbb{L}_0$ die Lösungsmenge des zugehörigen homogenen linearen Gleichungssystems mit Koeffizientenmatrix A . Ist $t \in \mathbb{L}$, so gilt

$$\mathbb{L} = t + \mathbb{L}_0.$$

BEWEIS. Dass x in $t + \mathbb{L}_0$ ist, ist äquivalent dazu, dass $x - t \in \mathbb{L}_0$ ist. Deshalb ist die behauptete Gleichheit von Mengen äquivalent zu der Behauptung

$$x \in \mathbb{L} \Leftrightarrow x - t \in \mathbb{L}_0$$

für alle $x = (x_1, \dots, x_n)^t \in K^n$. Nun bedeutet $x \in \mathbb{L}$ gerade, dass für alle $i = 1, \dots, m$ gilt:

$$\sum_{j=1}^n a_{ij}x_j = b_i,$$

und $x - t \in \mathbb{L}_0$ heißt, dass für alle i

$$\sum_{j=1}^n a_{ij}(x_j - t_j) = 0,$$

wobei wir mit t_j die Einträge des Vektors t bezeichnen. Nun ist nach Voraussetzung $t \in \mathbb{L}$, und deshalb gilt

$$\sum_{j=1}^n a_{ij}(x_j - t_j) = \sum_{j=1}^n a_{ij}x_j - \sum_{j=1}^n a_{ij}t_j = \sum_{j=1}^n a_{ij}x_j - b_j.$$

Die behauptete Äquivalenz ist damit klar. $\square$

Es ist klar, dass die Lösungsmenge eines (inhomogenen) linearen Gleichungssystems leer sein kann. Dann existiert gar kein t wie im Satz, und der Satz liefert über diesen Fall keine Informationen.

Zum Schluss noch eine Sprechweise: Die Einträge a_{ii} einer quadratischen Matrix $A = (a_{ij})_{i,j}$ heißen die *Diagonaleinträge*. Die »Felder« mit Indizes $(1, 1)$, $(2, 2)$, ... nennt man auch die (Haupt-)Diagonale einer Matrix. Die $(n \times n)$ -Matrix, deren Diagonaleinträge alle gleich 1, und deren andere Einträge alle gleich 0 sind, heißt die Einheitsmatrix (der Größe n) und wird mit E_n bezeichnet.

$$E_4 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

5.2. Der Gauß-Algorithmus

5.2.1. Elementare Zeilenumformungen. Die Standardmethode, um die Lösungsmenge eines linearen Gleichungssystems zu bestimmen, ist der *Gauß-Algorithmus*. Das Ziel des Algorithmus ist es, ein gegebenes lineares Gleichungssystem durch Äquivalenzumformungen auf ein lineares Gleichungssystem möglichst einfacher Gestalt zu bringen, das dieselbe Lösungsmenge hat wie das ursprüngliche lineare Gleichungssystem.

BEISPIEL 5.9. Um das Ziel zu illustrieren, hier zwei Beispiele.

- (1) Am einfachsten ist der Fall, dass die Koeffizientenmatrix des Gleichungssystems die Einheitsmatrix E_n ist (Einsen auf der Diagonale, Nullen überall sonst). Die Lösungsmenge des Gleichungssystems mit erweiterter Koeffizientenmatrix $(E_n \mid b)$ ist $\{b\}$.
- (2) Eine ganz so einfache Form wie in Teil (1) werden wir nicht immer erreichen können (jedenfalls kann das ja nur für eindeutig lösbare Gleichungssysteme die richtige Lösungsmenge geben).

Betrachten wir als weiteres Beispiel das Gleichungssystem mit erweiterter Koeffizientenmatrix (mit Elementen a und b_i im fixierten Körper K)

$$\left(\begin{array}{ccc|c} 1 & a & 0 & b_1 \\ 0 & 0 & 1 & b_2 \\ 0 & 0 & 0 & b_3 \end{array} \right).$$

In diesem Fall ist es immer noch einfach, die Lösungsmenge abzulesen. Ist $b_3 \neq 0$, so ist die Lösungsmenge leer.

Sei nun $b_3 = 0$. Die dritte Zeile besteht dann sämtlich aus Nullen und ist als Gleichung gesehen die Gleichung $0 = 0$, die immer erfüllt ist. Die zweite Zeile besagt dann als Gleichung ausgeschrieben, dass $x_3 = b_2$. Die erste Zeile beschreibt die Gleichung $x_1 + ax_2 = b_1$, also $x_1 = b_1 - ax_2$. Das bedeutet, dass wir für jede Wahl von x_2 in K den

Lösungsvektor $(b_1 - ax_2, x_2, b_2)^t$ erhalten, und alle Lösungsvektoren entstehen auf diese Art und Weise. Die Lösungsmenge ist

$$\{(b_1 - ax_2, x_2, b_2)^t; x_2 \in K\}.$$

Auch in diesem Fall braucht man also nicht weiter zu rechnen, sondern kann die Lösungsmenge aus der erweiterten Koeffizientenmatrix direkt ablesen. Das streben wir auch im allgemeinen Fall an.

◇

Wir beginnen nun damit, die Umformungen zu beschreiben, die wir erlauben werden, um das gegebene lineare Gleichungssystem zu verändern. Danach überlegen wir uns, dass diese Umformungen die Lösungsmenge nicht verändern.

DEFINITION 5.10 (Elementare Zeilenumformungen eines LGS). Gegeben ein lineares Gleichungssystem, so nennen wir die folgenden Umformungen *elementare Zeilenumformungen vom Typ I, II bzw. III*:

- (I) Addition eines Vielfachen (mit $a \in K$) einer Gleichung zu einer anderen Gleichung.
- (II) Vertauschung zweier Gleichungen.
- (III) Multiplikation einer Gleichung mit einem Skalar $a \in K^\times$.

⊢

In (I) kann man natürlich statt von der Addition des a -fachen auch von der Subtraktion des $(-a)$ -fachen sprechen.

BEISPIEL 5.11. Betrachten Sie das folgende lineare Gleichungssystem über $\mathbb{Q}$:

$$\begin{aligned} X - Y &= 3 \\ -2X + 3Y &= 0 \end{aligned}$$

Wir führen eine Zeilenumformung vom Typ I durch und zwar addieren wir das 2-fache der ersten Zeile zur zweiten Zeile. Wir erhalten das Gleichungssystem

$$\begin{aligned} X - Y &= 3 \\ Y &= 6. \end{aligned}$$

◇

Analog haben wir den Begriff der elementaren Zeilenumformungen vom Typ I, II, III von Matrizen:

- (I) Addition eines Vielfachen (mit $a \in K$) einer Zeile zu einer anderen Zeile.
- (II) Vertauschung zweier Zeilen.
- (III) Multiplikation einer Zeile mit einem Skalar $a \in K^\times$ (d.h. jeder Eintrag der Zeile wird mit a multipliziert).

Das bedeutet, dass eine elementare Zeilenumformung eines linearen Gleichungssystem genau der gleichen elementaren Zeilenumformung der erweiterten Koeffizientenmatrix entspricht.

(Entsprechend kann man auch von elementaren Spaltenumformungen von Matrizen sprechen. Das wird aber für uns erst später in der Vorlesung eine Rolle spielen.)

Wir benutzen für diese Zeilenumformungen die folgende Notation:

- Addition des a -fachen von Zeile j zu Zeile $i \neq j$: $Z_i \rightsquigarrow Z_i + aZ_j$,
- Vertauschen der Zeilen i und j : $Z_i \longleftrightarrow Z_j$,
- Multiplizieren von Zeile i mit $a \in K, a \neq 0$: $Z_i \rightsquigarrow aZ_i$.

Wenn ein lineares Gleichungssystem (oder analog eine Matrix) aus einer anderen Matrix durch eine elementare Zeilenumformung entsteht, dann kann man das ursprüngliche lineare Gleichungssystem (bzw. die ursprüngliche Matrix) durch eine elementare Zeilenumformung wieder zurückerhalten. Wir sagen deshalb, dass elementare Zeilenumformungen *umkehrbar* sind. In der Tat, entsteht A' aus A durch Addition des a -fachen von Zeile j zu Zeile i , so entsteht A aus A' durch Addition des $(-a)$ -fachen von Zeile j zu Zeile i . Entsteht A' aus A durch Vertauschen der Zeilen i und j , so entsteht auch A aus A' durch Vertauschen dieser Zeilen. Entsteht schließlich A' aus A durch Multiplikation der i -ten Zeile mit $a \neq 0$, so entsteht A aus A' durch Multiplikation der i -ten Zeile mit a^{-1} . (Hier ist es wichtig, dass wir für Typ (III) nur Elemente $\neq 0$ in K als Faktoren erlauben.)

LEMMA 5.12. *Zwei lineare Gleichungssysteme, die durch elementare Zeilenumformungen auseinander hervorgehen, haben dieselbe Lösungsmenge.*

BEWEIS. Sei $(A \mid b)$ die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems, und gehe das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A' \mid b')$ daraus durch eine elementare Zeilenumformung hervor. Wegen der Umkehrbarkeit der Operationen genügt es zu zeigen, dass jede Lösung von $(A \mid b)$ auch eine Lösung von $(A' \mid b')$ ist. Das ist offensichtlich für Umformungen vom Typ (II) und (III). Für Zeilenumformungen vom Typ (I), etwa $Z_i \rightsquigarrow Z_i + aZ_j$, haben wir für ein Element $x = (x_1, \dots, x_n)^t$ der Lösungsmenge des ursprünglichen Gleichungssystems die Ausgangssituation

$$\sum_{k=1}^n a_{ik}x_k = b_i, \quad \sum_{k=1}^n a_{jk}x_k = b_j,$$

und sehen, dass

$$\sum_{k=1}^n (a_{ik} + aa_{jk})x_k = \sum_{k=1}^n a_{ik}x_k + a \sum_{k=1}^n a_{jk}x_k = b_i + ab_j,$$

das bedeutet, dass x auch die neu erhaltene Gleichung erfüllt. (Wie man an der Rechnung sieht und wie auch direkt einsichtig ist, ist es natürlich wichtig, die Zeilenumformungen immer auch in der ganz rechten Spalte der erweiterten Koeffizientenmatrix anzuwenden.) $\square$

5.2.2. Die Zeilenstufenform. Wir beschreiben als nächstes die spezielle Form von Matrizen, die wir mit dem Gauß-Algorithmus erreichen wollen. (Danach werden wir uns überlegen, dass man jede Matrix durch eine Folge von elementaren Zeilenumformungen auf diese Form bringen kann; das ist der eigentliche Gauß-Algorithmus.)

DEFINITION 5.13. Sei A eine Matrix mit m Zeilen und n Spalten.

- (I) Wir sagen, die Matrix A habe *Zeilenstufenform*, wenn die folgenden Bedingungen erfüllt sind:
- (a) In jeder Zeile ist der erste Eintrag, der von Null verschieden ist, gleich 1. Diese erste Eins (von links gesehen), bezeichnen wir als die *führende Eins* der entsprechenden Zeile.
 - (b) Alle Einträge in der Spalte einer führenden Eins, die unter der führenden Eins liegen, sind gleich Null.

- (c) Die führende Eins einer Zeile liegt rechts von der führenden Eins der darüberliegenden Zeile.

(Es ist in Teil (a) erlaubt, dass eine Zeile nur aus Nullen besteht. Wegen Teil (c) müssen dann aber auch alle darunterliegenden Zeilen sämtlich aus Nullen bestehen. »Gleiche Höhe« ist in Bedingung (c) nicht ausreichend, die führende Eins einer Zeile darf nicht in derselben Spalte sein wie die der darüberliegenden Zeile.)

- (2) Wir sagen, die Matrix A habe *reduzierte Zeilenstufenform*, wenn sie Zeilenstufenform hat und zusätzlich *alle* Einträge in einer Spalte einer führenden Eins, abgesehen von der führenden Eins selbst, gleich Null sind (also auch die Einträge *über* der führenden Eins).

—

BEISPIEL 5.14. Schematisch dargestellt bedeutet Zeilenstufenform, dass die Matrix die folgende Gestalt hat:

$$\left(\begin{array}{cccccccccccc} \boxed{\text{I}} & * & * & * & * & * & * & * & * & * & * & * \\ & & & & \boxed{\text{I}} & * & * & * & * & * & * & * \\ & & & & & & & & \boxed{\text{I}} & * & * & * \\ & & & & & & & & & & \boxed{\text{I}} & * & * \end{array} \right)$$

Hier sind die »Stufen« eingezeichnet. Die führenden Einsen sind grau hinterlegt. An den mit Sternchen $*$ markierten Stellen dürfen beliebige Elemente von K stehen (selbstverständlich auch 0 und 1). Im leeren Bereich unten links stehen nur Nullen.

In der reduzierten Zeilenstufenform wird zusätzlich gefordert, dass über den *führenden Einsen* nur Nullen stehen:

$$\left(\begin{array}{cccccccccccc} \boxed{\text{I}} & * & * & & * & * & * & & * & * & & \\ & & & & \boxed{\text{I}} & * & * & * & & * & * & \\ & & & & & & & & \boxed{\text{I}} & * & * & \\ & & & & & & & & & & \boxed{\text{I}} & * & * \end{array} \right)$$

Über anderen, »nicht-führenden« Einsen (an den Sternchen-Stellen) müssen natürlich nicht unbedingt Nullen stehen. $\diamond$

THEOREM 5.15 (Gauß-Algorithmus). *Seien K ein Körper und m, n natürliche Zahlen. Jede Matrix $A \in M_{m \times n}(K)$ kann durch wiederholte Anwendung elementarer Zeilenumformungen in eine Matrix in reduzierter Zeilenstufenform überführt werden.*

BEWEIS. Wir erklären zuerst, wie sich eine Matrix A in Zeilenstufenform überführen lässt. Dazu führen wir Induktion nach Anzahl der Zeilen. Der Induktionsanfang (wahlweise für 0 Zeilen oder eine Zeile) ist klar.

Bringe nun eine der Zeilen, die mit einer minimalen Anzahl von Nullen beginnen, durch eine Zeilenvertauschung in die erste Zeile. Sei a der erste von Null verschiedene Eintrag der neuen ersten Zeile; er liege in Spalte j . Multipliziere die erste Zeile mit a^{-1} , so dass also der erste von Null verschiedene Eintrag der neuen ersten Zeile gleich 1 ist:

$$\begin{pmatrix} \mathbf{I} & * & * & * & * & * \\ * & * & * & * & * & * \\ * & * & * & * & * & * \\ * & * & * & * & * & * \end{pmatrix}$$

Durch dieses Vorgehen sind alle Einträge in Spalten mit Index $< j$ gleich Null. Gibt es in den Zeilen 2, 3, ... in der j -ten Spalte Einträge $\neq 0$, so ziehe nun geeignete Vielfache der ersten Zeile von diesen Zeilen ab, um eine Matrix zu erhalten, in denen die 1 in der ersten Zeile der einzige Eintrag in Spalte j ist, der von Null verschieden ist. Bezeichne die Matrix, die wir so erhalten, mit A' :

$$\begin{pmatrix} \mathbf{I} & * & * & * & * & * \\ * & * & * & * & * & * \\ * & * & * & * & * & * \\ * & * & * & * & * & * \end{pmatrix}$$

Nach Induktionsvoraussetzung wissen wir bereits, dass wir die Matrix, die durch die Zeilen 2, 3, ..., m von A' gegeben ist, durch elementare Zeilenumformungen in Zeilenstufenform bringen können. Wenn wir dieselben Umformungen auf die Matrix A' anwenden (und die erste Zeile von A' unverändert lassen), erhalten wir eine Matrix in Zeilenstufenform.

Es bleibt nun nur noch zu zeigen, dass sich eine Matrix in Zeilenstufenform durch elementare Zeilenumformungen in reduzierte Zeilenstufenform überführen lässt. Das können wir aber offensichtlich erreichen, indem wir jeweils geeignete Vielfache der Zeilen mit führenden Einsen von den Zeilen darüber abziehen. $\square$

BEISPIEL 5.16. Beispiele für die Durchführung des Gauß-Algorithmus.

(1) Sei $K = \mathbb{Q}$ und

$$A = \begin{pmatrix} 0 & 1 & 2 \\ 3 & 4 & 5 \\ 6 & 7 & 8 \end{pmatrix}.$$

Wir bringen A mit dem Gauß-Algorithmus zunächst auf Zeilenstufenform.

$$\begin{pmatrix} 0 & 1 & 2 \\ 3 & 4 & 5 \\ 6 & 7 & 8 \end{pmatrix} \xrightarrow{Z_1 \leftrightarrow Z_2} \begin{pmatrix} 3 & 4 & 5 \\ 0 & 1 & 2 \\ 6 & 7 & 8 \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow \frac{1}{3}Z_1} \begin{pmatrix} 1 & \frac{4}{3} & \frac{5}{3} \\ 0 & 1 & 2 \\ 6 & 7 & 8 \end{pmatrix} \xrightarrow{Z_3 \rightsquigarrow Z_3 - 6Z_1} \begin{pmatrix} 1 & \frac{4}{3} & \frac{5}{3} \\ 0 & 1 & 2 \\ 0 & -1 & -2 \end{pmatrix} \xrightarrow{Z_3 \rightsquigarrow Z_3 + Z_2} \begin{pmatrix} 1 & \frac{4}{3} & \frac{5}{3} \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{pmatrix}$$

Schließlich können wir noch die reduzierte Zeilenstufenform herstellen:

$$\begin{pmatrix} 1 & \frac{4}{3} & \frac{5}{3} \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow Z_1 - \frac{4}{3}Z_2} \begin{pmatrix} 1 & 0 & -1 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{pmatrix}.$$

(2) Sei $K = \mathbb{C}$, sei $z \in \mathbb{C}$ eine fixierte Zahl und sei

$$A = \begin{pmatrix} a & i \\ i & 1 \end{pmatrix}.$$

Wir wenden auf die Matrix A den Gauß-Algorithmus an. Da wir nicht wissen, ob $a \neq 0$ ist, beginnen wir damit, die Zeilen zu vertauschen.

$$\begin{pmatrix} a & i \\ i & 1 \end{pmatrix} \xrightarrow{Z_1 \leftrightarrow Z_2} \begin{pmatrix} i & 1 \\ a & i \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow -iZ_1} \begin{pmatrix} 1 & -i \\ a & i \end{pmatrix} \xrightarrow{Z_2 \rightsquigarrow Z_2 - aZ_1} \begin{pmatrix} 1 & -i \\ 0 & (a+1)i \end{pmatrix}.$$

An dieser Stelle müssen wir eine Fallunterscheidung machen. Ist $a = -1$, so hat A die reduzierte Zeilenstufenform

$$\begin{pmatrix} 1 & -i \\ 0 & 0 \end{pmatrix}.$$

Ist $a \neq -1$, so können wir die zweite Zeile durch $(a+1)i$ teilen und das i -fache der neuen zweiten Zeile zur ersten Zeile addieren. In diesem Fall hat A als reduzierte Zeilenstufenform die Einheitsmatrix E_2 .

◇

Wir sehen im ersten Beispiel, dass eine Matrix in aller Regel keine eindeutig bestimmte Zeilenstufenform hat (die Matrix, die wir durch die ersten Umformungen erhalten hatten, hat ja Zeilenstufenform, genau wie die Matrix in reduzierter Zeilenstufenform, die beiden Matrizen unterscheiden sich aber). Der folgende Satz sagt aber, dass die »Form«, d.h. die Anzahl und Position der Stufen bzw. der führenden Einsen für alle Matrizen in Zeilenstufenform, die man aus einer Matrix A durch elementare Zeilenumformungen erhalten kann, die gleiche ist. Und es gibt sogar nur eine einzige Matrix in reduzierter Zeilenstufenform, die man aus einer Matrix A so erhalten kann.

SATZ 5.17. Sei $A \in M_{m \times n}(K)$ eine Matrix.

- (1) Alle Matrizen in Zeilenstufenform, die man aus A durch elementare Zeilenumformungen erhalten kann, haben dieselbe Anzahl an führenden Einsen und haben die führenden Einsen in denselben Spalten.
- (2) Es gibt genau eine Matrix in reduzierter Zeilenstufenform, die aus A durch elementare Zeilenumformungen hervorgeht. Wir nennen diese auch die reduzierte Zeilenstufenform von A .

BEWEIS. Wir betrachten A als die Koeffizientenmatrix eines homogenen linearen Gleichungssystems. Da wir bereits wissen (Lemma 5.12), dass elementare Zeilenumformungen die Lösungsmenge nicht ändern, genügt es zu zeigen, dass die Anzahl und Position der führenden Einsen in einer Matrix in Zeilenstufenform durch die Lösungsmenge des zugehörigen homogenen Gleichungssystems eindeutig bestimmt ist. Und entsprechend für Teil (2), dass eine Matrix in reduzierter Zeilenstufenform durch die Lösungsmenge des zugehörigen homogenen linearen Gleichungssystems eindeutig bestimmt ist.

zu (1). Sei $\mathbb{L}$ die Lösungsmenge des durch A gegebenen homogenen linearen Gleichungssystems. Sei A' eine Matrix in Zeilenstufenform, die aus A durch elementare Zeilenumformungen entsteht.

Wir zeigen durch Induktion nach der Anzahl der Spalten von A , dass wir die Anzahl und Positionen der führenden Einsen in A' an $\mathbb{L}$ ablesen können. Wenn A nur eine Spalte hat, dann ist A' die Nullmatrix (wenn A selbst die Nullmatrix ist, also genau dann, wenn $\mathbb{L} = K$) oder die Matrix $(1, 0, \dots, 0)^t$ (in allen anderen Fällen, also genau dann, wenn $\mathbb{L} = \{0\}$). In diesem Fall ist die Zeilenstufenform eindeutig bestimmt.

Nun sei die Anzahl n der Spalten > 1 . Wir können die Induktionsvoraussetzung anwenden: Für jede Matrix in Zeilenstufenform sind auch die ersten $n-1$ Spalten in Zeilenstufenform, und elementare Zeilenumformungen verändern jede Spalte für sich, unabhängig von den anderen Spalten. Wir erhalten die Lösungsmenge des durch die ersten $n-1$ Spalten gegebenen homogenen Gleichungssystems als

$$\{(x_1, \dots, x_{n-1})^t; (x_1, \dots, x_{n-1}, 0)^t \in \mathbb{L}\} \subseteq K^{n-1},$$

diese ist also durch $\mathbb{L}$ eindeutig bestimmt. Es folgt per Induktion, dass die Anzahl und Lage der führenden Einsen in den ersten $n-1$ Spalten durch $\mathbb{L}$ eindeutig bestimmt ist.

Schließlich enthält die letzte Spalte genau dann eine führende 1, wenn der n -te Eintrag in allen Elementen der Lösungsmenge $= 0$ ist.

zu (2). Den Beweis von Teil (2) erläutern wir hier nur an einem Beispiel, das das Prinzip erklärt. Damit könnten Sie den Beweis auch an dieser Stelle schon in der allgemeinen Situation ausarbeiten. Wir verschieben das im Skript auf etwas später, weil es mit dem Produkt von Matrizen einfacher (und durchsichtiger) möglich sein wird, siehe Satz 5.59. Betrachten wir als Beispiel die Matrix

$$\begin{pmatrix} 1 & 0 & a & b & c \\ 0 & 1 & d & e & f \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

in reduzierter Zeilenstufenform. Wir wissen aus Teil (1) bereits, dass die Anzahl und Lage der Stufen bzw. der führenden Einsen durch die Lösungsmenge $\mathbb{L}$ des homogenen Gleichungssystems dazu eindeutig bestimmt sind. Es geht nun noch darum, die Koeffizienten $a, \dots, f$ durch $\mathbb{L}$ auszudrücken. Dafür können wir sagen, dass das einzige Element von $\mathbb{L}$ von der Form $(*, *, 1, 0, 0)^t$ (wobei die Sternchen wie üblich beliebige Elemente aus dem Grundkörper bezeichnen) der Vektor $(-a, -d, 1, 0, 0)^t$ ist. Genauso ist $(-b, -e, 0, 1, 0)^t$ der einzige Vektor in $\mathbb{L}$, dessen letzte drei Einträge $0, 1, 0$ sind, und $(-c, -f, 0, 0, 1)^t$ der einzige Vektor in $\mathbb{L}$, dessen letzte drei Einträge $0, 0, 1$ sind. Wir können also alle Einträge der Matrix in $\mathbb{L}$ »wiederfinden«. $\square$

SATZ 5.18 (Lösungsmenge eines LGS in (reduzierter) Zeilenstufenform). Sei $(A \mid b)$ die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems mit m Gleichungen und n Unbestimmten. Wir setzen voraus, dass die Matrix A Zeilenstufenform hat.

- (1) Wenn es ein i gibt, so dass die i -te Zeile von A eine Nullzeile, aber die i -te Zeile von $(A \mid b)$ keine Nullzeile ist, dann ist die Lösungsmenge leer. Hat sogar $(A \mid b)$ Zeilenstufenform, dann können wir die Bedingung formulieren als: die letzte Spalte von $(A \mid b)$ enthält eine führende Eins.
- (2) Nun gelte, dass in den Matrizen A und $(A \mid b)$ genau dieselben Zeilen Nullzeilen sind. Dann ist die Lösungsmenge nicht leer. Wir benennen die Spalten von A , die eine führende Eins enthalten, mit $1 \leq j_1 < \dots < j_r \leq n$. Dann besteht die Lösungsmenge aus den folgenden Elementen von K^n : Jede Wahl von Elementen $x_j \in K, j \notin \{j_1, \dots, j_r\}$ lässt sich durch die Vorschrift

$$x_{j_i} = b_i - \sum_{j > j_i} a_{ij} x_j, \quad i = 1, \dots, r$$

zu einem Lösungsvektor ergänzen, und auf diese Art und Weise erhält man alle Lösungsvektoren.

(Wenn A sogar reduzierte Zeilenstufenform hat, dann treten auf den rechten Seiten dieser Gleichungen nur x_j mit $j \notin \{j_1, \dots, j_r\}$ auf (die anderen a_{ij} sind Null). Sonst muss man die x_{j_i} beginnend mit x_{j_r} bestimmen und sich dann Schritt für Schritt zu x_{j_1} vorarbeiten.)

BEWEIS. zu (1). Wenn die angegebene Bedingung erfüllt ist, dann hat die i -te Gleichung die Form $0 = b_i$ mit $b_i \neq 0$ und ist daher nicht erfüllbar.

zu (2). Bei der oben angegebenen Darstellung für x_{j_i} ($i = 1, \dots, r$) handelt es sich einfach um eine äquivalente Formulierung der i -ten Gleichung des Systems, $i = 1, \dots, r$. Die Zeilen nach der r -ten Zeile sind Nullzeilen, die als Gleichung einfach $0 = 0$ aussagen – das ist immer erfüllt.

Daher beschreibt der Satz genau die Bedingungen, die ein Lösungsvektor erfüllen muss. $\square$

Schauen Sie sich auch die folgenden Beispiele an, die das etwas konkreter machen.

BEISPIEL 5.19. (1) Sei $K = \mathbb{Q}$ und

$$A = \begin{pmatrix} 0 & 1 & 2 \\ 3 & 4 & 5 \\ 6 & 7 & 8 \end{pmatrix}.$$

Wir haben in Beispiel 5.16 (1) gesehen, dass A die Zeilenstufenform

$$\begin{pmatrix} 1 & \frac{4}{3} & \frac{5}{3} \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{pmatrix}$$

hat. Daran können wir die Lösungsmenge des durch A gegebenen homogenen Gleichungssystems ablesen als

$$\mathbb{L} = \left\{ \begin{pmatrix} z \\ -2z \\ z \end{pmatrix} ; z \in \mathbb{Q} \right\},$$

denn $-\frac{4}{3} \cdot (-2z) - \frac{5}{3}z = z$. Statt diese Rechnung an dieser Stelle zu machen, kann man auch die Matrix auf die reduzierte Zeilenstufenform

$$\begin{pmatrix} 1 & 0 & -1 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{pmatrix}.$$

bringen und dann die Lösungsmenge noch direkter ablesen. Das Ergebnis ist natürlich dasselbe.

- (2) Sei K irgendein Körper und seien $a, b, c, d, e \in K$. Wir betrachten die erweiterte Koeffizientenmatrix

$$\left(\begin{array}{cccc|c} 1 & a & 0 & b & c \\ 0 & 0 & 1 & d & e \\ 0 & 0 & 0 & 0 & 0 \end{array} \right)$$

in reduzierter Zeilenstufenform. Die Lösungsmenge des zugehörigen linearen Gleichungssystems ist

$$\mathbb{L} = \left\{ \begin{pmatrix} c - ax - bx' \\ x \\ e - dx' \\ x' \end{pmatrix} ; x, x' \in K \right\}$$

Eine etwas andere Art, die Elemente der Lösungsmenge zu notieren, ist die folgende:

$$\mathbb{L} = \left\{ \begin{pmatrix} c \\ 0 \\ e \\ 0 \end{pmatrix} + x \begin{pmatrix} -a \\ 1 \\ 0 \\ 0 \end{pmatrix} + x' \begin{pmatrix} -b \\ 0 \\ -d \\ 1 \end{pmatrix} ; x, x' \in K \right\}$$

Damit sehen wir auch direkt die Darstellung von $\mathbb{L}$ in der Form $t + \mathbb{L}_0$, wobei $\mathbb{L}_0$ die Lösungsmenge des zugehörigen homogenen Systems ist:

$$\mathbb{L} = \begin{pmatrix} c \\ 0 \\ e \\ 0 \end{pmatrix} + \left\{ x \begin{pmatrix} -a \\ 1 \\ 0 \\ 0 \end{pmatrix} + x' \begin{pmatrix} -b \\ 0 \\ -d \\ 1 \end{pmatrix} ; x, x' \in K \right\}$$

◇

BEMERKUNG 5.20. Sei $(A \mid b)$ die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems. Wenn man in A Spalten vertauscht, entspricht das in Termen des Gleichungssystems einfach einem Umsortieren der Unbestimmten. Da alle Zeilenumformungen separat auf jede einzelne Spalte wirken, ist es zulässig, zunächst eine Spaltenvertauschung von A vorzunehmen, dann den Gauß-Algorithmus anzuwenden, und dann in der so erhaltenen Matrix die umgekehrte Spaltenumformung durchzuführen, und am Ergebnis die Lösungsmenge abzulesen. Allerdings muss man dabei besonders gut aufpassen, dass keine Flüchtigkeitsfehler passieren (zum Beispiel vergisst man leicht, am Ende wieder zurückzutauschen). Die letzte

Spalte von $(A \mid b)$ spielt selbstverständlich eine Sonderrolle und darf nicht mitvertauscht werden.

Bei den Gleichungssystemen, die Sie »per Hand« lösen müssen, sollte es nicht erforderlich sein, Spalten zu vertauschen. In der Praxis, bei der Berechnung der Lösungsmengen von Gleichungssystemen mit »vielen« Gleichungen und Unbestimmten mit einem Computerprogramm, wo auch Rundungsfehler ein Problem darstellen können, kann das aber sinnvoll sein (so wie es dann auch wichtiger ist, die Zeilenumformungen geschickt zu wählen). $\diamond$

Wir fassen unsere Ergebnisse noch einmal zusammen:

SATZ 5.21. Sei $(A \mid b)$ die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems. Sei $(A' \mid b')$ eine Matrix in Zeilenstufenform, die aus $(A \mid b)$ durch elementare Zeilenumformungen entsteht.

(1) Es sind äquivalent:

- (i) Das durch $(A \mid b)$ gegebene lineare Gleichungssystem besitzt eine Lösung.
- (ii) Die Matrizen A' und $(A' \mid b')$ haben gleich viele Stufen.

(2) Es sind äquivalent:

- (i) Das durch $(A \mid b)$ gegebene lineare Gleichungssystem besitzt eine eindeutige Lösung.
- (ii) Die Matrix A' hat in jeder Spalte eine Stufe, und für jede Nullzeile in A' ist der entsprechende Eintrag in b' ebenfalls gleich 0.

(3) Sei nun in der obigen Situation A eine quadratische Matrix. Dann sind äquivalent:

- (i) Das durch $(A \mid b)$ gegebene lineare Gleichungssystem besitzt eine eindeutige Lösung.
- (ii) Für die Matrix $A' = (a'_{ij})_{i,j}$ gilt $a'_{ii} = 1$ für alle i , d.h. die Diagonaleinträge sind alle = 1, und $a'_{ij} = 0$ für alle $i > j$, d.h. die »Einträge unterhalb der Diagonale« sind alle gleich Null.

Inbesondere ist diese Eigenschaft unabhängig von b .

Hat in dieser Situation A' sogar reduzierte Zeilenstufenform, dann beschreibt $(A \mid b)$ ein eindeutig lösbares lineares Gleichungssystem genau dann, wenn A' die Einheitsmatrix ist.

BEWEIS. Vorbemerkung: Wenn $(A' \mid b')$ durch elementare Zeilenumformungen aus $(A \mid b)$ entsteht und (reduzierte) Zeilenstufenform hat, so entsteht auch A' durch elementare Zeilenumformungen (nämlich durch genau dieselben Umformungen) aus A und hat (reduzierte) Zeilenstufenform.

zu (1). Das ist eine Umformulierung von Satz 5.18 (1).

zu (2). Dass das Gleichungssystem eindeutig lösbar ist, ist gleichbedeutend damit, dass es eine Lösung gibt, aber dass wir keine Elemente in K wie in Satz 5.18 (2) wählen können, denn sonst gäbe es mindestens zwei Lösungen. Das heißt genau, dass in jeder Spalte eine Stufe ist.

zu (3). Dies folgt aus Teil (2) mit der Bemerkung, dass eine quadratische Matrix in Zeilenstufenform genau dann in jeder Spalte eine Stufe hat, wenn es sich um eine Matrix mit der beschriebenen Form handelt. (Wir sprechen von einer »oberen Dreiecksmatrix« mit Einsen auf der Diagonale, siehe Abschnitt 5.3.2.)

Die Bemerkung am Ende ergibt sich daraus, dass die einzige quadratische Matrix in reduzierter Zeilenstufenform, die in jeder Spalte eine Stufe hat, die Einheitsmatrix ist. $\square$

Teil (2) zeigt, dass ein lineares Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$, $A \in M_{m \times n}(K)$ mit $m < n$ (mehr Unbestimmte als Gleichungen, man spricht auch von einem *unterbestimmten* linearen Gleichungssystem) nicht eindeutig lösbar sein kann. Dass in Teil (3) des Satzes die reduzierte Zeilenstufenform von A die Einheitsmatrix ist, ist eine Bedingung, die nicht von b abhängt. Wir können noch ein bisschen mehr zeigen:

KOROLLAR 5.22. Sei $A \in M_n(K)$ eine quadratische Matrix. Dann sind äquivalent:

- (i) Für alle $b \in K^n$ ist das durch $(A \mid b)$ gegebene Gleichungssystem eindeutig lösbar.
- (ii) Für alle $b \in K^n$ ist das durch $(A \mid b)$ gegebene Gleichungssystem lösbar.
- (iii) Es existiert $b \in K^n$, so dass das Gleichungssystem, das durch $(A \mid b)$ gegeben ist, eindeutig lösbar ist.
- (iv) Die reduzierte Zeilenstufenform von A ist E_n .

BEWEIS. Es ist offensichtlich, dass (i) $\Rightarrow$ (ii) gilt, und aus dem vorherigen Satz folgt (iii) $\Rightarrow$ (iv) $\Rightarrow$ (i). Es genügt daher, nun noch die Implikation (ii) $\Rightarrow$ (iii) zu zeigen. Wir zeigen, dass aus (ii) die Aussage (iii) für $b = 0$ folgt. Wenn das durch A gegebene homogene lineare Gleichungssystem mehrere Lösungen hat, dann hat die reduzierte Zeilenstufenform A' von A eine Nullzeile. Es ist dann klar, dass $b' \in K^n$ existiert, so dass das Gleichungssystem zu $(A' \mid b')$ nicht lösbar ist – wir müssen nur einen Eintrag von b' in einer Zeile, die in A' eine Nullzeile ist, auf einen Wert $\neq 0$ festsetzen.

Da A' aus A durch elementare Zeilenumformungen entsteht, können wir auch A aus A' durch elementare Zeilenumformungen erhalten. Wenden wir diese Zeilenumformungen auf die erweiterte Matrix $(A' \mid b')$ an, so erhalten wir eine Matrix $(A \mid b)$ mit $b \in K^n$, deren Gleichungssystem dieselbe Lösungsmenge hat wie das System $(A' \mid b')$, also die leere Menge. Das zeigt, dass (ii) falsch ist, wenn (iii) nicht gilt, oder mit anderen Worten (ii) $\Rightarrow$ (iii). $\square$

5.2.3. Teilräume von K^n . Um die Struktur der Lösungsmenge eines linearen Gleichungssystems noch besser beschreiben zu können (und um später die Frage zu beantworten, welche Teilmengen von K^n als Lösungsmenge eines linearen Gleichungssystems auftreten können), machen wir die folgende Definition:

DEFINITION 5.23. Eine Teilmenge $U \subseteq K^n$ heißt *Teilraum* von K^n , wenn die folgenden Bedingungen erfüllt sind:

- (a) $0 \in U$,
- (b) für alle $u, v \in U$ gilt $u + v \in U$, und
- (c) für alle $a \in K, u \in U$ gilt $au \in U$.

–

BEISPIEL 5.24. (1) Die Teilmengen $\{0\}$ und K^n sind Teilräume von K^n . Den Teilraum $\{0\}$ nennt man auch den *Nullraum*; manchmal bezeichnet man den Nullraum auch mit 0 (ohne Mengenklammern).

(2) Die einzigen Teilräume von $K^1 = K$ sind 0 und K . (Warum?)

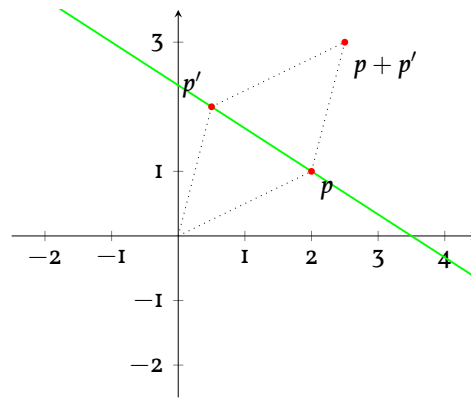
(3) Ist $x \in K^n$, so ist die Menge $\{ax; a \in K\}$ ein Teilraum von K^n . (Warum?)

(4) Jede Gerade in $\mathbb{R}^2$, die den Ursprung enthält, ist ein Teilraum (von der Form von Teil (3) dieses Beispiels). Wir werden sehen: Die Teilräume von $\mathbb{R}^2$ sind $\{0\}$, $\mathbb{R}^2$ und die Geraden in $\mathbb{R}^2$, die den Ursprung $(0, 0)^t$ enthalten. Können Sie zeigen, dass es keine anderen gibt?

(5) Die leere Menge $\emptyset$ ist kein Teilraum. Die Teilmenge $\{(x, y); y = x^2\}$ ist kein Teilraum von $\mathbb{R}^2$, denn sie enthält $(1, 1)^t$ und $(-1, 1)^t$, aber nicht die Summe $(1, 1)^t + (-1, 1)^t = (0, 2)^t$.

$\diamond$

BEISPIEL 5.25. Die grüne Gerade im Bild rechts erfüllt keine der Bedingungen (a), (b), (c): Zum Beispiel ist (b) verletzt, weil sie p und p' enthält, aber nicht die Summe $p + p'$. Finden Sie ein Beispiel, das zeigt, dass (c) ebenfalls nicht erfüllt ist. $\diamond$



Die (für den Moment) wichtigsten Beispiele für Teilräume liefert uns das folgende Lemma.

LEMMA 5.26. Sei $U \subseteq K^n$ die Lösungsmenge eines homogenen linearen Gleichungssystems. Dann ist U ein Teilraum von K .

BEWEIS. Zunächst liegt sicher der Nullvektor in der Lösungsmenge jedes homogenen linearen Gleichungssystems: $0 \in U$. Sind $v = (v_j)_j, w = (w_j)_j \in U$, erfüllen also $(v_1, \dots, v_n)$ und $(w_1, \dots, w_n)$ alle Gleichungen des gegebenen homogenen linearen Gleichungssystems, so gilt das auch für die Summe $(v_1 + w_1, \dots, v_n + w_n)$. In der Tat, betrachten wir eine Gleichung der Form

$$a_1 X_1 + \dots + a_n X_n = 0,$$

so sehen wir

$$\begin{aligned} a_1(v_1 + w_1) + \dots + a_n(v_n + w_n) &= a_1 v_1 + a_1 w_1 + \dots + a_n v_n + a_n w_n \\ &= a_1 v_1 + \dots + a_n v_n + a_1 w_1 + \dots + a_n w_n \\ &= 0 + 0 = 0. \end{aligned}$$

Ähnlich, aber noch einfacher, ist es zu zeigen, dass für jede Lösung v eines homogenen linearen Gleichungssystems und jedes Element $a \in K$ auch av eine Lösung ist. Damit sind alle Eigenschaften gezeigt, die wir überprüfen müssen, um zu beweisen, dass U ein Teilraum von K^n ist. $\square$

FRAGE 5.27. Wir wollen diesen Abschnitt damit beenden, zwei Fragen zu formulieren, die sich an dieser Stelle stellen, die wir aber erst mit etwas mehr Theorie werden befriedigend beantworten können.

- (1) Ist jeder Teilraum von K^n die Lösungsmenge eines homogenen linearen Gleichungssystems? (Das ist in der Tat der Fall, siehe Satz 7.39.)
- (2) In Korollar 5.22 haben wir gesehen, dass für eine quadratische Matrix A die Eigenschaft, dass das Gleichungssystem $(A|0)$ eindeutig lösbar ist, dazu äquivalent ist, dass alle Gleichungssysteme $(A|b)$ lösbar sind.

Können wir diese Beziehung zwischen der *Eindeutigkeit* der Lösung und der *Existenz* von Lösungen verallgemeinern? Kann man zum Beispiel »quantifizieren«, wie viele Lösungen ein lineares Gleichungssystem $(A | 0)$ hat, auch wenn es nicht eindeutig lösbar ist, und das in Beziehung dazu setzen, für »wie viele« b das Gleichungssystem $(A | b)$ lösbar ist?

Weil wir (wenn K unendlich viele Elemente hat) darüber sprechen, dass unendlich viele Lösungen existieren, ist nicht klar, was man unter »quantifizieren« verstehen soll. Wir werden diese Frage in einer sehr schönen (und nützlichen) Art und Weise beantworten; siehe Abschnitt 7.4.

5.3. Das Matrizenprodukt

5.3.1. Rechnen mit Matrizen. In diesem Abschnitt wollen wir Matrizen noch einmal losgelöst vom Begriff des linearen Gleichungssystems studieren und insbesondere Rechenoperationen auf der Menge der Matrizen definieren, die sich in vielen Situationen als sehr nützlich erweisen (auch für die Behandlung von linearen Gleichungssystemen). Wir wiederholen hier zur Erinnerung noch einmal die Definition 5.6.

DEFINITION 5.28. Seien m, n natürliche Zahlen. Unter einer *Matrix* der Größe $m \times n$ mit Einträgen in K (man spricht auch von einer $(m \times n)$ -Matrix über K) verstehen wir eine Familie $(a_{ij})_{i=1, \dots, m, j=1, \dots, n}$ von Elementen von K . Die Elemente a_{ij} heißen die *Einträge* oder *Koeffizienten* der Matrix. Eine Matrix stellen wir uns immer in der Form

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

vor, d.h. wir schreiben die Koeffizienten in einem rechteckigen Schema auf, in dem der erste Index die Zeile und der zweite Index die Spalte angibt.

Die Menge aller $(m \times n)$ -Matrizen über K bezeichnen wir mit $M_{m \times n}(K)$. Ist $m = n$, so schreiben wir manchmal $M_n(K)$ statt $M_{m \times n}(K)$. $\dashv$

Die *Nullmatrix* (der Größe $m \times n$) ist die Matrix $0 \in M_{m \times n}(K)$, deren Einträge alle gleich Null sind.

Als Spezialfall können wir die Elemente des K^n , die wir wie gehabt als Spaltenvektoren betrachten wollen, als Matrizen mit n Zeilen und einer Spalte auffassen, das heißt $K^n = M_{n \times 1}(K)$.

Auch mit Matrizen kann man rechnen. Um zum Beispiel zwei Matrizen (der selben Größe) zu addieren, addieren wir einfach auf jeder Position separat die Einträge. Wie man sinnvolles Produkt von zwei Matrizen definiert, ist nicht ganz so offensichtlich. Wir fassen die wichtigen Rechenoperationen in der folgenden Definition zusammen:

DEFINITION 5.29. Seien K ein Körper und l, m, n natürliche Zahlen. Die Einträge der Matrizen A und B bezeichnen wir mit a_{ij} beziehungsweise b_{ij} .

- (1) (Summe von Matrizen) Für Matrizen $A, B \in M_{m \times n}(K)$ derselben Größe definieren wir

$$A + B = (a_{ij} + b_{ij})_{i=1, \dots, m, j=1, \dots, n} \in M_{m \times n}(K).$$

- (2) (Produkt eines Skalars und einer Matrix) Für eine Matrix $A \in M_{m \times n}(K)$ und ein Element $a \in K$ definieren wir

$$aA = Aa = (aa_{ij})_{ij} \in M_{m \times n}(K).$$

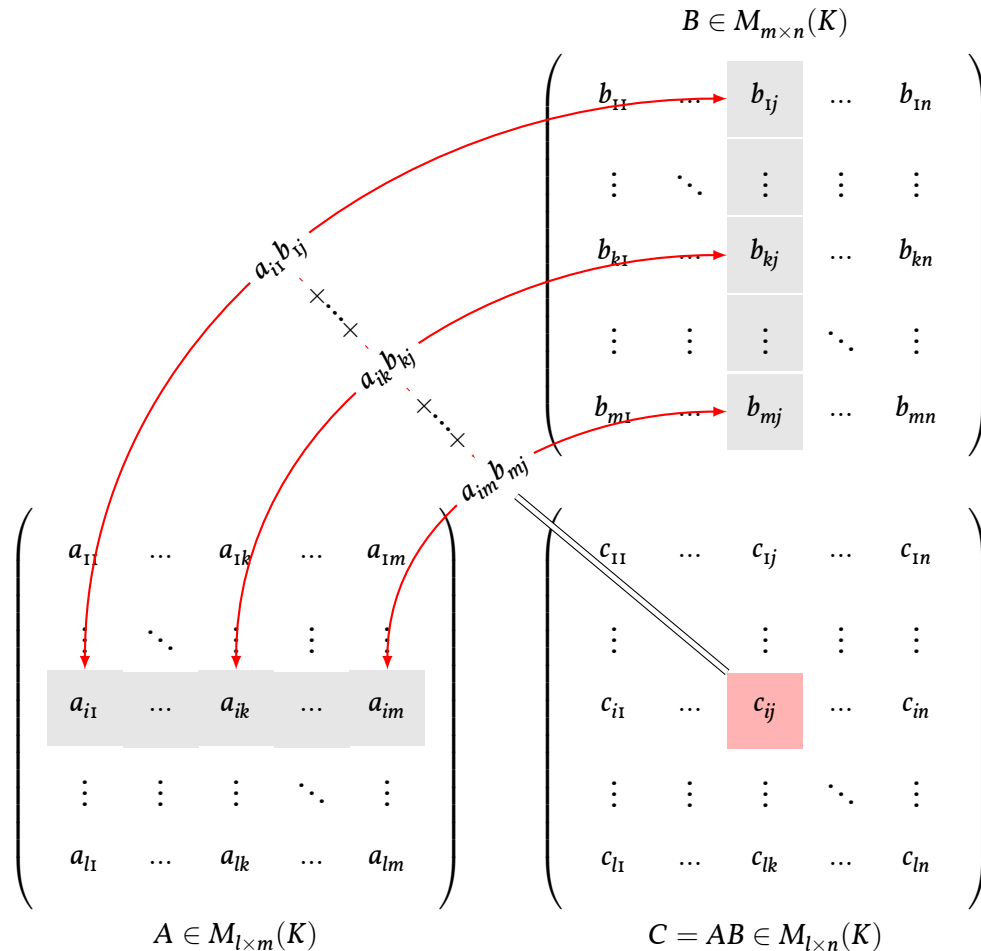
- (3) (Produkt zweier Matrizen) Sind $A \in M_{l \times m}(K)$ und $B \in M_{m \times n}(K)$ Matrizen, so definieren wir das *Produkt*

$$A \cdot B = \left(\sum_{j=1}^m a_{ij} b_{jk} \right)_{i=1, \dots, l, k=1, \dots, n} \in M_{l \times n}(K).$$

Den Multiplikationspunkt lässt man dabei oft aus: $AB = A \cdot B$.

$\dashv$

BEISPIEL 5.30. Es ist für die Berechnung des Matrizenprodukts manchmal hilfreich, die zu multiplizierenden Matrizen A und B so anzuordnen, wie in der Abbildung gezeigt, d.h. dass B nach oben verschoben wird, und das Ergebnis dann rechts von A und unter B zu stehen kommt. Um die Zelle c_{ij} von $C = AB$ auszurechnen, benutzt man die i -te Zeile von A und die j -te Spalte von B . Diese haben beide m Einträge. Man multipliziert die ersten Einträge, die zweiten Einträge, ..., und summiert alle diese Produkte auf.



(Illustration angepasst von texample.net¹, Autor: Alain Matthes.)

◇

Angesichts der Identifikation $K^n = M_{n \times 1}(K)$ können wir auch Elemente von K^n mit Matrizen multiplizieren (wenn die Größen zusammenpassen). Für quadratische Matrizen $A \in M_n(K)$ können wir dann auch die Potenzen A^n von A für $n \in \mathbb{N}$ definieren: $A^0 = E_n$, $A^1 = A$, $A^2 = AA$, und so weiter: $A^n = A^{n-1}A$ für $n > 0$.

SATZ 5.31. (Eigenschaften der Rechenoperationen mit Matrizen) Seien $l, m, n \in \mathbb{N}$.

(I) Die Addition von Matrizen definiert eine Verknüpfung

$$+ : M_{m \times n}(K) \times M_{m \times n}(K) \rightarrow M_{m \times n}(K),$$

die die folgenden Eigenschaften hat:

- (a) Assoziativität: $(A + B) + C = A + (B + C)$,
- (b) Kommutativität: $A + B = B + A$,
- (c) die Nullmatrix 0 ist neutrales Element: $0 + A = A = A + 0$,

¹<https://texample.net/tikz/examples/matrix-multiplication/>

- (d) *das additive Inverse der Matrix $(a_{ij})_{i,j}$ ist die Matrix $-A := (-a_{ij})_{i,j} (= (-I) \cdot A)$.*
- (2) *Die Multiplikation von Matrizen definiert eine assoziative Verknüpfung $M_{l \times m}(K) \times M_{m \times n}(K) \rightarrow M_{l \times n}(K)$.*
- (3) *Die Addition und das Produkt von Matrizen verhalten sich distributiv, d.h. es gilt*

$$(A + B)C = AC + BC, \quad A(B + C) = AB + AC,$$

wenn immer man die entsprechenden Summen und Produkte bilden kann.

BEWEIS. Die Eigenschaften der Addition sind klar, da die Addition eintragweise erfolgt. Seien nun A, B, C Matrizen mit Einträgen a_{ij}, b_{ij}, c_{ij} , so dass die Produkte AB , und BC existieren. Die Assoziativität des Produkts ergibt sich aus der Gleichungskette

$$\sum_k \left(\sum_j a_{ij} b_{jk} \right) c_{kl} = \sum_k \sum_j a_{ij} b_{jk} c_{kl} = \sum_j \sum_k a_{ij} b_{jk} c_{kl} = \sum_j a_{ij} \left(\sum_k b_{jk} c_{kl} \right),$$

in der links der Eintrag von $(AB)C$ in der i -ten Zeile und l -ten Spalte, und rechts der entsprechende Eintrag von $A(BC)$ steht. (Die Intervalle, über die die Indizes in den Summen laufen, bestimmen sich aus der Größe von B und sind hier ausgelassen.) Die äußeren Gleichheiten benutzen dabei das Distributivgesetz in K , die mittlere das Kommutativgesetz: Wir summieren die Ausdrücke $a_{ij} b_{jk} c_{kl}$ für alle j und k auf, und es spielt keine Rolle, wie wir sie anordnen.

Die Rechnung für das Distributivgesetz ist einfach und wir lassen sie hier aus. $\square$

Die Einheitsmatrix (der richtigen Größe) ist ein neutrales Element bezüglich der Multiplikation: $E_m A = A$ und $A E_n = A$ für alle $A \in M_{m \times n}(K)$.

BEISPIEL 5.32. Wie das folgende Beispiel zeigt, ist das Produkt von Matrizen *nicht kommutativ*:

$$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} \neq \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}.$$

$\diamond$

Die Multiplikation eines Skalars mit einer Matrix ist auch assoziativ und kommutativ: $a(AB) = (aA)B = A(aB)$, wie man unmittelbar nachprüft.

ERGÄNZUNG 5.33 (Der Strassen-Algorithmus). Interessanterweise lässt sich ein Produkt von zwei $(n \times n)$ -Matrizen mit weniger als den a priori erforderlichen n^3 Multiplikationen berechnen. Die bekannteste Methode ist der **Strassen-Algorithmus**². Für das Produkt von zwei (2×2) -Matrizen kommt er mit 7 Multiplikationen (statt 8 für die direkte Rechnung) aus; allerdings werden mehr Additionen als in der direkten Rechnung durchgeführt. Weil Computer für eine Multiplikation mehr Zeit benötigen als für eine Addition, wird der Algorithmus in der Praxis verwendet, aber er bringt einen signifikanten Vorteil erst für Matrizen mit mehreren Hundert Zeilen und Spalten. $\square$ Ergänzung 5.33

BEISPIEL 5.34. Eine weitere Eigenschaft des Matrizenprodukts, die einen Unterschied zur Multiplikation in einem Körper bedeutet, ist, dass das Produkt von zwei Matrizen $\neq 0$ die Nullmatrix sein kann:

$$\begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}^2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

$\diamond$

²https://en.wikipedia.org/wiki/Strassen_algorithm

Auch wenn die Berechnung des Matrizenprodukts nicht schwierig ist, gibt es darüber weitreichende strukturelle Aussagen, für deren Beweis wir erst im Laufe der Vorlesung die notwendigen Methoden entwickeln werden. Ein Beispiel für so eine Aussage ist der folgende

SATZ 5.35. *Seien K ein Körper, $n \geq 1$, $A \in M_{n \times n}(K)$. Wenn eine natürliche Zahl N existiert mit $A^N = 0$, dann gilt $A^n = 0$.*

Siehe Satz 6.56 für einen Beweis. In der Linearen Algebra 2 werden wir diese Aussage noch weiter verallgemeinern und besser verstehen können (der hier angegebene Satz ist ein Spezialfall des »Satzes von Cayley-Hamilton«).

BEMERKUNG 5.36. Sei $(A | b)$ die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems. Für $x \in K^n = M_{n \times 1}(K)$ ist dann das Matrizenprodukt Ax eine $(m \times 1)$ -Matrix, also ein Element von K^m , und zwar mit dem Eintrag

$$\sum_{j=1}^n a_{ij}x_j$$

in der i -ten Zeile; das ist genau die linke Seite der i -ten Gleichung des betrachteten Gleichungssystems. Das bedeutet: $x \in K^n$ ist ein Lösungsvektor des linearen Gleichungssystems genau dann, wenn $Ax = b$. Die Lösungsmenge des linearen Gleichungssystems $(A | b)$ ist also

$$\{x \in K^n; Ax = b\}.$$

◇

BEMERKUNG 5.37. Wir können elementare Zeilenumformungen als Multiplikation mit einer geeigneten Matrix von links ausdrücken. Sei A eine Matrix.

Typ (I). Sei A' die Matrix, die aus A durch Addition des a -fachen der j -ten Zeile zur i -ten Zeile entsteht. Dann gilt

$$A' = E_{ij}(a)A,$$

wobei wir für $i \neq j$ und $a \in K$ definieren:

$$E_{ij}(a) := \begin{pmatrix} \text{I} & & & \\ & \ddots & & \\ & & \ddots & a \\ & & & \ddots \\ & & & & \text{I} \end{pmatrix},$$

wobei der Eintrag a in der i -ten Zeile und j -ten Spalte ist, auf der Diagonale Einsen stehen und alle anderen Einträge Null sind. (Ist $i > j$ so befindet sich der Eintrag a , anders als in der Abbildung, unterhalb der Diagonalen.)

Typ (II). Sei A' die Matrix, die aus A durch Vertauschen der i -ten und j -ten Zeile entsteht. Dann gilt

$$A' = P_{ij}A,$$

wobei wir für $i \neq j$ definieren:

$$P_{ij} := \begin{pmatrix} \text{I} & & & \\ & \ddots & & \\ & & 0 & \text{I} \\ & & & \ddots \\ & & \text{I} & 0 \\ & & & & \ddots \\ & & & & & \text{I} \end{pmatrix}.$$

Hier befinden sich die beiden Einsen außerhalb der Diagonale in Zeile i und Spalte j , und in Zeile j und Zeile i . Auf der Diagonale stehen in den Zeilen i und j Nullen, sonst überall 1.

Typ (III). Sei A' die Matrix, die aus A durch Multiplikation der i -ten Zeile mit $a \in K^\times$ entsteht. Dann gilt

$$A' = \begin{pmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & & a & & \\ & & & & 1 & \\ & & & & & \ddots \\ & & & & & & 1 \end{pmatrix} A,$$

wobei sich der Eintrag a auf der Diagonalen an Position i befindet, alle anderen Einträge aus der Diagonalen sind 1, alle Einträge außerhalb der Diagonalen sind 0. Wir werden für diese »Diagonalmatrix« später die Notation $\text{diag}(1, \dots, 1, a, 1, \dots, 1)$ benutzen, siehe Abschnitt 5.3.2.

◇

Wir können nun auch die Schreibweise $-^t$, die wir für Zeilenvektoren schon eingeführt hatten, auf beliebige Matrizen ausdehnen:

DEFINITION 5.38. Sei $A = (a_{ij})_{i,j} \in M_{m \times n}(K)$ eine Matrix. Die *transponierte Matrix* A^t von A ist die Matrix

$$A^t = (a_{ji})_{i=1, \dots, n, j=1, \dots, m} \in M_{n \times m}.$$

—

Beim Übergang zur transponierten Matrix spiegeln wir also sozusagen an der Diagonalen. Die Anzahl der Zeilen in A^t ist die Anzahl der Spalten von A , und umgekehrt.

BEISPIEL 5.39. (1)

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}^t = \begin{pmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{pmatrix}.$$

(2)

$$\begin{pmatrix} 1 & 2 & 3 & 4 \end{pmatrix}^t = \begin{pmatrix} 1 \\ 2 \\ 3 \\ 4 \end{pmatrix}, \quad \begin{pmatrix} 1 \\ 2 \\ 3 \\ 4 \end{pmatrix}^t = \begin{pmatrix} 1 & 2 & 3 & 4 \end{pmatrix}$$

◇

LEMMA 5.40. (1) Sei A eine Matrix. Dann gilt $A^{tt} := (A^t)^t = A$.

(2) Seien A, B Matrizen, deren Produkt AB existiert. Dann existiert das Produkt $B^t A^t$ und es gilt

$$(AB)^t = B^t A^t.$$

BEWEIS. Teil (1) ist klar. Zu Teil (2). Die Existenz der Produkte heißt in beiden Fällen, dass B so viele Zeilen wie A Spalten hat. Der Eintrag von $(AB)^t$ in der j -ten Zeile und i -ten Spalte ist der Eintrag von AB in Zeile i und Spalte j , also

$$\sum_k a_{ik} b_{kj},$$

wobei wir $A = (a_{ik})_{i,k}$ und $B = (b_{kj})_{k,j}$ geschrieben haben. Aber das ist nichts anderes als

$$\sum_k b_{kj} a_{ik},$$

was der Eintrag von $B^t A^t$ in Zeile j und Spalte i ist. $\square$

5.3.2. Spezielle Matrizen. Wir führen noch einige Sprechweise über Matrizen mit speziellen Eigenschaften ein.

Sei $A = (a_{ij})_{i,j}$ eine quadratische Matrix der Größe $n \times n$. Wir nennen A eine *obere Dreiecksmatrix*, wenn $a_{ij} = 0$ für alle i, j mit $i > j$. (An den links mit * markierten Einträgen kann irgendein Element von K stehen, die unmarkierten Einträge müssen = 0 sein.)

$$\begin{pmatrix} * & * & \cdots & * \\ & * & \cdots & * \\ & & \ddots & \vdots \\ & & & * \end{pmatrix}$$

Analog kann man von *unteren Dreiecksmatrizen* sprechen (also wenn $a_{ij} = 0$ für alle i, j mit $i < j$).

Eine *Diagonalmatrix* ist eine quadratische Matrix $A = (a_{ij})_{i,j}$ mit $a_{ij} = 0$ für alle $i \neq j$.
 Mit $\text{diag}(c_1, \dots, c_n)$ bezeichnen wir die Diagonalmatrix der Größe n mit den Diagonaleinträgen $c_1, c_2, \dots, c_n$.
 Dann gilt $aA = \text{diag}(a, \dots, a)A = A \text{diag}(a, \dots, a)$ für $a \in K$ und $A \in M_{m \times n}$, wobei in der diag -Ausdruck auf der linken Seite von A dann m Einträge, der auf der rechten Seite n Einträge haben muss.

Die $(n \times n)$ -Matrix

$$E_n := \text{diag}(I, \dots, I) = \begin{pmatrix} I & & & \\ & I & & \\ & & \ddots & \\ & & & I \end{pmatrix}$$

heißt die *Einheitsmatrix* (der Größe n über dem Körper K). (Die Einträge, die gleich Null sind, sind hier weggelassen.)

Wir können Matrizen aus kleineren Matrizen (passender Größen) zusammensetzen; wir sprechen dann von der Notation als *Blockmatrizen*. Sind zum Beispiel $A \in M_{m_1 \times n_1}(K)$, $B \in M_{m_1 \times n_2}(K)$, $C \in M_{m_2 \times n_1}(K)$, $D \in M_{m_2 \times n_2}(K)$, so erhalten wir durch Zusammensetzen die $(m_1 + m_2) \times (n_1 + n_2)$ -Matrix $\begin{pmatrix} A & B \\ C & D \end{pmatrix}$.

Das ist verträglich mit der Multiplikation von Matrizen:

LEMMA 5.4I. Für $M = \begin{pmatrix} A & B \\ C & D \end{pmatrix}$ und $M' = \begin{pmatrix} A' & B' \\ C' & D' \end{pmatrix}$ mit $A \in M_{l_1 \times m_1}(K)$, $B \in M_{l_1 \times m_2}(K)$, $C \in M_{l_2 \times m_1}(K)$, $D \in M_{l_2 \times m_2}(K)$, $A' \in M_{m_1 \times n_1}(K)$, $B' \in M_{m_1 \times n_2}(K)$, $C' \in M_{m_2 \times n_1}(K)$, $D' \in M_{m_2 \times n_2}(K)$ gilt

$$MM' = \begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} A' & B' \\ C' & D' \end{pmatrix} = \begin{pmatrix} AA' + BC' & AB' + BD' \\ CA' + DC' & CB' + DD' \end{pmatrix} \in M_{(l_1+l_2) \times (n_1+n_2)}(K).$$

BEWEIS. Das lässt sich unmittelbar anhand der Definition des Matrixprodukts nachrechnen. $\square$

Weil das Matrizenprodukt (der kleineren Matrizen) nicht kommutativ ist, darf man natürlich in der Situation des Lemmas in den Produkten in den einzelnen Blockeinträgen von MM' die Faktoren nicht vertauschen. Ähnlich funktioniert das mit Blockmatrizen, deren Zeilen und/oder Spalten mehr als zwei Blöcke enthalten.

BEMERKUNG 5.42. Wir betrachten den Körper $\mathbb{C}$ der komplexen Zahlen (siehe Beispiel 4.5). Die Elemente von $\mathbb{C}$ schreiben wir als $a + bi$ mit $a, b \in \mathbb{R}$.

Wir haben eine Bijektion

$$\iota: \mathbb{C} \rightarrow \left\{ \begin{pmatrix} a & -b \\ b & a \end{pmatrix}; a, b \in \mathbb{R} \right\}$$

Dabei entsprechen die Addition und Multiplikation auf der linken Seite genau der Addition und Multiplikation von Matrizen auf der rechten Seite, das heißt für alle $z_1, z_2 \in \mathbb{C}$ gilt:

$$\iota(z_1 + z_2) = \iota(z_1) + \iota(z_2), \quad \iota(z_1 z_2) = \iota(z_1) \iota(z_2).$$

Dies rechnet man unmittelbar anhand der Definitionen von Addition und Multiplikation auf den beiden Seiten nach. Zum Beispiel für die Multiplikation (der schwierigere Fall), mit $z_\lambda = a_\lambda + b_\lambda i, \lambda = 1, 2$:

$$\iota(z_1 z_2) = \begin{pmatrix} a_1 a_2 - b_1 b_2 & -(a_1 b_2 + a_2 b_1) \\ a_1 b_2 + a_2 b_1 & a_1 a_2 - b_1 b_2 \end{pmatrix}$$

und

$$\iota(z_1) \iota(z_2) = \begin{pmatrix} a_1 & -b_1 \\ b_1 & a_1 \end{pmatrix} \begin{pmatrix} a_2 & -b_2 \\ b_2 & a_2 \end{pmatrix} = \begin{pmatrix} a_1 a_2 - b_1 b_2 & -a_1 b_2 - b_1 a_2 \\ b_1 a_2 + a_1 b_2 & -b_1 b_2 + a_1 a_2 \end{pmatrix}.$$

Es spielt dabei keine Rolle, ob wir schon wissen, dass $\mathbb{C}$ mit diesen Verknüpfungen ein Körper ist. Wir sehen auch, dass für zwei Matrizen von der oben betrachteten Form (d.h., die im Bild von ι liegen), auch die Summe und das Produkt diese Form haben.

Wir können die Bijektion ι (und die Gültigkeit der Assoziativgesetze und des Distributivgesetzes in $M_2(\mathbb{R})$) benutzen, um zu beweisen, dass die Assoziativgesetze für $+$ und $\cdot$ und das Distributivgesetz in $\mathbb{C}$ gelten. Dazu bezeichnen wir mit ι^{-1} die Umkehrabbildung.

Nun folgt zum Beispiel das Assoziativgesetz der Multiplikation:

$$\begin{aligned} (z_1 z_2) z_3 &= \iota^{-1}(\iota((z_1 z_2) z_3)) = \iota^{-1}(\iota(z_1 z_2) \iota(z_3)) = \iota^{-1}((\iota(z_1) \iota(z_2)) \iota(z_3)) \\ &= \iota^{-1}(\iota(z_1) (\iota(z_2) \iota(z_3))) = \iota^{-1}(\iota(z_1) \iota(z_2 z_3)) = \iota^{-1}(\iota(z_1 (z_2 z_3))) = z_1 (z_2 z_3). \end{aligned}$$

(Vergleichen Sie das mit der Methode, wie wir in Abschnitt 4.2.1 das Assoziativgesetz in $\mathbb{Z}$ und die kanonische Projektion $\mathbb{Z} \rightarrow \mathbb{Z}/n$ benutzt haben, um das Assoziativgesetz in $\mathbb{Z}/n$ zu beweisen.)

Für das Assoziativgesetz der Addition und das Distributivgesetz kann man genau dieselbe Methode anwenden, um diese auf die entsprechenden Aussagen in $M_2(\mathbb{R})$ zurückzuführen. $\diamond$

5.3.3. Bild und Kern einer Matrix. Wir wollen in diesem Abschnitt beschreiben, wie wir einer Matrix $A \in M_{m \times n}(K)$ eine Abbildung $f_A: K^n \rightarrow K^m, x \mapsto Ax$, zuordnen können, und erklären, wie wir damit einige Aussagen über lineare Gleichungssysteme noch einmal umformulieren können. Diese neue Sichtweise ist ein wichtiger Spezialfall der Theorie der linearen Abbildungen, die wir in Kapitel 7 entwickeln werden. Diese beiden Kapitel stellen einen sehr wichtigen Teil der Vorlesung dar, so dass es die Zeit wert ist, uns an dieser Stelle schon ein bisschen darauf vorzubereiten.

Das Grundprinzip ist sehr einfach:

DEFINITION 5.43. Seien K ein Körper und $A \in M_{m \times n}(K)$ eine Matrix über K . Wir nennen die Abbildung

$$f_A: K^n \rightarrow K^m, \quad x \mapsto Ax,$$

die durch A definierte Abbildung. $\dashv$

Zwei wichtige Begriffe über Abbildungen der Form f_A sind das Bild (wie für jede Abbildung) und der Kern:

DEFINITION 5.44. Sei $A \in M_{m \times n}(K)$. Das *Bild von A*, in Zeichen $\text{Im } A$, ist das Bild der Abbildung $f_A: K^n \rightarrow K^m$, d.h.

$$\text{Im } A = \{y \in K^m; \text{ es existiert } x \in K^n, \text{ so dass } Ax = y\}.$$

⊢

DEFINITION 5.45. Sei $A \in M_{m \times n}(K)$. Der *Kern von A*, in Zeichen $\text{Ker } A$, ist das Urbild des Elements $0 \in K^m$ unter der Abbildung f_A , d.h.

$$\text{Ker } A = \{x \in K^n; Ax = 0\}$$

⊢

Wir können die neuen Sprechweisen benutzen, um in prägnanter Form über lineare Gleichungssysteme zu sprechen (und später, um die Theorie der linearen Gleichungssysteme in natürlicher Weise in den allgemeinen Kontext einzuordnen):

SATZ 5.46. Sei $A \in M_{m \times n}(K)$ eine Matrix.

- (1) Der Teilraum $\text{Ker}(A)$ ist die Lösungsmenge des homogenen linearen Gleichungssystems mit Koeffizientenmatrix A .
- (2) Der Teilraum $\text{Im}(A)$ ist die Menge aller $b \in K^m$, so dass das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$ eine Lösung besitzt.

BEWEIS. Beide Teile sind direkte Übersetzungen der Definitionen $f_A(x) = Ax$, des Kerns und Bildes und der Tatsache, dass die Lösungsmenge des linearen Gleichungssystems mit erweiterter Koeffizientenmatrix $(A \mid b)$ die Menge $\{x \in K^n; Ax = b\}$ ist. $\square$

Ein weiteres Indiz dafür, dass die Abbildungen der Form f_A gut zu den Strukturen passen, auf die wir beim Studium linearer Gleichungssysteme gestoßen sind, ist das folgende Lemma:

LEMMA 5.47. Sei $A \in M_{m \times n}(K)$ eine Matrix. Dann gilt:

- (1) $\text{Ker}(A)$ ist ein Teilraum von K^n ,
- (2) $\text{Im}(A)$ ist ein Teilraum von K^m .

BEWEIS. zu (1). Das folgt mit dem vorherigen Satz aus Lemma 5.26. Natürlich lässt sich die Aussage auch leicht direkt nachrechnen.

zu (2). Wegen $f_A(0) = 0$ ist $0 \in \text{Im}(f_A)$. Sind $v, w \in \text{Im}(f_A)$, etwa $v = f_A(v')$, $w = f_A(w')$, so gilt

$$v + w = f_A(v') + f_A(w') = Av' + Aw' = A(v' + w') = f_A(v' + w') \in \text{Im}(f_A).$$

Für $v = f_A(v')$ und $a \in K$ gilt $av = aAv' = A(av') = f_A(av') \in \text{Im}(f_A)$. Damit haben wir alle Teilraumbedingungen überprüft. $\square$

Wir können jetzt Frage 5.27 (2) als die Frage formulieren, ob wir beschreiben können, »wie groß« die Teilräume $\text{Ker}(f_A)$ und $\text{Im}(f_A)$ sind, und einen Zusammenhang dazwischen herstellen können.

Die Verknüpfung von Abbildungen der Form f_A ist gerade durch das Matrizenprodukt gegeben. Das liefert eine weitere Rechtfertigung für die Definition des Matrizenprodukts, so wie wir sie gegeben haben.

SATZ 5.48. Seien $A \in M_{l \times m}(K)$, $B \in M_{m \times n}$. Dann gilt

$$f_A \circ f_B = f_{AB}.$$

BEWEIS. Dies folgt aus der Assoziativität des Matrizenprodukts, wie die folgende Rechnung zeigt:

$$(f_A \circ f_B)(x) = f_A(f_B(x)) = A(Bx) = (AB)x = f_{AB}(x)$$

für alle $x \in K^n$. □

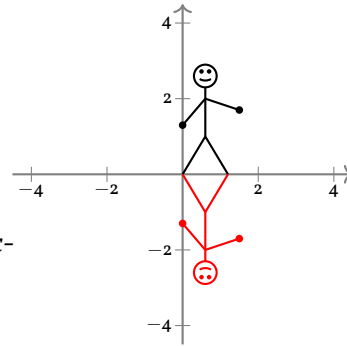
BEISPIEL 5.49. Sei $K = \mathbb{R}$. In diesem Fall kann man die Abbildungen f_A geometrisch interpretieren. Im Fall $n = 2$ (und mit gewissen Grenzen im Fall $n = 3$) kann man das auch graphisch gut veranschaulichen. In den Abbildungen unten ist jeweils die rote Figur das Bild der schwarzen Figur unter der betrachteten Abbildung.

(1)

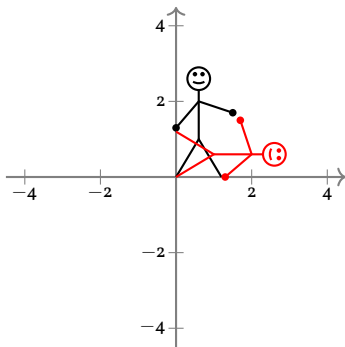
Sei $A = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. Dann gilt

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \\ -y \end{pmatrix}$$

und diese Abbildung ist einfach die *Spiegelung an der x-Achse*.



(2)

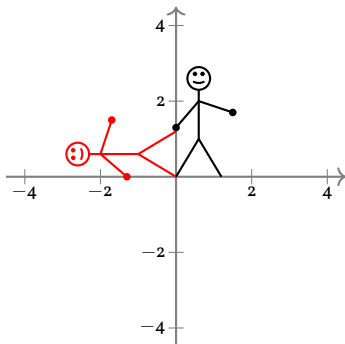


Sei $A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$. Dann gilt

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} y \\ x \end{pmatrix}$$

und diese Abbildung ist die *Spiegelung an der Diagonalen*. Jede Spiegelung an einer Ursprungsgeraden in $\mathbb{R}^2$ kann in der Form f_A dargestellt werden. Können Sie für eine gegebene Gerade $g = \{(x, y)^t \in \mathbb{R}^2; ax + by = 0\}$, $a, b \in \mathbb{R}$, die Matrix A ausrechnen, für die f_A die Spiegelung an g ist?

(3)



Sei $A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$. Dann gilt

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} -y \\ x \end{pmatrix}$$

und diese Abbildung ist die *Drehung um 90° gegen den Uhrzeigersinn*.

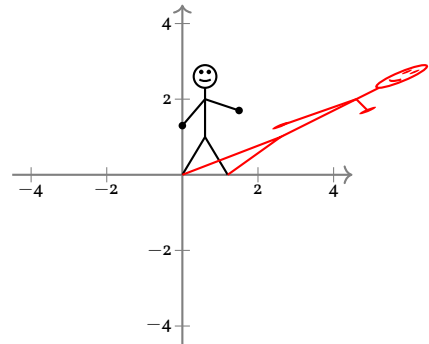
Können Sie auch die Drehung um 180° und um 270° gegen den Uhrzeigersinn in der Form f_A für geeignete Matrizen A darstellen? Wir werden später sehen, dass sich alle Drehungen um den Ursprung in der Form f_A darstellen lassen, siehe Beispiel 7.59.

(4)

Sei $A = \begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix}$. Dann gilt

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x + 2y \\ y \end{pmatrix}$$

und diese Abbildung nennt man eine *Scherung*.



(5)

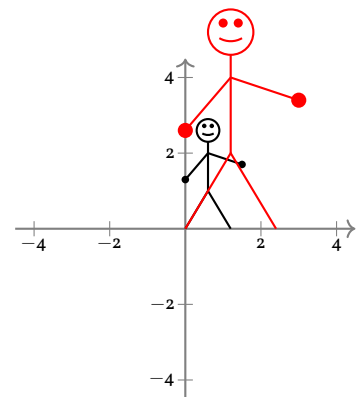
Sei $A = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix}$. Dann gilt

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 2x \\ 2y \end{pmatrix}.$$

Diese Abbildung nennt man eine *Streckung* (um den Faktor 2).

Wie würden Sie die Abbildung f_A beschreiben, wenn

$A = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix}$ mit $a < 0$ ist? Was macht f_A , wenn A eine Diagonalmatrix mit verschiedenen Werten auf der Diagonale ist? Was passiert, wenn eine Null auf der Diagonale steht?



◇

Auf [dieser Webseite](https://math.ug/applets/lineare-abbildungen-1.html)^a können Sie das ausprobieren.

^a<https://math.ug/applets/lineare-abbildungen-1.html>

Wir werden die Abbildungen der Form f_A in Kapitel 7 genauer untersuchen und dann unter anderem charakterisieren können, welche Abbildungen $K^n \rightarrow K^m$ in dieser Art beschrieben werden können.

5.3.4. Invertierbare Matrizen. Wir haben schon besprochen, dass elementare Zeilenumformungen sich als Multiplikation mit einer Matrix von links beschreiben lassen, und dass sie umkehrbar sind. Wenn S die Matrix zu einer elementaren Zeilenumformung von A ist, dann gibt es eine Matrix S' , die ebenfalls eine elementare Zeilenumformung beschreibt, so dass $S'SA = A$. Das führt uns auf die folgende Definition:

DEFINITION 5.50. Eine quadratische Matrix $A \in M_{n \times n}(K)$ heißt *invertierbar*, wenn eine Matrix $B \in M_{n \times n}(K)$ existiert mit

$$AB = BA = E_n.$$

⊢

LEMMA 5.51. Sei $A \in M_{n \times n}(K)$ eine invertierbare Matrix. Dann gibt es genau eine Matrix $B \in M_{n \times n}(K)$ mit $AB = BA = E_n$. Wir bezeichnen sie mit A^{-1} und nennen sie die zu A inverse Matrix. Es ist dann auch B invertierbar, und $B^{-1} = A$.

BEWEIS. Seien B und B' gegeben mit $AB = B'A = E_n$. (Wir brauchen also jeweils nur die eine Hälfte der definierenden Eigenschaft, allerdings einmal die Rechts- und einmal die Linksversion.) Dann gilt

$$B = E_n B = (B'A)B = B'(AB) = B'E_n = B'.$$

Der zweite Teil der Behauptung ist dann klar. $\square$

Wenn die Matrix A invertierbar ist, dann können wir aus dem Ergebnis $C = AB$ einer Matrizenmultiplikation die Matrix B zurückgewinnen: $B = A^{-1}C$. Insbesondere gilt: Ist A invertierbar und $Ax = b$ ein lineares Gleichungssystem, so ist dieses eindeutig lösbar und die Lösung ist $A^{-1}b$. Allerdings ist es selten sinnvoll, A^{-1} auszurechnen, wenn man nur ein lineares Gleichungssystem mit Koeffizientenmatrix A lösen möchte.

LEMMA 5.52. Seien $A, B \in M_n(K)$ invertierbar. Dann ist auch das Produkt AB invertierbar und es gilt

$$(AB)^{-1} = B^{-1}A^{-1}.$$

BEWEIS. Es gilt $(AB)(B^{-1}A^{-1}) = E_n = (B^{-1}A^{-1})(AB)$. $\square$

BEISPIEL 5.53. Wir benutzen die Notation von Bemerkung 5.37. Sei K ein Körper, $n \in \mathbb{N}$, $1 \leq i, j \leq n$, $i \neq j$, $a \in K$.

- (1) Es gilt $E_{ij}(a)^{-1} = E_{ij}(-a)$.
- (2) Es gilt $P_{ij}^{-1} = P_{ij}$.
- (3) Eine Diagonalmatrix $\text{diag}(a_1, \dots, a_n)$ mit $a_i \in K$ ist genau dann invertierbar, wenn $a_i \neq 0$ für alle $i = 1, \dots, n$ gilt. In diesem Fall ist

$$\text{diag}(a_1, \dots, a_n)^{-1} = \text{diag}(a_1^{-1}, \dots, a_n^{-1})$$

Insbesondere ist für $a \neq 0$: $\text{diag}(1, \dots, 1, a, 1, \dots, 1)^{-1} = \text{diag}(1, \dots, 1, a^{-1}, 1, \dots, 1)$.

Wir sehen erneut, dass elementare Zeilenumformungen die Lösungsmenge des zugehörigen homogenen linearen Gleichungssystems nicht ändern, denn für jede invertierbare Matrix S ist $Ax = 0$ äquivalent zu $SAX = 0$. $\diamond$

SATZ 5.54. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in M_n(K)$. Dann sind äquivalent:

- (i) Die Matrix A ist invertierbar.
- (ii) Die reduzierte Zeilenstufenform von A ist E_n .

Erinnern Sie sich, dass wir in Korollar 5.22 die Eigenschaft von A , reduzierte Zeilenstufenform E_n zu haben, in Termen von linearen Gleichungssystemen umgeschrieben haben. Dies werden wir im Beweis des Satzes benutzen.

BEWEIS. Wenn A invertierbar ist, dann ist $Ax = 0$ äquivalent zu $x = A^{-1}0 = 0$, das homogene lineare Gleichungssystem zu A ist daher eindeutig lösbar. Aus Korollar 5.22 folgt, dass A reduzierte Zeilenstufenform E_n hat.

Nun habe A reduzierte Zeilenstufenform E_n . Dann gibt es elementare Zeilenumformungen, die A in die Einheitsmatrix umformen. Wir können diese als Multiplikation mit Matrizen $S_1, S_2, \dots, S_l$ ausdrücken:

$$S_l S_{l-1} \cdots S_1 A = E_n.$$

Wir haben in Beispiel 5.53 gesehen, dass alle S_i invertierbar sind. Also können wir mit ihren Inversen von links multiplizieren und erhalten

$$A = S_1^{-1} S_2^{-1} \cdots S_l^{-1} E_n.$$

Die rechte Seite ist invertierbar und genauer erhalten wir

$$A^{-1} = S_l \cdots S_1.$$

□

Wenn wir am Schluss das Beweises $A^{-1} = S_l \cdots S_1 E_n$ schreiben, dann liefert uns dies direkt einen Algorithmus, um A^{-1} zu berechnen: Wenn wir A durch elementare Zeilenumformungen, die durch die Matrizen $S_1, \dots, S_l$ beschrieben werden, auf die Einheitsmatrix bringen können, dann bringen *genau dieselben* Zeilenumformungen die Einheitsmatrix auf die Matrix A^{-1} .

BEISPIEL 5.55. Um das Verfahren zum bestimmen des Inversen einer Matrix $A \in M_n(K)$ in der Praxis durchzuführen, schreibt man A und E_n nebeneinander in eine $(n \times 2n)$ -Matrix. Etwas übersichtlicher wird es, wenn man in der Mitte einen senkrechten Strich mitführt. Dann bringt man mit elementaren Zeilenumformungen die Matrix A auf reduzierte Zeilenstufenform und führt die Umformungen immer für die gesamte Matrix durch. Ist das Ergebnis der reduzierten Zeilenstufenform von A die Einheitsmatrix, dann ist A invertierbar, und die rechte Hälfte der so erhaltenen Matrix ist A^{-1} . Erhält man eine andere Matrix als E_n als reduzierte Zeilenstufenform von A , so ist die Matrix A nicht invertierbar.

Ein konkretes Beispiel: Sei $K = \mathbb{Q}$, $n = 3$ und

$$A = \begin{pmatrix} 1 & 1 & 0 \\ 2 & 0 & 1 \\ 3 & 1 & 0 \end{pmatrix}.$$

Wir bringen nun in der Matrix

$$(A \mid E_3) = \left(\begin{array}{ccc|ccc} 1 & 1 & 0 & 1 & 0 & 0 \\ 2 & 0 & 1 & 0 & 1 & 0 \\ 3 & 1 & 0 & 0 & 0 & 1 \end{array} \right)$$

die linke Hälfte durch elementare Zeilenumformungen auf reduzierte Zeilenstufenform:

$$\begin{aligned} & \left(\begin{array}{ccc|ccc} 1 & 1 & 0 & 1 & 0 & 0 \\ 2 & 0 & 1 & 0 & 1 & 0 \\ 3 & 1 & 0 & 0 & 0 & 1 \end{array} \right) \xrightarrow{Z_2 \rightsquigarrow Z_2 - 2Z_1, Z_3 \rightsquigarrow Z_3 - 3Z_1} \\ & \left(\begin{array}{ccc|ccc} 1 & 1 & 0 & 1 & 0 & 0 \\ 0 & -2 & 1 & -2 & 1 & 0 \\ 0 & -2 & 0 & -3 & 0 & 1 \end{array} \right) \xrightarrow{Z_3 \rightsquigarrow Z_3 - Z_2, Z_2 \rightsquigarrow -\frac{1}{2}Z_2} \\ & \left(\begin{array}{ccc|ccc} 1 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & -\frac{1}{2} & 1 & -\frac{1}{2} & 0 \\ 0 & 0 & -1 & -1 & -1 & 1 \end{array} \right) \xrightarrow{Z_3 \rightsquigarrow -Z_3, Z_2 \rightsquigarrow Z_2 + \frac{1}{2}Z_3, Z_1 \rightsquigarrow Z_1 - Z_2} \\ & \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & -\frac{1}{2} & 0 & \frac{1}{2} \\ 0 & 1 & 0 & \frac{3}{2} & 0 & -\frac{1}{2} \\ 0 & 0 & 1 & 1 & 1 & -1 \end{array} \right) \end{aligned}$$

Wir erhalten

$$A^{-1} = \begin{pmatrix} -\frac{1}{2} & 0 & \frac{1}{2} \\ \frac{3}{2} & 0 & -\frac{1}{2} \\ 1 & 1 & -1 \end{pmatrix}.$$

Es bietet sich bei dieser Rechnung an, die Probe durchzuführen: Das Produkt AA^{-1} muss die Einheitsmatrix ergeben. ◇

BEISPIEL 5.56. Seien K ein Körper und $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ eine 2×2 -Matrix.

Wir setzen $\delta(A) = ad - bc$. Wir haben in Abschnitt 2.5 gesehen, dass das homogene lineare Gleichungssystem zu A genau dann eindeutig lösbar ist, wenn $\delta(A) \neq 0$ gilt. Aus Satz 5.54 und Korollar 5.22 folgt damit, dass

$$A \text{ ist invertierbar} \iff \delta(A) \neq 0.$$

In diesem Fall gilt

$$A^{-1} = \frac{1}{\delta(A)} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

Um das zu sehen, muss man nur nachrechnen, dass $AA^{-1} = E_2$ gilt, und das ist eine leichte Rechnung.

Den Ausdruck $\delta(A)$ nennt man auch die *Determinante* der Matrix A (und später schreiben wir $\det(A)$ statt $\delta(A)$). Siehe Kapitel 9 für eine systematische Diskussion und die Verallgemeinerung auf quadratische Matrizen beliebiger Größe. $\diamond$

LEMMA 5.57. Sei $A \in M_{n \times n}(K)$ eine invertierbare Matrix. Dann ist auch die transponierte Matrix A^t invertierbar, und es gilt

$$(A^t)^{-1} = (A^{-1})^t.$$

BEWEIS. Wir haben $A^t(A^{-1})^t = (A^{-1}A)^t = E_n^t = E_n$ nach Lemma 5.40, analog $(A^{-1})^t A^t = E_n$, und zusammen das zeigt die Behauptung. $\square$

Das folgende Korollar ist tiefliegender als es erscheinen mag. Wir werden auf diese Aussage später wieder zurückkommen (und einen anderen Beweis geben), siehe Korollar 7.24.

KOROLLAR 5.58. Sei $A \in M_{n \times n}(K)$ eine Matrix.

- (1) Wenn eine Matrix $B \in M_{n \times n}(K)$ existiert mit $AB = E_n$, dann ist A invertierbar und $A^{-1} = B$.
- (2) Wenn eine Matrix $B \in M_{n \times n}(K)$ existiert mit $BA = E_n$, dann ist A invertierbar und $A^{-1} = B$.

BEWEIS. zu (1). Es genügt zu zeigen, dass A invertierbar ist; es ist dann klar, dass $B = A^{-1}$ gelten muss. Für die Invertierbarkeit genügt es wegen Satz 5.54 und Korollar 5.22 zu zeigen, dass das lineare Gleichungssystem $Ax = b$ für alle $b \in K^n$ lösbar ist. Das folgt aber unmittelbar aus der Voraussetzung, denn wir setzen einfach $x = Bb$ und erhalten $Ax = A(Bb) = E_nb = b$.

zu (2). Für diesen Teil können wir wieder Satz 5.54 und Korollar 5.22 anwenden und zeigen diesmal, dass das homogene lineare Gleichungssystem $Ax = 0$ eindeutig lösbar ist. In der Tat, ist $x \in K^n$ mit $Ax = 0$, so folgt $x = E_n x = BAx = 0$.

Alternativ kann man Teil (2) aus Teil (1) (oder auch umgekehrt) durch den Übergang zu den transponierten Matrizen erhalten. $\square$

Wir können nun auch noch einmal auf die Eindeutigkeit der reduzierten Zeilenstufenform zurückkommen:

SATZ 5.59. Die reduzierte Zeilenstufenform einer Matrix $A \in M_{m \times n}(K)$ ist eindeutig bestimmt, also unabhängig von der Wahl der elementaren Zeilenumformungen.

BEWEIS. Seien A' und A'' Matrizen in reduzierter Zeilenstufenform, die beide aus A durch elementare Zeilenumformungen erhalten werden können. Dann gibt es invertierbare Matrizen S' und S'' mit $A' = S'A$ und $A'' = S''A$, also

$$A' = S'(S'')^{-1}A''.$$

Wir bezeichnen mit B die Matrix, die aus den Spalten von A' besteht, die eine führende Eins enthalten. Alle anderen Spalten lassen wir weg. Ist r die Anzahl der führenden Einsen in A' , so ist $B \in M_{n \times r}(K)$. Wir wissen bereits, dass sich in A' und A'' die Spalten mit führenden Einsen an denselben Positionen befinden (Satz 5.17). Deshalb erhalten wir genau dasselbe Ergebnis, wenn wir B in der gleichen Art und Weise aus A'' statt aus A' konstruieren, und es folgt

$$B = S'(S'')^{-1}B.$$

Die Matrix B hat eine sehr einfache Form, es handelt sich um eine Blockmatrix, die aus der Einheitsmatrix E_r und der Nullmatrix der Größe $(n-r) \times r$ zusammengesetzt ist: $B = \begin{pmatrix} E_r \\ 0 \end{pmatrix}$ (und die Zeilen unterhalb der r -ten Zeilen sind auch in A' und A'' Nullzeilen).

Die Gleichheit $B = S'(S'')^{-1}B$ übersetzt sich damit in

$$S'(S'')^{-1} = \begin{pmatrix} E_r & * \\ 0 & * \end{pmatrix},$$

wobei das obere Sternchen eine Matrix in $M_{r \times (n-r)}(K)$ und das untere eine in $M_{(n-r) \times (n-r)}(K)$ bezeichnen, über die wir nichts Genauereres zu sagen brauchen.

Weil die Zeilen unterhalb der r -ten Zeile in A' und A'' sowieso Nullzeilen sind, genügt es zu zeigen, dass die ersten r Zeilen übereinstimmen. Wir schreiben (nur für diesen Beweis) $M_{\leq r}$ für die Matrix, die aus den ersten r Zeilen der Matrix M besteht. Dann haben wir

$$(A')_{\leq r} = (S'(S'')^{-1}A'')_{\leq r} = (S'(S'')^{-1})_{\leq r} A'' = \begin{pmatrix} E_r & * \end{pmatrix} \begin{pmatrix} (A'')_{\leq r} \\ 0 \end{pmatrix} = (A'')_{\leq r},$$

wobei wir ausgiebig von der Schreibweise als Blockmatrizen Gebrauch gemacht haben (Abschnitt 5.3.2 und Lemma 5.41). $\square$

5.3.5. Ergänzungen *

BEISPIEL 5.60. Wir betrachten die Folge $(F_n)_{n \geq 0}$ der Fibonacci Zahlen, die definiert ist durch

$$F_0 = 0, \quad F_1 = 1, \quad F_n = F_{n-1} + F_{n-2}, \quad n \geq 2.$$

Wir können die rekursive Definition durch die folgenden Matrixgleichungen formulieren:

$$\begin{pmatrix} F_1 \\ F_0 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad \begin{pmatrix} F_n \\ F_{n-1} \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} F_{n-1} \\ F_{n-2} \end{pmatrix}.$$

Setzen wir

$$A = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix},$$

so erhalten wir dadurch

$$\begin{pmatrix} F_{n+1} \\ F_n \end{pmatrix} = A^n \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

Dies kann man ausnutzen, um »schnell« eine Fibonacci-Zahl F_n für großes n zu berechnen, ohne alle vorherigen Fibonacci-Zahlen berechnen zu müssen. Denn die Potenz A^n kann man in wesentlich weniger als n Schritten ausrechnen: Wir schreiben $n = \sum_{i=0}^r a_i \cdot 2^i$ mit $a_i \in \{0, 1\}$, $a_r \neq 0$ (dies ist die Binärdarstellung der Zahl n). Zuerst berechnet man A^2 , dann $A^4 = A^2 \cdot A^2$, und alle weiteren Potenzen A^{2^i} für $i \leq r$. Dies erfordert nur r Matrixmultiplikationen. Dann berechnen wir A^{n-1} als Produkt derjenigen Matrizen A^{2^i} , für die $a_i = 1$ ist.

Ist zum Beispiel $n = 16781841 = 2^{24} + 2^{12} + 2^9 + 2^4 + 2^0$, so brauchen wir mit diesem Rezept für die Berechnung von F_n nur 28 Produkte von (2×2) -Matrizen auszurechnen. Das wäre notfalls noch per Hand machbar (im Gegensatz zu n Additionen, die bei einer Geschwindigkeit von einer Addition pro Sekunde fast 200 Tage dauern würden – wenn Sie rund um die Uhr addieren).

Die Darstellung mithilfe der Matrix A kann man auch benutzen, um Identitäten zwischen verschiedenen Gliedern der Fibonacci-Folge herzuleiten. Dazu bemerken wir zunächst, dass für $n \geq 1$

$$A^n = \begin{pmatrix} F_{n+1} & F_n \\ F_n & F_{n-1} \end{pmatrix},$$

wie man unmittelbar per Induktion nachprüft. Damit folgt zum Beispiel für $m, n \geq 1$ aus der Gleichheit $A^{m+n} = A^m A^n$, dass

$$\begin{pmatrix} F_{m+n+1} & F_{m+n} \\ F_{m+n} & F_{m+n-1} \end{pmatrix} = A^m A^n = \begin{pmatrix} F_{m+1} & F_m \\ F_m & F_{m-1} \end{pmatrix} \begin{pmatrix} F_{n+1} & F_n \\ F_n & F_{n-1} \end{pmatrix}.$$

Nehmen wir nun auf beiden Seiten den Eintrag in der zweiten Zeile und ersten Spalte, so erhalten wir die Formel

$$F_{m+n} = F_m F_{n+1} + F_{m-1} F_n.$$

Durch Spezialisierung erhält man viele verschiedene Identitäten zwischen den Fibonacci-Zahlen, zum Beispiel für $m = n$: $F_{2n} = (F_{n+1} + F_{n-1})F_n$. $\diamond$

ERGÄNZUNG 5.61 (Der Page-rank-Algorithmus, Fortsetzung). Dies ist die Fortsetzung der Diskussion aus Frage 2.7. Wir bezeichnen wie dort mit x_i , $i = 1, \dots, N$, die Zahlen die wir suchen, um die »Relevanz« von Webseiten zu beschreiben, mit L_i die Menge aller Seiten, die einen Link auf Seite i enthalten, und mit n_j die Zahl der von Seite j ausgehenden Links. Wir arbeiten über dem Körper $K = \mathbb{R}$.

Der aktuelle Diskussionsstand ist, dass die x_i das lineare Gleichungssystem

$$x_i = \sum_{j \in L_i} \frac{1}{n_j} x_j, \quad i = 1, \dots, N,$$

erfüllen sollen.

Wir suchen eine Lösung, in der $0 \leq x_i < 1$ für alle i und $\sum x_i = 1$ gilt, und zwar sollte es *genau eine* Lösung mit dieser Eigenschaft geben. Durch die Zusatzeigenschaften können wir die x_i auch als Wahrscheinlichkeit deuten, dass jemand, der im Internet surft und immer zufällig irgendeinem Link folgt, sich gerade auf der Seite i befindet.

Sei H die Matrix mit Einträgen

$$H_{ij} = \begin{cases} \frac{1}{n_j} & j \in L_i, \\ 0 & \text{sonst.} \end{cases}$$

Dann können wir das obige lineare Gleichungssystem schreiben als $(H - E_N)x = 0$, wobei $x = (x_i)_i \in K^N$.

Es gibt noch zwei »offensichtliche« Probleme mit diesem Ansatz:

- Es lassen sich leicht Beispiele finden, in denen dieses homogene lineare Gleichungssystem keine nicht-triviale Lösung hat, und zwar dann, wenn es Seiten gibt, die gar keine Links enthalten. Zum Beispiel ein »Internet« mit drei Seiten und den Links $1 \rightarrow 2$, $2 \rightarrow 1$, $2 \rightarrow 3$. In diesem Fall würde der Ansatz scheitern.

Wir wollen dieses Problem so lösen, dass wir in jeder Nullspalte der Matrix H alle Einträge durch $\frac{1}{N}$ ersetzen. Mit der Interpretation als Wahrscheinlichkeiten können wir uns das so vorstellen, dass die Surfer*in von einer Seite, die keine Links enthält, zufällig irgendeine Seite im Netz auswählt, und allen Seiten dieselbe Wahrscheinlichkeit zugewiesen wird.

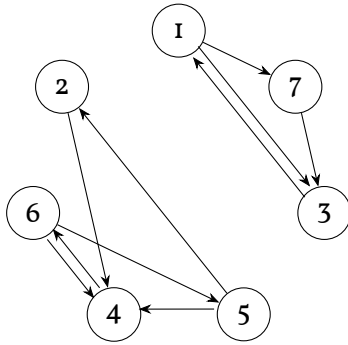
Wenn wir H so abändern:

$$H_{ij} = \begin{cases} \frac{1}{n_j} & j \in L_i, \\ \frac{1}{N} & \text{wenn kein } i' \text{ existiert mit } j \in L_{i'}, \\ 0 & \text{sonst,} \end{cases}$$

dann hat die Matrix H die Eigenschaft, dass alle Spaltensummen in H gleich 1 sind, d.h. $\sum_{i=1}^N H_{ij} = 1$ für alle j .

In den Spalten, die vorher Nullspalten waren, ist das klar, weil nun dort alle N Einträge den Wert $\frac{1}{N}$ haben. In den anderen Spalten ist die Spaltensumme $\sum \frac{1}{n_j}$, wobei über alle i mit $j \in L_i$ summiert wird, also über alle Seiten i , auf die die Seite j verlinkt, und nach Definition gibt es genau n_j solche Seiten.

- Stellen Sie sich vor, dass das »Netz« von Webseiten, das wir betrachten, verschiedene Komponenten hat, zwischen denen man nicht durch Links hin- und hergelangen kann:



In diesem Fall ist es klar, dass die gegebenen Daten keinen vernünftigen Anhaltspunkt liefern, um die Relevanz von Seiten, die in verschiedenen solchen Komponenten liegen, zu vergleichen. Das Internet hat sehr viele Komponenten, so dass man dieses Problem nicht einfach ignorieren kann.

Der Trick, mit dem wir uns (bzw. Google sich) behelfen, ist dass wir die Matrix H ersetzen durch die Matrix

$$G := dH + (1 - d)A,$$

wobei A die Matrix ist, deren Einträge alle gleich $\frac{1}{N}$ sind, und wobei d eine Zahl ist, die zwischen 0 und 1 liegt. Je näher d bei 1 liegt, desto geringer ist der Unterschied zwischen G und H . Interessanterweise ist es für die Qualität des Algorithmus aber sogar hilfreich, d nicht zu nahe bei 1 zu wählen. Google hat ursprünglich den Wert $d = 0,85$ benutzt.

Das passt auch gut zur Interpretation der x_i als Aufenthaltswahrscheinlichkeiten: Es bedeutet, dass die Surfer*in mit Wahrscheinlichkeit $1 - d$ (also im Beispiel 15%) nicht einem Link auf der Seite folgt, sondern auf irgendeine zufällige Seite im Netz springt.

Wir haben damit die endgültige mathematische Formulierung des Page-Rank-Algorithmus erhalten: Wir möchten das lineare Gleichungssystem $(G - E_N)x = 0$ lösen. In der Matrix G (der »Google-Matrix«) sind alle Einträge positiv, und die Spaltensummen sind alle gleich 1. Während das nützliche Eigenschaften der Matrix G sind, ist andererseits diese Matrix *extrem groß*, so groß, dass es aussichtslos ist, einfach den Gauß-Algorithmus auszuführen.

Wir werden später sehen, dass die Lösungsmenge des homogenen linearen Gleichungssystems $(G - E_N)x = 0$ aus den Vielfachen eines einzigen Vektors $x \neq 0$ besteht. Bis auf Skalieren gibt es also eine *eindeutig bestimmte Lösung*. Genauer werden wir den folgenden Satz beweisen (vergleiche Beispiel 2.8):

SATZ 5.62. *Es gibt genau eine Lösung $x \in \mathbb{R}^N$ des linearen Gleichungssystems $(G - E_N)x = 0$ mit der Eigenschaft $\sum_{i=1}^N x_i = 1$, und für diese Lösung gilt $0 \leq x_i \leq 1$ für alle i .*

(Ist $x' \in \mathbb{R}^N$ ein Element, das auch $(G - E_N)x' = 0$ erfüllt, so existiert $\lambda \in \mathbb{R}$ mit $x' = \lambda x$.)

Den Beweis geben wir wie gesagt später, siehe Korollar 7.69. Dass eine nicht-triviale Lösung existiert, ist jedenfalls leicht zu sehen. Weil die Spaltensummen von F alle $= 1$ sind, sind alle Spaltensummen von $G - E_N$ gleich Null. Addieren wir nacheinander die erste, zweite,



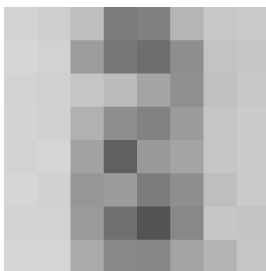
ABBILDUNG 1. Das Originalfoto

..., $(N - 1)$ -te Zeile zur letzten Zeile von $G - E_N$, so entsteht also dort eine Nullzeile. Dies sind elementare Zeilenumformungen, die die Lösungsmenge des zugehörigen linearen Gleichungssystems nicht ändern, und da eine Nullzeile existiert, muss es in der (reduzierten) Zeilenstufenform eine Spalte ohne führende Eins geben. □ Ergänzung 5.61

ERGÄNZUNG 5.63 (Bildkompression). In dieser Ergänzung besprechen wir eine Methode der Bildkompression mit dem sogenannten Haar-Wavelet. Wir beschränken uns darauf, die Grundidee zu erklären. In der Praxis kann man das Verfahren noch weiter verbessern. Es gibt auch andere sehr gute Kompressionsverfahren, die auf linearer Algebra beruhen, zum Beispiel auf der Singulärwertzerlegung, die wir in der Linearen Algebra 2 kennenlernen werden.

Um das Prinzip zu erläutern und zu illustrieren, beschränken wir uns auf ein Schwarz-Weiß-Bild. Man kann aber genau dieselbe Überlegung auf Farbbilder anwenden, indem man das Bild in mehrere Farben aufspaltet, beispielsweise in Rot, Grün und Blau (RGB), und dann die drei Farben separat behandelt.

Wir betrachten ein Bild (Abbildung 1) mit 512×512 Punkten (»Pixeln«), die jeweils einen Grauwert haben, der durch eine natürliche Zahl zwischen 0 und 255 gegeben ist. Dabei ist 0 gleich schwarz, und 255 ist gleich weiß. Um die Bildinformation in dieser vollständigen Form abzuspeichern, müssen also ca. 260 000 dieser Grauwerte abgespeichert werden, jeder Wert benötigt ein »Byte« an Speicherplatz. Insgesamt brauchen wir also ca. 260 Kilobyte für dieses Bild, wenn keine Kompression angewandt wird. (Nehmen Sie statt dieses kleinen Beispiels ein Bild mit beispielsweise 8MP und 3 Farben, so bräuchte man 24 Megabyte, um das Bild abzuspeichern.)



Die Grundidee aller Kompressionsverfahren ist es, auszunutzen, dass es in einem typischem Bild viele Bereiche gibt, die mehr oder weniger einfarbig sind, d.h. benachbarte Punkte haben sehr ähnliche Farbwerte bzw. in unserem Fall Grauwerte.

Für das hier vorgestellte Verfahren teilen wir das Bild in Blöcke von 8×8 Pixeln auf und behandeln jeden Block einzeln. Ein vergrößerter 8×8 -Block aus dem Originalbild ist hier dargestellt.

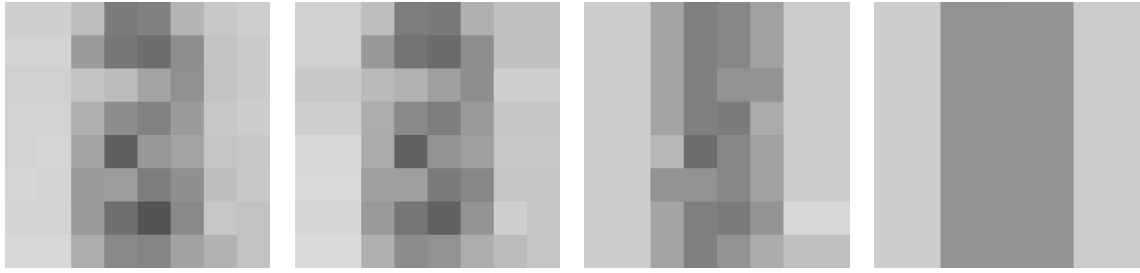


ABBILDUNG 2. Das Ergebnis, das aus dem vorher gezeigten Block von 8×8 Pixeln wird, wenn alle Werte < 1 , bzw. < 4 , bzw. < 10 , bzw. < 18 auf Null gesetzt werden.

Wir nehmen nun einen dieser Blöcke her und beginnen damit, die erste Zeile des Blocks folgendermaßen umzuformen: Wir fassen die 8 Werte als 4 Paare auf, d.h. die ersten beiden Einträge bilden das erste Paar, der dritte und vierte Eintrag das zweite Paar, usw. Für jedes dieser Paare (a, b) bilden wir den Durchschnitt $\frac{a+b}{2}$ und den Wert $\frac{a-b}{2}$. Aus diesen Werten kann man a und b zurückgewinnen als

$$a = \frac{a+b}{2} + \frac{a-b}{2}, \quad b = \frac{a+b}{2} - \frac{a-b}{2},$$

aber diese Ersetzung hat für uns den Vorteil, dass tendenziell oft $\frac{a-b}{2}$ eine kleine Zahl sein wird (sogar gleich Null, wenn $a = b$), weil typischerweise a und b nahe beieinander liegen. Wir schreiben das Ergebnis wieder als eine Zeile mit 8 Einträgen, wobei wir zuerst die 4 Durchschnittswerte aufschreiben, und dann die 4 Terme der zweiten Form.

Mit anderen Worten: Wir ersetzen die Zeile $v = (v_1, \dots, v_8)$ durch

$$\left(\frac{v_1 + v_2}{2}, \frac{v_3 + v_4}{2}, \frac{v_5 + v_6}{2}, \frac{v_7 + v_8}{2}, \frac{v_1 - v_2}{2}, \frac{v_3 - v_4}{2}, \frac{v_5 - v_6}{2}, \frac{v_7 - v_8}{2} \right)$$

Das können wir auch als ein Matrizenprodukt ausdrücken: Wir ersetzen den Zeilenvektor v durch das Produkt vH_1 , wobei

$$H_1 = \begin{pmatrix} 0,5 & & & & 0,5 & & & \\ 0,5 & & & & -0,5 & & & \\ & 0,5 & & & & 0,5 & & \\ & 0,5 & & & & -0,5 & & \\ & & 0,5 & & & & 0,5 & \\ & & 0,5 & & & & -0,5 & \\ & & & 0,5 & & & & 0,5 \\ & & & 0,5 & & & & -0,5 \end{pmatrix}$$

Im zweiten Schritt wenden wir dasselbe Verfahren noch einmal auf die ersten 4 Einträge der neuen Zeile an: Wir teilen diese in zwei Paare auf, und schreiben dann an die ersten beiden Stellen die Durchschnittswerte des ersten und zweiten Paares, und an die dritte und vierte Stelle die Hälfte der Differenzen. Auch dies können wir als Produkt mit einer Matrix

ausdrücken, und zwar mit

$$H_2 = \begin{pmatrix} 0,5 & & 0,5 & & & & & \\ 0,5 & & -0,5 & & & & & \\ & 0,5 & & 0,5 & & & & \\ & 0,5 & & -0,5 & & & & \\ & & & & \text{I} & & & \\ & & & & & \text{I} & & \\ & & & & & & \text{I} & \\ & & & & & & & \text{I} \end{pmatrix}$$

Zum Schluss machen wir das ganze noch einmal für die ersten beiden Einträge der neuen Zeile, das heißt wir multiplizieren mit

$$H_3 = \begin{pmatrix} 0,5 & 0,5 & & & & & & \\ 0,5 & -0,5 & & & & & & \\ & & \text{I} & & & & & \\ & & & \text{I} & & & & \\ & & & & \text{I} & & & \\ & & & & & \text{I} & & \\ & & & & & & \text{I} & \\ & & & & & & & \text{I} \end{pmatrix}$$

Insgesamt ersetzen wir also v durch $vH_1H_2H_3$. Setzen wir $H := H_1H_2H_3$, so wird v durch vH ersetzt. Für die anderen Zeilen machen wir dasselbe, das bedeutet, dass wir den Block A , aufgefasst als (8×8) -Matrix, ersetzen durch AH . Wie oben erwähnt, lässt sich diese Operation auch wieder umkehren (bis jetzt haben wir also keine Information verloren): die Umkehrung erhält man durch Multiplikation mit der inversen Matrix H^{-1} .

Zum Schluss führen wir dasselbe Verfahren für alle Spalten des neu erhaltenen Blocks durch. Das bedeutet, dass wir AH ersetzen durch H^tAH , wobei H^t die zu H transponierte Matrix bezeichnet. In derselben Weise formt man nun alle (8×8) -Blöcke der Bilddatei um.

Bis zum jetzigen Zeitpunkt haben wir die vorhandene Information nur in eine andere Form gebracht und können die Ausgangsdaten vollständig rekonstruieren. Trotzdem ergibt sich in der Regel auch dadurch schon die Möglichkeit, die Daten platzsparender abzuspeichern, weil die neue 512×512 -Matrix in der Regel viele Nulleinträge haben wird. Dann kann man »sich merken«, dass gewisse Bereiche gleich Null sind, statt jeden Eintrag einzeln zu speichern.

Seine volle Kraft entfaltet das Verfahren allerdings, wenn man in H^tAH alle Einträge auf Null setzt, deren Absolutbetrag kleiner ist als ein vorgegebener Wert, beispielweise alle Einträge, die betragsmäßig kleiner sind als 1. Dadurch erhält man normalerweise einen deutlich größeren Anteil an Nullen, so dass die Matrix wesentlich platzsparender abgespeichert werden kann, siehe die Abbildungen 2, 3. Trotzdem kann man, wenn man es nicht übertreibt, das ursprüngliche Bild ohne großen Qualitätsverlust zurückgewinnen, wenn man dasselbe Verfahren zur Rücktransformation anwendet.

Es ist zwar nicht zwingend, dieses Verfahren mithilfe des Matrizenprodukts zu beschreiben, und die ganze Sache ist auch mathematisch nicht sehr tiefgehend. Sie ist aber in der Praxis nützlich, lässt sich wie schon angedeutet auch noch weiter verbessern, und die Formulierung mit dem Matrizenprodukt ist einerseits eine klare und einfache Beschreibung, und ermöglicht es andererseits, auf die Theorie des Matrizenprodukts zurückzugreifen, und für die praktische Implementierung Programmbibliotheken zu benutzen, in denen schnelle Matrizenoperationen zur Verfügung gestellt werden.

Quelle: [VK].

□ Ergänzung 5.63



ABBILDUNG 3. Von links oben nach rechts unten die »komprimierten« Fotos, wo alle Werte < 1 , bzw. < 4 , bzw. < 10 , bzw. < 18 auf 0 gesetzt wurden. Damit wurden ca. 68%, bzw. 86%, bzw. 93%, bzw. 96% der Einträge in der beschreibenden Matrix zu Null.

Im dritten und vierten Bild sieht man bereits die deutlich niedrigere Qualität, an mehreren Stellen zeigen sich »Kompressionsartefakte«. Allerdings wurde die Datenmenge in diesen Fällen auch sehr stark verkleinert. Natürlich müsste man nun, wenn man die Qualität dieses Kompressionsverfahrens bewerten wollte, Vergleiche mit anderen Verfahren durchführen, insbesondere mit dem »naiven« Ansatz, die Datenmenge dadurch zu verkleinern, dass man das Bild auf eine kleinere Auflösung herunterrechnet. Das würde hier ein bisschen zu weit führen, aber ich stelle gerne den Python-Code zur Verfügung, wenn Sie diese Frage aufgreifen und selbst experimentieren möchten.

ERGÄNZUNG 5.64 (Quaternionen, Fortsetzung). Wir können nun auch noch einmal auf die Hamiltonschen Quaternionen zurückkommen (siehe Ergänzung 4.11) und erklären, wie man recht einfach zeigen kann, dass $\mathbb{H}$ ein Schiefkörper ist. Wir betrachten die injektive

(warum?) Abbildung

$$\iota: \mathbb{H} \rightarrow M_2(\mathbb{C}), \quad (a, b, c, d) \mapsto \begin{pmatrix} a + bi & -c - di \\ c - di & a - bi \end{pmatrix}.$$

Es ist leicht nachzurechnen, dass ι verträglich ist mit der Addition und Multiplikation in $\mathbb{H}$ und in $M_2(\mathbb{C})$:

$$\iota(x + y) = \iota(x) + \iota(y), \quad \iota(xy) = \iota(x)\iota(y).$$

Daraus folgen (ganz ähnlich wie in Bemerkung 5.42) das Assoziativgesetz der Multiplikation und das Distributivgesetz.

Es bleibt noch zu zeigen, dass jedes Element in $\mathbb{H} \setminus \{0\}$ ein multiplikatives Inverses besitzt. Dazu genügt es zu zeigen, dass jede Matrix $\neq 0$, die im Bild von ι liegt, invertierbar ist und dass die inverse Matrix ebenfalls im Bild von ι liegt. Dazu benutzen wir Beispiel 5.56 und die dort eingeführte *Determinante* einer (2×2) -Matrix. Es gilt

$$\delta(\iota((a, b, c, d))) = \delta \left(\begin{pmatrix} a + bi & -c - di \\ c - di & a - bi \end{pmatrix} \right) = a^2 + b^2 + c^2 + d^2,$$

und dies ist eine nicht-negative reelle Zahl, die genau dann $= 0$ ist, wenn $(a, b, c, d) = 0$ ist.

Ist der Wert von δ nicht Null, so gilt

$$\begin{aligned} \begin{pmatrix} a + bi & -c - di \\ c - di & a - bi \end{pmatrix}^{-1} &= (a^2 + b^2 + c^2 + d^2)^{-1} \begin{pmatrix} a - bi & c + di \\ -c + di & a + bi \end{pmatrix} \\ &= (a^2 + b^2 + c^2 + d^2)^{-1} \iota(a, -b, -c, -d). \end{aligned}$$

Es ist auch leicht, das Inverse von $(a, b, c, d) \neq 0$ direkt anzugeben:

$$(a, b, c, d)^{-1} = \left(\frac{1}{a^2 + b^2 + c^2 + d^2}, \frac{-b}{a^2 + b^2 + c^2 + d^2}, \frac{-c}{a^2 + b^2 + c^2 + d^2}, \frac{-d}{a^2 + b^2 + c^2 + d^2} \right).$$

Damit haben wir vollständig bewiesen, dass $\mathbb{H}$ ein Schiefkörper ist.

□ Ergänzung 5.64

KAPITEL 6

Vektorräume

Ähnlich wie der Begriff des Körpers die essenziellen Anforderungen an einen »Zahlbereich«, in dem die üblichen Rechenregeln für die Grundrechenarten gelten, zusammenfasst, macht dies der Begriff des Vektorraums, den wir in diesem Kapitel definieren, für die Eigenschaften von K^n und seiner Teilräume – insbesondere also für die Lösungsmengen homogener linearer Gleichungssysteme – und der Rechenoperationen der Addition und der Multiplikation mit Elementen (»Skalaren«) aus dem Körper K .

6.1. Vektorräume

DEFINITION 6.1. Sei K ein Körper. Ein $(K\text{-})$ Vektorraum oder *Vektorraum über K* ist eine Menge V zusammen mit Verknüpfungen $+$: $V \times V \rightarrow V$ (Addition) und $\cdot$: $K \times V \rightarrow V$ (Skalarmultiplikation), so dass gilt:

- (1) (a) Die Verknüpfung $+$ auf V ist *assoziativ*, d. h. für alle $u, v, w \in V$ gilt $(u + v) + w = u + (v + w)$.
- (b) Die Verknüpfung $+$ auf V besitzt ein eindeutig bestimmtes neutrales Element 0 . (Dieses Element nennt man den *Nullvektor* von V .)
- (c) Jedes Element $v \in V$ besitzt ein inverses Element bezüglich $+$, d.h. es existiert $w \in V$ mit $v + w = 0 = w + v$.
- (d) Die Verknüpfung $+$ auf V ist kommutativ: Für alle $v, w \in V$ gilt $v + w = w + v$.
- (2) (a) Für alle $a, b \in K, v \in V$ gilt: $a \cdot (b \cdot v) = (ab) \cdot v$.
- (b) Für alle $v \in V$ gilt $1 \cdot v = v$.
- (c) (Distributivgesetz 1) Für alle $a, b \in K, v \in V$ gilt: $(a + b) \cdot v = a \cdot v + b \cdot v$.
- (d) (Distributivgesetz 2) Für alle $a \in K, v, w \in V$ gilt: $a \cdot (v + w) = a \cdot v + a \cdot w$.

—

Man nennt K auch den *Grundkörper* des Vektorraums V . Die Skalarmultiplikation ist also keine Verknüpfung, die aus zwei Elementen von V ein weiteres produziert. Stattdessen verknüpft man ein »Skalar«, d.h. ein Element des Grundkörpers K , mit einem Element von V . Wie bei der Multiplikation von Elementen eines Körpers lassen wir bei der Skalarmultiplikation den Punkt $\cdot$ üblicherweise aus.

Wie im Fall eines Körpers (vergleiche Lemma 4.7) ist das additive Inverse eines Elements v eindeutig bestimmt. Wir bezeichnen es mit $-v$ und nennen dieses Element auch das *Negative* von v . Wir setzen $v - w := v + (-w)$ und können somit auch davon sprechen, Elemente von V zu subtrahieren.

BEISPIEL 6.2. Sei K ein Körper. In allen folgenden Fällen ist es leicht nachzurechnen, dass die Bedingungen der Vektorraumdefinition (die »Vektorraumaxiome«) erfüllt sind. (In den meisten Fällen haben wir das schon gemacht oder implizit ausgenutzt, ohne den Begriff Vektorraum zu benutzen.)

- (1) Ist $n \in \mathbb{N}$, so ist K^n mit der im vorherigen Kapitel definierten Addition und Skalarmultiplikation ein K -Vektorraum. Dies ist das prototypische Beispiel eines Vektorraums und man nennt K^n den *Standardvektorraum*.
- (2) Der *Nullvektorraum* (oder auch *Nullraum*) ist der K -Vektorraum, der nur ein einziges Element hat. Dies muss dann der Nullvektor des Vektorraums sein. Wir schreiben manchmal 0 für den Nullraum $\{0\}$. Eine andere Bezeichnung für denselben Vektorraum ist K^0 . (Ganz streng genommen könnten wir von unterschiedlichen Nullvektorräumen sprechen, je nachdem, wie wir das eine Element nennen. Da aber dieses Element wegen der Vektorraumaxiome immer die Rolle des Nullvektors in diesem Vektorraum haben muss, handelt es sich doch praktisch gesehen immer um denselben Raum. Wir werden das im folgenden Kapitel präzisieren, siehe Bemerkung 7.12.)
- (3) Der Raum der $(m \times n)$ -Matrizen $M_{m \times n}(K)$ ist mit der Addition von Matrizen und der Skalarmultiplikation (Definition 5.29) ein Vektorraum über K .
- (4) Sei $U \subseteq K^n$ ein Teilraum im Sinne von Definition 5.23. Dann sind die Einschränkungen der Addition und Skalarmultiplikation auf K^n Abbildungen $+: U \times U \rightarrow U$ und $\cdot: K \times U \rightarrow U$, mittels derer U zu einem Vektorraum über K wird.
- (5) Sei M eine Menge. Die Menge $\text{Abb}(M, K)$ aller Abbildungen $M \rightarrow K$ ist ein Vektorraum, wenn wir die Addition und Skalarmultiplikation für $f, g: M \rightarrow K, a \in K$ definieren durch

$$(f + g): M \rightarrow K, m \mapsto f(m) + g(m), \quad (af): M \rightarrow K, m \mapsto af(m).$$

- (6) Die Menge der Polynomfunktionen $K \rightarrow K$ ist ein K -Vektorraum, wenn die Addition und Skalarmultiplikation wie im vorherigen Punkt definiert werden.
- (7) Wir können das Beispiel $\text{Abb}(M, K)$ noch verallgemeinern: Sei W ein K -Vektorraum und sei X irgendeine Menge. Dann ist die Menge $\text{Abb}(X, W)$ mit den folgenden Verknüpfungen ein K -Vektorraum:

$$\begin{aligned} f + g &:= (X \rightarrow W, x \mapsto f(x) + g(x)), \\ a \cdot f &:= (X \rightarrow W, x \mapsto af(x)), \end{aligned}$$

wobei $f, g \in \text{Abb}(X, W), a \in K$. Alle Vektorraumaxiome sind leicht nachzuprüfen; es lässt sich alles auf die entsprechenden Bedingungen in W zurückführen. Der Nullvektor in diesem Vektorraum ist die konstante Abbildung mit Wert $0 \in W$.

- (8) Sei L ein Erweiterungskörper von K (siehe 4.4). Dann können wir die Multiplikation $L \times L \rightarrow L$ einschränken zu einer Abbildung $K \times L \rightarrow L$. Mit dieser Abbildung als Skalarmultiplikation und der Körperaddition wird L zu einem K -Vektorraum. Zum Beispiel können wir in dieser Weise $\mathbb{R}$ als $\mathbb{Q}$ -Vektorraum betrachten: Die Addition ist die gewöhnliche Addition reeller Zahlen, und für $a \in \mathbb{Q}, v \in \mathbb{R}$ ist die Skalarmultiplikation durch die Bildung des Produkts von $a \in \mathbb{Q} \subset \mathbb{R}$ und $v \in \mathbb{R}$ als reelle Zahlen gegeben.

Analog ist $\mathbb{C}$ ein $\mathbb{R}$ -Vektorraum.

So simpel diese Feststellung ist, so nützlich ist sie doch, später in der Algebra-Vorlesung, aber zum Beispiel auch in Ergänzung 6.57 und Ergänzung 7.64.

◇

ERGÄNZUNG 6.3. Da die Summe und das Produkt von stetigen Funktionen $\mathbb{R} \rightarrow \mathbb{R}$ wieder stetig ist, sieht man, dass die Menge aller stetigen Funktionen $\mathbb{R} \rightarrow \mathbb{R}$ (mit der Addition und Skalarmultiplikation für Abbildungen wie vorher) ein $\mathbb{R}$ -Vektorraum ist. Entsprechendes gilt für die Menge aller differenzierbaren Funktionen, und auch wenn man den Definitionsbereich durch eine geeignete Teilmenge von $\mathbb{R}$ ersetzt, beispielsweise durch ein offenes Intervall.

□ Ergänzung 6.3

BEMERKUNG 6.4 (Was ist ein Vektor). Man kann sich fragen, ob es nicht natürlicher wäre, die Definition eines Vektorraums dadurch zu ersetzen, dass man erst einmal definiert, was ein *Vektor* ist.

Es ist aber unmöglich, die gewünschten Eigenschaften von Vektoren in dieser Weise zu abstrahieren, weil sie eben nur im Zusammenspiel mit anderen Vektoren fassbar sind (der Addition von Vektoren und der Multiplikation mit Elementen des Grundkörpers). Deshalb stellt man in der linearen Algebra die Definition eines Vektorraums an den Beginn und definiert dann:

DEFINITION 6.5. Die Elemente eines Vektorraums heißen *Vektoren*. —

Die (etwas formalistische) Antwort auf die Frage *Was ist ein Vektor?* ist also: *Ein Vektor ist ein Element eines Vektorraums.*

”

In mathematics you don't understand things. You just get used to them.

John von Neumann,

Antwort auf die Bemerkung eines Physikers, der gesagt hatte
»I don't understand the method of characteristics.«

wiedergegeben in: G. Zukav, *The Dancing Wu Li Masters: An Overview of the New Physics*, Rider, London, 1990.

Man muss vielleicht schon etwas (mathematisch) abgebrüht sein, um sich mit dieser Definition anzufreunden, sie entspricht aber der Zielsetzung, jede Definition so anzulegen, dass sie möglichst vielseitig einsetzbar ist – in diesem Fall heißt das: über die geometrische Anschauung hinaus.

Aber wie stellt man die Verbindung zu der Anschauung her, dass ein Vektor als ein Pfeil interpretiert werden kann (der eine Länge und eine Richtung hat)? Selbst wenn wir uns auf den Vektorraum $\mathbb{R}^2$, also die reelle Zahlenebene, oder den $\mathbb{R}^3$, den wir anschaulich mit dem uns umgebenden Raum identifizieren können, einschränken – die Elemente sehen wir ja als Punkte in der Ebene beziehungsweise im Raum, und nicht als Pfeile. Diese Anschauung liegt schon der Wahl des Worts *Vektor*, von lateinisch *vectare*, »führen, tragen, bringen«, *vector*, »Fahrer, Träger«, zugrunde, das übrigens von Hamilton in Verbindung mit der Theorie der Quaternionen (siehe Abschnitt 2.3 und Ergänzung 4.II) als erstem benutzt wurde.

Was der Pfeil gut beschreibt, ist die *Operation*, v zu einem anderem Element zu addieren. Wenn wir den Startpunkt des Pfeils vom Ursprung nach v an den Punkt x legen, ist der Endpunkt des Pfeils bei $x + v$. In diesem Sinn ist dieser Pfeil eine Beschreibung der Abbildung $t_v: V \rightarrow V$, $x \mapsto x + v$.

Die Abbildung $V \rightarrow \text{Abb}(V, V)$, $v \mapsto t_v$ ist eine injektive Abbildung. Das bedeutet, dass der Vektor durch die Abbildung t_v eindeutig bestimmt ist. Man kann den Vektor v aus dem Pfeil »zurückgewinnen« – das ist ja klar, weil v gerade der Endpunkt des Pfeils ist, wenn wir ihn am Ursprung beginnen lassen.

Der Begriff des *affinen Raums unter einem Vektorraum* macht die Unterscheidung zwischen den Rollen der Elemente von V als »Punkte« bzw. als »Pfeile« noch deutlicher. Wir werden ihn in der Vorlesung aber allenfalls streifen. Siehe Kapitel II. ◇

LEMMA 6.6 (Rechenregeln). Seien K ein Körper und V ein K -Vektorraum. Für alle $a \in K, v \in V$ gilt:

- (1) $a \cdot v = 0$ genau dann, wenn $a = 0$ oder $v = 0$.
- (2) $-(av) = (-a)v = a(-v)$, und $(-1) \cdot v = -v$.
- (3) $a \sum v_i = \sum av_i$, $(\sum a_i)v = \sum a_i v$, $\sum a_i v_i + \sum b_i v_i = \sum (a_i + b_i) v_i$ (alle Summen seien endliche Summen).

BEWEIS. zu (1). Sei zuerst $a = 0$ und $v \in V$ beliebig. Dann gilt wegen der Distributivität und wegen $0 + 0 = 0$:

$$0v = (0 + 0)v = 0v + 0v,$$

und wenn wir auf beiden Seiten $-(0v)$ addieren (also das Inverse von $0v$ addieren), erhalten wir

$$0 = 0v.$$

Sei nun $v = 0$ und $a \in K$ beliebig. Dann gilt

$$a0 = a(0 + 0) = a0 + a0,$$

und wir folgern ähnlich wie vorher, dass $a0 = 0$.

Seien nun $a \in K^\times, v \in V \setminus \{0\}$. Um den Beweis von Teil (1) abzuschließen, müssen wir noch zeigen, dass $av \neq 0$. Wir wissen, dass

$$a^{-1}(av) = (a^{-1}a)v = 1 \cdot v = v \neq 0,$$

und deshalb muss $av \neq 0$ gelten (denn sonst wäre nach dem schon Gezeigten $a^{-1}(av) = a^{-1}0 = 0$).

zu (2). Wir zeigen zunächst $-(av) = (-a)v$, mit anderen Worten, dass $(-a)v$ das Negative von av ist. In der Tat gilt

$$av + (-a)v = (a - a)v = 0v = 0.$$

Die zweite Behauptung ergibt sich ganz ähnlich:

$$av + a(-v) = a(v - v) = a0 = 0.$$

Die letzte Behauptung folgt, wenn wir $a = 1$ setzen: $(-1) \cdot v = 1 \cdot (-v) = -v$.

zu (3). Diese Aussagen folgen mit vollständiger Induktion aus den Distributivgesetzen. $\square$

DEFINITION 6.7. Seien K ein Körper und V ein K -Vektorraum. Eine nichtleere Teilmenge $U \subseteq V$ heißt *Untervektorraum (oder Teilraum)* von V , wenn U abgeschlossen unter Addition und Skalarmultiplikation ist, das bedeutet: Sind $u, u' \in U$, so gilt auch $u + u' \in U$. Ist $u \in U$ und $a \in K$, so ist $au \in U$. $\dashv$

Wenn $U \subseteq V$ ein Untervektorraum ist, dann existiert ein Element $u \in U$, und aus den Untervektoreigenschaften folgt $0 = 0 \cdot u \in U$: die Menge U enthält den Nullvektor $0 \in V$ (alternativ kann man das auch sehen, indem man $0 = u + (-1) \cdot u \in U$ schreibt). Man könnte also in der Definition die Bedingung, dass U nicht leer ist, von vorneherein durch die Bedingung $0 \in U$ ersetzen.

Weil $-u = (-1) \cdot u$, impliziert die Abgeschlossenheit unter der Skalarmultiplikation, dass U zu jedem Element auch sein additives Inverses enthält. Deswegen ist U mit den Einschränkungen der Verknüpfungen von V selbst ein K -Vektorraum.

BEISPIEL 6.8. (1) In jedem Vektorraum V sind $\{0\}$ (oft geschrieben als 0 , und bezeichnet als der *triviale Untervektorraum*) und V Teilräume.

- (2) Für $V = K^n$ stimmt der hier definierte Begriff des Untervektorraums mit dem des Teilraums aus Definition 5.23 überein. Insbesondere erhalten wir so eine Vielzahl von Beispielen von Untervektorräumen: Jede Lösungsmenge eines homogenen Gleichungssystems mit n Unbestimmten ist ein Untervektorraum von K^n .
- (3) Sind $U_1, U_2 \subseteq V$ Untervektorräume, so ist $U_1 \cap U_2$ ein Untervektorraum von V . Allgemeiner gilt: Ist I irgendeine Menge und sind $U_i, i \in I$, Untervektorräume von V , so ist der Durchschnitt $\bigcap_{i \in I} U_i$ ein Untervektorraum von V .

◇

FRAGE 6.9. Sei K ein Körper. Was sind die Untervektorräume von K^2 ? Es ist nicht schwer zu sehen, dass neben 0 und K^2 auch alle Teilmengen der Form $\{ax; a \in K\}$ für ein festes Element $x \in K^2$ die Untervektorraumeigenschaften erfüllen. Es ist richtig, aber nicht offensichtlich, dass damit alle Untervektorräume von K^2 gefunden sind.

Versuchen Sie einmal, das zu zeigen! Spätestens am Ende dieses Kapitels, wenn wir den Begriff der Dimension behandelt haben, wird Ihnen der Beweis leicht fallen. □ Frage 6.9

Eine weitere Methode, um aus gegebenen Untervektorräumen einen weiteren zu konstruieren, ist die *Summe von Untervektorräumen*:

DEFINITION 6.10. Sei V ein K -Vektorraum. Sind $U, W \subseteq V$ Untervektorräume von V , so ist

$$U + W := \{u + w; u \in U, w \in W\}$$

ein Untervektorraum von V , den wir die *Summe* der Untervektorräume U und W nennen. $\dashv$

LEMMA 6.11. Sei V ein K -Vektorraum. Sind $U, W \subseteq V$ Untervektorräume von V , so ist $U + W$ der kleinste Untervektorraum von V , der U und W enthält.

Das bedeutet: $U + W$ ist ein Untervektorraum von V und wenn $V' \subseteq V$ irgendein Untervektorraum mit $U \subseteq V'$ und $W \subseteq V'$ ist, dann gilt $U + W \subseteq V'$.

BEWEIS. Es ist leicht nachzuprüfen, dass $U + W$ ein Untervektorraum von V ist.

Sei nun $V' \subseteq V$ ein Untervektorraum von V , der U und W enthält. Dann gilt für alle $u \in U$ und $w \in W$, dass $u \in V'$ und $w \in V'$, also $u + w \in V'$ (denn V' ist ein Untervektorraum). Also gilt $U + W \subseteq V'$. □

Zusammen mit der obigen Bemerkung, dass der Durchschnitt von Untervektorräumen wieder ein Untervektorraum ist, ergibt sich damit die Beschreibung von $U + W$ als dem Durchschnitt aller Untervektorräume V' von V , die U und W enthalten.

DEFINITION 6.12. Sei V ein K -Vektorraum.

- (1) Sind $U, W \subseteq V$ Untervektorräume mit $U \cap W = 0$, so schreiben wir statt $U + W$ auch $U \oplus W$ und sagen, U und W bilden (innerhalb V) eine *direkte Summe*.
- (2) Sei $U \subseteq V$ ein Untervektorraum von V . Wir nennen einen Untervektorraum W einen *Komplementärraum* von U (oder ein *Komplement* von U), wenn $U \oplus W = V$, mit anderen Worten: $U + W = V$ und $U \cap W = 0$.

 $\dashv$

Während die Summe von Untervektorräumen ein Verfahren ist, zu zwei beliebigen Untervektorräumen einen weiteren zu *konstruieren*, drückt das Symbol $\oplus$ an dieser Stelle eine *Eigenschaft* des Paares U, W von Untervektorräumen aus (nämlich, dass sie trivialen Durchschnitt haben). Siehe auch Abschnitt 6.6.

Wir werden unten sehen (Korollar 6.45), dass jeder Untervektorraum in einem Vektorraum einen Komplementärraum besitzt. In aller Regel ist dieser *nicht eindeutig* bestimmt, sondern es gibt viele (üblicherweise unendlich viele) verschiedene Komplementäräume zu einem gegebenen U .

BEISPIEL 6.13. Sei $K = \mathbb{R}$ und $V = \mathbb{R}^2$. Ist $U \subseteq V$ eine Gerade durch den Ursprung, so ist U ein Untervektorraum, wie wir bereits gesehen haben. Jede Ursprungsgerade $W \subseteq V$, die von U verschieden ist, ist dann ein Komplementärraum von U . Es ist klar, dass $U \cap W = \mathbf{o}$, da U und W verschieden sein sollen. Es ist dann nicht schwer zu sehen (aber auch nicht ganz offensichtlich), dass $U + W = V$ gilt. Machen Sie sich das geometrisch klar! $\diamond$

BEISPIEL 6.14. Sei $K = \mathbb{Q}$, $n \geq 1$. Seien

$$U = \left\{ (x_1, \dots, x_n) \in K^n; \sum_{i=1}^n x_i = \mathbf{o} \right\},$$

$$W = \{ (a, \dots, a) \in K^n; a \in K \}.$$

Dann sind U und W Untervektorräume von K^n und es gilt $U \oplus W = K^n$.

Wie ist es, wenn K ein anderer Körper ist? $\diamond$

LEMMA 6.15. Sei V ein K -Vektorraum, und seien U, W Untervektorräume von V . Dann sind äquivalent:

- (i) Es gilt $U \oplus W = V$.
- (ii) Jedes Element $v \in V$ lässt sich als Summe $v = u + w$ mit $u \in U$ und $w \in W$ schreiben und u und w sind dabei eindeutig bestimmt.

BEWEIS. (i) $\Rightarrow$ (ii). Weil $U + W = V$ gilt, ist klar, dass sich jedes Element $v \in V$ als Summe $v = u + w$ mit $u \in U$ und $w \in W$ schreiben lässt. Gilt $v = u + w = u' + w'$ mit $u, u' \in U$, $w, w' \in W$, so folgt $u - u' = w - w'$, und dieses Element liegt in U (als Differenz zweier Elemente in U), und auch in W . Weil $U \cap W = \mathbf{o}$, folgt $u - u' = \mathbf{o}$ und $w - w' = \mathbf{o}$, also $u = u'$, $w = w'$. Das beweist die Eindeutigkeit der Darstellung.

(ii) $\Rightarrow$ (i). Es ist klar, dass $U + W = V$ gilt. Sei nun $v \in U \cap W$. Wäre $v \neq \mathbf{o}$, so erhielten wir die beiden *verschiedenen* Darstellungen $v = \mathbf{o} + v = v + \mathbf{o}$ im Widerspruch zu (ii). Also muss $U \cap W = \mathbf{o}$ gelten. $\square$

6.2. Erzeugendensysteme

DEFINITION 6.16. Unter einer *Linearkombination* von Vektoren $v_1, \dots, v_n$ eines K -Vektorraums V verstehen wir einen Vektor der Form

$$a_1 v_1 + \dots + a_n v_n, \quad a_i \in K.$$

Wir können diese Sprechweise auch auf eine möglicherweise unendliche Familie $(v_i)_{i \in I}$ von Vektoren ausdehnen. Eine Linearkombination ist dann eine Summe der Form

$$\sum_{i \in I} a_i v_i, \quad a_i \in K, \quad \text{nur endlich viele } a_i \neq \mathbf{o}.$$

Eine Summe mit unendlich vielen Summanden können wir in einem Vektorraum nicht bilden, aber da alle bis auf endlich viele a_i verschwinden, handelt es sich hier nur um eine

endliche Summe. (In speziellen Konstellationen kann man natürlich auch dem Summen-symbol mit unendlich vielen Summanden $\neq 0$ einen Sinn geben, wie einer konvergenten Reihe in der Analysis. In der linearen Algebra kommen derartige Grenzwertbildungen aber nicht vor.) $\dashv$

Wir sagen, ein Vektor v lasse sich als Linearkombination der Familie $v_1, \dots, v_n$ darstellen, wenn $a_1, \dots, a_n \in K$ existieren mit $v = a_1 v_1 + \dots + a_n v_n$. Analog kann man davon sprechen, ob ein Vektor sich als Linearkombination einer (möglicherweise unendlichen) Familie $(v_i)_{i \in I}$ darstellen lässt.

Der Nullvektor lässt sich als Linearkombination jeder Familie $(v_i)_{i \in I}$ darstellen, indem man alle Koeffizienten $= 0$ wählt. (Mit der sinnvollen Konvention, dass die Summe über eine leere Indexmenge gleich Null sein soll, gilt das selbst für den Fall $I = \emptyset$.)

BEISPIEL 6.17. Der Vektor $(2, 4, 5)^t \in \mathbb{R}^3$ lässt sich nicht als Linearkombination der Vektoren $(1, 0, -1)^t$ und $(1, 1, 0)^t$ schreiben. (Warum nicht?)

Er lässt sich aber als Linearkombination von $(1, 0, 1)$ und $(0, 8, 6)$ schreiben. (Warum?) $\diamond$

DEFINITION 6.18. Seien K ein Körper und V ein K -Vektorraum. Eine Familie $B = (b_i)_{i \in I}$ von Elementen von V heißt *Basis* von V , falls jedes Element $v \in V$ in eindeutiger Weise (also für genau eine Wahl von Koeffizienten) als Linearkombination von Elementen aus B dargestellt werden kann. $\dashv$

BEISPIEL 6.19. Sei $n \geq 1$. Die Familie

$$e_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad e_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad e_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \quad \dots, \quad e_n = \begin{pmatrix} 0 \\ 0 \\ 0 \\ \vdots \\ 0 \\ 1 \end{pmatrix},$$

ist eine Basis von K^n und wird als die *Standardbasis* von K^n bezeichnet. Die Vektoren e_i nennen wir die *Standardbasisvektoren*. Manchmal spricht man auch von den *kanonischen Einheitsvektoren*.

Es ist klar, dass $e_1, \dots, e_n$ tatsächlich eine Basis von K^n bilden, denn ist $v = (x_i)_i \in K^n$ gegeben, so gilt $v = x_1 e_1 + \dots + x_n e_n$, und die x_i sind die einzigen Koeffizienten, für die die zugehörige Linearkombination den Vektor v liefert. $\diamond$

BEISPIEL 6.20. Ist V der Nullvektorraum über dem Körper K , so ist (per Definition) die leere Menge eine Basis von V , denn der Nullvektor kann als die »leere Linearkombination« geschrieben werden. $\diamond$

BEISPIEL 6.21. Sei K ein Körper, $V \neq 0$ ein K -Vektorraum und $v \in V$ irgendein Element $\neq 0$. Dann ist $U = \{av; a \in K\}$ ein Untervektorraum von V und jedes Element $\neq 0$ von U bildet eine Basis von U .

Wenn K mehr als zwei Elemente hat, dann hat U also mehrere verschiedene Basen. $\diamond$

Wie das Beispiel zeigt, ist eine Basis eines Teilraums fast nie eindeutig bestimmt.

Sei K ein Körper. Wir haben gesehen, dass die Lösungsmenge eines homogenen linearen Gleichungssystems in n Unbestimmten ein Teilraum von K^n ist. Ist das Gleichungssystem durch eine Matrix A in reduzierter Zeilenstufenform gegeben, so liefert die Methode, die wir kennengelernt haben, um die Lösungsmenge abzulesen, auch eine Basis dieses Vektorraums:

SATZ 6.22. Sei $A \in M_{m \times n}(K)$ die Koeffizientenmatrix eines homogenen linearen Gleichungssystems. Wir nehmen an, dass A reduzierte Zeilenstufenform hat. Seien r die Anzahl der Spalten von A , die eine führende Eins enthalten, und seien $j'_1 < \dots < j'_{n-r}$ die Indizes der anderen Spalten.

Sei für $\lambda = 1, \dots, n - r$ der Vektor $b_\lambda = (b_{\lambda j})_j \in K^n$ der eindeutig bestimmte Lösungsvektor mit

$$b_{\lambda j} = \begin{cases} 1 & j = j'_\lambda \\ 0 & j \in \{j'_1, j'_2, \dots, j'_{n-r}\} \setminus \{j'_\lambda\} \end{cases}$$

(die Werte $b_{\lambda j}$ für $j \notin \{j'_1, \dots, j'_{n-r}\}$ ergeben sich dann eindeutig aus der Bedingung, dass b_λ ein Lösungsvektor des gegebenen Gleichungssystems ist, siehe Satz 5.18).

Dann bilden die Vektoren $b_1, \dots, b_{n-r}$ eine Basis der Lösungsmenge des Gleichungssystems $Ax = 0$.

BEWEIS. Das folgt unmittelbar aus der Beschreibung der Lösungsmenge in Satz 5.18, vergleiche auch Beispiel 5.19. $\square$

Mit anderen Worten: Die Vektoren, mit denen wir nach der üblichen Methode die Lösungsmenge eines homogenen linearen Gleichungssystems aufschreiben (vergleiche das Ende von Beispiel 5.19 (2)) bilden eine Basis dieses Teilraums von K^n .

BEMERKUNG 6.23. Machen Sie sich an einem Beispiel klar, dass es zwar nach Satz 5.17 zu gegebener Lösungsmenge nur eine Matrix in reduzierter Zeilenstufenform gibt, die diese Lösungsmenge ergibt, und wir demnach aus dem vorherigen Satz eine Basis der Lösungsmenge erhalten, ohne weitere Wahlen zu treffen, dass aber die Lösungsmenge meistens noch viele andere Basen hat (die eben nicht auf diese Art und Weise entstehen). $\diamond$

BEMERKUNG 6.24. Wir können aus unseren Ergebnissen über lineare Gleichungssysteme auch ein Kriterium dafür ableiten, wann eine Familie $v_1, \dots, v_n \in K^n$ eine Basis bilden. Dass es so ist, heißt genau, dass die Gleichung

$$X_1 v_1 + X_2 v_2 + \dots + X_n v_n = b$$

für jedes $b \in K^n$ eine eindeutige Lösung $(x_i)_i \in K^n$ hat. Wenn wir die v_i als Spalten in eine Matrix A schreiben, so bedeutet das genau, dass für alle $b \in K^n$ das lineare Gleichungssystem $(A \mid b)$ eindeutig lösbar ist. Wir haben gesehen, dass das dazu äquivalent ist, dass die Matrix A invertierbar ist.

Es ist auch nicht schwierig, die Theorie der linearen Gleichungssysteme zu benutzen, um zu zeigen, dass eine Basis von K^n immer aus genau n Elementen bestehen muss. Weil wir dieses Ergebnis aber auch in Kürze aus der allgemeinen Theorie erhalten, verzichten wir hier auf eine genauere Analyse. Siehe auch Ergänzung 6.47. $\diamond$

DEFINITION 6.25. Sei V ein Vektorraum.

- (I) Seien $v_1, \dots, v_n \in V$. Der von $v_1, \dots, v_n$ aufgespannte Untervektorraum ist die Teilmenge aller derjenigen Vektoren von V , die sich als Linearkombination der v_i darstellen lassen. Er wird mit $\langle v_1, \dots, v_n \rangle$ bezeichnet:

$$\langle v_1, \dots, v_n \rangle = \left\{ \sum_{i=1}^n a_i v_i; a_i \in K \right\}.$$

- (2) Ist allgemeiner $M \subseteq V$ irgendeine (möglicherweise unendliche) Teilmenge, so ist der von M aufgespannte Untervektorraum von V die Teilmenge aller derjenigen Elemente aus V , die sich als Linearkombination von Elementen aus M darstellen lassen. Er wird mit $\langle M \rangle$ bezeichnet:

$$\langle M \rangle = \left\{ \sum_{i=1}^n a_i v_i; n \geq 1, a_i \in K, v_i \in M \right\}.$$

Man spricht auch vom von $v_1, \dots, v_n$ (bzw. von M) erzeugten Untervektorraum, von der linearen Hülle oder vom Spann der Vektoren $v_1, \dots, v_n$ bzw. der Menge M . $\dashv$

Wie in der Definition angedeutet, ist $\langle v_1, \dots, v_n \rangle$ und allgemeiner $\langle M \rangle$ ein Untervektorraum von V . Das prüft man leicht anhand der Definition nach.

SATZ 6.26 (Charakterisierung der linearen Hülle). Seien K ein Körper, V ein K -Vektorraum und $M \subseteq V$ eine Teilmenge. Dann gilt:

- (1) $\langle M \rangle$ ist der Durchschnitt aller Untervektorräume, die M enthalten.
- (2) $\langle M \rangle$ ist der kleinste Untervektorraum von V , der M enthält, das bedeutet: $\langle M \rangle$ ist ein Untervektorraum, und ist $U \subseteq V$ irgendein Untervektorraum mit $M \subseteq U$, so gilt $\langle M \rangle \subseteq U$.

BEWEIS. Zu (1). Sei U der Durchschnitt aller Untervektorräume, die M enthalten. Dies ist ein Untervektorraum von V (Beispiel 6.8). Er enthält M und damit auch alle Linearkombination von Elementen aus M , d.h. $\langle M \rangle \subseteq U$. Andererseits ist $\langle M \rangle$ ein Untervektorraum, der M enthält, also gilt auch $U \subseteq \langle M \rangle$.

Zu (2). Das folgt direkt aus Teil (1), denn der Durchschnitt aller Untervektorräume, die M enthalten, ist natürlich in jedem solchen Untervektorraum enthalten. $\square$

BEISPIEL 6.27. Wir betrachten den $\mathbb{R}$ -Vektorraum $V = \mathbb{R}^3$, den wir anschaulich mit dem uns umgebenden Raum identifizieren. Ist $v \in V$, $v \neq 0$, so ist $\langle v \rangle$ die eindeutig bestimmte Gerade durch den Ursprung, die v enthält. Ist $v' \in V$ ein Vektor, der nicht auf der Gerade $\langle v \rangle$ liegt, so ist $\langle v, v' \rangle$ die eindeutig bestimmte Ebene, die den Ursprung, v und v' enthält. $\diamond$

BEISPIEL 6.28. Sei K ein Körper, $n \geq 1$. Sei

$$U = \left\{ (x_1, \dots, x_n) \in K^n; \sum_{i=1}^n x_i = 0 \right\},$$

dies ist ein Untervektorraum von K^n . (Vergleiche Beispiel 6.14.) Dann gilt

$$U = \left\langle \begin{pmatrix} 1 \\ -1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ -1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 0 \\ 1 \\ -1 \end{pmatrix} \right\rangle$$

und diese Elemente bilden sogar eine Basis von U . $\diamond$

Um den Begriff der Basis besser zu verstehen, ist es nützlich, die zwei Aspekte der Definition – die Existenz einer Darstellung als Linearkombination einerseits, und die Eindeutigkeit andererseits – separat zu untersuchen. Der Begriff des Erzeugendensystems, den wir nun definieren, isoliert den Aspekt der Existenz solcher Darstellungen. Im nächsten Abschnitt behandeln wir dann mit dem Begriff der linearen Unabhängigkeit den Aspekt der Eindeutigkeit der Darstellung.

DEFINITION 6.29. Seien K ein Körper und V ein K -Vektorraum. Eine Teilmenge M von V heißt *Erzeugendensystem* von V , falls $\langle M \rangle = V$, d.h. wenn zu jedem $v \in V$ eine Zahl $N \geq 0$ und Elemente $m_1, \dots, m_N \in M$ und $a_1, \dots, a_N \in K$ existieren, so dass $x = \sum_{i=1}^N a_i m_i$. $\dashv$

DEFINITION 6.30. Ein K -Vektorraum V heißt *endlich erzeugt*, wenn ein Erzeugendensystem von V existiert, das nur endlich viele Elemente hat. $\dashv$

Die Vektorräume K^m , $m \in \mathbb{N}$ sind endlich erzeugt, denn es gibt ein Erzeugendensystem (sogar eine Basis) mit m Elementen. Seien $v_1, \dots, v_n \in K^m$ und sei $A \in M_{m \times n}(K)$ die Matrix mit den Spalten $v_1, \dots, v_n$. Dass $v_1, \dots, v_n$ ein Erzeugendensystem von K^m bilden, bedeutet genau, dass für alle $b \in K^m$ das lineare Gleichungssystem $(A \mid b)$ lösbar ist.

Die Begriffe *Erzeugendensystem* und *endlich erzeugt* können wir auch auf Untervektorräume eines Vektorraums anwenden (denn jeder Untervektorraum ist ja selbst auch ein K -Vektorraum).

BEISPIEL 6.31 (Beispiel eines nicht endlich erzeugten Vektorraums). Sei K ein Körper und sei V der Vektorraum aller Folgen $(a_i)_{i \in \mathbb{N}}$ mit $a_i \in K$, also $V = \prod_{i \in \mathbb{N}} K$, mit der komponentenweisen Addition und Skalarmultiplikation. Es ist leicht zu überprüfen, dass es sich tatsächlich um einen K -Vektorraum handelt. (Wir können diesen Vektorraum auch als den Vektorraum $\text{Abb}(\mathbb{N}, K)$ aus Beispiel 6.2 betrachten. Vergleiche auch Abschnitt 6.6.)

Die Teilmenge $U \subset V$, die aus allen Folgen mit nur endlich vielen Einträgen $\neq 0$ besteht, ist ein Untervektorraum. Auch das lässt sich leicht nachrechnen. (Mit der Notation von Abschnitt 6.6 ist $U = K^{(\mathbb{N})}$.)

Der Vektorraum U ist nicht endlich erzeugt. Denn sind $u_1, \dots, u_n \in U$, so haben nach Definition von U alle u_i nur endlich viele von 0 verschiedene Einträge. Sei $N \in \mathbb{N}$ der höchste Index, bei dem in irgendeinem der u_i ein Eintrag $\neq 0$ steht. Der Vektor, der an der Stelle $N+1$ eine 1 und sonst überall Nullen hat, ist ein Element von U , das nicht in $\langle u_1, \dots, u_n \rangle$ liegt.

Es ist richtig, aber weniger offensichtlich, dass auch der Vektorraum V nicht endlich erzeugt ist. Siehe Satz 6.44. $\diamond$

BEISPIEL 6.32. Sei K ein unendlicher Körper. Der K -Vektorraum U aller Polynomfunktionen $K \rightarrow K$ (Abschnitt 4.3 und Beispiel 6.2) ist nicht endlich erzeugt. Denn für jede Polynomfunktion $f: K \rightarrow K$, die nicht die Nullfunktion ist, ist die Darstellung $f(x) = \sum_{i=0}^n a_i x^i$ mit $a_n \neq 0$ eindeutig bestimmt. Die Zahl n bezeichnet man dann als den Grad von f . Sind $f_1, \dots, f_r$ Polynomfunktionen vom Grad $\leq n$, so hat auch jede Linearkombination der f_i Grad $\leq n$. Da Polynomfunktionen von beliebig großem Grad existieren, folgt aus dieser Überlegung, dass U nicht endlich erzeugt ist.

(Betrachtet man die Koeffizienten $a_0, \dots, a_n$ als eine unendliche Folge von Elementen in K , deren weitere Glieder alle $= 0$ sind, so kann man dieses Beispiel mit dem vorherigen identifizieren.)

Ist K endlich, so gibt es nur endlich viele Polynomfunktionen $K \rightarrow K$, und insbesondere ist der Vektorraum aller dieser Polynomfunktionen eine endliche Menge und erst recht endlich erzeugt. $\diamond$

6.3. Lineare Unabhängigkeit

Gegeben eine Familie $(v_i)_{i \in I}$ von Vektoren in einem Vektorraum V , dann bezeichnen wir $\sum_{i \in I} 0v_i$ (also die Linearkombination, in der alle Koeffizienten gleich Null sind) als die *triviale Linearkombination* der gegebenen Familie von Vektoren. Ihr Wert ist der Nullvektor. Diese Linearkombination ist natürlich völlig uninteressant, und sie hat nur deshalb einen eigenen Namen, damit man sie – zum Beispiel in der folgenden Definition – von der Betrachtung ausschließen kann.

DEFINITION 6.33. V ein K -Vektorraum. Eine Familie $(v_i)_{i \in I}$ von Vektoren $v_i \in V$ heißt *linear unabhängig* (l.u.), falls für jede endliche Teilmenge $J \subseteq I$ gilt: Sind $a_i \in K, i \in J$ mit

$$\sum_{i \in J} a_i v_i = 0,$$

so gilt für alle $i \in J: a_i = 0$.

Eine Familie von Vektoren aus V , die nicht linear unabhängig ist, heißt *linear abhängig* (l.a.). $\dashv$

In ähnlicher Weise definiert man die Begriffe *linear unabhängig* und *linear abhängig* für Teilmengen eines Vektorraums. Man beachte, dass in einer Familie von Vektoren derselbe Vektor mehrfach auftreten kann, jedoch nicht in einer Menge. Ist etwa $v \in V, v \neq 0$, so ist die Menge $\{v, v, \dots, v\}$ linear unabhängig, weil sie gleich der Menge $\{v\}$ ist. Die Familie $(v_i)_{i=1, \dots, n}$ mit $v_i := v$ ist jedoch für $n > 1$ linear abhängig.

Elemente $v_1, \dots, v_n \in K^m$ sind genau dann linear unabhängig, wenn das homogene lineare Gleichungssystem $Ax = 0$, in dem A die Matrix mit den Spalten $v_1, \dots, v_n$ ist, nur die triviale Lösung besitzt, oder anders ausgedrückt: eindeutig lösbar ist.

BEISPIEL 6.34. (1) Ist V ein Vektorraum, so ist jede Basis von V linear unabhängig, denn dann lässt sich jedes Element in eindeutiger Weise als Linearkombination schreiben, insbesondere der Nullvektor.

Zum Beispiel ist die Familie $e_1, \dots, e_n$ der Standardbasisvektoren in K^n linear unabhängig.

(2) Ist $(v_i)_{i \in I}$ eine linear unabhängige Familie von Elementen eines Vektorraums V und $J \subseteq I$ eine Teilmenge, so ist auch $(v_i)_{i \in J}$ linear unabhängig.

(3) Die leere Menge von Vektoren ist aus trivialen Gründen linear unabhängig.

(4) Jede Familie, in der der Nullvektor vorkommt, ist linear abhängig.

◇

LEMMA 6.35. Seien K ein Körper, V ein K -Vektorraum, $v_1, \dots, v_n \in V, n \geq 1$. Dann sind äquivalent:

- (i) Die Familie $v_1, \dots, v_n$ ist linear unabhängig
- (ii) Für alle $v \in \langle v_1, \dots, v_n \rangle$, existiert eine eindeutig bestimmte Darstellung von v als Linearkombination von $v_1, \dots, v_n$.
- (iii) Für alle i gilt: $v_i \notin \langle v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n \rangle$.
- (iv) Für alle i gilt: $\langle v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n \rangle \subsetneq \langle v_1, \dots, v_n \rangle$.
- (v) Es sind $v_1, \dots, v_{n-1}$ linear unabhängig und $v_n \notin \langle v_1, \dots, v_{n-1} \rangle$.

BEWEIS. Wir führen Induktion nach n . Für $n = 1$ ist die Sache klar: alle Bedingungen sind äquivalent zu $v_1 \neq 0$. Wir dürfen daher im folgenden annehmen, dass die 5 Aussagen für Familien von $n - 1$ Vektoren äquivalent sind.

Wir beweisen nun die folgenden Implikationen; daraus folgt die gewünschte Äquivalenz.

$$(i) \Leftrightarrow (ii), \quad (iii) \Leftrightarrow (iv), \quad (i) \Rightarrow (iii) \Rightarrow (v) \Rightarrow (i).$$

(i) $\Leftrightarrow$ (ii). Wenn es ein v gäbe, das zwei verschiedene Darstellungen

$$v = a_1 v_1 + \dots + a_n v_n = b_1 v_1 + \dots + b_n v_n$$

hätte, so wäre

$$(a_1 - b_1) v_1 + \dots + (a_n - b_n) v_n = 0$$

eine nichttriviale Linearkombination, die den Nullvektor darstellt, im Widerspruch zu (i).

(iii) $\Rightarrow$ (iv) ist klar und (iii) $\Leftarrow$ (iv) ist einfach.

(i) $\Rightarrow$ (iii). Beweis durch Widerspruch: Angenommen, es wäre $v_i \in \langle v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n \rangle$. Das hieße, dass sich v_i als Linearkombination von $v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n$ schreiben ließe:

$$v_i = a_1 v_1 + \dots + a_{i-1} v_{i-1} + a_{i+1} v_{i+1} + \dots + a_n v_n,$$

was wir umschreiben können zu

$$a_1 v_1 + \dots + a_{i-1} v_{i-1} + (-1) v_i + a_{i+1} v_{i+1} + \dots + a_n v_n = 0.$$

Dies ist eine Linearkombination, die den Nullvektor darstellt. Da der Koeffizient von v_i aber $\neq 0$ ist, können dann die Vektoren $v_1, \dots, v_n$ nicht linear unabhängig sein.

(iii) $\Rightarrow$ (v). Wenn (iii) gilt, so gilt erst recht $v_i \notin \langle v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_{n-1} \rangle$ für $i = 1, \dots, n-1$, also sind nach Induktionsvoraussetzung die Vektoren $v_1, \dots, v_{n-1}$ linear unabhängig. Damit ist klar, dass aus (iii) Bedingung (v) folgt.

(v) $\Rightarrow$ (i). Sei

$$a_1 v_1 + \dots + a_n v_n = 0$$

eine Linearkombination, die den Nullvektor darstellt. Wegen $v_n \notin \langle v_1, \dots, v_{n-1} \rangle$ muss dann $a_n = 0$ sein. Weil $v_1, \dots, v_{n-1}$ linear unabhängig sind, folgt daraus aber, dass alle a_i verschwinden. Also ist die Familie $v_1, \dots, v_n$ linear unabhängig. $\square$

6.4. Basen, Basissätze

SATZ 6.36. Seien K ein Körper und V ein K -Vektorraum. Sei B eine Teilmenge von V . Dann sind äquivalent:

- (i) B ist eine Basis von V .
- (ii) B ist ein linear unabhängiges Erzeugendensystem von V .
- (iii) B ist ein minimales Erzeugendensystem von V (d.h. B ist ein Erzeugendensystem von V , aber keine echte Teilmenge $B' \subsetneq B$ ist ein Erzeugendensystem von V).
- (iv) B ist eine maximale linear unabhängige Teilmenge in V (d.h. B ist linear unabhängig und für jedes $v \in V \setminus B$ ist $B \cup \{v\}$ linear abhängig).

BEWEIS. Einige der Argumente hier werden Ihnen aus dem Beweis von Lemma 6.35 bekannt vorkommen, aber weil diese Techniken so wichtig sind, und wir in dem Lemma erstmal nur den Fall einer endlichen Familie von Vektoren behandelt haben, wiederholen wir sie hier noch einmal.

(i) $\Rightarrow$ (ii). Sei B eine Basis von V . Nach Definition des Begriffs *Basis* ist klar, dass B ein Erzeugendensystem ist. Da der Nullvektor (wie jeder andere Vektor) in eindeutiger Weise als Linearkombination der Elemente von B dargestellt werden kann, folgt auch, dass B linear unabhängig ist.

(ii) $\Rightarrow$ (iii). Sei B ein linear unabhängiges Erzeugendensystem von V . Angenommen, es gäbe $b \in B$, so dass auch $B \setminus \{b\}$ ein Erzeugendensystem ist. Dann existieren $b_1, \dots, b_r \in B \setminus \{b\}$ und $a_1, \dots, a_r \in K$ mit

$$b = a_1 b_1 + \dots + a_r b_r,$$

also ist $a_1 b_1 + \dots + a_r b_r - b = 0$ eine Linearkombination, die den Nullvektor darstellt. Da $b_i \neq b$ für alle i gilt, ist das ein Widerspruch dazu, dass B linear unabhängig ist.

(iii) $\Rightarrow$ (iv). Sei B ein minimales Erzeugendensystem von V . Wir zeigen zuerst, dass B linear unabhängig ist. Sonst gäbe es $b_1, \dots, b_r \in B$ und eine Linearkombination $a_1 b_1 + \dots + a_r b_r = 0$ mit $a_1 \neq 0$, und dann wäre

$$b_1 = -\frac{1}{a_1} (a_2 b_2 + \dots + a_r b_r),$$

und daher auch $B \setminus \{b_1\}$ ein Erzeugendensystem von V , im Widerspruch zur Minimalität von B . Da B ein Erzeugendensystem ist, ist andererseits auch klar, dass für $v \in V \setminus B$ die Vereinigung $B \cup \{v\}$ nicht linear unabhängig sein kann.

(iv) $\Rightarrow$ (i). Sei B eine maximale linear unabhängige Teilmenge von V . Wegen der linearen Unabhängigkeit kann jeder Vektor $v \in V$ in höchstens einer Weise als Linearkombination von Elementen aus B dargestellt werden (denn gäbe es zwei verschiedene Darstellungen, so wäre die Differenz eine nicht-triviale Linearkombination, die den Nullvektor darstellt). Es ist daher nur noch zu zeigen, dass B ein Erzeugendensystem ist. Ist $v \in B$, so lässt sich v trivialerweise als Linearkombination von Elementen in B darstellen. Ist $v \in V \setminus B$, so ist $B \cup \{v\}$ linear abhängig, es gibt also eine *nicht-triviale* Linearkombination

$$av + a_1b_1 + \cdots + a_rb_r = 0$$

mit $a, a_i \in K, b_i \in B$. Weil B linear unabhängig ist, kann hier nicht $a = 0$ sein. Daher können wir durch a teilen und sehen dann leicht, dass v eine Linearkombination von Elementen aus B ist. $\square$

In den folgenden Ergebnissen dieses Abschnitts sind oftmals gewisse Endlichkeitsvoraussetzungen (wie »endlich erzeugt«) in eckigen Klammern angegeben. Das soll bedeuten, dass die Sätze zwar auch im allgemeinen Fall richtig sind, wir sie aber nur unter der genannten Voraussetzung beweisen. Siehe Ergänzung 6.48 für Hinweise zu den Beweisen im allgemeinen Fall.

SATZ 6.37. *Jeder [endlich erzeugte] K -Vektorraum besitzt eine Basis.*

BEWEIS. Wir geben hier den Beweis in dem Fall, dass V ein endliches Erzeugendensystem E besitzt. Ist das Erzeugendensystem E minimal, so haben wir schon eine Basis gefunden. Andernfalls ist die Menge E linear abhängig, und es existiert ein Element $v \in E$, das in der linearen Hülle der anderen Elemente aus E liegt. Daher ist auch $E \setminus \{v\}$ ein Erzeugendensystem von V . Indem wir diese Argumentation fortsetzen und gegebenenfalls weitere Elemente von E herausnehmen, finden wir in endlich vielen Schritten ein minimales Erzeugendensystem, also eine Basis. $\square$

Siehe auch Satz 6.40 für eine präzisere Form des Satzes. Für den Fall des Nullvektorraums ist die leere Menge eine (und die einzige) Basis.

SATZ 6.38 (Basisaustauschsatz). *Seien K ein Körper, V ein K -Vektorraum, $v_1, \dots, v_n \in V$ eine Basis, $w_1, \dots, w_i \in V$ eine linear unabhängige Familie. Dann existiert eine Teilmenge $I \subseteq \{1, \dots, n\}$, $\#I = i$, so dass die n Elemente $w_1, \dots, w_i, v_j, j \notin I$, eine Basis von V bilden.*

BEWEIS. Wir können das Ergebnis des Satzes auch so formulieren: Nach Umnummerieren der v_j (falls erforderlich), ist $w_1, \dots, w_i, v_{i+1}, \dots, v_n$ eine Basis von V .

Wir beweisen diese Behauptung durch vollständige Induktion nach i . *Induktionsanfang:* $i = 0$. In diesem Fall ist nichts zu beweisen: Die Behauptung fällt nämlich mit der Annahme zusammen, dass $v_1, \dots, v_n$ eine Basis ist. (Wenn Ihnen mit der Optimierung, für den Induktionsanfang $i = 0$ zu wählen, unwohl ist, dann überlegen Sie sich den Fall $i = 1$ zusätzlich direkt. Dann können Sie die Induktion von dort fortsetzen.)

Induktionsschritt: $i - 1 \rightarrow i$. Sei nun $i \geq 1$ und sei der Satz für linear unabhängige Familien mit $i - 1$ Elementen bereits bewiesen. Sei eine linear unabhängige Familie $w_1, \dots, w_i$ gegeben. Dann ist erst recht die Familie $w_1, \dots, w_{i-1}$ linear unabhängig, nach Induktionsvoraussetzung können wir die v_j also so umnummerieren, dass $w_1, \dots, w_{i-1}, v_i, \dots, v_n$ eine Basis bildet. Dann können wir w_i als Linearkombination

$$w_i = a_1w_1 + \cdots + a_{i-1}w_{i-1} + a_iv_i + \cdots + a_nv_n$$

schreiben. Da $w_1, \dots, w_i$ als linear unabhängig vorausgesetzt wurde, kann es keine Linearkombination von $w_1, \dots, w_{i-1}$ geben, wie w_i darstellt, es muss also mindestens einen Koeffizienten a_j mit $j \geq i$ geben, der nicht gleich 0 ist (wir sehen an dieser Stelle auch, dass automatisch $i \leq n$ gelten muss). Nach Umnummerieren können wir annehmen, dass $a_i \neq 0$ gilt. Wir können dann die obige Gleichheit umschreiben als

$$v_i = -\frac{1}{a_i}(a_1 w_1 + \dots + a_{i-1} w_{i-1} - w_i + a_{i+1} v_{i+1} + \dots + a_n v_n).$$

Es ist dann leicht zu sehen, dass $w_1, \dots, w_i, v_{i+1}, \dots, v_n$ wieder ein linear unabhängiges Erzeugendensystem von V ist, also eine Basis. $\square$

Direkt aus dem Satz ergibt sich (weil eine Teilmenge von $\{1, \dots, n\}$ höchstens n Elemente haben kann):

KOROLLAR 6.39. *Seien K ein Körper, V ein K -Vektorraum, $v_1, \dots, v_n \in V$ eine Basis, $w_1, \dots, w_i \in V$ eine linear unabhängige Familie. Dann gilt $i \leq n$.*

SATZ 6.40 (Basisergänzungssatz). *Seien K ein Körper und V ein K -Vektorraum. Sei M eine linear unabhängige Teilmenge von V , und sei $E \subseteq V$ ein [endliches] Erzeugendensystem von V , das M enthält. Dann existiert eine Basis B von V mit $M \subseteq B \subseteq E$.*

Insbesondere: Jede linear unabhängige Teilmenge von V lässt sich zu einer Basis ergänzen, [sofern V endlich erzeugt ist.]

BEWEIS. Wir betrachten den Fall, dass E endlich ist. Wir können dann E zu einem minimalen Erzeugendensystem, also einer Basis, verkleinern. Mit dem Basisaustauschsatz können wir gegebenenfalls die Elemente von M , die bei diesem Prozess herausgefallen sind, wieder hinzunehmen (und dafür andere Elemente fortlassen).

Für den zweiten Teil ergänzen wir die gegebene linear unabhängige Teilmenge zunächst in beliebiger Weise zu einem endlichen Erzeugendensystem (zum Beispiel, indem wir irgendein endliches Erzeugendensystem hinzufügen). $\square$

THEOREM 6.41. *Sei K ein Körper und sei V ein [endlich erzeugter] K -Vektorraum. Je zwei Basen von V besitzen dieselbe Mächtigkeit (d.h. im endlich erzeugten Fall einfach, dass je zwei Basen gleich viele Elemente haben; im allgemeinen Fall kann man sagen, dass es zwischen je zwei Basen eine Bijektion gibt).*

Diese Zahl wird als Dimension $\dim V$ von V bezeichnet (zumindest sofern sie endlich ist, sonst sagen wir einfach, die Dimension von V sei unendlich; präziser kann man die Dimension als Kardinalzahl verstehen).

BEWEIS. Wir beweisen den Satz in dem Fall, dass V endlich erzeugt ist. In diesem Fall folgt die Behauptung unmittelbar aus Korollar 6.39. $\square$

Nach dem Theorem hat jeder endlich erzeugte Vektorraum endliche Dimension. Wir sprechen daher auch von *endlich-dimensionalen Vektorräumen*.

BEISPIEL 6.42. (1) Es gilt $\dim(K^n) = n$, denn die Standardbasisvektoren $e_1, \dots, e_n$ bilden eine Basis. (Ist $n = 0$, so ist $K^n = 0$ und die leere Menge eine Basis, also ist auch $\dim K^0 = 0$.)

(2) Es gilt $\dim(M_{m \times n}(K)) = mn$. Die mn Matrizen, in denen ein einziger Eintrag 1 und alle anderen Einträge 0 sind, bilden eine Basis.

(3) Ist $A \in M_{m \times n}(K)$ die Koeffizientenmatrix eines homogenen linearen Gleichungssystems, so ist die Zahl r der Spalten mit führenden Einsen in jeder Matrix in Zeilenstufenform, die aus A durch elementare Zeilenumformungen entsteht, gleich $n - \dim \operatorname{Ker} A$. Siehe Satz 6.22.

- (4) Der $\mathbb{R}$ -Vektorraum $\mathbb{C}$ hat Dimension $\dim_{\mathbb{R}} \mathbb{C} = 2$. (Natürlich können wir $\mathbb{C}$ auch als Vektorraum über sich selbst auffassen und es ist $\dim_{\mathbb{C}} \mathbb{C} = 1$.)

◇

Der Dimensionsbegriff gibt uns die Möglichkeit, die »Größe« eines Vektorraums (oder Untervektorraums in einem Vektorraum) zu messen/zu quantifizieren. Das ist der erste wichtige Schritt zur Beantwortung von Frage 5.27 (2).

THEOREM 6.43. *Seien K ein Körper und V ein n -dimensionaler Vektorraum, $n \in \mathbb{N}$.*

- (1) *Jedes linear unabhängige System von n Vektoren in V ist eine Basis.*
- (2) *Jedes Erzeugendensystem von V , das aus n Elementen besteht, ist eine Basis von V .*

BEWEIS. zu (1). Sei ein linear unabhängiges System gegeben, das aus n Vektoren besteht. Nach dem Basisergänzungssatz können wir dieses System zu einer Basis ergänzen. Da alle Basen aus genau n Elementen bestehen, kann eine echte Ergänzung in diesem Fall aber gar nicht erforderlich/möglich sein. Das gegebene System muss selbst schon eine Basis sein.

zu (2). Ähnlich wie Teil (1). □

Mit diesem Ergebnis kann man sich auch leicht überlegen, dass man im Basisergänzungssatz die Voraussetzung, dass E endlich sei, ersetzen kann durch die Bedingung, dass V endlich erzeugt ist, also überhaupt irgendein endliches Erzeugendensystem besitzt. (Mit dem Lemma von Zorn kann man natürlich, wie durch die eckigen Klammern angedeutet, die Endlichkeitsvoraussetzung ohnehin vollständig fallenlassen.)

SATZ 6.44. *Sei V ein endlich erzeugter K -Vektorraum und $U \subseteq V$ ein Untervektorraum.*

- (1) *Der Vektorraum U ist endlich erzeugt und $\dim U \leq \dim V$.*
- (2) *Gilt $\dim U = \dim V$, so folgt $U = V$.*

BEWEIS. Wäre U nicht endlich erzeugt, so müsste U zu jeder natürlichen Zahl r eine linear unabhängige Teilmenge mit r Elementen enthalten (denn sonst könnten wir eine endliche maximale linear unabhängige Teilmenge, also eine endliche Basis, finden). Es kann aber nicht sein, dass U eine linear unabhängige Teilmenge mit mehr als $\dim V$ Elementen enthält. Also ist U endlich erzeugt. Jede Basis von U ist linear unabhängig, auch als Teilmenge von V betrachtet, und wir können sie zu einer Basis von V ergänzen. Also muss $\dim(U) \leq \dim(V)$ gelten.

Gilt $\dim U = \dim V$, so folgt aus Theorem 6.43, dass jede Basis von U auch eine Basis von V ist. Das impliziert $U = V$. □

Wir können nun auch beweisen, dass jeder Untervektorraum in einem Vektorraum ein Komplement (Def. 6.12) besitzt.

KOROLLAR 6.45. *Sei V ein [endlich erzeugter] K -Vektorraum und sei $U \subseteq V$ ein Untervektorraum. Dann besitzt U einen Komplementärraum.*

BEWEIS. Wir wählen eine Basis $b_1, \dots, b_m$ von U . Diese Vektoren sind linear unabhängig, egal, ob wir sie als Elemente von U oder von V betrachten. Wir können daher die Familie der b_i nach dem Basisergänzungssatz zu einer Basis $b_1, \dots, b_n$ von V ergänzen. Dann ist $W := \langle b_{m+1}, \dots, b_n \rangle$ ein Komplement von U . □

Wir sehen am Beweis auch, dass für einen Untervektorraum $U \subseteq V$ in einem endlich erzeugten Vektorraum mit Komplement W gilt, dass $\dim U + \dim W = \dim V$. Diese Formel verallgemeinert der folgende Satz:

SATZ 6.46 (Dimensionsformel für den Durchschnitt von zwei Untervektorräumen). *Seien V ein [endlich erzeugter] K -Vektorraum und seien $U, W \subseteq V$ Untervektorräume. Dann gilt*

$$\dim(U + W) + \dim(U \cap W) = \dim(U) + \dim(W).$$

BEWEIS. Wir geben den Beweis im Fall dass V und damit auch alle Untervektorräume von V endlich erzeugt sind und demnach endliche Dimension haben.

Sei $v_1, \dots, v_r$ eine Basis von $U \cap W$. Wir ergänzen diese einerseits durch Vektoren $u_1, \dots, u_s$ zu einer Basis $v_1, \dots, v_r, u_1, \dots, u_s$ von U , andererseits durch Vektoren $w_1, \dots, w_t$ zu einer Basis $v_1, \dots, v_r, w_1, \dots, w_t$ von W .

Es ist klar, dass es dann genügt, die folgende Behauptung zu zeigen:

Behauptung. $v_1, \dots, v_r, u_1, \dots, u_s, w_1, \dots, w_t$ ist eine Basis von $U + W$.

Begründung. Offenbar handelt es sich um ein Erzeugendensystem von $U + W$, weil die gegebene Familie sowohl ein Erzeugendensystem von U als auch eines von W enthält. Es bleibt also nur noch die lineare Unabhängigkeit zu zeigen. Sei dazu

$$a_1 v_1 + \dots + a_r v_r + b_1 u_1 + \dots + b_s u_s + c_1 w_1 + \dots + c_t w_t = 0$$

eine Linearkombination, die den Nullvektor darstellt. Dann gilt $-\sum c_i w_i = \sum a_i v_i + \sum b_i u_i \in U \cap W$. Ein Vektor aus $U \cap W$ lässt sich (in eindeutiger Weise) als Linearkombination der Vektoren $v_1, \dots, v_r$ darstellen, und da $v_1, \dots, v_r, u_1, \dots, u_s$ eine Basis von U bilden, auch *in eindeutiger Weise* als Linearkombination von $v_1, \dots, v_r, u_1, \dots, u_s$. Es folgt, dass $b_1 = \dots = b_s = 0$. Damit erhalten wir dann

$$a_1 v_1 + \dots + a_r v_r + c_1 w_1 + \dots + c_t w_t = 0,$$

und da $v_1, \dots, v_r, w_1, \dots, w_t$ (als Basis von W) linear unabhängig ist, folgt, dass alle Koeffizienten verschwinden. $\square$

Um sich die Formel zu merken, vergleichen Sie den Fall endlicher Mengen: Sind U und W Teilmengen einer endlichen Menge X , so gilt für die Anzahlen der Elemente $\#(U \cup W) - \#(U \cap W) = \#U + \#W$. Das natürliche Analogon im Vektorraumfall ist es, die Vereinigung durch die Summe von Untervektorräumen und die Anzahl der Elemente durch die Dimension zu ersetzen. Man erhält dann genau die obige Formel. Eine ähnliche Analogie kann man für Satz 6.44 herstellen: Teilmengen endlicher Mengen sind endlich, und hat eine Teilmenge ebenso viele Elemente wie die umgebende (endliche!) Menge, so sind die beiden notwendigerweise gleich.

ERGÄNZUNG 6.47. Es ist möglich, den Begriff der Dimension noch direkter auf der Theorie der linearen Gleichungssysteme aufzubauen. Dazu würde man zeigen, dass für jeden Vektorraum V mit einem endlichen Erzeugendensystem E und jede linear unabhängige Teilmenge $L \subseteq V$ die Abschätzung $\#L \leq \#E$ gilt.

Um das zu zeigen, schreiben wir die Elemente von E als $e_1, \dots, e_m$ und die von L als $l_1, \dots, l_n$. Da E ein Erzeugendensystem ist, gibt es $a_{ij} \in K$ mit

$$l_j = a_{1j} e_1 + \dots + a_{mj} e_m, \quad j = 1, \dots, n.$$

Sei $A = (a_{ij})_{ij} \in M_{m \times n}(K)$. Für jede Lösung $(x_1, \dots, x_n)^t$ des linearen Gleichungssystems $Ax = 0$ gilt dann $\sum_j a_{ij} x_j = 0$, also

$$\sum_{j=1}^n x_j l_j = \sum_{i,j} a_{ij} x_j e_i = 0,$$

also $x_1 = \dots = x_n = 0$ wegen der linearen Unabhängigkeit von L . Das homogene Gleichungssystem $Ax = 0$ hat also nur die triviale Lösung, die Zeilenstufenform von A hat in jeder

Spalte eine Stufe. Daher muss A mindestens so viele Zeilen wie Spalten haben, es gilt somit $m \geq n$. Siehe auch [So] I.7. □ Ergänzung 6.47

ERGÄNZUNG 6.48 (Existenz von Basen, allgemeiner Fall). Wir haben in Satz 6.37 behauptet, aber nicht bewiesen, dass jeder Vektorraum eine Basis besitzt. Es ist klar, dass der Beweis, den wir für den endlich erzeugten Fall gegeben haben, nicht ohne weiteres auf den allgemeinen Fall übertragen werden kann. Es ist zwar klar, dass jeder Vektorraum V ein Erzeugendensystem besitzt (zum Beispiel die Menge V selbst). Aber wenn dieses unendlich ist, kann man nicht erwarten, es durch Entfernen endlich vieler Vektoren zu einem minimalen Erzeugendensystem zu machen.

Im allgemeinen Fall ist es daher günstiger zu zeigen, dass in V eine maximale linear unabhängige Teilmenge existiert. Wie wir gesehen haben, ist dies eine Basis von V .

Dafür verwenden wir das Lemma von Zorn, siehe Abschnitt B.1. Wir betrachten die Menge $\mathcal{U}$ aller linear unabhängigen Teilmengen von V mit der Inklusion als »partieller Ordnung«. Das Lemma von Zorn liefert uns die Existenz eines maximalen Elements in $\mathcal{U}$ bezüglich $\subseteq$, also einer Basis von V , wenn wir die folgende Behauptung zeigen können:

Behauptung. Sei $\mathcal{U}' \subseteq \mathcal{U}$ eine Teilmenge von $\mathcal{U}$ mit der Eigenschaft, dass für alle $L_1, L_2 \in \mathcal{U}'$ gilt, dass $L_1 \subseteq L_2$ oder $L_2 \subseteq L_1$. Dann existiert ein Element $L \in \mathcal{U}$, so dass $L' \subseteq L$ für alle $L' \in \mathcal{U}'$.

(Die Bedingung an $\mathcal{U}'$ beschreibt man, indem man sagt, $\mathcal{U}'$ sei bezüglich $\subseteq$ *total* geordnet. Das Element L von $\mathcal{U}$ nennt man eine *obere Schranke* von $\mathcal{U}'$.)

Begründung. Sei also $\mathcal{U}'$ wie in der Behauptung gegeben. Wir definieren L als die Vereinigung aller Teilmengen von $\mathcal{U}'$. Es ist klar, dass $L' \subseteq L$ für alle $L' \in \mathcal{U}'$ gilt. Aber wir müssen noch zeigen, dass überhaupt $L \in \mathcal{U}$ gilt, dass also L eine linear unabhängige Teilmenge von V ist.

Natürlich ist im allgemeinen die Vereinigung von linear unabhängigen Teilmengen nicht linear unabhängig; aber hier sind wir ja in einer speziellen Situation.

Wir beginnen mit dem folgenden Lemma:

LEMMA 6.49. *Seien V ein K -Vektorraum und $L \subseteq V$ eine Teilmenge. Die Teilmenge L ist genau dann linear unabhängig, wenn jede endliche Teilmenge von L linear unabhängig ist.*

BEWEIS. Wenn L linear unabhängig ist, dann ist erst recht jede Teilmenge von L linear unabhängig. Wenn L linear abhängig ist, dann gibt es eine nicht-triviale Linearkombination von Elementen aus L , die den Nullvektor darstellt, und in jeder Linearkombination treten nur endlich viele Vektoren auf. Daher existiert dann eine endliche Teilmenge von L , die linear abhängig ist. □

Nun zurück zum Beweis der obigen Behauptung. Es genügt also zu zeigen, dass jede endliche Teilmenge $\{v_1, \dots, v_r\}$ von L linear unabhängig ist. Da L die Vereinigung der Elemente von $\mathcal{U}'$ ist, existieren $L_1, \dots, L_r \in \mathcal{U}'$ mit $v_i \in L_i$. Da je zwei Elemente von $\mathcal{U}'$ ineinander enthalten sind, können wir, wenn wir die v_i und L_i umnummerieren, falls erforderlich, annehmen, dass

$$L_1 \subseteq L_2 \subseteq \dots \subseteq L_r.$$

Dann liegen aber $v_1, \dots, v_r$ alle in L_r , und da L_r linear unabhängig ist, folgt auch die lineare Unabhängigkeit der Familie $v_1, \dots, v_r$. Damit ist die Behauptung bewiesen und wir können das Lemma von Zorn anwenden.

Es ist auch leicht, den Beweis so anzupassen, dass die Basis ein vorgegebenes linear unabhängiges System enthält und in einem vorgegebenen Erzeugendensystem von V enthalten ist, indem man die Definition von $\mathcal{U}$ entsprechend modifiziert.

Durch ein geschicktes Vorgehen kann man (mithilfe des Lemmas von Zorn) auch den Basisaustauschsatz auf unendliche Mengen erweitern und zeigen, dass für jeden Vektorraum V , jedes Erzeugendensystem E von V und jede linear unabhängige Teilmenge M von V eine injektive Abbildung $\iota: M \rightarrow E$ existiert, so dass auch $(E \setminus \iota(M)) \cup M$ ein Erzeugendensystem von V ist. Es werden also sozusagen die Elemente in $\iota(M)$ durch die Elemente in M ausgetauscht. Dazu versieht man die Menge aller injektiven Abbildungen $\iota': M' \rightarrow E$, so dass $M' \subseteq M$ und $(E \setminus \iota'(M')) \cup M'$ ein Erzeugendensystem von V ist, mit einer geeigneten partiellen Ordnung, benutzt eine ähnliche Methode wie oben, um zu zeigen, dass die Voraussetzungen des Lemmas von Zorn erfüllt sind, und überlegt sich unter Ausnutzung des »endlichen« Basisaustauschprinzips, dass ein maximales Element dieser Menge notwendigerweise M als Definitionsbereich haben muss.

Zusammen mit dem Satz von Schröder-Bernstein (Theorem 3.85) erhält man dann, dass es zwischen je zwei Basen eines Vektorraums eine Bijektion gibt. Siehe auch [So-AZT] 5.3 für eine etwas ausführlichere Diskussion. □ Ergänzung 6.48

6.5. Wie berechne ich ...?

Ich habe etwas gezögert, ob es überhaupt sinnvoll ist, diesen Abschnitt ins Skript aufzunehmen, denn gerade für diese »Rechenverfahren« gilt: Es ist besser, wenn Sie sich selbst überlegen, wie Sie die uns zur Verfügung stehenden Verfahren (und das ist im Moment eigentlich nur der Gauß-Algorithmus) auf die zu bewältigende Aufgabe anwenden. Insbesondere rate ich davon ab, die Varianten, die unten diskutiert werden, auswendig zu lernen. Letztlich habe ich mich aber doch entschieden, einige Aufgabentypen hier zu sammeln, auch um illustrieren zu können, dass es am Ende von Kapitel 7 möglich sein wird, einiges noch klarer darzustellen. Konkrete Beispiele zu diesen Aufgabentypen finden Sie in den Online-Aufgaben.

Sei K ein Körper. Um Basen von Untervektorräumen in K^n , $n \in \mathbb{N}$, zu berechnen und damit zusammenhängende Fragen über Untervektorräume zu beantworten, erweist sich der Gauß-Algorithmus als Universalwerkzeug.

Wir haben bereits erwähnt, dass die Spalten einer Matrix $A \in M_{m \times n}(K)$ genau dann eine linear unabhängige Familie in K^m bilden, wenn das homogene Gleichungssystem $Ax = 0$ eindeutig lösbar ist, also nur die triviale Lösung besitzt. Ein Vektor $b \in K^m$ liegt genau dann im Untervektorraum von K^m , der von den Spalten von A erzeugt wird, wenn das Gleichungssystem $Ax = b$ lösbar ist. Die Spalten von A bilden genau dann eine Basis von K^m , wenn das lineare Gleichungssystem $Ax = b$ für alle b eindeutig lösbar ist. Das ist nur dann möglich, wenn $m = n$ ist, d.h. wenn A quadratisch ist; es ist dann äquivalent dazu, dass die Matrix reduzierte Zeilenstufenform E_n hat, und auch dazu, dass A invertierbar ist. Siehe Korollar 5.22 und Satz 5.54.

6.5.1. Basis der Lösungsmenge eines homogenen linearen Gleichungssystems. Diese Aufgabenstellung haben wir bereits besprochen, siehe Satz 6.22.

Eine äquivalente Formulierung des Problems ist, eine Basis vom Kern $\text{Ker}(A)$ einer Matrix A zu bestimmen, denn $\text{Ker}(A)$ ist genau die Lösungsmenge des homogenen Gleichungssystems $Ax = 0$.

6.5.2. Eine Basis innerhalb eines gegebenen Erzeugendensystems. Sei $U \subseteq K^m$ ein Untervektorraum und sei $u_1, \dots, u_n \in U$ ein Erzeugendensystem von U . Wir wissen, dass es eine Basis von U gibt, die aus Vektoren dieses Erzeugendensystems besteht. Um eine solche zu finden, betrachten wir die Matrix $A \in M_{m \times n}$, deren j -te Spalte der Spaltenvektoren $u_j \in K^m$ ist. Sei B eine Matrix in Zeilenstufenform, die aus A durch elementare Zeilenumformungen entsteht. Dann bilden die u_j für diejenigen Indizes j , die in B zu einer Spalte mit einer führenden Eins korrespondieren, eine Basis von U .

Begründung des Verfahrens. Sind $s_1, \dots, s_n \in K^m$ die Spalten einer Matrix S , geht die Matrix S' mit den Spalten $s'_1, \dots, s'_n$ aus S durch elementare Zeilenumformungen hervor und sind $c_j \in K$, so gilt $\sum_{j=1}^n c_j s_j = 0$ genau dann, wenn $\sum_{j=1}^n c_j s'_j = 0$. Das ist für alle drei Typen von elementaren Zeilenumformungen leicht einzusehen.

Das bedeutet, dass für eine Teilmenge $J \subseteq \{1, \dots, n\}$ die Spalten $s_j, j \in J$, genau dann linear unabhängig sind, wenn es $s'_j, j \in J$ sind.

Eine Basis innerhalb des gegebenen Erzeugendensystems ist eine maximale linear unabhängige Teilmenge darin. Für eine Matrix in Zeilenstufenform ist klar, dass die Spalten mit führenden Einsen eine solche maximale linear unabhängige Teilmenge bilden.

Eine äquivalente Formulierung des Problems ist, eine Basis vom Bild einer Matrix A zu bestimmen, denn das Bild von A ist genau der von den Spalten von A erzeugte Untervektorraum.

6.5.3. Eine linear unabhängige Teilmenge zu einer Basis ergänzen. In ähnlicher Weise kann man eine linear unabhängige Teilmenge L in einem Untervektorraum $U \subseteq K^m$ zu einer Basis ergänzen, wenn man irgendein Erzeugendensystem E von U kennt. Denn dann ist $L \cup E$ auch ein Erzeugendensystem, auf das man das Verfahren des vorherigen Abschnitts anwenden kann. Wenn man die Elemente von L als die linken Spalten der Matrix A schreibt, führt die lineare Unabhängigkeit dazu, dass diese Spalten in der Zeilenstufenform jedenfalls führende Einsen enthalten werden. An den zusätzlich auftretenden führenden Einsen kann man ablesen, welche Elemente von E man noch hinzunehmen kann, um eine Basis zu erhalten.

6.5.4. Eine Basis vom Durchschnitt und der Summe von Untervektorräumen. Seien $U, W \subseteq K^m$ Untervektorräume mit Basen $u_1, \dots, u_r$ und $w_1, \dots, w_s$. Wir wollen Basen von $U + W$ und $U \cap W$ finden.

Da $u_1, \dots, u_r, w_1, \dots, w_s$ die Summe $U + W$ erzeugen, können wir das vorher besprochene Verfahren anwenden, um innerhalb dieses Erzeugendensystems eine Basis von $U + W$ zu finden.

Um eine Basis von $U \cap W$ zu finden, verfahren wir folgendermaßen. Ein Element $v \in K^m$ liegt genau dann in $U \cap W$, wenn es sowohl als Linearkombination der u_j , als auch als Linearkombination der w_j darstellbar ist, wenn also a_j und b_j existieren mit

$$a_1 u_1 + \dots + a_r u_r = v = b_1 w_1 + \dots + b_s w_s.$$

Um alle Elemente von $U \cap W$ zu finden, müssen wir also alle Tupel $(a_1, \dots, a_r, b_1, \dots, b_s)$ finden mit

$$a_1 u_1 + \dots + a_r u_r - b_1 w_1 - \dots - b_s w_s = 0,$$

oder mit anderen Worten die Lösungsmenge des linearen Gleichungssystems $Ax = 0$, wo A die Matrix mit Spalten $u_1, \dots, u_r, -w_1, \dots, -w_s$ ist. Ist $(x_j)_j \in K^{r+s}$ ein Element der Lösungsmenge, so ist $\sum_{j=1}^r x_j u_j = \sum_{j=r+1}^{r+s} x_j w_j \in U \cap W$. Wir bezeichnen für den Moment für

$x = (x_j)_j \in K^{r+s}$ mit $\bar{x}$ den Vektor $\sum_{j=1}^r x_j u_j$. Bilden dann $v_1, \dots, v_t$ eine Basis der Lösungsmenge, so ist $\bar{v}_1, \dots, \bar{v}_t$ ein Erzeugendensystem von $U \cap W$, und sogar eine Basis, wie wir gleich begründen werden.

Zunächst die folgende Bemerkung: Weil mit $w_1, \dots, w_s$ auch $-w_1, \dots, -w_s$ eine Basis ist, kann man auch die beiden Verfahren in einem Schritt durchführen und aus der Zeilenstufenform der Matrix A mit Spalten $u_1, \dots, u_r, -w_1, \dots, -w_s$ eine Basis von $U + W$ ablesen.

Insbesondere folgt, dass die Anzahl der führenden Einsen in der Zeilenstufenform von A genau $\dim(U + W)$ ist. Die Dimension t der Lösungsmenge ist also $r + s - \dim(U + W) = \dim(U) + \dim(W) - \dim(U + W) = \dim(U \cap W)$, siehe Satz 6.46. Weil das Erzeugendensystem $\bar{v}_1, \dots, \bar{v}_t$ von $U \cap W$, das wir oben gefunden haben, genau $\dim(U \cap W)$ Elemente hat, handelt es sich um eine Basis. (Man kann sich die lineare Unabhängigkeit auch direkt überlegen, ohne die Dimensionsformel für den Durchschnitt von Untervektorräumen zu benutzen. Sehen Sie, wie?)

6.5.5. Eine möglichst einfache Basis eines Untervektorraums finden – Spaltenumformungen. Seien K ein Körper, V ein K -Vektorraum und seien $u_1, \dots, u_n \in V$. Wir setzen $U := \langle u_1, \dots, u_n \rangle$, so dass $\{u_1, \dots, u_n\}$ ein Erzeugendensystem des Untervektorraums U ist. Es gibt aber in aller Regel viele andere Erzeugendensysteme. Ist die Familie $u_1, \dots, u_n$ nicht linear unabhängig, so können wir sie geeignet verkleinern, ohne ihr Erzeugnis zu ändern. Wir können auch den Austauschsatz auf den Vektorraum U anwenden.

Hier wollen wir besprechen, wie wir das gegebene Erzeugendensystem durch direkte Umformungen abändern können, ohne die lineare Hülle zu verändern. Sind $1 \leq i, j \leq n$, $i \neq j$ und $a \in K$, so gilt

$$\langle u_1, \dots, u_n \rangle = \langle u_1, \dots, u_{i-1}, u_i + au_j, u_{i+1}, \dots, u_n \rangle.$$

Ist $1 \leq i \leq n$ und $a \in K^\times$, so gilt

$$\langle u_1, \dots, u_n \rangle = \langle u_1, \dots, u_{i-1}, au_i, u_{i+1}, \dots, u_n \rangle.$$

Beide Behauptungen prüft man unmittelbar nach, indem man zeigt, dass jeweils die linke Seite in der rechten enthalten ist, und umgekehrt.

Sei nun $V = K^m$, und $u_1, \dots, u_n \in K^m$. Sei $A \in M_{m \times n}(K)$ die Matrix mit den Spalten $u_1, \dots, u_n$. Dann können wir, ähnlich wie wir in Abschnitt 5.2.1 elementare Zeilenumformungen definiert haben, elementare Spaltenumformungen betrachten: (I) Addieren eines Vielfachen einer Spalte zu einer anderen Spalte, (II) Vertauschen zweier Spalten, (III) Multiplizieren einer Spalte mit $a \in K^\times$. Die Überlegungen des vorherigen Absatzes können wir dann formulieren als

LEMMA 6.50. Sei $A \in M_{m \times n}(K)$ eine Matrix, und sei A' eine Matrix, die aus A durch elementare Spaltenumformungen entsteht. Sei $U \subseteq K^m$ der Untervektorraum, der von den Spalten von A erzeugt wird, und sei $U' \subseteq K^m$ der Untervektorraum, der von den Spalten von A' erzeugt wird. Dann gilt $U = U'$.

Das ermöglicht es, durch *Spaltenumformungen* eine besonders einfache Basis eines Untervektorraums von K^m zu finden, der durch ein Erzeugendensystem gegeben ist: Analog zum gewöhnlichen Gauß-Algorithmus kann man eine Matrix durch elementare Spaltenumformungen auch »reduzierte Spaltenstufenform« bringen. (Wir können definieren, dass eine Matrix A reduzierte Spaltenstufenform hat, wenn die transponierte Matrix A^t reduzierte Zeilenstufenform hat.)

Machen Sie sich die Unterschiede zwischen Zeilen- und Spaltenumformungen klar: Spaltenumformungen verändern nicht den Untervektorraum, der von den Spaltenvektoren erzeugt wird; schon allein, weil beliebige Spaltenvertauschungen erlaubt sind, ist aber klar, dass man keine Rückschlüsse der Art ziehen kann, welche der ursprünglich gegebenen Vektoren linear

unabhängig, ein Erzeugendensystem oder eine Basis sind. Bei Zeilenumformungen bleibt hingegen die Information, ob eine Spalte eine Linearkombination von anderen Spalten ist, erhalten; es verändert sich aber der Untervektorraum, der von den Spalten erzeugt wird.

Videos von 3Blue1Brown zum Thema Vektoren, Vektorräume, lineare Hülle, ...

[Vectors, what even are they^a](https://youtu.be/fNk_zzaMoSs)

[Linear combinations, span, and basis vectors^b](https://youtu.be/k7RM-ot2NWY)

^a https://youtu.be/fNk_zzaMoSs

^b <https://youtu.be/k7RM-ot2NWY>

6.6. Produkt und direkte Summe von Vektorräumen

Sei K ein Körper. Sind $V_1, \dots, V_n$ Vektorräume über K , so ist das kartesische Produkt $V_1 \times \dots \times V_n$, also die Menge aller n -Tupel $(v_1, \dots, v_n)$ mit $v_i \in V_i, i = 1, \dots, n$ mit der *komponentenweisen Addition*

$$(v_1, \dots, v_n) + (w_1, \dots, w_n) := (v_1 + w_1, \dots, v_n + w_n)$$

und der *komponentenweisen Skalarmultiplikation*

$$a \cdot (v_1, \dots, v_n) = (av_1, \dots, av_n)$$

ein K -Vektorraum. Dies prüft man leicht nach, weil sich alle Bedingungen in den einzelnen Einträgen separat überprüfen lassen und man nur benutzen muss, dass alle V_i Vektorräume sind. Wir nennen diesen Vektorraum das *Produkt* der Vektorräume $V_1, \dots, V_n$.

Genauso definiert man das Produkt $\prod_{i \in I} V_i$ einer (möglicherweise unendlichen) Familie $(V_i)_{i \in I}$ von Vektorräumen über K . Ist $V = V_i$ für alle i , so schreiben wir auch $V^I = \prod_{i \in I} V$ für diesen Vektorraum. Die Vektorraumstruktur auf dem Produkt ist dieselbe wie die, die wir in Beispiel 6.2 auf $\text{Abb}(I, V)$ definiert haben.

Die Teilmenge

$$\bigoplus_{i \in I} V_i := \left\{ (v_i)_{i \in I} \in \prod_{i \in I} V_i; \text{ höchstens endlich viele } v_i \text{ sind } \neq 0 \right\}$$

von $\prod_{i \in I} V_i$ ist ein Untervektorraum und heißt die *direkte Summe* der Vektorräume V_i . Man spricht manchmal auch vom *Koprodukt* der Vektorräume V_i .

Wenn die Indexmenge I endlich ist, dann gilt natürlich $\bigoplus_{i \in I} V_i = \prod_{i \in I} V_i$.

BEMERKUNG 6.51. Wir hatten das Symbol $\oplus$ in Definition 6.12 schon in einer etwas anderen Weise definiert, und zwar hatten wir für Untervektorräume U, W eines Vektorraums V die Summe von U und W als $U \oplus W$ geschrieben, wenn $U \cap W = 0$ gilt. Wir können das leicht mit der neuen Schreibweise zusammenbringen, denn unter der Voraussetzung $U \cap W = 0$ können wir den Vektorraum $U \oplus W$ (im neu definierten Sinne) identifizieren mit der Summe $U + W$; das ist genau die Aussage von Lemma 6.15.

Wenn es notwendig ist, den kleinen Unterschied zwischen den beiden Sichtweisen zu betonen, dann bezeichnet man manchmal die in diesem Abschnitt definierte direkte Summe als die *äußere direkte Summe*, und die vorherige Bezeichnung als die *innere direkte Summe*. Die äußere direkte Summe ist also eine Konstruktion, die aus zwei (oder mehr) Vektorräumen einen neuen konstruiert; wir können $U \oplus W$ für beliebige Untervektorräume, sogar $V \oplus V$, und $V \oplus V'$ für Vektorräume V, V' , die gar nicht Untervektorräume desselben Vektorraums sind, betrachten. Eine innere direkte Summe zu bilden, ist eine Eigenschaft, die Untervektorräume eines Vektorraums haben können, oder eben nicht – es ist gerade die Eigenschaft,

dass die äußere direkte Summe dieser Untervektorräume mit der Summe innerhalb des umgebenden Vektorraums identifiziert werden kann. Siehe auch Bemerkung 7.13. $\diamond$

Im Fall $V_i = V$ für alle i , dass also alle »Summanden« übereinstimmen, schreibt man auch $V^{(I)}$ statt $\bigoplus_{i \in I} V$. Es ist also $V^{(I)}$ der Vektorraum aller durch I indizierter Familien von Elementen aus V , in denen höchstens endlich viele Vektoren vom Nullvektor verschieden sind, mit der komponentenweisen Addition und Skalarmultiplikation.

BEISPIEL 6.52. Seien K ein Körper, I eine Menge, und $V = K^{(I)} = \bigoplus_{i \in I} K$, die durch I indizierte direkte Summe von Kopien von K . Ist $I = \{1, \dots, n\}$, so ist $K^{(I)} = K^I = K^n$ der Standardvektorraum von n -Tupeln von Elementen von K .

In jedem Fall können wir für $i \in I$ mit $e_i \in K^{(I)}$ das Element bezeichnen, das an der Stelle i den Eintrag 1, und überall sonst den Eintrag 0 hat. Ist $v = (v_i)_{i \in I}$ irgendein Element von $K^{(I)}$, so gilt

$$v = \sum_{i \in I} v_i e_i,$$

wobei wir die Summe bilden können, weil nur endlich viele v_i von Null verschieden sind, und es sich daher letztlich um eine endliche Summe handelt. Dies ist die einzige Linearkombination der e_i , die den Vektor v darstellt, und wir sehen so, dass $(e_i)_{i \in I}$ eine Basis von $K^{(I)}$ ist. Diese Basis nennt man auch die Standardbasis von $K^{(I)}$, und bezeichnet die e_i als die Standardbasisvektoren.

Da diese Familie für unendliche Mengen I aus unendlich vielen Elementen besteht, sehen wir erneut, dass der Vektorraum $K^{(I)}$ nicht endlich erzeugt ist. $\diamond$

Ist I unendlich, so hat K^I den nicht-endlich-erzeugten Untervektorraum $K^{(I)}$, ist also wegen Satz 6.44 nicht endlich erzeugt. Es ist aber nicht möglich, ähnlich explizit wie im vorherigen Beispiel eine Basis von K^I anzugeben.

Sind $V_1, \dots, V_n$ Vektorräume über dem Körper K , so können wir aus Basen der V_i eine Basis der direkten Summe »zusammensetzen«:

LEMMA 6.53. Seien K ein Körper und seien $V_1, \dots, V_n$ Vektorräume über K . Für $i = 1, \dots, n$ sei $\iota_i: V_i \rightarrow \bigoplus_{i=1}^n V_i$ die Abbildung, die $v \in V_i$ abbildet auf das Element $(0, \dots, 0, v, 0, \dots, 0)$ (mit v an der i -ten Stelle).

Ist $\mathcal{B}_i \subset V_i$ eine Basis, $i = 1, \dots, n$, so ist $\bigcup_{i=1}^n \iota_i(\mathcal{B}_i)$ eine Basis von $\bigoplus_{i=1}^n V_i$.

BEWEIS. Per Induktion können wir uns auf den Fall der direkten Summe $V \times W$ von zwei Vektorräumen beschränken. Die Abbildungen ι sind dann $v \mapsto (v, 0)$ und $w \mapsto (0, w)$. Sei $v_1, \dots, v_m$ eine Basis von V und $w_1, \dots, w_n$ eine Basis von W .

Für $(v, w) \in V \oplus W$ schreiben wir

$$v = \sum_{i=1}^m a_i v_i, \quad w = \sum_{j=1}^n b_j w_j.$$

Dann ist

$$(v, w) = \sum_{i=1}^m a_i (v_i, 0) + \sum_{j=1}^n b_j (0, w_j),$$

und es ist leicht zu sehen, dass diese Darstellung eindeutig ist. $\square$

Das Lemma lässt sich auch auf den Fall von unendlichen direkten Summen verallgemeinern. Im endlich erzeugten Fall erhalten wir als Folgerung:

KOROLLAR 6.54. Seien K ein Körper und seien $V_1, \dots, V_n$ endlich erzeugte Vektorräume über K . Dann gilt

$$\dim \bigoplus_{i=1}^n V_i = \sum_{i=1}^n \dim(V_i).$$

6.7. Ergänzungen *

BEISPIEL 6.55. Sei $V = \mathbb{R}$, betrachtet als $\mathbb{Q}$ -Vektorraum. Dieser Vektorraum ist nicht endlich erzeugt. Am einfachsten ist es, das mit einem Mächtigkeitsargument zu sehen, vergleiche Abschnitt 3.15.

Es ist nämlich $\mathbb{Q}$ eine abzählbar unendliche Menge, und es ist nicht sehr schwer zu sehen, dass dann auch alle Mengen der Form $\mathbb{Q}^n$ abzählbar sind. Deshalb kann keine dieser Mengen in Bijektion zu der überabzählbaren Menge $\mathbb{R}$ stehen.

Genauer zeigt ein ähnliches Argument sogar, dass für jede Basis $(b_i)_{i \in I}$ von $\mathbb{R}$ als $\mathbb{Q}$ -Vektorraum die Indexmenge I selbst überabzählbar sein muss, weil ein $\mathbb{Q}$ -Vektorraum mit abzählbarer Basis ebenfalls eine abzählbare Menge ist. $\diamond$

Wir können nun den folgenden Satz beweisen, den wir als Satz 5.35 als ein Beispiel für eine Aussage über Matrizen formuliert hatten, deren Beweis eine Weiterentwicklung der Theorie erfordert. In der Linearen Algebra 2 werden wir sehen, wie sich der Satz in die allgemeine Theorie der Eigenwerte einer Matrix einfügt.

SATZ 6.56. Seien K ein Körper, $n \geq 1$, $A \in M_{n \times n}(K)$. Wenn eine natürliche Zahl N existiert mit $A^N = 0$, dann gilt $A^n = 0$.

BEWEIS. Um zu zeigen, dass $A^n = 0$, zeigen wir $\text{Ker}(A^n) = K^n$. Das impliziert $A^n = 0$, sonst würde für mindestens ein i gelten, dass $A^n e_i \neq 0$. Wir betrachten $U_i := \text{Ker}(A^i)$. Dann ist $U_i \subset K^n$ ein Untervektorraum, und nach Voraussetzung gilt $U_N = K^n$. Wenn $A^i v = 0$, dann gilt erst recht $A^{i+1} v = A A^i v = 0$, die U_i sind also jeweils ineinander enthalten:

$$0 = U_0 \subseteq U_1 \subseteq \dots \subseteq U_N = K^n.$$

Behauptung. Gilt $U_i = U_{i+1}$, so gilt $U_i = U_j$ für alle $j \geq i$.

Begründung. Es genügt zu zeigen, dass aus $U_i = U_{i+1}$ auch $U_{i+1} = U_{i+2}$ folgt. Danach kann man induktiv fortfahren. Wir wissen bereits, dass $U_{i+1} \subseteq U_{i+2}$. Sei nun $v \in U_{i+2}$, also $A^{i+2} v = 0$. Das bedeutet $Av \in U_{i+1} = U_i$, also $A^i(Av) = 0$. Wir sehen, dass $v \in U_{i+1}$, wie behauptet.

Um den Beweis abzuschließen, zeigen wir, dass aus der Behauptung die Gleichheit $U_n = K^n$ folgt. Weil $U_N = K^n$ gilt, zeigt die Behauptung, dass $U_i = U_{i+1}$ nur gelten kann, wenn $U_i = K^n$ gilt. Im Fall einer echten Inklusion $U_i \subsetneq U_{i+1}$ gilt $\dim(U_i) < \dim(U_{i+1})$ (Satz 6.44). Solange wir in der Kette der U_i nicht bei K^n »angekommen« sind, muss also in jedem Schritt die Dimension um mindestens Eins ansteigen. Deshalb gilt $\dim U_i \geq i$ für alle $i = 0, \dots, n$. Aus $\dim U_n \geq n = \dim K^n$ folgt $U_n = K^n$, wobei wir wieder Satz 6.44 anwenden. $\square$

ERGÄNZUNG 6.57 (Die Kardinalität endlicher Körper). Sei K ein endlicher Körper. Wir haben in Abschnitt 4.2.2 gesehen, dass die Charakteristik von K eine Primzahl p ist. Es gilt dann $p_K = 1 + \dots + 1 = 0$ in K (mit p Summanden in der Summe), und für alle $1 \leq n < p$ ist $n_K \neq 0$.

LEMMA 6.58. Der Körper K wird durch die Körperaddition auf K und die Skalarmultiplikation

$$\cdot: \mathbb{F}_p \times K \rightarrow K, \quad n_{\mathbb{F}_p} \cdot x = n_K x = \underbrace{x + \dots + x}_{n \text{ Summanden}}$$

zu einem $\mathbb{F}_p$ -Vektorraum.

BEWEIS. Das ist eine leichte Rechnung. Genauer kann man sagen, dass die Abbildung $\mathbb{F}_p \rightarrow K, n_{\mathbb{F}_p} \mapsto n_K$, den Körper K zu einem Erweiterungskörper von $\mathbb{F}_p$ macht (und wir hatten in Beispiel 6.2 angemerkt, dass dann der Erweiterungskörper auch ein Vektorraum über dem kleineren Körper ist). $\square$

KOROLLAR 6.59. Es gibt $r \in \mathbb{N}_{\geq 1}$ mit $\#K = p^r$.

BEWEIS. Da K nach Voraussetzung endlich ist, ist K als $\mathbb{F}_p$ -Vektorraum endlich erzeugt. Daher besitzt K eine Basis $b_1, \dots, b_r, r \in \mathbb{N}$. Weil K als Körper mehr als ein Element haben muss, kann nicht $r = 0$ gelten.

Wir können jedes Element von K in eindeutiger Weise als $\sum_{i=1}^r a_i b_i$ mit $a_i \in \mathbb{F}_p$ schreiben. Die Abbildung

$$\mathbb{F}_p^r \rightarrow K, \quad (a_1, \dots, a_r)^t \mapsto \sum_{i=1}^r a_i b_i,$$

ist also eine Bijektion, und wir sehen $\#K = \#\mathbb{F}_p^r = p^r$. $\square$

$\square$ Ergänzung 6.57

ERGÄNZUNG 6.60 (Die Fibonacci-Zahlen, Fortsetzung). In dieser Ergänzung beweisen wir die folgende Formel für die n -te Fibonacci-Zahl (vgl. Frage 2.1, Beispiel 5.60). Siehe auch Beispiel 10.19 für einen ganz anderen Beweis für diese Formel. Der Beweis, den wir hier geben, lässt sich auch auf andere, in ähnlicher Weise rekursiv definierte Folgen anpassen.

SATZ 6.61. Sei $(F_n)_n$ die Folge der Fibonacci-Zahlen. Dann gilt

$$F_n = \frac{1}{\sqrt{5}} \left(\left(\frac{1 + \sqrt{5}}{2} \right)^n - \left(\frac{1 - \sqrt{5}}{2} \right)^n \right).$$

Auch wenn die Formel nicht nützlich ist, um F_n auszurechnen, handelt es sich dennoch um ein interessantes Ergebnis. Es ist ja nicht einmal völlig offensichtlich, dass es sich bei der rechten Seite der Formel überhaupt um eine ganze Zahl handelt.

BEWEIS. Wir betrachten den $\mathbb{R}$ -Vektorraum $\mathbb{R}^{\mathbb{N}} = \prod_{n \in \mathbb{N}} \mathbb{R}$. Als Menge ist dies die Menge aller Folgen $(a_n)_{n \in \mathbb{N}}$ von reellen Zahlen a_n , und wir versehen diese Menge mit der komponentenweisen Addition und Skalarmultiplikation:

$$(a_n)_n + (b_n)_n = (a_n + b_n)_n, \quad a \cdot (a_n)_n = (a a_n)_n.$$

Es ist klar, dass wir so einen Vektorraum erhalten. (In der Sprechweise von Abschnitt 6.6 handelt es sich einfach um das Produkt des Vektorraums $\mathbb{R}$ mit sich selbst mit Indexmenge $\mathbb{N}$.)

Sei W der Untervektorraum von V , der aus denjenigen Folgen $(a_n)_n$ besteht, für die $a_{n+2} = a_{n+1} + a_n$ für alle $n \in \mathbb{N}$ gilt. Die Fibonacci-Folge ist ein Element von W . Ein anderes Element ist die Folge

$$(G_n)_n = (1, 0, 1, 1, 2, 3, 5, 8, 13, \dots).$$

Weil ein Element aus W durch seine ersten beiden Einträge vollständig bestimmt ist, bilden $(F_n)_n$ und $(G_n)_n$ ein Erzeugendensystem von W . Da sie nicht Vielfache von einander sind, handelt es sich sogar um eine Basis, und wir sehen, dass $\dim W = 2$.

Um den gesuchten Ausdruck für F_n zu finden, betrachten wir noch eine andere Basis von W , und zwar eine von einer besonders einfachen Form. Der Raum W enthält genau zwei Elemente der Form

$$(1, \varphi, \varphi^2, \varphi^3, \dots)$$

mit $\varphi \in \mathbb{R}^\times$. Denn die Folge aller Potenzen der reellen Zahl $\varphi \neq 0$ ist genau dann in W , wenn für alle n

$$\varphi^{n+2} = \varphi^{n+1} + \varphi^n$$

gilt, äquivalent:

$$\varphi^2 - \varphi - 1 = 0.$$

Die beiden Lösungen dieser quadratischen Gleichung in $\mathbb{R}$ sind

$$\varphi_1 = \frac{1 + \sqrt{5}}{2}, \quad \varphi_2 = \frac{-1 + \sqrt{5}}{2}.$$

Wir schreiben

$$w_1 = (\varphi_1^n)_n, \quad w_2 = (\varphi_2^n)_n.$$

Dann bilden auch w_1 und w_2 eine Basis von W , weil diese beiden Elemente offenbar linear unabhängig sind.

Es gilt dann $(F_n)_n = \frac{1}{\sqrt{5}}w_1 - \frac{1}{\sqrt{5}}w_2$, denn es genügt, das für die ersten beiden Einträge zu überprüfen. Es folgt daraus, dass

$$F_n = \frac{1}{\sqrt{5}}(\varphi_1^n - \varphi_2^n) = \frac{1}{\sqrt{5}} \left(\left(\frac{1 + \sqrt{5}}{2} \right)^n - \left(\frac{1 - \sqrt{5}}{2} \right)^n \right).$$

für alle $n \in \mathbb{N}$ – genau die Formel, die wir zeigen wollten. $\square$

Für den Grenzwert des Verhältnisses F_{n+1} / F_n im Sinne der Analysis erhalten wir damit das folgende Ergebnis:

KOROLLAR 6.62. *Es gilt*

$$\lim_{n \rightarrow \infty} \frac{F_{n+1}}{F_n} = \frac{1 + \sqrt{5}}{2}.$$

BEWEIS. Wir benutzen die Notation φ_1, φ_2 wie im Beweis des Satzes und die grundlegenden Rechenregeln für Grenzwerte. Dann gilt $|\varphi_2 / \varphi_1| = (\sqrt{5} - 1)^2 / 4 < 1$, also

$$\lim_{n \rightarrow \infty} \frac{\varphi_2^n}{\varphi_1^n} = 0,$$

und daher

$$\lim_{n \rightarrow \infty} F_{n+1} / F_n = \lim_{n \rightarrow \infty} \frac{\varphi_1^{n+1} - \varphi_2^{n+1}}{\varphi_1^n - \varphi_2^n} = \lim_{n \rightarrow \infty} \frac{\varphi_1 - \varphi_2 \frac{\varphi_2^n}{\varphi_1^n}}{1 - \frac{\varphi_2^n}{\varphi_1^n}} = \varphi_1 = \frac{1 + \sqrt{5}}{2}.$$

$\square$

$\square$ Ergänzung 6.60

ERGÄNZUNG 6.63 (Das quadratische Sieb, Fortsetzung). Wir können nun Frage 2.2 beantworten:

SATZ 6.64. *Gegeben seien eine natürliche Zahl $n \geq 1$ und n verschiedene Primzahlen $p_1, \dots, p_n$. Wenn $a_1, \dots, a_{n+1}$ natürliche Zahlen > 1 sind, in deren Primfaktorzerlegungen nur die Primzahlen $p_1, \dots, p_n$ vorkommen, dann gibt es eine Möglichkeit, einige der Zahlen a_i so auszuwählen, dass ihr Produkt eine Quadratzahl ist.*

BEWEIS. Wir betrachten zu jedem a_i die eindeutige Primfaktorzerlegung, in der nach Voraussetzung nur die Primzahlen $p_1, \dots, p_n$ vorkommen, etwa

$$a_i = p_1^{a_{i1}} \cdot p_2^{a_{i2}} \cdot \dots \cdot p_n^{a_{in}}.$$

Die Primfaktorzerlegung eines Produkts von Zahlen a_i erhalten wir dann, indem wir die Exponenten addieren. Wir möchten ein Produkt von a_i 's finden, in dem alle Exponenten in der Primfaktorzerlegung gerade sind, denn das bedeutet gerade, dass es sich um eine Quadratzahl handelt.

Wir können das auch folgendermaßen ausdrücken: Setze

$$v_{ij} = \begin{cases} 0 & a_{ij} \text{ gerade} \\ 1 & a_{ij} \text{ ungerade} \end{cases} \in \mathbb{F}_2$$

(es ist also v_{ij} die Restklasse von a_{ij} in $\mathbb{F}_2$) und sei $v_i = (v_{i1}, \dots, v_{in})^t \in \mathbb{F}_2^n$.

Aus den $n+1$ Zahlen $a_1, \dots, a_{n+1}$ erhalten wir so $n+1$ Vektoren $v_1, \dots, v_{n+1}$ in dem n -dimensionalen $\mathbb{F}_2$ -Vektorraum $\mathbb{F}_2^n$. Diese Vektoren müssen also linear abhängig sein. Da die einzigen Elemente von $\mathbb{F}_2$ die 0 und 1 sind, hat eine nicht-triviale Linearkombination die Form

$$v_{i_1} + \dots + v_{i_r} = 0,$$

für eine nicht-leere Teilmenge $\{i_1, \dots, i_r\} \subseteq \{1, \dots, n+1\}$.

Das bedeutet aber gerade, dass die Summen der entsprechenden Zahlen a_{i_j} gerade sind, also dass das Produkt $a_{i_1} \cdot \dots \cdot a_{i_r}$ eine Quadratzahl ist. $\square$

$\square$ Ergänzung 6.63

ERGÄNZUNG 6.65 (Schiefkörper, Fortsetzung). Wir hatten in Ergänzung 4.9 den Begriff des Schiefkörpers und der Divisionsalgebra definiert. Für den Begriff der Divisionsalgebra lässt sich äquivalent die folgende Definition geben: Eine Divisionsalgebra über einem Körper K ist ein Schiefkörper D , der zugleich ein K -Vektorraum ist, und so dass die Schiefkörperaddition und die Vektorraumaddition übereinstimmen und die Multiplikation des Schiefkörpers D und die Skalarmultiplikation $K \times D \rightarrow D$ im folgenden Sinne kompatibel sind: $a(xy) = (ax)y = x(ay)$ für alle $a \in K, x, y \in D$.

Dann ist die Abbildung $K \rightarrow D, a \mapsto a \cdot 1_D$, eine injektive Abbildung (warum?), so dass wir K mit einer Teilmenge von D identifizieren können. Es ist dann nicht schwierig, die Äquivalenz der beiden Definitionen zu überprüfen.

Wir können nun mithilfe des Dimensionsbegriffs die folgenden Sätze formulieren. Dabei nennen wir eine Divisionsalgebra endlich-dimensional, wenn sie als Vektorraum betrachtet endliche Dimension hat.

SATZ 6.66. *Es gibt keine endlich-dimensionale Divisionsalgebra über dem Körper $\mathbb{R}$, die als Vektorraum ungerade Dimension > 1 hat.*

Wir werden diesen Satz in Ergänzung 10.20 beweisen. Genauer kann man zeigen, dass die einzigen möglichen Dimensionen für eine Divisionsalgebra über $\mathbb{R}$ die Zahlen 1, 2 und 4 sind. (Dass es für diese Dimensionen tatsächlich eine Divisionsalgebra gibt, ist mit dem, was wir wissen, klar: $\mathbb{R}$, die komplexen Zahlen $\mathbb{C}$ und die Hamiltonschen Quaternionen $\mathbb{H}$.)

SATZ 6.67. *Es gibt keine endlich-dimensionale Divisionsalgebra über den komplexen Zahlen $\mathbb{C}$ von Vektorraum-Dimension > 1 .*

Auch diesen Satz werden wir in Ergänzung 10.20 beweisen. (Allerdings müssen wir benutzen, dass jede nicht-konstante Polynomfunktion über den komplexen Zahlen eine Nullstelle besitzt. Dieses Ergebnis nennt man den Fundamentalsatz der Algebra.) $\square$ Ergänzung 6.65

Lineare Abbildungen

Wir fixieren wie üblich einen Körper K .

7.1. Lineare Abbildungen

In diesem Kapitel definieren und untersuchen wir den Begriff der *linearen Abbildung*. Eine lineare Abbildung ist eine Abbildung $f: V \rightarrow W$ zwischen Vektorräumen V und W über demselben Körper K , die »mit den Vektorraum-Strukturen verträglich« ist (Definition 7.1). Zwei Punkte, die diesen Begriff motivieren, sind die folgenden:

- Es ist normalerweise nicht sehr interessant, beliebige Abbildungen zwischen Vektorräumen zu betrachten. Genauso wie man zum Beispiel in der Analysis nicht alle, sondern in der Regel nur stetige oder sogar nur differenzierbare Abbildungen betrachtet, ist es im Kontext der linearen Algebra vernünftig, sich auf lineare Abbildungen einzuschränken.

Etwas konkreter: Sei K ein Körper und seien V und W endlich-dimensionale K -Vektorräume. Ist $f: V \rightarrow W$ eine *bijektive lineare* Abbildung, dann folgt $\dim V = \dim W$. Andererseits gibt es für unendliche Körper K und endlich-dimensionale Vektorräume V und W der Dimension > 0 immer eine Bijektion $V \rightarrow W$ – ohne die Linearität ist es aber unmöglich, damit die Vektorräume V und W in nützlicher Weise zueinander in Beziehung zu setzen.

- Wir haben in Abschnitt 5.3.3 zu einer Matrix $A \in M_{m \times n}(K)$ über einem Körper K die Abbildung $f_A: K^n \rightarrow K^m, x \mapsto Ax$, betrachtet. Es ist wichtig, eine einfache Charakterisierung dafür zu haben, welche Abbildungen $K^n \rightarrow K^m$ von dieser Form für eine geeignete Matrix A sind. Es ist offensichtlich, dass alle Abbildungen f_A lineare Abbildungen im Sinne von Definition 7.1 sind, und nicht sehr schwierig zu sehen, dass genau die *linearen* Abbildungen $K^n \rightarrow K^m$ die Form f_A haben. Mit dieser Korrespondenz zwischen Matrizen und linearen Abbildungen werden wir uns ausführlich in Abschnitt 7.3 befassen.

DEFINITION 7.1. Sei K ein Körper. Seien V, W Vektorräume über K . Wir nennen eine Abbildung $f: V \rightarrow W$ eine *lineare Abbildung* (oder *K -lineare Abbildung*, oder (*Vektorraum*-)*Homomorphismus*) von V nach W , falls

- (a) für alle $v, v' \in V$ gilt $f(v + v') = f(v) + f(v')$, und
- (b) für alle $a \in K, v \in V$ gilt $f(av) = af(v)$.

–

BEISPIEL 7.2. (1) Sind V, W Vektorräume über K , so verstehen wir unter der *Nullabbildung* $f: V \rightarrow W$ die konstante Abbildung $v \mapsto 0$. Dies ist eine lineare Abbildung.

(2) Ist V ein K -Vektorraum, so ist die identische Abbildung id_V eine lineare Abbildung.

(3) Ist V ein K -Vektorraum und $U \subseteq V$, so ist die Inklusionsabbildung $U \rightarrow V, u \mapsto u$, eine lineare Abbildung.

- (4) Seien $m, n \geq 0$ und $A \in M_{m \times n}(K)$ eine Matrix. Wir haben in Definition 5.43 die Abbildung $f_A: K^n \rightarrow K^m, x \mapsto Ax$, definiert. Weil das Matrizenprodukt das Distributivgesetz erfüllt und mit der Multiplikation mit Skalaren aus K kompatibel ist ($A(ax) = a(Ax)$), ist dies eine lineare Abbildung. In Satz 7.26 werden wir sehen, dass jede lineare Abbildung $K^n \rightarrow K^m$ die Form f_A für eine eindeutig bestimmte Matrix $A \in M_{m \times n}(K)$ hat.

Insbesondere kennen wir mit den Abbildungen aus Beispiel 5.49 schon einige lineare Abbildungen.

◇

BEMERKUNG 7.3. Wir haben – siehe Bemerkung 5.5 – die Addition und Skalarmultiplikation in $\mathbb{R}^2$ geometrisch interpretiert. In diesem Sinne können wir die beiden Bedingungen, dass eine Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ (und analog Abbildungen $\mathbb{R}^n \rightarrow \mathbb{R}^m$) linear ist, folgendermaßen in eine geometrische Sprache übersetzen. Wir setzen der Einfachheit halber voraus, dass f bijektiv ist. (Sonst muss man in (b) auch zulassen, dass die Ecken des »Parallelogramms« alle auf einer Geraden liegen, und es ist dann lästiger zu formulieren, was man genau als Bedingung verlangen möchte. Und es geht hier ja nicht darum, eine äquivalente Definition der Linearität zu geben, sondern den Begriff mit der geometrischen Anschauung zusammenzubringen.)

- (a) Die Abbildung f bildet den Ursprung o auf den Ursprung ab,
- (b) sind o, u, v, w vier Ecken eines Parallelogramms, so sind auch $o = f(o), f(u), f(v), f(w)$ vier Ecken eines Parallelogramms, und
- (c) f bildet jede Ursprungsgerade g auf eine Ursprungsgerade ab: $f(\langle v \rangle) = \langle f(v) \rangle$, und die induzierte Abbildung zwischen diesen Geraden ist eine »Streckung«.

◇

LEMMA 7.4. Die Verkettung von linearen Abbildungen ist wieder eine lineare Abbildung.

BEWEIS. Der Beweis ist eine leichte Rechnung, zum Beispiel gilt für $f: W \rightarrow U, g: V \rightarrow W, v \in V, a \in K$:

$$(f \circ g)(av) = f(g(av)) = f(ag(v)) = af(g(v)) = a(f \circ g)(v).$$

□

BEMERKUNG 7.5. Ist f eine lineare Abbildung, so gilt $f(o) = o$ (benutze Eigenschaft (b) mit $a = o$ und irgendeinem $v \in V$, zum Beispiel $v = o$).

Manchmal ist die folgende einfache Umformulierung nützlich. Sei $f: V \rightarrow W$ eine Abbildung zwischen K -Vektorräumen V und W . Die Abbildung f ist genau dann eine lineare Abbildung, wenn $f(av + bv') = af(v) + bf(v')$ für alle $v, v' \in V, a, b \in K$. Der Beweis der Äquivalenz ist eine einfache Übungsaufgabe.

Induktiv zeigt man für jede lineare Abbildung $f: V \rightarrow W$: Für Elemente $v_1, \dots, v_n \in V, a_1, \dots, a_n \in K$ gilt $f\left(\sum_{i=1}^n a_i v_i\right) = \sum_{i=1}^n a_i f(v_i)$. ◇

Dass für jede lineare Abbildung f im Sinne unserer Definition $f(o) = o$ gelten muss, bedeutet insbesondere, dass wir diesen Begriff mit einer etwas *anderen Bedeutung* benutzen, als es in der Schule (und teilweise auch in anderen Bereichen der Mathematik) üblich ist. Ist zum Beispiel $K = \mathbb{R}$, so sind die einzigen Vektorraum-Homomorphismen $\mathbb{R}^1 \rightarrow \mathbb{R}^1$ die Abbildungen $x \mapsto ax, a \in \mathbb{R}$. Die Abbildungen $\mathbb{R} \rightarrow \mathbb{R}, x \mapsto mx + b$ für $m, b \in \mathbb{R}$, deren Funktionsgraph eine Gerade in $\mathbb{R}^2$ ist, bezeichnen wir für $b \neq 0$ in der linearen Algebra nicht als lineare Abbildungen.

DEFINITION 7.6. Seien K ein Körper und seien V und W Vektorräume über K . Wir bezeichnen mit $\text{Hom}_K(V, W)$ die Menge aller K -linearen Abbildungen $V \rightarrow W$. $\dashv$

Dann ist $\text{Hom}_K(V, W)$ ein Untervektorraum des K -Vektorraums $\text{Abb}(V, W)$ (siehe Beispiel 6.2). In der Tat liegt offenbar die Nullabbildung in $\text{Hom}_K(V, W)$, und sind f und g lineare Abbildungen $V \rightarrow W$ und ist $a \in K$, so prüft man leicht nach, dass die Abbildungen

$$(f + g): V \rightarrow W, v \mapsto f(v) + g(v), \quad \text{und} \\ (af): V \rightarrow W, v \mapsto a \cdot f(v)$$

ebenfalls linear sind. Zum Beispiel gilt

$$(f + g)(v + w) = f(v + w) + g(v + w) = f(v) + f(w) + g(v) + g(w) \\ = f(v) + g(v) + f(w) + g(w) = (f + g)(v) + (f + g)(w).$$

BEISPIEL 7.7. (I) Seien K ein Körper, I eine Menge, und sei für jedes $i \in I$ ein K -Vektorraum V_i gegeben. Dann sind für $j, j' \in I$ die Abbildungen

$$V_j \longrightarrow \bigoplus_{i \in I} V_i \longrightarrow \prod_{i \in I} V_i \longrightarrow V_{j'},$$

die gegeben sind durch

$$v_j \mapsto (0, \dots, 0, \underbrace{v_j}_{\text{bei Index } j}, 0, \dots, 0), \quad (v_i)_{i \in I} \rightarrow (v_i)_{i \in I}, \quad (v_i)_{i \in I} \rightarrow v_{j'}$$

lineare Abbildungen.

(2) Der Grundkörper K spielt in der Definition von linearen Abbildungen durch die Skalarmultiplikation eine Rolle: Die Abbildung $\mathbb{C} \rightarrow \mathbb{C}, a + bi \mapsto a - bi$ (mit $a, b \in \mathbb{R}$) ist ein Homomorphismus von $\mathbb{R}$ -Vektorräumen (mit anderen Worten: es handelt sich um eine lineare Abbildung, wenn wir $\mathbb{C}$ als $\mathbb{R}$ -Vektorraum betrachten), aber kein Homomorphismus von $\mathbb{C}$ -Vektorräumen (denn die Abbildung ist nicht kompatibel mit Skalarmultiplikation mit beliebigen Elementen aus $\mathbb{C}$).

(3) Seien K ein Körper, V ein K -Vektorraum und $v \in V, v \neq 0$. Die Abbildung $t_v: V \rightarrow V, w \mapsto w + v$, heißt die *Verschiebung* oder *Translation* mit (oder um) v .

Dies ist *keine* lineare Abbildung, wie man zum Beispiel schon daran sieht, dass $t_v(0) = v \neq 0$.

◇

Für den besonders wichtigen Fall, dass Definitions- und Wertebereich einer linearen Abbildung übereinstimmen, vereinbaren wir eine eigene Benennung:

DEFINITION 7.8. Sei V ein K -Vektorraum. Ein *Endomorphismus* von V ist eine lineare Abbildung $V \rightarrow V$. $\dashv$

Ein anderer wichtiger Fall ist derjenige, dass zu einer linearen Abbildung $f: V \rightarrow W$ eine Umkehrabbildung g existiert. Wir setzen in der folgenden Definition voraus, dass die Umkehrabbildung g ebenfalls linear ist. Man kann dann mit f und g sozusagen ohne Informationsverlust zwischen V und W hin- und herwechseln; siehe Lemma 7.11 für einige beispielhafte konkrete Aussagen dazu. Dieses Konzept des Isomorphismus ist in sehr vielen Situationen von großer Wichtigkeit, nicht nur bei Vektorräumen. Im Fall von Vektorräumen sehen wir in Lemma 7.10, dass jede bijektive lineare Abbildung automatisch eine *lineare* Abbildung als Umkehrabbildung hat, also ein Isomorphismus ist.

DEFINITION 7.9. Sei $f: V \rightarrow W$ ein Vektorraum-Homomorphismus. Wir nennen f einen (Vektorraum-)Isomorphismus, wenn ein Vektorraum-Homomorphismus $g: W \rightarrow V$ existiert mit $g \circ f = \text{id}_V, f \circ g = \text{id}_W$.

Wir sagen, K -Vektorräume V und W seien *isomorph*, wenn ein Isomorphismus $f: V \rightarrow W$ existiert. $\dashv$

Um anzuzeigen, dass eine lineare Abbildung $V \rightarrow W$ ein Isomorphismus ist, macht man manchmal eine Tilde über den Pfeil: $V \xrightarrow{\sim} W$. Um anzugeben, dass Vektorräume V und W isomorph sind, ohne dem Isomorphismus zwischen V und W eine Bezeichnung zu geben, schreibt man manchmal $V \cong W$. Es ist leicht zu sehen, dass die Verkettung von zwei Isomorphismen wieder ein Isomorphismus ist. Ist f ein Isomorphismus, so bezeichnen wir die Umkehrabbildung von f oft mit f^{-1} .

Ist V ein Vektorraum, so nennt man einen Isomorphismus $V \rightarrow V$ von V mit sich selbst auch einen *Automorphismus* von V .

LEMMA 7.10. Ein Vektorraum-Homomorphismus $f: V \rightarrow W$ ist genau dann ein Vektorraum-Isomorphismus, wenn er bijektiv ist.

BEWEIS. Wenn f ein Vektorraum-Isomorphismus ist, dann existiert ein Umkehrhomomorphismus g wie in der Definition. Dies ist insbesondere eine Umkehrabbildung der Abbildung f (in dem Sinne, dass wir die Vektorraum-Struktur vergessen und nur die Mengen V und W betrachten). Daher ist f bijektiv (vgl. Satz 3.31).

Sei nun f ein bijektiver Vektorraum-Homomorphismus. Nach Satz 3.31 hat die Abbildung f (von Mengen) eine Umkehrabbildung g . Wir müssen zeigen, dass g ein Vektorraum-Homomorphismus ist. Seien $w_1, w_2 \in W, v_i := g(w_i)$. Dann gilt $f(v_i) = w_i$, da f und g Umkehrabbildungen voneinander sind. Also folgt

$$g(w_1 + w_2) = g(f(v_1) + f(v_2)) = g(f(v_1 + v_2)) = v_1 + v_2 = g(w_1) + g(w_2)$$

und für alle $a \in K$

$$g(aw_1) = g(af(v_1)) = g(f(av_1)) = av_1 = ag(w_1).$$

Somit ist g tatsächlich eine lineare Abbildung. $\square$

Der Begriff des Isomorphismus ist von großer Bedeutung für die lineare Algebra, weil Isomorphismen die allermeisten Eigenschaften von Vektorräumen, für die wir uns interessieren, erhalten. Zum Beispiel:

LEMMA 7.11. Sei $f: V \rightarrow W$ ein Isomorphismus von K -Vektorräumen.

- (1) Ist $(v_i)_{i \in I}$ ein Erzeugendensystem von V , so ist $(f(v_i))_{i \in I}$ ein Erzeugendensystem von W .
- (2) Ist $(v_i)_{i \in I}$ eine linear unabhängige Familie von Vektoren in V , so ist $(f(v_i))_{i \in I}$ eine linear unabhängige Familie von Vektoren in W .
- (3) Ist $(v_i)_{i \in I}$ eine Basis von V , so ist $(f(v_i))_{i \in I}$ eine Basis von W .
- (4) Es gilt $\dim(V) = \dim(W)$.

BEWEIS. Teil (3) folgt aus (1) und (2), und Teil (4) folgt aus (3). Für den Beweis der Teile (1) und (2) benutzt man, dass $f(\sum_i a_i v_i) = \sum_i a_i f(v_i)$ (für endliche Summen, bzw. Summen in denen alle bis auf endlich viele Terme = 0 sind). In Teil (1) benutzt man dann noch die Surjektivität von f , und für Teil (2) die Injektivität.

Als Beispiel führen wir Teil (2) aus: Sei $\sum a_i f(v_i) = 0$. Dann ist

$$f\left(\sum a_i v_i\right) = \sum a_i f(v_i) = 0.$$

Wegen der Injektivität und weil jedenfalls $f(0) = 0$ gilt, folgt $\sum a_i v_i = 0$. Da die v_i nach Voraussetzung linear unabhängig sind, folgt, dass alle Koeffizienten a_i gleich Null sein müssen. $\square$

In den Teilen (1), (2) und (3) des Lemmas gilt jeweils auch die Umkehrung, wie wir sehen, wenn wir das Lemma auf den Umkehrhomomorphismus f^{-1} anwenden. Eine Abbildung der Form $f_A: K^n \rightarrow K^n$ für $A \in M_n(K)$ ist genau dann bijektiv, wenn das lineare Gleichungssystem $Ax = b$ für alle $b \in K^n$ eindeutig lösbar ist, also wenn die Matrix A invertierbar ist. Wir sehen also: f_A ist genau dann ein Isomorphismus, wenn A invertierbar ist. Die Umkehrabbildung ist dann durch $f_{A^{-1}}$ gegeben (siehe Satz 5.48).

BEMERKUNG 7.12. Wir hatten in Beispiel 6.2 gesagt, dass wir »den« Vektorraum mit nur einem einzigen Element als »den« Nullvektorraum bezeichnen. Aber streng genommen gibt es ganz viele K -Vektorräume, die nur ein einziges Element haben – wir sind ja völlig frei, was wir als zugrundeliegende Menge wählen, solange nur die Verknüpfungen, die wir definieren, die Vektorraumaxiome erfüllen. Wir können nun präzisieren, warum diese Vieldeutigkeit im Fall des Nullvektorraums unproblematisch ist und es deshalb zulässig ist, von *dem einen* Nullvektorraum zu sprechen.

Seien V und V' zwei K -Vektorräume mit jeweils nur einem Element. (Mit anderen Worten: Es sei $\dim V = \dim V' = 0$.) Dann gibt es einen *eindeutig bestimmten* Isomorphismus $V \rightarrow V'$.

Die Aussage selbst ist klar, denn die einzige Abbildung zwischen V und V' bildet das eine Element von V auf das eine Element von V' ab; dies ist offenbar ein Vektorraumisomorphismus. Wegen dieser Aussage können wir alle K -Vektorräume der Dimension 0 miteinander identifizieren, und zwar auf eindeutige Art und Weise. $\diamond$

BEMERKUNG 7.13. Wir können an dieser Stelle auch die Bemerkung 6.51 noch einmal besser formulieren, weil wir besser ausdrücken können, was wir dort unter *identifizieren* verstehen: Untervektorräume $U, W \subseteq V$ bilden genau dann eine (innere) direkte Summe, wenn der Vektorraum-Homomorphismus $U \oplus W \rightarrow U + W$, $(u, w) \mapsto u + w$, von der (äußeren) direkten Summe in die Summe innerhalb von V ein Isomorphismus ist. $\diamond$

SATZ 7.14. Seien V und W Vektorräume über dem Körper K , und sei $b_1, \dots, b_n$ eine Basis von V , $n = \dim V \in \mathbb{N}$.

Sind $w_1, \dots, w_n \in W$ beliebige Elemente, dann gibt es genau eine lineare Abbildung $f: V \rightarrow W$, so dass $f(b_i) = w_i$ für alle $i = 1, \dots, n$.

Wir können also die Gleichheit von zwei linearen Abbildungen $V \rightarrow W$ dadurch zeigen, dass wir überprüfen, dass sie auf allen Basisvektoren einer fixierten Basis von V dasselbe Bild haben. Andererseits können wir auf den Basisvektoren beliebige Elemente von W als Bild vorgeben und diese Auswahl mit einer linearen Abbildung $V \rightarrow W$ »realisieren«.

BEWEIS. Wir zeigen zuerst die Eindeutigkeit. Seien $f, g: V \rightarrow W$ lineare Abbildungen mit $f(b_i) = g(b_i)$ für alle i . Ist $v \in V$ irgendein Element, so existieren $a_i \in K$ mit $v = \sum_{i=1}^n a_i b_i$. Dann gilt

$$f(v) = f\left(\sum_{i=1}^n a_i b_i\right) = \sum_{i=1}^n a_i f(b_i) = \sum_{i=1}^n a_i g(b_i) = g\left(\sum_{i=1}^n a_i b_i\right) = g(v)$$

wegen der Linearität von f und g .

Seien nun $w_1, \dots, w_n \in W$ gegeben. Wir definieren die Abbildung $f: V \rightarrow W$ durch

$$f\left(\sum_{i=1}^n a_i b_i\right) = \sum_{i=1}^n a_i w_i.$$

Da $b_1, \dots, b_n$ eine Basis von V bilden, besitzt jedes Element eine *eindeutige Darstellung* der Form $\sum_{i=1}^n a_i b_i$, wir erhalten in dieser Art und Weise also tatsächlich eine Abbildung $V \rightarrow W$. Es ist klar, dass $f(b_j) = w_j$ für alle j gilt, denn in diesem Fall sind die Koeffizienten a_i Null für $i \neq j$, und $a_j = 1$.

Man prüft unmittelbar nach, dass die Abbildung linear ist, zum Beispiel für die Additivität:

$$\begin{aligned} f\left(\sum_{i=1}^n a_i b_i + \sum_{i=1}^n a'_i b_i\right) &= f\left(\sum_{i=1}^n (a_i + a'_i) b_i\right) = \sum_{i=1}^n (a_i + a'_i) w_i \\ &= \sum_{i=1}^n a_i w_i + \sum_{i=1}^n a'_i w_i = f\left(\sum_{i=1}^n a_i b_i\right) + f\left(\sum_{i=1}^n a'_i b_i\right). \end{aligned}$$

□

Weil der Satz so wichtig ist, haben wir ihn in dieser »konkreten« Form angegeben, die für uns die relevanteste ist. Er gilt aber auch für Basen mit unendlich vielen Elementen, und der Beweis ist praktisch derselbe:

SATZ 7.15. Seien V und W Vektorräume über dem Körper K , und sei $(b_i)_{i \in I}$ eine Basis von V .

Sind $w_i \in W, i \in I$, beliebige Elemente, dann gibt es genau eine lineare Abbildung $f: V \rightarrow W$, so dass $f(b_i) = w_i$ für alle $i \in I$ gilt.

BEWEIS. Den Beweis erhält man, indem man im vorherigen Satz alle Summen der Form $\sum_{i=1}^n$ durch Summen $\sum_{i \in I}$ ersetzt, in denen immer nur endlich viele Terme ungleich Null sind. □

BEISPIEL 7.16. Sei V ein Vektorraum, und sei $(v_i)_{i \in I}$ eine (endliche oder unendliche) Familie von Vektoren in V . Sei $K^{(I)}$ der Vektorraum $\bigoplus_{i \in I} K$. Wie wir in Beispiel 6.52 gesehen haben, bilden die Vektoren $e_i, i \in I$, die jeweils an der Stelle i eine 1, und in allen anderen Einträgen eine Null haben, eine Basis. Nach dem Satz erhalten wir eine lineare Abbildung

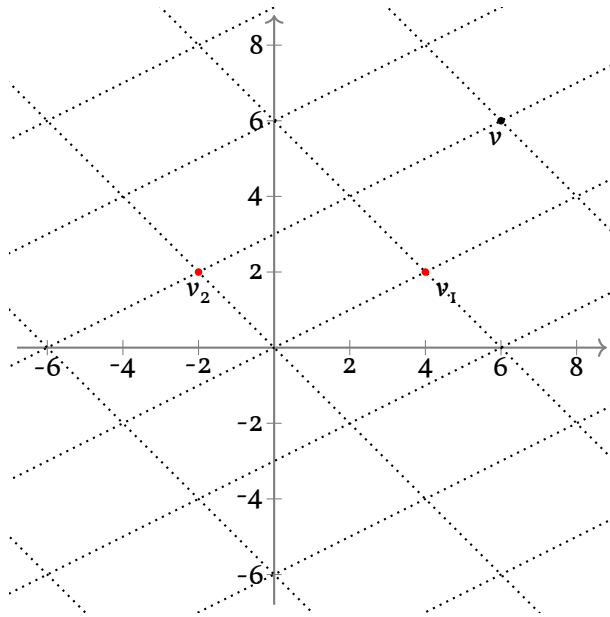
$$f: K^{(I)} \rightarrow V, \quad (x_i)_i \mapsto \sum_{i \in I} x_i v_i,$$

wobei wir die Summe auf der rechten Seite bilden können, weil alle bis auf endlich viele x_i gleich Null sind.

Wir können Eigenschaften der Familie $(v_i)_{i \in I}$ durch Eigenschaften der Abbildung f ausdrücken:

- (1) Die Familie $(v_i)_{i \in I}$ ist genau dann ein Erzeugendensystem von V , wenn die Abbildung f surjektiv ist.
- (2) Die Familie $(v_i)_{i \in I}$ ist genau dann linear unabhängig, wenn die Abbildung f injektiv ist.
- (3) Die Familie $(v_i)_{i \in I}$ ist genau dann eine Basis von V , wenn die Abbildung f bijektiv, und damit ein Isomorphismus, ist.

Ist andererseits V ein Vektorraum und $f: K^{(I)} \rightarrow V$ ein Isomorphismus, so bilden die Elemente $f(e_i), i \in I$, nach Lemma 7.11 eine Basis von V .



Die Wahl einer Basis in V ist insofern »dasselbe« wie die Wahl eines Isomorphismus zwischen einem Standardvektorraum $K^{(I)}$ und V . Für endlich erzeugte Vektorräume V kommen wir mit den Indexmengen der Form $\{1, \dots, n\}$ aus (und $n = \dim V$). Wir können die Wahl einer Basis $v_1, \dots, v_n$ in V als die Wahl eines *Koordinatensystems* sehen. Der Vektor $v = \sum_{i=1}^n a_i v_i$ hat die Koordinaten $(a_1, \dots, a_n)$.

Das ist nicht selten auch interessant, wenn V selbst ein Standardvektorraum K^n ist, man aber eine andere Basis als die Standardbasis betrachtet; das bedeutet, das Koordinatensystem zu wechseln. In der Abbildung hat der Vektor v bezüglich der Basis (v_1, v_2) von $\mathbb{R}^2$ die Koordinaten $(2, 1)$ – wir sagen, der »Koordinatenvektor« von v sei $\begin{pmatrix} 2 \\ 1 \end{pmatrix}$ (Definition 7.18).

◇

Wir können den Satz noch etwas abstrakter umformulieren, und schließen dabei gleich den Fall von Basen mit unendlich vielen Elementen ein. Für einen Vektorraum W und eine Menge I bezeichnet W^I wie üblich das Produkt $\prod_{i \in I} W$. Für den Fall $I = \{1, \dots, n\}$ ist $W^I = W^n$ der Vektorraum aller n -Tupel von Elementen aus W .

SATZ 7.17. Sei V ein K -Vektorraum mit Basis $(v_i)_{i \in I}$. Dann ist die Abbildung

$$\Phi: \text{Hom}_K(V, W) \rightarrow W^I, \quad f \mapsto (f(v_i))_{i \in I}$$

ein Isomorphismus von K -Vektorräumen.

BEWEIS. Die Bijektivität der Abbildung f ist eine Umformulierung von Satz 7.15. Wegen Lemma 7.10 bleibt nur noch zu zeigen, dass Φ linear ist. Für $f, g \in \text{Hom}_K(V, W)$ und $a, b \in K$ gilt aber

$$\Phi(af + bg) = ((af + bg)(v_i))_i = (af(v_i) + bg(v_i))_i = a\Phi(f) + b\Phi(g),$$

wobei die erste Gleichheit aus der Definition von Φ , die zweite aus der Definition von Addition und Skalarmultiplikation in $\text{Hom}_K(V, W)$, und die dritte aus der Definition der Vektorraumstruktur auf W^I folgt. □

Wir haben in Beispiel 7.16 gesehen, dass für einen Vektorraum V und eine Basis $v_1, \dots, v_n$ von V die Abbildung $K^n \rightarrow V, (a_1, \dots, a_n)^t \mapsto \sum_{i=1}^n a_i v_i$ ein Isomorphismus von Vektorräumen ist. Die Umkehrabbildung dieses Isomorphismus nennen wir die Koordinatenabbildung:

DEFINITION 7.18. Sei V ein K -Vektorraum mit Basis $\mathcal{B} = (v_1, \dots, v_n)$. Die eindeutig bestimmte lineare Abbildung $c_{\mathcal{B}}: V \rightarrow K^n$ mit $c_{\mathcal{B}}(v_i) = e_i$ heißt die *Koordinatenabbildung* zur Basis $\mathcal{B}$.

Für $v \in V$ heißt der Vektor $c_{\mathcal{B}}(v) \in K^n$ der *Koordinatenvektor* von v bezüglich der Basis $\mathcal{B}$. ▯

Wie gehabt bezeichnen wir mit $e_1, \dots, e_n$ die Standardbasisvektoren des Vektorraums K^n . Die Existenz und Eindeutigkeit der Koordinatenabbildung folgt aus Satz 7.14.

KOROLLAR 7.19. *Seien K ein Körper und V, W endlich erzeugte K -Vektorräume. Dann sind äquivalent:*

- (i) V und W sind isomorph.
- (ii) $\dim V = \dim W$.

BEWEIS. Die Implikation (i) $\Rightarrow$ (ii) haben wir im Lemma 7.11 gesehen. Sei nun $n = \dim V = \dim W$. (Da V und W nach Voraussetzung endlich erzeugt sind, handelt es sich hier um eine natürliche Zahl.) Es gibt also Basen $\mathcal{B} = (v_1, \dots, v_n)$ von V und $\mathcal{C} = (w_1, \dots, w_n)$ von W . Die zugehörigen Koordinatenabbildungen $c_{\mathcal{B}}: V \rightarrow K^n$ und $c_{\mathcal{C}}: W \rightarrow K^n$ sind Isomorphismen, und die Verkettung $c_{\mathcal{C}}^{-1} \circ c_{\mathcal{B}}$ ist ein Isomorphismus $V \rightarrow W$.

Alternativ kann man direkt Satz 7.14 anwenden, um Abbildungen $f: V \rightarrow W$ mit $f(v_i) = w_i$ und $g: W \rightarrow V$ mit $g(w_i) = v_i$ zu definieren. Es gilt dann $g \circ f = \text{id}_V, f \circ g = \text{id}_W$, denn es genügt, das jeweils auf einer Basis zu überprüfen, also sind f und g Isomorphismen. $\square$

Insbesondere ist die Frage, ob zwei Vektorräume isomorph sind, relativ leicht zu entscheiden (und daher oft nicht so interessant). Aber: Natürlich ist nicht jede lineare Abbildung zwischen Vektorräumen derselben Dimension ein Isomorphismus, und einen konkreten Isomorphismus zwischen zwei Vektorräumen zu »haben« ist eine wesentlich stärkere/nützlichere Information, als nur zu wissen, dass zwei Vektorräume isomorph sind.

7.2. Kern und Bild einer linearen Abbildung

Sei $f: V \rightarrow W$ ein K -Vektorraumhomomorphismus.

DEFINITION 7.20. Der *Kern* von f ist definiert als

$$\text{Ker}(f) := f^{-1}(\{0\}) = \{v \in V; f(v) = 0\}.$$

—

Wie bei jeder Abbildung, so haben wir auch für die lineare Abbildung f den Begriff des Bildes $\text{Im}(f)$: $\text{Im}(f) = \{f(v); v \in V\} \subseteq W$.

LEMMA 7.21. *Für jede lineare Abbildung $f: V \rightarrow W$ ist $\text{Ker}(f)$ ein Untervektorraum von V und $\text{Im}(f)$ ein Untervektorraum von W .*

BEWEIS. Weil $f(0) = 0$ ist, ist $0 \in \text{Ker}(f)$. Sind $v, v' \in \text{Ker}(f)$, so gilt $f(v + v') = f(v) + f(v') = 0 + 0 = 0$, also $v + v' \in \text{Ker}(f)$. Sind $v \in \text{Ker}(f)$ und $a \in K$, so gilt $f(av) = af(v) = a \cdot 0 = 0$, also $av \in \text{Ker}(f)$.

Wir zeigen nun die Behauptung für $\text{Im}(f)$. Es gilt $f(0) = 0$, also $0 \in \text{Im}(f)$. Sind $w, w' \in \text{Im}(f)$, so existieren $v, v' \in V$ mit $w = f(v), w' = f(v')$. Dann gilt

$$w + w' = f(v) + f(v') = f(v + v') \in \text{Im}(f)$$

wegen der Linearität von f . Für $w = f(v) \in \text{Im}(f)$ und $a \in K$ erhalten wir entsprechend $aw = af(v) = f(av) \in \text{Im}(f)$. $\square$

SATZ 7.22. *Die lineare Abbildung $f: V \rightarrow W$ ist genau dann injektiv, wenn $\text{Ker}(f) = \{0\}$.*

BEWEIS. Wenn f injektiv ist, kann es höchstens ein Element von V geben, das auf $0 \in W$ abgebildet wird. Weil jedenfalls $f(0) = 0$ gilt, folgt $\text{Ker}(f) = \{0\}$. Ist andererseits $\text{Ker}(f) = \{0\}$ und gilt $f(v) = f(v')$, so folgt $f(v - v') = f(v) - f(v') = 0$, also $v - v' \in \text{Ker}(f) = \{0\}$, das heißt $v = v'$. $\square$

Eine *injektive lineare* Abbildung $V \rightarrow W$ nennt man auch einen *Monomorphismus*. Eine *surjektive lineare* Abbildung $V \rightarrow W$ nennt man auch einen *Epimorphismus*.

Für eine Matrix A gilt $\text{Ker}(A) = \text{Ker}(f_A)$, $\text{Im}(A) = \text{Im}(f_A)$. Nun ist $\text{Ker}(A)$ gerade die Lösungsmenge des durch A gegebenen linearen Gleichungssystems, und $\text{Im}(A)$ ist der Teilraum derjenigen Vektoren b , für die das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$ lösbar ist. Wir können also die hier gegebenen Definitionen von Kern und Bild einer linearen Abbildung als (weitreichende) Verallgemeinerungen dieser Konzepte aus der Theorie der linearen Gleichungssysteme betrachten. Andererseits liefert die abstrakte Sichtweise auch Erkenntnisse über lineare Gleichungssysteme: Das folgende Theorem, die *Dimensionsformel für lineare Abbildungen*, gibt eine präzise und sehr elegante Antwort auf die in Frage 5.27 (2) formulierte Frage, siehe auch Abschnitt 7.4.

THEOREM 7.23 (Dimensionsformel für lineare Abbildungen). *Sei $f: V \rightarrow W$ eine lineare Abbildung zwischen K -Vektorräumen und sei V endlich-dimensional. Dann gilt:*

$$\dim V = \dim \text{Ker } f + \dim \text{Im } f.$$

Die Zahl $\dim \text{Im } f$ heißt auch der *Rang von f* , in Zeichen: $\text{rg}(f)$.

BEWEIS. Sei $U \subseteq V$ ein Komplementärraum von $\text{Ker}(f)$. Wir bezeichnen die Einschränkung von f auf U mit $f|_U$. Ihr Bild liegt natürlich in $\text{Im}(f)$. Wir zeigen gleich, dass $f|_U: U \rightarrow \text{Im}(f)$ ein Isomorphismus ist. Daraus folgt jedenfalls der Satz, denn es folgt $\dim(U) = \dim \text{Im}(f)$ und damit $\dim V = \dim \text{Ker}(f) + \dim U = \dim \text{Ker}(f) + \dim \text{Im}(f)$ (benutze Satz 6.46 oder Korollar 6.54 und Lemma 7.II).

Um zu zeigen, dass $f|_U: U \rightarrow \text{Im}(f)$ ein Isomorphismus ist, zeigen wir die Injektivität und die Surjektivität.

Injektivität. Ist $u \in U$, $f|_U(u) = 0$, so gilt $u \in U \cap \text{Ker}(f) = 0$, also $u = 0$.

Surjektivität. Sei $w \in \text{Im}(f)$. Dann existiert $v \in V$ mit $f(v) = w$. Wir schreiben $v = v' + u$ mit $v' \in \text{Ker}(f)$, $u \in U$ und erhalten

$$f|_U(u) = f(v - v') = f(v) - f(v') = w.$$

□

KOROLLAR 7.24. *Seien V, W endlich-dimensionale K -Vektorräume mit $\dim V = \dim W$. Ferner sei $f: V \rightarrow W$ eine lineare Abbildung. Dann sind äquivalent:*

- (i) f ist ein Isomorphismus,
- (ii) f ist injektiv,
- (iii) f ist surjektiv.

BEWEIS. Wir schreiben $d = \dim(V) = \dim(W)$, $d' = \dim \text{Ker}(f)$ und $d'' = \dim \text{Im}(f)$. Dann gilt $0 \leq d', d'' \leq d$ und die Dimensionsformel besagt $d' + d'' = d$. Daraus folgt die Äquivalenz

$$d' = 0 \text{ und } d'' = d \iff d' = 0 \iff d'' = d.$$

Das Korollar folgt nun daraus, dass $d' = 0$ gleichbedeutend damit ist, dass $\text{Ker}(f) = 0$, also dass f injektiv ist, und dass $d'' = d$ bedeutet, dass $\text{Im}(f) = W$, also dass f surjektiv ist. □

Beachten Sie die Analogie zu Satz 3.64 der besagt, dass eine Abbildung zwischen endlichen Mengen mit gleich vielen Elementen genau dann injektiv ist, wenn sie surjektiv ist.

Wir skizzieren noch einen etwas anderen Beweis des Korollars, der direkt Theorem 6.43 und das folgende einfache Lemma benutzt.

LEMMA 7.25. Sei $f: V \rightarrow W$ ein Vektorraum-Homomorphismus. Seien $v_1, \dots, v_n \in V$ linear unabhängig. Wir schreiben $w_i := f(v_i)$. Dann sind äquivalent:

- (i) Die Abbildung f ist injektiv.
- (ii) Die Familie $w_1, \dots, w_n$ ist linear unabhängig.

Sei nun $f: V \rightarrow W$ wie im Korollar ein Homomorphismus zwischen Vektorräumen derselben Dimension n , und sei $v_1, \dots, v_n$ eine Basis. Ist f injektiv, so sind die Bilder $f(v_i)$ nach dem Lemma ebenfalls linear unabhängig, bilden also nach Theorem 6.43 eine Basis. Damit enthält $\text{Im}(f)$ ein Erzeugendensystem, f ist folglich surjektiv. Ist andererseits f surjektiv, so bilden die $f(v_i)$, die offenbar das Bild von f erzeugen, ein Erzeugendensystem von W , das aus $\dim(W)$ Elementen besteht, also eine Basis. Nach dem Lemma ist f injektiv.

Für Abbildungen der Form f_A für eine Matrix A folgt der Satz auch unmittelbar aus Korollar 5.22 (und andersherum erhalten wir mit dem obigen Satz einen neuen Beweis dieses Korollars).

7.3. Der Zusammenhang zwischen linearen Abbildungen und Matrizen

In diesem Abschnitt studieren wir den Zusammenhang zwischen linearen Abbildungen und Matrizen. Das sieht auf den ersten Blick nach furchtbar viel Rechnerei aus (besonders im zweiten Teil des Abschnitts, wenn wir die Matrizen $M_{\mathcal{C}}^{\mathcal{B}}(f)$ definieren und untersuchen, was bei »Basiswechsel« passiert). Diese Überlegungen sind aber wichtig für das strukturelle Verständnis: Siehe Bemerkung 7.34.

Zuerst betrachten wir lineare Abbildungen $K^n \rightarrow K^m$ zwischen Standardvektorräumen über einem Körper K . Wir kennen schon die Abbildungen der Form $f_A, x \mapsto Ax$ für Matrizen $A \in M_{m \times n}(K)$. Mit Satz 7.14 sehen wir leicht, dass jede lineare Abbildung $f: K^n \rightarrow K^m$ die Form f_A für eine eindeutig bestimmte Matrix $A \in M_{m \times n}(K)$ hat. Wir formulieren das als den folgenden

SATZ 7.26. Seien $m, n \geq 0$. Die Abbildung

$$M_{m \times n}(K) \rightarrow \text{Hom}_K(K^n, K^m), \quad A \mapsto (f_A: x \mapsto Ax),$$

ist ein Isomorphismus mit Umkehrabbildung

$$f \mapsto M(f) := (a_{ij})_{i,j},$$

wobei

$$f(e_j) = (a_{1j}, \dots, a_{mj})^t \in K^m.$$

Hier bezeichnet das e_j links den j -ten Standardbasisvektor in K^n , dem Definitionsbereich von f .

Die j -te Spalte von $M(f)$ ist also der Vektor $f(e_j) \in K^m$.

BEWEIS. Der Beweis ist im Grunde nicht schwierig, sondern besteht eher darin, zu »organisieren«, was gezeigt werden muss, und diese Liste dann abzuarbeiten – Sie sollten versuchen, einen eigenen Beweis aufzuschreiben, bevor Sie den Beweis hier lesen.

Wir überprüfen, dass die Abbildung $A \mapsto f_A$ linear ist, und dass die beiden gegebenen Abbildungen zueinander invers sind. Die Linearität der Abbildung $f \mapsto M(f)$ folgt dann aus Lemma 7.10, sobald wir gezeigt haben, dass die Abbildungen invers zueinander sind. (Es ist natürlich auch möglich, die Linearität dieser Abbildung direkt nachzurechnen.)

Die Tatsache, dass $A \mapsto f_A$ eine lineare Abbildung ist, schreibt man unmittelbar aus; für $A, B \in M_{m \times n}(K)$, $a, b \in K$ und $x \in K^n$ gilt

$$(aA + bB)x = a(Ax) + b(Bx),$$

also $f_{aA+bB} = af_A + bf_B$.

Nun überprüfen wir, dass die Abbildungen $A \mapsto f_A$ und $f \mapsto M(f)$ zueinander invers sind. Sei zunächst $A \in M_{m \times n}(K)$. Wir wollen zeigen, dass $A = M(f_A)$. Wir schreiben dazu $A = (a_{ij})_{i,j}$ und $M(f_A) = (a'_{ij})_{i,j}$. Wir haben dann nach Definition von $M(f_A)$, dass

$$(a_{1j}, \dots, a_{mj})^t = Ae_j = f_A(e_j) = (a'_{1j}, \dots, a'_{mj})^t,$$

wobei wir noch für die erste Gleichheit benutzen, dass Ae_j gerade die j -te Spalte von A ist, und für die zweite Gleichheit die Definition von f_A benutzen. Es folgt also $a_{ij} = a'_{ij}$ für alle i und j .

Es bleibt noch zu zeigen, dass $f_{M(f)} = f$ für jede Abbildung $f: K^n \rightarrow K^m$ gilt. Nach Satz 7.14 genügt es zu zeigen, dass $f_{M(f)}(e_j) = f(e_j)$ für alle $j = 1, \dots, n$ gilt. In der Tat haben wir, wenn wir $M(f) = (a_{ij})_{i,j}$ schreiben, dass

$$f_{M(f)}(e_j) = M(f)e_j = (a_{1j}, \dots, a_{mj})^t = f(e_j)$$

nach Definition der a_{ij} . □

Man kann den Beweis der Bijektivität der Abbildung $f \mapsto M(f)$ auch »direkter« führen, indem man mithilfe von Satz 7.14 die Injektivität und Surjektivität nachweist. Es ist aber auch nützlich, die Umkehrabbildung anzugeben, wie wir es getan haben.

LEMMA 7.27. Seien $f: K^m \rightarrow K^l$ und $g: K^n \rightarrow K^m$ lineare Abbildungen und seien $M(f) \in M_{l \times m}(K)$ und $M(g) \in M_{m \times n}(K)$ die zugehörigen Matrizen im Sinne von Satz 7.26.

Dann gilt $M(f \circ g) = M(f)M(g)$.

BEWEIS. Weil die in Satz 7.26 angegebenen Abbildungen zwischen dem Raum der linearen Abbildungen und dem Raum der Matrizen bijektiv sind, ist das eine rein formale Folgerung aus Satz 5.48, der besagt, dass für Matrizen $A \in M_{l \times m}(K)$, $B \in M_{m \times n}(K)$ gilt, dass $f_A \circ f_B = f_{AB}$.

In der Tat erhalten wir so

$$f \circ g = f_{M(f)} \circ f_{M(g)} = f_{M(f)M(g)},$$

also $M(f \circ g) = M(f_{M(f)M(g)}) = M(f)M(g)$, wie gewünscht. □

Mit dem Lemma sehen wir erneut (vergleiche den Kommentar nach Lemma 7.10), dass eine Abbildung $f: K^n \rightarrow K^n$ genau dann ein Isomorphismus ist, wenn $M(f)$ eine invertierbare Matrix ist.

Wir wollen nun Satz 7.26 auf den Fall beliebiger Vektorräume V und W der Dimension n und m übertragen. Dabei fixieren wir Basen $\mathcal{B} = (v_1, \dots, v_n)$ von V und $\mathcal{C} = (w_1, \dots, w_m)$ von W . Die Koordinatenabbildungen $c_{\mathcal{B}}$ und $c_{\mathcal{C}}$ sind Isomorphismen $V \cong K^n$ und $W \cong K^m$, die wir uns als Identifizierungen vorstellen können; siehe Bemerkung 7.29.

Satz 7.26 erhalten wir zurück, wenn wir $V = K^n$, $W = K^m$ und als Basen die Standardbasen fixieren. Da wir für eine lineare Abbildung $f: V \rightarrow W$ aus den Koordinatenvektoren der $f(v_i)$ bezüglich $\mathcal{C}$ eine Matrix konstruieren, ist es hier wichtig, dass wir die Anordnung (also die Nummerierung) der Basisvektoren in den Basen $\mathcal{B}$ und $\mathcal{C}$ fixieren. Wenn wir unten sagen »Sei $\mathcal{B}$ eine Basis des endlich-dimensionalen Vektorraums V .«, so ist damit ein Tupel von Elementen aus V (mit fixierter Reihenfolge, also nicht nur eine Teilmenge von V) gemeint. Man spricht auch von einer (an-)geordneten Basis.

SATZ 7.28. Seien V, W Vektorräume der Dimensionen $n = \dim V, m = \dim W$, und seien $\mathcal{B} = (v_1, \dots, v_n), \mathcal{C} = (w_1, \dots, w_m)$ Basen von V bzw. W . Dann sind die Abbildungen

$$M_{m \times n}(K) \longrightarrow \text{Hom}_K(V, W), \quad A \mapsto c_{\mathcal{C}}^{-1} \circ f_A \circ c_{\mathcal{B}}$$

und

$$\text{Hom}_K(V, W) \longrightarrow M_{m \times n}(K), \quad f \mapsto M_{\mathcal{C}}^{\mathcal{B}}(f) := M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1})$$

zueinander inverse Vektorraumisomorphismen.

BEWEIS. Wir können die Aussage des Satzes in die folgenden Punkte aufteilen:

- (1) Die gegebene Abbildung $M_{m \times n}(K) \longrightarrow \text{Hom}_K(V, W)$ ist linear.
- (2) Die gegebene Abbildung $\text{Hom}_K(V, W) \longrightarrow M_{m \times n}(K)$ ist linear.
- (3) Die Verkettung $M_{m \times n}(K) \longrightarrow \text{Hom}_K(V, W) \longrightarrow M_{m \times n}(K)$ ist die Identitätsabbildung.
- (4) Die Verkettung $\text{Hom}_K(V, W) \longrightarrow M_{m \times n}(K) \longrightarrow \text{Hom}_K(V, W)$ ist die Identitätsabbildung.

Jeder dieser Punkte lässt sich durch einfache Überlegungen zeigen bzw. auf Satz 7.26 zurückführen, und es ist besser (und vielleicht sogar einfacher), wenn Sie sich selbst durch das Gestrüpp von Bezeichnungen (das sich an dieser Stelle leider nicht ganz vermeiden lässt) schlagen, als den Beweis hier nachzulesen.

Wegen Lemma 7.10 genügt es, einen der beiden Punkte (1), (2) zu zeigen, wenn man zuerst (3) und (4) abhandelt.

zu (3). Es gilt

$$M_{\mathcal{C}}^{\mathcal{B}}(c_{\mathcal{C}}^{-1} \circ f_A \circ c_{\mathcal{B}}) = M(c_{\mathcal{C}} \circ c_{\mathcal{C}}^{-1} \circ f_A \circ c_{\mathcal{B}} \circ c_{\mathcal{B}}^{-1}) = M(f_A) = A.$$

zu (4). Es gilt

$$c_{\mathcal{C}}^{-1} \circ f_{M_{\mathcal{C}}^{\mathcal{B}}(f)} \circ c_{\mathcal{B}} = c_{\mathcal{C}}^{-1} \circ f_{M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1})} \circ c_{\mathcal{B}} = c_{\mathcal{C}}^{-1} \circ c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1} \circ c_{\mathcal{B}} = f.$$

zu (1). Sind $\varphi: V_1 \rightarrow V_2$ und $\psi, \psi': V_2 \rightarrow V_3$ Vektorraum-Homomorphismen, so gilt $(\psi + \psi') \circ \varphi = \psi \circ \varphi + \psi' \circ \varphi$. Denn auf beiden Seiten steht die Abbildung $v_1 \mapsto \psi(\varphi(v_1)) + \psi'(\varphi(v_1))$. Entsprechend ist es, wenn die Summe zweier Abbildungen auf der rechten Seite einer Verkettung steht. Mithilfe dieser Überlegung und der Linearität der Abbildung $A \mapsto f_A$ (Satz 7.26) folgt die Kompatibilität der Abbildung mit der Addition. Die Kompatibilität mit der Skalarmultiplikation kann man in analoger (aber noch einfacherer) Weise zeigen. (Stattdessen bzw. zusätzlich kann man mit ganz ähnlichen Argumenten die Linearität in (2) zeigen.) $\square$

Mit anderen Worten: Für alle $v \in V$ gilt

$$M_{\mathcal{C}}^{\mathcal{B}}(f)c_{\mathcal{B}}(v) = c_{\mathcal{C}}(f(v)).$$

Man nennt $M_{\mathcal{C}}^{\mathcal{B}}(f)$ auch die *darstellende Matrix* von f , oder sagt, f werde bezüglich der Basen $\mathcal{B}$ und $\mathcal{C}$ durch die Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ dargestellt. Für ein Element v_j der Basis $\mathcal{B}$ ist $c_{\mathcal{B}}(v_j) = e_j \in K^n$ der j -te Standardbasisvektor, also ist $M_{\mathcal{C}}^{\mathcal{B}}(f)c_{\mathcal{B}}(v_j) = c_{\mathcal{C}}(f(v_j))$ die j -te Spalte von $M_{\mathcal{C}}^{\mathcal{B}}(f)$. Das ist die Beschreibung von $M_{\mathcal{C}}^{\mathcal{B}}(f)$, die Sie sich merken sollten:

Die j -te Spalte von $M_{\mathcal{C}}^{\mathcal{B}}(f)$ ist der Koordinatenvektor (bezüglich $\mathcal{C}$) des Bildes $f(v_j)$ vom j -ten Vektor in $\mathcal{B}$.

Mit dieser Beschreibung können wir auch die Bijektivität der Abbildung $f \mapsto M_{\mathcal{C}}^{\mathcal{B}}(f)$ noch einmal auf etwas anderem Wege begründen, und zwar indem wir die Aussage direkt auf Satz 7.14 zurückführen: Die Abbildung ist injektiv, weil die lineare Abbildung f durch die Bilder $f(v_j)$ der Vektoren in $\mathcal{B}$ eindeutig bestimmt ist. Sie ist surjektiv, weil wir als Bilder der v_j beliebige Elemente in W (und damit beliebige Koordinatenvektoren) vorgeben können.

Im Fall $V = K^n$, $W = K^m$ und wenn $\mathcal{B}$ und $\mathcal{C}$ die Standardbasen bezeichnen, dann ist $M_{\mathcal{C}}^{\mathcal{B}}(f) = M(f)$.

BEMERKUNG 7.29. Wir können die Abbildungen, die hier eine Rolle spielen, wie folgt aufschreiben:

$$\begin{array}{ccc} V & \xrightarrow{f} & W \\ \downarrow c_{\mathcal{B}} & & \downarrow c_{\mathcal{C}} \\ K^n & & K^m \end{array}$$

Die vertikalen Abbildungen sind Isomorphismen. Wir können diese Abbildungen und ihre Umkehrabbildungen benutzen um V und K^n einerseits, und W und K^m andererseits zu identifizieren. Wenn wir diese Identifizierung vornehmen, ist eine Abbildung $V \rightarrow W$ »praktisch dasselbe« wie eine Abbildung $K^n \rightarrow K^m$. Formal gesprochen können wir das obige »Diagramm« vervollständigen, indem wir φ als die Verkettung $c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1}$ definieren:

$$\begin{array}{ccc} V & \xrightarrow{f} & W \\ \downarrow c_{\mathcal{B}} & & \downarrow c_{\mathcal{C}} \\ K^n & \xrightarrow{\varphi} & K^m \end{array}$$

Die beiden Abbildungen $V \rightarrow K^m$ (die Verkettung $V \rightarrow W \rightarrow K^m$ und die Verkettung $V \rightarrow K^n \rightarrow K^m$) stimmen überein:

$$c_{\mathcal{C}} \circ f = c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1} \circ c_{\mathcal{B}} = \varphi \circ c_{\mathcal{B}}.$$

Diesen Sachverhalt drückt man auch dadurch aus, dass man von einem *kommutativen Diagramm* spricht.

Da $c_{\mathcal{C}}$ und $c_{\mathcal{B}}$ Isomorphismen sind, erhalten wir aus der Kommutativität des Diagramms auch die Gleichungen

$$f = c_{\mathcal{C}}^{-1} \circ \varphi \circ c_{\mathcal{B}}, \quad \varphi = c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1},$$

die zeigen, dass die Abbildung

$$\mathrm{Hom}_K(V, W) \rightarrow \mathrm{Hom}_K(K^n, K^m), \quad f \mapsto c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1},$$

eine Bijektion ist. Da die Abbildung auch linear ist, wie man leicht nachprüft, ist es ein Isomorphismus. Es gilt dann $M_{\mathcal{C}}^{\mathcal{B}}(f) = M(\varphi)$.

Insgesamt sehen wir, dass wir den Isomorphismus aus Satz 7.28 auch als Verkettung

$$\mathrm{Hom}_K(V, W) \xrightarrow{\sim} \mathrm{Hom}_K(K^n, K^m) \xrightarrow{\sim} M_{m \times n}(K)$$

des obigen Isomorphismus mit dem Isomorphismus aus Satz 7.26 erhalten können. $\diamond$

BEMERKUNG 7.30. Als »toy example« ist vielleicht die folgende Bemerkung nützlich: Dass man nach Wahl von Basen eine lineare Abbildung zwischen endlich-dimensionalen Vektorräumen durch eine Matrix beschreiben kann, ist im weitesten Sinne dazu analog, dass man

nach Durchnummerieren der Elemente eine Abbildung $f: X \rightarrow Y$ zwischen endlichen Mengen (mit n und m Elementen) durch eine Tabelle $\frac{1 \quad 2 \quad \dots \quad n}{a_1 \quad a_2 \quad \dots \quad a_n}$ mit $a_j \in \{1, \dots, m\}$ beschreiben kann. Das Durchnummerieren entspricht der Wahl von Bijektionen $c_X: X \rightarrow \{1, \dots, n\}$, $c_Y: \{1, \dots, m\} \rightarrow Y$. Dann ist $a_j = c_Y(f(c_X^{-1}(j))) = (c_Y \circ f \circ c_X^{-1})(j)$. Vergleichen Sie diese Vorschrift mit der Definition von $M_{\mathcal{C}}^{\mathcal{B}}(f)$. $\diamond$

Videos von 3Blue1Brown zum Thema Lineare Abbildungen und Matrizen

[Linear transformations and matrices^a](https://www.youtube.com/watch?v=kYB8IZa5AuE)

[Matrix multiplication as composition^b](https://www.youtube.com/watch?v=XkY2DOUCWMU)

[Three-dimensional linear transformations^c](https://www.youtube.com/watch?v=rHLEWRxRGiM)

^a<https://youtu.be/kYB8IZa5AuE>

^b<https://youtu.be/XkY2DOUCWMU>

^c<https://youtu.be/rHLEWRxRGiM>

SATZ 7.31. Seien $f: V \rightarrow W, g: W \rightarrow U$ Homomorphismen endlich-dimensionaler Vektorräume V, W, U mit Basen $\mathcal{B}, \mathcal{C}, \mathcal{D}$. Dann gilt

$$M_{\mathcal{D}}^{\mathcal{C}}(g)M_{\mathcal{C}}^{\mathcal{B}}(f) = M_{\mathcal{D}}^{\mathcal{B}}(g \circ f).$$

BEWEIS. Wir haben

$$\begin{aligned} M_{\mathcal{D}}^{\mathcal{C}}(g)M_{\mathcal{C}}^{\mathcal{B}}(f) &= M(c_{\mathcal{D}} \circ g \circ c_{\mathcal{C}}^{-1})M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1}) \\ &= M(c_{\mathcal{D}} \circ g \circ c_{\mathcal{C}}^{-1} \circ c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1}) = M_{\mathcal{D}}^{\mathcal{B}}(g \circ f). \end{aligned}$$

Dabei haben wir für die Gleichheit in der Mitte Lemma 7.27 benutzt. $\square$

KOROLLAR 7.32. Seien V und W endlich-dimensionale Vektorräume derselben Dimension n mit Basen $\mathcal{B}$ und $\mathcal{C}$, und sei $f: V \rightarrow W$ eine lineare Abbildung. Die Abbildung f ist genau dann ein Isomorphismus, wenn die Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ invertierbar ist, und in diesem Fall gilt

$$M_{\mathcal{C}}^{\mathcal{B}}(f)^{-1} = M_{\mathcal{C}}^{\mathcal{B}}(f^{-1}).$$

BEWEIS. Ist f invertierbar mit Umkehrhomomorphismus f^{-1} , so folgt aus dem Satz, dass

$$M_{\mathcal{C}}^{\mathcal{B}}(f)M_{\mathcal{B}}^{\mathcal{C}}(f^{-1}) = M_{\mathcal{C}}^{\mathcal{C}}(f \circ f^{-1}) = M_{\mathcal{C}}^{\mathcal{C}}(\text{id}) = E_n.$$

Ist $A := M_{\mathcal{C}}^{\mathcal{B}}(f)$ invertierbar, so erhalten wir eine Umkehrabbildung von $f = c_{\mathcal{C}}^{-1} \circ f_A \circ c_{\mathcal{B}}$ als $c_{\mathcal{B}}^{-1} \circ f_{A^{-1}} \circ c_{\mathcal{C}}$. $\square$

KOROLLAR 7.33 (Basiswechsel). Sei $f: V \rightarrow W$ ein Homomorphismus endlich-dimensionaler K -Vektorräume, seien $\mathcal{B}, \mathcal{B}'$ Basen von V und $\mathcal{C}, \mathcal{C}'$ Basen von W . Dann gilt

$$M_{\mathcal{C}'}^{\mathcal{B}'}(f) = M_{\mathcal{C}'}^{\mathcal{C}}(\text{id}_W)M_{\mathcal{C}}^{\mathcal{B}}(f)M_{\mathcal{B}}^{\mathcal{B}'}(\text{id}_V).$$

BEWEIS. Das folgt unmittelbar durch zweifache Anwendung des vorhergehenden Satzes. $\square$

Wir schreiben auch $M_{\mathcal{B}}^{\mathcal{B}'} := M_{\mathcal{B}}^{\mathcal{B}'}(\text{id}_V)$ und nennen diese Matrix die *Basiswechselmatrix*, die den Basiswechsel von $\mathcal{B}'$ nach $\mathcal{B}$ beschreibt. Es gilt $M_{\mathcal{B}}^{\mathcal{B}'}c_{\mathcal{B}'}(v) = c_{\mathcal{B}}(v)$. Die j -te Spalte von $M_{\mathcal{B}}^{\mathcal{B}'}$ ist also $c_{\mathcal{B}}(b'_j)$ (wenn wir $\mathcal{B}' = (b'_1, \dots, b'_{\dim(V)})$ schreiben).

BEMERKUNG 7.34. Die als *Basiswechsel* bezeichnete Aussage von Korollar 7.33 lässt sich zwar auch ge- bzw. missbrauchen, um Rechenaufgaben zu stellen, ihr liegt aber ein wichtiges Konzept zugrunde. Um zu verstehen, was eine lineare Abbildung »wirklich macht«, ist es oft sehr hilfreich, ein geeignetes Koordinatensystem zu wählen. Sich zu erlauben, so einen Koordinatenwechsel durchzuführen bedeutet gerade, dass man für eine lineare Abbildung einen Basiswechsel durchführen kann.

Dies ist sowohl für die geometrische Anschauung, als auch »algebraisch«, also für die Durchführung konkreter Rechnungen, und auch für weitergehende theoretische Überlegungen nützlich.

Siehe Satz 7.36 und Beispiel 7.37 für konkretere Hinweise dazu. $\diamond$

KOROLLAR 7.35.

$$M_{\mathcal{B}}^{\mathcal{B}'} M_{\mathcal{B}'}^{\mathcal{B}} = E_n.$$

Insbesondere gilt: Jede Basiswechselmatrix ist invertierbar. Umgekehrt kann jede invertierbare Matrix $A \in M_n(K)$ als Basiswechselmatrix gesehen werden. Genauer gilt: Ist V ein n -dimensionaler K -Vektorraum, $\mathcal{B}$ eine Basis von V , und $A = (a_{ij})_{i,j} \in M_n(K)$ invertierbar, so existiert genau eine Basis $\mathcal{B}' = (b'_1, \dots, b'_n)$ von V , so dass $A = M_{\mathcal{B}}^{\mathcal{B}'}$, nämlich

$$b'_j = \sum_{i=1}^n a_{ij} b_i.$$

Sei $f: V \rightarrow W$ eine lineare Abbildung zwischen endlich-dimensionalen Vektorräumen V und W . Wenn wir uns erlauben, in V und W geeignete Basen zu wählen, dann können wir f immer durch eine Matrix einer sehr einfachen Form darstellen:

SATZ 7.36 (Smith-sche Normalform). Sei $f: V \rightarrow W$ ein Homomorphismus zwischen endlich-dimensionalen Vektorräumen, $n = \dim V$, $m = \dim W$. Dann gibt es Basen $\mathcal{B}$ von V und $\mathcal{C}$ von W , so dass

$$M_{\mathcal{C}}^{\mathcal{B}}(f) = \begin{pmatrix} E_r & 0 \\ 0 & 0 \end{pmatrix}$$

(als Blockmatrix verstanden). Dabei ist $r = \text{rg}(f)$ der Rang von f .

BEWEIS. Vielleicht können Sie den Beweis selber finden? Versuchen Sie das erst einmal. Es gibt (mindestens) zwei Möglichkeiten, entweder mit der »Basiswechseltheorie«, die wir gerade behandeln, oder mit einem ziemlich elementaren Argument mit Matrizen.

Wir suchen Basen $b_1, \dots, b_n$ und $c_1, \dots, c_m$ von V bzw. W , so dass

$$(2) \quad f(b_j) = \begin{cases} c_j & j \leq r \\ 0 & j > r \end{cases}.$$

Für $j > r$ soll also $b_j \in \text{Ker}(f)$ gelten, und das führt auf das folgende Vorgehen:

Sei $U \subseteq V$ ein Komplement von $\text{Ker}(f)$, sei $r = \dim(U)$, sei $b_1, \dots, b_r$ eine Basis von U und sei $b_{r+1}, \dots, b_n$ eine Basis von $\text{Ker}(f)$. Dann ist $\mathcal{B} = (b_1, \dots, b_n)$ eine Basis von V . Für $j = 1, \dots, r$ sei $c_j = f(b_j)$. Nach Definition von U gilt $\text{Ker}(f|_U) = U \cap \text{Ker}(f) = 0$, also ist die Einschränkung $f|_U: U \rightarrow \text{Im}(f)$ ein Isomorphismus. Insbesondere ist das System $c_1, \dots, c_r$ linear unabhängig, und wir können es zu einer Basis $\mathcal{C} = (c_1, \dots, c_m)$ von W ergänzen.

Offenbar gilt dann die Beschreibung (2). Außerdem ist $r = \dim(U) = \dim(V) - \dim \text{Ker}(f) = \text{rg}(f)$ nach der Dimensionsformel für lineare Abbildungen (Theorem 7.23).

Eine Alternative, den Satz zu beweisen, ist zu zeigen, dass man jede Matrix auf die im Satz gegebene Form bringen kann, wenn man elementare Zeilenumformungen und elementare

Spaltenumformungen durchführt. Zeilenumformungen entsprechen, wie wir gesehen haben, der Multiplikation mit gewissen invertierbaren Matrizen von links (Bemerkung 5.37). Analog entsprechen Spaltenumformungen der Multiplikation mit invertierbaren Matrizen von rechts (auch hier kann man die Matrizen $E_{ij}(a)$, P_{ij} und invertierbare Diagonalmatrizen der entsprechenden Größe verwenden). Diese invertierbaren Matrizen kann man wie oben besprochen als Basiswechselmatrizen interpretieren. $\square$

Eine oberflächlich ähnliche Frage, die aber wesentlich schwieriger zu beantworten ist und die uns in der Linearen Algebra 2 für längere Zeit beschäftigen wird, ist die folgende: Gegeben ein Endomorphismus $f: V \rightarrow V$, was ist eine möglichst einfache Gestalt von $M_{\mathcal{B}}^{\mathcal{B}}(f)$, die man für geeignetes $\mathcal{B}$ erreichen kann? (Hier soll beide Male *dieselbe* Basis von V verwendet werden.) Es ist für gegebenes f im allgemeinen nicht möglich, eine Basis $\mathcal{B}$ zu finden, so dass $M_{\mathcal{B}}^{\mathcal{B}}(f)$ eine Diagonalmatrix ist.

BEISPIEL 7.37. Sei $A = \frac{1}{3} \begin{pmatrix} 4 & -2 \\ -1 & 5 \end{pmatrix}$ und $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die zugehörige Abbildung. Sei $\mathcal{E} = (e_1, e_2)$ die Standardbasis und sei

$$\mathcal{B} = (b_1, b_2) = \left(\begin{pmatrix} 4 \\ 2 \end{pmatrix}, \begin{pmatrix} -2 \\ 2 \end{pmatrix} \right).$$

Dies ist ebenfalls eine Basis von $\mathbb{R}^2$. Als Basiswechselmatrizen erhalten wir

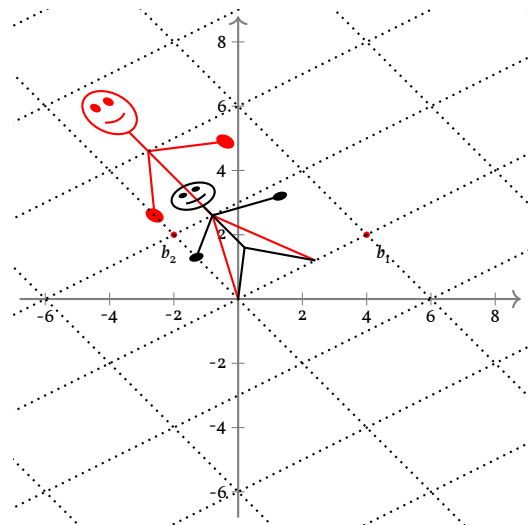
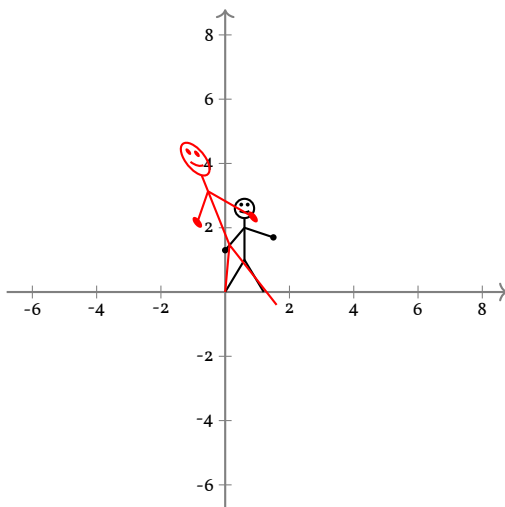
$$M_{\mathcal{E}}^{\mathcal{B}} = M_{\mathcal{E}}^{\mathcal{B}}(\text{id}) = \begin{pmatrix} 4 & -2 \\ 2 & 2 \end{pmatrix}, \quad M_{\mathcal{B}}^{\mathcal{E}} = (M_{\mathcal{E}}^{\mathcal{B}})^{-1} = \frac{1}{6} \begin{pmatrix} 1 & 1 \\ -1 & 2 \end{pmatrix}.$$

Die Matrix $M_{\mathcal{E}}^{\mathcal{B}}$ erhalten wir also, indem wir einfach die Vektoren von $\mathcal{B}$ als Spalten in eine Matrix schreiben. Die Matrix $M_{\mathcal{B}}^{\mathcal{E}}$ können wir dann als die dazu inverse Matrix auf dem üblichen Wege berechnen.

Dann gilt

$$\begin{aligned} M_{\mathcal{B}}^{\mathcal{B}}(f) &= M_{\mathcal{B}}^{\mathcal{E}} A M_{\mathcal{E}}^{\mathcal{B}} \\ &= \frac{1}{18} \begin{pmatrix} 1 & 1 \\ -1 & 2 \end{pmatrix} \begin{pmatrix} 4 & -2 \\ -1 & 5 \end{pmatrix} \begin{pmatrix} 4 & -2 \\ 2 & 2 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}. \end{aligned}$$

Bezüglich des durch $\mathcal{B}$ gegebenen Koordinatensystems ist die Abbildung f also einfach die Abbildung, die die b_1 -Achse punktweise fixiert und die die b_2 -Achse mit dem Faktor 2 streckt.



In beiden Bildern wird dieselbe Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ (wie oben definiert) gezeigt, und in beiden Bildern ist die rote Figur das Bild der schwarzen Figur. Während man links »nichts erkennen kann«, sieht man rechts, dass die gegebene Abbildung bezüglich des durch b_1, b_2 gegebenen »Koordinatensystems« (das durch die gepunkteten Linien angedeutet wird) einfach eine Streckung um den Faktor 2 in Richtung der b_2 -Achse ist.

Darüber, wie man systematisch so eine »schöne« Basis $\mathcal{B}$ zu einer gegebenen Abbildung f findet, werden wir in Kapitel 10 und in der Linearen Algebra 2 mehr lernen. $\diamond$

Wir haben gesehen, dass für jede lineare Abbildung $V \rightarrow W$ der Kern ein Untervektorraum von V ist. Andererseits gilt

LEMMA 7.38. *Sei U ein Untervektorraum des [endlich erzeugten] K -Vektorraums V . Dann gibt es eine lineare Abbildung $f: V \rightarrow W$ mit $\text{Ker}(f) = U$.*

BEWEIS. Sei $U' \subseteq V$ ein Komplement von U . Jedes $v \in V$ lässt sich dann eindeutig als $u + u'$ mit $u \in U, u' \in U'$ schreiben. Wir definieren die Abbildung $f: V \rightarrow U'$ durch $v = u + u' \mapsto u'$. (Man nennt diese Abbildung die *Projektion* auf U' . Sie hängt aber von der Zerlegung $V = U \oplus U'$, also auch von U ab.)

Für $u \in U$ ist die obige Zerlegung $u = u + 0$, also gilt $U \subseteq \text{Ker}(f)$. Ist umgekehrt $v = u + u' \in \text{Ker}(f)$, so gilt $u' = 0$, also $v \in U$. $\square$

Da wir die Lösungsmenge eines linearen Gleichungssystems als einen Kern sehen können, erhalten wir mit wenig Mehraufwand den Satz, dass jeder Untervektorraum von K^n die Lösungsmenge eines linearen Gleichungssystems ist. Siehe auch Bemerkung 7.56 und Abschnitt 7.6.4.

SATZ 7.39. *Jeder Untervektorraum $U \subseteq K^n$ ist Lösungsmenge eines homogenen linearen Gleichungssystems. Jedes lineare Gleichungssystem mit Lösungsmenge U hat $\geq n - \dim(U)$ Gleichungen, und es ist möglich, mit $n - \dim(U)$ Gleichungen auszukommen.*

BEWEIS. Sei $U \subseteq K^n$ ein Untervektorraum. Wir müssen eine Matrix A finden, so dass $\text{Ker}(A) = U$. Wegen des Zusammenhangs zwischen Matrizen und linearen Abbildungen ist es äquivalent, eine lineare Abbildung $f: K^n \rightarrow K^m$ mit $\text{Ker}(f) = U$ zu finden.

Sei $U' \subseteq K^n$ ein Komplement von U und sei $g: K^n \rightarrow U'$ die Projektion. Sei $h: U' \rightarrow K^m$ ein Isomorphismus (mit $m := \dim U' = n - \dim U$). Dann ist $h \circ g$ ein Vektorraum-Homomorphismus $K^n \rightarrow K^m$ mit Kern U .

Die Anzahl der Gleichungen ist gleich $m = n - \dim(U)$. Aus der Dimensionsformel für lineare Abbildungen, Theorem 7.23, folgt, dass es nicht möglich ist, mit weniger Gleichungen auszukommen. $\square$

7.4. Der Rang einer Matrix

DEFINITION 7.40. Sei $A \in M_{m \times n}(K)$.

- (1) Der *Rang der Matrix* A , geschrieben $\text{rg } A$, ist der Rang der linearen Abbildung f_A , also die Dimension von $\text{Im}(f_A)$. Mit anderen Worten: $\text{rg } A$ ist die Dimension des Untervektorraums von K^m , der von den Spalten von A erzeugt wird. Man spricht daher auch vom *Spaltenrang* von A .
- (2) Der *Zeilenrang* von A ist die Dimension des Untervektorraums von K^n , der von den Zeilen von A erzeugt wird.

Wir wie gleich sehen werden, stimmen für jede Matrix A ihr Spaltenrang und ihr Zeilenrang überein. Diese – nicht offensichtliche – Tatsache erlaubt es uns, für den Zeilenrang kein eigenes Symbol einzuführen.

THEOREM 7.41. *Sei $A \in M_{m \times n}(K)$ eine Matrix. Dann stimmen der Spaltenrang und der Zeilenrang von A überein.*

BEWEIS. Sei $A \in M_{m \times n}(K)$ gegeben, und sei $f = f_A$ die zugehörige lineare Abbildung $K^n \rightarrow K^m$. Der Spaltenrang von A ist dann gerade $\dim \operatorname{Im}(A)$, also der Rang von f . Nach der Dimensionsformel ist diese Zahl gleich $n - \dim \operatorname{Ker}(A)$. Es genügt also zu zeigen, dass der Zeilenrang von A mit $n - \dim \operatorname{Ker}(A)$ übereinstimmt.

Nun verändern sich weder der Zeilenrang noch die Zahl $n - \dim \operatorname{Ker}(A)$, wenn wir A ersetzen durch eine Matrix, die aus A durch eine elementare Zeilenumformung hervorgeht. (Was den Zeilenrang angeht, bleibt sogar der von den Zeilen erzeugte Untervektorraum gleich. Außerdem ändert sich der Kern von A nicht, siehe Lemma 5.12). Durch solche Umformungen können wir A letztlich auf Zeilenstufenform bringen. Für eine Matrix in Zeilenstufenform ist die Aussage aber klar: Denn der Zeilenrang ist dann gerade die Anzahl der nicht-verschwindenden Zeilen, also die Anzahl der führenden Einsen, und die Dimension des Kerns ist die Anzahl der Spalten ohne führende Eins (eben die Anzahl der »frei wählbaren Variablen«). Die Differenz $n - \dim(\operatorname{Ker}(A))$ ist also ebenfalls die Anzahl der führenden Einsen. $\square$

Wohlgemerkt muss die Matrix A hier nicht quadratisch sein! Die von den Zeilen bzw. Spalten erzeugten Untervektorräume befinden sich im nicht-quadratischen Fall in *unterschiedlichen* umgebenden K -Vektorräumen. Das Theorem sagt, dass sie aber immer dieselbe Dimension haben (und insbesondere ist diese $\leq \min(m, n)$).

Wir geben noch einen alternativen Beweis (der insofern einfacher ist, als dass die Dimensionsformel für lineare Abbildungen nicht benutzt werden muss, sondern nur die Tatsache, dass isomorphe K -Vektorräume dieselbe Dimension haben (Lemma 7.11)). Siehe auch Bemerkung 7.55 für noch einen anderen Beweis.

Sei $A \in M_{m \times n}(K)$. Wir können aus Satz 7.36 folgern, dass invertierbare Matrizen $S \in M_m(K)$, $T \in M_n(K)$ existieren, so dass

$$SAT = \begin{pmatrix} E_r & 0 \\ 0 & 0 \end{pmatrix}$$

(als Blockmatrix verstanden), wobei $r = \operatorname{rg}(A)$. In der Tat können wir den Satz anwenden auf die Abbildung f_A und finden damit Basen $\mathcal{B}, \mathcal{C}$ von K^n und K^m , so dass $M_{\mathcal{C}}^{\mathcal{B}}(f_A)$ die obige Form hat. Wir definieren dann $S = M_{\mathcal{C}}^{\mathcal{E}_m}$, $T = M_{\mathcal{E}_n}^{\mathcal{B}}$, wobei $\mathcal{E}_n$ und $\mathcal{E}_m$ die Standardbasen von K^n und K^m bezeichnen.

Für die Matrix $\begin{pmatrix} E_r & 0 \\ 0 & 0 \end{pmatrix}$ ist klar, dass Zeilenrang und Spaltenrang übereinstimmen (nämlich gleich r sind). Es genügt daher zu zeigen, dass Zeilenrang und Spaltenrang sich nicht verändern, wenn wir eine Matrix von links und rechts mit invertierbaren Matrizen (der richtigen Größen) multiplizieren. Für den Spaltenrang folgt das aus Teil (3) des folgenden Lemmas. Für den Zeilenrang erhalten wir das Ergebnis, wenn wir das Lemma auf die transponierte Matrix anwenden (denn der Zeilenrang einer Matrix A ist gleich dem Spaltenrang von A^t , und mit S und T sind auch S^t und T^t invertierbar, und $(SAT)^t = T^t A^t S^t$).

LEMMA 7.42. *Sei K ein Körper.*

- (I) *Sei $f: V \rightarrow W$ eine lineare Abbildung zwischen endlich-dimensionalen K -Vektorräumen und seien $c: W \rightarrow W'$ und $c': V' \rightarrow V$ Isomorphismen. Dann gilt $\operatorname{rg}(c \circ f \circ c') = \operatorname{rg}(f)$.*

- (2) Sei $f: V \rightarrow W$ eine lineare Abbildung zwischen endlich-dimensionalen K -Vektorräumen. Seien $\mathcal{B}$ eine Basis von V und $\mathcal{C}$ eine Basis von W . Dann gilt $\operatorname{rg}(M_{\mathcal{C}}^{\mathcal{B}}(f)) = \operatorname{rg}(f)$.
- (3) Sei $A \in M_{m \times n}(K)$. Seien $S \in M_m(K)$, $T \in M_n(K)$ invertierbare Matrizen. Dann gilt $\operatorname{rg}(SAT) = \operatorname{rg}(A)$.

BEWEIS. zu (1). Diese Aussage belegt erneut die Philosophie, dass Vektorraumstrukturen durch Isomorphismen erhalten bleiben. Es ist $\operatorname{Im}(c') = V$, weil c' ein Isomorphismus, also insbesondere surjektiv ist; also ist $\operatorname{Im}(c \circ f \circ c') = \operatorname{Im}(c \circ f)$. Weil c ein Isomorphismus ist, ist auch die Einschränkung $c|_{\operatorname{Im}(f)}: \operatorname{Im}(f) \rightarrow W'$ injektiv, induziert somit einen Isomorphismus $\operatorname{Im}(f) \cong \operatorname{Im}(c \circ f)$. Es folgt $\operatorname{rg}(f) = \dim \operatorname{Im}(c \circ f) = \dim \operatorname{Im}(c \circ f \circ c') = \operatorname{rg}(c \circ f \circ c')$.

zu (2). Nach Definition ist

$$M_{\mathcal{C}}^{\mathcal{B}}(f) = M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1}),$$

und die Aussage folgt aus Teil (1), weil wir bereits wissen, dass $\operatorname{rg}(M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1})) = \operatorname{rg}(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1})$ ist.

zu (3). Man kann diese Aussage auch durch Argumente »mit Matrizen« beweisen, aber wir wollen einen kurzen und sehr transparenten Beweis geben, der Teil (1) und die Entsprechung von Matrizen und linearen Abbildungen ausnutzt:

Es gilt

$$\operatorname{rg}(SAT) = \operatorname{rg}(f_S \circ f_A \circ f_T) = \operatorname{rg}(A),$$

wobei wir bei der zweiten Gleichheit Teil (1) nutzen, und in beiden Schritten die Gleichheit $\operatorname{rg}(f_M) = \operatorname{rg}(M)$, die für beliebige Matrizen M richtig ist. $\square$

Man sagt manchmal, zwei Matrizen $A, B \in M_{m \times n}(K)$ seien *äquivalent*, wenn invertierbare Matrizen $S \in M_m(K)$ und $T \in M_n(K)$ existieren, so dass $B = SAT$. Aus dem Lemma sehen wir, dass äquivalente Matrizen denselben Rang haben. Aus der oben angesprochenen »Matrizenversion« der Smith-schen Normalform folgt, dass Matrizen mit demselben Rang äquivalent sind.

7.4.1. Zusammenfassung. Wir wollen noch einmal verschiedene Ergebnisse, die wir inzwischen bewiesen haben, zusammentragen.

Seien V, W endlich erzeugte Vektorräume der Dimensionen $n = \dim V$, $m = \dim W$, und seien $\mathcal{B} = (v_1, \dots, v_n)$, $\mathcal{C} = (w_1, \dots, w_m)$ Basen von V bzw. W . Sei $f: V \rightarrow W$ eine lineare Abbildung und $A = M_{\mathcal{C}}^{\mathcal{B}}(f) \in M_{m \times n}(K)$ die zugehörige Matrix.

Es ist dann $\operatorname{Ker}(A) = \operatorname{Ker}(f_A)$ die Lösungsmenge des durch A gegebenen homogenen linearen Gleichungssystems und der Kern der Abbildung $x \mapsto Ax$. Die Koordinatenabbildung $c_{\mathcal{B}}$ induziert einen Isomorphismus $\operatorname{Ker}(f) \cong \operatorname{Ker}(A)$.

Andererseits ist $\operatorname{Im}(A) = \operatorname{Im}(f_A)$ die Menge aller $b \in K^m$, für die das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$ lösbar ist. Die Koordinatenabbildung $c_{\mathcal{C}}$ induziert einen Isomorphismus $\operatorname{Im}(f) \cong \operatorname{Im}(A)$. Insbesondere gilt $\operatorname{rg}(f) = \operatorname{rg}(A)$.

(1) Es sind äquivalent:

(i) Die Abbildung f ist injektiv.

(ii) $\operatorname{Ker}(A) = 0$.

(iii) $\operatorname{rg}(A) = n$.

(iv) Die Spalten von A sind linear unabhängig.

(v) Das durch A gegebene homogene lineare Gleichungssystem ist eindeutig lösbar (d.h. es existiert nur die triviale Lösung).

(2) Es sind äquivalent:

- (i) Die Abbildung f ist surjektiv.
 - (ii) $\operatorname{rg} A = m$.
 - (iii) Die Zeilen von A sind linear unabhängig.
 - (iv) Das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$ ist für alle $b \in K^m$ lösbar.
- (3) Ist $m = n$, so sind alle Punkte in (1) und (2) äquivalent, und äquivalent zu
- (i) Die Abbildung f ist ein Isomorphismus.
 - (ii) Die Matrix A ist invertierbar.
 - (iii) Die Spalten von A bilden eine Basis von K^m .
 - (iv) Das lineare Gleichungssystem mit erweiterter Koeffizientenmatrix $(A \mid b)$ ist für alle $b \in K^m$ eindeutig lösbar.

7.5. Der Dualraum eines Vektorraums

Sei K ein Körper.

DEFINITION 7.43. Sei V ein K -Vektorraum. Dann heißt der K -Vektorraum $V^\vee := \operatorname{Hom}_K(V, K)$ der *Dualraum* von V (oder der *duale Vektorraum* zu V). $\dashv$

Die Vektorraumstruktur auf $\operatorname{Hom}_K(V, K)$ ist dabei wie folgt gegeben (siehe Definition 7.6 und die der Definition folgende Erklärung):

$$\begin{aligned}(\lambda + \mu): V &\rightarrow K, & v &\mapsto \lambda(v) + \mu(v), \\(a\lambda): V &\rightarrow K, & v &\mapsto a\lambda(v).\end{aligned}$$

Oft wird der Dualraum auch mit V^* bezeichnet. Seine Elemente, die linearen Abbildungen $V \rightarrow K$, nennt man auch *Linearformen*. Sie werden oft mit den griechischen Kleinbuchstaben $\lambda, \mu, \dots$ bezeichnet.

Diese (eigentlich sehr einfache) Konstruktion ist abstrakter als etwa ein Teilraum von K^n , der durch explizite Gleichungen gegeben ist. Gleichzeitig liefert sie Beispiele für K -Vektorräume, die nicht ein Untervektorraum eines Standard-Vektorraums K^n sind.

Die Theorie des Dualraums lässt sich direkt mit linearen Gleichungssystemen verbinden, denn die Lösungsmenge in K^n einer einzigen linearen Gleichung in n Unbestimmten ist der Kern einer geeigneten Linearform $K^n \rightarrow K$. So liefern die Ergebnisse dieses Abschnitts auch einen weiteren Beweis der Tatsache, dass jeder Untervektorraum von K^n Lösungsmenge eines homogenen linearen Gleichungssystems ist (Bemerkung 7.56).

Interessant wird der Dualraum (unter anderem), weil man zu einem Vektorraum-Homomorphismus $f: V \rightarrow W$ eine Abbildung zwischen den Dualräumen definieren kann, und zwar die sogenannte *duale Abbildung*

$$f^\vee: W^\vee \rightarrow V^\vee, \quad \lambda \mapsto \lambda \circ f,$$

wir bilden also ein Element $\lambda \in W^\vee$, d.h. eine lineare Abbildung $\lambda: W \rightarrow K$, ab auf die Verkettung $\lambda \circ f$, die eine lineare Abbildung $V \rightarrow K$ ist, also ein Element von $V^\vee$. Man beachte, dass sich sozusagen die Richtung des Abbildungspfeils umkehrt: Beginnen wir mit einer Abbildung $V \rightarrow W$, so ist die duale Abbildung eine Abbildung von $W^\vee$ nach $V^\vee$.

Die Abbildung $f^\vee$ ist eine lineare Abbildung: Für $a, b \in K, \lambda, \mu \in W^\vee$ haben wir

$$f^\vee(a\lambda + b\mu) = (a\lambda + b\mu) \circ f,$$

und diese Abbildung hat die Abbildungsvorschrift

$$v \mapsto (a\lambda + b\mu)(f(v)) = a\lambda(f(v)) + b\mu(f(v)),$$

genauso wie die Abbildung

$$a(\lambda \circ f) + b(\mu \circ f) = af^\vee(\lambda) + bf^\vee(\mu).$$

Die Konstruktion der dualen Abbildung ist kompatibel mit der Verkettung von Abbildungen: Sind $f: U \rightarrow V$ und $g: V \rightarrow W$ lineare Abbildungen, und sind $f^\vee: V^\vee \rightarrow U^\vee$ und $g^\vee: W^\vee \rightarrow V^\vee$ die dazu dualen Abbildungen, so gilt

$$(g \circ f)^\vee = f^\vee \circ g^\vee.$$

In der Tat ist $(g \circ f)^\vee(\lambda) = \lambda \circ (g \circ f) = (\lambda \circ g) \circ f = g^\vee(\lambda) \circ f = f^\vee(g^\vee(\lambda))$. Außerdem gilt natürlich $\text{id}_V^\vee = \text{id}_{V^\vee}$.

Diese abstrakten Eigenschaften haben die folgende nützliche Konsequenz: Ist $f: V \rightarrow W$ ein Vektorraum-Isomorphismus, dann ist auch $f^\vee: W^\vee \rightarrow V^\vee$ ein Isomorphismus. Denn ist g der Umkehrhomomorphismus zu f , so gilt $f \circ g = \text{id}_W$, $g \circ f = \text{id}_V$, also mit dem eben gesagten: $g^\vee \circ f^\vee = \text{id}_{W^\vee}$, $f^\vee \circ g^\vee = \text{id}_{V^\vee}$, d.h. die Abbildung $g^\vee$ ist die Umkehrabbildung von $f^\vee$.

Es gilt auch der folgende Satz, der Injektivität und Surjektivität von f und seiner dualen Abbildung miteinander in Verbindung bringt.

SATZ 7.44. *Seien V und W [endlich erzeugte] Vektorräume über K . Sei $f: V \rightarrow W$ ein Vektorraum-Homomorphismus, und $f^\vee$ die dazu duale Abbildung.*

- (1) *Die Abbildung f ist genau dann injektiv, wenn $f^\vee$ surjektiv ist.*
- (2) *Die Abbildung f ist genau dann surjektiv, wenn $f^\vee$ injektiv ist.*

BEWEIS. Die Einschränkung in eckigen Klammern erklärt sich wie üblich dadurch, dass wir im Beweis vom Basisergänzungssatz Gebrauch machen.

zu (1). Sei zunächst f injektiv und sei $U = \text{Im}(f)$ das Bild der Abbildung f . Da f injektiv ist, ist die Abbildung $f: V \rightarrow U$ injektiv und surjektiv, also ein Isomorphismus. Wir bezeichnen mit $g: U \rightarrow V$ die Umkehrabbildung. Sei nun U' ein Komplement von U in W .

Wir wollen zeigen, dass $f^\vee: W^\vee \rightarrow V^\vee$ surjektiv ist. Sei also $\lambda \in V^\vee$. Wir definieren eine lineare Abbildung $\mu: W \rightarrow K$ (also ein Element von $W^\vee$) wie folgt:

$$\mu: W = U \oplus U' \rightarrow K, \quad (u, u') \mapsto \lambda(g(u)).$$

Dann gilt $(\mu \circ f)(v) = \mu(f(v)) = \lambda(v)$, also $\mu \circ f = \lambda$. Das bedeutet gerade $f^\vee(\mu) = \lambda$.

Sei nun $f^\vee$ surjektiv. Wir wollen zeigen, dass f injektiv ist. Angenommen, das wäre nicht der Fall, so dass der Untervektorraum $\text{Ker}(f)$ von V nicht der Nullvektorraum ist. Dann gibt es eine lineare Abbildung $\lambda: V \rightarrow K$, deren Einschränkung auf $\text{Ker}(f)$ nicht die Nullabbildung ist. (Denn wir können eine Basis von $\text{Ker}(f)$ zu einer Basis von V ergänzen und dann die Abbildung λ definieren, indem wir auf den Basisvektoren beliebige Werte vorgeben.) Dann kann aber λ nicht die Form $\mu \circ f$ für irgendein μ haben, denn jede Abbildung dieser Form ist auf $\text{Ker}(f)$ die Nullabbildung.

zu (2). Sei zunächst f surjektiv. Sind $\mu, \mu' \in W^\vee$ mit $f^\vee(\mu) = f^\vee(\mu')$, also mit $\mu \circ f = \mu' \circ f$, so folgt $\mu(w) = \mu'(w)$ für alle $w \in W$, denn jedes w lässt sich als $f(v)$ schreiben. Somit gilt $\mu = \mu'$.

Es bleibt noch der Fall zu betrachten, dass $f^\vee$ injektiv ist. Um zu zeigen, dass dann f surjektiv ist, nehmen wir an, das wäre nicht der Fall und wählen einen Komplementärraum U von $\text{Im}(f)$ in W . Unsere Annahme bedeutet, dass $\text{Im}(f) \neq W$ ist, also dass $U \neq 0$. Mit einem ähnlichen Argument wie im ersten Teil des Beweises sehen wir, dass eine lineare Abbildung

$\mu: W \rightarrow K$ existiert mit $\mu_{\text{Im}(f)} = 0$, aber $\mu|_U \neq 0$. Dann gilt aber $f^\vee(\mu) = \mu \circ f = 0 = f^\vee(0)$ (die letzte 0 bezeichnet die Nullabbildung $W \rightarrow K$). Das steht im Widerspruch zur Injektivität von $f^\vee$. $\square$

ERGÄNZUNG 7.45. In dieser Ergänzung soll noch ein etwas allgemeinerer Zugang zu dem vorherigen Ergebnis skizziert werden. Wir beginnen mit dem folgenden

LEMMA 7.46. Seien $f: U \rightarrow V, g: V \rightarrow W$ lineare Abbildungen zwischen [endlich erzeugten] K -Vektorräumen, und seien $g^\vee: W^\vee \rightarrow V^\vee, f^\vee: V^\vee \rightarrow U^\vee$ die dualen Abbildungen von g und f . Dann gilt

$$\text{Im}(f) = \text{Ker}(g) \iff \text{Im}(g^\vee) = \text{Ker}(f^\vee).$$

BEWEIS. Es gelte zunächst $\text{Im}(f) = \text{Ker}(g)$. Wir zeigen $\text{Im}(g^\vee) \subseteq \text{Ker}(f^\vee)$ und $\text{Im}(g^\vee) \supseteq \text{Ker}(f^\vee)$.

Die Inklusion $\text{Im}(g^\vee) \subseteq \text{Ker}(f^\vee)$ ist gleichbedeutend damit, dass $f^\vee \circ g^\vee = 0$ ist. Weil $g \circ f = 0$ (denn $\text{Im}(f) \subseteq \text{Ker}(g)$), weil der Übergang zur dualen Abbildung kompatibel mit der Verkettung von Abbildungen ist und weil die zur Nullabbildung duale Abbildung die Nullabbildung ist, gilt tatsächlich $f^\vee \circ g^\vee = 0$.

Die umgekehrte Inklusion zeigt man mit ähnlichen Argumenten, wie wir sie im Beweis des Satzes verwendet haben. Sei $V' \subseteq V$ ein Komplementärraum von $f(U) = \text{Ker}(g)$. Die Einschränkung von g auf V' ist dann injektiv, induziert also einen Isomorphismus $h: V' \rightarrow g(V')$. Sei $W' \subseteq W$ ein Komplement von $g(V')$. Ist $\lambda \in \text{Ker}(f^\vee)$, das heißt $\lambda \circ f = 0$, so definieren wir $\mu: W \rightarrow K$ wie folgt. Für $w \in g(V')$ sei $\mu(w) = \lambda(h^{-1}(w))$. Für $w \in W'$ sei $\mu(w) = 0$. Weil $W = W' \oplus g(V')$ gilt, definiert dies tatsächlich eine lineare Abbildung, also ein Element von $W^\vee$. Es ist dann leicht zu sehen, dass $\lambda = \mu \circ g$ gilt (es genügt, das für Elemente in $f(U)$ und Elemente in V' zu überprüfen; erstere werden unter beiden Abbildungen auf 0 abgebildet, und für $v \in V'$ gilt $\mu(g(v)) = \lambda(v)$ nach Definition von μ). Es folgt $\lambda = g^\vee(\mu) \in \text{Im}(g^\vee)$.

Nun setzen wir voraus, dass $\text{Im}(g^\vee) = \text{Ker}(f^\vee)$ gilt. Um die Inklusion $\text{Im}(f) \subseteq \text{Ker}(g)$ zu zeigen, müssen wir zeigen, dass $g \circ f = 0$. Angenommen, es gibt ein Element $u \in U$, so dass $w := g(f(u)) \neq 0$. Dann finden wir eine Linearform $\lambda \in W^\vee$ mit $\lambda(w) \neq 0$. Da $\text{Im}(g^\vee) \subseteq \text{Ker}(f^\vee)$, gilt aber $\lambda \circ g \circ f = 0$, ein Widerspruch.

Es bleibt nun noch, die Inklusion $\text{Ker}(g) \subseteq \text{Im}(f)$ zu zeigen. Sei dazu $v \in \text{Ker}(g)$. Wäre $v \notin \text{Im}(f)$, so finden wir mit ähnlichen Argumenten wie oben eine Linearform $\lambda: V \rightarrow K$ mit $\lambda(\text{Im}(f)) = 0$ und $\lambda(v) \neq 0$. Dass λ auf $\text{Im}(f)$ verschwindet, bedeutet genau, dass $\lambda \in \text{Ker}(f^\vee) = \text{Im}(g^\vee)$, etwa $\lambda = \mu \circ g$. Dann gilt aber $\lambda(v) = \mu(g(v)) = 0$, weil wir $v \in \text{Ker}(g)$ gewählt hatten – ein Widerspruch. Es muss also doch $v \in \text{Im}(f)$ gelten, und wir sind fertig. $\square$

Zusammen mit dem folgenden Lemma (und der offensichtlichen Bemerkung, dass der Dualraum des Nullvektorraums der Nullvektorraum ist) erhalten wir einen neuen Beweis von Satz 7.44.

LEMMA 7.47. (1) Seien $f: 0 \rightarrow V$ die Nullabbildung und $g: V \rightarrow W$ eine lineare Abbildung. Es gilt genau dann $\text{Im}(f) = \text{Ker}(g)$, wenn g injektiv ist.

(2) Seien $f: U \rightarrow V$ eine lineare Abbildung und $g: V \rightarrow 0$ die Nullabbildung. Es gilt genau dann $\text{Im}(f) = \text{Ker}(g)$, wenn f surjektiv ist.

Der Beweis ist einfach.

$\square$ Ergänzung 7.45

BEISPIEL 7.48 (Dualraum von K^n). Ein Element $\lambda \in \text{Hom}_K(K^n, K)$, also eine lineare Abbildung $\lambda: K^n \rightarrow K$, ist bestimmt durch die Werte $\lambda(e_i)$ auf den Standardbasisvektoren, also durch den Vektor $(\lambda(e_i))_i^t \in K^n$. Wir erhalten einen Isomorphismus

$$(K^n)^\vee = \text{Hom}_K(K^n, K) \rightarrow K^n, \quad \lambda \mapsto (\lambda(e_i))_i^t,$$

von K -Vektorräumen. ◇

Dieses Beispiel können wir auch als Spezialfall des Isomorphismus $\text{Hom}_K(K^n, K^m) \cong M_{m \times n}(K)$ aus Satz 7.26 für $m = 1$ betrachten.

SATZ 7.49. Sei V ein K -Vektorraum mit Basis $\mathcal{B} = (b_1, \dots, b_n)$. Sei $\lambda_i \in V^\vee$ das eindeutig bestimmte Element mit

$$\lambda_i(b_j) = \begin{cases} 1 & i = j, \\ 0 & i \neq j. \end{cases}$$

Dann ist $\lambda_1, \dots, \lambda_n$ eine Basis von $V^\vee$. Wir bezeichnen diese Basis als die zu $\mathcal{B}$ duale Basis.

BEWEIS. Eine andere Beschreibung der im Satz definierten Abbildungen λ_i ist, dass $\lambda_i(v) = c_{\mathcal{B}}(v)_i$, der i -te Eintrag des Koordinatenvektors von v bezüglich der Basis $\mathcal{B}$, ist. Weil sowohl λ_i als auch $v \mapsto c_{\mathcal{B}}(v)_i$ lineare Abbildungen sind, genügt es, das für $v = b_1, \dots, b_n$ nachzuprüfen, und in diesen Fällen ist die Aussage klar.

Es ist nicht schwierig, direkt zu überprüfen, dass sich jedes Element von $V^\vee$ in eindeutiger Weise als Linearkombination von $\lambda_1, \dots, \lambda_n$ darstellen lässt. In der Tat, ist $\mu: V \rightarrow K$ eine lineare Abbildung, so gilt

$$\mu(v) = \mu\left(\sum_{i=1}^n c_{\mathcal{B}}(v)_i b_i\right) = \sum_{i=1}^n c_{\mathcal{B}}(v)_i \mu(b_i) = \sum_{i=1}^n \mu(b_i) \lambda_i(v),$$

also $\mu = \sum_{i=1}^n \mu(b_i) \lambda_i$, und die $\mu(b_i)$ sind die einzig möglichen Koeffizienten, wie man sieht, wenn man die Abbildung bei $b_1, \dots, b_n$ auswertet.

Alternativ kann man wie folgt argumentieren. Zur Basis $b_1, \dots, b_n$ korrespondiert der Koordinatenisomorphismus $c_{\mathcal{B}}: V \rightarrow K^n$. Durch Übergang zur dualen Abbildung erhalten wir einen Isomorphismus

$$K^n = (K^n)^\vee \rightarrow V^\vee,$$

wobei wir für die erste Gleichheit die Identifikation aus Beispiel 7.48 benutzen. Das Bild von $e_i \in K^n$ in $V^\vee$ unter dieser Abbildung ist λ_i , wie man leicht anhand der Definitionen nachrechnet. Daraus folgt der Satz. □

Als Folgerungen erhalten wir:

KOROLLAR 7.50. Sei V ein endlich-dimensionaler K -Vektorraum. Dann gilt $\dim V = \dim V^\vee$.

KOROLLAR 7.51. Sei $f: V \rightarrow W$ eine lineare Abbildung zwischen endlich-dimensionalen Vektorräumen. Dann gilt $\text{rg}(f) = \text{rg}(f^\vee)$.

BEWEIS. Wir schreiben f als die Verkettung $V \rightarrow \text{Im}(f) \rightarrow W$. Durch Übergang zur dualen Abbildung sehen wir, dass $f^\vee$ die Verkettung $W^\vee \rightarrow \text{Im}(f)^\vee \rightarrow V^\vee$ ist. Hier ist nun nach Satz 7.44 die Abbildung $W^\vee \rightarrow \text{Im}(f)^\vee$ surjektiv und die Abbildung $\text{Im}(f)^\vee \rightarrow V^\vee$ injektiv. Wenn wir diese injektive Abbildung nutzen, um den Vektorraum $\text{Im}(f)^\vee$ mit seinem Bild in $V^\vee$ identifizieren, können wir das formulieren als

$$\text{Im}(f)^\vee = \text{Im}(f^\vee).$$

Insbesondere erhalten wir $\text{rg}(f) = \dim \text{Im}(f) = \dim \text{Im}(f)^\vee = \dim \text{Im}(f^\vee) = \text{rg}(f^\vee)$, wie behauptet. □

Wir können auch den Dualraum des Dualraums bilden:

$$V^{\vee\vee} := (V^\vee)^\vee = \text{Hom}_K(V^\vee, K).$$

Man hat dann die folgende Abbildung

$$V \rightarrow V^{\vee\vee}, \quad v \mapsto (V^\vee \rightarrow K, \lambda \mapsto \lambda(v)).$$

Man nennt diese Abbildung die *kanonische Abbildung von V in seinen Doppeldualraum*.

KOROLLAR 7.52. *Sei V ein [endlich erzeugter] K -Vektorraum.*

- (1) *Die kanonische Abbildung $V \rightarrow V^{\vee\vee}$ von V in seinen Doppeldualraum ist injektiv.*
- (2) *Ist V endlich-dimensional, dann ist die kanonische Abbildung $V \rightarrow V^{\vee\vee}$ von V in seinen Doppeldualraum ein Isomorphismus.*

BEWEIS. zu (1). Sei $v \in V$, so dass die Abbildung $V^\vee \rightarrow K, \lambda \mapsto \lambda(v)$, die Nullabbildung ist. Wir müssen zeigen, dass $v = 0$ ist. Andernfalls ist die Menge $\{v\}$ linear unabhängig und wir können sie zu einer Basis ergänzen, in der v vorkommt. Das erlaubt uns (Satz 7.14), eine lineare Abbildung $\lambda: V \rightarrow K$ zu definieren mit $\lambda(v) \neq 0$ (zum Beispiel können wir $\lambda(v) = 1$ und $\lambda(w) = 0$ für alle die anderen Basisvektoren setzen). Das ist ein Widerspruch zur Voraussetzung, es gilt also $v = 0$.

zu (2). Dies folgt direkt aus Teil (1), weil wir bereits gezeigt haben, dass V und $V^\vee$ dieselbe Dimension haben. $\square$

Teil (2) des Korollars (mit der kleinen Präzisierung, dass auch für jede Abbildung f zwischen endlich-dimensionalen Vektorräumen $f^{\vee\vee}$ unter den Identifizierungen der Doppeldualräume mit den ursprünglichen Vektorräumen mit f identifiziert wird) würde es erlauben, im Beweis von Satz 7.44 den zweiten Teil aus dem ersten durch Übergang zum Dual zu folgern (oder umgekehrt); man könnte so den Beweis auf die Hälfte kürzen. Eine entsprechende Bemerkung gilt auch für Lemma 7.46.

ERGÄNZUNG 7.53 (Der Dualraum eines unendlich-dimensionalen Vektorraums). Ist V unendlich-dimensional, so ist $V^\vee$ ebenfalls unendlich-dimensional, hat aber niemals dieselbe Dimension wie V , wenn man die Dimension als Kardinalzahl interpretiert; das bedeutet: Sind $\mathcal{B}$ eine Basis von V und $\mathcal{C}$ eine Basis von $V^\vee$, so existiert keine Bijektion zwischen $\mathcal{B}$ und $\mathcal{C}$.

Insbesondere ist Teil (2) des Korollars im unendlich-dimensionalen Fall falsch (das *endlich-dimensional* dort steht also bewusst nicht in eckigen Klammern). $\square$ Ergänzung 7.53

Den folgenden Satz können wir einerseits als eine Methode sehen, die darstellende Matrix der dualen Abbildung zu einer Abbildung f auszurechnen. Wir können ihn aber auch »andersherum« lesen: Wir erhalten eine konzeptionelle Interpretation des Übergangs von einer Matrix A zu ihrer Transponierten.

SATZ 7.54. *Seien V und W endlich-dimensionale K -Vektorräume mit Basen $\mathcal{B}$ und $\mathcal{C}$. Sei $\mathcal{B}^\vee$ die duale Basis von $\mathcal{B}$ und $\mathcal{C}^\vee$ die duale Basis von $\mathcal{C}$. Ist $f: V \rightarrow W$ ein Homomorphismus, so gilt*

$$M_{\mathcal{B}^\vee}^{\mathcal{C}^\vee}(f^\vee) = M_{\mathcal{C}}^{\mathcal{B}}(f)^t.$$

BEWEIS. Sei $n = \dim V$, $m = \dim W$, $\mathcal{B} = (b_1, \dots, b_n)$, $\mathcal{B}^\vee = (\lambda_1, \dots, \lambda_n)$, $\mathcal{C}^\vee = (\mu_1, \dots, \mu_m)$. Wir schreiben $M_{\mathcal{C}}^{\mathcal{B}}(f) = (a_{ij})_{i,j} \in M_{m \times n}(K)$ und $M_{\mathcal{B}^\vee}^{\mathcal{C}^\vee}(f^\vee) = (b_{ij})_{i,j} \in M_{n \times m}(K)$. Wir wollen zeigen, dass $a_{ij} = b_{ji}$ für alle $i = 1, \dots, m, j = 1, \dots, n$.

Es ist

$$(b_{1i}, \dots, b_{ni})^t = M_{\mathcal{B}^\vee}^{\mathcal{C}^\vee}(f^\vee)e_i = c_{\mathcal{B}^\vee}(f^\vee(\mu_i)) = c_{\mathcal{B}^\vee}(\mu_i \circ f).$$

Daraus folgern wir

$$\mu_i \circ f = \sum_{k=1}^m b_{ki} \lambda_k.$$

Setzen wir auf beiden Seiten b_j ein, so erhalten wir

$$\mu_i(f(b_j)) = b_{ji}.$$

Weil $\mu_i(f(b_j)) = c_{\mathcal{C}}(f(b_j))_i = a_{ji}$ gilt (vergleiche den Beginn des Beweises von Satz 7.49), erhalten wir die Behauptung. $\square$

BEMERKUNG 7.55. Wir haben mit Korollar 7.51 und Satz 7.54 einen neuen Beweis erhalten, dass für jede Matrix A der Zeilenrang und der Spaltenrang übereinstimmen. Denn ist $f = f_A$ die durch A gegebene Abbildung, so ist der Spaltenrang von A der Rang von f . Der Zeilenrang von A ist der Spaltenrang der transponierten Matrix A^t , also nach Satz 7.54 der Rang von $f^\vee$.

Andersherum könnte man auch Korollar 7.51 aus Satz 7.54 und der Gleichheit von Zeilen- und Spaltenrang folgern. Aber hier ging es ja gerade darum, einen neuen Zugang zu zeigen. $\diamond$

BEMERKUNG 7.56. Kommen wir noch einmal auf Satz 7.39 zurück. Mit der Theorie des Dualraums können wir dieses Ergebnis aus einem anderen Blickwinkel betrachten. Allerdings ist diese Bemerkung keine »leichte Kost«. Wir benutzen praktisch alles, was wir bisher über den Dualraum und die duale Abbildung einer linearen Abbildung gelernt haben.

Wir beginnen mit einem Untervektorraum $U \subseteq K^n$. Sei $r = \dim U$. Wählen wir eine Basis von U , so können wir U als das Bild einer linearen Abbildung $f: K^r \rightarrow K^n$ betrachten (die $e_i \in K^r$ abbildet auf den i -ten Basisvektor der gewählten Basis von U). Die Spalten der Matrix $M(f)$ sind die gewählten Basisvektoren von U .

Wir wollen U als Lösungsmenge eines linearen Gleichungssystems darstellen, wir suchen also eine Abbildung $g: K^n \rightarrow K^m$ mit $\text{Ker}(g) = U = \text{Im}(f)$. Dann wird m die Anzahl der Gleichungen sein. Wir wollen versuchen, mit möglichst wenig Gleichungen auszukommen, streben also $m = n - r$ an.

Wir betrachten nun die zu f duale Abbildung $f^\vee: (K^n)^\vee \rightarrow (K^r)^\vee$. Sei $\tilde{W} = \text{Ker}(f^\vee)$. Wir haben die Inklusionsabbildung $\tilde{W} \rightarrow (K^n)^\vee$. Weil $f^\vee$ nach Satz 7.44 surjektiv ist, folgt $\dim \tilde{W} = n - r$.

Von der Verkettung

$$\tilde{W} \rightarrow (K^n)^\vee \rightarrow (K^r)^\vee$$

gehen wir nun *wieder* zu den dualen Vektorräumen und dualen Abbildungen über. Wir bekommen

$$(K^r)^{\vee\vee} \rightarrow (K^n)^{\vee\vee} \rightarrow \tilde{W}^\vee,$$

und wenn wir entsprechend Korollar 7.52 $(K^r)^{\vee\vee} = K^r$ (und entsprechend für K^n) identifizieren,

$$K^r \rightarrow K^n \rightarrow \tilde{W}^\vee,$$

wobei die Abbildung $K^r \rightarrow K^n$ wieder die Abbildung f ist, und die Abbildung $K^n \rightarrow \tilde{W}^\vee$ eine surjektive Abbildung ist. Weil der Übergang zur dualen Abbildung verträglich mit der Verkettung von Abbildungen ist und die zur Nullabbildung duale Abbildung die Nullabbildung ist, folgt $U = \text{Im}(f) \subseteq \text{Ker}(K^n \rightarrow \tilde{W}^\vee)$. Aus Dimensionsgründen folgt die Gleichheit.

Das bedeutet: Ist $K^{n-r} \rightarrow \tilde{W}$ ein Isomorphismus (wie wir ihn durch Wahl einer Basis von $\tilde{W}$ erhalten), so bekommen wir dual einen Isomorphismus $\tilde{W}^\vee \rightarrow (K^{n-r})^\vee = K^{n-r}$ – dieser entspricht der dualen Basis. Als Verkettung erhalten wir eine Abbildung $g: K^n \rightarrow \tilde{W}^\vee \rightarrow K^{n-r}$ mit Kern U .

Wir können das auch zu einem konkreten Rechenrezept machen: Nach Definition ist $\tilde{W}$ der Kern von $f^\vee$, also die Lösungsmenge des durch $M(f)^t$ gegebenen homogenen linearen Gleichungssystems. Mit dem Gauß-Algorithmus können wir eine Basis von $\tilde{W}$ finden. Schreiben wir die Basisvektoren als die Spalten in eine $n \times (n - r)$ -Matrix B , so ist die Abbildung $g: K^n \rightarrow K^r$ durch die Matrix B^t gegeben. Dann ist U die Lösungsmenge des homogenen linearen Gleichungssystems zu B^t .

Siehe auch Abschnitt 7.6.4. ◇

Das Thema *Dualraum* passt gut in dieses Kapitel über lineare Abbildungen, und die Ergebnisse, die wir erhalten haben, belegen hoffentlich, dass der Begriff von einem gewissen Interesse ist. Für den weiteren Verlauf der Vorlesung Lineare Algebra 1 wird der Begriff allerdings keine große Rolle spielen. In der Linearen Algebra 2 kommen wir dann wieder darauf zurück, wenn wir »Bilinearformen« auf einem Vektorraum studieren.

7.6. Wie berechne ich ...?

7.6.1. Basiswechsel. Zum Basiswechsel gibt es verschiedene Aufgabentypen, die aber größtenteils durch die Formeln aus Satz 7.31 und Korollar 7.33 abgedeckt sind.

Entscheidend ist, sich zu merken, wie die Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ aufgebaut ist: In der j -ten Spalte stehen die Koeffizienten, die $f(b_j)$ in der Darstellung als Linearkombination der Basis $\mathcal{C}$ hat, mit anderen Worten der Koordinatenvektor $c_{\mathcal{C}}(f(b_j))$. (Hier ist $\mathcal{B} = (b_1, \dots, b_n)$.)

Sind $\mathcal{B}, \mathcal{B}'$ Basen von V , so sind die Spalten der Basiswechselmatrix $M_{\mathcal{B}}^{\mathcal{B}'} = M_{\mathcal{B}}^{\mathcal{B}'}(\text{id})$ die Koordinatenvektoren $c_{\mathcal{B}}(b'_j)$, wobei wir $\mathcal{B}' = (b'_1, \dots, b'_n)$ schreiben. Dementsprechend gilt $M_{\mathcal{B}}^{\mathcal{B}'} c_{\mathcal{B}'}(v) = c_{\mathcal{B}}(v)$. So kann man also zwischen Koordinatenvektoren bezüglich verschiedener Basen umrechnen.

Die Abbildung $c_{\mathcal{B}}^{-1}: K^n \rightarrow V$ (für eine Basis $\mathcal{B} = (b_1, \dots, b_n)$ von V) ist gegeben durch die Vorschrift $(a_1, \dots, a_n)^t \mapsto \sum_{j=1}^n a_j b_j$.

7.6.2. Der Kern einer Abbildung. Der Kern einer linearen Abbildung $f_A: K^n \rightarrow K^m$, mit anderen Worten der Kern von A , ist die Lösungsmenge des durch A gegebenen homogenen linearen Gleichungssystems, und wir können mit dem Gauß-Algorithmus eine Basis bestimmen.

Ist $f: V \rightarrow W$ eine lineare Abbildung, deren darstellende Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ bezüglich endlicher Basen $\mathcal{B}$ von V und $\mathcal{C}$ von W wir kennen, so erhalten wir ein kommutatives Diagramm (vergleiche Bemerkung 7.29)

$$\begin{array}{ccccc} \text{Ker}(f) & \longrightarrow & V & \xrightarrow{f} & W \\ \downarrow & & \downarrow c_{\mathcal{B}} & & \downarrow c_{\mathcal{C}} \\ \text{Ker}(M_{\mathcal{C}}^{\mathcal{B}}(f)) & \longrightarrow & K^n & \longrightarrow & K^m, \end{array}$$

wobei die linke vertikale Abbildung, die Einschränkung von $c_{\mathcal{B}}$ auf $\text{Ker}(f)$, ein Isomorphismus $\text{Ker}(f) \rightarrow \text{Ker}(M_{\mathcal{C}}^{\mathcal{B}}(f))$ ist. Sind $u_1, \dots, u_s \in K^n$ eine Basis von $\text{Ker}(M_{\mathcal{C}}^{\mathcal{B}}(f))$, so bilden also $c_{\mathcal{B}}^{-1}(u_1), \dots, c_{\mathcal{B}}^{-1}(u_s)$ eine Basis von $\text{Ker}(f)$.

7.6.3. Das Bild einer Abbildung. Sei zunächst $A \in M_{m \times n}(K)$ und $f = f_A$ die durch A gegebene Abbildung. Das Bild von A ist der von den Spalten von A erzeugte Untervektorraum von K^m , und wir können eine Basis finden, indem wir innerhalb dieses Erzeugendensystems ein linear unabhängiges Erzeugendensystem auswählen (siehe Abschnitt 6.5.2).

Ist $f: V \rightarrow W$ eine lineare Abbildung, deren darstellende Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ bezüglich endlicher Basen $\mathcal{B}$ von V und $\mathcal{C}$ von W wir kennen, so erhalten wir ein kommutatives Diagramm (mit $n = \dim V, m = \dim W$; vergleiche Bemerkung 7.29)

$$\begin{array}{ccccc} V & \xrightarrow{f} & \operatorname{Im}(f) & \longrightarrow & W \\ \downarrow c_{\mathcal{B}} & & \downarrow & & \downarrow c_{\mathcal{C}} \\ K^n & \longrightarrow & \operatorname{Im}(M_{\mathcal{C}}^{\mathcal{B}}(f)) & \longrightarrow & K^m, \end{array}$$

wobei die mittlere vertikale Abbildung, die Einschränkung von $c_{\mathcal{C}}$ auf den Unterraum $\operatorname{Im}(f)$, ein Isomorphismus $\operatorname{Im}(f) \rightarrow \operatorname{Im}(M_{\mathcal{C}}^{\mathcal{B}}(f))$ ist. Bilden $u_1, \dots, u_r \in K^m$ eine Basis von $\operatorname{Im}(M_{\mathcal{C}}^{\mathcal{B}}(f))$, dann bilden $c_{\mathcal{C}}^{-1}(u_1), \dots, c_{\mathcal{C}}^{-1}(u_r)$ eine Basis von $\operatorname{Im}(f)$.

7.6.4. Untervektorraum als Lösungsmenge. Sei $U \subseteq K^n$ ein Untervektorraum. Wir wollen ein lineares Gleichungssystem mit Lösungsmenge U und möglichst wenigen Gleichungen finden. Siehe auch Satz 7.39. Sei $r = \dim U$. Wir wissen, dass wir mindestens $n - r$ Gleichungen benötigen. Ist $u_1, \dots, u_r$ eine Basis von U , so suchen wir eine Matrix $A \in M_{(n-r) \times n}(K)$, so dass $Au_i = 0$ für alle i , und mit $\operatorname{rg}(A) = n - r$. Es folgt dann $U \subset \operatorname{Ker}(A)$, und weil $\dim \operatorname{Ker}(A) = n - \operatorname{rg}(A) = r$, sogar die Gleichheit $U = \operatorname{Ker}(A)$.

Wir schreiben die u_i als die Spalten einer Matrix B . Wir suchen also A mit der Eigenschaft $AB = 0$ und $\operatorname{rg}(A) = n - r$. Die Eigenschaft $AB = 0$ ist äquivalent zu $B^t A^t = 0$. Weil B Rang r hat, hat auch $B^t \in M_{n \times (n-r)}(K)$ den Rang r (Theorem 7.41), also gilt $\dim \operatorname{Ker}(B^t) = n - r$. Wir können eine Basis $c_1, \dots, c_{n-r}$ von $\operatorname{Ker}(B^t)$ mit dem Gauß-Algorithmus berechnen (das bedeutet einfach, das durch B^t gegebene homogene lineare Gleichungssystem zu lösen). Die Matrix C mit den Spalten $c_1, \dots, c_{n-r}$ ist eine Matrix vom Rang $n - r$ mit $B^t C = 0$. Wir können also $A = C^t$ setzen.

Siehe Bemerkung 7.56 für eine theoretische »Aufarbeitung« dieses Rezepts. (Dass das Rezept funktioniert, folgt aber aus dem, was wir hier gesagt haben.)

In dem speziellen Fall, dass U der von den Spalten von $\begin{pmatrix} E_r \\ M \end{pmatrix}$ erzeugte Untervektorraum ist (mit $M \in M_{(n-r) \times n}(K)$), dann ist ein lineares Gleichungssystem mit Lösungsmenge U gegeben durch $(-A \ E_{n-r})x = 0$. Dieses Rezept lässt sich auch leicht auf den Fall verallgemeinern, dass U eine Basis hat, deren Spalten eine Matrix ergeben, die die Einheitsmatrix E_r in anderen Zeilen als den Zeilen $1, \dots, r$ enthält. Eine solche Basis lässt sich immer finden: Startet man mit irgendeiner Basis, so kann man die entsprechende Matrix durch elementare Spaltenumformungen auf »Spaltenstufenform« bringen. Die Spaltenumformungen liefern eine andere Basis *desselben* Untervektorraums (Abschnitt 6.5.5).

Es gibt auch (wie bei vielen der hier diskutierten Aufgabentypen) andere Möglichkeiten, an die Sache heranzugehen: Man kann auch den Beweis von Satz 7.39 zu einem expliziten Rechenrezept machen. Das würde bedeuten, eine Basis von U zu einer Basis von V zu ergänzen; das entspricht der Wahl eines Komplements U' von U . Die Basis von U' können wir als Isomorphismus $U' \rightarrow K^{n-r}$ betrachten, und man muss dann die Abbildung

$$K^n = U \oplus U' \rightarrow U' \xrightarrow{\sim} K^{n-r}$$

bezüglich der Standardbasen beschreiben, um eine Matrix mit Kern U zu finden.

7.6.5. Basis von Durchschnitt und Summe von Untervektorräumen. Wir kommen hier auf die Berechnung von Basen von Durchschnitt und Summe zweier Untervektorräume $U, W \subseteq V$ zurück, und beschreiben das in Abschnitt 6.5.4 vorgestellte Rechenverfahren noch einmal etwas anders. Wir betrachten die »äußere« direkte Summe $U \oplus W$ (also den Vektorraum aller Paare (u, w) mit den komponentenweisen Rechenoperationen) und die lineare Abbildung

$$f: U \oplus W \rightarrow V, \quad (u, w) \mapsto u - w.$$

Dann ist $\text{Im}(f) = U + W$ (wir können also nach dem oben beschriebenen Verfahren eine Basis von $U + W$ finden) und

$$\text{Ker}(f) = \{(v, v) \in U \oplus W; v \in U \cap W\}.$$

Die Abbildung $\text{Ker}(f) \rightarrow U \cap W, (v, v) \mapsto v$, ist also ein Isomorphismus. Ist $(v_1, v_1), \dots, (v_t, v_t)$ eine Basis von $\text{Ker}(f)$, so ist $v_1, \dots, v_t$ eine Basis von $U \cap W$.

Sind $u_1, \dots, u_r$ und $w_1, \dots, w_s$ Basen von U bzw. W , so ist $\mathcal{B} = (u_1, \dots, u_r, w_1, w_s)$ eine Basis von $U \oplus W$ (genau genommen sollten wir $(u_i, 0)$ bzw. $(0, w_i)$ schreiben). Wir fixieren auch eine Basis $\mathcal{C}$ von V (in Abschnitt 6.5.4 hatten wir den Fall $V = K^m$ betrachtet und könnten an dieser Stelle die Standardbasis betrachten). Die Spalten der Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ sind dann die Vektoren

$$c_{\mathcal{C}}(u_1), \dots, c_{\mathcal{C}}(u_r), -c_{\mathcal{C}}(w_1), \dots, -c_{\mathcal{C}}(w_s)$$

Damit können wir nach dem oben beschriebenen Verfahren eine Basis von $\text{Ker}(f)$ finden.

7.7. Ergänzungen *

BEISPIEL 7.57. Sei n eine natürliche Zahl. Wir betrachten den Körper $K = \mathbb{R}$ und darüber den Vektorraum V der Polynomfunktionen »vom Grad $\leq n$ «, d.h. den Vektorraum

$$V = \left\{ f: \mathbb{R} \rightarrow \mathbb{R}; \text{ es existieren } a_0, \dots, a_n, \text{ so dass } f(x) = \sum_{i=0}^n a_i x^i \text{ für alle } x \in \mathbb{R} \right\}.$$

Hier ist n von vorneherein festgelegt, aber es ist erlaubt, dass einige (oder sogar alle) a_i gleich Null sind.

Die übliche Addition und Skalarmultiplikation von Abbildungen in einen Vektorraum (in diesem Fall $\mathbb{R}$) liefern uns eine Vektorraumstruktur auf V (vergleiche Beispiel 6.2).

Die Abbildung

$$\Phi: \mathbb{R}^{n+1} \rightarrow V, \quad (a_0, \dots, a_n)^t \mapsto \left(f(x) = \sum_{i=0}^n a_i x^i \right)$$

ist dann eine lineare Abbildung. Nach Definition von V ist klar, dass sie surjektiv ist. Wir haben gesehen (Korollar 4.28), dass die Koeffizienten a_i durch die Abbildung $f(x)$ eindeutig bestimmt sind. Das bedeutet genau, dass die Abbildung Φ injektiv ist. Es handelt sich also um einen Isomorphismus, und wir sehen insbesondere, dass $\dim V = n + 1$, und genauer, dass $1, x, \dots, x^n$ eine Basis von V bilden (wobei wir hier x^i als Abkürzung für die Polynomfunktion $\mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^i$, schreiben).

Seien $c_0, \dots, c_n$ paarweise verschiedene reelle Zahlen. Die Abbildung

$$\Psi: V \rightarrow \mathbb{R}^{n+1}, \quad f \mapsto (f(c_0), \dots, f(c_n))^t$$

ist linear, und sie hat trivialen Kern, weil die Nullfunktion die einzige Polynomfunktion vom Grad $\leq n$ mit $n + 1$ Nullstellen ist (Satz 4.25). Daher ist die Abbildung injektiv und aus Dimensionsgründen auch surjektiv. Wir können damit schließen:

LEMMA 7.58. *Es gibt für jedes Element $(y_0, \dots, y_n)^t \in \mathbb{R}^{n+1}$ genau eine Polynomfunktion $f: \mathbb{R} \rightarrow \mathbb{R}$ vom Grad $\leq n$ mit $f(c_i) = y_i$ für alle i .*

(Andersherum könnte man das Lemma direkt beweisen – ist $(y_0, \dots, y_n)^t$ einer der Standardbasisvektoren, dann kann man die entsprechende Funktion direkt hinschreiben. Dann wüsste man, dass Ψ ein Isomorphismus ist, und könnte aus der Surjektivität der Abbildung Φ auf deren Injektivität schließen. Das würde einen neuen Beweis von Korollar 4.28 liefern.)

Bis hierhin könnte man in der Diskussion $\mathbb{R}$ durch irgendeinen unendlichen Körper (oder sogar einen Körper der mindestens $n + 1$ Elemente hat) ersetzen. Zum Schluss des Beispiels wollen wir noch erwähnen, dass im Fall der reellen Zahlen auch die Ableitung und das Integral über ein Intervall $[a, b]$ lineare Abbildungen definieren: Die Abbildungen

$$V \rightarrow V, \quad f(x) \mapsto f'(x),$$

und (für $a, b \in \mathbb{R}$)

$$V \rightarrow \mathbb{R}, \quad f(x) \mapsto \int_a^b f(x) dx,$$

sind Homomorphismen von $\mathbb{R}$ -Vektorräumen. Das folgt aus den Rechenregeln, die beschreiben, wie sich Ableitung und Integral für die Summe von Funktionen bzw. das Produkt einer Funktion mit einer reellen Zahl verhalten. Die Aussage über die Linearität des Ableitungsoperators und des Integrals über ein fest gewähltes Intervall gilt natürlich nicht nur für Polynomfunktionen, sondern allgemeiner für differenzierbare bzw. integrierbare Funktionen. $\diamond$

ERGÄNZUNG 7.59. Wir hatten in Beispiel 5.49 schon einige Beispiele für Endomorphismen von $\mathbb{R}^2$ gesehen. In Bemerkung 7.3 haben wir die Bedingungen an eine (bijektive) Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^m$, linear zu sein, geometrisch formuliert. Viele Arten von Abbildungen, die wir aus der Anschauung kennen, sind also lineare Abbildungen, zum Beispiel für den Fall $\mathbb{R}^2 \rightarrow \mathbb{R}^2$: Drehungen um den Ursprung, Spiegelungen an einer Ursprungsgeraden, Scherungen, Streckungen, und Verkettung von Abbildungen dieser Art. (Und ähnlich für den allgemeinen Fall $\mathbb{R}^n \rightarrow \mathbb{R}^m$.) Wir sollten also in der Lage sein, alle diese Abbildungen in unserem Formalismus wiederzufinden und sie insbesondere durch Matrizen zu beschreiben.

Wie üblich, wenn man Anschauung und mathematischen Formalismus zusammenbringen möchte, stehen wir hier allerdings vor dem Problem, dass wir – beispielsweise – den Begriff der Drehung ja bisher gar nicht definiert haben. Da gibt es nun mehrere Möglichkeiten:

- (1) Ich könnte Ihnen einfach eine Definition vorsetzen: Eine Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ heißt *Drehung*, wenn die darstellende Matrix $M(f)$ die Form

$$M(f) = \begin{pmatrix} a & -b \\ b & a \end{pmatrix}$$

für reelle Zahlen a, b mit $a^2 + b^2 = 1$ hat.

- (2) Wir könnten die Theorie der analytischen Geometrie entwickeln und geometrische Begriffe wie Abstand, Winkel, ..., Drehung, Spiegelung, ...definieren, und danach die Verbindung zur linearen Algebra herstellen.

Das würde letzten Endes *zur selben Definition* wie in Teil (1) führen, und würde diese Definition tatsächlich geometrisch motivieren und erklären.

- (3) Oder wir wählen einen naiven Ansatz, nehmen den Standpunkt ein, dass wir ja wissen, was ein Winkel und eine Drehung sind, dass man den Umfang eines Kreises und die Länge eines Kreisbogens messen kann, und was die Funktionen Sinus und Kosinus (im Sinne der »elementaren« Definition als Verhältnis der Längen von Gegenkathete und Hypotenuse bzw. von Ankathete und Hypotenuse im rechtwinkligen Dreieck) sind.

Ansatz (1) lässt sich schnell durchführen, aber Sie würden sich zu recht fragen, warum ausgerechnet Matrizen dieser Form Drehungen heißen sollen; das ist also nicht sehr befriedigend. Andererseits gibt es ein offensichtliches Problem mit Ansatz (2): Es würde einige Zeit kosten,

ihn tatsächlich durchzuführen. Und zwar auch deshalb, weil man für einen befriedigenden Aufbau der Theorie eine ganze Menge Mathematik braucht, die in den Anfängervorlesungen Lineare Algebra I & 2 und Analysis I & 2 gar nicht behandelt wird. (Das soll nicht heißen, dass es so schwierig wäre, dass man es nicht machen *könnte*. Aber im üblichen Aufbau der Anfängervorlesungen werden diese Sachen – etwa die Begriffe des Flächeninhalts/Volumens von Teilmengen von $\mathbb{R}^2$ und der Länge von (nicht geradlinigen) »Kurven« in $\mathbb{R}^n$ – üblicherweise zurückgestellt. In der Maß- und Integrationstheorie, wie sie typischerweise in der Analysis 3 behandelt wird, werden diese Probleme dann (neben weitreichenden Verallgemeinerungen) miterledigt.)

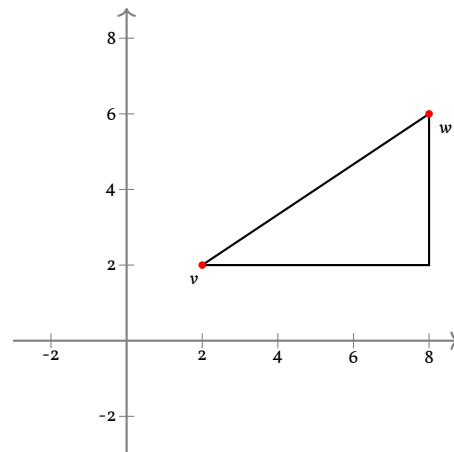
Mehr Informationen dazu finden Sie in Kapitel II.

Ansatz (3) ist insofern unbefriedigend, als ja das Ziel in den Mathematik-Vorlesungen ist, die Theorie rigoros und so vollständig wie möglich aufzubauen, und sich eben nicht auf eine Anschauung zu verlassen, die nicht hinterfragt wird. An dieser Stelle sollen Sie aber wenigstens die Möglichkeit haben, ein bisschen mehr als in (1) oben zu erfahren. Ich wähle hier eine Kompromisslösung, die nur wenig Geometrie benötigt, mit der aber immerhin einiges motiviert werden kann. Wir beginnen recht rigoros, und gehen am Ende im Sinne von Ansatz (3) noch ein bisschen weiter, ohne vollständige Herleitungen der verwendeten Begriffe zu geben.

Sei zunächst $n \geq 1$ eine natürliche Zahl. Wir definieren den *Abstand* von Punkten $(v_i)_i, (w_i)_i$ in $\mathbb{R}^n$ als

$$d(v, w) = \sqrt{\sum_{i=1}^n (w_i - v_i)^2}.$$

Veranschaulichen Sie sich diese Definition, wenigstens im Fall $n = 2$, anhand des Satzes des Pythagoras. Vergleiche die nebenstehende Abbildung: Die Länge der waagerechten Kathete des eingezeichneten rechtwinkligen Dreiecks ist $|w_1 - v_1|$, die der senkrechten ist $|w_2 - v_2|$. Siehe Abschnitt II.2.1 für eine ausführlichere Diskussion.



Wir wollen zunächst Endomorphismen f von $\mathbb{R}^n$ betrachten, die alle Abstände erhalten, das bedeutet, dass für je zwei Punkte $v, w \in \mathbb{R}^n$ gelten soll: $d(f(v), f(w)) = d(v, w)$. Beispiele dafür sind die Beispiele von Drehungen und Spiegelungen, die wir in Beispiel 5.49 gesehen haben. Jedenfalls ist jede abstandserhaltende Abbildung $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ ein Isomorphismus, da für alle $v \neq 0$ gilt, dass $d(f(v), 0) = d(v, 0) \neq 0$, also $f(v) \neq 0$ ist.

SATZ 7.60. Ein Endomorphismus $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ ist genau dann abstandserhaltend, wenn die darstellende Matrix von f die Form

$$\begin{pmatrix} a & -b \\ b & a \end{pmatrix} \quad \text{oder} \quad \begin{pmatrix} a & b \\ b & -a \end{pmatrix}$$

mit $a^2 + b^2 = 1$ hat.

BEWEIS. Für die Standardbasisvektoren e_1, e_2 gilt $d(0, e_1) = 1, d(0, e_2) = 1$, also müssen auch ihre Bilder unter f Abstand 1 vom Ursprung haben. Außerdem gilt $d(e_1, e_2) = \sqrt{2}$, also auch $d(f(e_1), f(e_2)) = \sqrt{2}$.

Es ist anschaulich klar, dass es zu vorgegebenem $f(e_1)$, also zu vorgegebenen Zahlen $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$, genau zwei Möglichkeiten gibt, $f(e_2)$ so zu wählen, dass $f(e_2)$ Abstand 1 zum Ursprung hat (also auf dem »Einheitskreis« liegt) und Abstand $\sqrt{2}$ zu $f(e_1)$ hat.

Es ist auch nicht schwer (wenn auch etwas lästig), die obige Bedingung damit auszurechnen, und ich überlasse für den Moment diese Rechnung Ihnen. Sobald wir darüber gesprochen haben werden, wann zwei Geraden bzw. Vektoren senkrecht zueinander sind, lässt sich das Ganze noch transparenter machen. Siehe Satz stz:isometrie-r2. $\square$

Während anschaulich die eine der beiden Möglichkeiten im Satz eine *Drehung* darstellt, handelt es sich bei der anderen um die *Spiegelung* an einer Ursprungsgeraden. Wir wollen diese Unterscheidung im folgenden etwas genauer diskutieren.

SATZ 7.61. (1) Seien $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$. Sei $f = f_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit

$$A = \begin{pmatrix} a & b \\ b & -a \end{pmatrix}.$$

Dann ist f abstandserhaltend und es gilt $f \circ f = \text{id}$. Ferner existieren $v, v' \in \mathbb{R}^2 \setminus \{0\}$ mit

$$f(v) = v, \quad f(v') = -v'.$$

(2) Sei $v \in \mathbb{R}^2 \setminus \{0\}$. Dann existiert genau eine abstandserhaltende Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ so dass $f(v) = v$ und dass ein Vektor $v' \in \mathbb{R}^2 \setminus \{0\}$ existiert mit $f(v') = -v'$, und die Matrix $M(f)$ hat die Form der Matrix A in Teil (1).

Wir nennen f die Spiegelung an der Geraden $\langle v \rangle$.

BEWEIS. Wir betrachten zunächst Teil (1). Dass $f \circ f = \text{id}$ gilt, ist gleichbedeutend mit $A^2 = E_2$; eine leichte Rechnung zeigt, dass dies tatsächlich der Fall ist. Die Bedingungen an v und v' können wir formulieren als

$$(A - E_2)v = 0, \quad (A + E_2)v' = 0.$$

Die Existenz solcher Vektoren zu zeigen, heißt also einfach zu beweisen, dass diese homogenen linearen Gleichungssysteme nicht nur die triviale Lösung besitzen. Wir können die Determinante von (2×2) -Matrizen benutzen (siehe Beispiel 5.56 und Abschnitt 2.5). Es gilt $\delta(A - E_2) = (a - 1)(-a - 1) - b^2 = -a^2 + 1 - b^2 = 0$, und ähnlich berechnet man $\delta(A + E_2) = 0$.

zu (2). Zunächst zur Existenz: Ist $v = (x, y)^t$, so setzen wir $v' = (-y, x)^t$. Da $v \neq 0$ gilt, überprüft man leicht, dass $\mathcal{B} = (v, v')$ eine Basis von $\mathbb{R}^2$ ist. Es gibt daher eine eindeutig bestimmte lineare Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f(v) = v, f(v') = -v'$. Mit anderen Worten:

Wir können f durch $M_{\mathcal{B}}^{\mathcal{B}}(f) = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ definieren. Es bleibt dann nur nachzurechnen, dass das so definierte f abstandserhaltend ist, und durch eine Matrix wie in Teil (1) beschrieben wird.

Um die Eindeutigkeit zu sehen, schreiben wir wieder $v = (x, y)^t$ und betrachten eine Abbildung f mit den gewünschten Eigenschaften. Aus der Existenz von v' mit $f(v') = -v'$ kann man mit einer ähnlichen Betrachtung ein in Teil (1) folgern, dass die darstellende Matrix von f die Form

$$\begin{pmatrix} a & b \\ b & -a \end{pmatrix}$$

haben muss (und nicht die andere Möglichkeit aus Satz 7.60 eintreten kann). Dass $f(v) = v$ ist, bedeutet also

$$ax + by = x, \quad bx - ay = y.$$

Das können wir als Gleichungssystem in den *Unbestimmten* a und b (und mit Koeffizienten x, y) betrachten. Die Koeffizientenmatrix $\begin{pmatrix} x & y \\ -y & x \end{pmatrix}$ ist invertierbar, weil $v \neq 0$ gilt. Das bedeutet, dass a und b durch v eindeutig bestimmt sind. $\square$

Es ist auch nicht schwierig, aus der Existenzaussage in Teil (2) eine Formel für die darstellende Matrix dieser Abbildung zu gewinnen. Mit den Notationen aus dem Beweis, und wenn wir die Basis (v, v') mit $\mathcal{B}$ und die Standardbasis mit $\mathcal{E}$ bezeichnen, gilt:

$$M_{\mathcal{E}}^{\mathcal{B}} = \begin{pmatrix} x & -y \\ y & x \end{pmatrix}, \quad M(f) = M_{\mathcal{E}}^{\mathcal{B}} \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix} M_{\mathcal{B}}^{\mathcal{E}} = \frac{1}{x^2 + y^2} \begin{pmatrix} y^2 - x^2 & -2xy \\ -2xy & x^2 - y^2 \end{pmatrix}.$$

Allgemeiner nennt man eine abstandserhaltende Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine Spiegelung, wenn es einen Untervektorraum $U \subset \mathbb{R}^n$ der Dimension $n - 1$ gibt, so dass $f(u) = u$ für alle $u \in U$ gilt, und es einen Vektor $v \in \mathbb{R}^n$ mit $f(v) = -v$ gibt.

Matrizen der Form $\begin{pmatrix} a & -b \\ b & a \end{pmatrix}$ (mit $a^2 + b^2 = 1$) definieren eine Drehung. Eine ausführlichere Diskussion dieses Begriffs benötigt eigentlich, wie auch oben schon erwähnt, den Begriff des Winkels. Jedenfalls stimmt es mit der Anschauung überein, dass die Drehung, die e_1 auf den Punkt $\begin{pmatrix} a \\ b \end{pmatrix}$ abbildet, den Punkt e_2 auf $\begin{pmatrix} -b \\ a \end{pmatrix}$ abbildet. Eine Drehung erfüllt sicher die geometrische Beschreibung linearer Abbildungen in Bemerkung 7.3, so dass die folgende Definition sinnvoll ist.

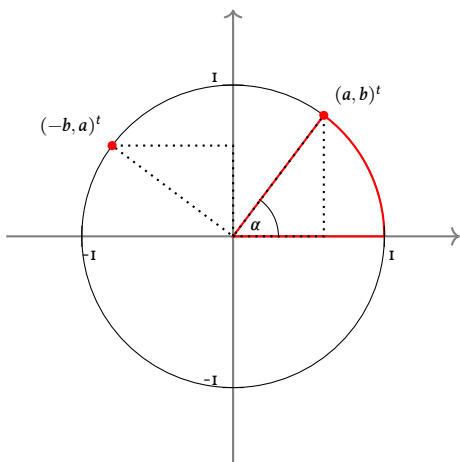
DEFINITION 7.62. Seien $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$. Sei $f = f_A: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit

$$A = \begin{pmatrix} a & -b \\ b & a \end{pmatrix}.$$

Dann ist f abstandserhaltend. Wenn $a = 1$, dann ist $A = E_2$ und $f = \text{id}$. Andernfalls existiert kein $v \in \mathbb{R}^2 \setminus \{0\}$ mit $f(v) = v$.

Wir nennen f eine *Drehung*. Matrizen A der obigen Form nennt man auch *Drehmatrizen*. $\dashv$

Die Sachen, die hier noch zu beweisen sind, zeigt man mit ähnlichen Argumenten wie bei der Betrachtung von Spiegelungen.



Die Abbildung zeigt die Drehung ρ um den Winkel α . Es gilt dann $\rho(e_1) = \begin{pmatrix} a \\ b \end{pmatrix}$ und $\rho(e_2) = \begin{pmatrix} -b \\ a \end{pmatrix}$. Es ist dann $a = \cos(\alpha)$, $b = \sin(\alpha)$.

Auf diesen Webseiten können Sie mit Drehungen und Spiegelungen »herumspielen«:

<https://math.ug/applets/drehungen.html>^a,

<https://math.ug/applets/spiegelungen.html>^b.

^aDrehungen

^bSpiegelungen

Allgemein kann man (allerdings etwas behelfsmäßig, siehe Kapitel 11 für eine ausführlichere Diskussion) definieren: Eine lineare Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt *Drehung*, wenn sie abstandserhaltend ist und eine lineare Abbildung $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit $f = g \circ g$ existiert. (Anschaulich

kann man jede Drehung in zwei gleiche Schritte aufteilen; in $\mathbb{R}^2$ heißt das einfach, dass man zweimal um den halben Drehwinkel dreht. Eine Spiegelung kann man hingegen nie in der Form $g \circ g$ schreiben, und diese Bedingung dient dazu, die Spiegelungen auszuschließen.)

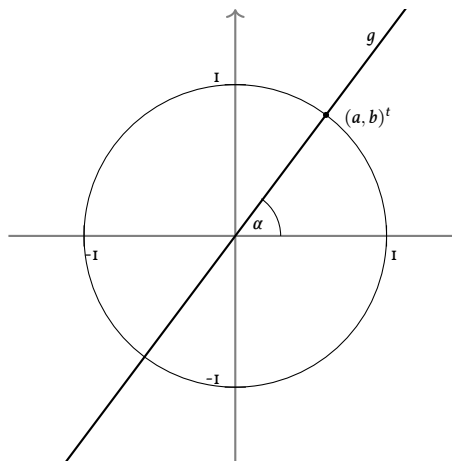
Selbst im Fall $n = 2$ ist es aber nicht-trivial, für eine Matrix $\begin{pmatrix} a & b \\ -b & a \end{pmatrix}$ die Matrix für die Drehung »um den halben Drehwinkel« zu berechnen. Hier kommen die trigonometrischen Funktionen Sinus und Kosinus ins Spiel. Wir wollen die Diskussion noch ein bisschen unter der Annahme fortführen, dass wir wissen, was der Winkel zwischen zwei sich schneidenden Geraden ist, und was die Sinus- und

Sei $(a, b)^t$ ein Punkt auf dem Einheitskreis, also mit Abstand 1 zum Ursprung, sei g die Gerade durch $(0, 0)^t$ und $(a, b)^t$, und sei α der Winkel den die waagerechte Koordinatenachse und g einschließen, wobei wir von der waagerechten Koordinatenachse gegen den Uhrzeigersinn bis zum Punkt $(a, b)^t$ laufen (vergleiche die Abbildung). Kosinus-Funktion sind.

Es gilt dann $a = \cos(\alpha)$ und $b = \sin(\alpha)$. Umgekehrt ist für jeden Winkel α der Punkt $(\cos(\alpha), \sin(\alpha))^t$ ein Punkt auf dem Einheitskreis, denn $\cos(\alpha)^2 + \sin(\alpha)^2 = 1$. Die Drehung um den Winkel α gegen den Uhrzeigersinn wird also durch die Matrix

$$\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}$$

beschrieben.



Dass für die Drehungen ρ_α und ρ_β um den Winkel α bzw. β die Verkettung $\rho_\alpha \circ \rho_\beta$ die Drehung um den Winkel $\alpha + \beta$ ist, bedeutet

$$\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix} \begin{pmatrix} \cos(\beta) & -\sin(\beta) \\ \sin(\beta) & \cos(\beta) \end{pmatrix} = \begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix},$$

oder für die Einträge ausgeschrieben:

$$\cos(\alpha + \beta) = \cos(\alpha) \cos(\beta) - \sin(\alpha) \sin(\beta),$$

$$\sin(\alpha + \beta) = \sin(\alpha) \cos(\beta) + \cos(\alpha) \sin(\beta).$$

Diese Gleichungen nennt man die *Additionstheoreme* für die Sinus- und Kosinus-Funktionen. Die weitere Diskussion verschieben wir ins Kapitel 11.

In der Linearen Algebra 2 werden wir die Theorie der Bilinearformen kennenlernen, die einige in der analytischen Geometrie besonders wichtige Konzepte in noch allgemeinerem Rahmen bereitstellt. Dann wird sich auch in natürlicher Weise die Gelegenheit ergeben, auf die hier genannten Beispiele noch einmal zurückzukommen.

Zum Schluss sei noch die Verbindung zur Matrixdarstellung der komplexen Zahlen (Bemerkung 5.42) erwähnt: Wenn wir $\mathbb{C} = \mathbb{R}^2$ identifizieren, dann können wir als Ergebnis der obigen Diskussion sagen, dass die (lineare) Abbildung $\mathbb{C} \rightarrow \mathbb{C}, x \mapsto zx$, die durch Multiplikation mit einer komplexen Zahl $z = a + bi$ mit $a^2 + b^2 = 1$ gegeben ist, geometrisch gesehen eine Drehung um den Ursprung ist.

Für $z = a + bi$ nennt man die Zahl $\sqrt{a^2 + b^2}$ auch den Betrag von z . Im Sinne der Identifikation $\mathbb{C} = \mathbb{R}^2$ ist dies genau der Abstand des Punktes z zum Ursprung. □ Ergänzung 7.59

BEISPIEL 7.63. Das folgende Beispiel spielt in der Vorlesung Lineare Algebra weiter keine Rolle, soll aber exemplarisch die neu eingeführten Begriffe illustrieren und zeigen, dass der Begriff der linearen Abbildung es oft erlaubt, mathematische Sachverhalte prägnant auszudrücken.

Sei $V = \mathbb{R}^{\mathbb{N}}$ der $\mathbb{R}$ -Vektorraum aller reellen Zahlenfolgen $(a_n)_{n \in \mathbb{N}}$, $a_n \in \mathbb{R}$. Sei U die Teilmenge der konvergenten Folgen. Diese bildet einen Untervektorraum, denn die gliedweise Summe zweier konvergenter Folgen ist wieder eine konvergente Folge, und ebenso erhält die Multiplikation aller Terme einer Folge mit einer reellen Zahl a die Konvergenzeigenschaft.

Wir erhalten eine Abbildung $L: U \rightarrow \mathbb{R}$, $(a_n)_n \mapsto \lim_{n \rightarrow \infty} a_n$. Dass für alle konvergenten Folgen $(a_n)_n$ und $(b_n)_n$ und für $a \in \mathbb{R}$ die Grenzwertsätze

$$\lim_{n \rightarrow \infty} (a_n + b_n) = \lim_{n \rightarrow \infty} a_n + \lim_{n \rightarrow \infty} b_n, \quad \lim_{n \rightarrow \infty} aa_n = a \lim_{n \rightarrow \infty} a_n$$

gelten, ist äquivalent dazu, dass L eine lineare Abbildung ist.

Der Beweis der Grenzwertsätze oder äquivalent der Tatsache, dass L linear ist, wird durch die neue Terminologie natürlich nicht einfacher.

(Der – ebenfalls richtige – Grenzwertsatz

$$\lim_{n \rightarrow \infty} (a_n b_n) = \lim_{n \rightarrow \infty} a_n \cdot \lim_{n \rightarrow \infty} b_n$$

für konvergente Folgen $(a_n)_n$, $(b_n)_n$ wird durch den Begriff der linearen Abbildung nicht abgedeckt. Das zeigt, dass es in dieser Situation noch passender wäre, stattdessen über Abbildungen zu sprechen, die auch mit der komponentenweisen Multiplikation auf $\mathbb{R}^{\mathbb{N}}$ und der Multiplikation auf $\mathbb{R}$ verträglich sind. Wir werden in der Linearen Algebra 2 dieser Art von Abbildungen den Namen *Ringhomomorphismus* geben.) $\diamond$

ERGÄNZUNG 7.64 (Rechtecke mit Quadraten überdecken). Eine Anwendung von Satz 7.14 ist die folgende geometrische Tatsache: Wir betrachten ein Rechteck und die Frage, ob es durch (endlich viele) Quadrate lückenlos überdeckt werden kann.



Es ist leicht zu sehen, dass das nur möglich ist, wenn die Seiten aller Quadrate parallel sind zu den Seiten des Rechtecks, denn jedenfalls müssen die Quadrate, die in den vier Ecken des Rechtecks liegen, diese Eigenschaft haben, und von dort kann man sich weiter »vorarbeiten«.

Wenn $a, b \in \mathbb{Q}$ sind, dann lässt sich das Rechteck mit Seitenlängen a und b durch Quadrate überdecken. (Das ist »offensichtlich«. Warum? Lassen Sie sich nicht von der Zeichnung verwirren ...)

SATZ 7.65. Sei R ein Rechteck mit Seitenlängen $a, b \in \mathbb{R}_{>0}$. Dabei sei $a \in \mathbb{Q}$ und $b \in \mathbb{R} \setminus \mathbb{Q}$. Dann lässt sich R nicht vollständig durch (endlich viele) Quadrate überdecken, die sich nicht überlappen.

BEWEIS. Angenommen, R ließe sich durch Quadrate $Q_1, \dots, Q_r$ vollständig und ohne Überlappung überdecken. Sei q_i die Seitenlänge von Q_i . Insbesondere ist ab (der Flächeninhalt von R) gleich der Summe $q_1^2 + \dots + q_r^2$. Das allein ist aber kein Widerspruch – natürlich lässt sich ab als Summe von Quadratzahlen schreiben, da wir in $\mathbb{R}$ aus allen positiven Zahlen die Quadratwurzel ziehen können. Um einen Widerspruch zu erhalten, müssen wir die zusätzliche geometrische Information benutzen, dass die Quadrate so in das Rechteck gelegt werden können, dass sie es genau passend überdecken.

Wir betrachten dazu den $\mathbb{Q}$ -Vektorraum $\mathbb{R}$. (Als $\mathbb{Q}$ -Vektorraum ist $\mathbb{R}$ »riesig«, insbesondere nicht endlich erzeugt (Beispiel 6.55), aber das muss uns an dieser Stelle nicht stören. Insbesondere brauchen wir nicht zu benutzen, dass der $\mathbb{Q}$ -Vektorraum $\mathbb{R}$ eine Basis besitzt.)

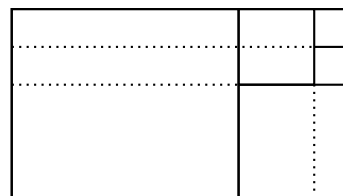
Sei $V \subseteq \mathbb{R}$ der $\mathbb{Q}$ -Untervektorraum, der erzeugt wird von a, b und den Zahlen q_i .

Dann sind a und b linear unabhängig, denn wäre $\alpha a + \beta b = 0$ mit $\alpha, \beta \in \mathbb{Q}$ eine nicht-triviale Linearkombination, so muss $\beta \neq 0$ sein. Dann folgt aber $b = -\frac{\alpha}{\beta} a \in \mathbb{Q}$, ein Widerspruch zu unserer Voraussetzung $b \notin \mathbb{Q}$.

Wir können das Paar (a, b) also zu einer Basis von V ergänzen. Nach Satz 7.14 existiert eine lineare Abbildung $f: V \rightarrow V$ mit $f(a) = 1, f(b) = -1$. Wir definieren für jedes Rechteck P mit Kantenlängen $x, y \in V$: $\varphi(P) := f(x)f(y)$. Zum Beispiel gilt $\varphi(R) = -1$, und für das Quadrat Q_i gilt $\varphi(Q_i) = f(q_i)^2 \geq 0$.

Wir zeigen nun: φ verhält sich additiv für Überdeckungen, d.h. wenn P irgendein Rechteck ist, das durch andere Rechtecke $P_1, \dots, P_r$ vollständig und ohne Überlappung überdeckt wird, und alle auftretenden Kantenlängen in V liegen, dann gilt $\varphi(P) = \sum_{i=1}^r \varphi(P_i)$. Daraus folgt jedenfalls der Satz, denn wenn es eine Überdeckung des Rechtecks R durch Quadrate $Q_1, \dots, Q_r$ gäbe, so hätten wir $-1 = \varphi(R) = \sum \varphi(Q_i) \geq 0$.

Um die Behauptung über die Additivität von φ einzusehen, sei also eine Überdeckung eines Rechtecks P durch Rechtecke P_i gegeben. Wir zerlegen das Rechteck noch weiter, indem wir alle Seiten der P_i über die ganze Länge des Rechtecks P verlängern, wie durch die gepunkteten Linien in der Abbildung am Beispiel der oben gezeigten Überdeckung eines Rechtecks durch Quadrate gezeigt. In dieser feineren Zerlegung sind wieder alle Kantenlängen in V .



Da sich jedes der P_i und auch P durch eine geeignete Auswahl von diesen kleinen Rechtecken überdecken lässt, genügt es, die Additivität der Abbildung φ für Überdeckungen dieser einfachen Form zu zeigen, also für solche, die dadurch entstehen, dass in P einige zu den Seiten parallele Strecken über die ganze Seitenlänge eingezeichnet werden. Sind $x_1, \dots, x_r$ die Längen der Abschnitte, in die die eine Seite unterteilt wird, und $y_1, \dots, y_s$ die Längen der Abschnitte der anderen Seite (die Unterteilung besteht also aus rs kleinen Rechtecken), so gilt

$$\varphi(P) = f\left(\sum_{i=1}^r x_i\right)f\left(\sum_{j=1}^s y_j\right) = \sum_{ij} f(x_i)f(y_j),$$

weil die Abbildung f linear ist. Die Summe auf der rechten Seite ist genau die Summe aller $\varphi(P_i)$ für die Rechtecke P_i der verfeinerten Zerlegung von P . $\square$

Quelle: [Ma] Ch. 12; dort gibt es auch weiterführende Literaturverweise. $\square$ Ergänzung 7.64

ERGÄNZUNG 7.66 (Der Page-rank-Algorithmus, Fortsetzung). Dies ist die Fortsetzung der Diskussion aus Frage 2.7 und Ergänzung erg-page-rank-1. Wir hatten dort die folgende Situation erreicht.

Gegeben ist eine Matrix $G \in M_N(\mathbb{R})$, deren Einträge alle positiv sind, und deren Spaltensummen alle $= 1$ sind. Wir möchten das lineare Gleichungssystem $(G - E_N)x = 0$ lösen, und genauer eine eindeutig bestimmten Lösungsvektor $x \in \mathbb{R}^N$ finden, in dem alle Einträge zwischen 0 und 1 liegen und so dass die Summe der Einträge $= 1$ ist. (In der Praxis ist zudem die Zahl N sehr groß.)

Das folgende Lemma wird sich als der Schlüssel erweisen, um die Eindeutigkeit der Lösung (mit den oben genannten Eigenschaften) zu zeigen.

LEMMA 7.67. Seien $N \geq 2$ und $x, y \in \mathbb{R}^N$ linear unabhängige Vektoren. Dann gibt es eine Linearkombination $\alpha x + \beta y$, die sowohl positive als auch negative Einträge hat.

BEWEIS. Ist $N = 2$, so gilt $\langle x, y \rangle = \mathbb{R}^2$, und die Behauptung ist klar. Den allgemeinen Fall können wir darauf zurückführen; entweder durch eine elementare Überlegung, oder mit dem folgenden Argument: Wir betrachten x und y als die Spalten einer $(N \times 2)$ -Matrix. Nach Voraussetzung ist der Spaltenrang $= 2$, also auch der Zeilenrang. Es gibt also zwei Zeilen, die linear unabhängig sind. Auf diese beiden Zeilen können wir den schon behandelten Fall $N = 2$ anwenden. Es gibt also eine Linearkombination von x und y , die in diesen beiden Zeilen je einen positiven und einen negativen Eintrag hat. $\square$

SATZ 7.68. Sei G eine quadratische Matrix der Größe N über $\mathbb{R}$, deren Einträge alle positiv und deren Spaltensummen alle gleich 1 sind.

Dann gilt $\operatorname{rg}(G - E_N) = N - 1$.

Mit anderen Worten: das homogene lineare Gleichungssystem $(G - E_N)x = 0$ hat eine nicht-triviale Lösung, und je zwei Lösungen unterscheiden sich nur durch Multiplikation mit einem Skalar.

Ferner gilt: Ist $x \in \operatorname{Ker}(G - E_N)$ ein Element $\neq 0$ der Lösungsmenge, dann sind alle Einträge von x positiv, oder alle Einträge von x negativ.

BEWEIS. Da alle Spaltensummen von G gleich 1 sind, sind die Spaltensummen von $G - E_N$ alle gleich 0, oder mit anderen Worten: Die Summe aller Zeilen von $G - E_N$ ist der Nullvektor. Dies liefert uns eine nicht-triviale Linearkombination der Zeilen, die den Nullvektor darstellt, und es folgt $\operatorname{rg}(G - E_N) < N$, also $\operatorname{Ker}(G - E_N) \neq 0$.

Für die nächsten Schritte müssen wir benutzen, dass die Einträge von G positiv sind. (Insbesondere ist dieser Teil der Aussage nicht über beliebigen Körpern sinnvoll.) Zuerst zeigen wir die Aussage am Ende des Satzes. Sei also $x \in \operatorname{Ker}(G - E_N)$, $x \neq 0$. Wenn x sowohl positive als auch negative Einträge hätte, dann wäre

$$\left| \sum_{j=1}^N G_{ij} x_j \right| < \sum_{j=1}^N G_{ij} |x_j|,$$

denn alle G_{ij} sind > 0 . Die linke Seite ist gleich $|x_i|$, denn nach Voraussetzung ist $Gx = E_N x = x$. Wenn wir diese Ungleichungen über alle i aufsummieren, dann erhalten wir

$$\sum_{i=1}^N |x_i| < \sum_{i,j=1}^N G_{ij} |x_j| = \sum_{j=1}^N \left(|x_j| \sum_{i=1}^N G_{ij} \right) = \sum_{j=1}^N |x_j|,$$

ein Widerspruch!

Nun können wir auch beweisen, dass $\operatorname{Ker}(G - E_N)$ eindimensional ist. Dazu benutzen wir Lemma 7.67, das besagt, dass zu zwei linear unabhängigen Vektoren $x, y \in \mathbb{R}^N$ immer eine Linearkombination $\alpha x + \beta y$ existiert, die sowohl positive als auch negative Einträge hat. Deshalb zeigt die gerade bewiesene Aussage, dass es in $\operatorname{Ker}(G - E_N)$ keine zwei linear unabhängigen Vektoren geben kann. $\square$

Als unmittelbare Folgerung erhalten wir:

KOROLLAR 7.69. Es gibt genau eine Lösung $x \in \mathbb{R}^N$ des linearen Gleichungssystems $(G - E_N)x = 0$ mit der Eigenschaft $\sum_{i=1}^N x_i = 1$, und für diese Lösung gilt $0 \leq x_i \leq 1$ für alle i .

Wir werden später noch etwas mehr über Verallgemeinerungen des Satzes und darüber sagen, wie man in der Praxis versuchen könnte, diese Lösung zu berechnen. (Denken Sie daran, dass das Verfahren auch funktionieren sollte, wenn die Matrix G mehrere Milliarden Zeilen hat ...). Das passt aber besser in den Abschnitt über Eigenwerte, siehe Ergänzung 10.24.

□ Ergänzung 7.66

ERGÄNZUNG 7.70 (Das Vier-Unterraum-Problem). Für den Moment ist diese Ergänzung nur ein Platzhalter: Das *Vier-Unterraum-Problem* ist eine Frage, die recht einfach zu formulieren ist und erst (relativ) kürzlich (um 1970 von Gelfand und Ponomarev [GP] und unabhängig davon von Nazarova [Na1], [Na2]) beantwortet wurde. Es gibt ähnliche Klassifikationsprobleme, die immer noch unbeantwortet sind. Hier in der Vorlesung wird die Lineare Algebra zwar überwiegend in einer extrem »optimierten« und »geglätteten« Form präsentiert, wo alle Bausteine gut zusammenpassen (und die historische Entwicklung, die natürlich auch Schwierigkeiten und Sackgassen beinhaltet, ausgeblendet). Aber auch in diesem Gebiet der Mathematik gibt es offene Fragen und wird auch aktuell geforscht.

Besonders interessant ist es, und auch das ist nicht selten, wenn solche Fragen der Linearen Algebra durch konkrete Probleme aus anderen Bereichen der Mathematik, oder durch mögliche Anwendungen, motiviert sind.

Bis ich dazu komme, etwas mehr zum Vier-Unterraum-Problem zu schreiben, können Sie – wenn Sie der Sache schon jetzt nachgehen möchten – in die folgenden Quellen schauen (die oben zitierten Arbeiten sind vermutlich schwerer aufzutreiben und auch schwieriger zu lesen).

Recht zugänglich ist die Arbeit

G. Medina, A. Zavadskiy, *The four subspace problem: An elementary solution*, Linear Algebra and its Applications **392** (2004), 11–23,
<https://doi.org/10.1016/j.laa.2004.05.010>

Ein anderer Ansatz wird in der Masterarbeit [The 4 Subspace Problem](#)^a von T. Forbregd (Norwegian University of Science and Technology 2008) vorgestellt.

^ahttps://ntnuopen.ntnu.no/ntnu-xmlui/bitstream/handle/11250/258466/348722_FULLTEXT01.pdf

□ Ergänzung 7.70

KAPITEL 8

Gruppen

In diesem Kapitel lernen wir die Struktur der Gruppe kennen. Ähnlich wie bei Körpern und Vektorräumen wird mit dem Begriff der Gruppe ein »Rechenbereich« beschrieben, in dem gewisse Rechenregeln gelten. Anders als bei Körpern und Vektorräumen ist in einer Gruppe allerdings nur eine einzige Verknüpfung gegeben. Die Bedingungen an diese Verknüpfung werden Sie sofort an Teile der Körper- und Vektorraumdefinitionen erinnern und in der Tat kann man diese Definitionen noch ein bisschen kürzer formulieren, wenn der Gruppenbegriff schon zur Verfügung steht.

In einem abstrakten Sinn kann man sagen, dass die Theorie der Gruppen die mathematische Art und Weise ist, über Symmetrie zu sprechen. Insofern ist es nicht erstaunlich, dass Gruppen fast überall in der Mathematik anzutreffen sind. In der Algebra-Vorlesung wird die Gruppentheorie weiter vertieft.

8.1. Grundlagen

8.1.1. Definition und einfache Beispiele. Wir beginnen mit der Definition einer Gruppe.

DEFINITION 8.1. Eine *Gruppe* ist ein Paar $(G, \cdot)$ bestehend aus einer Menge G und einer Verknüpfung $\cdot : G \times G \rightarrow G$, $(g, h) \mapsto g \cdot h$, so dass gilt:

- (a) (Assoziativgesetz). Für alle $g, h, j \in G$ gilt: $(g \cdot h) \cdot j = g \cdot (h \cdot j)$.
- (b) (Neutrales Element). Es gibt ein Element $e \in G$, so dass für alle $g \in G$ gilt: $e \cdot g = g, g \cdot e = g$.
- (c) (Inverse Elemente). Für jedes $g \in G$ existiert ein Element $h \in G$ mit $g \cdot h = h \cdot g = e$.

—

BEMERKUNG 8.2. (1) Das neutrale Element e der Verknüpfung wie in (b) ist eindeutig bestimmt (das Argument haben wir schon bei den Körpern und Vektorräumen gesehen; sind e, e' Elemente von G mit der Eigenschaft des neutralen Elements, so gilt $e = ee' = e'$). Insofern hängt Teil (c) der Definition nicht von einer Wahl ab.

- (2) Zu jedem $g \in G$ ist das inverse Element h eindeutig bestimmt. Denn sind h, h' Elemente mit dieser Eigenschaft, so folgt $h = h(gh') = (hg)h' = h'$. Auch dieses Argument haben wir schon gesehen, als wir über Körper gesprochen haben. Wenn die Verknüpfung als Multiplikation, also mit dem Punkt $\cdot$ geschrieben wird, dann bezeichnen wir das Inverse von g mit g^{-1} .

◇

Die Verknüpfung kann man statt mit $\cdot$ natürlich auch mit irgendeinem anderen Symbol bezeichnen, und wir werden unten Beispiele sehen, wo sich das anbietet. Wenn man den Punkt $\cdot$ verwendet, dann sagt man, die Gruppe G werde *multiplikativ* geschrieben, und man wendet die üblichen Konventionen an: Das neutrale Element wird auch mit 1 (oder 1_G) statt mit e

bezeichnet, und der Punkt wird zwischen Buchstaben, die Gruppenelemente bezeichnen, (und/oder Klammern) oft ausgelassen.

Aus der Existenz inverser Elemente in einer Gruppe folgt:

LEMMA 8.3 (Kürzungsregel). Seien G eine Gruppe, $g, h, h' \in G$.

- (1) Wenn $gh = gh'$ gilt dann folgt $h = h'$.
- (2) Wenn $hg = h'g$ gilt dann folgt $h = h'$.

In einigen der Gruppen, mit denen wir zu tun haben, gilt zusätzlich das Kommutativgesetz, und wir definieren:

DEFINITION 8.4. Eine Gruppe G heißt *kommutativ* oder *abelsch*, wenn das Kommutativgesetz gilt:

$$g \cdot h = h \cdot g \quad \text{für alle } g, h \in G.$$

—

Die Bezeichnung *abelsch* bezieht sich auf den Mathematiker [Niels Abel](https://de.wikipedia.org/wiki/Niels_Henrik_Abel)¹ (1802–1829; trotz seines frühen Todes mit nur 26 Jahren hat er sehr bedeutende Beiträge zur Mathematik geleistet).

Für kommutative Gruppen schreibt man die Verknüpfung oft als Addition $+$ (und dann das Inverse von g als $-g$, und man definiert $g - h := g + (-h)$). Auch multiplikativ geschriebene Gruppen können natürlich kommutativ sein, aber jedenfalls verwendet man das Zeichen $+$ nur für kommutative Verknüpfungen. Wird die Verknüpfung mit $+$ bezeichnet, dann bezeichnet man das neutrale Element üblicherweise als 0 . Im Gegensatz zum Punkt $\cdot$ kann das $+$ aber nicht ausgelassen werden.

BEISPIEL 8.5. (1) Die Gruppe $\{1\}$ (multiplikativ geschrieben) (bzw. die Gruppe $\{0\}$ (additiv geschrieben)) nennen wir die *triviale Gruppe*.

- (2) $(\mathbb{Z}, +)$ ist eine kommutative Gruppe.
- (3) Für jedes $n \geq 1$ ist $(\mathbb{Z}/n, +)$ eine kommutative Gruppe.
- (4) Ist K ein Körper, so sind $(K, +)$ und $(K^\times, \cdot)$ kommutative Gruppen. Das erklärt den Namen *multiplikative Gruppe* für $K^\times$.
- (5) Ist K ein Körper und V ein K -Vektorraum, so ist $(V, +)$ eine kommutative Gruppe. Spezialfälle sind K^n und $M_{n \times m}(K)$.
- (6) Sei X eine Menge. Die Menge $\text{Bij}(X)$ der bijektiven Abbildungen $X \rightarrow X$ mit der Verkettung $\circ$ von Abbildungen als Verknüpfung ist eine Gruppe.

Das neutrale Element ist die identische Abbildung $\text{id}_X: X \rightarrow X, x \mapsto x$. Das Inverse einer Bijektion f ist ihre Umkehrabbildung.

Hat X mehr als zwei Elemente, so ist die Gruppe $\text{Bij}(X)$ nicht kommutativ. Im Fall $X = \{1, \dots, n\}$, $n \geq 1$, schreibt man auch $S_n := \text{Bij}(X)$ und nennt diese Gruppe die *symmetrische Gruppe*. Die Elemente dieser Gruppe bezeichnet man als *Permutationen (der Menge $\{1, \dots, n\})$* . Siehe Abschnitt 8.3.

- (7) Sei V ein Vektorraum über dem Körper K , und sei $\text{Aut}(V)$ die Menge aller Automorphismen $V \xrightarrow{\sim} V$. Diese bildet mit der Verknüpfung von Abbildungen eine Gruppe. Das neutrale Element ist die Identitätsabbildung, das zu $f \in \text{Aut}(V)$ inverse Element ist der Umkehrhomomorphismus f^{-1} von f .

¹https://de.wikipedia.org/wiki/Niels_Henrik_Abel

- (8) Sind G_1, G_2 Gruppen, so ist $G_1 \times G_2$ mit der komponentenweisen Verknüpfung eine Gruppe, das *Produkt* der Gruppen G_1 und G_2 . Allgemeiner kann man für jede Familie $G_i, i \in I$, von Gruppen das Produkt $\prod_{i \in I} G_i$ (mit der komponentenweisen Verknüpfung $(g_i)_i \cdot (h_i)_i = (g_i h_i)_i$) betrachten.

◇

Ein besonders wichtiges Beispiel ist für uns die *allgemeine lineare Gruppe*. Die Bezeichnung GL_n ist als Abkürzung des englischen Begriffs *general linear group* zu verstehen.

DEFINITION 8.6 (Allgemeine Lineare Gruppe). Seien K ein Körper und $n \in \mathbb{N}$. Die Teilmenge $GL_n(K) \subset M_{n \times n}(K)$ der invertierbaren Matrizen ist eine Gruppe bezüglich des Matrizenprodukts und heißt die *allgemeine lineare Gruppe (vom Rang (oder Grad) n über K)*. $\dashv$

Für $n = 1$ können wir $GL_1(K)$ und $K^\times$ identifizieren. Für $n > 1$ ist die Gruppe $GL_n(K)$ nicht kommutativ.

ERGÄNZUNG 8.7 (Ringe und Körper). Mit dem Begriff der Gruppe können wir die Definition eines Körpers, Definition 4.1, kürzer fassen: Ein Körper ist eine Menge K zusammen mit Verknüpfungen $+: K \times K \rightarrow K$ und $\cdot: K \times K \rightarrow K$, so dass gilt:

- (a) $(K, +)$ ist eine abelsche Gruppe (mit neutralem Element 0),
- (b) $(K \setminus \{0\}, \cdot)$ ist eine abelsche Gruppe,
- (c) Für $+$ und $\cdot$ gilt das Distributivgesetz: $a(b + c) = ab + ac$ für alle $a, b, c \in K$.

(Aus der ersten und dritten Bedingung folgt, dass $0a = 0$ für alle $a \in K$; es ist daher kein Problem, dass in der zweiten Bedingung das Element 0 ganz ausgeschlossen wird.)

Ähnlich definiert man den Begriff des *Rings* (der im kommenden Semester eine größere Rolle spielen wird): Ein Ring ist eine Menge R zusammen mit Verknüpfungen $+$ und $\cdot$, so dass gilt:

- (a) $(R, +)$ ist eine abelsche Gruppe (mit neutralem Element 0),
- (b) Die Multiplikation $\cdot$ ist assoziativ.
- (c) Für $+$ und $\cdot$ gilt das Distributivgesetz: $a(b + c) = ab + ac$, $(a + b)c = ac + bc$ für alle $a, b, c \in R$.

Besitzt die Multiplikation auf R ein neutrales Element, so nennt man R auch einen *Ring mit Eins*. Ist sie kommutativ, so nennt man R einen *kommutativen Ring*.

Zum Beispiel ist $\mathbb{Z}$ mit der üblichen Addition und Multiplikation ein kommutativer Ring mit Eins. Man spricht daher auch vom *Ring der ganzen Zahlen*. Für alle $n \geq 1$ ist $\mathbb{Z}/n$ ein kommutativer Ring mit Eins, den man auch als den *Restklassenring (von $\mathbb{Z}$) modulo n* bezeichnet.

Der *Nullring* ist der Ring $R = \{0\}$ (mit den einzig möglichen Verknüpfungen $0 + 0 = 0$, $0 \cdot 0 = 0$). Dies ist ebenfalls ein kommutativer Ring mit Eins. Ist umgekehrt R irgendein Ring, in dem Null- und Einselement übereinstimmen, so folgt, dass R der Nullring ist.

Ein Körper ist demnach ein kommutativer Ring mit Eins, der mehr als ein Element hat und in dem jedes Element $\neq 0$ ein inverses Element bezüglich der Multiplikation besitzt.

Ein *Schiefkörper* (Definition 4.10) ist ein (nicht notwendig kommutativer) Ring mit Eins, der mehr als ein Element hat und in dem jedes Element $\neq 0$ ein inverses Element bezüglich der Multiplikation besitzt. Siehe Ergänzungen 4.9, 6.65, 10.20, 4.11, 5.64.

BEISPIEL 8.8. Seien K ein Körper und $n \in \mathbb{N}$. Dann ist die Menge $M_n(K)$ der quadratischen Matrizen der Größe n über K mit der üblichen Addition und Multiplikation von Matrizen ein Ring mit Eins. Ist $n > 1$, so ist dieser Ring nicht kommutativ. $\diamond$

□ Ergänzung 8.7

ERGÄNZUNG 8.9 (Die Einheitengruppe von $\mathbb{Z}/n$). Sei $n > 1$ eine natürliche Zahl. Dann ist $\mathbb{Z}/n$ wie oben bemerkt eine Gruppe bezüglich der Addition. Natürlich ist $\mathbb{Z}/n$ keine Gruppe bezüglich der Multiplikation, weil 0 kein inverses Element besitzt. Die Menge $\mathbb{Z}/n \setminus \{0\}$ ist genau dann eine Gruppe bezüglich der Multiplikation, wenn jedes Element ein Inverses besitzt, und wir haben gesehen, dass das genau dann der Fall ist, wenn n eine Primzahl ist (siehe Korollar 4.17 und die Diskussion davor).

Für beliebiges $n > 1$ definieren wir

$$(\mathbb{Z}/n)^\times := \{a \in \mathbb{Z}/n \setminus \{0\}; \text{ es existiert } b \in \mathbb{Z}/n \text{ mit } ab = ba = 1\}.$$

Es ist dann leicht zu sehen, dass $(\mathbb{Z}/n)^\times$ eine Gruppe ist – das einzige Problem, die Existenz multiplikativer Inverser, haben wir sozusagen wegdefiniert. Man nennt $(\mathbb{Z}/n)^\times$ die *Einheitengruppe* von $\mathbb{Z}/n$. (Diese Konstruktion kann man analog in jedem Ring R (Ergänzung 8.7) durchführen; man erhält dann die Einheitengruppe des Rings $R^\times$.)

Wir können Satz 4.16 umformulieren als

SATZ 8.10. Sei $n > 1$. Die Restklasse in $\mathbb{Z}/n$ einer ganzen Zahl $a \in \mathbb{Z}$ liegt genau dann in $(\mathbb{Z}/n)^\times$, wenn $\text{ggT}(a, n) = 1$, d.h. wenn a und n teilerfremd sind.

□ Ergänzung 8.9

8.1.2. Gruppenhomomorphismen. Sind G, H Gruppen, so ist es meistens – ähnlich wie im Fall von Vektorräumen – sinnvoller, nicht beliebige Abbildungen von G nach H zu betrachten, sondern solche, die mit der Gruppenstruktur verträglich sind:

DEFINITION 8.11. Seien G, H Gruppen. Ein *Gruppenhomomorphismus* von G nach H ist eine Abbildung $f: G \rightarrow H$, so dass für alle $x, y \in G$ gilt:

$$f(xy) = f(x)f(y).$$

⊥

Wir haben hier beide Gruppen G und H multiplikativ geschrieben. Auf der linken Seite werden x und y in G multipliziert (»verknüpft«), auf der rechten Seite werden $f(x)$ und $f(y)$ in H multipliziert. Wird beispielsweise die Verknüpfung auf G mit $+$, die auf H mit $\circ$ bezeichnet, so ist die Bedingung umzuschreiben als $f(x + y) = f(x) \circ f(y)$. Die Menge aller Gruppenhomomorphismen von G nach H bezeichnen wir mit $\text{Hom}(G, H)$.

Ist $f: G \rightarrow H$ ein Gruppenhomomorphismus, so gilt $f(e_G) = e_H$, d.h. f bildet das neutrale Element von G auf das neutrale Element von H ab. Dies folgt aus der Rechnung

$$f(e_G) = f(e_G e_G) = f(e_G)f(e_G),$$

indem man beide Seiten mit $f(e_G)^{-1}$ multipliziert. Ferner gilt $f(g^{-1}) = f(g)^{-1}$, denn $f(g^{-1})f(g) = f(g^{-1}g) = f(e_G) = e_H$.

Sind $f: G_1 \rightarrow G_2$ und $g: G_2 \rightarrow G_3$ Gruppenhomomorphismen, so ist die Verkettung $g \circ f: G_1 \rightarrow G_3$ ebenfalls ein Gruppenhomomorphismus, wie man leicht nachprüft. Für jede Gruppe G ist die identische Abbildung $\text{id}: G \rightarrow G$ ein Gruppenhomomorphismus.

BEISPIEL 8.12. Die Abbildung $\mathbb{Z} \rightarrow \mathbb{Z}/n$, die eine ganze Zahl a auf ihre Restklasse in $\mathbb{Z}/n$ abbildet, ist ein Gruppenhomomorphismus. Das ist genau die Aussage $\pi(x+x') = \pi(x) + \pi(x')$ aus Lemma 4.13. $\diamond$

Ähnlich wie bei Vektorräumen definieren wir, aufbauend auf dem Begriff des Homomorphismus, den Begriff des Isomorphismus zwischen Gruppen:

DEFINITION 8.13. Seien G, H Gruppen. Ein *Gruppenisomorphismus* zwischen G und H ist ein Gruppenhomomorphismus $f: G \rightarrow H$, so dass ein Gruppenhomomorphismus $g: H \rightarrow G$ existiert mit $g \circ f = \text{id}_G, f \circ g = \text{id}_H$, d.h. f und g sind Umkehrabbildungen voneinander. $\dashv$

Wir sagen, zwei Gruppen G und H seien *isomorph*, wenn ein Isomorphismus zwischen G und H existiert. Wir schreiben dann $G \cong H$. Wie bei Vektorräumen schreibt man im Fall von Isomorphismen auch $G \xrightarrow{\sim} H$ statt $G \rightarrow H$.

Wir können dann die oben erwähnte »Identifikation« von $GL_1(K)$ mit $K^\times$ (für einen Körper K) vornehmer ausdrücken als: Die Abbildung $GL_1(K) \rightarrow K^\times, (a) \mapsto a$, ist ein Gruppenisomorphismus mit Umkehrhomomorphismus $a \mapsto (a)$. (Hier bezeichnet (a) die (1×1) -Matrix mit dem einzigen Eintrag a .)

LEMMA 8.14. Ein Gruppenhomomorphismus $f: G \rightarrow H$ ist genau dann ein Isomorphismus, wenn er bijektiv ist.

BEWEIS. Ist f ein Isomorphismus, so existiert eine Umkehrabbildung, also ist f bijektiv. Ist f bijektiv, so existiert eine Umkehrabbildung $g: H \rightarrow G$ zwischen den Mengen G und H (es gilt also $g \circ f = \text{id}_G, f \circ g = \text{id}_H$). Wir müssen zeigen, dass g ein Gruppenhomomorphismus ist.

Sind $x', y' \in H$ und $x := g(x'), y := g(y')$, also $x' = f(x), y' = f(y)$, dann gilt

$$g(x'y') = g(f(x)f(y)) = g(f(xy)) = xy = g(x')g(y').$$

□

BEISPIEL 8.15. (1) Sei K ein Körper, $n \in \mathbb{N}$. Dann ist die Abbildung

$$\text{Aut}(K^n) \longrightarrow GL_n(K), \quad f \mapsto M(f),$$

ein Gruppenisomorphismus mit Umkehrabbildung $A \mapsto f_A$. Das folgt aus Satz 7.26 und Lemma 7.27.

(2) Seien K ein Körper und V ein K -Vektorraum der Dimension $n \in \mathbb{N}$ mit Basis $\mathcal{B} = (v_1, \dots, v_n)$. Wir haben dann die Koordinatenabbildung $c_{\mathcal{B}}: V \xrightarrow{\sim} K^n$, ein Isomorphismus von K -Vektorräumen. Daraus erhalten wir einen Isomorphismus

$$\text{Aut}(V) \longrightarrow \text{Aut}(K^n), \quad f \mapsto c_{\mathcal{B}} \circ f \circ c_{\mathcal{B}}^{-1}.$$

Durch Verkettung mit dem Isomorphismus aus Teil (1) erhalten wir den Gruppenisomorphismus

$$\text{Aut}(V) \longrightarrow GL_n(K), \quad f \mapsto M_{\mathcal{B}}^{\mathcal{B}}(f).$$

□

BEISPIEL 8.16. (1) Sei G eine Gruppe und $g \in G$. Die Abbildung $c_g: G \rightarrow G, x \mapsto gxg^{-1}$ ist ein Gruppenhomomorphismus (denn $gxyg^{-1} = gxg^{-1}gyg^{-1}$) und sogar ein Isomorphismus: Die Umkehrabbildung ist gegeben durch $x \mapsto g^{-1}xg$. Man nennt die Abbildung c_g die *Konjugation* mit dem Element g .

- (2) Ist speziell $G = GL_n(K)$ für einen Körper K und $n \in \mathbb{N}$, so bezeichnen wir für $A \in GL_n(K)$ die Abbildung $GL_n(K) \rightarrow GL_n(K)$, $M \mapsto AMA^{-1}$, als die Konjugation mit A .

Etwas allgemeiner sagt man, Matrizen $M, M' \in M_n(K)$ seien *zueinander konjugiert*, wenn $A \in GL_n(K)$ existiert mit $M' = AMA^{-1}$. Eine andere häufig verwandte Bezeichnung für diesen Sachverhalt ist, dass M und M' *ähnlich* seien.

Wir haben diese Situation schon bei der Diskussion des Basiswechsels gesehen. In der Tat folgt aus Korollar 7.33 die folgende Aussage: Seien $M, M' \in M_n(K)$. Dann sind äquivalent:

- (i) Die Matrizen M und M' sind zueinander konjugiert.
- (ii) Es existieren Basen $\mathcal{B}, \mathcal{B}'$ von K^n , so dass

$$M' = M_{\mathcal{B}'}^{\mathcal{B}} M M_{\mathcal{B}}^{\mathcal{B}'}$$

- (iii) Es existieren eine lineare Abbildung $f: K^n \rightarrow K^n$ und Basen $\mathcal{B}, \mathcal{B}'$ von K^n , so dass

$$M = M_{\mathcal{B}}^{\mathcal{B}}(f), \quad M' = M_{\mathcal{B}'}^{\mathcal{B}'}(f).$$

Die Frage, wie man entscheidet, ob Matrizen $M, M' \in M_n(K)$ zueinander konjugiert sind, wird uns in der Linearen Algebra 2 für einige Zeit beschäftigen.

◇

8.1.3. Untergruppen. Ähnlich wie wir für Vektorräume den Begriff des Untervektorraums definiert haben, können wir auch von Untergruppen sprechen.

DEFINITION 8.17. Sei G eine Gruppe (multiplikativ geschrieben). Eine *Untergruppe* von G ist eine Teilmenge $H \subseteq G$ mit den folgenden Eigenschaften:

- (a) $1 \in H$,
- (b) für alle $g, g' \in H$ gilt $gg' \in H$,
- (c) für alle $g \in H$ gilt $g^{-1} \in H$,

⊢

Ist $H \subseteq G$ eine Untergruppe, so ist H mit der Einschränkung der Verknüpfung von G selbst eine Gruppe.

LEMMA 8.18. Sei G eine Gruppe. Seien $H, H' \subseteq G$ Untergruppen. Dann ist $H \cap H'$ eine Untergruppe.

Allgemeiner gilt: Der Durchschnitt einer beliebigen Familie von Untergruppen einer Gruppe G ist eine Untergruppe.

BEWEIS. Wir beweisen direkt die allgemeine Form. Seien $H_i \subseteq G$, $i \in I$, Untergruppen von G , und sei $H := \bigcap_{i \in I} H_i$.

Das neutrale Element von G liegt in allen H_i , also auch in H . Sind $x, y \in H$, so gilt $x, y \in H_i$ und demnach $xy \in H_i$ für alle i , da H_i nach Voraussetzung eine Untergruppe ist. Es folgt $xy \in H$. Ähnlich sieht man, dass für jedes $x \in H$ das inverse Element $x^{-1} \in G$ in H liegen muss. □

DEFINITION 8.19. Sei G eine Gruppe, und sei $M \subseteq G$ eine Teilmenge von G . Wir bezeichnen mit $\langle M \rangle$ den Durchschnitt aller Untergruppen von G , die M enthalten. Dies ist eine Untergruppe von G (nach Lemma 8.18) und ist die kleinste Untergruppe von G , die M enthält. Wir nennen $\langle M \rangle$ die *von M erzeugte Untergruppe* von G . ⊢

Für $M = \{g_1, \dots, g_n\}$ schreibt man auch $\langle g_1, \dots, g_n \rangle$ statt $\langle \{g_1, \dots, g_n\} \rangle$.

BEISPIEL 8.20. (1) Wir betrachten $\mathbb{Z}$ als Gruppe bezüglich der Addition. Ist $a \in \mathbb{Z}$, so ist $\langle a \rangle$, die von a erzeugte Untergruppe, die Menge aller Vielfachen von a :

$$\langle a \rangle = \{ka; k \in \mathbb{Z}\}.$$

Denn die rechte Seite ist eine Untergruppe von $\mathbb{Z}$, wie man leicht nachprüft. Offensichtlich enthält sie a und damit auch $\langle a \rangle$. Andererseits liegen mit a auch $-a$ und alle Summen der Form $a + \dots + a$ und $(-a) + \dots + (-a)$ in $\langle a \rangle$.

(2) Wir betrachten nochmals $\mathbb{Z}$ als Gruppe bezüglich der Addition und wählen nun zwei ganze Zahlen $a, b \in \mathbb{Z}$. Dann ist

$$\langle a, b \rangle = \{k \operatorname{ggT}(a, b); k \in \mathbb{Z}\} = \langle \operatorname{ggT}(a, b) \rangle$$

die Menge aller Vielfachen des größten gemeinsamen Teilers von a und b . Jedenfalls ist klar, dass sowohl a als auch b Vielfache von $\operatorname{ggT}(a, b)$ sind, also gilt $\subseteq$. Für die umgekehrte Inklusion kann man Lemma 3.53 benutzen (das hatten wir nur innerhalb einer Ergänzung behandelt, aber inzwischen sind Sie soweit, dass Sie den Beweis in ein paar Minuten nachvollziehen können).

(3) Wir können Teil (1) folgendermaßen verallgemeinern: Sei G irgendeine Gruppe (die wir nun multiplikativ schreiben wollen), und sei $g \in G$. Wir schreiben g^n für das n -fache Produkt von g mit sich selbst (für $n \in \mathbb{N}$, wie üblich sei g^0 das neutrale Element von G), und $g^{-n} = (g^{-1})^n$. Dann gilt

$$\langle g \rangle = \{g^k; k \in \mathbb{Z}\}.$$

Zur Begründung muss man sich, ähnlich wie in Teil (1), eigentlich nur überlegen, dass die rechte Seite eine Untergruppe von G ist. Daraus folgt die Gleichheit dann leicht.

Ist G eine Gruppe, die ein Element g enthält mit $\langle g \rangle = G$, so sagt man, G sei von *einem* Element erzeugt oder G sei eine *zyklische Gruppe*. Es ist klar, dass jede zyklische Gruppe kommutativ ist. Beispiele sind $\mathbb{Z} = \langle 1 \rangle$ und $\mathbb{Z}/n = \langle \bar{1} \rangle$ (jeweils mit der Addition als Verknüpfung). Es ist nicht sehr schwer zu zeigen, dass dies die einzigen zyklischen Gruppen sind in dem Sinne, dass jede zyklische Gruppe G zu einer der genannten Gruppen isomorph ist.

(4) Die Untergruppe von $GL_n(K)$, die erzeugt wird von allen Matrizen der Form $E_{ij}(a)$, P_{ij} und $\operatorname{diag}(1, \dots, 1, a, 1, \dots, 1)$, $a \in K^\times$ (vergleiche Bemerkung 5.37), ist $GL_n(K)$. Denn jede invertierbare Matrix hat als reduzierte Zeilenstufenform die Einheitsmatrix, und das bedeutet gerade, dass sie sich als Produkt von Matrizen der oben genannten Typen schreiben lässt.

◇

ERGÄNZUNG 8.21 (Die von einer Teilmenge erzeugte Gruppe). Als wir den von einer Teilmenge M eines Vektorraums V erzeugten Untervektorraum $\langle M \rangle$ eingeführt haben, haben wir gezeigt, dass $\langle M \rangle$ gleich dem Durchschnitt aller Untervektorräume von V ist, die M enthalten. Im Vektorraumfall ist es aber auch einfach, die Elemente von $\langle M \rangle$ (einigermaßen) explizit anzugeben: es sind alle Linearkombinationen von Elementen aus M . Das war unsere Definition von $\langle M \rangle$. Aber wir hätten natürlich auch die Charakterisierung als Durchschnitt zur Definition machen können.

Im Fall von Gruppen ist es etwas komplizierter und auch weniger nützlich, eine explizite Beschreibung der Elemente einer Untergruppe der Form $\langle M \rangle$ zu geben. Deswegen haben wir oben die von M erzeugte Untergruppe als den Durchschnitt aller Untergruppen, die M enthalten, definiert.

Eine »konkrete« Beschreibung ist die folgende: Sei G eine (multiplikativ geschriebene) Gruppe, $M \subseteq G$. Da $\langle M \rangle$ eine Untergruppe ist, sind neben den Elementen von M auch deren Inverse und alle Produkte von Elementen dieser Form in $\langle M \rangle$. Es gilt also

$$\left\{ \prod_{i=1}^N m_i^{\varepsilon_i}; N \in \mathbb{N}, m_1, \dots, m_N \in M, \varepsilon_1, \dots, \varepsilon_N \in \{1, -1\} \right\} \subseteq \langle M \rangle.$$

Andererseits ist die linke Seite eine Untergruppe von G , die M enthält. Deshalb stimmen diese beiden Mengen überein. Man beachte, dass man die Anzahl N der Faktoren in dem Produkt im allgemeinen nicht beschränken kann. Selbst wenn M endlich ist, muss man (wenn immer $\langle M \rangle$ unendlich ist) beliebig große N zulassen, um alle Elemente in der obigen Form darzustellen.

Je nachdem, in welcher Form die Gruppe G gegeben ist und was über ihre Struktur bekannt ist, kann es sehr schwierig sein zu entscheiden, ob Elemente in G , die als solche Produkte gegeben sind, gleich sind. □ Ergänzung 8.21

8.1.4. Kern und Bild von Gruppenhomomorphismen. Auch für Gruppenhomomorphismen definiert man Kern und Bild. Die Definitionen sind analog zu den Definitionen, die wir vom Vektorraumkontext schon kennen.

DEFINITION 8.22. Sei $f: G \rightarrow H$ ein Gruppenhomomorphismus.

(1) Der *Kern* von f ist die Untergruppe

$$\text{Ker}(f) := \{g \in G; f(g) = 1\}$$

von G . (Hier bezeichnet $1 \in H$ das neutrale Element.)

(2) Das *Bild* von f ist das Bild $\text{Im}(f) = f(G)$ von f im Sinne des Bildes einer Abbildung zwischen Mengen. Dies ist eine Untergruppe von H .

⊢

Wir überlassen es der Leser*in als Übung nachzuprüfen, dass Kern und Bild eines Gruppenhomomorphismus tatsächlich Untergruppen sind.

ERGÄNZUNG 8.23 (Normalteiler). Im allgemeinen ist nicht jede Untergruppe von G der Kern eines Gruppenhomomorphismus in irgendeine andere Gruppe. Hier verhalten sich Gruppen also anders als Vektorräume über einem Körper. Dieses Phänomen wird uns später noch beschäftigen (jedenfalls in der Algebra-Vorlesung). Man kann zeigen, dass jede Untergruppe einer kommutativen Gruppe auch der Kern irgendeines Gruppenhomomorphismus ist. Im allgemeinen Fall sind die Untergruppen, die auch als Kern auftreten, gerade die sogenannten *Normalteiler*. □ Ergänzung 8.23

LEMMA 8.24. Sei $f: G \rightarrow H$ ein Gruppenhomomorphismus. Die Abbildung f ist genau dann injektiv, wenn $\text{Ker}(f) = \{1\}$ ist. (Man sagt, der Kern von f sei trivial.)

BEWEIS. Der Beweis ist einfach und verläuft genau wie im Vektorraumfall. Wenn f injektiv ist, kann höchstens ein Element auf $1 \in H$ abgebildet werden, und aus $f(1) = 1$ folgt $\text{Ker}(f) = \{1\}$. Ist andererseits der Kern von f trivial und sind $g, g' \in G$ mit $f(g) = f(g')$, so folgt

$$f(g^{-1}g') = f(g)^{-1}f(g') = 1,$$

also $g^{-1}g' = 1$, das heißt $g = g'$. □

8.1.5. Permutationsmatrizen.

Sei K ein Körper. Wir definieren

DEFINITION 8.25. Eine Matrix $P \in M_n(K)$ heißt *Permutationsmatrix*, wenn in jeder Zeile und jeder Spalte genau ein Eintrag von Null verschieden ist, und alle diese nicht-verschwindenden Einträge gleich 1 sind. $\dashv$

BEISPIEL 8.26. (1) Die Einheitsmatrix E_n ist eine Permutationsmatrix.

(2) Die Permutationsmatrizen der Größe 3 sind

$$\begin{pmatrix} 1 & & \\ & 1 & \\ & & 1 \end{pmatrix}, \begin{pmatrix} & 1 & \\ 1 & & \\ & & 1 \end{pmatrix}, \begin{pmatrix} 1 & & \\ & & 1 \\ & 1 & \end{pmatrix}, \begin{pmatrix} & & 1 \\ 1 & 1 & \\ & 1 & \end{pmatrix}, \begin{pmatrix} 1 & & \\ & & \\ 1 & 1 & \end{pmatrix}, \begin{pmatrix} & 1 & \\ & & 1 \\ 1 & & \end{pmatrix},$$

wobei wir die Nullen der Übersichtlichkeit halber weggelassen haben.

$\diamond$

Wir bezeichnen mit W_n die Menge aller Permutationsmatrizen. Eine äquivalente Definition ist die folgende: Eine Matrix $P \in M_n(K)$ ist genau dann eine Permutationsmatrix, wenn $f_P(\{e_1, \dots, e_n\}) = \{e_1, \dots, e_n\}$, das heißt, f_P bildet jeden Standardbasisvektor e_j auf einen der Standardbasisvektoren ab, und jeder Standardbasisvektor tritt dabei als Bild auf. Die Standardbasisvektoren werden also untereinander vertauscht oder *permutiert* – daher der Name Permutationsmatrix. Insbesondere bildet f_P eine Basis auf eine Basis ab, und ist daher ein Isomorphismus. (Konkret gilt $P^{-1} = P^t$.) Es gilt also $W_n \subseteq GL_n(K)$ und darüberhinaus:

LEMMA 8.27. Die Teilmenge $W_n \subseteq GL_n(K)$ der Permutationsmatrizen ist eine Untergruppe.

BEWEIS. Es ist klar, dass die Eigenschaft, dass f_P eine Bijektion der Menge $\{e_1, \dots, e_n\}$ mit sich selbst induziert, sich auf Verkettungen und die Umkehrabbildung überträgt. $\square$

Ist $\sigma: \{1, \dots, n\} \rightarrow \{1, \dots, n\}$ eine Bijektion, also ein Element der symmetrischen Gruppe $S_n = \text{Bij}(\{1, \dots, n\})$, so definieren wir einen Isomorphismus $f_\sigma: K^n \rightarrow K^n$ durch

$$f(e_i) = e_{\sigma(i)}, \quad i = 1, \dots, n.$$

Dann ist die Abbildung $S_n \rightarrow GL_n(K)$, $\sigma \mapsto M(f_\sigma)$, ein injektiver Gruppenhomomorphismus mit Bild W_n , induziert also einen Isomorphismus $S_n \cong W_n$. Zum Beweis, dass dies ein Homomorphismus ist, muss man nur überprüfen, dass für $\sigma, \tau \in S_n$ gilt, dass $f_{\sigma\tau} = f_\sigma \circ f_\tau$, und das ist einfach (denn es genügt, das auf den Vektoren der Standardbasis nachzurechnen). Dass die Abbildung Bild in W_n hat und eine Bijektion $S_n \rightarrow W_n$ induziert, ist nicht schwer zu sehen. Wir schreiben auch $P_\sigma := M(f_\sigma)$. Die j -te Spalte von P_σ ist also $e_{\sigma(j)}$.

8.1.6. Symmetriegruppen. Seien K ein Körper und V ein K -Vektorraum. Für eine Teilmenge $M \subseteq V$ nennen wir

$$\text{Stab}(M) = \{f \in \text{Aut}(V); f(M) = M\}$$

die *Symmetriegruppe* von M (man spricht auch vom *Stabilisator* von M in $\text{Aut}(V)$). Dies ist, wie der Name andeutet, eine Untergruppe von $\text{Aut}(V)$: Offenbar gilt $\text{id}_V \in \text{Stab}(M)$, für $f, g \in \text{Stab}(M)$ gilt $f(g(M)) = f(M) = M$, also $f \circ g \in \text{Stab}(M)$, und aus $f(M) = M$ folgt durch Anwenden von f^{-1} auf beiden Seiten, dass $M = f^{-1}(M)$, also $f^{-1} \in \text{Stab}(M)$.

Um einige konkrete Beispiele angeben zu können, betrachten wir den Fall $K = \mathbb{R}$ und $V = \mathbb{R}^n$, $n \in \mathbb{N}$. Wir können dann $\text{Aut}(V) = \text{Aut}(\mathbb{R}^n)$ mit der Gruppe $GL_n(\mathbb{R})$ der invertierbaren Matrizen der Größe n identifizieren, wie wir in Beispiel 8.15 gesehen haben.

BEISPIEL 8.28 (Symmetriegruppe eines Quadrats in $\mathbb{R}^2$). Die folgenden linearen Abbildungen bilden das Quadrat $Q \subset \mathbb{R}^2$ mit Eckpunkten $(1, 1)^t, (1, -1)^t, (-1, 1)^t, (-1, -1)^t$ auf sich selbst ab:

Beschreibung	Darstellende Matrix
Identität	$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$
Drehung um 90° gegen den Uhrzeigersinn	$\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$
Drehung um 180° gegen den Uhrzeigersinn	$\begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix}$
Drehung um 270° gegen den Uhrzeigersinn	$\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$
Spiegelung an der waagerechten Koordinatenachse	$\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$
Spiegelung an der Diagonalen $\langle (1, 1)^t \rangle$	$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$
Spiegelung an der senkrechten Koordinatenachse	$\begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$
Spiegelung an der Diagonalen $\langle (1, -1)^t \rangle$	$\begin{pmatrix} 0 & -1 \\ -1 & 0 \end{pmatrix}$

Es ist leicht zu sehen, dass es keine weiteren Automorphismen f von $\mathbb{R}^2$ mit $f(Q) = Q$ gibt: Jede lineare Abbildung, die das Quadrat fixiert, muss auch jeden Eckpunkt des Quadrates auf einen Eckpunkt abbilden. Weil $(1, 1)^t$ und $(1, -1)^t$ linear unabhängig sind, ist eine lineare Abbildung durch ihre Werte auf diesen beiden Punkten bereits eindeutig bestimmt. Weil benachbarte Eckpunkte auf benachbarte Eckpunkte abgebildet werden müssen, wie man sich leicht überlegt, folgt, dass es höchstens 8 Möglichkeiten gibt.

Wenn wir die Ecken durchnummerieren, etwa $E_1 = (1, 1)^t, E_2 = (1, -1)^t, E_3 = (-1, 1)^t, E_4 = (-1, -1)^t$, liefert jede der obigen Abbildungen eine Bijektion $\{E_1, E_2, E_3, E_4\} \rightarrow \{E_1, E_2, E_3, E_4\}$, und ist durch diese Bijektion eindeutig bestimmt. Wir erhalten so einen injektiven Gruppenhomomorphismus

$$\text{Stab}(Q) \rightarrow \text{Bij}(\{E_1, E_2, E_3, E_4\}).$$

Die Gruppe $\text{Stab}(Q)$ nennt man auch die *Diedergruppe mit 8 Elementen*. Wir bezeichnen sie mit D_8 ; in einigen Quellen wird stattdessen die Bezeichnung D_4 benutzt. (Aussprache: Diedergruppe, das erste e wird getrennt vom i ausgesprochen und trägt auch die Betonung.)

Die Gruppe $\text{Stab}(Q)$ ist nicht abelsch. (Können Sie ein Beispiel von zwei Elementen finden, für die das Kommutativgesetz nicht gilt?) Man kann leicht nachrechnen, oder durch eine geometrische Überlegung begründen, dass die Gruppe $\text{Stab}(Q)$ erzeugt wird von $\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$

und $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ (es gibt natürlich auch noch andere Möglichkeiten, die Gruppe durch zwei Elemente zu erzeugen). Insbesondere lässt sich $\text{Stab}(Q)$ durch Spiegelungen erzeugen.

Die vier Drehungen (um $0^\circ, 90^\circ, 180^\circ, 270^\circ$) bilden eine Untergruppe von $\text{Stab}(Q)$. $\diamond$

ERGÄNZUNG 8.29 (Symmetriegruppe des regelmäßigen n -Eckes in $\mathbb{R}^2$). Das vorhergehende Beispiel können wir verallgemeinern und statt des Quadrats ein regelmäßiges n -Eck in $\mathbb{R}^2$

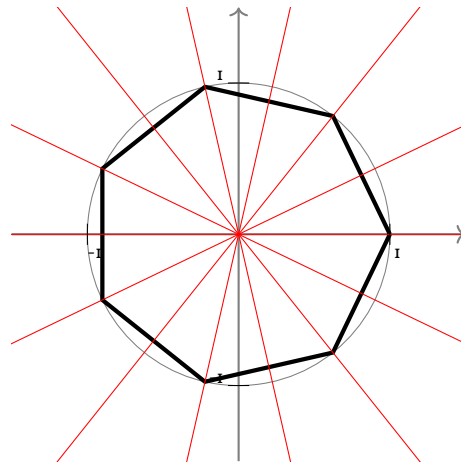


ABBILDUNG 1. Ein regelmäßiges Siebeneck, dessen Ecken auf dem Einheitskreis liegen. Die Symmetrieachsen sind in rot eingezeichnet. Die Symmetriegruppe besteht aus den Spiegelungen an diesen Geraden und den sieben Drehungen um $i \cdot 360 / 7$ Grad, $i = 0, \dots, 6$.

betrachten, dessen Zentrum der Ursprung $(0, 0)^t$ ist. (Hier sei $n \geq 3$ eine natürliche Zahl. Unter einem regelmäßigen n -Eck verstehen wir einen geschlossenen Streckenzug, der aus n Strecken (und folglich n Eckpunkten, an denen je zwei Strecken zusammenkommen) besteht, wobei alle Strecken gleich lang und alle Innenwinkel gleich groß sind.)

Die Symmetriegruppe des regelmäßigen n -Eckes hat $2n$ Elemente, und zwar n Drehungen (um den $\frac{i}{n}$ -ten Teil des Vollwinkels, $i = 0, \dots, n-1$) und n Spiegelungen (an den Ursprungsgeraden durch die Eckpunkte und an den Ursprungsgeraden durch die Seitenmittelpunkte; je nachdem, ob n ungerade ist, enthält jede dieser Geraden einen Eckpunkt und einen Seitenmittelpunkt, oder zwei Punkte desselben Typs). Sie lässt sich durch zwei (geeignet gewählte) Spiegelungen erzeugen. Man nennt diese Gruppe auch die *Diedergruppe mit $2n$ Elementen*; wir bezeichnen sie mit D_{2n} . In einigen Quellen wird sie stattdessen mit D_n bezeichnet. □ Ergän-

zung 8.29

ERGÄNZUNG 8.30 (Symmetriegruppe des Würfels in $\mathbb{R}^3$). Natürlich ist es auch interessant, die Symmetriegruppen von Teilmengen von $\mathbb{R}^3$ zu studieren. Betrachten wir zum Beispiel den Würfel mit Eckpunkten $(1, 1, 1)^t, (1, 1, -1)^t, (1, -1, 1)^t, (-1, 1, 1)^t, (-1, -1, 1)^t, (-1, 1, -1)^t, (1, -1, -1)^t, (-1, -1, -1)^t$.

Die Symmetriegruppe hat 48 Elemente, darunter 24 Drehungen und 24 Spiegelungen (vergleiche Ergänzung 7.59).

Die Drehungen bilden eine Untergruppe der Symmetriegruppe des Würfels. Jede Drehung induziert eine Permutation der Raumdiagonalen des Würfels (als Raumdiagonale bezeichnet man eine Gerade durch zwei Eckpunkte des Würfels, die auch durch den Ursprung geht). Man kann zeigen, dass jede Permutation der Raumdiagonalen durch eine Drehung realisiert werden kann, und dass eine Drehung durch die entsprechende Permutation der Raumdiagonalen eindeutig bestimmt ist. Daraus folgt, dass die Untergruppe der Drehungen des Würfels isomorph ist zur symmetrischen Gruppe S_4 . □ Ergänzung 8.30

8.2. Die spezielle lineare Gruppe

Wir haben in Bemerkung 5.37 die Matrizen $E_{ij}(a) \in GL_n(K)$ definiert (für $1 \leq i \neq j \leq n$, $a \in K$). Wir wiederholen hier die Definition:

DEFINITION 8.31. Seien K ein Körper, $n \in \mathbb{N}$, $1 \leq i \neq j \leq n$, $a \in K$. Die Matrix $E_{ij}(a) = (a_{\lambda\mu})_{\lambda,\mu} \in GL_n(K)$ ist definiert durch

$$a_{\lambda\mu} = \begin{cases} 1 & \lambda = \mu, \\ a & \lambda = i, \mu = j, \\ 0 & \text{sonst.} \end{cases}$$

Die Matrizen der Form $E_{ij}(a)$ nennen wir *Elementarmatrizen*. ←

Der Begriff *Elementarmatrix* wird teilweise auch mit einer etwas anderen Bedeutung benutzt; bei manchen Autoren heißt jede Matrix, die aus der Einheitsmatrix durch eine einzige elementare Zeilenumformung hervorgeht, eine Elementarmatrix. Bei uns sind es nur diejenigen, die durch eine einzige elementare Zeilenumformung vom Typ I aus der Einheitsmatrix hervorgehen.

Die Elementarmatrizen sind tatsächlich invertierbar; es gilt $E_{ij}(a)^{-1} = E_{ij}(-a)$, wie man leicht nachrechnet. Genauer gilt $E_{ij}(a)E_{ij}(b) = E_{ij}(a+b)$. Das können wir auch formulieren als: $\varphi_{ij}: K \rightarrow GL_n(K)$, $a \mapsto E_{ij}(a)$ ist ein Gruppenhomomorphismus von der additiven Gruppe K in die allgemeine lineare Gruppe.

Wir hatten in Abschnitt 5.3 besprochen, dass für jede Matrix $A \in M_n(K)$ (und sogar allgemeiner für jede Matrix mit n Zeilen) $E_{ij}(a)A$ die Matrix ist, die aus A durch Addition des a -fachen der j -ten Zeile zur i -ten Zeile hervorgeht. Analog ist $AE_{ij}(a)$ die Matrix, die aus A durch Addition der i -ten Spalte zur j -ten Spalte hervorgeht.



If people do not believe that mathematics is simple, it is only because they do not realize how complicated life is.

John von Neumann.

Fundort: <https://mathoverflow.net/a/7306>

DEFINITION 8.32 (Spezielle lineare Gruppe). Sei K ein Körper. Die spezielle lineare Gruppe $SL_n(K)$ ist die von allen Elementarmatrizen $E_{ij}(a)$ erzeugte Untergruppe von $GL_n(K)$. ←

Natürlich liegen alle Elementarmatrizen $E_{ij}(a)$ selbst in $SL_n(K)$, und damit auch alle Produkte von solchen, und damit kann man eine ganze Menge Matrizen »finden«. Weil das Inverse einer Elementarmatrix wieder eine Elementarmatrix ist, besteht die Gruppe $SL_n(K)$ genau aus allen Produkten (mit endlich vielen Faktoren) von Elementarmatrizen.

Weil wir Multiplikation mit Elementarmatrizen auch als elementare Zeilenumformungen vom Typ I deuten können, können wir die Zugehörigkeit zur Gruppe $SL_n(K)$ auch so beschreiben: Eine invertierbare Matrix A der Größe n liegt genau dann in $SL_n(K)$, wenn sich die Einheitsmatrix durch wiederholte Anwendung von elementaren Zeilenumformungen vom Typ I in die Matrix A umformen lässt. Oder äquivalent: Wenn sich A durch wiederholte Anwendung von elementaren Zeilenumformungen vom Typ I in die Einheitsmatrix überführen lässt.

Wir wollen als nächstes einige Beispiele von Elementen der speziellen linearen Gruppe sammeln.

BEISPIEL 8.33. Sei K ein Körper.

- (1) Es gilt $\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \in SL_2(K)$. Äquivalent: Wir können die Einheitsmatrix E_2 durch elementare Zeilenumformungen vom Typ I auf die genannte Matrix bringen. In der Tat:

$$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \xrightarrow{Z_2 \rightsquigarrow Z_2 - Z_1} \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow Z_1 + Z_2} \begin{pmatrix} 0 & 1 \\ -1 & 1 \end{pmatrix} \xrightarrow{Z_2 \rightsquigarrow Z_2 - Z_1} \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

Wir können natürlich diese Umformungen als Matrizenmultiplikation umschreiben und erhalten damit die Darstellung

$$\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix}.$$

- (2) Seien $n \in \mathbb{N}$ und $a_1, \dots, a_n \in K$ mit $\prod_{i=1}^n a_i = 1$. Dann ist $\text{diag}(a_i) \in SL_n(K)$.

Der entscheidende Fall ist der Fall $n = 2$, den wir als erstes erklären werden. Die Behauptung ist dazu äquivalent, dass für jedes $a \in K^\times$ die Einheitsmatrix E_2 durch elementare

Zeilenumformungen vom Typ I auf die Matrix $\begin{pmatrix} a & 0 \\ 0 & a^{-1} \end{pmatrix}$ gebracht werden kann. Wir rechnen dazu

$$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \xrightarrow{Z_2 \rightsquigarrow Z_2 + aZ_1} \begin{pmatrix} 1 & 0 \\ a & 1 \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow Z_1 - \frac{1-a}{a}Z_2} \begin{pmatrix} 1 - (1-a) & -\frac{1-a}{a} \\ a & 1 \end{pmatrix} = \begin{pmatrix} a & -\frac{1-a}{a} \\ a & 1 \end{pmatrix} \xrightarrow{Z_2 \rightsquigarrow Z_2 - Z_1} \begin{pmatrix} a & -\frac{1-a}{a} \\ 0 & 1 + \frac{1-a}{a} \end{pmatrix} = \begin{pmatrix} a & -\frac{1-a}{a} \\ 0 & a^{-1} \end{pmatrix} \xrightarrow{Z_1 \rightsquigarrow Z_1 + (1-a)Z_2} \begin{pmatrix} a & 0 \\ 0 & a^{-1} \end{pmatrix}.$$

Für den allgemeinen Fall ergibt sich, indem wir die entsprechenden Zeilenumformungen auf die Zeilen $i, i+1$ anwenden, dass alle Matrizen der Form $\text{diag}(1, \dots, 1, a, a^{-1}, 1, \dots, 1)$ für $a \in K$ in $SL_n(K)$ liegen. Für Elemente $a_i \in K^\times$ mit $\prod_{i=1}^n a_i = 1$ gilt $(a_1 \cdots a_{n-1})^{-1} = a_n$ und wir erhalten, dass

$$\text{diag}(a_1, \dots, a_n) =$$

$$\text{diag}(a_1, a_1^{-1}, 1, \dots, 1) \text{diag}(1, a_1 a_2, (a_1 a_2)^{-1}, 1, \dots, 1) \cdots \text{diag}(1, \dots, 1, a_1 \cdots a_{n-1}, (a_1 \cdots a_{n-1})^{-1})$$

in $SL_n(K)$ liegt.

- (3) Ist $B \in SL_n(K)$, so gilt auch $B^t \in SL_n(K)$. In der Tat überprüft man leicht, dass für jede Untergruppe $H \subseteq GL_n(K)$ auch die Teilmenge $H^t := \{h^t; h \in H\}$ eine Untergruppe von $GL_n(K)$ ist. Weil $E_{ij}(a)^t = E_{ji}(a)$ gilt, liegen alle $E_{ij}(a)$ in $SL_n(K)^t = \{g^t; g \in SL_n(K)\}$, es folgt also $SL_n(K) \subseteq SL_n(K)^t$. Weil $(SL_n(K)^t)^t = SL_n(K)$ gilt, folgt daraus durch Anwenden von $-^t$ auch die umgekehrte Inklusion.

◇

Wir hatten in Beispiel 8.20 bemerkt, dass die Gruppe $GL_n(K)$ von den Matrizen der Form $E_{ij}(a)$, P_{ij} und $\text{diag}(1, \dots, 1, a, 1, \dots, 1)$, $a \in K^\times$ erzeugt wird. Insofern »ist der Unterschied zwischen $SL_n(K)$ und $GL_n(K)$ nicht sehr groß«. Das zeigt auch der folgende Satz. Wir werden die Situation in Kapitel 9 noch genauer verstehen und unter anderem zeigen, dass für jeden Körper, der mehr als 2 Elemente hat, und alle $n \geq 1$ die Untergruppe $SL_n(K)$ eine echte Untergruppe von $GL_n(K)$ ist, also dass es invertierbare Matrizen gibt, die nicht in $SL_n(K)$ liegen. Interessanterweise ist es mit den Methoden, die uns momentan zur Verfügung stehen, aber nicht so leicht zu zeigen, dass irgendeine konkrete Matrix *nicht* in $SL_n(K)$ liegt. (Eine Ausnahme ist der Fall $n = 1$; dann gibt es gar keine Matrizen der Form $E_{ij}(a)$, so dass $SL_1(K) = \{1\}$ die triviale Gruppe ist.)

SATZ 8.34. Seien K ein Körper und $n \in \mathbb{N}$. Sei $A \in GL_n(K)$. Dann existieren $d, d' \in K$ und $B, C \in SL_n(K)$, so dass

$$A = B \operatorname{diag}(1, \dots, 1, d), \quad A = \operatorname{diag}(1, \dots, 1, d')C.$$

BEWEIS. Wir zeigen zuerst die Existenz von B , bzw. äquivalent die Existenz einer Matrix $B' \in SL_n(K)$, so dass $B'A$ eine Diagonalmatrix ist, deren erste $n - 1$ Einträge alle gleich 1 sind. Haben wir B' gefunden, so setzen wir $B = (B')^{-1}$.

In Beispiel 8.33 (1) haben wir gesehen, dass $\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$ in der Gruppe $SL_2(K)$ liegt. Eine Matrix $A \in M_2(K)$ von links mit dieser Matrix zu multiplizieren, bewirkt *bis auf Vorzeichen* eine Zeilenvertauschung. Es ist klar, dass wir das Beispiel auf die Gruppe $SL_n(K)$ übertragen können. Wir können daher durch Multiplikation mit einer geeigneten Matrix in $SL_n(K)$ von links, die Matrix A bis auf Vorzeichen durch Zeilenvertauschungen abändern.

Damit können wir – mit kleinen Änderungen – den Gauß-Algorithmus in dem Sinne durchführen, dass wir A auf eine Diagonalmatrix bringen können. In der Tat kann A als invertierbare Matrix keine Nullspalte haben. Im ersten Schritt ändern wir A durch Multiplikation mit einer Matrix in $SL_n(K)$ so ab, dass der Eintrag in der linken oberen Ecke $\neq 0$ ist. Wir können zwar nicht mit einer Zeilenumformung vom Typ III diesen Eintrag auf 1 bringen, weil wir nur mit Matrizen aus $SL_n(K)$ multiplizieren dürfen, aber dennoch können wir geeignete Vielfache der ersten Zeile von den anderen Zeilen abziehen, um zu erreichen, dass alle anderen Einträge der ersten Spalte gleich 0 sind.

Wir erhalten so eine Matrix der Form $\begin{pmatrix} a & * \\ 0 & A' \end{pmatrix}$, als Blockmatrix geschrieben, d.h. $A' \in M_{n-1}(K)$. Weil A invertierbar ist, ist auch A' invertierbar. Eine Möglichkeit, das zu begründen, ist, dass A als reduzierte Zeilenstufenform die Einheitsmatrix E_n hat; wäre A' nicht invertierbar, so wäre das nicht der Fall. Alternativ könnte man sagen, dass A Rang n hat und dass deshalb A' Rang $n - 1$ haben muss, also invertierbar ist.

Wir fahren dann induktiv mit den anderen Spalten so fort, dass wir A auf die Form einer oberen Dreiecksmatrix bringen, und zwar – wieder wegen der Invertierbarkeit – mit Einträgen $\neq 0$ auf der Diagonale. Als nächstes bringen wir alle Einträge oberhalb der Diagonalen durch Zeilenumformungen vom Typ I auf 0. Wir erhalten so eine Diagonalmatrix $\operatorname{diag}(c_1, \dots, c_n)$.

Zum Schluss multiplizieren wir von links mit der Diagonalmatrix

$$\operatorname{diag}(c_1, \dots, c_{n-1}, (c_1 \cdots c_{n-1})^{-1})^{-1}$$

(siehe Beispiel 8.33 (2)) und erhalten die gewünschte Form.

Um die zweite Version des Satzes zu zeigen, schreiben wir A^t in der Form $B \operatorname{diag}(1, \dots, 1, d)$, so dass $A = (B \operatorname{diag}(1, \dots, 1, d))^t = \operatorname{diag}(1, \dots, 1, d) B^t$ ist. Weil mit B auch B^t in $SL_n(K)$ liegt (Beispiel 8.33 (3)), sind wir damit fertig. Alternativ kann man für Teil (2) das Argument aus Teil (1) mit Spalten- statt Zeilenumformungen wiederholen. $\square$

Wir werden in Kapitel 9 – wenn wir die »Determinante« $\det(A)$ von A eingeführt haben – sehen, dass d und d' im obigen Satz stets gleich sind, und auch nicht von der Wahl von B und C abhängen, nämlich $d = d' = \det(A)$. Daraus folgt dann auch, dass für gegebenes A die Matrizen B und C eindeutig bestimmt sind.

8.3. Permutationen

In diesem Abschnitt beschäftigen wir uns genauer mit der Struktur der symmetrischen Gruppe S_n , also der Gruppe aller bijektiven Abbildungen $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$, wo die Gruppenstruktur durch die Verknüpfung von Abbildungen gegeben ist. Wir schreiben meistens $\sigma\tau$

statt $\sigma \circ \tau$. Die symmetrischen Gruppen sind wichtige Beispiele von endlichen Gruppen. Zudem ist das Material aus diesem Abschnitt auch wichtig für das Studium der Determinante im Kapitel 9.

Anstelle von Bijektionen von $\{1, \dots, n\}$ kann man natürlich ebenso gut Bijektionen irgendeiner n -elementigen Menge mit sich selbst betrachten. (Ganz formal gesagt erhält man dann eine zu S_n isomorphe Gruppe.) Man spricht daher auch von Permutationen »von n Symbolen« oder von der »symmetrischen Gruppe auf n Buchstaben«. Die Anzahl der Elemente von S_n ist $n! = 1 \cdot 2 \cdot \dots \cdot n$.

DEFINITION 8.35. Eine bijektive Abbildung $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$, also ein Element von S_n , nennen wir auch *Permutation*. $\dashv$

Um eine Permutation σ konkret anzugeben, können wir die Werte $\sigma(1)$ bis $\sigma(n)$ der Reihe nach auflisten. Dazu schreibt man die Permutation in der Regel als zweizeilige Matrix, in deren erster Zeile die Zahlen 1 bis n , und in deren zweiten Zeile der jeweilige Wert eingetragen ist. Zum Beispiel können wir für $n = 7$ die folgende Permutation betrachten:

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ 5 & 3 & 2 & 4 & 7 & 1 & 6 \end{pmatrix}$$

(Also $\sigma(1) = 5, \sigma(2) = 3$, usw.)

Eine andere Möglichkeit, eine Permutation darzustellen, ist mit Diagrammen der nebenstehenden Form. Die eingezeichneten Pfeile geben die Zuordnungsvorschrift der Abbildung $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$ an.

Die gezeigten Beispiele entsprechen den Permutationen

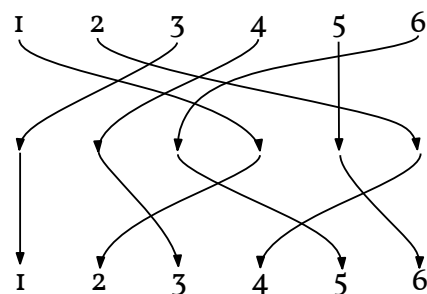
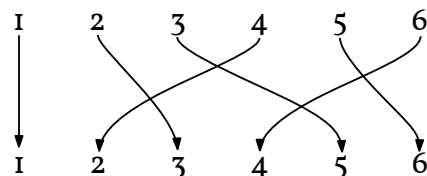
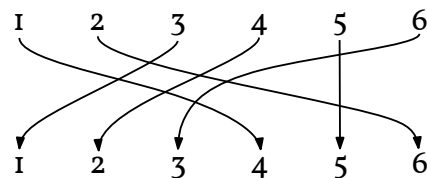
$$\tau = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 4 & 6 & 1 & 2 & 5 & 3 \end{pmatrix}$$

und

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 1 & 3 & 5 & 2 & 6 & 4 \end{pmatrix}.$$

Wenn man die Permutationen so »aufschreibt«, dann erhält man die Verkettung $\tau\sigma$, indem man die beiden Abbildungen einfach direkt untereinander zeichnet.

$$\sigma\tau = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 2 & 4 & 1 & 3 & 6 & 5 \end{pmatrix}$$



DEFINITION 8.36. Ein *Zykel* der Ordnung l (oder: ein l -Zykel) ist eine Permutation $\sigma \in S_n$, so dass eine Teilmenge $I = \{i_1, \dots, i_l\} \subseteq \{1, \dots, n\}$ mit l Elementen existiert, so dass

$$\sigma(i_1) = i_2, \quad \sigma(i_2) = i_3, \quad \dots, \quad \sigma(i_{l-1}) = \sigma(i_l), \quad \sigma(i_l) = i_1$$

und $\sigma(j) = j$ für alle $j \notin \{i_1, \dots, i_l\}$ gilt. Die Menge I heißt dann der *Träger* des Zyklus σ .

Wir schreiben dann $\sigma = (i_1, i_2, \dots, i_l)$ (wenn keine Missverständnisse zu befürchten sind, schreibt man die Einträge i_λ auch ohne Kommata nebeneinander). $\dashv$

Man beachte, dass zum Beispiel $(1234) = (2341) = (3412) = (4123)$ alle denselben Zykel der Ordnung 4 bezeichnen.

Sind $\sigma, \tau \in S_n$ Zykel, deren Träger disjunkt sind, dann gilt $\sigma\tau = \tau\sigma$. Die folgende »Rechenregel« ist oft nützlich und leicht zu überprüfen:

LEMMA 8.37. Sei $\pi \in S_n$ und sei $(i_1, \dots, i_l) \in S_n$ ein Zykel der Länge l . Dann gilt

$$\pi(i_1, \dots, i_l)\pi^{-1} = (\pi(i_1), \dots, \pi(i_l)).$$

ERGÄNZUNG 8.38 (Zerlegung in Zykel mit disjunkten Trägern).

SATZ 8.39. Jede Permutation $\sigma \in S_n$ lässt sich als Produkt von Zykeln mit paarweise disjunkten Trägern schreiben. Diese Darstellung ist bis auf die Reihenfolge der Faktoren eindeutig.

Wir verstehen den Satz so, dass die Identität als das leere Produkt geschrieben wird.

BEWEIS. Sei $\sigma \in S_n$. Wir nennen eine Teilmenge $I \subseteq \{1, \dots, n\}$ stabil unter σ , wenn $\sigma(I) = I$ gilt und nennen (nur in diesem Beweis) eine minimale nicht-leere σ -stabile Teilmenge eine *Komponente* von σ .

Dass eine Komponente aus einem einzigen Element i besteht, bedeutet einfach, dass $\sigma(i) = i$ gilt. Komponenten mit nur einem Element nennen wir *trivial*.

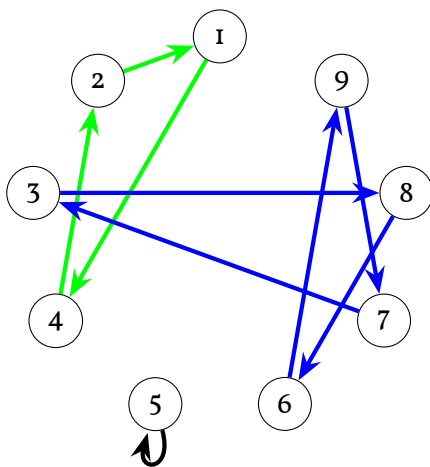
Ist $i \in \{1, \dots, n\}$, so ist die Menge $\{\sigma^r(i); r \in \mathbb{N}\}$ eine Komponente. Das ist klar, wenn $\sigma(i) = i$ ist. Andernfalls argumentieren wir wie folgt. Da S_n endlich ist, handelt es sich natürlich um eine endliche Menge: Es gibt $r < r'$ mit $\sigma^r(i) = \sigma^{r'}(i)$, also $\sigma^{r'-r}(i) = i$. Daran sehen wir, dass die Menge $\{\sigma^r(i); r \in \mathbb{N}\}$ eine σ -stabile Menge ist, und auch, dass es sich um eine Komponente von σ handelt. Denn ist j irgendein Element dieser Menge, so folgt $\{\sigma^r(j); r \in \mathbb{N}\} = \{\sigma^r(i); r \in \mathbb{N}\}$. Die einzige echte σ -stabile Teilmenge ist also die leere Menge. Es ist dann klar, dass alle Komponenten diese Form haben.

Wir führen nun Induktion nach der Anzahl der nicht-trivialen Komponenten von σ . Gibt es keine nicht-triviale Komponente, dann ist $\sigma = \text{id}$, und dies ist das leere Produkt von Zykeln.

Sei nun I eine Komponente von σ mit $l > 1$ Elementen, und sei $i \in I$. Sei τ der Zykel $(i, \sigma(i), \sigma^2(i), \dots, \sigma^{l-1}(i))$. Dass τ ein Zykel ist, folgt aus der Beschreibung der Komponenten von σ , die wir oben gegeben haben.

Dann gilt $\tau^{-1}\sigma(j) = \sigma(j)$ für $j \notin I$ und $\tau^{-1}\sigma(i) = i$ für $i \in I$. Daher hat $\tau^{-1}\sigma(j)$ eine nicht-triviale Komponente weniger als σ , ist also nach Induktionsvoraussetzung das Produkt $\sigma_1 \cdots \sigma_k$ von Zykeln mit paarweise disjunkten Trägern. Dann ist $\sigma = \tau\sigma_1 \cdots \sigma_k$. Weil $\tau^{-1}\sigma(i) = i$ für $i \in I$, ist I disjunkt zu den Trägern der σ_i . Daraus folgt die Existenz der Zerlegung als Produkt von Zykeln mit disjunkten Trägern.

Die Eindeutigkeit ergibt sich mit ähnlichen Argumenten. □



Die Zerlegung in Zykel mit disjunktem Träger im Beispiel:

$$\begin{aligned} \sigma &= \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 \\ 4 & 1 & 8 & 2 & 5 & 9 & 3 & 6 & 7 \end{pmatrix} \\ &= (142)(38697). \end{aligned}$$

Ein besonders wichtiger Fall ist der der Zykel der Ordnung 2.

DEFINITION 8.40. Eine *Transposition* in S_n ist ein Zykel der Länge 2, d.h. eine Permutation, die zwei Zahlen $i \neq j$ vertauscht und alle anderen Zahlen festlässt.

Eine *elementare Transposition* in S_n ist eine Transposition der Form $(i, i+1)$, $i \in \{1, \dots, n-1\}$. $\dashv$

Die Matrix P_{ij} aus Bemerkung 5.37 ist die Permutationsmatrix P_τ (Definition 8.25) für die Transposition $\tau = (ij)$.

LEMMA 8.41. Jede Permutation lässt sich als Produkt von Transpositionen schreiben. (Sogar: als Produkt von elementaren Transpositionen. Mit anderen Worten: Die elementaren Transpositionen erzeugen die Gruppe S_n .)

Die Aussage des Lemmas ist so zu verstehen, dass dieselbe (elementare) Transposition auch mehrfach vorkommen darf. Zum Beispiel gilt $(13) = (12)(23)(12)$.

BEWEIS. Es genügt natürlich, die Behauptung über elementare Transpositionen zu beweisen. Wir tun das durch Induktion nach n . Für $n = 1$ (und $n = 2$) ist die Sache klar. Nun nehmen wir an, dass sich jede Permutation in S_{n-1} als Produkt elementarer Transpositionen schreiben lässt. Es ist dann klar, dass sich auch jede Permutation $\sigma \in S_n$ mit der Eigenschaft $\sigma(n) = n$ als Produkt der elementaren Transpositionen $(1, 2), \dots, (n-2, n-1)$ schreiben lässt.

Ist $\sigma \in S_n$ eine beliebige Transposition, so können wir σ von links mit elementaren Transpositionen multiplizieren um zu erreichen, dass n auf sich selbst abgebildet wird. Damit sind wir in dem bereits abgehandelten Fall. $\square$

DEFINITION 8.42. Sei $\sigma \in S_n$ eine Permutation.

- (1) Ein Paar (i, j) , $1 \leq i, j \leq n$, heißt *Fehlstand* der Permutation σ , wenn $i < j$ und $\sigma(i) > \sigma(j)$.
- (2) Die *Länge* $l(\sigma)$ von σ ist die Anzahl der Fehlstände:

$$l(\sigma) = \#\{(i, j); 1 \leq i < j \leq n, \sigma(i) > \sigma(j)\}.$$

- (3) Das *Signum* $\text{sgn}(\sigma)$ von σ ist definiert als

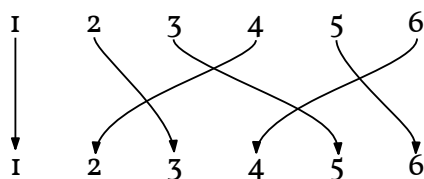
$$\text{sgn}(\sigma) := (-1)^{l(\sigma)},$$

mit anderen Worten: Ist $l(\sigma)$ gerade, so ist $\text{sgn}(\sigma) = 1$ (und man nennt σ dann auch eine *gerade Permutation*). Ist $l(\sigma)$ ungerade, so ist $\text{sgn}(\sigma) = -1$ (und man nennt σ dann auch eine *ungerade Permutation*).

$\dashv$

Es gilt $l(\text{id}) = 0$, und die Identität ist die einzige Permutation der Länge 0. Wir verstehen das Signum einer Permutation, also 1 bzw. -1 , als ganze Zahl.

Man beachte: Ist σ ein Zykel, so ist die Länge $l(\sigma)$ in der Regel *verschieden* von der Zykelordnung, und auch nicht allein durch die Zykelordnung bestimmt! Zum Beispiel gilt $l((12)) = 1$, $l((13)) = 3$.



Die Länge oder Anzahl der Fehlstände einer Permutation ist in einer visuellen Darstellung wie hier gezeigt einfach die Anzahl von Schnittpunkten der Verbindungslinien. Die hier gezeigte Permutation hat also Länge 4, ihr Signum ist folglich gleich -1 .

ERGÄNZUNG 8.43 (Länge und reduzierte Darstellungen). Man kann zeigen (und so erklärt sich der Name *Länge*), dass für eine Permutation σ die Länge $l(\sigma)$ die minimale Anzahl von Faktoren ist, die man braucht, um σ als Produkt von elementaren Transpositionen darzustellen. □ Ergänzung 8.43

LEMMA 8.44. Ist σ eine Transposition, $\tau \in S_n$ irgendeine Permutation, so gilt $\text{sgn}(\tau\sigma) = -\text{sgn}(\tau)$. Insbesondere gilt $\text{sgn}(\sigma) = -1$: alle Transpositionen sind ungerade Permutationen.

BEWEIS. Sei zunächst σ eine elementare Transposition $(i, i+1)$. Für $j = 1, \dots, i-1$ und $j = i+2, \dots, n$ gilt $\tau\sigma(j) = \tau(j)$. Je nachdem, ob $\tau(i) < \tau(i+1)$ oder $\tau(i) > \tau(i+1)$ hat also $\tau\sigma$ gerade einen Fehlstand mehr oder einen Fehlstand weniger als τ . In beiden Fällen folgt mithin $\text{sgn}(\tau\sigma) = -\text{sgn}(\tau)$.

Nun behandeln wir den Fall einer beliebigen Transposition $\sigma = (ij)$. Nach dem vorherigen Fall genügt es zu zeigen, dass wir σ als ein Produkt von elementaren Transpositionen mit einer ungeraden Anzahl von Faktoren schreiben können. Sei ohne Einschränkung $i < j$. Dann gilt für $\pi = (i, j-1)$ (wegen Lemma 8.37):

$$(ij) = \pi(j-1, j)\pi^{-1}$$

Schreiben wir π als Produkt von elementaren Transpositionen, so ist die Anzahl der Faktoren in dem Produkt $\pi(j-1, j)\pi^{-1}$ offenbar ungerade (denn jede Transposition ist zu sich selbst invers).

Der zweite Teil folgt, indem wir den ersten Teil mit $\tau = \text{id}$ anwenden. □

SATZ 8.45. Sind $\sigma, \tau \in S_n$, so gilt

$$\text{sgn}(\sigma\tau) = \text{sgn}(\sigma)\text{sgn}(\tau).$$

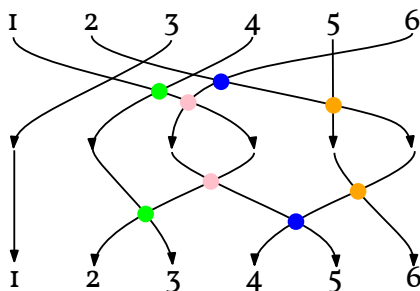
BEWEIS. Wir haben gesehen, dass wir τ als ein Produkt von Transpositionen schreiben können. Ist r die Anzahl der Faktoren, so gilt $\text{sgn}(\tau) = (-1)^r$. Aus dem vorherigen Lemma folgt $\text{sgn}(\sigma\tau) = (-1)^r \text{sgn}(\sigma)$. □

Betrachten wir $\{-1, 1\}$ als Gruppe bezüglich der Multiplikation, so können wir den vorherigen Satz auch wie folgt ausdrücken:

KOROLLAR 8.46. Die Abbildung $\text{sgn}: S_n \rightarrow \{1, -1\}$ ist ein Gruppenhomomorphismus.

Wir sehen damit: Schreibt man eine ungerade Permutation als ein Produkt von Transpositionen, so ist die Anzahl der Faktoren ungerade. Schreibt man eine gerade Permutation als ein Produkt von Transpositionen, so ist die Anzahl der Faktoren gerade. Weil $\text{sgn}(\sigma)\text{sgn}(\sigma^{-1}) = \text{sgn}(\text{id}) = 1$ ist, folgt auch $\text{sgn}(\sigma^{-1}) = \text{sgn}(\sigma)$ für alle $\sigma \in S_n$.

BEMERKUNG 8.47. Anhand des obigen Beispiels kann man sich die Tatsache, dass das Signum mit der Verkettung kompatibel ist, auch folgendermaßen plausibel machen. Wir wollen wissen, ob die Anzahl der Fehlstände einer Verkettung gerade oder ungerade ist. Wenn wir die Verkettung wie nebenstehend zeichnen, gibt es zwar unnötige Überschneidungen; die Länge einer Verkettung $\sigma\tau$ ist üblicherweise nicht gleich der Summe der Längen von σ und τ . Wenn man die »Sache aber geradezieht«, dann fallen die unnötigen Überschneidungen immer jeweils in Paaren weg (wie im Bild hier farbig markiert). Ob die Anzahl der Überschneidungen gerade oder ungerade ist, verändert sich also nicht. ◇



Einen Zykel $(i_1, i_2, \dots, i_l)$ kann man als Produkt von $l - 1$ Transpositionen schreiben. (Wie?) Das Signum ist daher $(-1)^{l-1}$.

ERGÄNZUNG 8.48 (Zum Signum einer Permutation). Alternativ kann man das Signum einer Permutation σ definieren als

$$\operatorname{sgn}(\sigma) = \prod_{1 \leq i < j \leq n} \frac{\sigma(j) - \sigma(i)}{j - i}.$$

Es ist nicht schwer zu begründen, dass diese Definition mit unserer Definition übereinstimmt. Um die Multiplikativität zu zeigen, kann man dann wie folgt rechnen:

$$\begin{aligned} \operatorname{sgn}(\sigma\tau) &= \prod_{1 \leq i < j \leq n} \frac{\sigma\tau(j) - \sigma\tau(i)}{j - i} \\ &= \prod_{1 \leq i < j \leq n} \left(\frac{\sigma\tau(j) - \sigma\tau(i)}{\tau(j) - \tau(i)} \cdot \frac{\tau(j) - \tau(i)}{j - i} \right) \\ &= \left(\prod_{1 \leq i < j \leq n} \frac{\sigma(j) - \sigma(i)}{j - i} \right) \cdot \left(\prod_{1 \leq i < j \leq n} \frac{\tau(j) - \tau(i)}{j - i} \right) \\ &= \operatorname{sgn}(\sigma) \operatorname{sgn}(\tau), \end{aligned}$$

wobei im dritten Schritt eine »Indexverschiebung« stattfindet, die ausnutzt, dass für die Permutation τ die Werte $\tau(1), \dots, \tau(n)$ die Zahlen $1, \dots, n$ sind. Dabei ist zu beachten, dass

$$\frac{\sigma(j) - \sigma(i)}{j - i} = \frac{\sigma(i) - \sigma(j)}{i - j}$$

gilt. Es ist daher egal, ob $i < j$ oder $j < i$ gilt; Hauptsache, jede zweielementige Teilmenge $\{i, j\}$ kommt im Produkt genau einmal vor. □ Ergänzung 8.48

8.4. Die Bruhat-Zerlegung *

Seien K ein Körper und $n \geq 1$.

Sei $B \subset GL_n(K)$ die Teilmenge der invertierbaren oberen Dreiecksmatrizen. Es ist nicht schwer zu sehen, dass B eine Untergruppe von $GL_n(K)$ ist: Das Produkt von oberen Dreiecksmatrizen ist eine obere Dreiecksmatrix, und das Inverse einer invertierbaren oberen Dreiecksmatrix ist eine obere Dreiecksmatrix. Dass die Einheitsmatrix E_n in B liegt, ist offensichtlich.

Die Menge U aller Matrizen in B , deren Diagonaleinträge sämtlich $= 1$ sind, ist eine Untergruppe.

Wir bezeichnen mit $W \subset GL_n(K)$ die Untergruppe der Permutationsmatrizen (siehe Definition 8.25).

SATZ 8.49 (Bruhat-Zerlegung). *Sei $A \in GL_n(K)$. Dann existiert eine Permutationsmatrix $w \in GL_n(K)$, so dass sich A in der Form bwb' mit $b, b' \in B$ schreiben lässt. Ferner gilt: die Matrix w ist durch A eindeutig bestimmt. Außerdem lässt sich bei geeigneter Wahl von b und b' erreichen, dass alle Diagonaleinträge von b gleich 1 sind.*

BEMERKUNG 8.50. Wir können den Satz umformulieren als

$$GL_n(K) = \bigsqcup_{w \in W} UwB.$$

Hier schreiben wir

$$UwB = \{uwb; u \in U, b \in B\}.$$

Es gilt

$$UwB = BwB = \{bwb'; b, b' \in B\}.$$

Das Symbol $\sqcup$ bedeutet, dass es sich um eine *disjunkte Vereinigung* handelt, das heißt $GL_n(K) = \bigcup_w UwB$ und $UwB \cap Uw'B = \emptyset$ für alle $w \neq w'$. Die Gruppe $GL_n(K)$ als eine disjunkte Vereinigung zu beschreiben, nennt man auch eine Zerlegung von $GL_n(K)$. Der Name *Bruhat-Zerlegung* geht zurück auf den französischen Mathematiker François Bruhat (1929–2007).

Es ist wichtig zu beachten, dass b und b' im Satz in der Regel nicht eindeutig bestimmt sind.

◇

BEWEIS. I. Existenz der Zerlegung. Wir führen den sogenannten Gauß-Bruhat-Algorithmus durch und bringen A im ersten Schritt durch geeignete Zeilen- und Spaltenumformungen vom Typ I auf eine Matrix, die in jeder Zeile und in jeder Spalte genau einen Eintrag $\neq 0$ hat.

Sei $A = (a_{ij})_{i,j} \in GL_n(K)$. Sei i_1 maximal mit $a_{i_1,1} \neq 0$. Bringe alle Einträge der ersten Spalte in den Zeilen $i = 1, \dots, i_1 - 1$ durch geeignete Zeilenumformungen vom Typ I auf 0. Danach ist nur noch ein Eintrag in der ersten Spalte vorhanden. Bringe nun alle Einträge der i_1 -ten Zeile in den Spalten $2, 3, \dots, n$ durch geeignete Spaltenumformungen vom Typ I auf 0. Danach ist auch in Zeile i_1 nur noch ein einziger Eintrag (nämlich $a_{i_1,1}$) ungleich 0.

Bezeichne die Einträge der neuen Matrix wieder mit a_{ij} .

Sei nun i_2 maximal mit $a_{i_2,2} \neq 0$. Nach dem ersten Schritt ist jedenfalls $i_2 \neq i_1$ (und insbesondere $a_{i_2,1} = 0$). Bringe alle Einträge der zweiten Spalte in den Zeilen $i = 1, \dots, i_2 - 1$ durch geeignete Zeilenumformungen vom Typ I auf 0. Danach ist nur noch ein Eintrag in der zweiten Spalte vorhanden. Bringe nun alle Einträge der i_2 -ten Zeile in den Spalten $3, \dots, n$ durch geeignete Spaltenumformungen vom Typ I auf 0. Danach ist auch in Zeile i_2 nur noch ein einziger Eintrag (nämlich $a_{i_2,2}$) ungleich 0.

Fahre entsprechend fort mit den Spalten $3, \dots, n$. Alle diese elementaren Zeilen- und Spaltenumformungen lassen sich beschreiben durch Multiplikation von A mit Elementarmatrizen $E_{ij}(a)$ von links bzw. rechts. Weil wir immer nur Vielfache einer Zeile zu einer *darüberliegenden* Zeile bzw. Vielfache einer Spalte zu einer *davon rechts liegenden* Spalte addieren, gilt stets $i < j$, d.h. $E_{ij}(a) \in U$. Es gibt also $u, u' \in U$, so dass uAu' in jeder Zeile und jeder Spalte genau einen Eintrag $\neq 0$ hat. Sei d die eindeutig bestimmte Diagonalmatrix, so dass $w := uAu'd$ eine Permutationsmatrix ist.

Wir setzen $b := u^{-1} \in U$, $b' = d^{-1}(u')^{-1} \in B$ und haben $A = bwb'$ wie gewünscht.

II. Eindeutigkeit von w . Sei $b_1wb'_1 = b_2vb'_2$, wo $b_i, b'_i \in B$, $v, w \in W$. Sei σ die zu v , und τ die zu w gehörige Permutation. Zu zeigen ist $w = v$. Jedenfalls gilt für $b := b_2^{-1}b_1$, $b' := b'_2(b'_1)^{-1}$:

$$bw = vb'.$$

Vergleiche nun die ersten Spalten dieser Matrizen. Die erste Spalte von bw ist die $\tau(1)$ -te Spalte von b . Jedenfalls an der Stelle $\tau(1)$ befindet sich hier ein Eintrag $\neq 0$. Die erste Spalte von vb' ist das b'_1 -fache der ersten Spalte von v . Höchstens (und sogar genau) an der Stelle $\sigma(1)$ befindet sich hier ein Eintrag $\neq 0$. Es folgt $\tau(1) \in \{\sigma(1)\}$, also $\tau(1) = \sigma(1)$.

Vergleiche nun die zweiten Spalten der Matrizen. Die zweite Spalte von bw ist die $\tau(2)$ -te Spalte von b . Jedenfalls an der Stelle $\tau(2)$ befindet sich hier ein Eintrag $\neq 0$. Die zweite Spalte von vb' ist die Summe des b'_{12} -fachen der ersten und des b'_{22} -fachen der zweiten Spalte. Höchstens an den Stellen $\sigma(1)$ und $\sigma(2)$ befinden sich hier Einträge $\neq 0$. Es folgt $\tau(2) \in \{\sigma(1), \sigma(2)\}$. Da τ bijektiv ist und $\tau(1) = \sigma(1)$, folgt $\tau(2) = \sigma(2)$.

Induktiv folgt $\tau = \sigma$, also $w = v$, wie gewünscht. □

BEMERKUNG 8.51. Bei der praktischen Durchführung des Algorithmus ist es unter Umständen sinnvoller, die Einträge $a_{i,v}$ direkt durch Multiplikation der entsprechenden Spalte mit $a_{i,v}^{-1}$ auf 1 zu bringen. $\diamond$

ERGÄNZUNG 8.52 (Die LR-Zerlegung). Die Bruhat-Zerlegung ist eng verwandt mit der LR-Zerlegung, die auch (neben der direkten Anwendung des Gauß-Verfahrens) bei der Lösung von Gleichungssystemen in der Praxis eine Rolle spielt (für Gleichungssysteme in der Größenordnung von $\leq 10\,000$ Gleichungen). Darunter versteht man den folgenden Satz:

SATZ 8.53. Sei $A \in GL_n(K)$ eine invertierbare Matrix. Dann existieren eine Permutationsmatrix P , eine untere Dreiecksmatrix L (mit Einsen auf der Diagonale) und eine obere Dreiecksmatrix R , so dass

$$A = PLR.$$

Wie der Beweis zeigen wird, ist P genau dieselbe Permutationsmatrix, die auch in der Bruhat-Zerlegung von A auftritt. (Eine Variante ist, den Satz auf $w_o A$ anzuwenden; gilt $w_o A \in Bw_o B$ – in gewissem Sinne der Normalfall – so erhält man dann $w_o A = w_o LR$, also $A = LR$ für eine untere Dreiecksmatrix L und eine obere Dreiecksmatrix R .)

BEWEIS. Sei $w_o = P_\sigma$, wobei die Permutation σ definiert ist durch $\sigma(i) = n - i + 1$. Dann gilt $w_o = w_o^{-1}$ und für eine obere Dreiecksmatrix M ist $w_o M w_o^{-1}$ eine untere Dreiecksmatrix.

Sei nun $A = bwb'$ mit $b, b' \in B$, $w \in W$, die Bruhat-Zerlegung von A . Es folgt dann leicht aus dem Beweis der Bruhat-Zerlegung, dass $w_o w^{-1} A \in Bw_o B$ ist, d.h. die Permutationsmatrix in der Bruhat-Zerlegung der Matrix $w_o w^{-1} A$ ist w_o . Schreiben wir

$$w_o w^{-1} A = L' w_o R$$

mit $L', R \in B$, so erhalten wir

$$A = w(w_o L' w_o) R,$$

wir können also $P := w$, $L := w_o L' w_o$ setzen. $\square$

Die LR-Zerlegung heißt manchmal auch LU-Zerlegung und lässt sich auch auf den Fall verallgemeinern, dass A nicht invertierbar ist (vergleiche [LM] Satz 5.4).

Ist A invertierbar und ein lineares Gleichungssystem $Ax = b$ gegeben, und ist $A = PLR$ wie oben die LR-Zerlegung von A , so haben wir für die eindeutig bestimmte Lösung x :

$$x = A^{-1}b = R^{-1}L^{-1}P^{-1}b.$$

Der Vektor $b' = P^{-1}b$ entsteht durch Vertauschung der Einträge von b . Der Vektor $x' = L^{-1}b'$ ist die eindeutige Lösung des linearen Gleichungssystems $Lx' = b'$, die unmittelbar »abgelesen« werden kann, weil L eine Dreiecksmatrix ist. Dann ist x die Lösung des Gleichungssystems $Rx = x'$, und auch diese lässt sich direkt ablesen, weil R eine Dreiecksmatrix ist.

Wenn man das Verfahren in der Praxis in Fällen anwendet, die man nicht per Hand lösen kann, sind natürlich Fragen der Effizienz und vor allem der Rechengenauigkeit zu beachten, insbesondere Strategien, um Rundungsfehler zu vermeiden und/oder zu kontrollieren. Diese Aspekte lassen wir hier völlig außer Acht. $\square$ Ergänzung 8.52

8.5. Ergänzungen *

Zur Gruppentheorie lassen sich ganze Bücher schreiben, und auch wenn dieser Abschnitt mit Ergänzungen sehr lang ist, kratzen wir im Grunde nur an der Oberfläche. Immerhin spannen wir einen weiten Bogen, beginnend mit kleinen Ergänzungen zur allgemeinen Theorie, über Anwendungen der Gruppentheorie auf abstrakte Fragen der Zahlentheorie mit dem quadratischen Reziprozitätsgesetz als einem besonders schönen Ergebnis, zu alltagsnahen (wenn vielleicht auch nicht sehr wichtigen) Themen wie dem 15-Puzzle und Rubiks Zauberwürfel. Vielleicht weckt das ja Ihr Interesse, das eine oder andere Thema weiterzuverfolgen.

Wenn Sie mehr »Gruppentheorie« lernen möchten, als wir in der linearen Algebra behandeln (und benötigen):

Ein Kapitel über Gruppentheorie gibt es praktisch in jedem Algebra-Buch, und was dort erklärt wird, wird Ihnen vermutlich erstmal ausreichen. Darüberhinaus hat das Thema Gruppentheorie (und damit eng verbunden: die Darstellungstheorie, siehe Abschnitt 8.5.2 für ein paar Bemerkungen dazu) so viele Facetten, dass man erstmal genauer eingrenzen sollte, über welchen Teil man etwas lernen möchte.

Speziell zum Thema Gruppentheorie und Symmetrie in Verbindung mit geometrischen Aspekten, wie es gut zur linearen Algebra passt, können Sie zum Beispiel in die folgenden Bücher schauen:

M. Armstrong, *Groups and Symmetry*, Springer 1988.

S. Rosebrock, *Geometrische Gruppentheorie*, Vieweg+Teubner 2010

<https://doi.org/10.1007/978-3-8348-9648-3>

8.5.1. Elementare Zahlentheorie. Viele Ergebnisse der »elementaren Zahlentheorie« lassen sich leicht aus der Theorie der Gruppen herleiten. Wir beginnen damit, ein kleines bisschen mehr Gruppentheorie zu behandeln. Ein wichtiger Begriff ist die Ordnung eines Gruppenelements, die im folgenden Lemma definiert wird.

LEMMA 8.54. Sei G eine endliche Gruppe. Für jedes $g \in G$ existiert $n \in \mathbb{N}_{>0}$, so dass $g^n = 1$. Das kleinste solche n nennen wir die Ordnung des Elements g , und wir bezeichnen diese Zahl mit $\text{ord}(g)$.

Die Ordnung von g ist gleich der Anzahl der Elemente der von g erzeugten Untergruppe $\langle g \rangle$ von G .

BEWEIS. Sei $g \in G$. Weil G endlich ist, ist klar, dass $i < j \in \mathbb{N}$ existieren mit $g^i = g^j$, und es folgt $g^{j-i} = 1$.

Es ist noch zu zeigen, dass $\langle g \rangle = \text{ord}(g)$ gilt. Wir schreiben $n = \text{ord}(g)$. Mit dem obigen Argument, angewandt auf die Gruppe $\langle g \rangle$, ist dann klar, dass $n \leq \# \langle g \rangle$ ist. Andererseits ist $\{g^i; i = 0, \dots, n-1\}$ eine Gruppe, die g enthält (denn $g^{-1} = g^{n-1}$, und allgemeiner $g^i = g^j$ für alle $i, j \in \mathbb{Z}$ mit $n \mid (i-j)$ – man kann also jeden Exponenten durch seinen Rest bei Division durch n ersetzen). Da $\langle g \rangle$ keine echten Untergruppen hat, die g enthalten, folgt die Behauptung. $\square$

Die Ordnung des neutralen Elements einer endlichen Gruppe G ist 1, und alle anderen Elemente haben Ordnung > 1 . Der folgende Satz zeigt, dass die Ordnung eines Elements $g \in G$ immer ein Teiler von $\#G$ ist. Die Zahl $\#G$, die Anzahl der Elemente von G , bezeichnet man auch als die *Ordnung der Gruppe* G .

SATZ 8.55 (Satz von Lagrange). Sei G eine endliche Gruppe und H eine Untergruppe von G . Dann ist $\#H$ ein Teiler von $\#G$.

Insbesondere gilt: Ist $g \in G$, so ist die Ordnung von g ein Teiler von $\#G$.

BEWEIS. Wir verwenden für den Beweis den Begriff der Äquivalenzrelation aus Abschnitt 3.14.2.

Für $g, g' \in G$ setzen wir

$$g \sim g' :\iff g^{-1}g' \in H.$$

Da H eine Untergruppe ist, ist dies eine Äquivalenzrelation auf G . Die Reflexivität gilt, weil H das neutrale Element enthält, die Symmetrie folgt daraus, dass für jedes Element aus H auch sein Inverses in H liegt, und die Transitivität ergibt sich daraus, dass H unter der Bildung von Produkten abgeschlossen ist.

Also ist G die disjunkte Vereinigung der Äquivalenzklassen. Die Behauptung des Satzes folgt nun daraus, dass in dieser speziellen Situation jede Äquivalenzklasse genau $\#H$ Elemente haben. Denn ist $X \subseteq G$ irgendeine der Äquivalenzklassen und $x \in X$, so ist die Abbildung

$$H \longrightarrow X, \quad h \mapsto xh,$$

eine Bijektion. Wegen $x^{-1}xh = h \in H$ gilt $x \sim xh$, also liegt xh in X . Weil G eine Gruppe ist, ist klar, dass die Abbildung injektiv ist. Ist schließlich $x' \in X$ irgendein Element, dann ist $x \sim x'$, also $h := x^{-1}x' \in H$, und $x' = xh$ im Bild der obigen Abbildung. $\square$

Als direkte Folgerung erhalten wir:

KOROLLAR 8.56 (Satz von Euler). Sei $n \geq 1$ eine natürliche Zahl. Wir definieren $\varphi(n) = \#(\mathbb{Z}/n)^\times$, die Anzahl der Elemente von $\mathbb{Z}/n$, die ein multiplikatives Inverses besitzen. Dies ist die Anzahl der zu n teilerfremden Zahlen zwischen 1 und $n - 1$. Siehe Ergänzung 8.9.

Die Funktion $\varphi: \mathbb{N}_{\geq 1} \rightarrow \mathbb{N}_{\geq 1}$ heißt die Eulersche φ -Funktion.

Dann gilt für alle $a \in \mathbb{Z}$, die zu n teilerfremd sind, dass

$$a^{\varphi(n)} \equiv 1 \pmod{n},$$

mit anderen Worten: $a^{\varphi(n)}$ hat Rest 1 bei Division durch n .

Ist n eine Primzahl, so ist $\varphi(n) = n - 1$, und die Aussage des Korollars ist genau die des kleinen Fermatschen Satzes (Satz 4.21), für den wir hier einen neuen Beweis erhalten.

SATZ 8.57. Seien $m, n \in \mathbb{N}$, $m, n \geq 1$, so dass $\text{ggT}(m, n) = 1$.

(1) Die Abbildung

$$\mathbb{Z}/(mn) \rightarrow \mathbb{Z}/m \times \mathbb{Z}/n, \quad a \mapsto (a, a),$$

ist ein Isomorphismus von (additiven) Gruppen.

(2) Die Abbildung aus Teil (1) liefert durch Einschränkung einen Isomorphismus

$$\mathbb{Z}/(mn)^\times \rightarrow (\mathbb{Z}/m)^\times \times (\mathbb{Z}/n)^\times$$

von (multiplikativen) Gruppen. Insbesondere gilt $\varphi(mn) = \varphi(m)\varphi(n)$ für alle teilerfremden natürlichen Zahlen $m, n \geq 1$.

Die Aussage in Teil (1) ist auch unter dem Namen *chinesischer Restsatz* bekannt. Die letzte Aussage von Teil (2) bezeichnet man auch als die *Multiplikativität der φ -Funktion*. Es ist aber zu beachten, dass diese Multiplikativität nur für *teilerfremde* Zahlen m, n gilt, wie man leicht an Beispielen sieht.

BEWEIS. zu (1). Das Bild eines Elements $a \in \mathbb{Z}/(mn)$ ist das Paar (a, a) , wo hier der linke Eintrag als Restklasse in $\mathbb{Z}/m$, der rechte als Restklasse in $\mathbb{Z}/n$ zu verstehen ist. Die Abbildung ist unabhängig davon, wie wir ein Element von $\mathbb{Z}/(mn)$ durch eine ganze Zahl repräsentieren, also wohldefiniert, und es ist klar, dass es sich um einen Gruppenhomomorphismus handelt. Da beide Seiten mn Elemente haben, genügt es, die Injektivität zu zeigen. Diese ist aber klar, denn wenn a auf $(0, 0)$ abgebildet wird, so bedeutet das gerade, dass $m \mid a$

und $n \mid a$; da m und n teilerfremd sind, impliziert das $mn \mid a$, also ist das Element a in $\mathbb{Z}/(mn)$ gleich 0.

zu (2). Ist f die Abbildung aus Teil (1), so gilt jedenfalls $f(\mathbb{Z}/(mn)^\times) \subseteq (\mathbb{Z}/m)^\times \times (\mathbb{Z}/n)^\times$, denn aus $ab = 1$ folgt $f(a)f(b) = f(ab) = f(1) = 1$. Die Injektivität überträgt sich direkt, und es bleibt nur noch die Surjektivität zu zeigen.

Sei dazu $(b, c) \in (\mathbb{Z}/m)^\times \times (\mathbb{Z}/n)^\times$. Wir zeigen, dass eine ganze Zahl a existiert, die teilerfremd ist zu mn , und deren Restklasse in $\mathbb{Z}/m$ (bzw. $\mathbb{Z}/n$) gleich b (bzw. gleich c) ist. Da die Abbildung f bijektiv ist, gibt es sowieso nur einen Kandidaten: Wir bezeichnen mit $a \in \mathbb{Z}/(mn)$ das eindeutig bestimmte Urbild von (b, c) unter f wie in Teil (1). Es ist dann zu zeigen, dass $\text{ggT}(a, mn) = 1$. Da die Bilder b und c von a in $\mathbb{Z}/m$ und $\mathbb{Z}/n$ nach Voraussetzung »teilerfremd« zu m bzw. n sind, folgt aber, dass a zu m und zu n teilerfremd ist. Da auch m und n unter sich teilerfremd sind, liefert das die Behauptung.

Dass in dieser Situation $\varphi(mn) = \varphi(m)\varphi(n)$ gilt, folgt, indem wir die Anzahl der Elemente in der linken und rechten Seite vergleichen. $\square$

BEMERKUNG 8.58 (Das RSA-Verfahren). Die Multiplikativität der φ -Funktion zusammen mit dem Satz von Euler (Korollar 8.56) sind die Grundlage des **RSA-Verfahrens**² (vergleiche Bemerkung 4.22), dessen Prinzip das folgende ist:

Alice wählt als *privaten Schlüssel* zwei verschiedene, große Primzahlen p und q aus, berechnet das Produkt $N = pq$ und sucht eine zu $\varphi(N) = (p-1)(q-1)$ teilerfremde Zahl e . Weil sie $\varphi(N)$ kennt, kann Sie $a, b \in \mathbb{Z}$ mit $ae + b\varphi(N) = 1$ finden.

Als *öffentlichen Schlüssel* veröffentlicht sie die Zahlen N und e .

Bob, der Alice eine Nachricht schicken möchte, zerlegt die Nachricht in Teile, die sich als Elemente $M \in \mathbb{Z}/N$ schreiben lassen, und schickt Alice jeweils die Zahl $C = M^e \in \mathbb{Z}/N$.

Empfängt Alice die Nachricht C , so berechnet sie

$$C^a = M^{ea} = M^{ae+b\varphi(N)} = M \quad \text{in } \mathbb{Z}/N.$$

Im zweiten Schritt wird der Satz von Euler verwendet, der impliziert, dass $M^{\varphi(N)} = 1$ ist. Dass Alice M aus C rekonstruieren kann, zeigt die Durchführbarkeit des Verfahrens.

Die Sicherheit von RSA beruht darauf, dass es nicht mit akzeptablem Aufwand möglich ist, aus den Informationen N und e allein eine Zahl a mit $ae = 1 \in \mathbb{Z}/N^\times$ zu berechnen. Eine offensichtliche Möglichkeit, dies zu tun, wäre, die Zahl N in ihre Primfaktoren zu zerlegen und daraus $\varphi(N)$ zu berechnen. Danach ist es einfach, a zu finden (und Alice muss ja zu Beginn diese Rechnung durchführen). Man kann (und muss) N so groß wählen, dass die Faktorisierung von N mit allen bekannten Verfahren zu lange dauern würde. $\diamond$

ERGÄNZUNG 8.59 (Primitivwurzeln). Die multiplikative Gruppe eines endlichen Körper $\mathbb{F}_p$ ist immer eine zyklische Gruppe, d.h. sie lässt sich von einem einzigen Element erzeugen. Allgemeiner gilt das folgende Resultat.

THEOREM 8.60. Sei $G \subseteq K^\times$ eine endliche Untergruppe der multiplikativen Gruppe eines Körpers K . Dann ist G zyklisch.

Wir benutzen im Beweis das folgende Lemma:

LEMMA 8.61. Sei G eine endliche kommutative Gruppe.

(1) Seien $a, b \in G$ Elemente mit $\text{ord}(a) = m$, $\text{ord}(b) = n$ und $\text{ggT}(m, n) = 1$. Dann gilt $\text{ord}(ab) = mn$.

²<https://de.wikipedia.org/wiki/RSA-Kryptosystem>

- (2) Sei n die größte Zahl, die als Ordnung eines Elements von G auftritt. Dann gilt für alle $a \in G$, dass $\text{ord}(a) \mid n$.

BEWEIS. Wir schreiben die Gruppe G multiplikativ.

zu (1). Es ist klar, dass $(ab)^{mn} = 1$ gilt.

Zunächst bemerken wir, dass $a^i = b^j$ für $i, j \in \mathbb{Z}$ nur gelten kann, wenn $a^i = 1 = b^j$ ist (also i von m und j von n geteilt werden). Denn es folgt $a^{in} = b^{jn} = 1$, also $m \mid in$, wegen der Teilerfremdheit von m und n also $m \mid i$, und analog $n \mid j$.

Ist nun $(ab)^i = 1$ für eine ganze Zahl i , so haben wir $a^i b^i = (ab)^i = 1$, also folgt aus der Vorbemerkung, dass $m \mid i$ und $n \mid i$. Weil m und n teilerfremd sind, impliziert das $mn \mid i$. Insgesamt folgt $\text{ord}(ab) = mn$.

zu (2). Sei n wie im Lemma beschrieben, sei $b \in G$ ein Element mit $\text{ord}(b) = n$ und sei $a \in G$. Sei $k = \text{ggT}(n, \text{ord}(a))$, $i = \text{ord}(a) / k$. Dann hat a^k die Ordnung i , und $\text{ggT}(i, n) = 1$. Aus Teil (1) folgt, dass das Element ab die Ordnung in hat. Nach Definition von n folgt $i = 1$, also $k = \text{ord}(a)$, und das bedeutet genau, dass $\text{ord}(a) \mid n$. $\square$

BEWEIS DES THEOREMS. Neben dem Lemma ist die entscheidende Zutat für den hier gegebenen Beweis, dass eine Polynomfunktion der Form $x \mapsto \sum_{i=0}^d a_i x^i$ höchstens d Nullstellen in K haben kann (Satz 4.25).

Sei also $G \subseteq K^\times$ eine endliche Untergruppe und sei n die maximale Ordnung eines Elements von G . Wir wollen zeigen, dass $n = \#G$ gilt. Aus Teil (2) des Lemmas folgt, dass für alle Elemente $a \in G$ gilt, dass $\text{ord}(a) \mid n$, also $a^n = 1$. Alle Elemente von G sind also Nullstellen der Polynomfunktion $x \mapsto x^n - 1$. Es folgt $\#G \leq n$. Andererseits wissen wir wegen Satz 8.55, dass $n \leq \#G$ gilt.

Es gibt [verschiedene andere Beweise dieses Resultats](#)³. (In dieser Quelle wird speziell der Fall $K = \mathbb{F}_p$ betrachtet; einige der Beweise liefern aber das Ergebnis in der allgemeinen Form.) $\square$

Wenn die Restklasse von $a \in \mathbb{Z}$ die Gruppe $\mathbb{F}_p^\times$ erzeugt, dann sagt man auch, a sei eine *Primitivwurzel* modulo p . Der Satz sagt aus, dass stets eine Primitivwurzel modulo p existiert. Es ist aber nicht einfach, systematisch eine solche zu finden. Eine [Vermutung](#)⁴ von [E. Artin](#)⁵ (1898–1962) besagt, dass für jede ganze Zahl $a \neq -1$, die keine Quadratzahl ist, die Restklasse von a für unendlich viele Primzahlen p eine Primitivwurzel ist. Es gibt keine einzige Zahl a , für die man die Vermutung bisher beweisen konnte! Es ist aber bekannt, dass die Artinsche Vermutung aus der »verallgemeinerten Riemannschen Vermutung«, einer Variante der berühmten [Riemannschen Vermutung](#)⁶, folgt. $\square$ Ergänzung 8.59

ERGÄNZUNG 8.62 (Quadratische Reste). Sei p eine ungerade Primzahl. Es ist eine interessante Frage, wie man feststellen kann, ob eine Gleichung der Form $X^2 = a$ in $\mathbb{F}_p$ lösbar ist. Hier kann man a als Element in $\mathbb{F}_p$ betrachten, oder als ganze Zahl und dann implizit zur Restklasse in $\mathbb{F}_p$ übergehen. Mit anderen Worten: Für $a \in \mathbb{Z}$ wird gefragt, ob eine ganze Zahl $n \in \mathbb{Z}$ existiert, so dass $n^2 - a$ durch p teilbar ist.

Wir betrachten die Abbildung $\mathbb{F}_p^\times \rightarrow \mathbb{F}_p^\times, x \mapsto x^2$. Wie man unmittelbar nachprüft, handelt es sich um einen Gruppenhomomorphismus. Dieser ist nicht injektiv, denn für a im Bild,

³<https://kconrad.math.uconn.edu/blurbs/grouptheory/cyclicmodp.pdf>

⁴https://en.wikipedia.org/wiki/Artin%27s_conjecture_on_primitive_roots

⁵https://de.wikipedia.org/wiki/Emil_Artin

⁶https://de.wikipedia.org/wiki/Riemannsche_Vermutung

etwa $a = x^2$ gilt auch $a = (-x)^2$. Allerdings sind dann x und $-x$ die einzigen Urbilder von a , da quadratische Gleichungen über einem Körper höchstens zwei Lösungen haben. Da nach Voraussetzung $p > 2$ ist, gilt $x \neq -x$. Alle Elemente des Bildes haben also genau zwei Urbilder, und wir sehen, dass das Bild dieser Abbildung $(p-1)/2$ Elemente umfasst, mit anderen Worten, dass für genau $(p-1)/2$ Elemente $a \in \mathbb{F}_p^\times$ die Gleichung $X^2 - a$ lösbar ist.

Nach Satz 4.21 ist $x^{p-1} = 1$ für alle $x \in \mathbb{F}_p^\times$, also $x^{\frac{p-1}{2}} \in \{1, -1\}$. Ist $a = x^2$, so folgt $a^{\frac{p-1}{2}} = x^{p-1} = 1$.

Andererseits haben wir in Theorem 8.60 gesehen, dass die Gruppe $\mathbb{F}_p^\times$ zyklisch ist. Es gibt also ein Element $x \in \mathbb{F}_p^\times$ von Ordnung $p-1$. Für dieses kann nicht $x^{\frac{p-1}{2}} = 1$ gelten. Das bedeutet, dass der Gruppenhomomorphismus $\mathbb{F}_p^\times \rightarrow \{1, -1\}$, $x \mapsto x^{\frac{p-1}{2}}$, surjektiv ist. Für alle Elemente im Bild eines Gruppenhomomorphismus stimmen die Anzahlen der Urbildmengen überein; es werden also $\frac{p-1}{2}$ Elemente auf 1, und ebenso $\frac{p-1}{2}$ Elemente auf -1 abgebildet. Wir haben bereits gesehen, dass genau die Hälfte der Elemente von $\mathbb{F}_p^\times$ ein Quadrat ist, und es folgt:

LEMMA 8.63 (Kriterium von Euler). *Sei $a \in \mathbb{F}_p^\times$. Es gibt genau dann ein Element $x \in \mathbb{F}_p^\times$ mit $x^2 = a$, wenn $a^{\frac{p-1}{2}} = 1$.*

DEFINITION 8.64 (Legendre-Symbol). Wir definieren für eine ungerade Primzahl p und eine zu p teilerfremde ganze Zahl a das *Legendre-Symbol*

$$\left(\frac{a}{p}\right) := \begin{cases} 1 & \text{falls die Restklasse von } a \text{ in } \mathbb{F}_p \text{ ein Quadrat ist,} \\ -1 & \text{sonst.} \end{cases}$$

⊥

Wir betrachten hier 1 und -1 als ganze Zahlen, damit wir die Legendre-Symbole für verschiedene Primzahlen vergleichen können. Die Restklasse von $\left(\frac{a}{p}\right)$ in $\mathbb{F}_p^\times$ ist nach dem Eulerschen Kriterium gleich $a^{\frac{p-1}{2}}$. Das Legendre-Symbol ist »als Ganzes« zu verstehen, der waagerechte Strich ist kein Bruchstrich (und dementsprechend kann man auch die Klammern nicht weglassen).

Aus dem Euler-Kriterium folgt, dass die Abbildung $a \mapsto \left(\frac{a}{p}\right)$ ein Gruppenhomomorphismus $\mathbb{F}_p^\times \rightarrow \{1, -1\}$ ist, mit anderen Worten: Sind $a, b \in \mathbb{Z}$ zu p teilerfremd, so gilt

$$\left(\frac{ab}{p}\right) = \left(\frac{a}{p}\right) \left(\frac{b}{p}\right).$$

Sind p und q verschiedene ungerade Primzahlen, dann gibt es einen außerordentlich überraschenden Zusammenhang zwischen $\left(\frac{p}{q}\right)$ und $\left(\frac{q}{p}\right)$, das sogenannte quadratische Reziprozitätsgesetz. Es ist alles andere als offensichtlich, dass die beiden Eigenschaften – ob p ein Quadrat in $\mathbb{F}_q^\times$ ist einerseits, und ob q ein Quadrat in $\mathbb{F}_p^\times$ ist andererseits – etwas miteinander zu tun haben. Und doch kann man diese Eigenschaften in verblüffender Weise verbinden. Die Aussage wurde bereits von Euler und Legendre vermutet. Den ersten Beweis hat Gauß um 1800 gegeben; er hat später noch mehrere andere Beweise gefunden und diesen Satz sein »goldenes Theorem« (*theorema aureum*) genannt.

THEOREM 8.65 (Quadratisches Reziprozitätsgesetz). *Seien p und q verschiedene ungerade Primzahlen. Dann gilt*

$$\left(\frac{p}{q}\right) \left(\frac{q}{p}\right) = (-1)^{\frac{p-1}{2} \frac{q-1}{2}}.$$

Mit anderen Worten: Die Legendre-Symbole $\left(\frac{p}{q}\right)$ und $\left(\frac{q}{p}\right)$ sind gleich, es sei denn, sowohl p als auch q haben bei Division durch 4 den Rest 3, und dann sind sie verschieden.

Es gibt [viele Beweise](#)⁷ des quadratischen Reziprozitätsgesetzes. Wir skizzieren einen »trickreichen« Beweis nach G. Rousseau (J. Austral. Math. Soc. Ser. A **51** no. 3 (1991), 423–425.) Einen besonders durchsichtigen Beweis kann man mit den Methoden der Galois-Theorie geben, wie sie in der Algebra-Vorlesung erarbeitet werden.

BEWEIS. Wir betrachten die Gruppe

$$G = (\mathbb{Z}/p)^\times \times (\mathbb{Z}/q)^\times \cong (\mathbb{Z}/pq)^\times$$

(Satz 8.57).

Sei

$$P_1 = (\mathbb{Z}/p)^\times \times \{1, \dots, \frac{q-1}{2}\} \subset G,$$

$$P_2 = \{1, \dots, \frac{pq-1}{2}\} \subset (\mathbb{Z}/pq)^\times \cong G$$

In beiden Fällen liegt für x in G genau eines der Elemente $x, -x$ in P_i . (Wir benutzen hier, dass der Isomorphismus $G \cong (\mathbb{Z}/pq)^\times$ aus Satz 8.57 verträglich ist mit der Bildung des Negativen.)

Das Produkt über alle Elemente in P_1 und das Produkt über alle Elemente in P_2 unterscheiden sich also nur um das Vorzeichen.

Nun haben wir

$$\prod_{x \in P_1} x = \left(((p-1)!)^{\frac{q-1}{2}}, \left(\frac{q-1}{2}\right)^{p-1} \right) = \left(((p-1)!)^{\frac{q-1}{2}}, ((q-1)!)^{\frac{p-1}{2}} (-1)^{\frac{p-1}{2} \frac{q-1}{2}} \right)$$

Um das Produkt über alle Elemente in P_2 auszurechnen, müssen wir alle die Zahlen von 1 bis $\frac{pq-1}{2}$ multiplizieren, die zu p und q teilerfremd sind. Wir können das folgendermaßen schreiben: $\prod_{x \in P_2} x$ ist die Restklasse der ganzen Zahl

$$\frac{1 \cdots (p-1) \cdot (p+1) \cdots (2p-1) \cdots \left(\frac{q-3}{2}p+1\right) \cdots \left(\frac{q-1}{2}p-1\right) \left(\frac{q-1}{2}p+1\right) \cdots \left(\frac{q-1}{2}p+\frac{p-1}{2}\right)}{q \cdot 2q \cdots \frac{p-1}{2}q}$$

(im Zähler multiplizieren wir alle Zahlen von 1 bis $\frac{pq-1}{2}$ bis auf die Vielfachen von p , und teilen dann durch alle Vielfachen von q ; weil p und q verschiedene Primzahlen sind, gibt es keine Überschneidung). Die Restklasse dieses Elements in $\mathbb{Z}/p$ ist

$$\frac{((p-1)!)^{\frac{q-1}{2}}}{q^{\frac{p-1}{2}}} = ((p-1)!)^{\frac{q-1}{2}} \left(\frac{q}{p}\right),$$

wobei wir das Euler-Kriterium (Lemma 8.63) sowie die Tatsache, dass $\left(\frac{q}{p}\right) \in \{1, -1\}$ ist, benutzt haben. Streng genommen steht hier die Restklasse des Legendre-Symbols in $\mathbb{Z}/p$. Analog erhalten wir als die Restklasse von $\prod_{x \in P_2} x$ in $\mathbb{Z}/q$ das Element

$$\frac{((q-1)!)^{\frac{p-1}{2}}}{p^{\frac{q-1}{2}}} = ((q-1)!)^{\frac{p-1}{2}} \left(\frac{p}{q}\right).$$

In G können wir also schreiben

$$\prod_{x \in P_2} x = \left(((p-1)!)^{\frac{q-1}{2}} \left(\frac{q}{p}\right), ((q-1)!)^{\frac{p-1}{2}} \left(\frac{p}{q}\right) \right).$$

⁷<http://www.rzuser.uni-heidelberg.de/~hb3/fchrono.html>

Weil sich $\prod_{x \in P_1} x$ und $\prod_{x \in P_2} x$ höchstens um das Vorzeichen unterscheiden, erhalten wir

$$\left(((p-1)!)^{\frac{q-1}{2}}, ((q-1)!)^{\frac{p-1}{2}} (-1)^{\frac{p-1}{2} \frac{q-1}{2}} \right) = \varepsilon \left(((p-1)!)^{\frac{q-1}{2}} \left(\frac{q}{p} \right), ((q-1)!)^{\frac{p-1}{2}} \left(\frac{p}{q} \right) \right)$$

für $\varepsilon \in \{1, -1\}$. Wir multiplizieren diese Gleichung mit dem Inversen von $((p-1)!)^{\frac{q-1}{2}}, ((q-1)!)^{\frac{p-1}{2}}$ und erhalten das quadratische Reziprozitätsgesetz

$$\left(\frac{p}{q} \right) \left(\frac{q}{p} \right) = (-1)^{\frac{p-1}{2} \frac{q-1}{2}}.$$

□

Das quadratische Reziprozitätsgesetz wird vervollständigt durch die sogenannten Ergänzungssätze:

SATZ 8.66 (Ergänzungssätze zum quadratischen Reziprozitätsgesetz). *Sei p eine ungerade Primzahl.*

(1) *Es gilt*

$$\left(\frac{-1}{p} \right) = (-1)^{\frac{p-1}{2}} = \begin{cases} 1 & p \equiv 1 \pmod{4}, \\ -1 & p \equiv 3 \pmod{4}. \end{cases}$$

(2) *Es gilt*

$$\left(\frac{2}{p} \right) = (-1)^{\frac{p^2-1}{8}} = \begin{cases} 1 & p \equiv 1 \pmod{8} \text{ oder } p \equiv 7 \pmod{8}, \\ -1 & p \equiv 3 \pmod{8} \text{ oder } p \equiv 5 \pmod{8}. \end{cases}$$

BEWEIS. Teil (1) ist ein Spezialfall des Eulerschen Kriteriums. Der Beweis des zweiten Teils ist nicht besonders schwierig (jedenfalls wesentlich leichter als der Beweis des Reziprozitätsgesetzes selbst), aber wir lassen ihn an dieser Stelle aus. □

Mit dem Reziprozitätsgesetz, den Ergänzungssätzen und der Multiplikativität des Legendre-Symbols kann man leicht entscheiden, ob eine Zahl ein quadratischer Rest in $\mathbb{F}_p$ ist:

BEISPIEL 8.67. Um festzustellen, ob die Gleichung $X^2 - 137$ in $\mathbb{F}_{211}$ eine Lösung hat, rechnen wir

$$\begin{aligned} \left(\frac{137}{211} \right) &= \left(\frac{211}{137} \right) = \left(\frac{74}{137} \right) = \left(\frac{2}{137} \right) \left(\frac{37}{137} \right) = \left(\frac{2}{137} \right) \left(\frac{137}{37} \right) = \left(\frac{26}{37} \right) = \left(\frac{2}{37} \right) \left(\frac{13}{37} \right) \\ &= - \left(\frac{37}{13} \right) = - \left(\frac{11}{13} \right) = - \left(\frac{13}{11} \right) = - \left(\frac{2}{11} \right) = 1, \end{aligned}$$

die obige Gleichung ist also lösbar. ◇

Mehr Informationen auf Wikipedia: [deutsch](https://de.wikipedia.org/wiki/Quadratisches_Reziprozit%C3%A4tsgesetz)⁸, [englisch](https://en.wikipedia.org/wiki/Quadratic_reciprocity)⁹ (und wesentlich ausführlicher, auch zur Geschichte des quadratischen Reziprozitätsgesetzes).

Eine weitreichende Verallgemeinerung des quadratischen Reziprozitätsgesetzes ist das Artinsche Reziprozitätsgesetz, ein Kernstück der sogenannten [Klassenkörpertheorie](https://de.wikipedia.org/wiki/Klassenk%C3%B6rpertheorie)¹⁰. Die Suche nach Verallgemeinerungen der Klassenkörpertheorie bestimmt maßgeblich die heutige algebraische Zahlentheorie.

⁸https://de.wikipedia.org/wiki/Quadratisches_Reziprozit%C3%A4tsgesetz

⁹https://en.wikipedia.org/wiki/Quadratic_reciprocity

¹⁰<https://de.wikipedia.org/wiki/Klassenk%C3%B6rpertheorie>

Mehr Informationen zum quadratischen Reziprozitätsgesetz und zur »elementaren Zahlentheorie« finden Sie zum Beispiel im Buch

A. Schmidt, *Einführung in die algebraische Zahlentheorie*, Springer 2007,
<https://doi.org/10.1007/978-3-540-45974-3>

Dort finden Sie auch weitere Literaturhinweise.

□ Ergänzung 8.62

ERGÄNZUNG 8.68 (Die Fibonacci-Zahlen modulo p). Für einen »spielerischen« Abschluss des Abschnitts über Zahlentheorie kommen wir noch einmal auf die Fibonacci-Zahlen (Beispiel 5.60) zu sprechen und beweisen:

SATZ 8.69. Sei p eine Primzahl. Dann teilt p die Fibonacci-Zahl $F_{2p(p^2-1)}$.

BEWEIS. Für $p = 2$ kann man die Behauptung direkt verifizieren ($F_{12} = 144$ ist gerade). Wir setzen im folgenden voraus, dass p ungerade ist. Wir betrachten

$$G = \left\{ A \in GL_2(\mathbb{F}_p); \det(A) \in \{1, -1\} \right\},$$

eine Untergruppe von $GL_2(\mathbb{F}_p)$. Hier benutzen wir die Determinante einer (2×2) -Matrix,

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = ad - bc,$$

siehe Beispiel 5.56 oder Kapitel 9 für eine systematische Behandlung. Um zu sehen, dass G eine Untergruppe von $GL_2(\mathbb{F}_p)$ ist, muss man wissen, dass sich die Determinante multiplikativ verhält: $\det(AB) = \det(A)\det(B)$. Im (2×2) -Fall kann man das natürlich anhand der gegebenen Formel direkt nachrechnen.

Die Gruppe G hat $2p(p^2 - 1)$ Elemente: Die erste Spalte der Matrix kann irgendein Vektor aus $\mathbb{F}_p^2$ sein, der nicht Null ist. Die zweite Spalte muss linear unabhängig zur ersten sein, alle p Vielfache der ersten Spalte fallen damit schon einmal heraus. Von den verbleibenden $p^2 - p$ Elementen führen genau $\frac{p^2-p}{(p-1)/2} = 2p$ Elemente zu einer Determinante 1 oder -1 (denn wenn wir die zweite Spalte um ein Skalar abändern, ändert sich die Determinante um denselben Faktor).

Nun ist $A := \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \in G$, und aus Satz 8.55 folgt, dass $A^{2p(p^2-1)} = E_2$ in G gilt. Andererseits folgt aus den Überlegungen in Beispiel 5.60, dass

$$\begin{pmatrix} \bar{F}_{n+1} \\ \bar{F}_n \end{pmatrix} = A^n \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

ist, wobei $\bar{F}_n$ die Restklasse der Fibonacci-Zahl F_n in $\mathbb{F}_p$ bezeichnet. Es folgt $\bar{F}_{2p(p^2-1)} = 0$, und das ist genau die Behauptung. □

Quelle: <https://mathoverflow.net/a/53643>

□ Ergänzung 8.68

8.5.2. Darstellungstheorie. Unter einer *Darstellung* einer Gruppe G versteht man einen Gruppenhomomorphismus $\rho: G \rightarrow \text{Aut}(V)$, wo V ein Vektorraum (über einem Körper K , den man zu Beginn fixiert) ist. Ist V ein endlich-dimensionaler Vektorraum, so liefert die Wahl einer Basis von V einen Isomorphismus $\text{Aut}(V) \cong GL_n(K)$ (Beispiel 8.15). Wenn man sich auf endlich-dimensionale Vektorräume einschränkt, kann man also genauso gut Gruppenhomomorphismen $G \rightarrow GL_n(K)$ studieren.

Jedem Element $g \in G$ wird also ein Automorphismus $\rho(g): V \rightarrow V$ zugeordnet. Man sagt auch, dass die Gruppe G mittels dieser Automorphismen auf dem Vektorraum V *wirke*.

Ist die Abbildung $\rho: G \rightarrow \text{Aut}(V)$ injektiv, so kann man G mit einer Untergruppe von $\text{Aut}(V)$ identifizieren und die Gruppe G so sehr explizit realisieren. Manchmal kann man G dann sogar als die Symmetriegruppe einer Teilmenge (wie in Abschnitt 8.1.6) sehen.

Je nachdem, welcher Art die Gruppe G ist, sind unterschiedliche Darstellungen von Interesse. Bei Gruppen wie $G = GL_n(\mathbb{R})$ wird man in der Regel nur solche ρ betrachten, die stetig sind oder ähnliche Bedingungen erfüllen, die die »analytische« Natur der Gruppe G reflektieren. Bei endlichen Gruppen G ist es ein großer Unterschied, ob die Charakteristik des Grundkörpers K ein Teiler der Gruppenordnung $\#G$ ist (sofern die Charakteristik von K nicht 0 ist).

8.5.3. Verschiedenes.

BEISPIEL 8.70 (Das 15-Puzzle). Das »15-Puzzle^{II}« ist ein Schiebepuzzle mit 15 quadratischen Teilen, die in 4 Zeilen mit je 4 Teilen angeordnet sind; ein Feld bleibt dabei frei. Man kann dann die Position verändern, indem man eines der Teile, die zu dem freien Feld benachbart sind, dort hin verschiebt.



Während eine offensichtliche Aufgabe ist, die Ausgangsstellung wiederherzustellen, nachdem sie durch mehrere Züge der oben beschriebenen Art durcheinander gebracht wurde, ist eine interessantere Frage die folgende: Stellen wir uns die 15 Teile durchnummeriert von 1 bis 15 und aufsteigend angeordnet (1, 2, ...) vor, so dass das Feld rechts unten frei ist. Ist es möglich, durch eine Folge von Zügen dieselbe Position mit dem einzigen Unterschied, dass die Teile 14 und 15 vertauscht sind, zu erreichen?

Allgemeiner kann man fragen, ob alle Permutationen der Teile durch eine Folge der oben beschriebenen Züge erreichbar sind; oder – wenn nicht – welche erreichbar sind.

Das Puzzle fand ab ca. 1880 eine große Verbreitung in den USA; es wurde sogar ein Preisgeld für eine positive Antwort auf die oben gestellte Frage ausgesetzt. Allerdings zeigt die folgende Überlegung, dass das Problem nicht lösbar ist:

Wir definieren für jede mögliche Position P des Puzzles eine Zahl $i(P) \in \{1, -1\}$ wie folgt: Sei $\sigma \in S_{15}$ die Permutation, so dass (von oben nach unten, und von links nach rechts gelesen) die Zahlen 1, ..., 15 in der Reihenfolge $\sigma(1), \sigma(2), \dots, \sigma(15)$ auftreten. Hierbei wird das leere Feld übersprungen. Dann setzen wir

$$i(P) = \begin{cases} -\text{sgn}(\sigma) & \text{wenn sich das leere Feld in Reihe 1 oder Reihe 3 befindet,} \\ \text{sgn}(\sigma) & \text{wenn sich das leere Feld in Reihe 2 oder Reihe 4 befindet.} \end{cases}$$

^{II} <https://de.wikipedia.org/wiki/15-Puzzle>



ABBILDUNG 2. Lässt sich die Ausgangsstellung (rechts) aus der links dargestellten Position durch Verschieben der Steine herstellen?

Es ist dann leicht zu sehen, dass jeder mögliche Verschiebe-Zug die Zahl $i(P)$ nicht ändert: In der Tat, verschiebt man innerhalb einer Zeile, dann ändern sich weder σ noch der Zeilenindex des leeren Feldes. Verschiebt man innerhalb einer Spalte, dann ändert sich der Zeilenindex des leeren Feldes, und σ ändert sich um eine Transposition.

Da für die Position P , wo alle Teile 1 bis 15 in aufsteigender Reihenfolge angeordnet sind, $i(P) = 1$ gilt, aber für die Position $i(P')$, wo im Vergleich dazu nur 14 und 15 vertauscht sind, $i(P') = -1$ ist, kann es keine Zugfolge geben, die diese beiden Positionen ineinander überführt.

Man kann auch zeigen, dass man zwei Positionen P_1 und P_2 genau dann ineinander überführen kann, wenn $i(P_1) = i(P_2)$ gilt. Das kann man auch mit ganz elementaren Überlegungen zeigen, lässt sich aber nicht in einem Absatz abhandeln. Allerdings ist auch diese Tatsache schon lange bekannt, siehe [W. Johnson, W. Story, [Notes on the "15" Puzzle](#)¹², Amer. J. Math. **2**, no. 4 (1879), 397–404]. Siehe auch [Jo] 7.4. $\diamond$

ERGÄNZUNG 8.71 (Rubiks Zauberwürfel). Die Theorie der Gruppen ist nützlich, um [Rubiks Zauberwürfel](#)¹³ mathematisch zu analysieren. Dazu stellen wir uns den Würfel selbst im Raum fixiert vor (d.h. dass die sechs Mittelplättchen des Würfels sich nicht bewegen, sondern höchstens »mitdrehen«; in den Bildern hier zeigt das grüne Mittelplättchen immer nach vorne, das gelbe nach oben, das rote nach links, usw.). Wir nummerieren alle anderen Plättchen durch (von 1 bis 48). Jede Verdrehung, die man mit dem Würfel dann machen kann, entspricht einer Permutation der Zahlen von 1 bis 48.

Wir verwenden die übliche Notation für die möglichen »Züge«, die man ausführen kann:

- (1) **F** – Drehung der vorderen Ebene (front) um 90° im Uhrzeigersinn,
- (2) **B** – Drehung der hinteren Ebene (back) um 90° im Uhrzeigersinn,
- (3) **R** – Drehung der rechten Ebene (right) um 90° im Uhrzeigersinn,
- (4) **L** – Drehung der linken Ebene (left) um 90° im Uhrzeigersinn,
- (5) **U** – Drehung der oberen Ebene (up) um 90° im Uhrzeigersinn,

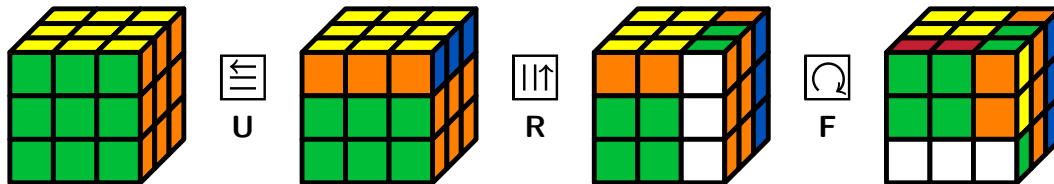
¹²<https://doi.org/10.2307/2369492>

¹³<https://de.wikipedia.org/wiki/Zauberw%C3%BCrfel>

(6) **D** – Drehung der unteren Ebene (down) um 90° im Uhrzeigersinn,

Es hat sich eingebürgert, statt F^{-1} kürzer F' zu schreiben, und wir folgen dieser Konvention. Also ist F' die Drehung der vorderen Ebene um 90° gegen den Uhrzeigersinn. Wie in jeder multiplikativ geschriebenen Gruppe steht F^2 für FF .

Wenn mehrere Züge nacheinander ausgeführt werden, lesen wir die Zugfolge $\mathbf{F R'}$ von links nach rechts: Drehe erst die vordere Ebene einmal im Uhrzeigersinn, dann die rechte Ebene einmal gegen den Uhrzeigersinn. Die Folge $\mathbf{U R F}$ ergibt also:



Weil wir bei Permutationen die entgegengesetzte Konvention verwenden – $\sigma\tau$ bedeutet, erst die Permutation τ anzuwenden, und dann σ , muss man ein bisschen aufpassen, welche Permutationen in S_{48} zu den oben beschriebenen Elementen F , B , R , ... gehören sollen. Wenn die Drehung F das Plättchen i des Würfels auf die Position j bewegt, dann soll für F als Permutation betrachtet gelten, dass $F(j) = i$ ist; das ist gerade das Inverse der Permutation, die auf den Plättchen induziert wird. Diese Definition führt dazu, dass die Konvention, Züge auf den Würfel von links nach rechts anzuwenden, kompatibel ist mit der Multiplikation von Permutationen in der symmetrischen Gruppe.

Wir können damit die Zauberwürfelgruppe G als die Untergruppe von $\subset S_{48}$ definieren, die von den Permutationen erzeugt wird, die zu den oben aufgelisteten Zauberwürfeldrehungen gehören.

Natürlich lassen sich bei weitem nicht alle Permutationen aus der S_{48} durch Verdrehungen des Zauberwürfels realisieren. Man kann zeigen, dass die Anzahl der Elemente von G gleich

$$43\,252\,003\,274\,489\,856\,000$$

ist, also ungefähr $43 \cdot 10^{36}$, 43 Trillionen. Dies ist also die Anzahl möglicher Positionen des Zauberwürfels. Das ist eine ganze Menge, aber viel weniger als $\#S_{48} = 48! \approx 12 \cdot 10^{60}$, in Ziffern ausgeschrieben:

$$12\,413\,915\,592\,536\,072\,670\,862\,289\,047\,373\,375\,038\,521\,486\,354\,677\,760\,000\,000\,000.$$

Den Würfel lösen zu können, bedeutet sozusagen, ein Erzeugendensystem der Gruppe G zu kennen, das kontrolliert einzelne Steine austauscht/bewegt, so dass man sich schrittweise von einer beliebigen Ausgangsposition zum Ursprungszustand des Würfels vorarbeiten kann. Je nachdem, wie viele Zugfolgen sich zu merken man bereit ist, kann man die Sache ziemlich schnell durchführen: Der Weltrekord liegt bei unter 4 Sekunden (für das Drehen des Würfels; eine (kurze) Zeit, in der sich der Löser den verdrehten Würfel vorher anschauen durfte, nicht mitgerechnet).

Es ist offensichtlich, dass die Gruppe G nicht kommutativ, also insbesondere nicht zyklisch ist. Das bedeutet, dass es »leider« nicht ein Gruppenelement, also eine einzige Zugfolge gibt, die aus jeder Position, geeignet oft angewandt, die Ausgangsstellung wiederherstellen würde.

Hier ein einfaches »Kochrezept«, um den Zauberwürfel aus jeder Position wieder in die Ausgangsstellung zu »drehen«. Im wesentlichen handelt es sich um die Lösung die 1981 von der Zeitschrift DER SPIEGEL veröffentlicht wurde (Nr. 4/1981, S. 183/184¹⁴). Siehe auch

¹⁴ <http://magazin.spiegel.de/EpubDelivery/spiegel/pdf/14319775>

cube3x3.com¹⁵. Es gibt natürlich viele Alternativen. Wenn es schneller gehen soll: Siehe die Links zu Speed-cubing-Verfahren weiter unten.

Wir verwenden die folgende Notation: ${}^gh := ghg^{-1}$ ist die Konjugation von h mit g . Mit $[g, h] := ghg^{-1}h^{-1}$ bezeichnen wir den sogenannten Kommutator der Elemente g, h , der sozusagen misst, ob g und h miteinander kommutieren. Zum Beispiel ist also ${}^F[R, U]$ eine Kurzschreibweise für $FRUR'U'F'$, und ${}^{L'}UR'$ eine andere Art, das Element $L'UR'U'L$ zu schreiben. Mit diesen Bezeichnungen lassen sich alle Zugfolgen, die in dem folgenden Lösungsrezept benutzt werden, ziemlich kurz und strukturiert angeben. (Vergleiche zum Beispiel Lemma 8.37.)

Aus mathematischer Sicht interessanter als das Lösungsrezept an sich sind strukturelle Eigenschaften der Gruppe G . Dafür führen wir jetzt noch einige Beispiele an. Lemma 8.54 besagt in diesem Kontext, dass jede Zugfolge, wenn man sie oft genug wiederholt, wieder zur Ausgangsposition führt. Das kann eine ganze Weile dauern: Die maximale Ordnung eines Elements in G ist 1260, ein Beispiel für ein Element dieser Ordnung ist $RU^2D'BD'$.

Erst im Jahr 2014 wurde gezeigt, dass es in jeder Position des Zauberwürfels möglich ist, die Ausgangsstellung mit 26 oder weniger einzelnen Drehungen (also $F, F', R, \dots$) wiederherzustellen. Wenn man auch Halbdrehungen (d.h. F^2 , usw.) zulässt, dann genügen 20 Züge. Diese Zahl wird von »Cubern« auch als *God's number* bezeichnet.

Wenn wir die »Orientierung« der einzelnen Würfelteile vernachlässigen und nur die Position der Teile betrachten, haben wir es nur mit den 8 Eckteilen und 12 Kanten Teilen zu tun. Wir nummerieren die 20 Positionen dieser Teile durch und können dann jedem Würfelzug eine Permutation dieser 20 Positionen zuordnen. Wir erhalten damit einen Gruppenhomomorphismus $G \rightarrow S_{20}$ (wobei wir, ähnlich wie zu Beginn, die Elemente von G auf die inverse Permutation abbilden).

Das Bild dieses Homomorphismus liegt in der Untergruppe A_{20} . Diese ist definiert als der Kern des Signum-Homomorphismus $\text{sgn}: S_{20} \rightarrow \{1, -1\}$ und heißt die *alternierende Gruppe*.

Um zu zeigen, dass das Bild des Homomorphismus $G \rightarrow S_{20}$ in A_{20} liegt, genügt es zu zeigen, dass die Erzeuger F, B, L, R, U, B auf Elemente in A_{20} abgebildet werden. Die Rechnung dafür ist für alle diese Elemente im Prinzip dieselbe; das Bild in S_{20} ist jeweils ein Produkt von 2 Zykeln der Ordnung 4 (4 Ecksteine und 4 Kantensteine werden jeweils zyklisch vertauscht), und dieses hat Signum 1.

Da alle Transpositionen Signum -1 haben, enthält A_{20} keine Transpositionen. Das genannte Ergebnis impliziert, dass es keine Zugfolge gibt, die zwei Würfelteile vertauscht, aber alle anderen Würfelteile am selben Platz belässt. Die in diesem Sinne einfachsten Zugfolgen sind also Zugfolgen, die alle bis auf drei Würfelteile an ihrem Platz lassen; ein Beispiel dafür ist das Element ${}^UR'L'UR'$ von G , das in Schritt (3c) des Lösungsrezepts vorkommt und (nur) drei der Ecksteine der oberen Ebene bewegt. (Die Zugfolgen in den Schritten (2) und (3b) vertauschen zwei Kantensteine, bewegen aber auch noch andere Steine.)

Der erwähnte Homomorphismus $G \rightarrow S_{20}$ ist nicht injektiv: Es gibt Zugfolgen, die den Würfel verändern, aber alle Steine an ihrem Platz lassen. Es wird also nur bei einigen Steinen die Orientierung verändert. Man kann auch zeigen, dass es nicht möglich ist, alle Steine an ihrem Platz zu lassen und nur bei einem einzigen die Orientierung zu verändern.

Man kann die Struktur der Zauberwürfelgruppe G ganz explizit beschreiben, sie ist isomorph zur Gruppe

$$((\mathbb{Z}/3)^7 \times (\mathbb{Z}/2)^{11}) \rtimes ((A_8 \times A_{12}) \rtimes \mathbb{Z}/2),$$

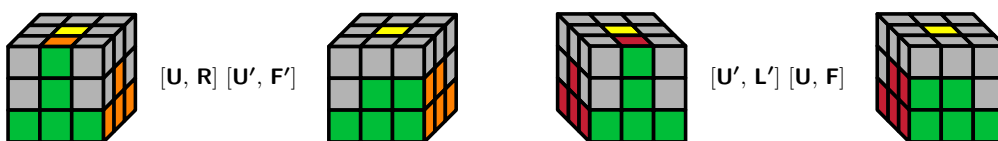
wobei $\times$ das Produkt von Gruppen bezeichnet (Beispiel 8.5) und $\rtimes$ das halbdirekte (oder semi-direkte) Produkt – eine dem Produkt ähnliche Konstruktion, die wir hier nicht näher

¹⁵<https://cube3x3.com/wie-man-einen-zauberwurfel-rubiks-cube-lost-de/>

Lösung des Zauberwürfels

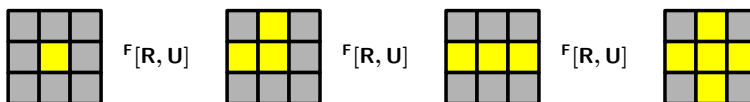
Mit den hier gezeigten Sequenzen kann man den Würfel lösen. Oft müssen die Zugfolgen aber mehrfach angewendet werden, und man muss den Würfel zwischendurch auch mal drehen (zum Beispiel in Schritt (2), aber auch in Schritt (3)). Die grau gefärbten Quadrate sind für den jeweiligen Schritt irrelevant (und ihre Farbe kann sich durch die angegebene Zugfolge ändern.)

- (1) Man beginnt damit, eine Ebene (im Bild: die weiße Ebene) in Ordnung zu bringen. Weil man nicht darauf achten muss, was mit den anderen Ebenen passiert, ist das nicht so schwierig. Für das weitere Vorgehen nehmen wir an, dass der Würfel so gedreht wird, dass die **untere Ebene** bereits gelöst ist, und zwar so, dass alle Kanten- und Eckenteile richtig orientiert sind (also zu den Mittelsteinen der angrenzenden Seiten passen).
- (2) Um die **mittlere Ebene** fertigzustellen müssen gegebenenfalls Kantenteile aus der oberen Ebene in die mittlere Ebene gebracht werden:

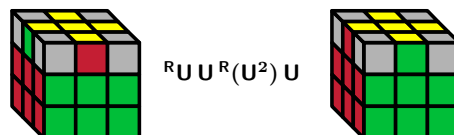


- (3) Die **obere Ebene** in Ordnung zu bringen, ist am aufwändigsten, weil die anderen Ebenen dabei nicht zerstört werden sollen.

(a) Das erste Ziel ist, in der oberen Ebene das Kreuz aus Mittel- und Kantensteinen zu erhalten. Dazu kann man die Zugfolge $F[R, U]$ – gegebenenfalls mehrfach – verwenden. Ob die Kantensteine zu den Mittelsteinen der anderen Seiten passen, ist erstmal egal.



(b) Falls nötig, werden die Kantensteine der oberen Ebene vertauscht. Diese Folge vertauscht die Kanten vorne und links.



(c) Jetzt werden die Ecksteine der oberen Ebene in die richtige Position gebracht. Die Orientierung ist in diesem Schritt egal.

Die angegebene Zugfolge vertauscht drei Ecksteine der oberen Ebene gegen den Uhrzeigersinn, der Stein vorne rechts bleibt in seiner Position. Je nach Lage der Ecksteine führt man die Folge mehrfach aus, oder benutzt das inverse Element.

(d) Zum Schluss muss noch die richtige Orientierung der Ecken der oberen Ebene hergestellt werden. Die Zugfolge $[R', D']^2$ dreht die Ecke oben-vorne-rechts im Uhrzeigersinn – wende diese Folge ein- oder zweimal an, so dass die Ecke vorne rechts in der gewünschten Position ist. Danach führe man U so oft aus, bis die nächste Ecke der oberen Ebene, die noch in sich gedreht werden muss, vorne rechts ist, und mache entsprechend weiter.

Achtung: Zwischendurch werden in Schritt (d) die unteren Ebenen durcheinandergebracht, aber wenn alle Ecken der oberen Ebene gelb oben haben, ist wieder alles in Ordnung! (Deshalb ist es wichtig, jeweils mit U die nächste Ecke »auszuwählen«, die gedreht werden muss, und nicht den ganzen Würfel zu drehen.)

besprechen wollen. Die Gruppen A_8 und A_{12} sind die alternierenden Gruppen, also der Kern der Signum-Abbildung auf S_8 bzw. auf S_{12} .

Literaturverweise und Links zur mathematischen Seite des Zauberwürfels.

[Webseite von D. Bump \(Stanford\) zur Mathematik des Zauberwürfels^a](#)

C. Bandelow, *Inside Rubik's cube and beyond*, Birkhäuser 1982.

A. Frey, D. Singmaster, *Handbook of cubic math*, Enslow Publ. 1982.

D. Singmaster, *Notes on Rubik's magic cube*, Enslow Publ. 1981.

Weitere Links bei [J. Burke^b](#), [Lehrstuhl Geometrie, Uni Rostock^c](#).

Das Buch [Jo] von Joyner enthält Material zum Zauberwürfel und zu anderen Spiel(zeug)en, die sich gruppentheoretisch analysieren lassen.

Die [Fridrich-Methode^d](#), einer der populärsten Speed-Cubing-Algorithmen, entwickelt von Jessica Fridrich.

Die Zauberwürfelgruppe wurde im Computeralgebra-Programm SAGE implementiert. Dort kann man viele der obigen Behauptungen direkt am Rechner überprüfen: [Zauberwürfelgruppe in SAGE^e](#)

^a<http://sporadic.stanford.edu/bump/match/rubik.html>

^b<https://maths-people.anu.edu.au/burkej/cube/cube.html>

^c<https://www.mathematik.uni-rostock.de/struktur/professuren-apl-prof/geometrie/rubik/>

^d<http://www.ws.binghamton.edu/fridrich/cube.html>

^ehttps://doc.sagemath.org/html/en/reference/groups/sage/groups/perm_gps/cubegroup.html

□ Ergänzung 8.71

8.5.4. Gruppentheorie in anderen Disziplinen. Gruppentheorie spielt auch in vielen anderen Wissenschaften, besonders natürlich in den Naturwissenschaften, eine wichtige Rolle. Wir belassen es hier bei einigen Andeutungen und Literaturhinweisen:

Chemie. Symmetrieüberlegungen sind in mehreren Bereichen der Chemie essenziell, beispielsweise bei der Untersuchung von Molekülstrukturen, oder von Kristallen.

Literatur:

G. James, M. Liebeck, *Representations and Characters of Groups*, Cambridge Univ. Press 2005. Kapitel 32 (nach einer Menge Mathematik zur Vorbereitung) hat den Titel *An application of representation theory to molecular vibration*.

J. P. Serre, *Linear representations of finite groups*, Springer 1977. Dieser Klassiker (ursprünglich auf französisch, und es gab auch eine deutsche Übersetzung) ist aus einer Vorlesung entstanden, die für Chemiker*innen gehalten wurde (allerdings wird die Verbindung zur Chemie im Buch selbst kaum thematisiert).

Physik. Aus der modernen Physik ist die Gruppentheorie nicht wegzudenken. Auch Matrixgruppen wie die $GL_n(\mathbb{R})$ oder $GL_n(\mathbb{C})$, und dazu verwandte Gruppen wie die orthogonale Gruppe (Definition 11.25) und »unitäre Gruppen«, die wir in der Linearen Algebra 2 genauer kennenlernen werden, spielen eine wichtige Rolle.

Anthropologie. Zum Abschluss noch ein etwas kurioses Beispiel, Gruppentheorie in der Anthropologie: <https://mathoverflow.net/a/25743> und wesentlich ausführlicher: <https://ncatlab.org/nlab/show/kinship>. (Es geht um die Erklärung/Formalisierung von Systemen von Heiratsregeln und resultierenden Verwandtschaftsbeziehungen.)

Die Determinante

Die *Determinante* einer Matrix bzw. eines Endomorphismus kann man unter mehreren Aspekten sehen. Ursprünglich wurde sie entwickelt als Ausdruck, der bestimmt, ob ein lineares Gleichungssystem eindeutig lösbar ist, und später, um eine Art »Lösungsformel« für solche Gleichungssysteme zu haben, vergleiche Abschnitt 2.5 und Beispiel 5.56. Mit der Cramerschen Regel Satz 9.32 werden wir dieses Ziel erreichen. Allerdings direkt an dieser Stelle die Bemerkung: Zur praktischen Berechnung der Lösung eines linearen Gleichungssystems ist diese Methode viel zu aufwändig; der Gauß-Algorithmus ist ihr *immer* überlegen. Für theoretische Überlegungen (wie Korollar 9.33) hat sie aber ihren Nutzen.

Von ihrem Ursprung abgesehen hat die Determinante in der heutigen linearen Algebra eine wesentlich größere, auch theoretische Bedeutung, die wir nach und nach kennenlernen werden. An dieser Stelle erwähnen wir:

- Die Determinante $\det(A)$ einer quadratischen Matrix $A \in M_n(K)$ ist ein Element von K , das in Termen der Koeffizienten von A angegeben/ausgerechnet werden kann. Es gilt genau dann $\det(A) \neq 0$, wenn A invertierbar ist. Die Determinante liefert also eine sowohl praktisch wie auch theoretisch nützliche Art und Weise zu entscheiden, ob eine Matrix invertierbar ist.
- Für eine quadratische Matrix $A \in M_n(\mathbb{R})$ über den reellen Zahlen hat die Determinante von A die folgende geometrische Interpretation: Ist $M \subseteq \mathbb{R}^n$ eine Teilmenge mit Volumen m (im Sinne von Abschnitt 11.6.2 oder Abschnitt 11.6.3), dann hat ihr Bild $f_A(M)$ unter der linearen Abbildung $f_A: \mathbb{R}^n \rightarrow \mathbb{R}^n$ das Volumen $|\det(A)| \cdot m$. Siehe Satz 11.66. Insbesondere ist $|\det(A)|$ das Volumen des »Parallelotops«, das das Bild des Einheitswürfels unter f_A ist.

Video von [3blue1brown](https://www.youtube.com/watch?v=Ip3X9L0h2dk)^a zur Determinante.

^a<https://www.youtube.com/watch?v=Ip3X9L0h2dk>

9.1. Determinantenfunktionen

Seien K ein Körper, V ein K -Vektorraum, $n \geq 0$. Wir beginnen damit, einige Eigenschaften festzuhalten, die die Abbildung $A \mapsto \det(A)$ sinnvollerweise haben sollte. Auch wenn wir nicht jede dieser Eigenschaften einzeln motivieren, passen sie alle zu Operationen, die wir schon vorher betrachtet haben (speziell: elementare Spaltenumformungen). Wie wir später sehen werden, passen sie auch gut zu den in der Einleitung des Kapitels beschriebenen Eigenschaften, die wir letztlich für die Determinante einer Matrix haben möchten.

Danach werden wir uns als erstes überlegen, dass es bis auf Skalierung höchstens eine Determinantenfunktion geben kann. Diese Eindeutigkeitsaussage wird uns später mehrfach von Nutzen sein. Die Existenz beweisen wir dann, indem wir mit der Leibniz-Formel konkret eine Determinantenfunktion angeben.

Wir betrachten den Vektorraum $V^n = V \times \cdots \times V$ (n Kopien) mit der komponentenweisen Addition und Skalarmultiplikation. Im Fall $V = K^n$, der für uns der wichtigste Fall sein

wird, betrachten wir Elemente von K^n wie üblich als Spaltenvektoren und identifizieren $(K^n)^n$, also die Menge aller n -Tupel von solchen Spaltenvektoren, mit dem Raum $M_{n \times n}(K)$ der $(n \times n)$ -Matrizen über K , indem wir die n Spaltenvektoren nebeneinander in eine Matrix schreiben.

9.1.1. Eindeutigkeit und Existenz der Determinante.

DEFINITION 9.1. Seien $n \in \mathbb{N}$ und seien V und W Vektorräume über K .

- (1) Eine Abbildung $\delta: V^n \rightarrow W$ heißt *multilinear*, falls für alle $i \in \{1, \dots, n\}$ und alle $v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n \in V$ die Abbildung

$$V \rightarrow W, \quad v \mapsto \delta(v_1, \dots, v_{i-1}, v, v_{i+1}, \dots, v_n)$$

eine lineare Abbildung ist.

- (2) Eine multilineare Abbildung $V^n \rightarrow W$ heißt *alternierend*, falls für alle $v_1, \dots, v_n \in V$, so dass $i \neq j$ existieren mit $v_i = v_j$, gilt:

$$\delta(v_1, \dots, v_n) = 0.$$

—

Die Menge aller multilinearen Abbildungen sowie die Menge aller alternierenden multilinearen Abbildungen $V^n \rightarrow W$ sind K -Vektorräume (mit der üblichen Addition und Skalarmultiplikation für Abbildungen, d.h. diese Mengen sind Untervektorräume von $\text{Abb}(V^n, W)$).

Für $n > 1$ sind multilineare Abbildungen $V^n \rightarrow W$ abgesehen von der Nullabbildung nicht linear!

Die Bezeichnung *alternierend* weist darauf hin, dass multilineare Abbildungen mit dieser Eigenschaft ihr Vorzeichen ändern, wenn man zwei ihrer Argumente vertauscht. Genauer gilt das folgende Lemma:

LEMMA 9.2 (Charakterisierung alternierender Abbildungen). *Wir verwenden dieselben Bezeichnungen wie oben. Für die folgenden Eigenschaften*

- (i) δ ist alternierend,
- (ii) für alle linear abhängigen Familien $v_1, \dots, v_n \in V$ gilt: $\delta(v_1, \dots, v_n) = 0$,
- (iii) bei Vertauschung zweier Einträge ändert sich der Wert von δ um das Vorzeichen, d.h. für alle $i \neq j$ und alle $v_1, \dots, v_n \in V$ gilt

$$\delta(v_1, \dots, v_n) = -\delta(v_1, \dots, v_j, \dots, v_i, \dots, v_n).$$

- (iv) für alle $\sigma \in S_n$, $v_1, \dots, v_n \in V$ gilt

$$\delta(v_{\sigma(1)}, \dots, v_{\sigma(n)}) = \text{sgn}(\sigma) \delta(v_1, \dots, v_n).$$

gilt (i) $\Leftrightarrow$ (ii) $\Rightarrow$ (iii) $\Leftrightarrow$ (iv). Wenn $1 + 1 \neq 0$ in K ist, so sind alle vier Eigenschaften äquivalent.

BEWEIS. Es gelte (i), und es sei $v_1, \dots, v_n \in V$ eine linear abhängige Familie von Vektoren in V . Wir wollen zeigen, dass $\delta(v_1, \dots, v_n) = 0$. Sei dazu $\sum_{i=1}^n a_i v_i = 0$ eine nicht-triviale Darstellung des Nullvektors als Linearkombination der v_i . Um die Notation zu vereinfachen, nehmen wir an, dass $a_1 \neq 0$ ist. In den anderen Fällen argumentiert man analog. Dann gilt

$$\begin{aligned} a_1 \delta(v_1, v_2, \dots, v_n) &= \\ a_1 \delta(v_1, v_2, \dots, v_n) + a_2 \delta(v_2, v_2, \dots, v_n) + \dots + a_n \delta(v_n, v_2, \dots, v_n) &= \\ \delta\left(\sum_i a_i v_i, v_2, \dots, v_n\right) &= \delta(0, v_2, \dots, v_n) = 0, \end{aligned}$$

wobei wir im ersten Schritt Eigenschaft (i) und im zweiten und letzten Schritt die Linearität in der ersten Komponente ausgenutzt haben. Es folgt $\delta(v_1, v_2, \dots, v_n) = 0$.

Dass umgekehrt (i) aus (ii) folgt, ist klar.

Aus (iv) folgt Eigenschaft (iii), weil die Transposition (ij) das Signum -1 hat. Weil man jede Permutation als Produkt von Transpositionen schreiben kann (Lemma 8.41), folgt auch (iv) aus (iii).

Schließlich zeigen wir noch, dass $(iii) \Rightarrow (i)$ gilt, wenn im Körper K das Element $1 + 1$ von 0 verschieden ist. Sind $v_1, \dots, v_n$ und $i \neq j$ mit $v_i = v_j$ gegeben, so haben wir wegen (iii):

$$\delta(v_1, \dots, v_n) = -\delta(v_1, \dots, v_n),$$

weil wegen $v_i = v_j$ die Vertauschung des i -ten und j -ten Eintrags das gegebene Tupel von Vektoren nicht ändert. Es folgt $(1 + 1) \cdot \delta(v_1, \dots, v_n) = 0$, mit unserer Voraussetzung an K also $\delta(v_1, \dots, v_n) = 0$. $\square$

Als unmittelbare Konsequenz aus der Definition erhalten wir

LEMMA 9.3. Sei $\delta: V^n \rightarrow W$ multilinear und alternierend. Dann gilt

$$\delta(v_1, \dots, v_{i-1}, v_i + av_j, v_{i+1}, \dots, v_n) = \delta(v_1, \dots, v_n)$$

für alle $v_1, \dots, v_n \in V, a \in K$ und $i \neq j$.

BEWEIS. Es gilt

$$\delta(v_1, \dots, v_i + av_j, \dots, v_n) = \delta(v_1, \dots, v_n) + a\delta(v_1, \dots, v_{i-1}, v_j, v_{i+1}, \dots, v_n) = \delta(v_1, \dots, v_n)$$

wegen der Linearität im i -ten Eintrag, und weil δ alternierend und nach Voraussetzung $i \neq j$ ist. $\square$

Auch wenn diese Feststellung einfach ist, ist sie sehr wichtig, und wir halten den Spezialfall $V = K^n$ noch einmal gesondert fest:

KOROLLAR 9.4. Sei $V = K^n$ und $\delta: (K^n)^n = M_n(K) \rightarrow W$ multilinear und alternierend. Seien $A, A' \in M_n(K)$.

- (1) Geht A' aus A durch eine elementare Spaltenumformung vom Typ I hervor (Addieren des Vielfachen einer Spalte zu einer anderen Spalte), so gilt $\delta(A') = \delta(A)$.
- (2) Geht A' aus A durch eine elementare Spaltenumformung vom Typ II hervor (Vertauschen zweier Spalten), so gilt $\delta(A') = -\delta(A)$.
- (3) Geht A' aus A durch eine elementare Spaltenumformung vom Typ III hervor (Multiplikation einer Spalte mit $a \in K^\times$), so gilt $\delta(A') = a\delta(A)$.

BEWEIS. Teil (1) folgt direkt aus dem vorherigen Lemma, Teil (2) folgt aus Lemma 9.2, und Teil (3) folgt aus der Multilinearität. $\square$

Wir schränken und nun auf den Fall ein, dass $n = \dim(V)$ und $W = K$ ist. Der wichtigste Fall ist, wie schon erwähnt, der Fall des Standardvektorraums $V = K^n$, wo wir V als den Matrizenraum $M_{n \times n}(K)$ betrachten.

DEFINITION 9.5. Seien K ein Körper und V ein endlich-dimensionaler K -Vektorraum. Sei $n = \dim V$. Eine Determinantenfunktion auf V ist eine alternierende multilineare Abbildung $V^n \rightarrow K$.

Eine Determinantenfunktion Δ heißt *nicht-trivial*, wenn Δ nicht die Nullabbildung ist.

Wir bezeichnen den Vektorraum aller Determinantenfunktionen auf V mit $\mathcal{D}_V$. $\dashv$

Die zu Beginn erwähnte Eindeutigkeitsaussage können wir nun formulieren als

SATZ 9.6. *Sei V ein endlich-dimensionaler K -Vektorraum und $\mathcal{D}_V$ der Vektorraum der Determinantenfunktionen $V^n \rightarrow K$. Dann gilt $\dim \mathcal{D}_V \leq 1$.*

Mit anderen Worten: Ist Δ eine nicht-triviale und Δ' irgendeine Determinantenfunktion $V^n \rightarrow K$, so existiert ein eindeutig bestimmtes Element $a \in K$ mit $\Delta' = a\Delta$. (Der Satz lässt auch die Möglichkeit zu, dass es gar keine nicht-triviale Determinantenfunktion gibt; wir werden aber sehen, dass $\dim \mathcal{D}_V = 1$ gilt.)

BEWEIS. Wir können ohne Einschränkung den Fall $V = K^n$ betrachten. Denn im allgemeinen Fall liefert die Wahl einer Basis von V einen Koordinatenisomorphismus $c: V \xrightarrow{\sim} K^n$. Dieser induziert einen Isomorphismus

$$\mathcal{D}_{K^n} \longrightarrow \mathcal{D}_V, \quad \Delta \mapsto ((v_1, \dots, v_n) \mapsto \Delta(c(v_1), \dots, c(v_n))),$$

von K -Vektorräumen.

Wir betrachten nun die Abbildung $\varepsilon: \mathcal{D}_V \rightarrow K, \Delta \mapsto \Delta(E_n)$. Es genügt dann zu zeigen, dass ε injektiv ist. Die Abbildung ε ist ein Homomorphismus von K -Vektorräumen, also müssen wir nur zeigen, dass ε trivialen Kern hat. Sei dazu $\Delta \in \text{Ker}(\varepsilon)$, d.h. Δ ist eine Determinantenfunktion mit $\Delta(E_n) = 0$. Wir wollen zeigen, dass $\Delta(A) = 0$ für alle $A \in M_{n \times n}(K)$ gilt.

Ist A nicht invertierbar, so folgt das aus Lemma 9.2.

Sei nun A invertierbar: Wir können dann A schreiben als Produkt einer Diagonalmatrix D der Form $D = \text{diag}(1, \dots, 1, d')$ und einer Matrix $C \in SL_n(K)$, siehe Satz 8.34. Die Matrix DC entsteht aus D durch elementare Spaltenumformungen (durch Multiplikation von *rechts* erhalten wir die Spaltenversion von Bemerkung 5.37), daher folgt aus Lemma 9.3, dass

$$\Delta(A) = \Delta(DC) = \Delta(D) = d' \Delta(E_n),$$

wobei wir im letzten Schritt ausgenutzt haben, dass Δ im letzten Eintrag linear ist. Nach Voraussetzung gilt $\Delta(E_n) = 0$, also auch $\Delta(A) = 0$, wie gewünscht. $\square$

SATZ 9.7 (Leibniz-Formel). *Die Abbildung $\det: M_{n \times n}(K) \rightarrow K$, die gegeben ist durch*

$$A = (a_{ij})_{i,j} \mapsto \sum_{\sigma \in S_n} \text{sgn}(\sigma) a_{1,\sigma(1)} \cdots a_{n,\sigma(n)}$$

ist eine nichttriviale Determinantenfunktion auf K^n (vermöge der oben angegebenen Identifikation $(K^n)^n = M_{n \times n}(K)$). Insbesondere gilt $\dim \mathcal{D}_V = 1$ für jeden endlich-dimensionalen K -Vektorraum V .

Wir nennen $\det(A) \in K$ die Determinante der Matrix A .

BEWEIS. Der »Insbesondere ...«-Teil folgt, weil eine Basiswahl von V einen Isomorphismus $\mathcal{D}_V \cong \mathcal{D}_{K^n}$ induziert, vergleiche den Beginn des Beweises von Satz 9.6.

Wir prüfen nach, dass $\det$ die angegebenen Eigenschaften hat. Jedenfalls gilt $\det(E_n) = 1$, da in diesem Fall nur die Permutation id zu der angegebenen Summe beiträgt. Also ist $\det$ nicht-trivial.

Um die Multilinearität zu zeigen, erinnern wir an die Bemerkung, dass die Summe multilinearer Funktionen $V^n \rightarrow K$ wieder multilinear ist. Es genügt also einzusehen, dass alle Abbildungen

$$A = (a_{ij})_{i,j} \mapsto \text{sgn}(\sigma) a_{1,\sigma(1)} \cdots a_{n,\sigma(n)}$$

für eine fixierte Permutation σ multilinear sind, und das folgt direkt aus dem Distributivgesetz in K .

Es bleibt zu zeigen, dass die Abbildung $\det$ alternierend ist. (Zwar ist die Summe alternierender Funktionen wieder alternierend, allerdings sind die einzelnen Summanden keine

alternierenden Funktionen, sondern nur die ganze Summe; wir müssen daher alle Summanden zusammen betrachten.) Sei also eine Matrix $A = (a_{ij})_{i,j}$ gegeben, in der die Spalten j und j' (für $j \neq j'$) übereinstimmen. Wir wollen zeigen, dass $\det(A) = 0$ ist.

Dazu betrachten wir die Transposition $\tau = (jj') \in S_n$. Sei $A_n \subset S_n$ die Teilmenge aller Permutationen, die Signum 1 haben (mit anderen Worten: A_n ist der Kern des Signum-Homomorphismus und daher sogar eine Untergruppe, die sogenannte alternierende Gruppe). Wir zerlegen S_n als die disjunkte Vereinigung von A_n und dem Komplement $S_n \setminus A_n$. Die Abbildung

$$A_n \longrightarrow S_n \setminus A_n, \quad \sigma \mapsto \tau\sigma,$$

ist eine Bijektion. Denn $\operatorname{sgn}(\tau\sigma) = \operatorname{sgn}(\tau) \operatorname{sgn}(\sigma) = -\operatorname{sgn}(\sigma)$, so dass die Abbildung tatsächlich Elemente aus A_n in die Menge $S_n \setminus A_n$ abbildet, und eine Umkehrabbildung ist gegeben durch $\sigma \mapsto \tau\sigma$ – schließlich ist $\tau^2 = \operatorname{id}$ die triviale Permutation. Wir können die Summe in der Leibniz-Formal für $\det(A)$ also zerlegen als

$$\det(A) = \sum_{\sigma \in A_n} a_{1,\sigma(1)} \cdots a_{n,\sigma(n)} - \sum_{\sigma \in A_n} a_{1,\tau\sigma(1)} \cdots a_{n,\tau\sigma(n)}.$$

Da wir hier die Permutationen nach dem Signum sortiert haben, tritt die Signumsfunktion in diesem Ausdruck nicht mehr explizit auf; wir müssen nur die zweite Summe von der ersten abziehen.

Weil $a_{i,\tau\sigma(i)} = a_{i,\sigma(i)}$ für alle $\sigma \in S_n$ und alle i gilt (für $\sigma(i) \notin \{j, j'\}$ ist ja sogar $\sigma(i) = \tau\sigma(i)$, und im anderen Fall benutzen wir unsere Voraussetzung, dass die j -te und j' -te Spalte von A übereinstimmen), sind diese Summen gleich, die Differenz ist also $= 0$. $\square$

BEISPIEL 9.8. (1) $n = 1$. In diesem Fall hat die Matrix $A = (a)$ nur einen einzigen Eintrag, und das ist ihre Determinante: $\det(A) = a$. Wenn wir $M_1(K) = K$ identifizieren, ist also $\det = \operatorname{id}_K$.

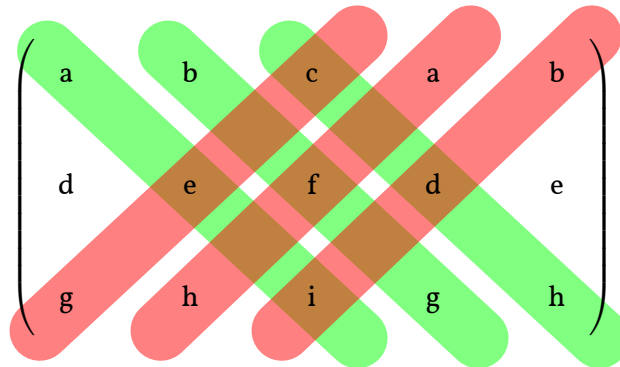
(2) $n = 2$. Für (2×2) -Matrizen finden wir den Ausdruck aus Abschnitt 2.5 bzw. Beispiel 5.56 wieder:

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} = ad - bc.$$

(3) $n = 3$. Im Fall von (3×3) -Matrizen kann man auch noch alle Permutationen auflisten und kommt zu der Formel

$$\det \begin{pmatrix} a & b & c \\ d & e & f \\ g & h & i \end{pmatrix} = aei + bfg + cdh - ceg - afh - bdi,$$

die auch als die *Regel von Sarrus* (nach dem französischen Mathematiker Pierre-Frédéric Sarrus (1798–1861)) bezeichnet wird. Man kann die Regel von Sarrus visuell darstellen, wenn man die ersten beiden Spalten der Matrix noch einmal rechts anfügt. Dann werden die Produkte über die grünen Balken addiert, die über die roten Balken werden subtrahiert.



Für $n \geq 4$ ist es nicht mehr so einfach (denn dann gibt es schon 24 Permutationen, und nicht mehr so ein explizites Schema, um sie alle aufzulisten)! **Auch für $n = 3$ ist diese Formel normalerweise nicht die beste Möglichkeit, die Determinante auszurechnen!**

- (4) Ist $A = (a_{ij})_{ij}$ eine obere Dreiecksmatrix (d.h. $a_{ij} = 0$ für alle $i > j$), so ist die Determinante von A das Produkt der Diagonaleinträge: $\det(A) = a_{11} \cdots a_{nn}$. Genauso ist es natürlich für untere Dreiecksmatrizen. In beiden Fällen ist $\sigma = \text{id}$ die einzige Permutation, die einen Beitrag zur Summe in der Leibniz-Formel liefert. Alternativ (und fast noch einfacher) kann man elementare Spaltenumformungen anwenden. Ist A nicht invertierbar, so ist sowieso $\det(A) = 0$. Andernfalls kann man A durch Spaltenumformungen vom Typ I auf Diagonalgestalt bringen; dadurch ändert sich die Determinante nicht. Dann wendet man Spaltenumformungen vom Typ III an und erhält $\det(A)$ als das Produkt aller Diagonalelemente von A mit $\det(E_n) = 1$.

◇

Vielleicht sollte der Vollständigkeit halber auch der Fall $n = 0$ erwähnt werden; dieser ist natürlich nur von rein formalem Interesse. Weil es genau eine (bijektive) Abbildung $\emptyset \rightarrow \emptyset$ gibt, hat die Summe der Leibniz-Formel einen Summanden, das leere Produkt hat den Wert 1, also ist die Determinante der einzigen Matrix der Größe 0×0 das Einselement von K .

Wie berechnet man eine Determinante?

Fast immer ist der Gauß-Algorithmus die beste Methode, um eine Determinante konkret auszurechnen. Genauer kann und sollte man sowohl Zeilen- als auch Spaltenumformungen benutzen, um die gegebene Matrix auf eine obere oder untere Dreiecksmatrix zu bringen (Lemma 9.3, Satz 9.18). Bei Umformungen vom Typ II und III ist natürlich zu beachten, wie sich die Determinante jeweils verändert. Auch die anderen Methoden, die wir später kennenlernen (zum Beispiel der Laplace'sche Entwicklungssatz) eignen sich nur in Ausnahmefällen, um eine Determinante tatsächlich auszurechnen.

Anders ist es (manchmal), wenn man theoretische Überlegungen über eine Determinante anstellt und/oder dann, wenn die betrachtete Matrix Variablen enthält, die die Anwendung des Gauß-Algorithmus erschweren.

BEISPIEL 9.9. Für eine Permutation $\sigma \in S_n$ mit zugehöriger Permutationsmatrix P_σ ist $\det(P_\sigma) = \text{sgn}(\sigma)$ (wobei hier $\text{sgn}(\sigma)$ als Element des Grundkörpers K aufzufassen ist).

Begründung. Wir schreiben σ als Produkt von Transpositionen, etwa mit r Faktoren. Dementsprechend können wir P_σ als Produkt der zugehörigen Permutationsmatrizen schreiben. Eine Matrix von rechts mit der Permutationsmatrix zu einer Transposition zu multiplizieren, bewirkt die Vertauschung zweier Spalten. Das bedeutet: Wir können $P_\sigma = E_n P_\sigma$ aus der Einheitsmatrix erhalten, indem wir r -mal zwei Spalten vertauschen. Es folgt $\det(P_\sigma) = (-1)^r = \text{sgn}(\sigma)$. Alternativ könnte man die Leibniz-Formel für P_σ anwenden; nur für die Permutation σ^{-1} erhält man einen Beitrag, und zwar $\text{sgn}(\sigma^{-1}) \cdot 1 = \text{sgn}(\sigma)$. ◇

SATZ 9.10. Sei V ein n -dimensionaler K -Vektorraum, und sei $\mathcal{B}$ eine Basis von V . Dann ist die Abbildung $\Delta_{\mathcal{B}}: V^n \rightarrow K$, $(v_1, \dots, v_n) \mapsto \det A$, wobei A die Matrix mit den Spalten $c_{\mathcal{B}}(v_1), \dots, c_{\mathcal{B}}(v_n)$ ist, eine nichttriviale Determinantenfunktion auf V .

Ist $\mathcal{C}$ eine weitere Basis von V , so gilt

$$\Delta_{\mathcal{C}} = \det(M_{\mathcal{C}}^{\mathcal{B}}) \Delta_{\mathcal{B}}.$$

BEWEIS. Es ist klar, dass $\Delta_{\mathcal{B}}$ eine nicht-triviale alternierende multilineare Abbildung ist.

Weil $\dim \mathcal{D}_V = 1$ ist, genügt es, den zweiten Teil für ein einziges Tupel $(v_1, \dots, v_n)$ zu überprüfen, auf dem $\Delta_{\mathcal{B}}$ nicht verschwindet. Schreiben wir $\mathcal{B} = (b_1, \dots, b_n)$, so ist $\Delta_{\mathcal{B}}(b_1, \dots, b_n) = \det(E_n) = 1$, und

$$\Delta_{\mathcal{C}}(b_1, \dots, b_n) = \det(c_{\mathcal{C}}(b_1), \dots, c_{\mathcal{C}}(b_n)) = \det(M_{\mathcal{C}}^{\mathcal{B}}),$$

und das liefert die Behauptung. $\square$

9.1.2. Eigenschaften der Determinante.

SATZ 9.11 (Produktsatz). Seien $A, B \in M_n(K)$. Dann gilt

$$\det(AB) = \det(A) \det(B).$$

BEWEIS. Wir fixieren die Matrix A und betrachten die Abbildung $M_n(K) \rightarrow K, B \mapsto \det(AB)$. Diese Abbildung ist multilinear und alternierend in den Spalten von B . Denn ist eine Spalte von B eine Linearkombination $ab + a'b'$ mit $a, a' \in K, b, b' \in K^n$, dann ist die entsprechende Spalte von AB die Linearkombination $aAb + a'Ab'$. Die entsprechenden Eigenschaften der Determinante übertragen sich deshalb auf die hier betrachtete Funktion.

Es gibt also nach Satz 9.6 eine Zahl $a \in K$, so dass

$$\det(AB) = a \det(B) \text{ für alle } B \in M_{n \times n}(K).$$

Setzen wir $B = E_n$ ein, so sehen wir, dass $a = \det(A)$ gelten muss. Das liefert genau die Behauptung des Satzes. $\square$

ERGÄNZUNG 9.12 (Cauchy-Binet-Formel). Alternativ kann man den Produktsatz durch eine explizite Rechnung mit der Leibniz-Formel beweisen (das ist nicht so elegant wie der obige Beweis (finde ich), aber weniger lästig, als man denken würde).

Allgemeiner gilt die **Formel von Cauchy-Binet**¹. Seien $A \in M_{m \times n}(K), B \in M_{n \times m}(K)$, sei $\mathcal{M}$ die Menge aller m -elementigen Teilmengen von $\{1, \dots, n\}$, und sei für $M \in \mathcal{M}$ einerseits $A_M \in M_m(K)$ die Matrix, die aus den Spalten von A mit Spaltenindex in M , und andererseits $B_M \in M_m(K)$ die Matrix, die aus den Zeilen von B mit Zeilenindex in M besteht. Dann gilt

$$\det(AB) = \sum_{M \in \mathcal{M}} \det(A_M) \det(B_M).$$

Im Fall $m = n$ hat die Summe nur ein Element und die Formel ist genau der Produktsatz für die Determinante eines Produkts quadratischer Matrizen.

Übungsaufgabe: Was sagt die Cauchy-Binet-Formel im Fall $m > n$ aus, und warum ist sie auch in diesem Fall gültig? $\square$ Ergänzung 9.12

BEMERKUNG 9.13. In ähnlicher Weise wie Satz 9.11 kann man zeigen, dass für eine Blockmatrix der Form $\begin{pmatrix} A & B \\ 0 & D \end{pmatrix}$ für $A \in M_r(K), B \in M_{r \times s}(K), D \in M_s(K)$ gilt, dass

$$\det \begin{pmatrix} A & B \\ 0 & D \end{pmatrix} = \det(A) \det(D).$$

Zum Beweis fixieren wir B und D und betrachten die Abbildung

$$A \mapsto \det \begin{pmatrix} A & B \\ 0 & D \end{pmatrix},$$

¹https://en.wikipedia.org/wiki/Cauchy%E2%80%93Binet_formula

die multilinear und alternierend in den Spalten von A ist. Für $A = E_r$ ist

$$\det \begin{pmatrix} E_r & B \\ 0 & D \end{pmatrix} = \det \begin{pmatrix} E_r & 0 \\ 0 & D \end{pmatrix}.$$

Es ist noch zu zeigen, dass diese Matrix Determinante $\det(D)$ hat. Dazu kann man auch den Satz über die Eindeutigkeit von Determinantenfunktionen bis auf Vielfache anwenden und nun die Abbildung

$$D \mapsto \det \left(\begin{pmatrix} E_r & 0 \\ 0 & D \end{pmatrix} \right)$$

betrachten, die multilinear und alternierend in den Spalten von D und mithin ein Vielfaches der Abbildung $D \mapsto \det(D)$ ist. Durch Einsetzen von $D = E_s$ erhalten wir das gewünschte Ergebnis. $\diamond$

KOROLLAR 9.14. Ist $A \in GL_n(K)$, so gilt $\det(A) \neq 0$ und

$$\det(A^{-1}) = \det(A)^{-1}.$$

BEWEIS. Es gilt $AA^{-1} = E_n$, also nach Satz 9.11:

$$\det(A) \det(A^{-1}) = \det(E_n) = 1.$$

Das ist genau die Behauptung. $\square$

SATZ 9.15. Sei $A \in M_{n \times n}(K)$. Die Matrix A ist genau dann invertierbar, wenn $\det(A) \neq 0$.

BEWEIS. Wir haben bereits gesehen, dass $\det(A) = 0$, wenn A nicht invertierbar ist. Aus Korollar 9.14 folgt andererseits, dass für eine invertierbare Matrix A die Determinante nicht verschwinden kann. $\square$

Insbesondere sehen wir, dass für jede quadratische Matrix $A \in M_n(K)$ und invertierbare Matrix $S \in GL_n(K)$ gilt, dass $\det(SAS^{-1}) = \det(A)$ ist; zueinander konjugierte Matrizen haben dieselbe Determinante.

KOROLLAR 9.16. Die Einschränkung von $\det$ auf die Gruppe $GL_n(K)$ ist ein Gruppenhomomorphismus $GL_n(K) \rightarrow K^\times$. Sein Kern ist die Gruppe $SL_n(K)$.

BEWEIS. Wir wissen wegen Satz 9.15, dass $\det(GL_n(K)) \subseteq K^\times$. Satz 9.11 liefert dann, dass $\det: GL_n(K) \rightarrow K^\times$ ein Gruppenhomomorphismus ist.

Es ist klar, dass die Elementarmatrizen $E_{ij}(a)$ alle Determinante $= 1$ haben (da es sich um obere bzw. untere Dreiecksmatrizen handelt, deren Diagonaleinträge alle $= 1$ sind), sie liegen folglich im Kern des Homomorphismus $\det$. Da die Gruppe $SL_n(K)$ nach Definition von diesen Matrizen erzeugt wird, gilt $SL_n(K) \subseteq \text{Ker}(\det)$.

Ist andererseits A eine Matrix mit Determinante 1 , so können wir A nach Satz 8.34 schreiben als Produkt $B \text{ diag}(1, \dots, 1, d)$ mit $B \in SL_n(K)$. Es gilt dann $1 = \det(A) = \det(B) \cdot d = d$, also $A = B \in SL_n(K)$. $\square$

Insbesondere zeigt der Produktsatz auch, dass $d = d' = \det(A)$ in Satz 8.34 gilt; daher sind auch $B = A \text{ diag}(1, \dots, 1, \det(A)^{-1})$ und $C = \text{diag}(1, \dots, 1, \det(A)^{-1})A$ eindeutig durch A bestimmt. Oft wird die Gruppe $SL_n(K)$ definiert als die Gruppe aller invertierbaren Matrizen mit Determinante 1 . Der obige Satz zeigt, dass das zu unserer Definition (für die es auch gute Gründe gibt) äquivalent ist.

SATZ 9.17. Sei $A \in M_n(K)$. Dann gilt $\det(A^t) = \det(A)$.

BEWEIS. Ist $\text{rg}(A) < n$, so gilt auch $\text{rg}(A^t) < n$, und in diesem Fall ist $\det(A) = 0 = \det(A^t)$. Andernfalls ist A invertierbar, und wir können $A = CD$ mit $C \in \text{SL}_n(K)$ und einer Diagonalmatrix D schreiben. Dann gilt $\det(C) = 1 = \det(C^t)$, weil auch C^t in $\text{SL}_n(K)$ liegt (Beispiel 8.33), und $\det(D) = \det(D^t)$, weil $D = D^t$ gilt. Also erhalten wir $\det(A^t) = \det(D^t C^t) = \det(D^t) \det(C^t) = \det(D) = \det(A)$ aus dem Produktsatz.

Alternativ kann man die im Satz behauptete Gleichheit auch direkt anhand der Leibniz-Formel nachrechnen, indem man ausnutzt, dass $S_n = \{\sigma^{-1}; \sigma \in S_n\}$ und dass $\text{sgn}(\sigma^{-1}) = \text{sgn}(\sigma)$ für alle σ gilt. $\square$

SATZ 9.18. Die Determinante $\det: M_{n \times n}(K) \rightarrow K$ ist multilinear und alternierend in den Zeilen der Matrix.

Sie verhält sich unter elementaren Zeilenumformungen ebenso wie unter elementaren Spaltenumformungen, das heißt:

- (1) Geht A' aus A durch eine elementare Zeilenumformung vom Typ I hervor (Addieren des Vielfachen einer Zeile zu einer anderen Zeile), so gilt $\det(A') = \det(A)$.
- (2) Geht A' aus A durch eine elementare Zeilenumformung vom Typ II hervor (Vertauschen zweier Zeilen), so gilt $\det(A') = -\det(A)$.
- (3) Geht A' aus A durch eine elementare Zeilenumformung vom Typ III hervor (Multiplikation einer Zeile mit $a \in K^\times$), so gilt $\det(A') = a \det(A)$.

BEWEIS. Weil $A \mapsto \det(A)$ multilinear und alternierend in den Spalten von A ist, ist $A \mapsto \det(A^t)$ multilinear und alternierend in den Zeilen von A . Der erste Teil des Satzes folgt also daraus, dass $\det(A) = \det(A^t)$, wie wir in Satz 9.17 gezeigt haben.

Dass sich die Determinante unter elementaren Zeilenumformungen ebenso verhält, wie unter Spaltenumformungen, folgt direkt aus dem Produktsatz, weil wir elementare Zeilenumformungen als Multiplikation von links mit Matrizen der Form $E_{ij}(a)$, P_i , $\text{diag}(1, \dots, 1, a, 1, \dots, 1)$ ausdrücken können, deren Determinante 1 bzw. -1 bzw. a ist. $\square$

BEISPIEL 9.19. Den Formalismus der Determinante und speziell den Produktsatz können wir auch benutzen, um Identitäten zwischen den Fibonacci-Zahlen (siehe Beispiel 5.60) herzuleiten. Wir hatten in dem genannten Beispiel für die Fibonacci-Zahlen F_n gezeigt, dass

$$\begin{pmatrix} F_{n+1} & F_n \\ F_n & F_{n-1} \end{pmatrix} = A^n$$

gilt, wobei A die Matrix $\begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}$ bezeichnet. Bilden wir auf beiden Seiten die Determinante, so erhalten wir

$$F_{n+1}F_{n-1} - F_n^2 = \det(A)^n = (-1)^n,$$

das ist die Identität von Cassini. $\diamond$

BEISPIEL 9.20 (Die Vandermonde-Determinante). Als ein Beispiel einer Determinantenberechnung, die häufig nützlich ist, betrachten wir die Vandermonde-Matrix

$$A := \begin{pmatrix} 1 & x_1 & x_1^2 & \cdots & x_1^{n-1} \\ 1 & x_2 & x_2^2 & \cdots & x_2^{n-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & x_n^2 & \cdots & x_n^{n-1} \end{pmatrix} \in M_n(K)$$

für $x_1, \dots, x_n \in K$.

Behauptung. $\det(A) = \prod_{1 \leq i < j \leq n} (x_j - x_i)$.

Wir beweisen die Behauptung durch vollständige Induktion nach n . Der Induktionsanfang $n = 1$ ist klar. Zunächst ziehen wir das x_1 -fache der $(n-1)$ -ten Spalte von der n -ten Spalte ab, dann das x_1 -fache der $(n-2)$ -ten Spalte von der $(n-1)$ -ten Spalte, und so weiter, schließlich das x_1 -fache der ersten von der zweiten Spalte. Die Determinante verändert sich dadurch nicht. Wir haben nun die Matrix

$$\begin{pmatrix} 1 & 0 & 0 & \cdots & 0 \\ 1 & x_2 - x_1 & x_2(x_2 - x_1) & \cdots & x_2^{n-2}(x_2 - x_1) \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n - x_1 & x_n(x_n - x_1) & \cdots & x_n^{n-2}(x_n - x_1) \end{pmatrix}$$

erhalten. Durch oder eine Variante von Bemerkung 9.13 für untere Block-Dreiecksmatrizen (oder durch Entwicklung nach der ersten Zeile, siehe Satz 9.29) erhalten wir

$$\begin{aligned} \det(A) &= \det \begin{pmatrix} x_2 - x_1 & x_2(x_2 - x_1) & \cdots & x_2^{n-2}(x_2 - x_1) \\ \vdots & \vdots & \ddots & \vdots \\ x_n - x_1 & x_n(x_n - x_1) & \cdots & x_n^{n-2}(x_n - x_1) \end{pmatrix} \\ &= \prod_{i=2}^n (x_i - x_1) \det \begin{pmatrix} 1 & x_2 & \cdots & x_2^{n-2} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & \cdots & x_n^{n-2} \end{pmatrix}. \end{aligned}$$

Wir können nun die Induktionsvoraussetzung einsetzen und erhalten die Behauptung.

Eine Anwendung ist ein neuer Beweis von Satz 4.25, der (äquivalent umformuliert) besagt, dass die Koeffizienten einer Polynomfunktion $K \rightarrow K, x \mapsto \sum_{i=0}^{n-1} a_i x^i$, die n paarweise verschiedene Nullstellen $x_1, \dots, x_n \in K$ hat, alle gleich 0 sein müssen.

Weil die x_i paarweise verschieden sind, hat die zu diesen Elementen x_i wie oben gebildete Vandermonde-Matrix A eine von 0 verschiedene Determinante, also ist A invertierbar. Für den aus den Koeffizienten gebildeten Vektor $v = (a_0, \dots, a_{n-1})^t \in K^n$ gilt aber $Av = 0$; das ist einfach eine Umformulierung der Voraussetzung, dass alle x_i Nullstellen dieser Polynomfunktion sind. Es folgt $v = 0$, wie gewünscht.

◇

9.2. Die Determinante eines Endomorphismus

Wir betrachten nun einen endlich-dimensionalen Vektorraum V und einen Endomorphismus $V \rightarrow V$.

DEFINITION 9.21. Seien V ein n -dimensionaler Vektorraum, Δ eine nicht-triviale Determinantenfunktion auf V und $f: V \rightarrow V$ ein Endomorphismus von V . Dann ist $\Delta_f: V^n \rightarrow K, (v_1, \dots, v_n) \mapsto \Delta(f(v_1), \dots, f(v_n))$ eine Determinantenfunktion und das Element $a \in K$ mit $\Delta_f = a\Delta$ ist unabhängig von der Wahl von Δ . Es heißt die *Determinante des Endomorphismus* f , in Zeichen: $\det(f)$. ⊥

Es ist klar, dass die in der Definition angegebene Abbildung tatsächlich multilinear (weil f linear ist) und alternierend ist, da Δ diese Eigenschaften hat, und daher eine Determinantenfunktion ist. Auch die Unabhängigkeit von Δ folgt daraus, dass der Vektorraum $\mathcal{D}_V$ der Determinantenfunktionen eindimensional ist. Denn ist Δ' eine andere Determinantenfunktion, so existiert ein (eindeutig bestimmtes) $c \in K$ mit $\Delta' = c\Delta$, also ist $(\Delta')_f = c\Delta_f = ca\Delta = a\Delta'$.

SATZ 9.22. Ist $\mathcal{B}$ eine Basis von V , so gilt $\det(f) = \det(M_{\mathcal{B}}^{\mathcal{B}}(f))$. Insbesondere ist $\det(M_{\mathcal{B}}^{\mathcal{B}}(f))$ unabhängig von der Wahl von $\mathcal{B}$.

BEWEIS. Wir können $\det(f)$ mit der Determinantenfunktion $\Delta: V^n \rightarrow M_n(K) \rightarrow K$ berechnen, die als Verkettung der Abbildungen $V^n \rightarrow M_n(K), (v_1, \dots, v_n) \mapsto (c_{\mathcal{B}}(v_1), \dots, c_{\mathcal{B}}(v_n)) \in (K^n)^n = M_n(K)$ und $\det: M_n(K) \rightarrow K$ entsteht. Schreiben wir $\mathcal{B} = (b_1, \dots, b_n)$, so erhalten wir $\Delta(b_1, \dots, b_n) = \det(E_n) = 1$ und mit der Notation wie in Definition 9.21, dass $\Delta_f(b_1, \dots, b_n) = \det(M_{\mathcal{B}}^{\mathcal{B}}(f))$ ist. Daraus folgt die Behauptung. $\square$

Dass $\det(M_{\mathcal{B}}^{\mathcal{B}}(f))$ unabhängig von der Wahl von $\mathcal{B}$ ist, kann man natürlich auch anhand der Basiswechselformel aus dem Produktsatz für Determinanten folgern. Es ist eine Umformulierung der Tatsache, dass zueinander konjugierte Matrizen dieselbe Determinante haben.

Die Eigenschaften der Determinante von Matrizen übertragen sich in naheliegender Weise auf die Determinante von Endomorphismen.

SATZ 9.23. Sei V ein endlich-dimensionaler Vektorraum.

- (1) Es gilt $\det(\text{id}_V) = 1$.
- (2) Seien f und g Endomorphismen von V . Dann gilt $\det(f \circ g) = \det(f) \det(g)$.
- (3) Sei f ein Endomorphismus von V . Es gilt genau dann $\det(f) \neq 0$, wenn f ein Isomorphismus ist, und dann ist $\det(f^{-1}) = \det(f)^{-1}$.
- (4) Für die Determinante des zu f dualen Endomorphismus $f^\vee: V^\vee \rightarrow V^\vee$ gilt $\det(f^\vee) = \det(f)$.

BEWEIS. Die Aussagen in (1) und (2) folgen leicht aus der Definition, und (3) kann man dann aus (2) erhalten.

Alternativ kann man die Beschreibung von $\det(f)$ als die Determinante einer darstellenden Matrix benutzen. Die Teile (1), (2), (3) folgen dann direkt aus den entsprechenden Eigenschaften der Determinante von Matrizen. Für Teil (4) benutzt man zusätzlich Satz 7.54 und Satz 9.17. $\square$

”

We [he and Halmos] share a philosophy about linear algebra: we think basis-free, we write basis-free, but when the chips are down we close the office door and compute with matrices like fury.

Kaplansky, Irving, in: P. Halmos, *Celebrating 50 Years of Mathematics*
Gefunden auf <http://math.furman.edu/~mwoodard/mqs/data.html>

”

There is hardly any theory which is more elementary [than linear algebra], in spite of the fact that generations of professors and textbook writers have obscured its simplicity by preposterous calculations with matrices.

J. Dieudonné, in: *Foundations of Modern Analysis*, Vol. 1

ERGÄNZUNG 9.24. Mit der Determinante können wir die »richtige« Definition des Begriffs der *Drehung* geben. Wir arbeiten über dem Körper der reellen Zahlen.

DEFINITION 9.25. Sei $n \in \mathbb{N}$. Ein Endomorphismus $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt *Drehung*, wenn für die darstellende Matrix $A = M(f)$ gilt:

- (a) $AA^t = E_n$, und
 (b) $\det(A) = 1$.

⊥

Die Bedingung (1) ist damit äquivalent, dass f abstandserhaltend für den üblichen (»euklidischen«) Abstand $d(v, w)$ von Punkten $v, w \in \mathbb{R}^n$ ist, es gilt also $d(f(v), f(w)) = d(v, w)$ für alle $v, w \in \mathbb{R}^n$. Siehe Ergänzung 7.59 und die Abschnitte 11.2.1, 11.2.5. Daraus folgt $1 = \det(E_n) = \det(A) \det(A^t) = \det(A)^2$, also $\det(A) = 1$ oder $\det(A) = -1$.

Die Äquivalenz dieser Definition mit der Definition in Ergänzung 7.59 werden wir für allgemeines n erst in der Linearen Algebra 2 sehen. (Es ist klar, dass ein Endomorphismus der Form $g \circ g$ nicht Determinante -1 haben kann, weil die Determinante von $g \circ g$ gleich $\det(g)^2$ ist. Um zu zeigen, dass man eine Drehung im Sinne der hier gegebenen Definition immer als die Verkettung eines Endomorphismus mit sich selbst schreiben kann, muss man aber etwas mehr arbeiten, und wir verschieben das noch. Für $n = 2$ ist es nicht so schwer, vergleiche Satz 11.27 und Abschnitt 11.5.)

Da eine Matrix A mit der Eigenschaft $AA^t = E_n$ Determinante 1 oder -1 hat, ist es äquivalent, $\det(A) = 1$ oder $\det(A) > 0$ zu fordern. Eine abstandserhaltende Abbildung ist also genau dann eine Drehung, wenn sie »orientierungserhaltend« im Sinne der folgenden Ergänzung ist. □ Ergänzung 9.24

ERGÄNZUNG 9.26 (Orientierung). Über den reellen Zahlen haben wir das besondere Phänomen, dass die Determinante einer invertierbaren Matrix A entweder positiv oder negativ ist. Das ermöglicht es uns zu definieren, wann zwei Basen eines reellen Vektorraums *dieselbe Orientierung* haben:

DEFINITION 9.27. (1) Sei V ein endlich-dimensionaler $\mathbb{R}$ -Vektorraum. Wir sagen, dass zwei (angeordnete) Basen $\mathcal{B}$ und $\mathcal{B}'$ von V *dieselbe Orientierung* haben, wenn die Determinante der Basiswechselmatrix $M_{\mathcal{B}}^{\mathcal{B}'}$ positiv ist. Andernfalls sagen wir, dass die Basen die *entgegengesetzte Orientierung* haben.

(2) Ist speziell $V = \mathbb{R}^n$ der Standardvektorraum, so sagen wir auch, eine Basis sei *positiv orientiert*, wenn sie dieselbe Orientierung wie die Standardbasis hat.

⊥

Weil $M_{\mathcal{B}'}^{\mathcal{B}} = (M_{\mathcal{B}}^{\mathcal{B}'})^{-1}$ ist, haben die Determinanten der beiden Basiswechselmatrizen dasselbe Vorzeichen. Genauer ist die Eigenschaft, dieselbe Orientierung zu haben, eine Äquivalenzrelation auf der Menge aller Basen von V (Abschnitt 3.14.2).

Für $\mathbb{R}^2$ besitzt eine Basis v, w genau dann dieselbe Orientierung wie die Standardbasis e_1, e_2 , wenn »der Winkel von v nach w gegen den Uhrzeigersinn« kleiner als 180° ist. Zum Beispiel haben die Basen

$$\left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right), \quad \left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right), \quad \left(\begin{pmatrix} 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right), \quad \left(\begin{pmatrix} -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right)$$

alle dieselbe Orientierung wie die Standardbasis. Wenn man die Reihenfolge der beiden Basisvektoren vertauscht, wechselt die Orientierung.

DEFINITION 9.28. Ein Automorphismus $f: V \rightarrow V$ eines endlich-dimensionalen $\mathbb{R}$ -Vektorraums heißt *orientierungserhaltend*, wenn er die folgenden äquivalenten Bedingungen erfüllt:

- (i) Es gibt eine Basis $\mathcal{B} = (v_1, \dots, v_n)$, so dass die Basis $(f(v_1), \dots, f(v_n))$ dieselbe Orientierung hat wie $\mathcal{B}$.
- (ii) Für jede Basis $\mathcal{B} = (v_1, \dots, v_n)$ hat die Basis $(f(v_1), \dots, f(v_n))$ dieselbe Orientierung wie $\mathcal{B}$.
- (iii) Es gilt $\det(f) > 0$.

⊢

Seien $a_{ij}: I \rightarrow \mathbb{R}$ stetige Funktionen, die auf einem Intervall $I \subseteq \mathbb{R}$ definiert sind. Dann ist auch die Funktion $I \rightarrow \mathbb{R}, t \mapsto d(t) := \det((a_{ij}(t))_{i,j})$ stetig, weil die Determinante ein polynomieller Ausdruck in den Einträgen der Matrix ist und Summen und Produkte stetiger Funktionen wieder stetig sind. Daraus folgt: Ist $x_0 \in I$ ein innerer Punkt, so dass $d(t) \neq 0$ ist, so existiert $\varepsilon > 0$, so dass $I' := (x_0 - \varepsilon, x_0 + \varepsilon)$ in I enthalten ist, und dass $d(t) \neq 0$ für alle $t \in I'$ ist. Aus dem Zwischenwertsatz folgt, dass die Determinanten $d(t)$ dann für alle $t \in I'$ dasselbe Vorzeichen haben. Wenn wir zu der Matrix $(a_{ij}(t))_{i,j}$ die von ihren Spalten gebildete Basis betrachten, haben diese Basen für alle $t \in I'$ dieselbe Orientierung. Wir sagen, dass sich »die Orientierung in einer stetigen Familie von Basen nicht ändert«. □ Ergänzung 9.26

9.3. Die Cramersche Regel

Wir lernen in diesem Abschnitt mit dem Entwicklungssatz von Laplace eine Möglichkeit kennen, die Determinante einer Matrix A in Termen von Determinanten von »Untermatrizen« von A auszudrücken. Das führt dann auch auf die Cramersche Regel, mit der wir eine Formel für die inverse Matrix einer invertierbaren Matrix erhalten (die zwar für konkrete Berechnungen viel zu aufwändig ist, aber einen theoretischen Nutzen hat).

Sei $A = (a_{ij})_{i,j} \in M_{n \times n}(K)$ eine quadratische Matrix über dem Körper K . Für $i, j \in \{1, \dots, n\}$ bezeichnen wir mit $A_{ij} \in M_{n \times n}$ die Matrix, die aus A durch Ersetzen der i -ten Zeile durch den j -ten Standard-Zeilenvektor $e_j^t = (0, \dots, 0, 1, 0, \dots, 0)$ (1 an der j -ten Stelle) und der j -ten Spalte durch den i -ten Standardbasisvektor e_i entsteht,

$$A_{ij} = \begin{pmatrix} a_{11} & \cdots & a_{1,j-1} & 0 & a_{1,j+1} & \cdots & a_{1n} \\ \vdots & & \vdots & \vdots & \vdots & & \vdots \\ a_{i-1,1} & \cdots & a_{i-1,j-1} & 0 & a_{i-1,j+1} & \cdots & a_{i-1,n} \\ 0 & \cdots & 0 & 1 & 0 & \cdots & 0 \\ a_{i+1,1} & \cdots & a_{i+1,j-1} & 0 & a_{i+1,j+1} & \cdots & a_{i+1,n} \\ \vdots & & \vdots & \vdots & \vdots & & \vdots \\ a_{n1} & \cdots & a_{n,j-1} & 0 & a_{n,j+1} & \cdots & a_{nn} \end{pmatrix}.$$

Ferner bezeichnen wir mit $A'_{ij} \in M_{(n-1) \times (n-1)}(K)$ die Matrix, die aus A (oder äquivalent aus A_{ij}) durch Streichen der i -ten Zeile und der j -ten Spalte hervorgeht,

$$A'_{ij} = \begin{pmatrix} a_{11} & \cdots & a_{1,j-1} & a_{1,j+1} & \cdots & a_{1n} \\ \vdots & & \vdots & \vdots & & \vdots \\ a_{i-1,1} & \cdots & a_{i-1,j-1} & a_{i-1,j+1} & \cdots & a_{i-1,n} \\ a_{i+1,1} & \cdots & a_{i+1,j-1} & a_{i+1,j+1} & \cdots & a_{i+1,n} \\ \vdots & & \vdots & \vdots & & \vdots \\ a_{n1} & \cdots & a_{n,j-1} & a_{n,j+1} & \cdots & a_{nn} \end{pmatrix}.$$

Es gilt dann

$$\det A_{ij} = (-1)^{i+j} \det \begin{pmatrix} 1 & 0 & \cdots & \cdots & \cdots & \cdots & 0 \\ 0 & a_{11} & \cdots & a_{1,j-1} & a_{1,j+1} & \cdots & a_{1n} \\ \vdots & \vdots & & \vdots & \vdots & & \vdots \\ 0 & a_{i-1,1} & \cdots & a_{i-1,j-1} & a_{i-1,j+1} & \cdots & a_{i-1,n} \\ 0 & a_{i+1,1} & \cdots & a_{i+1,j-1} & a_{i+1,j+1} & \cdots & a_{i+1,n} \\ \vdots & \vdots & & \vdots & \vdots & & \vdots \\ 0 & a_{n1} & \cdots & a_{n,j-1} & a_{n,j+1} & \cdots & a_{nn} \end{pmatrix} = (-1)^{i+j} \det A'_{ij}.$$

wobei wir für die erste Gleichheit die j -te Spalte mit $j-1$ Vertauschungen in die erste Spalte bringen, und die i -te Zeile mit $i-1$ Vertauschungen in die erste Zeile bringen. Das bewirkt die Änderung des Vorzeichens um $(-1)^{j-1+i-1} = (-1)^{i+j}$. Für die zweite Gleichheit können wir Beispiel 9.13 benutzen.

Damit können wir den Laplaceschen Entwicklungssatz formulieren (und der Beweis benötigt weniger Platz, als wir für die Erklärung der Notation verwendet haben).

SATZ 9.29 (Laplacescher Entwicklungssatz). *Mit den obigen Notationen gilt für alle i die »Entwicklung von $\det A$ nach der i -ten Zeile«:*

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A'_{ij}$$

und für alle j die »Entwicklung von $\det A$ nach der j -ten Spalte«:

$$\det A = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det A'_{ij}.$$

BEWEIS. Wir beweisen die Entwicklung nach der i -ten Zeile in der Form

$$\det A = \sum_{j=1}^n a_{ij} \det A_{ij}.$$

Nach der obigen Diskussion ist klar, dass diese Behauptung zur Formulierung im Satz äquivalent ist. Die Entwicklung nach einer Spalte erhält man durch ein analoges Argument, oder indem man die Entwicklung nach einer Zeile auf die transponierte Matrix anwendet und Satz 9.17 benutzt.

Die zu beweisende Formel erhalten wir direkt aus der Leibniz-Formel, indem wir die dort auftretende Summe über $\sigma \in S_n$ nach dem Wert $\sigma(i)$ sortieren:

$$\begin{aligned} \det(A) &= \sum_{j=1}^n \sum_{\sigma \in S_n, \sigma(i)=j} \operatorname{sgn}(\sigma) a_{ij} \prod_{l \neq j} a_{l, \sigma(l)} \\ &= \sum_{j=1}^n a_{ij} \det(A_{ij}). \end{aligned}$$

Im zweiten Schritt benutzen wir, dass der Summand zu einer Permutation σ in der Leibniz-Formel für A_{ij} im Fall $\sigma(i) = j$ gleich $\prod_{l \neq j} a_{l, \sigma(l)}$ ist, und im Fall $\sigma(i) \neq j$ verschwindet. $\square$

Es ist auch möglich (und nicht schwieriger), den Satz ohne Verwendung der Leibniz-Formel zu beweisen, indem man ausnutzt, dass die Determinante alternierend und multilinear ist. Siehe Bemerkung 9.42.

BEMERKUNG 9.30. Zur praktischen Berechnung von Determinanten sollte man den Entwicklungssatz (oder gar die Leibniz-Formel) nach Möglichkeit vermeiden. Stattdessen ist es fast immer besser, Zeilen- und Spaltenumformungen vorzunehmen, um die Matrix in eine einfachere Form (vorzugsweise natürlich in eine Dreiecksmatrix) umzuformen. Nur wenn die betrachtete Matrix Zeilen oder Spalten hat, in denen nur ein einziger Eintrag vorkommt oder wenn komplizierte Fallunterscheidungen/Divisionen erforderlich sind (weil in der Matrix Variablen oder Funktionen o.ä. auftreten), ist der Entwicklungssatz eine sinnvolle Alternative. $\diamond$

DEFINITION 9.31. Wir definieren die *Komplementärmatrix* A^{ad} zu A durch

$$(A^{ad})_{ij} = (-1)^{i+j} \det A'_{ji}.$$

In der Literatur wird A^{ad} manchmal als die *adjungierte* Matrix bezeichnet; daher auch die Notation. Wir vermeiden diese Sprechweise aber, weil der Begriff der adjungierten Abbildung bzw. adjungierten Matrix in der Linearen Algebra 2 bei der Behandlung von Bilinearformen mit einer völlig anderen Bedeutung auftritt. $\dashv$

Wenn wir den Entwicklungssatz für alle Spalten bzw. alle Zeilen »zusammenfassen«, erhalten wir die nach [Gabriel Cramer](https://de.wikipedia.org/wiki/Gabriel_Cramer)² (1704–1752) benannte Regel. Genau genommen gab es damals weder den Begriff der Determinante noch den der Matrix in der heutigen Form, so dass die Formulierung im folgenden Satz als eine Verallgemeinerung/Umformulierung zu verstehen ist. Cramer ging es seinerzeit um eine Lösungsformel für eindeutig lösbare lineare Gleichungssysteme, vergleiche Korollar 9.34.

SATZ 9.32 (Cramersche Regel). *Es gilt*

$$AA^{ad} = A^{ad}A = \det(A)E_n.$$

Ist A invertierbar, so gilt

$$A^{-1} = \det(A)^{-1}A^{ad}.$$

BEWEIS. Die zweite Aussage folgt direkt aus der ersten. Für die erste Aussage betrachten wir den Eintrag in Zeile i und Spalte j des Produkts AA^{ad} , den wir nach Definition des Matrizenprodukts und der Matrix A^{ad} wie folgt ausdrücken können:

$$\sum_{l=1}^n a_{il} \cdot (-1)^{l+j} \det A'_{jl}$$

Ist $i = j$, so ist dies die Entwicklung von $\det(A)$ nach der i -ten Zeile von A und wir erhalten nach dem Laplaceschen Satz

$$\sum_{l=1}^n a_{il} \cdot (-1)^{l+i} \det A'_{il} = \det(A).$$

Ist andererseits $i \neq j$, so können wir den Ausdruck – wieder nach dem Entwicklungssatz – als die Determinante der Matrix verstehen, die aus A dadurch entsteht, dass wir die j -te Zeile durch die i -te Zeile ersetzen (die Matrizen A'_{jl} ändern sich durch diese Ersetzung nicht, weil dort die j -te Zeile von A gestrichen wurde). Da diese Matrix zwei gleiche Zeilen hat, ist ihre Determinante gleich 0. Insgesamt haben wir damit bewiesen, dass $AA^{ad} = \det(A)E_n$ gilt.

Um $A^{ad}A = \det(A)E_n$ zu zeigen, argumentiert man analog mit der Entwicklung nach einer Spalte oder überlegt sich, dass $(A^t)^{ad} = (A^{ad})^t$ gilt und wendet die vorherige Überlegung auf A^t an. $\square$

² https://de.wikipedia.org/wiki/Gabriel_Cramer

KOROLLAR 9.33. Sei $A \in GL_n(\mathbb{Q})$. Wir nehmen an, dass alle Einträge von A in $\mathbb{Z}$ liegen. Genau dann liegen auch alle Einträge von A^{-1} in $\mathbb{Z}$, wenn $\det(A) \in \{1, -1\}$.

BEWEIS. Weil alle Einträge von A in $\mathbb{Z}$ liegen, gilt auch $\det(A) \in \mathbb{Z}$. Das sehen wir beispielsweise aus der Leibniz-Formel. Es ist auch klar, dass $\det(A) \in \{1, -1\}$ folgt, wenn auch alle Einträge von A^{-1} in $\mathbb{Z}$ sind. Denn nach dem Produktsatz gilt $\det(A) \det(A^{-1}) = 1$, und beide Faktoren sind ganze Zahlen.

Die umgekehrte, schwierigere Implikation folgt nun aus der Cramerschen Regel. Denn die Einträge der Komplementärmatrix A^{ad} sind polynomielle Ausdrücke in den Einträgen von A , sind also wegen unserer Voraussetzung an A sämtlich ganze Zahlen. Wenn $\det(A) \in \{1, -1\}$ ist, folgt, dass auch die Matrix $A^{-1} = \det(A)^{-1} A^{\text{ad}}$ nur ganzzahlige Einträge hat. $\square$

Wir erhalten aus der obigen abstrakten Form der Cramerschen Regel auch die oben schon erwähnte »Lösungsformel« für lineare Gleichungssysteme. Um ein konkret gegebenes Gleichungssystem zu lösen, ist der Gauß-Algorithmus in aller Regel besser geeignet, aber manchmal ist die Formel nützlich.

KOROLLAR 9.34. Sei $Ax = b$ ein lineares Gleichungssystem mit n Unbestimmten und n Gleichungen (d.h. die Koeffizientenmatrix A ist eine quadratische Matrix der Größe $n \times n$), das genau eine Lösung besitzt (d.h. A ist invertierbar).

Für $i = 1, \dots, n$ sei A_i die Matrix, die aus A entsteht, indem die i -te Spalte von A durch den Vektor b ersetzt wird. Dann ist die eindeutig bestimmte Lösung des obigen linearen Gleichungssystems gegeben durch den Vektor mit den Einträgen

$$\left(\frac{\det(A_1)}{\det(A)}, \dots, \frac{\det(A_n)}{\det(A)} \right)^t.$$

BEWEIS. Da A invertierbar ist, ist $\det(A) \neq 0$, so dass wir durch $\det(A)$ teilen können. Die eindeutige Lösung des gegebenen linearen Gleichungssystems ist

$$A^{-1}b = \frac{1}{\det(A)} A^{\text{ad}} b,$$

wobei wir die Cramersche Regel in der Form von Satz 9.32 benutzt haben.

Der i -te Eintrag des Vektors $A^{\text{ad}} b$ ist mit den obigen Bezeichnungen

$$\sum_{j=1}^n (-1)^{i+j} \det(A'_{ji}) b_j = \det(A_i),$$

wobei wir den Laplaceschen Entwicklungssatz für die Matrix A_i benutzen (und dass $(A_i)'_{ji} = A'_{ji}$ gilt, weil für die Bildung dieser Matrix die i -te Spalte ohnehin gestrichen wird). $\square$

9.4. Die Spur einer Matrix

Zum Abschluss des Kapitels definieren wir noch die sogenannte *Spur* einer quadratischen Matrix. Die Spur von A ist, wie die Determinante, ein Element von K . Einerseits ist die Spur einfacher zu definieren als die Determinante, andererseits ist es zu Beginn weniger klar, was man aus der Spur von A über die Matrix A »ablesen« kann. Jedenfalls werden wir sehen, dass zueinander konjugierte Matrizen dieselbe Spur haben. Nachdem wir die Theorie der Eigenwerte einer Matrix in der Linearen Algebra 2 noch weiter vertieft haben, werden wir die Spur auch noch durchsichtiger interpretieren können.

DEFINITION 9.35. Sei $A = (a_{ij})_{i,j} \in M_{n \times n}(K)$. Wir definieren die *Spur* von A als

$$\text{Spur}(A) = \sum_{i=1}^n a_{ii},$$

das heißt als die Summe der Diagonalelemente von A . $\dashv$

Es gilt $\text{Spur}(A + B) = \text{Spur}(A) + \text{Spur}(B)$ und $\text{Spur}(aA) = a \text{ Spur}(A)$ für $A, B \in M_n(K)$, $a \in K$, wir erhalten also eine lineare Abbildung $\text{Spur}: M_n(K) \rightarrow K$.

SATZ 9.36. Seien $A \in M_{m \times n}(K)$ und $B \in M_{n \times m}(K)$ Matrizen. Dann gilt $\text{Spur}(AB) = \text{Spur}(BA)$.

BEWEIS. Wir bezeichnen die Einträge von A mit a_{ij} , die von B mit b_{ji} . Dann gilt

$$\text{Spur}(AB) = \sum_{i=1}^m \sum_{j=1}^n a_{ij} b_{ji} = \sum_{j=1}^n \sum_{i=1}^m b_{ji} a_{ij} = \text{Spur}(BA).$$

□

KOROLLAR 9.37. Sei $A \in M_{n \times n}(K)$, $S \in GL_n(K)$. Dann gilt $\text{Spur}(A) = \text{Spur}(SAS^{-1})$, d.h. zueinander konjugierte Matrizen haben dieselbe Spur.

BEWEIS. Nach dem Satz gilt $\text{Spur}(SAS^{-1}) = \text{Spur}(S^{-1}SA) = \text{Spur}(A)$. $\square$

Das bedeutet, dass für zueinander konjugierte Matrizen A und B gilt: $\text{Spur}(A^j) = \text{Spur}(B^j)$ für alle $j \geq 0$. (Denn wenn $B = SAS^{-1}$, dann gilt $B^j = SA^jS^{-1}$ für alle $j \geq 0$.) Das kann man manchmal benutzen, um zu zeigen, dass zwei gegebene Matrizen A und B nicht zueinander konjugiert sind: Es genügt, ein j zu finden, so dass A^j und B^j nicht dieselbe Spur haben.

Weil konjugierte Matrizen dieselbe Spur haben, gilt für jeden Endomorphismus f eines endlich-dimensionalen K -Vektorraums V , dass die Spur der Matrix $M_{\mathcal{B}}^{\mathcal{B}}(f)$ unabhängig ist von der Wahl der Basis $\mathcal{B}$. Daher ist die folgende Definition sinnvoll.

DEFINITION 9.38. Seien K ein Körper und f ein Endomorphismus des endlich-dimensionalen K -Vektorraums V . Sei $\mathcal{B}$ eine Basis von V . Dann heißt $\text{Spur}(f) := \text{Spur}(M_{\mathcal{B}}^{\mathcal{B}}(f))$ auch die *Spur des Endomorphismus f* . $\dashv$

9.5. Ergänzungen *

BEMERKUNG 9.39 (Eindeutigkeit der Determinante). Ergänzend hier noch einige zusätzliche Anmerkungen dazu, inwiefern die Determinante »eindeutig bestimmt« ist bzw. eine Abbildung $M_n(K) \rightarrow K$ ist, die sich in »natürlicher Weise« ergibt. Um sich zu überzeugen, dass es sich lohnt, die Determinante in die Theorie der linearen Abbildungen mit aufzunehmen, benötigt man diese Überlegungen natürlich nicht. Die Determinante wird uns beim weiteren Aufbau der Theorie an verschiedenen Stellen nutzen. Natürlich ist die Determinante auch nicht eindeutig bestimmt als Abbildung $M_n(K) \rightarrow K$ – es ist ja klar, dass es viele andere Abbildungen $M_n(K) \rightarrow K$ gibt (sowohl interessante wie die Spurabbildung aus Abschnitt 9.4, als auch viele uninteressante). Es ist also eher die Frage, durch welche ihrer Eigenschaften man die Determinante eindeutig charakterisieren kann. Wir haben gesehen, dass die Determinante die eindeutig bestimmte Abbildung $M_n(K) \rightarrow K$ ist, die alternierend und multilinear in den Spalten ist und die Einheitsmatrix auf das Element $1 \in K$ abbildet. Aus den weiteren Überlegungen (insbesondere Satz 9.17) ergibt sich, dass sie auch die einzige solche Abbildung ist, die alternierend und multilinear in den Zeilen ist, und die Einheitsmatrix auf 1 abbildet. Eine andersartige Charakterisierung erlaubt nun der folgende Satz, in dem wir zunächst über die Einschränkung von $\det$ auf den Definitionsbereich $GL_n(K)$ und Wertebereich $K^\times$ sprechen; vergleiche aber auch das dem Satz folgende Lemma.

SATZ 9.40. Einer invertierbaren Matrix ihre Determinante zuzuordnen ist der eindeutig bestimmte Gruppenhomomorphismus $GL_n(K) \rightarrow K^\times$, der jede Diagonalmatrix der Form $\text{diag}(1, \dots, 1, a)$, $a \in K^\times$, auf a abbildet.

Mit anderen Worten: Ist $\Delta: GL_n(K) \rightarrow K^\times$ eine Abbildung mit $\Delta(AB) = \Delta(A)\Delta(B)$ für alle $A, B \in GL_n(K)$ und mit $\Delta(\text{diag}(1, \dots, 1, a)) = a$ für alle $a \in K^\times$, so gilt $\Delta(A) = \det(A)$ für alle $A \in GL_n(K)$; und die Determinante hat diese beiden Eigenschaften.

BEWEIS. Siehe Hausaufgabe 11.4. □

Mit Methoden, die über das hinausgehen, was uns im Moment zur Verfügung steht, kann man das folgende Lemma beweisen. Es zeigt über unendlichen Körpern und unter der vernünftigen Zusatzvoraussetzung, dass die Determinante als ein polynomialer Ausdruck in den Koeffizienten der betrachteten Matrix gegeben sei, dass die Determinante für invertierbare Matrizen schon die Werte auf allen Matrizen festlegt.

LEMMA 9.41. Seien K ein unendlicher Körper und $\Delta: M_n(K) \rightarrow K$ eine Abbildung, wobei $\Delta(A)$ durch einen polynomialen Ausdruck in den Einträgen von A gegeben sei. Es gelte $\Delta(A) = \det(A)$ für alle $A \in GL_n(K)$.

Dann gilt $\Delta(A) = \det(A)$.

Um das Lemma zu beweisen, benutzt man ein Argument aus der Topologie, nämlich, dass die Menge der invertierbaren Matrizen in der Menge aller Matrizen »dicht« liegt. Über einem allgemeinen Körper verwendet man dafür die sogenannte Zariski-Topologie, die in der algebraischen Geometrie eingeführt wird. Über dem Körper der reellen Zahlen (und ähnlich über dem Körper der komplexen Zahlen) kann man die üblichen Begriffe des Grenzwerts und der Stetigkeit benutzen. Die Menge $GL_n(\mathbb{R})$ ist dicht in $M_n(\mathbb{R})$; das bedeutet, dass es zu jeder Matrix $A \in M_n(K)$ eine Folge $(A_i)_{i \geq 0}$ invertierbarer Matrizen mit $\lim_{i \rightarrow \infty} A_i = A$ gilt. Der Limes von Matrizen soll hier einfach eintragsweise gebildet werden. Wir lassen den Beweis hier aus. Die Voraussetzung an Δ impliziert, dass diese Abbildung ebenso wie die Determinante in den Einträgen der Matrix stetig ist, es gilt also

$$\Delta(A) = \lim_{i \rightarrow \infty} \Delta(A_i) = \lim_{i \rightarrow \infty} \det(A_i) = \det(A)$$

wegen der Voraussetzung, dass Δ und $\det$ für invertierbare Matrizen übereinstimmen. ◇

BEMERKUNG 9.42. Hierher passt auch gut eine Erläuterung, wie sich die Leibniz-Formel aus den Eigenschaften der Determinante sozusagen zwingend ergibt, wie man also »auf die Formel kommen kann«. Natürlich wissen wir schon, dass die Determinante die *eindeutig* bestimmte in den Spalten alternierende multilineare Abbildung $M_n(K) \rightarrow K$ ist, die E_n auf 1 abbildet, und deshalb mit dem durch die Leibniz-Formel gegebenen Ausdruck übereinstimmen muss. Man kann den Zusammenhang folgendermaßen noch expliziter herstellen und, wenn man möchte, auch den Existenzbeweis für die Determinante führen, ohne die Leibniz-Regel zu benutzen.

Wir benutzen dafür wieder die Notation, die wir in Abschnitt 9.3 definiert haben, um den Laplaceschen Entwicklungssatz zu formulieren. Sei $A \in M_n(K)$ und A_{ij} also die Matrix, die aus A entsteht, indem die i -te Zeile und j -te Spalte durch die jeweiligen (transponierten) Standardbasisvektoren ersetzt werden. Sei A_{ij}° die Matrix, die aus A entsteht, indem die j -te Zeile durch den Standardbasisvektor e_i ersetzt wird. Die anderen Einträge der i -ten Zeile sollen nicht verändert werden. Dann lässt sich A_{ij}° durch elementare Spaltenumformungen vom Typ I in die Matrix A_{ij} umformen; es folgt $\Delta(A_{ij}^\circ) = \Delta(A_{ij})$ für jede Determinantenfunktion $\Delta: M_n(K) \rightarrow K$.

Mit dieser Bemerkung erhalten wir die Aussage des Laplaceschen Satzes für die j -te Spalte für jede Determinantenfunktion Δ , also

$$\Delta(A) = \sum_{i=1}^n a_{ij} \Delta(A_{ij}^{\circ}) = \sum_{i=1}^n a_{ij} \Delta(A_{ij}),$$

wobei wir im ersten Schritt die Multilinearität von Δ , konkret die Linearität in der j -ten Spalte benutzen.

Die Leibniz-Formel wird in diesem Argument nicht benötigt. Stattdessen können wir nun die Leibniz-Formel aus dem Entwicklungssatz induktiv »herleiten«. Nehmen wir also an, dass wir die Leibniz-Formel für Determinanten der Größe $(n-1) \times (n-1)$ bereits kennen. Wir erhalten durch Entwicklung nach der letzten Spalte, dass

$$\det(A) = \sum_{i=1}^n (-1)^{i+n} a_{in} \det(A'_{in}) = \sum_{i=1}^n (-1)^{i+n} a_{in} \sum_{\sigma \in S_{n-1}} \operatorname{sgn}(\sigma) \prod_{k=1}^{n-1} a_{k(i), \sigma(k)}$$

wobei wir (nur an dieser einen Stelle) $k^{(i)} = \begin{cases} k & k < i \\ k+1 & k \geq i \end{cases}$ setzen. Die letzte Summe können wir mit der Leibniz-Formel für $\det(A)$ identifizieren, indem wir die entsprechende Überlegung im Beweis von Satz 9.29 in die andere Richtung als dort durchführen.

Wie schon bemerkt, kann man in dieser Weise auch die Existenz einer Determinantenfunktion Δ mit $\Delta(E_n) = 1$ durch Induktion nach n zeigen. Insbesondere wird die Leibniz-Formel an keiner Stelle zwingend für den Aufbau der Theorie benötigt. $\diamond$

BEISPIEL 9.43. Die Grundidee des »global positioning system« GPS, das Navigationssysteme benutzen, um den aktuellen Aufenthaltsort festzustellen, ist die folgende (wir ignorieren mehrere wichtige Verbesserungen und »Details«):

GPS basiert darauf, dass Satelliten die Erde umkreisen und es zu jedem Zeitpunkt möglich ist, anhand der Signale, die diese Satelliten aussenden, für mindestens drei Satelliten S_i , $i = 1, 2, 3$, die Position $(a_i, b_i, c_i)^t \in \mathbb{R}^3$ und die Entfernung d_i von S_i zum Punkt P , dessen Position bestimmt werden soll, zu ermitteln. (Wie man die Satelliten entsprechende Signale aussenden lässt, hat ebenfalls mehrere interessante mathematische Aspekte. Es kann ja keine Kommunikation in die andere Richtung stattfinden.)

Das bedeutet, dass die gesuchten Koordinaten x, y, z des Punktes P die folgenden Gleichungen erfüllen. (Siehe Ergänzung 7.59 und Abschnitt II.2.1.)

$$\begin{aligned} (x - a_1)^2 + (y - b_1)^2 + (z - c_1)^2 &= d_1^2 \\ (x - a_2)^2 + (y - b_2)^2 + (z - c_2)^2 &= d_2^2 \\ (x - a_3)^2 + (y - b_3)^2 + (z - c_3)^2 &= d_3^2. \end{aligned}$$

Auch wenn diese quadratischen Gleichungen nicht direkt im Einflußbereich der linearen Algebra liegen, können wir uns folgendermaßen behelfen. Wir ziehen die dritte Gleichung von der ersten ab, und ziehen die dritte Gleichung von der zweiten Gleichung ab. Dadurch erhalten wir das äquivalente Gleichungssystem

$$\begin{aligned} 2(a_3 - a_1)x + 2(b_3 - b_1)y + 2(c_3 - c_1)z &= d_1^2 - d_3^2 + (a_3 - a_1)^2 + (b_3 - b_1)^2 + (c_3 - c_1)^2 \\ 2(a_3 - a_2)x + 2(b_3 - b_2)y + 2(c_3 - c_2)z &= d_2^2 - d_3^2 + (a_3 - a_2)^2 + (b_3 - b_2)^2 + (c_3 - c_2)^2 \\ (x - a_3)^2 + (x - b_3)^2 + (x - c_3)^2 &= d_3^2. \end{aligned}$$

Eine weitere Eigenschaft des Systems ist, dass die drei Satelliten S_1, S_2, S_3 sich niemals gleichzeitig auf einer Geraden befinden. Mit anderen Worten: Nach der Verschiebung, die S_1 in den Ursprung verschiebt, befinden sich S_2 und S_3 auf verschiedenen Ursprungsgeraden. Die Matrix

$$\begin{pmatrix} a_2 - a_1 & a_3 - a_1 \\ b_2 - b_1 & b_3 - b_1 \\ c_2 - c_1 & c_3 - c_1 \end{pmatrix}$$

hat also Rang 2. Daher (Zeilenrang=Spaltenrang) hat wenigstens eine der Matrizen

$$\begin{pmatrix} a_2 - a_1 & a_3 - a_1 \\ b_2 - b_1 & b_3 - b_1 \end{pmatrix}, \quad \begin{pmatrix} a_2 - a_1 & a_3 - a_1 \\ c_2 - c_1 & c_3 - c_1 \end{pmatrix}, \quad \begin{pmatrix} b_2 - b_1 & b_3 - b_1 \\ c_2 - c_1 & c_3 - c_1 \end{pmatrix}$$

eine Determinante $\neq 0$. Nehmen wir ohne Einschränkung an, dass die erste dieser Matrizen eine nicht-verschwindende Determinante hat.

Wir schreiben die ersten beiden Gleichungen des obigen Gleichungssystems um als

$$\begin{aligned} 2(a_3 - a_1)x + 2(b_3 - b_1)y &= -2(c_3 - c_1)z + d_1^2 + (a_3 - a_1)^2 + (b_3 - b_1)^2 + (c_3 - c_1)^2 =: r_1(z) \\ 2(a_3 - a_2)x + 2(b_3 - b_2)y &= -2(c_3 - c_2)z + d_1^2 + (a_3 - a_2)^2 + (b_3 - b_2)^2 + (c_3 - c_2)^2 =: r_2(z) \end{aligned}$$

und betrachten das Ergebnis als *lineares* Gleichungssystem in den Unbestimmten x und y . Unsere Voraussetzung, dass die Koeffizientenmatrix dieses Gleichungssystems invertierbar ist, hat zur Folge, dass wir die eindeutig bestimmte Lösung (für fixiertes z) mit der Cramerschen Regel (hier genügen auch die Formeln aus Abschnitt 2.5) schreiben können als

$$x = \frac{\det \begin{pmatrix} r_1(z) & 2(b_3 - b_1)y \\ r_2(z) & 2(b_3 - b_2)y \end{pmatrix}}{\det \begin{pmatrix} 2(a_2 - a_1) & 2(a_3 - a_1) \\ 2(b_2 - b_1) & 2(b_3 - b_1) \end{pmatrix}}, \quad y = \frac{\det \begin{pmatrix} 2(a_3 - a_1) & r_1(z) \\ 2(a_3 - a_1) & r_2(z) \end{pmatrix}}{\det \begin{pmatrix} 2(a_2 - a_1) & 2(a_3 - a_1) \\ 2(b_2 - b_1) & 2(b_3 - b_1) \end{pmatrix}}.$$

Wir erhalten damit Ausdrücke für x und y , die wir als Funktionen im Parameter z verstehen und in die dritte Gleichung unseres Gleichungssystems einsetzen können. Das liefert eine quadratische Gleichung für z .

Diese Gleichung hat zwei Lösungen für z , und durch Einsetzen in die soeben hergeleiteten Formel für x und y erhalten wir zwei Punkte als Lösung des aufgestellten Gleichungssystems. Eine entspricht dem gesuchten Punkt, die andere einem Punkt, der »offensichtlich« nicht auf der Erdoberfläche liegen kann und daher verworfen werden kann.

Siehe C. Rousseau, Y. Saint-Aubin, *Mathematik und Technologie*, Springer 2012, <https://doi.org/10.1007/978-3-642-30092-9> für viele weitere interessante Ausführungen zu GPS. $\diamond$

ERGÄNZUNG 9.44 (Quaternionen und Drehungen). An dieser Stelle möchte ich bei Gelegenheit noch etwas über den Zusammenhang zwischen den Quaternionen und Drehungen von $\mathbb{R}^3$ und $\mathbb{R}^4$ ergänzen. Für den Moment zwei Referenzen:

[Wikipedia \(Englisch\)](#)³

W. Klingenberg, *Lineare Algebra und Geometrie* [K1], Abschnitt 8.4.

□ Ergänzung 9.44

Auf der Webseite mathe-vital.de^a sind mehrere geometrisch-anschauliche »Anwendungen« der Determinante erklärt.

^a<https://mathe-vital.de/LinAlg1/18-topic.html>

³https://en.wikipedia.org/wiki/Quaternion_and_spatial_rotation

Eigenwerte

Seien K ein Körper und V ein endlichdimensionaler K -Vektorraum. Wir greifen nun noch einmal die Frage auf, wie man Endomorphismen $f: V \rightarrow V$ gut verstehen bzw. gut beschreiben kann.

Ein Beispiel für eine besonders einfache Situation ist die, dass der Endomorphismus durch eine Diagonalmatrix beschrieben werden kann, also im Sinne der folgenden Definition *diagonalisierbar* ist.

DEFINITION 10.1. Seien K ein Körper und V ein endlichdimensionaler K -Vektorraum. Ein Endomorphismus f von V heißt *diagonalisierbar*, wenn eine Basis $\mathcal{B}$ von V existiert, so dass $M_{\mathcal{B}}^{\mathcal{B}}(f)$ eine Diagonalmatrix ist. $\dashv$

Es ist entscheidend, dass hier *nur eine* Basis von V verwendet wird. (Es ist immer möglich, Basen $\mathcal{B}$ und $\mathcal{C}$ von V zu finden, so dass $M_{\mathcal{C}}^{\mathcal{B}}(f)$ eine Diagonalmatrix ist, siehe Satz 7.36; das ist also weit weniger interessant.) Dass f diagonalisierbar ist, bedeutet geometrisch, dass eine Basis (ein »Koordinatensystem«) existiert, so dass f eine Streckung entlang der Koordinatenachsen ist (möglicherweise mit unterschiedlichen Streckfaktoren). Siehe Beispiel 7.37.

Analog haben wir den Begriff der Diagonalisierbarkeit von Matrizen.

DEFINITION 10.2. Eine Matrix $A \in M_{n \times n}(K)$ heißt *diagonalisierbar*, wenn A zu einer Diagonalmatrix konjugiert ist, d.h. wenn $S \in GL_n(K)$ existiert, so dass SAS^{-1} eine Diagonalmatrix ist. $\dashv$

Es ist also A genau dann diagonalisierbar, wenn $f_A: K^n \rightarrow K^n, x \mapsto Ax$, diagonalisierbar ist. Ein Endomorphismus $f: V \rightarrow V$ eines endlichdimensionalen Vektorraums V ist genau dann diagonalisierbar, wenn für irgendeine (äquivalent: für jede) Basis $\mathcal{B}$ die Matrix $M_{\mathcal{B}}^{\mathcal{B}}(f)$ diagonalisierbar ist. Wir sehen auch, dass es für zueinander konjugierte Matrizen A und B äquivalent ist, dass einerseits A und andererseits B diagonalisierbar ist.

Dass Diagonalmatrizen einfacher zu behandeln sind als »allgemeine« Matrizen, ist recht offensichtlich. Aber auch die Diagonalisierbarkeit ist eine Eigenschaft, die unter anderem einige konkrete Rechnungen einfacher macht – zum Beispiel, Potenzen einer Matrix zu berechnen. Denn für Diagonalmatrizen ist das einfach, und ist $D = SAS^{-1}$ eine Diagonalmatrix, so gilt $A = S^{-1}DS$ und damit $A^n = S^{-1}D^nS$. In Beispiel 10.19 leiten wir mithilfe dieser Beobachtung nochmals eine explizite Formel für die n -te Fibonacci-Zahl her. Siehe Ergänzung 10.28 für eine Skizze, wie man ein System linearer gewöhnlicher Differentialgleichungen mit einer diagonalisierbaren konstanten Koeffizientenmatrix löst.

10.1. Eigenwerte, Eigenvektoren

Um Hilfsmittel zu entwickeln, die uns erlauben zu entscheiden, ob ein Vektorraum-Endomorphismus oder eine Matrix diagonalisierbar sind, und letztlich auch nicht-diagonalisierbare Endomorphismen besser verstehen zu können, nähern wir uns der Sache schrittweise. Ist $f: V \rightarrow V$ ein diagonalisierbarer Endomorphismus und $\mathcal{B}$ eine Basis von V , so dass $M_{\mathcal{B}}^{\mathcal{B}}(f)$

eine Diagonalmatrix ist (Satz 7.36), so gilt für jeden der Basisvektoren b_i aus $\mathcal{B}$, dass $f(b_i)$ ein Vielfaches von b_i ist. Dieser wichtigen Eigenschaft geben wir in der folgenden Definition einen Namen.

DEFINITION 10.3. Seien K ein Körper, V ein K -Vektorraum und f ein Endomorphismus von V .

- (1) Ein Vektor $v \in V \setminus \{0\}$ heißt *Eigenvektor von f zum Eigenwert $\lambda \in K$* , falls $f(v) = \lambda v$. Ein Element $\lambda \in K$ heißt *Eigenwert der linearen Abbildung f* , falls ein Vektor $v \in V \setminus \{0\}$ existiert, der Eigenvektor von f zum Eigenwert λ ist.
- (2) Ist λ ein Eigenwert von f , so heißt die Menge

$$V_\lambda = V_\lambda(f) = \{v \in V; f(v) = \lambda v\}$$

aller Eigenvektoren zum Eigenwert λ zusammen mit dem Nullvektor der *Eigenraum* von f zum Eigenwert λ . (Dies ist ein Untervektorraum von V .)

—

Gelegentlich verwenden wir die Schreibweise $V_\lambda(f) = \{v \in V; f(v) = \lambda v\}$ auch dann, wenn λ kein Eigenwert von f sei (und dann handelt es sich einfach um den Nullvektorraum).

Dass die hier definierten Eigenräume tatsächlich Untervektorräume sind, prüft man unmittelbar nach, indem man ausnutzt, dass die Abbildung f linear ist. Dass $v \in V \setminus \{0\}$ ein Eigenvektor eines Endomorphismus f von V sei, können wir auch dadurch ausdrücken, dass f die Gerade (d.h. den eindimensionalen Unterraum) $\langle v \rangle$ in sich abbildet. Die Eigenvektoren von f zum Eigenwert 0 sind genau die Elemente von $\text{Ker}(f) \setminus \{0\}$. Ein geometrisches Beispiel: Ist $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ eine Drehung, also eine abstandserhaltende Abbildung mit $\det(f) = 1$ (siehe Ergänzung 9.24), und ist v ein Eigenvektor von f zum Eigenwert 1, so ist die von v erzeugte Gerade eine/die Drehachse der Drehung. In Beispiel 10.11 sehen wir, dass jede Drehung von $\mathbb{R}^3$ einen Eigenvektor vom Eigenwert 1, also eine Drehachse, besitzt.

Der Zusammenhang zum Begriff der Diagonalisierbarkeit ist der folgende.

SATZ 10.4. Ein Endomorphismus f eines endlichdimensionalen Vektorraums V ist genau dann diagonalisierbar, wenn eine Basis $\mathcal{B}$ von V existiert, die aus Eigenvektoren von f besteht.

BEWEIS. Dass eine Basis $\mathcal{B}$ aus Eigenvektoren von f besteht, ist offenbar dazu äquivalent, dass $M_{\mathcal{B}}^{\mathcal{B}}(f)$ eine Diagonalmatrix ist. Daraus folgt direkt die Behauptung. $\square$

Analog definieren wir wieder die Begriffe des Eigenwerts, Eigenvektors und Eigenraums einer quadratischen Matrix A . Es handelt sich dabei genau um die Eigenwerte, Eigenvektoren und Eigenräume des Endomorphismus f_A von K^n , der durch $v \mapsto Av$ definiert ist. Explizit ausgeschrieben bedeutet das:

DEFINITION 10.5. Sei $n \in \mathbb{N}$ und $A \in M_n(K)$.

- (1) Ein Vektor $v \in K^n \setminus \{0\}$ heißt *Eigenvektor von A zum Eigenwert $\lambda \in K$* , falls $Av = \lambda v$. Ein Element $\lambda \in K$ heißt *Eigenwert der Matrix A* , falls ein Vektor $v \in K^n \setminus \{0\}$ existiert, der Eigenvektor von A zum Eigenwert λ ist.
- (2) Ist λ ein Eigenwert von A , so heißt die Menge

$$V_\lambda = V_\lambda(A) = \{v \in K^n; Av = \lambda v\}$$

aller Eigenvektoren zum Eigenwert λ zusammen mit dem Nullvektor der *Eigenraum* von A zum Eigenwert λ . (Dies ist ein Untervektorraum von K^n .)

—

BEMERKUNG 10.6. Die Einführung der Begriffe *Eigen*-wert, -vektor, ... schreibt man [David Hilbert](#)¹ (1862 – 1943, einer der wichtigsten Mathematiker seiner Zeit) zu, der in der Einleitung seiner Arbeit *Grundzüge einer allgemeinen Theorie der linearen Integralgleichungen*, Nachr. Königl. Ges. der Wissensch. zu Göttingen (Math.-physik. Kl.) (1904), 49–91, schreibt: »...gelange ich zu Formeln, die die Entwicklung einer willkürlichen Funktion nach gewissen ausgezeichneten Funktionen, die ich Eigenfunktionen nenne, liefern ...«. Vielleicht war er vom Begriff des Eigentons aus der Akustik beeinflusst. Jedenfalls geht es darum, Eigenschaften eines »Systems« (für uns: einer linearen Abbildung) zu finden bzw. zu benennen, die besonders charakteristisch dafür und sozusagen »diesem System eigen« sind. Auch das Wort *Eigenschaft* beginnt ja, aus dem gleichen Grund, mit diesem Wortteil.

Viele Anwendungsprobleme lassen sich als Eigenwertprobleme formulieren, d.h. als die Aufgabenstellung, Eigenwerte und Eigenvektoren einer linearen Abbildung zu finden, unter anderem in der Physik, Chemie, Biologie, in den Ingenieurwissenschaften, den Wirtschaftswissenschaften und anderen Gebieten. Siehe den Abschnitt *Praktische Beispiele* auf der Wikipedia-Seite zum Begriff [Eigenwertproblem](#)². Siehe auch [Wikipedia \(Englisch\)](#)³. Siehe Abschnitt 10.3 für einige konkrete Beispiele.

Oft handelt es sich bei den dabei betrachteten Vektorräumen um Vektorräume von Funktionen. Als einfaches Beispiel betrachten wir den Vektorraum $C^\infty(\mathbb{R})$ aller beliebig oft differenzierbaren Funktionen $\mathbb{R} \rightarrow \mathbb{R}$, und darauf die Ableitung $D: C^\infty(\mathbb{R}) \rightarrow C^\infty(\mathbb{R})$ als Endomorphismus. Aus den Ableitungsregeln (hier benötigen wir nur $(f + g)' = f' + g'$ und $(af)' = af'$ für differenzierbare Funktionen f, g und $a \in \mathbb{R}$) folgt, dass D eine lineare Abbildung ist. Die Exponentialfunktion $\exp$ ist ein Eigenvektor von D mit Eigenwert 1, denn es gilt $D(\exp) = \exp' = \exp$. Die Sinusfunktion ist ein Eigenvektor von der Verkettung $D \circ D$, also der zweifachen Ableitung, und zwar mit Eigenwert -1 .

”

Physics is very interesting. There are many, many interesting theorems. Unfortunately, there are no definitions.

David Kazhdan

...it is impossible to explain honestly the beauties of the laws of nature in a way that people can feel, without their having some deep understanding of mathematics. I am sorry, but this seems to be the case.

Richard Feynman

Fundort: <https://www.jmilne.org/math/>

◇

Für einen Endomorphismus eines endlichdimensionalen K -Vektorraums V stimmen die Eigenwerte von f und der darstellenden Matrix $M_{\mathcal{B}}^{\mathcal{B}}(f)$ bezüglich irgendeiner Basis $\mathcal{B}$ überein, wie das folgende Lemma zeigt.

LEMMA 10.7. (I) Seien V, W Vektorräume über dem Körper K , f ein Endomorphismus von V und $\iota: V \rightarrow W$ ein Isomorphismus. Dann ist $g := \iota \circ f \circ \iota^{-1}$ ein Endomorphismus von W . Es haben f und g dieselben Eigenwerte, und ist $\lambda \in K$ ein Eigenwert von f und g , so induziert ι einen Isomorphismus $V_\lambda(f) \xrightarrow{\sim} V_\lambda(g)$. Insbesondere gilt $\dim V_\lambda(f) = \dim V_\lambda(g)$.

¹https://de.wikipedia.org/wiki/David_Hilbert

²<https://de.wikipedia.org/wiki/Eigenwertproblem>

³https://en.wikipedia.org/wiki/Eigenvalues_and_eigenvectors#Applications

- (2) Sei f ein Endomorphismus des endlichdimensionalen K -Vektorraums V und sei $\mathcal{B}$ eine Basis von V . Genau dann ist ein Element $\lambda \in K$ ein Eigenwert von f , wenn λ ein Eigenwert der Matrix $M_{\mathcal{B}}^{\mathcal{B}}(f)$ ist, und in diesem Fall induziert der Koordinatenisomorphismus $c_{\mathcal{B}}: V \rightarrow K^n$ einen Isomorphismus zwischen den zugehörigen Eigenräumen von f und $M_{\mathcal{B}}^{\mathcal{B}}(f)$. Insbesondere haben die Eigenräume zu λ dieselbe Dimension.
- (3) Zueinander konjugierte Matrizen $A \in M_n(K)$, SAS^{-1} (mit $S \in GL_n(K)$) haben dieselben Eigenwerte. Die Eigenräume von A und SAS^{-1} zu einem gemeinsamen Eigenwert λ haben dieselben Dimensionen.

BEWEIS. zu (1). Ist $v \in V$, $v \neq 0$, mit $f(v) = \lambda v$, so folgt

$$g(\iota(v)) = \iota(f(\iota^{-1}(\iota(v)))) = \iota(f(v)) = \lambda \iota(v),$$

mithin ist $\iota(v)$ ein Eigenvektor von g und insbesondere λ ein Eigenwert von g . Genauer folgt, dass ι eine Abbildung $V_{\lambda}(f) \rightarrow V_{\lambda}(g)$ induziert. Analog induziert ι^{-1} eine Abbildung $V_{\lambda}(g) \rightarrow V_{\lambda}(f)$. Die Verkettung dieser beiden Abbildungen ist die Identität, und damit ist Teil (1) bewiesen.

zu (2). Die Behauptung folgt direkt aus Teil (1), wenn man für ι den Koordinatenisomorphismus $c_{\mathcal{B}}$ hernimmt.

zu (3). Dies können wir als Folgerung aus Teil (2) erhalten, denn daran sehen wir, dass die Eigenwerte der darstellenden Matrix eines Endomorphismus von der Wahl der Basis unabhängig sind. Weil A und SAS^{-1} Matrizen sind, die beide dieselbe lineare Abbildung $K^n \rightarrow K^n$ (nämlich, beispielsweise, f_A) bezüglich unterschiedlicher Basen darstellen, folgt die Behauptung. $\square$

Aus dem Lemma und den uns bekannten Charakterisierungen der Eigenschaft eines Endomorphismus bzw. einer Matrix, (nicht) invertierbar zu sein, erhalten wir den folgenden Satz.

SATZ 10.8 (Charakterisierung von Eigenwerten). Seien K ein Körper und V ein endlichdimensionaler K -Vektorraum. Sei $f: V \rightarrow V$ ein Endomorphismus, $n = \dim V$, $\mathcal{B}$ eine Basis von V , $A = M_{\mathcal{B}}^{\mathcal{B}}(f)$. Sei $\lambda \in K$. Die folgenden Aussagen sind äquivalent:

- (i) λ ist Eigenwert von f ,
- (ii) λ ist Eigenwert von A ,
- (iii) $\text{Ker}(f - \lambda \cdot \text{id}_V) \neq \{0\}$,
- (iv) $\text{Ker}(A - \lambda E_n) \neq \{0\}$,
- (v) $\det(f - \lambda \cdot \text{id}_V) = 0$,
- (vi) $\det(A - \lambda E_n) = 0$.

In diesem Fall ist $\text{Ker}(f - \lambda \cdot \text{id}_V)$ der Eigenraum von f zum Eigenwert λ und $\text{Ker}(A - \lambda E_n)$ der Eigenraum von A zum Eigenwert λ .

KOROLLAR 10.9. Die Eigenwerte einer oberen Dreiecksmatrix sind genau die Diagonaleinträge der Matrix.

BEWEIS. Eine obere Dreiecksmatrix ist genau dann invertierbar, wenn alle ihre Diagonaleinträge $\neq 0$ sind. Für eine obere Dreiecksmatrix A ist also $\det(A - \lambda E_n) = 0$ dazu äquivalent, dass λ auf der Diagonale von A vorkommt. Die Behauptung folgt nun aus Satz 10.8. $\square$

Im allgemeinen Fall ist die Bestimmung der Eigenwerte (und Eigenräume) einer Matrix aufwändiger. Die Eigenwerte von A sind die Nullstellen der Funktion $K \rightarrow K$, $\lambda \mapsto \det(A - \lambda E_n)$. Die Leibniz-Formel zeigt, dass es sich um eine Polynomfunktion in λ , die höchste

Potenz von λ , die hier auftritt ist n . Für $n = 2$ lassen sich die Nullstellen durch Lösen einer quadratischen Gleichung bestimmen, aber für größere n gibt es keine praktikablen (und für $n \geq 5$: überhaupt keine) allgemeinen Lösungsformeln. In konkret gegebenen Fällen lassen sich die Nullstellen aber oft bestimmen. Äquivalent kann man stattdessen mit $\det(\lambda E_n - A)$ arbeiten; das hat den (kleinen) Vorteil, dass dann λ^n mit dem Koeffizient 1 auftritt.

Der Eigenraum $\text{Ker}(A - \lambda E_n)$ zum Eigenwert λ ist die Lösungsmenge des homogenen linearen Gleichungssystems mit Koeffizientenmatrix $A - \lambda E_n$. Ist λ bekannt, so lässt sich also mit dem Gauß-Algorithmus eine Basis bestimmen. Die Dimension zu errechnen, ist wegen $\dim(\text{Ker}(A - \lambda E_n)) = n - \text{rg}(A - \lambda E_n)$ dazu gleichwertig, den Rang dieser Matrix zu finden.

Mit dem Satz können wir auch beweisen, dass eine Matrix A und ihre transponierte Matrix dieselben Eigenwerte haben.

LEMMA 10.10. Sei K ein Körper und sei $A \in M_{n \times n}(K)$ eine quadratische Matrix. Die Matrizen A und A^t haben dieselben Eigenwerte, und ist $\lambda \in K$ ein Eigenwert, so gilt $\dim V_\lambda(A) = \dim V_\lambda(A^t)$.

BEWEIS. Dass $\lambda \in K$ ein Eigenwert von A ist, ist dazu äquivalent, dass $\text{Ker}(A - \lambda E_n) \neq 0$ ist, und in diesem Fall ist $\text{Ker}(A - \lambda E_n)$ der zugehörige Eigenraum. Es gilt aber $(A - \lambda E_n)^t = A^t - \lambda E_n$, und wegen Theorem 7.41 (Zeilenrang und Spaltenrang einer Matrix stimmen überein) gilt $\dim(\text{Ker}(A - \lambda E_n)) = \dim(\text{Ker}(A^t - \lambda E_n))$. $\square$



Der Arme ist unter die Dichter gegangen. Für die Mathematik hatte er nicht genug Fantasie!

... hat angeblich D. Hilbert gesagt, als ihm berichtet wurde, dass einer seiner früheren Studenten von der Mathematik in die Germanistik gewechselt sei.

BEISPIEL 10.11. Wir haben in Beispiel 7.37 schon ein Beispiel einer diagonalisierbaren Abbildung (bzw. Matrix) gesehen. Einfache Klassen von weiteren Beispielen sind die folgenden.

- (1) Sei $K = \mathbb{R}$. Für reelle Zahlen $-1 < a, b < 1$ mit $a^2 + b^2 = 1$ hat die Matrix $\begin{pmatrix} a & -b \\ b & a \end{pmatrix}$ keine Eigenwerte, ist also insbesondere nicht diagonalisierbar. Das stimmt mit der geometrischen Anschauung überein, dass die zugehörige lineare Abbildung eine Drehung um den Ursprung ist (Ergänzung 7.59). Diese Drehung bildet keine Ursprungsgerade auf sich selbst ab, weil wir die Drehungen um 0° und um 180° ausgeschlossen haben (sie entsprechen den Werten $a = 1$ bzw. $a = -1$).
- (2) Sei $K = \mathbb{R}$, und $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ eine Drehung, d.h. $\det(f) = 1$ und für die darstellende Matrix $A = M(f)$ gilt $A^{-1} = A^t$ (Ergänzung 9.24).

Es gilt dann, weil $\det(A) = 1$ und die Determinante einer Matrix mit der Determinante ihrer Transponierten übereinstimmt,

$$\det(A - E_n) = \det(A - AA^t) = \det(A) \det(E_n - A^t) = \det(E_n - A) = -\det(A - E_n).$$

Im letzten Schritt haben wir dabei ausgenutzt, dass $n = 3$ ungerade ist.

Es folgt, dass $\det(A - E_n) = 0$, somit dass 1 ein Eigenwert von f ist, es gibt also einen Vektor $v \in \mathbb{R}^3$ mit $f(v) = v$. Die von v erzeugte Gerade kann man sich anschaulich als Drehachse vorstellen, da ihre Elemente sämtlich festgelassen werden.

Aus unserer Definition von Drehungen folgt leicht, dass die Verkettung von zwei (oder mehr) Drehungen wieder eine Drehung ist. Insbesondere hat die Verkettung von zwei Drehungen wieder eine Drehachse.

- (3) Sei $K = \mathbb{R}$, n eine ungerade natürliche Zahl und f ein Endomorphismus eines n -dimensionalen $\mathbb{R}$ -Vektorraums. Dann hat f einen Eigenwert in $\mathbb{R}$. Denn die Abbildung $\varphi: \mathbb{R} \rightarrow \mathbb{R}$, $\lambda \mapsto \det(A - \lambda E_n)$, ist eine Polynomfunktion vom Grad n , wie wir anhand der Leibniz-Formel sehen. Jede solche Polynomfunktion (von *ungeradem* Grad!) hat mindestens eine Nullstelle. Um dies zu zeigen, kann und muss man *analytische* Eigenschaften des Körpers der reellen Zahlen benutzen, etwa den Zwischenwertsatz. Daraus folgt die Behauptung, weil $\lim_{x \rightarrow -\infty} \varphi(x) = \infty$ und $\lim_{x \rightarrow \infty} \varphi(x) = -\infty$ gilt.

Ist λ eine Nullstelle von φ , so ist $\det(A - \lambda E_n) = 0$, also λ ein Eigenwert von f .

- (4) Sei nun K irgendein Körper und $A = \begin{pmatrix} a & b \\ 0 & d \end{pmatrix}$ eine obere Dreiecksmatrix der Größe 2×2 . Die Eigenwerte von A sind dann a und d , es hat also A zwei Eigenwerte oder einen Eigenwert, je nachdem, ob $a \neq d$ oder $a = d$ ist.

Ist $a \neq d$ und sind v_1, v_2 Eigenvektoren von A , so können diese offenbar nicht Vielfache voneinander sein, weil die zugehörigen Eigenwerte sonst gleich wären. Folglich bilden v_1 und v_2 eine aus Eigenvektoren bestehende Basis. Die Matrix A ist dann diagonalisierbar, und zwar konjugiert zu $\text{diag}(a, d)$. Genauer können wir $v_1 = e_1$ wählen. Um auch für v_2 konkret eine Möglichkeit anzugeben, müsste man eine weitere Fallunterscheidung treffen.

Ist andererseits $a = d$, so ist der Eigenraum $V_a(A)$ die Lösungsmenge des homogenen linearen Gleichungssystems mit Koeffizientenmatrix $\begin{pmatrix} 0 & b \\ 0 & 0 \end{pmatrix}$. Dieser Raum stimmt genau dann mit K^2 überein, wenn $b = 0$ ist; genau in diesem Fall ist also A diagonalisierbar.

Insbesondere sehen wir, dass über jedem Körper K die Matrix $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ nicht diagonalisierbar ist.

◇

BEMERKUNG 10.12. Zum Abschluss dieses Abschnitts noch eine wichtige Bemerkung über ein Phänomen, das bisher keine Rolle gespielt hat, aber nun berücksichtigt werden muss: Die Frage, ob eine Matrix diagonalisierbar ist, ist im allgemeinen davon abhängig, *über welchem Körper* wir sie betrachten. Mit anderen Worten: Eine Matrix $A \in M_n(K)$, die nicht diagonalisierbar ist, kann – aufgefasst als Matrix in $M_n(L)$ für einen geeigneten Erweiterungskörper L von K – diagonalisierbar sein. Ähnlich verhält es sich mit der Frage, ob eine Matrix überhaupt einen Eigenwert besitzt.

Als konkretes Beispiel betrachten wir wie in Beispiel 10.11 (1) eine Drehmatrix $\begin{pmatrix} a & -b \\ b & a \end{pmatrix} \in M_2(\mathbb{R})$ mit $-1 < a, b < 1$, $a^2 + b^2 = 1$. Aufgefasst als Matrix in $M_2(\mathbb{C})$, wenn wir also den Grundkörper durch den Körper $\mathbb{C}$ der komplexen Zahlen ersetzen, ist diese Matrix diagonalisierbar. Ihre Eigenwerte sind die beiden Lösungen der quadratischen Gleichung $x^2 - 2ax + 1$, das sind $a + ib$ und $a - ib$. Unsere Voraussetzungen implizieren, dass $b \neq 0$, also handelt es sich um komplexe Zahlen, die nicht in $\mathbb{R}$ liegen. ◇

10.2. Eigenräume

Wir haben gesehen, dass ein Endomorphismus f eines endlichdimensionalen Vektorraums V genau dann diagonalisierbar ist, wenn eine Basis aus Eigenvektoren besteht. Insbesondere zerfällt dann V in die direkte Summe der Eigenräume von f (siehe Satz 10.16, Korollar 10.17). Im allgemeinen können wir das nicht erwarten, aber die Eigenvektoren zu verschiedenen Eigenwerten sind in jedem Fall voneinander unabhängig, und zwar in dem folgenden präzisen Sinn.

SATZ 10.13. Seien $f: V \rightarrow V$ ein Endomorphismus eines K -Vektorraums V und $v_1, \dots, v_m$ Eigenvektoren von f zu paarweise verschiedenen Eigenwerten von f . Dann sind $v_1, \dots, v_m$ linear unabhängig.

BEWEIS. Wir führen Induktion nach m . Für $m = 1$ ist die Sache klar, da ein Eigenvektor nach Definition nicht der Nullvektor ist. Für den Induktionsschritt dürfen wir annehmen, dass sich der Nullvektor nur in trivialer Weise als Linearkombination von $m-1$ Eigenvektoren zu paarweise verschiedenen Eigenvektoren darstellen lässt.

Seien $a_i \in K$ mit

$$a_1 v_1 + \dots + a_m v_m = 0.$$

Wir wollen zeigen, dass alle a_i verschwinden. Anwendung von f liefert

$$\lambda_1 a_1 v_1 + \dots + \lambda_m a_m v_m = 0.$$

Wenn wir die ursprüngliche Linearkombination mit λ_1 durchmultiplizieren und das Ergebnis von der vorherigen Gleichung abziehen, erhalten wir:

$$(\lambda_2 - \lambda_1) a_2 v_2 + \dots + (\lambda_m - \lambda_1) a_m v_m = 0.$$

Die Induktionsvoraussetzung liefert uns nun

$$(\lambda_2 - \lambda_1) a_2 = \dots = (\lambda_m - \lambda_1) a_m = 0$$

und damit

$$a_2 = \dots = a_m = 0,$$

denn die Differenzen $\lambda_i - \lambda_1, i = 2, \dots, m$ sind nach Voraussetzung alle verschieden von 0. Dann folgt aber $a_1 v_1 = 0$, also muss auch $a_1 = 0$ gelten. $\square$

ERGÄNZUNG 10.14. Wir erklären einen alternativen Beweis von Satz 10.13, der die Vandermonde-Determinante (Beispiel 9.20) benutzt.

Wir betrachten mit derselben Notation wie im Beweis des Satzes eine Linearkombination

$$a_1 v_1 + \dots + a_m v_m = 0.$$

Wir wollen zeigen, dass alle Koeffizienten a_i verschwinden.

Dazu wenden wir $f, f^2 := f \circ f, f^3, \dots, f^{m-1}$ auf diese Gleichung an. Weil die v_i nach Voraussetzung Eigenvektoren von f sind, erhalten wir insgesamt die m Gleichungen

$$\begin{array}{ccccccc} v_1 & + & \dots & + & v_m & = & 0 \\ \lambda_1 v_1 & + & \dots & + & \lambda_m v_m & = & 0 \\ \lambda_1^2 v_1 & + & \dots & + & \lambda_m^2 v_m & = & 0 \\ \lambda_1^3 v_1 & + & \dots & + & \lambda_m^3 v_m & = & 0 \\ \vdots & & & & \vdots & & \vdots \\ \lambda_1^{m-1} v_1 & + & \dots & + & \lambda_m^{m-1} v_m & = & 0. \end{array}$$

Weil die λ_i nach Voraussetzung paarweise verschieden sind, ist die zugehörige Vandermonde-Determinante von Null verschieden, die zugrundeliegende Vandermonde-Matrix also invertierbar. Die Zeilen dieser Matrix sind also linear unabhängig und bilden eine Basis von K^m . Unter der Abbildung

$$g: K^m \longrightarrow V, \quad (x_1, \dots, x_m)^t \mapsto x_1 a_1 v_1 + \dots + x_m a_m v_m,$$

werden alle Vektoren dieser Basis auf 0 abgebildet; genau das sagen die obigen Gleichungen aus. Da alle Vektoren einer Basis auf 0 abgebildet werden, handelt es sich bei g um die Nullabbildung. Andererseits liegen – als die Bilder der Standardbasisvektoren – alle $a_i v_i$ im Bild von g . Es gilt also $a_i v_i = 0$ für alle i , und weil Eigenvektoren per Definition von Null verschieden sind, $a_i = 0$ für alle i . $\square$ Ergänzung 10.14

KOROLLAR 10.15. Sei $f: V \rightarrow V$ ein Endomorphismus eines endlichdimensionalen Vektorraums, $n = \dim(V)$.

- (1) Der Endomorphismus f hat höchstens n verschiedene Eigenwerte.
- (2) Hat f genau n verschiedene Eigenwerte, so ist f diagonalisierbar.

BEWEIS. Der erste Teil ist nach Satz 10.13 klar. Hat f genau n verschiedene Eigenwerte, so ist nach demselben Satz jede Familie, die aus je einem Eigenvektor zu jedem der Eigenwerte besteht, linear unabhängig, bildet also eine Basis. $\square$

Wir können die Situation gut beschreiben, indem wir den Begriff der direkten Summe von zwei Untervektorräumen folgendermaßen auf den Fall mehrerer Untervektorräume verallgemeinern:

SATZ 10.16. Seien V ein Vektorraum und $U_1, \dots, U_m \subseteq V$ Untervektorräume. Dann sind äquivalent:

- (i) Es ist $V = \sum_{i=1}^m U_i$ und für jedes $v \in V$ ist die Darstellung $v = \sum_{i=1}^m u_i$, $u_i \in U_i$, eindeutig.
- (ii) Es ist $V = \sum_{i=1}^m U_i$ und die einzige Möglichkeit, den Nullvektor von V als Summe $0 = \sum_{i=1}^m u_i$, $u_i \in U_i$, darzustellen, ist mit $u_i = 0$ für alle i .
- (iii) Es ist $V = \sum_{i=1}^m U_i$ und für alle i gilt $U_i \cap \left(\sum_{j \neq i} U_j\right) = \{0\}$.
- (iv) Der natürliche Homomorphismus $\bigoplus_{i=1}^m U_i \rightarrow V$, $(u_1, \dots, u_m) \mapsto \sum_{i=1}^m u_i$, von der »äußeren« direkten Summe der U_i (Abschnitt 6.6) ist ein Isomorphismus.

In diesem Fall schreiben wir $V = \bigoplus_{i=1}^m U_i$ und sagen, V sei die direkte Summe der U_i .

BEWEIS. Der Beweis ist eine leichte Übung. $\square$

Ist V endlichdimensional, so folgt in der Situation des Satzes, dass $\dim(V) = \sum_{i=1}^m \dim(U_i)$ ist, und wenn diese Dimensionsgleichheit schon bekannt ist, genügt es dazu noch entweder $V = \sum_{i=1}^m U_i$ oder die Eindeutigkeit der Darstellung wie in (i) oder (ii) – mit anderen Worten: entweder die Surjektivität oder die Injektivität in (iv) zu fordern.

Mit dieser Sprechweise können wir aus dem obigen Satz die Folgerung ableiten:

KOROLLAR 10.17. Sei f ein Endomorphismus eines endlichdimensionalen Vektorraums V , seien $\lambda_1, \dots, \lambda_m \in K$ die paarweise verschiedenen Eigenwerte von f und sei $V' \subseteq V$ der von den Eigenräumen $V_{\lambda_1}, \dots, V_{\lambda_m}$ von f erzeugte Untervektorraum. Dann gilt $V' = \bigoplus_{i=1}^m V_{\lambda_i}$. Der Endomorphismus f ist genau dann diagonalisierbar, wenn $V' = V$.

Wir werden die Theorie der Eigenwerte einer linearen Abbildung in der Vorlesung Lineare Algebra 2 weiter studieren und dann noch wesentlich besser verstehen. Beispiele für Aussagen, die wir im kommenden Semester beweisen werden, sind die folgenden.

- (1) Sei $A \in M_n(\mathbb{R})$ symmetrisch, d.h. es gelte $A = A^t$. Dann ist A diagonalisierbar.
- (2) Sei K der Körper der rationalen Zahlen (oder allgemeiner irgendein Körper, in dem $1 + 1 + 1 \neq 0$ gilt) und sei $n \in \mathbb{N}$. Sei $A \in M_n(K)$ eine Matrix mit $A^3 = A$. Dann ist A diagonalisierbar.

Wichtiger als dieses konkrete Beispiel ist natürlich, woran es liegt, dass die obige Aussage richtig ist, und wie/ob sie sich verallgemeinern lässt.

- (3) Wir werden mit der *Jordanschen Normalform* (und der *rationalen Normalform*) auch für nicht-diagonalisierbare Endomorphismen eines endlichdimensionalen Vektorraums V ein Ergebnis der Art beweisen, dass sich immer eine Basis von V finden lässt, so dass die zugehörige Darstellungsmatrix eine »relativ einfache« Form hat. Und zwar konkret, dass es sich um eine Block-Diagonalmatrix handelt, so dass in den Blöcken auf der Diagonale nur an ganz bestimmten, wenigen Stellen Einträge $\neq 0$ stehen.

10.3. Ergänzungen *

Das Thema Eigenwerte ist (wieder einmal ...) ein Thema, zu dem man ganze Bücher füllen könnte. Es folgt eine kleine Auswahl der möglichen Themen für Ergänzungen – hoffentlich finden Sie etwas, was Sie interessiert. In der Linearen Algebra 2 werden wir dann nochmals die Gelegenheit haben, uns diesem Thema zu widmen.

BEMERKUNG 10.18. Wir kommen noch einmal auf die Spur einer quadratischen Matrix zu sprechen, die wir als die Summe ihrer Diagonaleinträge definiert hatten. Ist A eine obere Dreiecksmatrix, so sind die Diagonaleinträge genau die Eigenwerte von A , und die Spur ist also die Summe der Eigenwerte, jedenfalls dann, wenn die Diagonaleinträge paarweise verschieden sind. Wenn derselbe Wert mehrfach auf der Diagonale auftritt, müssen wir ihn bei der Berechnung der Spur entsprechend mehrfach berücksichtigen. Wir sprechen davon, dass der Wert mit einer gewissen »Vielfachheit« auftritt.

Allgemeiner den Eigenwerten einer beliebigen Matrix eine »Vielfachheit« zuzuordnen, wird ein wichtiger Schritt in der weiteren Entwicklung der Eigenwerttheorie sein (dann in der Linearen Algebra 2, wenn wir das charakteristische Polynom einführen).

Sei nun K ein Körper der Charakteristik 0, das heißt, dass die natürliche Abbildung $\mathbb{N} \rightarrow K$, $n \mapsto 1 + \dots + n$ (n Summanden), injektiv ist. Dies ist für den Körper der rationalen Zahlen und somit für alle seine Erweiterungskörper der Fall. Dann kann man die folgende nicht-offensichtliche Aussage zeigen. Sind A, B quadratische Matrizen, für die $\text{Spur}(A^j) = \text{Spur}(B^j)$ für alle $j \geq 0$ gilt, dann haben A und B dieselben Eigenwerte.

Weil für einen Körper der Charakteristik $p > 0$ gilt, dass $\text{Spur}(E_p) = 0$ ist, ist dieses Ergebnis im Fall positiver Charakteristik nicht richtig: E_p und die Nullmatrix haben natürlich nicht dieselben Eigenwerte. $\diamond$

BEISPIEL 10.19. Wir betrachten noch einmal die Fibonacci-Folge $(F_n)_{n \geq 0}$, siehe Beispiel 5.60.

Wir setzen wieder $A = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}$. Wir hatten gesehen, dass

$$\begin{pmatrix} F_{n+1} \\ F_n \end{pmatrix} = A^n \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

für alle $n \geq 0$ gilt. Indem wir die Matrix A diagonalisieren, können wir daraus eine Formel für die n -te Fibonacci-Zahl herleiten. (Vergleiche Ergänzung 6.60 für einen völlig anderen Beweis dafür.)

Eine Zahl λ ist genau dann ein Eigenwert von A (in $\mathbb{R}$), wenn es einen Vektor $v \in \mathbb{R}^2$ gibt mit $Av = \lambda v$ und $v \neq 0$, also wenn das homogene lineare Gleichungssystem mit Koeffizientenmatrix

$$\begin{pmatrix} 1 - \lambda & 1 \\ 1 & -\lambda \end{pmatrix}$$

eine nichttriviale Lösung hat. Die Determinante dieser Matrix ist $-(1 - \lambda)\lambda - 1 = \lambda^2 - \lambda - 1$. Sie verschwindet für

$$\lambda = \frac{1 + \sqrt{5}}{2} \quad \text{und} \quad \lambda = \frac{1 - \sqrt{5}}{2}.$$

Wir setzen $\varphi = \frac{1 + \sqrt{5}}{2}$. Dann ist $1 - \varphi = \frac{1 - \sqrt{5}}{2}$. Die beiden Eigenwerte von A sind also φ und $1 - \varphi$.

Da A also 2 Eigenwerte hat, muss die Matrix diagonalisierbar sein. Für $\lambda \in \{\varphi, 1 - \varphi\}$ sind die Zeilen der obigen Matrix linear abhängig, im Sinne des linearen Gleichungssystems

stellen sie äquivalente Gleichungen dar. Wir können daher Eigenvektoren finden, indem wir beispielsweise die Gleichung zur zweiten Zeile lösen, und kommen auf

$$b_1 = \begin{pmatrix} \varphi \\ 1 \end{pmatrix}, \quad b_2 = \begin{pmatrix} 1 - \varphi \\ 1 \end{pmatrix}$$

Diese beiden Vektoren bilden eine Basis von $\mathbb{R}^2$ aus Eigenvektoren von A .

Durch den entsprechenden Basiswechsel erhalten wir die Darstellung

$$A = \frac{1}{\sqrt{5}} \begin{pmatrix} \varphi & 1 - \varphi \\ 1 & 1 \end{pmatrix} \begin{pmatrix} \varphi & \\ & 1 - \varphi \end{pmatrix} \begin{pmatrix} 1 & \varphi - 1 \\ -1 & \varphi \end{pmatrix}$$

(denn $\begin{pmatrix} \varphi & 1 - \varphi \\ 1 & 1 \end{pmatrix}^{-1} = \frac{1}{\sqrt{5}} \begin{pmatrix} 1 & \varphi - 1 \\ -1 & \varphi \end{pmatrix}$).

Damit bekommen wir

$$A^n = \frac{1}{\sqrt{5}} \begin{pmatrix} \varphi & 1 - \varphi \\ 1 & 1 \end{pmatrix} \begin{pmatrix} \varphi^n & \\ & (1 - \varphi)^n \end{pmatrix} \begin{pmatrix} 1 & \varphi - 1 \\ -1 & \varphi \end{pmatrix},$$

also

$$\begin{aligned} \begin{pmatrix} F_{n+1} \\ F_n \end{pmatrix} &= A^n \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \frac{1}{\sqrt{5}} \begin{pmatrix} \varphi & 1 - \varphi \\ 1 & 1 \end{pmatrix} \begin{pmatrix} \varphi^n & \\ & (1 - \varphi)^n \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \\ &= \frac{1}{\sqrt{5}} \begin{pmatrix} \varphi & 1 - \varphi \\ 1 & 1 \end{pmatrix} \begin{pmatrix} \varphi^n \\ -(1 - \varphi)^n \end{pmatrix} \\ &= \frac{1}{\sqrt{5}} \begin{pmatrix} \varphi^{n+1} - (1 - \varphi)^{n+1} \\ \varphi^n - (1 - \varphi)^n \end{pmatrix}, \end{aligned}$$

das heißt

$$F_n = \frac{\varphi^n - (1 - \varphi)^n}{\sqrt{5}} = \frac{1}{\sqrt{5}} \left(\left(\frac{1 + \sqrt{5}}{2} \right)^n - \left(\frac{1 - \sqrt{5}}{2} \right)^n \right).$$

◇

ERGÄNZUNG 10.20 (Divisionsalgebren, Fortsetzung). Wir kommen nun auf die Frage nach der Existenz von endlich-dimensionalen Divisionsalgebren über den reellen und den komplexen Zahlen zurück. Siehe Ergänzung 6.65, wo wir die folgenden Sätze bereits angekündigt hatten.

SATZ 10.21. *Die einzige Divisionsalgebra über den reellen Zahlen, die als $\mathbb{R}$ -Vektorraum ungerade endliche Dimension hat, ist der Körper $\mathbb{R}$ selbst.*

BEWEIS. Sei A eine Divisionsalgebra über $\mathbb{R}$, so dass $\dim(A)$ eine ungerade Zahl ist. Wir wollen zeigen, dass $A = \mathbb{R}$, mit anderen Worten, dass die Inklusion $\mathbb{R} \subseteq A$ eine Gleichheit ist. Sei dazu $a \in A$. Die Abbildung $m_a: A \rightarrow A, x \mapsto ax$ ist ein $\mathbb{R}$ -Vektorraum-Homomorphismus. Da die Vektorraumdimension von A ungerade ist und jede Polynomfunktion von ungeradem Grad eine Nullstelle in $\mathbb{R}$ hat, sehen wir (Beispiel 10.11 (3)), dass m_a einen Eigenvektor x zu einem Eigenwert $\lambda \in \mathbb{R}$ besitzt, also $ax = \lambda x$. Es folgt $(a - \lambda)x = 0$, und da x als Eigenvektor von 0 verschieden ist und A ein Schiefkörper ist, dass $a - \lambda = 0$, also $a = \lambda \in \mathbb{R}$. $\square$

Der entscheidende Punkt im Beweis war, dass jede Polynomfunktion der Form $\lambda \mapsto \det(m_a - \lambda \text{id})$ eine Nullstelle hat. Dasselbe Argument zeigt also den folgenden Satz (wo wir »sicherheitshalber« voraussetzen, dass K unendlich sei, damit wir vernünftig über den Grad einer Polynomfunktion sprechen können; in der Linearen Algebra 2 werden wir die Situation noch besser aufklären).

SATZ 10.22. *Sei K ein unendlicher Körper, so dass jede Polynomfunktion vom Grad > 0 eine Nullstelle in K hat. Dann gibt es außer K selbst keine endlich-dimensionalen Divisionsalgebren über K .*

Einen Körper, der die Voraussetzungen des Satzes erfüllt, nennt man *algebraisch abgeschlossen*. Es gilt

THEOREM 10.23 (Fundamentalsatz der Algebra). *Der Körper der komplexen Zahlen ist algebraisch abgeschlossen.*

Auch wenn der Satz den Namen Fundamentalsatz der *Algebra* trägt, ist es natürlicher, ihn mit Mitteln der Analysis (oder der *Funktionentheorie*, der Analysis einer komplexen Veränderlichen) zu beweisen.

Zum Abschluss noch eine Bemerkung zur Situation über $\mathbb{R}$, wo wir oben ja nur die Existenz von Divisionsalgebren mit ungerader Dimension > 1 ausgeschlossen haben. Wir haben ja mit den komplexen Zahlen und den Quaternionen schon zwei Divisionsalgebren über $\mathbb{R}$ kennengelernt, wobei $\mathbb{C}$ sogar kommutativ ist, der Schiefkörper der Quaternionen jedoch nicht. Der Satz von Frobenius besagt, dass es neben diesen keine weiteren endlichdimensionalen Divisionsalgebren über $\mathbb{R}$ gibt. Lässt man die Bedingung der Assoziativität fallen, kommen noch als einzige weitere Möglichkeit die *Oktonionen* hinzu, die über $\mathbb{R}$ die Vektorraumdimension 8 haben. □ Ergänzung 10.20

ERGÄNZUNG 10.24 (Der Page-rank-Algorithmus, Fortsetzung). Wir setzen die Diskussion von Ergänzung 7.66 fort. Wir hatten dort die »Google-Matrix« G gefunden, eine $(N \times N)$ -Matrix, deren Einträge positive reelle Zahlen sind und deren Spaltensummen alle $= 1$ sind. Unser Ziel ist es, eine Lösung des linearen Gleichungssystems $(G - E_N)x = 0$ zu finden, so dass alle Einträge von x positiv sind und die Summe der Einträge $= 1$ ist. Hat man einmal eine Lösung, in der alle Einträge positiv sind, kann man die Summe natürlich durch eine geeignete Skalierung auf 1 bringen. In der neu eingeführten Sprache der Eigenwerte sagen wir, dass x ein Eigenvektor von G zum Eigenwert 1 ist. Wir haben schon gezeigt, dass genau ein solches x existiert. Es bleibt aber die Frage, wie man versuchen könnte, x zu berechnen, denn der Gauß-Algorithmus ist für eine Matrix dieser Größe nicht praktisch durchführbar. Wie wir sehen werden, ist die Sichtweise der Eigenwerttheorie an dieser Stelle hilfreich. Wir beginnen mit dem folgenden Satz, der noch einmal zusammenfasst, was wir schon bewiesen haben, und eine wichtige Zusatzinformation über die Eigenwerte von G liefert.

SATZ 10.25. *Sei $G = (G_{ij})_{i,j} \in M_{N \times N}(\mathbb{R})$ eine Matrix mit $G_{ij} > 0$ für alle i, j , und so dass alle Spaltensummen $= 1$ sind.*

- (1) *Es ist $\dim V_1(G) = 1$. Ist $x \in V_1(G)$, so sind entweder alle Einträge positiv, oder alle Einträge negativ.*
- (2) *Für alle Eigenwerte $\lambda \neq 1$ von G gilt $|\lambda| < 1$.*

BEWEIS. zu (1). Diese Aussagen haben wir bereits gezeigt (Korollar 7.69), sie sind hier lediglich in die Terminologie der Eigenräume übersetzt.

zu (2). Um die anderen Eigenwerte abzuschätzen, betrachten wir die transponierte Matrix G^t . Wir wissen (Lemma 10.10), dass G und G^t dieselben Eigenwerte haben, und dass die Dimensionen der entsprechenden Eigenräume gleich sind. Also gilt $\dim V_1(G^t) = 1$. Für die transponierte Matrix (deren Zeilensummen alle gleich 1 sind), können wir direkt einen Eigenvektor zum Eigenwert 1 angeben $(1, \dots, 1)^t$. (Natürlich wird dies praktisch nie ein Eigenvektor von G sein, und es gibt keine Möglichkeit, durch einen Trick daraus einen Eigenvektor von G zu berechnen!)

Sei nun $\lambda \neq 1$ ein Eigenwert von G und G^t und sei v ein Eigenvektor von G^t zum Eigenwert λ . Weil $V(1, G^t) = \langle (1, \dots, 1)^t \rangle$ ist, sehen wir, dass nicht alle Komponenten von v gleich sein können.

Sei $i \in \{1, \dots, N\}$, so dass $|v_i|$ maximal ist. Dann gilt

$$|\lambda v_i| = |(G^t v)_i| = \left| \sum_{j=1}^N (G^t)_{ij} v_j \right| \leq \sum_{j=1}^N (G^t)_{ij} |v_j| < |v_i| \sum_{j=1}^N (G^t)_{ij} = |v_i|.$$

Es folgt $|\lambda| < 1$, wie gewünscht.

Übrigens: Wie man sieht, gilt dieses Argument sogar für alle Eigenwerte von G in $\mathbb{C}$. (Und im allgemeinen können Matrizen G mit den obigen Eigenschaften Eigenwerte in $\mathbb{C} \setminus \mathbb{R}$ haben.) $\square$

Nun zu den Methoden, mit denen man x berechnen kann. Für eine Folge $(v_{(n)})_n$ von Vektoren schreiben wir $x = \lim_{n \rightarrow \infty} v_{(n)}$, wenn für alle i gilt: $x_i = \lim_{n \rightarrow \infty} v_{(n),i}$. (Wir schreiben $v_{(n)}$ statt v_n um eine Verwechslung mit der Notation für den n -ten Eintrag eines Vektors v zu vermeiden. Jedes $v_{(n)}$ ist ein Vektor mit Einträgen $v_{(n),i}$.)

SATZ 10.26. Sei $G = (G_{ij})_{i,j} \in M_{N \times N}(\mathbb{R})$ eine Matrix mit $G_{ij} > 0$ für alle i, j , und so dass alle Spaltensummen $= 1$ sind. Sei $x \in V_1(G)$ der eindeutig bestimmte Vektor mit $\sum_i x_i = 1$.

Sei v irgendein Vektor mit positiven Einträgen und $\sum_i v_i = 1$. Dann gilt

$$x = \lim_{n \rightarrow \infty} G^n v.$$

BEWEIS. Wir führen den Beweis, um das Beweisprinzip transparent darzustellen, unter der Zusatzannahme aus, dass die Matrix G diagonalisierbar ist. Diese Annahme ist allerdings in der Praxis typischerweise nicht erfüllt. Siehe den in Frage 2.7 zitierten Artikel von Bryan und Leise (Prop. 4, Prop. 5) für weitere Informationen, wie man den allgemeinen Fall behandelt.

Ist G diagonalisierbar, dann können wir eine Basis $\mathcal{B} = (b_1, \dots, b_N)$ aus Eigenvektoren von G finden. Dabei sei $b_1 = x$ (mit Eigenwert 1). Sei λ_i der Eigenwert zu b_i . Es gilt dann $|\lambda_i| < 1$ für alle $i > 1$. Wenn wir

$$v = a_1 b_1 + a_2 b_2 + \dots + a_N b_N$$

schreiben, so folgt

$$G^n v = a_1 b_1 + a_2 \lambda_2^n b_2 + \dots + a_N \lambda_N^n b_N.$$

Weil $\lim_{n \rightarrow \infty} \lambda_i = 0$ gilt, folgt

$$\lim_{n \rightarrow \infty} G^n v = a_1 b_1.$$

Da v und damit alle $G^n v$ ebenso wie $b_1 = x$ als Summe aller Einträge 1 haben, folgt auch $a_1 = 1$. (In der Tat kann man auch leicht direkt sehen, dass für alle b_i mit $i > 1$ die Summe der Einträge gleich 0 ist. Das liegt daran, dass G den $(N-1)$ -dimensionalen Unterraum aller Vektoren mit dieser Eigenschaft in sich selbst abbildet und $\langle x \rangle$ ein Komplement dazu ist, das ebenfalls in sich selbst abgebildet wird.) $\square$

Es zeigt sich, dass die Folge der $G^n v$ wie im Satz ziemlich schnell konvergiert, zwischen 50 und 100 Iterationen waren angeblich für Google ausreichend, um den gesuchten Vektor x hinreichend genau anzunähern.

BEISPIEL 10.27. Wenn man dieses Verfahren mit dem Mini-Internet auf Beispiel 2.8 durchführt, kommt man auf das folgende Ergebnis:

Mit $v = \frac{1}{6} \cdot (1, 1, 1, 1, 1, 1)^t$ ist der Unterschied zwischen x und $G^5 v$ in jedem Eintrag < 0.0058 und die Reihenfolge bezüglich Relevanz ist für diese beiden Vektoren (und ebenso für alle $G^n v$ mit $n > 5$) dieselbe. Der Unterschied zwischen x und $G^{14} v$ ist in jedem Eintrag kleiner als 0.0001. $\diamond$

Es bleibt natürlich ein nicht-triviales Problem, diesen Eigenvektor der Google-Matrix wirklich auszurechnen, und sicher hat Google eine Kombination von verschiedenen Methoden (und Tricks) benutzt, um das zu machen. Diese Rechnung wurde auch bei weitem nicht täglich, sondern eher im Abstand von mehreren Monaten durchgeführt.

Eigenwerte von (großen) Matrizen zu berechnen, ist auch an vielen anderen Stellen ein Problem. Die hier vorgestellte Potenzmethode ist eine von mehreren Möglichkeiten, die man dafür kennt. □ Ergänzung 10.24

ERGÄNZUNG 10.28. Die Eigenwerte und Eigenvektoren einer diagonalisierbaren Matrix $A \in M_n(\mathbb{R})$ spielen eine Rolle bei der Lösung des durch A gegebenen Systems von linearen Differentialgleichungen $y' = Ay$, was die Kurzschreibweise dafür ist, dass n (beliebig oft) differenzierbare Funktionen $y_i: \mathbb{R} \rightarrow \mathbb{R}$ gesucht sind, für deren Ableitungen

$$y'_i(x) = \sum_{j=1}^n a_{ij} y_j(x)$$

ist. Hier seien $a_{ij} \in \mathbb{R}$ die Einträge von A .

Ist $v = (v_1, \dots, v_n)^t \in \mathbb{R}^n$ ein Eigenvektor von A zum Eigenwert λ , so ist durch

$$y_i(x) = v_i \exp(\lambda x)$$

eine Lösung des Systems gegeben, weil

$$y'_i(x) = v_i \lambda \exp(\lambda x) = \sum_{j=1}^n a_{ij} v_j \exp(\lambda x),$$

denn $Av = \lambda v$ impliziert $\lambda v_i = \sum_{j=1}^n a_{ij} v_j$ für alle i . Der Ausdruck auf der rechten Seite ist $\sum_{j=1}^n a_{ij} y_j(x)$, wie gewünscht.

Man kann zeigen, dass in dem Fall, dass A diagonalisierbar ist, die Lösungen, die man so aus einer Basis von Eigenvektoren von A erhält, eine Basis vom Vektorraum aller Lösungen dieses Systems von Differentialgleichungen bilden.

Mit der Erweiterung der Theorie, die wir in der Linearen Algebra 2 erreichen werden, speziell dem Satz über die Jordansche Normalform, kann man auch den Fall von Matrizen A behandeln, die nicht diagonalisierbar sind. Dies ist die Situation, in der die lineare Algebra ihre Kraft wirklich ausspielen kann.

Weitere Informationen (aus Sicht der linearen Algebra) finden sich in [K1], Abschnitte 5.5 und 5.7. Aus Sicht der Theorie der Differentialgleichungen wird dieser Sachverhalt wohl in jedem Buch über gewöhnliche Differentialgleichungen dargestellt. □ Ergänzung 10.28

ERGÄNZUNG 10.29 (SEIR-Modell für die Ausbreitung einer Epidemie). Aus aktuellem Anlass ... sei erwähnt, dass auch die Standardmodelle für die Ausbreitung einer Epidemie die Eigenwertsprache benutzen. Die [Basisreproduktionszahl](https://de.wikipedia.org/wiki/Basisreproduktionszahl)⁴ R_0 , die im Moment im Fall des Corona-Virus eine wichtige Kennzahl für die erwartete weitere Entwicklung der Epidemie ist, kann als Eigenwert einer Matrix (der sogenannten *next generation matrix*) definiert werden.

Alle gängigen Modelle basieren auf Systemen von Differentialgleichungen, die nicht unbedingt die einfache Form aus Ergänzung 10.28 haben, aber oft durch ein solch einfaches System gut angenähert werden können.

⁴ <https://de.wikipedia.org/wiki/Basisreproduktionszahl>

”

Auf Englisch heißen Eigenwerte, Eigenvektoren und Eigenräume üblicherweise *eigenvalues*, *eigenvectors* und *eigenspaces*. Das sind die *scary German terms*, auf die im folgenden Zitat angespielt wird.

It turns out that there is a straightforward extension of the theory for structured epidemic models. The mathematics behind this theory is not especially difficult, but it does involve scary German terms that are not familiar to the non-engineers in our midst ...

J. H. Jones, [Notes on \$\mathcal{R}_0\$](#) ^a, May 2007

^a<https://web.stanford.edu/~jhj1/teachingdocs/Jones-on-R0.pdf>

Ein für die Corona-Pandemie wichtiges Modell ist das sogenannte [SEIR-Modell](#)⁵. Dazu hoffentlich demnächst an dieser Stelle noch etwas mehr ...

□ Ergänzung 10.29

⁵<https://de.wikipedia.org/wiki/SEIR-Modell>

Analytische Geometrie *

In diesem Kapitel werden wir die Verbindung zwischen der »algebraischen« Theorie und geometrischen Fragen herstellen. Es soll hier vor allem darum gehen zu erklären, dass die Theorie, so wie wir sie in den Kapiteln 6 bis 10 entwickelt haben, auch geometrisch interpretiert werden und dazu benutzt werden *kann*, geometrische Sachverhalte zu beschreiben und zu beweisen. Es würde hier zu weit führen, dies tatsächlich in nennenswertem Umfang durchzuführen.

Was ist Geometrie?

Vom Wortursprung – griechisch *γεωμετρία* – bedeutet Geometrie »Vermessung der Erde«, und die Vermessung (von Teilen) der Erde und des Universums war eine Aufgabe, die maßgeblich zur Entwicklung der mathematischen Theorie der Geometrie und der Mathematik insgesamt beigetragen hat. Und zwar schon vor mehreren Tausend Jahren; aber auch Gauß hat sich intensiv mit der Landvermessung (konkret der Vermessung des Königreichs Hannover) beschäftigt, und wurde dadurch zu wichtigen Ergebnissen der Differentialgeometrie von Flächen motiviert, insbesondere zu dem Begriff der Krümmung, der heute nach ihm benannt ist. Die Berechnung der Planetenbahnen ist ein anderes Beispiel für ein historisch bedeutsames Einsatzgebiet der Geometrie. Auch heutzutage sind ähnliche Fragen von praktischer Relevanz, Stichwort (zum Beispiel): GPS.

Es ist allerdings nicht ganz einfach, in ein oder zwei Sätzen zu beschreiben, was »Geometrie« im mathematischen Sinne eigentlich ist. Vielleicht auch deshalb, weil es mehrere Gebiete der Geometrie gibt, die zwar gewisse Gemeinsamkeiten haben, sich aber zum Teil auch sehr unterscheiden. Hier eine – unvollständige – Liste:

Analytische Geometrie – die analytische Geometrie ist die »Geometrie des Koordinatensystems«, d.h. Punkte, Geraden, Ebenen ... werden in Termen des Koordinatensystems in $\mathbb{R}^n$ (oder allgemeiner: eines geeigneten Koordinatensystems) beschrieben. Man kann dann Abstände und Winkel definieren, geometrische Zusammenhänge »algebraisch« ausdrücken und mit diesen Objekten »rechnen«.

Affine, projektive Geometrie – unter affiner Geometrie versteht man zunächst die Geometrie des Raums $\mathbb{R}^n$ (oder unter Umständen des K^n für irgendeinen Körper K). Man kann außerdem aus den essenziellen Eigenschaften dieser Räume ein Axiomensystem ableiten und dann darin arbeiten, ohne noch auf den konkreten Raum $\mathbb{R}^n$ (oder überhaupt auf einen Grundkörper) Bezug zu nehmen. Man spricht dann auch von synthetischer (oder: axiomatischer) Geometrie. In der projektiven Geometrie wird dabei der »affine Raum« ersetzt durch den größeren »projektiven Raum«, der zwar nicht mehr direkt dem Anschauungsraum entspricht, sich aber aus geometrischer Sicht in mancherlei Hinsicht besser (und einfacher) verhält. Zum Beispiel gilt, dass sich je zwei verschiedene Geraden in der projektiven Ebene in genau einem Punkt schneiden; den Sonderfall der parallelen Geraden gibt es also nicht mehr. Man erreicht das, indem man der affinen Ebene in geeigneter Weise zusätzliche Punkte hinzufügt, die man die Punkte »im Unendlichen« nennt.



Der Leser soll gleichsam in dem großen Garten der Geometrie spazieren geführt werden, und jeder soll sich einen Strauß pflücken können, wie er ihm gefällt.

David Hilbert, in *Anschauliche Geometrie*, Springer 1932, Vorwort, S. VI.

Die spezielle Problematik dieses Kapitels ist, dass viele Begriffe der Geometrie uns »im Prinzip« bekannt sind, wir aber eine mathematische Definition erst geben müssen (zum Beispiel: Abstand, senkrecht, Drehung, Winkel). Ich möchte hier besonderen Wert darauf legen, zu erklären, warum genau die gegebenen Definitionen die »richtigen« sind. Mindestens sollte das bedeuten, dass sich die so definierten Begriffe so verhalten, wie wir es vom anschaulichen Verständnis her erwarten. Noch überzeugender ist es in der Regel, wenn man sogar begründen kann, dass es sich um die einzig mögliche Definition handelt, die gewisse Eigenschaften, die wir jedenfalls fordern wollen, hat.

Was ist Geometrie (Fortsetzung)

Differentialgeometrie – in der Differentialgeometrie untersucht man »differenzierbare Mannigfaltigkeiten«, also geometrische Objekte, die lokal aussehen wie offene Teilmengen von $\mathbb{R}^n$, und die »aus solchen Stücken durch differenzierbare Übergangsabbildungen zusammengeklebt« werden können. (Wenn man zusätzliche Strukturen betrachtet, kommt man zu Varianten wie der Riemannschen Geometrie oder der symplektischen Geometrie.)

Komplexe Geometrie – die komplexe Geometrie ist das Analogon der Differentialgeometrie über den komplexen Zahlen, d.h. der »übliche« Begriff der differenzierbaren Abbildung wird ersetzt durch den Begriff der komplex differenzierbaren (oder: holomorphen) Funktion.

Algebraische Geometrie – in der algebraischen Geometrie beschäftigt man sich, grob gesprochen, mit geometrischen Objekten, die als Lösungsmengen von Gleichungssystemen von Polynomgleichungen (in mehreren Unbestimmten) beschrieben werden können. Über den reellen und komplexen Zahlen lassen sich die Techniken der Differentialgeometrie und der komplexen Geometrie anwenden, aber die Klasse der »algebraischen« Objekte, die hier betrachtet werden, ist wesentlich kleiner. Weil man Polynome über jedem Körper betrachten kann (auch ohne einen Grenzwert- und Ableitungsbegriff zur Verfügung zu haben), ist die algebraische Geometrie auch nützlich, um Fragen über (beispielsweise) endliche Körpern zu studieren und eröffnet insbesondere die Möglichkeit, zahlentheoretische Fragen mit geometrischen Methoden zu untersuchen. Dieser Einfluß hat die Zahlentheorie in den letzten 60 Jahren wesentlich geprägt.

Rigide Geometrie – die rigide (oder: rigid-analytische) Geometrie ist ein Analogon der komplexen Geometrie über gewissen anderen Grundkörpern als den komplexen Zahlen, insbesondere über dem Körper $\mathbb{Q}_p$ der » p -adischen Zahlen«. Die von Scholze eingeführten »perfektoiden Räume«, die den Grundstein für die Ergebnisse bilden, für die er 2018 die Fields-Medaille bekommen hat, sind Objekte der rigiden Geometrie.

Wir werden in der Linearen Algebra 2 einige der Konzepte, die wir hier einführen, verallgemeinern und noch ausführlicher studieren. Dort werden wir dann – wie in der Linearen Algebra 1 zum Beispiel mit den Begriffen des Körpers und des Vektorraums – die wesentlichen Eigenschaften, beispielsweise des Abstands- oder Winkelbegriffs – in abstrakte Definitionen fassen, die uns erlauben, eine Theorie zu entwickeln, die gleichzeitig auf die Geometrie

der Ebene $\mathbb{R}^2$ und des Raums $\mathbb{R}^3$ (und höherdimensionale $\mathbb{R}^n$), wie wir sie kennen, und auf andere Räume angewandt werden kann, deren »geometrische« Eigenschaften nicht unserer gewöhnlichen Anschauung entsprechen, deren Studium aber trotzdem nützlich ist, sowohl in der Mathematik als auch beispielsweise in der Physik. Dabei können wir die Theorie auch zum Teil auf beliebige Grundkörper, und zum größten Teil immerhin auf die Grundkörper $\mathbb{R}$ und $\mathbb{C}$ ausdehnen.

Einige Referenzen zum Thema Analytische Geometrie, und speziell zu Büchern, die die Verbindung zwischen linearer Algebra und analytischer Geometrie herstellen.

J. Dieudonné, *Linear Algebra and Geometry*, Hermann 1969

Mathematisch sehr rigoros und gründlich geschrieben, im berühmt-berüchtigten Stil des Autors Dieudonné^a. Auch wenn das Wort »geometry« im Titel des Buches und in den Überschriften von 5 der 7 Kapitel vorkommt: Es gibt in diesem Buch keine einzige Abbildung. (Dennoch handelt es sich um ein nützliches Buch, in dem zum Beispiel sehr klar herausgearbeitet wird, welche Eigenschaften des Körpers der reellen Zahlen genau benutzt/benötigt werden, um den Winkelbegriff einzuführen.)

G. Fischer, *Analytische Geometrie*, [Fi-AG]

W. Klingenberg, *Lineare Algebra und Geometrie*, [Kl]

M. Koecher, *Lineare Algebra und analytische Geometrie*, Springer, 4. Aufl., 2003

<https://doi.org/10.1007/978-3-642-59056-6>

J. Richter-Gebert, *Perspectives on Projective Geometry*, Springer, 2011

<https://doi.org/10.1007/978-3-642-17286-1>

^ahttps://de.wikipedia.org/wiki/Jean_Dieudonn%C3%A9

Für uns soll »Geometrie« hier bedeuten, dass wir Vektorräume, und zwar speziell die Vektorräume $\mathbb{R}^n$, die natürlichen Verallgemeinerungen der Räume $\mathbb{R}^1$, $\mathbb{R}^2$, $\mathbb{R}^3$, die direkt unserer Anschauung zugänglich sind, in Hinsicht auf »geometrische Eigenschaften« besser verstehen möchten. Die geometrischen Eigenschaften, die wir betrachten werden, sind

- Der Abstand zwischen zwei Punkten,
- der Winkel zwischen zwei sich schneidenden Geraden (bzw. zwischen zwei Vektoren),
- der Flächeninhalt von (geeigneten) Teilmengen von $\mathbb{R}^n$.

Wir werden beispielsweise den Satz des Pythagoras, den Satz des Thales und den Kosinussatz beweisen und zeigen, dass der Flächeninhalt des Einheitskreises gleich der Hälfte seines Umfangs ist.

Im ersten Abschnitt des Kapitels führen wir den Begriff des affinen Raums ein. Dies ist der natürliche Kontext, um über Geraden, Ebenen, ... und ihre Lage zueinander (zum Beispiel, ob zwei Geraden parallel sind) zu sprechen, noch ohne Abstände oder Flächeninhalte zu messen.

Danach führen wir zunächst den Abstand zwischen zwei Punkten ein. Dies gibt auch einen direkten Zugang zur Definition, wann zwei Geraden bzw. Vektoren aufeinander senkrecht stehen. Das Skalarprodukt zweier Vektoren, das wir an dieser Stelle definieren werden, wird uns auch ermöglichen, den Winkel zwischen zwei Vektoren zu definieren. Um das in einer Art und Weise tun zu können, die sowohl mathematisch-formalen Ansprüchen genügt als auch die verwendeten Hilfsmittel (insbesondere die trigonometrischen Funktionen Sinus und Kosinus) einigermaßen motiviert, besprechen wir zunächst, wie man die Länge einer

»Kurve« messen kann. Wir werden das dann benutzen, um die Länge von Teilen (»Kreisbögen«) des Einheitskreises, des Kreises mit Radius 1 um den Ursprung $(0, 0)^t \in \mathbb{R}^2$, messen zu können.

Das erlaubt uns, einen Winkel in der Ebene, also den »Sektor«, der zwischen zwei am Ursprung beginnenden Halbgeraden eingeschlossen ist, durch die Länge des Kreisbogens zu messen, der durch die beiden Halbgeraden vom Einheitskreis ausgeschnitten wird.

In Abschnitt II.6 werden wir dann über den Begriff des Flächeninhalts von Teilmengen von $\mathbb{R}^2$ und allgemeiner über das Volumen (oder das *Maß*) von »messbaren« Teilmengen von $\mathbb{R}^n$ sprechen und den Zusammenhang zwischen der Determinante eines Endomorphismus und der dadurch bewirkten Volumenveränderung herstellen.

II.1. Affine Räume

Sei K ein Körper. Wir haben gesehen, dass die Lösungsmenge eines homogenen linearen Gleichungssystems $Ax = 0$ für $A \in M_{m \times n}(K)$ ein Untervektorraum von K^n ist. Für inhomogene lineare Gleichungssysteme erhalten wir, sofern die Lösungsmenge nicht leer ist, eine Lösungsmenge der Form $t + U$, wobei t irgendeine feste Lösung und U die Lösungsmenge des zugehörigen linearen Gleichungssystems ist. Wenn das Gleichungssystem nicht homogen ist, dann enthält $t + U$ nicht den Nullvektor, kann also kein Untervektorraum sein. Der Begriff des affinen Raums (bzw. des affinen Teilraums) erlaubt es uns, die Beschreibung von »verschobenen« Untervektorräumen, wie sie hier auftreten, zu formalisieren.

Als Kurzversion kann man definieren:

DEFINITION. Sei V ein Vektorraum. Ein *affiner Teilraum* (oder *affiner Unterraum*) von V ist eine Teilmenge von V , die entweder die leere Menge ist, oder die Form $v + U$ für $v \in V$ und einen Untervektorraum $U \subseteq V$ hat.

Dann können wir sagen, dass die Lösungsmenge eines linearen Gleichungssystems in n Unbestimmten ein affiner Unterraum von K^n ist. Ist $Z \subseteq V$ ein nicht-leerer affiner Unterraum, so ist der zugehörige Untervektorraum (das U aus der Definition) durch Z eindeutig bestimmt. Es ist nämlich, für irgendeinen fixierten Punkt $v \in Z$,

$$U = \{v - w; w \in Z\}.$$

Man kann dann Z auch eine Dimension (nämlich $\dim Z := \dim U$) zuordnen. Man kann leicht sehen, dass der Durchschnitt von affinen Unterräumen wieder ein affiner Unterraum ist. So kann man auch vom affinen Unterraum sprechen, der von einer Teilmenge erzeugt wird.

Wir holen noch ein kleines bisschen weiter aus und definieren allgemeiner den Begriff des affinen Raums:

DEFINITION II.1. Ein *affiner Raum* unter einem K -Vektorraum V ist eine Menge A zusammen mit einer Abbildung $t: V \rightarrow \text{Abb}(A, A)$, $v \mapsto t_v$, so dass gilt:

- (a) Für alle $v \in V$ ist die Abbildung t_v bijektiv, und es gilt $t_{v+w} = t_v \circ t_w$.
- (b) Für alle $a, a' \in A$ existiert ein eindeutig bestimmtes $v \in V$ mit $t_v(a) = a'$.

⊥

Eigenschaft (a) können wir auch dadurch ausdrücken, dass wir sagen, die Abbildung t sei ein Gruppenhomomorphismus $V \rightarrow \text{Bij}(A)$. Es folgt daraus, dass $t_0 = \text{id}_A$ ist. Der »ausgeartete« Fall $A = \emptyset$ ist auch erlaubt.

Wir schreiben statt $t_v(a)$ manchmal auch $a + v$. Sind $a, a' \in A$, dann bezeichnen wir das eindeutig bestimmte Element $v \in V$ in Teil (b) der Definition auch als $a' - a$.

BEISPIEL II.2. Ist V ein Vektorraum, so können wir $A := V$ als affinen Raum unter V auffassen, indem wir $t_v(w) := v + w$ definieren. Die Abbildung $t_v: V \rightarrow V, w \mapsto v + w$ heißt die *Verschiebung* oder *Translation* um den Vektor v .

Wenn wir in diesem Sinne die Rollen von V und A trennen, dann sind die Elemente von A die »Punkte«, und die Elemente von V operieren durch Translation auf A . Die Anwendung von v verschiebt alle Punkte um den *Vektor* v ; in diesem Modell ist es sinnvoll, sich Vektoren als Pfeile vorzustellen. $\diamond$

Wir können dann allgemein den Begriff des affinen Teilraums definieren. Im Fall $A = V$ des Beispiels erhalten wir die provisorische Definition vom Anfang dieses Abschnitts zurück.

DEFINITION II.3. Sei A ein affiner Raum über dem Vektorraum V . Eine Teilmenge $B \subseteq A$ ist ein *affiner Teilraum* (oder: *affiner Unterraum*), wenn $B = \emptyset$ ist, oder ein $b \in B$ existiert, so dass $\{b' - b; b' \in B\}$ ein Untervektorraum von V ist. $\dashv$

Ist $B \subseteq A$ ein Teilraum, dann ist für jedes $b \in B$ die Menge $\{b' - b; b' \in B\}$ ein Untervektorraum von V . Andererseits ist für jeden Untervektorraum $U \subseteq V$ und $a \in A$ die Teilmenge $a + U = \{t_u(a); u \in U\}$ ein affiner Teilraum von A .

Die natürlichen Abbildungen zwischen affinen Räumen sind die affin-linearen Abbildungen:

DEFINITION II.4. Seien A und B affine Räume über dem Körper K , mit zugrundeliegenden Vektorräumen V und W . Eine affin-lineare Abbildung zwischen A und B ist eine Abbildung $f: A \rightarrow B$, so dass für $a \in A$ die Abbildung

$$g: V \rightarrow W, \quad v \mapsto f(t_v(a)) - f(a),$$

eine lineare Abbildung ist und so dass diese Abbildung unabhängig von a ist, d.h. für alle $a, a' \in A, v \in V$ gilt $f(t_v(a)) - f(a) = f(t_v(a')) - f(a')$.

Es gilt dann also $f(a') = f(a) + g(a' - a)$ für alle $a, a' \in A$. $\dashv$

BEISPIEL II.5. Sind V und W Vektorräume über K und betrachten wir V und W jeweils als affinen Raum über sich selbst, dann ist eine affin-lineare Abbildung $f: V \rightarrow W$ eine Abbildung von der Form $v \mapsto g(v) + w$ für eine lineare Abbildung $g: V \rightarrow W$ und ein Element $w \in W$. $\diamond$

Eine *Gerade* in einem affinen Raum ist ein eindimensionaler affiner Teilraum. Es ist leicht zu überprüfen, dass das Bild einer Geraden unter einer affin-linearen Abbildung eine Gerade oder ein einziger Punkt ist. Umgekehrt gilt das folgende Theorem, das aber nicht ganz einfach zu beweisen ist:

THEOREM II.6 (Fundamentalsatz der affinen Geometrie). Seien A und B affine Räume derselben Dimension über $K = \mathbb{R}$ und sei $f: A \rightarrow B$ eine bijektive Abbildung, die jede Gerade in A auf eine Gerade in B abbildet. Dann ist f eine affin-lineare Abbildung.

Siehe [Fi-AG] 1.3.4. Über anderen Körpern als den reellen Zahlen ist die Lage etwas komplizierter. Das wird auch in der angegebenen Quelle diskutiert. Der Punkt ist, dass die reellen Zahlen nur einen einzigen »Körperautomorphismus« haben (nämlich die Identität).

Wir erhalten auch das folgende Korollar. Mit Geraden sind in Teil (b) alle Geraden (nicht nur Ursprungsgeraden) gemeint, also alle affinen Teilräume der Form $v + U$ für eindimensionale Untervektorräume U .

KOROLLAR II.7. Sei $f: V \rightarrow W$ eine bijektive Abbildung zwischen $\mathbb{R}$ -Vektorräumen derselben Dimension, die die folgenden Eigenschaften hat:

- (a) $f(0) = 0$ und
- (b) f bildet jede Gerade in V auf eine Gerade in W ab.

Dann ist f eine lineare Abbildung.

Bekannte Sätze, die man in diesem Kontext zeigen kann sind der [Satz von Pappos](#)¹ und der [Satz von Desargues](#)².

II.2. Das Standard-Skalarprodukt auf $\mathbb{R}^n$

Wir betrachten nun als Grundkörper den Körper der reellen Zahlen, und fixieren $n \in \mathbb{N}$. Zuerst definieren wir, was wir unter dem Abstand zwischen zwei Punkten in $\mathbb{R}^n$ verstehen, und untersuchen diesen Begriff.

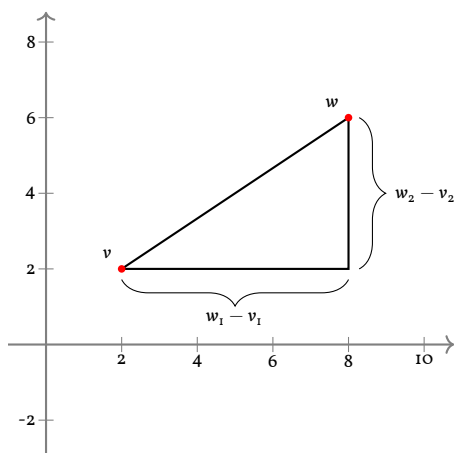
II.2.1. Der Abstand zwischen Punkten in $\mathbb{R}^n$.

DEFINITION II.8. Seien $v = (v_1, \dots, v_n)^t$ und $w = (w_1, \dots, w_n)^t$ Elemente von $\mathbb{R}^n$. Der Abstand zwischen den Punkten v und w (oder die Länge der Strecke zwischen v und w) ist

$$d(v, w) = \sqrt{\sum_{i=1}^n (w_i - v_i)^2}.$$

⊥

Da Quadrate reeller Zahlen niemals negativ sind, steht unter der Wurzel eine nicht-negative Zahl, und wir können die Wurzel (in $\mathbb{R}_{\geq 0}$) ziehen. Dieser Begriff spiegelt unsere Anschauung vom Abstand zwischen zwei Punkten wider – einerseits hat er die »üblichen Eigenschaften« einer Abstandsfunktion, andererseits ergibt sich der Begriff in natürlicher Weise aus dem Satz des Pythagoras.



Hier ergibt sich natürlich wieder das konzeptionelle Problem, dass wir schon von der Länge einer Strecke sprechen müssen, um den Satz des Pythagoras formulieren zu können. Zum jetzigen Zeitpunkt kann uns dieser also nur als Heuristik dienen. Immerhin zeigen Beweise wie in Abschnitt 3.3.1, dass der Satz des Pythagoras nur wenige Voraussetzungen an die zugrundeliegende »Geometrie« erfordert.

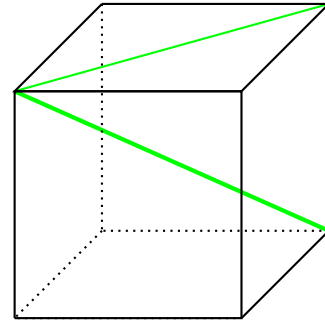
Dass der Abstand von Punkten $v, w \in \mathbb{R}^2$ durch die oben gegebene Formel berechnet/definiert werden kann, ist genau die Aussage des Satzes des Pythagoras für Dreiecke, deren Katheten zu den Koordinatenachsen parallel sind.

Für höherdimensionale Räume $\mathbb{R}^n$ lässt sich der Abstand schrittweise mithilfe des Satzes des Pythagoras in der gewünschten Form ausdrücken:

¹https://de.wikipedia.org/wiki/Satz_von_Pappos

²https://de.wikipedia.org/wiki/Satz_von_Desargues

Um den Abstand der »gegenüberliegenden« Punkte des hier abgebildeten Würfels, also die Länge der dick gezeichneten grünen Strecke auszudrücken, zeichnen wir die dünner gezeichneten grünen Hilfslinien ein. Zusammen bilden die grünen Strecken ein rechtwinkliges Dreieck; wir erwarten, dass für dieses der Satz des Pythagoras gelten soll. Die Länge der Diagonale durch die obere Seite des Würfels können wir wiederum mit dem Satz des Pythagoras ausdrücken. Kombiniert man alles, so ergibt sich die Formel, die wir oben zur Definition des Abstands benutzt haben.



SATZ 11.9. Die Abstandsfunktion d hat die folgenden Eigenschaften:

- (1) Für $v, w \in \mathbb{R}^n$ gilt $d(v, w) \in \mathbb{R}_{\geq 0}$ und $d(v, w) = 0$ genau dann, wenn $v = w$.
- (2) Für alle $v, w \in \mathbb{R}^n$ gilt $d(v, w) = d(w, v)$.
- (3) (Dreiecksungleichung) Für alle $u, v, w \in \mathbb{R}^n$ gilt

$$d(u, v) + d(v, w) \geq d(u, w)$$

BEWEIS. Die ersten beiden Eigenschaften sind klar. Die Dreiecksungleichung besagt anschaulich, dass in dem Dreieck mit Ecken u, v und w die Länge der Kante mit Endpunkten u und w höchstens so groß ist wie die Summe der beiden anderen Kanten (und man kann sich auch davon überzeugen, dass die Gleichheit nur dann eintreten kann, wenn u, v und w auf einer Geraden liegen). Das ist nicht schwierig (aber auch nicht ganz offensichtlich – versuchen Sie es einmal!). Das ganze ist etwas angenehmer aufzuschreiben, nachdem wir den Begriff des Skalarprodukts zweier Vektoren (und die entsprechende Notation) eingeführt haben. Wir verschieben daher den Beweis auf Abschnitt 11.2.3. $\square$

BEMERKUNG 11.10. Man sagt, eine Menge X zusammen mit einer Funktion $d: X \times X \rightarrow \mathbb{R}_{\geq 0}$, so dass

- (a) $d(x, y) = 0 \iff x = y$,
- (b) $d(x, y) = d(y, x)$ und
- (c) $d(x, y) + d(y, z) \geq d(x, z)$

für alle $x, y, z \in X$ sei ein *metrischer Raum*. Die Funktion d heißt dann eine *Metrik* oder *Abstandsfunktion* auf X .

Es gibt noch (viele...) andere Abstandsfunktionen auf $\mathbb{R}^n$. Die oben definierte Metrik heißt die *euklidische Metrik* auf $\mathbb{R}^n$. Da sie die Geometrie des Raums, in dem wir uns befinden, definiert, ist es die wichtigste Abstandsfunktion auf $\mathbb{R}^n$. $\diamond$

Der Beweis der folgenden Aussage ist offensichtlich.

LEMMA 11.11. Für alle $u, v, w \in \mathbb{R}^n$ gilt

$$d(u + w, v + w) = d(u, v).$$

Man sagt, der Abstand d sei *translationsinvariant*. Wir können also $d(v, w)$ ausdrücken als $d(0, w - v)$; es genügt daher, für alle $v \in \mathbb{R}^n$ den Abstand $d(0, v)$ zu kennen, den wir auch mit $\|v\|$ bezeichnen.

DEFINITION II.12. Für $v = (v_1, \dots, v_n)^t \in \mathbb{R}^n$ ist

$$\|v\| := \sqrt{\sum_{i=1}^n v_i^2}$$

die Norm (oder: die Länge) von v . +

Von der Länge zu sprechen, ist sinnvoll, wenn man sich einen Vektor als Pfeil (mit einer Richtung und einer Länge) vorstellt. Wie oben bemerkt können wir die Abstandsfunktion als $d(v, w) = \|w - v\|$ aus der Norm zurückgewinnen. Dementsprechend lassen sich auch die Eigenschaften des Abstands in Eigenschaften der Norm übersetzen:

SATZ II.13. Die Norm $\|\cdot\|: \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ hat die folgenden Eigenschaften:

- (1) $\|v\| = 0 \iff v = 0$,
- (2) $\|v + w\| \leq \|v\| + \|w\|$,
- (3) $\|av\| = |a| \|v\|$

für $v, w \in \mathbb{R}^n$ und $a \in \mathbb{R}$.

BEWEIS. Eigenschaften (1) und (3) sind klar. Eigenschaft (2), die wieder als Dreiecksungleichung bezeichnet wird, ist äquivalent zur Dreiecksungleichung für den euklidischen Abstand (die wir allerdings noch beweisen müssen). □

II.2.2. Orthogonalität und das Skalarprodukt. Aus dem Abstandsbegriff können wir auch ableiten, wann zwei Vektoren (bzw. die davon erzeugten Geraden) einen rechten Winkel bilden. Anschaulich sollte nämlich ein Dreieck durch seine Seitenlängen bis auf Kongruenz eindeutig bestimmt sein, und ein Dreieck mit Seitenlängen a, b, c genau dann ein rechtwinkliges Dreieck sein, wenn für die Seitenlängen die Gleichung $a^2 + b^2 = c^2$ aus dem Satz des Pythagoras gilt (wobei c hier die längste der drei Seiten bezeichne).

Wir machen daraus die folgende provisorische Definition (siehe Definition II.15): Vektoren $v, w \in \mathbb{R}^n$ stehen *senkrecht* aufeinander, wenn »für das Dreieck mit Ecken $0, v, w$ der Satz des Pythagoras gilt«, also genau dann, wenn

$$\|v\|^2 + \|w\|^2 = \|w - v\|^2$$

gilt.

Für Vektoren, die nicht zueinander senkrecht sind, können wir diese Überlegung dazu benutzen, die Abweichung davon zu »messen«. Wir definieren dafür:

DEFINITION II.14. Das (Standard-)Skalarprodukt auf $\mathbb{R}^n$ ist die Abbildung $\mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$, $(v, w) \mapsto v \cdot w$ mit

$$v \cdot w = \frac{1}{2}(\|v\|^2 + \|w\|^2 - \|w - v\|^2).$$

+

Für das Skalarprodukt sind auch viele andere Schreibweisen gebräuchlich, unter anderem vw , (v, w) , $\langle v, w \rangle$, $(v | w)$.

Sind $v = (v_1, \dots, v_n)^t$, $w = (w_1, \dots, w_n)^t$, so gilt konkret

$$v \cdot w = v^t w = \sum_{i=1}^n v_i w_i,$$

wobei wir für den Ausdruck in der Mitte v und w als $(n \times 1)$ -Matrizen verstehen und das Matrizenprodukt von v^t und w bilden.

Wir können also auch die Norm von v mithilfe des Skalarprodukts ausdrücken, es gilt

$$\|v\| = \sqrt{v \cdot v}.$$

Damit können wir die Definition, wann zwei Vektoren senkrecht aufeinander stehen, nun in der endgültigen Form angeben:

DEFINITION 11.15. Wir sagen, Vektoren $v, w \in \mathbb{R}^n$ seien *senkrecht* (oder: *orthogonal*) zueinander, wenn $v \cdot w = 0$ ist. $\dashv$

LEMMA 11.16. Das Skalarprodukt auf $\mathbb{R}^n$ hat die folgenden Eigenschaften.

- (1) $v \cdot w = w \cdot v$
- (2) $(av + a'v') \cdot w = a(v \cdot w) + a'(v' \cdot w),$
- (3) $v \cdot (aw + a'w') = a(v \cdot w) + a'(v \cdot w').$
- (4) $v \cdot v \geq 0$, und $v \cdot v = 0 \Leftrightarrow v = 0$.

BEWEIS. Alle Eigenschaften sind leicht nachzurechnen. Man sagt wegen Eigenschaft (1) auch, das Skalarprodukt sei *symmetrisch*. Eigenschaften (2) und (3) bedeuten, dass die Abbildung $\mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ multilinear im Sinne von Definition 9.1 ist. Da der Definitionsbereich aus zwei Faktoren besteht, spricht man auch von einer *bilinearen* Abbildung, bzw. von einer *Bilinearform*, weil der Wertebereich der Grundkörper $\mathbb{R}$ ist. Schließlich sagt man wegen Eigenschaft (4), dass diese Bilinearform *positiv definit* sei. Wir können das Lemma also zusammenfassen, indem wir sagen, das Standard-Skalarprodukt sei eine positiv definite symmetrische Bilinearform. In der Linearen Algebra werden wir die Theorie der Bilinearformen auf einem Vektorraum genauer studieren. $\square$

Wir sagen, dass zwei affine Geraden g, h (also Geraden, die nicht notwendig durch den Ursprung verlaufen), die sich in einem Punkt P schneiden, senkrecht aufeinander stehen, wenn die um P verschobenen Geraden $-P + g$ und $-P + h$ (hier handelt es sich um Ursprungsgeraden) von Vektoren erzeugt werden, die aufeinander senkrecht stehen, mit anderen Worten, wenn

$$g = P + \langle v \rangle, \quad h = P + \langle w \rangle, \quad v \cdot w = 0.$$

Die Vektoren v und w (die »Richtungsvektoren«) sind durch die Geraden g und h bis auf Vielfache eindeutig bestimmt und unabhängig vom Punkt P . Es ist also an dieser Stelle nicht erforderlich, den Schnittpunkt der Geraden zu verwenden (aber wir verlangen, dass sich die Geraden überhaupt schneiden).

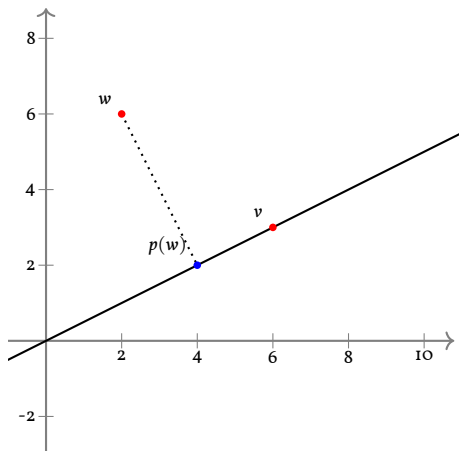
11.2.3. Die Ungleichung von Cauchy-Schwarz. Wir liefern nun den Beweis der Dreiecksungleichung nach (Satz 11.9).

Zunächst fassen wir noch einmal die geometrische Bedeutung des Skalarprodukts $v \cdot w$ zweier Vektoren v und w zusammen: Das Skalarprodukt $v \cdot w$ ist genau dann $= 0$, wenn v und w (bzw. die von diesen Vektoren erzeugten Geraden $\langle v \rangle, \langle w \rangle$) zueinander senkrecht sind. Im allgemeinen Fall misst das Skalarprodukt die Abweichung davon. Konkreter können wir das in der Ebene folgendermaßen beschreiben.

LEMMA 11.17. Sei $v \in \mathbb{R}^2, v \neq 0$. Sei $v' \in \mathbb{R}^2$ ein Vektor $\neq 0$, so dass $v \cdot v' = 0$. Dann bilden v, v' eine Basis von $\mathbb{R}^2$. Sei $p: \mathbb{R}^2 \rightarrow \langle v \rangle$ die eindeutig bestimmte lineare Abbildung mit $p(v) = v, p(v') = 0$, also die Projektion der Ebene auf die Gerade $\langle v \rangle$.

Dann gilt für alle $w \in \mathbb{R}^2$:

$$v \cdot w = v \cdot p(w) = \begin{cases} \|v\| \|p(w)\| & \text{falls } p(w) = av \text{ mit } a \geq 0, \\ -\|v\| \|p(w)\| & \text{falls } p(w) = av \text{ mit } a < 0. \end{cases}$$



Speziell gibt in dem Fall, dass $\|v\| = 1$ ist, das Skalarprodukt $v \cdot w$ bis auf das Vorzeichen einfach die Länge des Vektors an, der durch Projektion von w auf die Gerade $\langle v \rangle$ entsteht. Wir sehen daran auch schon, dass wir das Skalarprodukt der Vektoren v und w benutzen können, um den Winkel zwischen v und w (bzw. zwischen den Geraden $\langle v \rangle$ und $\langle w \rangle$) zu messen, siehe Abschnitt II.5.

BEWEIS. Wir schreiben $w = av + a'v'$. Dann gilt $p(w) = av$ und wegen $v \cdot v' = 0$, dass

$$v \cdot w = v \cdot p(w) = a \|v\|^2$$

und

$$\|p(w)\| = |a| \|v\|.$$

Daraus folgt die Behauptung. □

Zur Vorbereitung für den Beweis der Dreiecksungleichung beweisen wir das folgende etwas technische, aber sehr nützliche Resultat.

SATZ II.18 (Cauchy-Schwarzsche Ungleichung). Seien $v, w \in \mathbb{R}^n$. Dann gilt

$$|v \cdot w| \leq \|v\| \|w\|.$$

Die Gleichheit gilt genau dann, wenn v, w linear abhängig sind.

BEWEIS. Ist $v = 0$ oder $w = 0$, so sind beide Seiten der Ungleichung 0, und die Aussage des Satzes ist klar. Seien nun v und w verschieden von 0. Wenn wir v oder w mit einem Skalar aus $\mathbb{R}_{>0}$ multiplizieren, dann verändern sich beide Seiten um diesen Faktor, wir können daher ohne Einschränkung annehmen, dass $\|v\| = \|w\| = 1$ gilt. Indem wir gegebenenfalls v durch $-v$ ersetzen, können wir zusätzlich annehmen, dass $v \cdot w \geq 0$ gilt. Es genügt dann zu zeigen, dass $v \cdot w \leq 1$ ist, mit Gleichheit genau dann, wenn v und w linear abhängig sind.

Nun haben wir

$$0 \leq (v - w) \cdot (v - w) = v \cdot v - 2v \cdot w + w \cdot w = 2 - 2v \cdot w,$$

und daraus folgt $v \cdot w \leq 1$, wie gewünscht. Zudem gilt die Gleichheit nur dann, wenn $v = w$ ist. Nach den obigen Reduktionsschritten ist das äquivalent dazu, dass v und w linear abhängig sind. Jedenfalls ist klar, dass die Vektoren in diesem Fall linear abhängig sind. Dass für linear abhängige v, w die Gleichheit $\|v \cdot w\| \leq \|v\| \|w\|$ gilt, sieht man auch leicht direkt. □

Wir beweisen nun die Dreiecksungleichung.

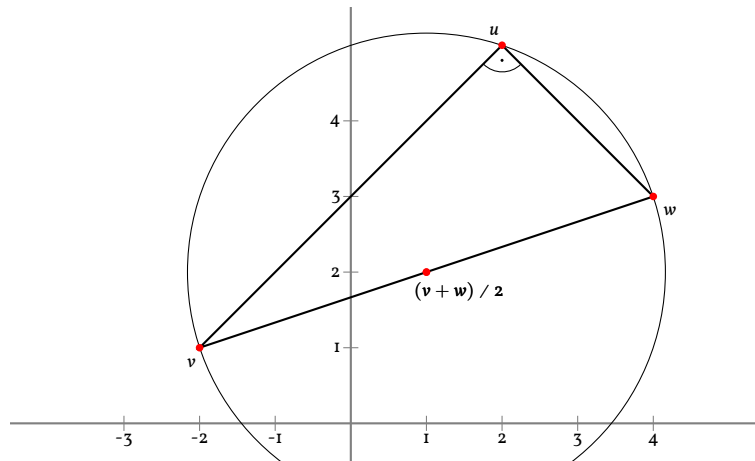


ABBILDUNG 1. Der Satz des Thales. Bei u ist ein rechter Winkel.

BEWEIS DER DREIECKSUNGLEICHUNG, SATZ 11.9. Seien u, v, w in $\mathbb{R}^n$ gegeben. Wir wollen zeigen, dass

$$d(u, v) + d(v, w) \geq d(u, w)$$

gilt.

Es genügt, die Abschätzung für die Quadrate der beiden Seiten zu zeigen, da es sich um nicht-negative reelle Zahlen handelt. Das Quadrat der linken Seite ist

$$\sum_{i=1}^n u_i^2 + \sum_{i=1}^n w_i^2 + 2d(u, v)d(v, w).$$

Das Quadrat der rechten Seite ist

$$\sum_{i=1}^n (w_i - u_i)^2.$$

Wegen der Translationsinvarianz des Abstands können wir ohne Einschränkung annehmen, dass $v = 0$ ist. Zu zeigen ist dann $2\|u\| \|w\| \geq -2w \cdot u$. Das folgt direkt aus der Cauchy-Schwarzschen Ungleichung. $\square$

11.2.4. Die Sätze von Pythagoras und von Thales. Wir wollen überprüfen, dass mit unseren Definitionen der Satz des Pythagoras gilt, und als weiteres einfaches Beispiel den Satz des Thales beweisen.

SATZ 11.19 (Satz des Pythagoras). Seien u, v, w Punkte in $\mathbb{R}^n$, so dass das Dreieck mit den Eckpunkten an der Ecke u einen rechten Winkel hat, das bedeutet $(v - u) \cdot (w - u) = 0$. Dann gilt

$$d(v, u)^2 + d(w, u)^2 = d(v, w)^2.$$

BEWEIS. Wegen der Translationsinvarianz des Abstands können wir ohne Einschränkung annehmen, dass $u = 0$ ist. Die Voraussetzung ist dann $v \cdot w = 0$, und wir können wie folgt rechnen.

$$\|v\|^2 + \|w\|^2 = v \cdot v + w \cdot w = v \cdot v - 2v \cdot w + w \cdot w = (v - w) \cdot (v - w) = \|w - v\|^2.$$

Das ist genau die Behauptung des Satzes. $\square$

Nun kommen wir zum Satz des Thales. Für $v, w \in \mathbb{R}^n$ ist $(v + w) / 2$ der Mittelpunkt der Strecke zwischen v und w (also der eindeutig bestimmte Punkt auf der Geraden durch v und w , der zu v und zu w denselben Abstand hat).

SATZ II.20 (Satz des Thales). Seien $u, v, w \in \mathbb{R}^2$, so dass u auf dem Kreis

$$\{x \in \mathbb{R}^n; d(x, (v+w)/2) = d(v, w)/2\}$$

mit Mittelpunkt $(v+w)/2$ und Radius $d(v, w)/2$ liegt. Dann bilden u, v, w ein rechtwinkliges Dreieck, mit rechtem Winkel an der Ecke u .

BEWEIS. Wir müssen zeigen, dass $(v-u) \cdot (w-u) = 0$ gilt. Die Voraussetzung besagt, dass

$$d(u, (v+w)/2) = d(v, w)/2,$$

also, nach Quadrieren beider Seiten,

$$\frac{1}{4}(v+w-2u) \cdot (v+w-2u) = \frac{1}{4}(w-v) \cdot (w-v),$$

das bedeutet

$$(v-u) \cdot (w-u) = u \cdot u + v \cdot w - u \cdot v - u \cdot w = 0,$$

wie gewünscht. □

II.2.5. Abstandserhaltende Abbildungen.

DEFINITION II.21. Eine lineare Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt *abstandserhaltend* (oder: eine *Isometrie*), wenn für alle $v, w \in \mathbb{R}^n$ gilt:

$$d(f(v), f(w)) = d(v, w).$$

⊥

Jede Isometrie f ist injektiv, denn aus $f(v) = 0$ folgt $0 = d(f(v), 0) = d(f(v), f(0)) = d(v, 0)$, also $v = 0$, und ein injektiver Endomorphismus eines endlich-dimensionalen Vektorraums ist notwendigerweise ein Isomorphismus.

BEMERKUNG II.22. Wir setzen in der Definition einer Isometrie voraus, dass f linear ist, weil wir vor allem lineare Abbildungen betrachten möchten. Teilweise wird eine Isometrie etwas allgemeiner als eine Abbildung g definiert, für die $d(g(v), g(w)) = d(v, w)$ für alle v, w gilt.

Der Unterschied ist allerdings nicht groß. Denn ist g eine Abbildung, die alle Abstände zwischen zwei Punkten erhält, dann ist g eine *affine* Abbildung, lässt sich also schreiben als Verkettung einer linearen abstandserhaltenden Abbildung f und einer Verschiebung $t_u: x \mapsto x + u$, d.h. $g = t_u \circ f (= f \circ t_{f^{-1}(u)})$, wobei $u = g(0)$ ist.

Mit anderen Worten: Ist g abstandserhaltend, aber nicht notwendig linear, so ist $f := t_{-g(0)} \circ g$ eine Isometrie in unserem Sinne. Geometrisch lässt sich das folgendermaßen begründen (allerdings an dieser Stelle etwas skizzenhaft). Die Strecke zwischen zwei Punkten ist die kürzeste Verbindungslinie; das kann man zu einem Beweis ausbauen, dass f Geraden (nicht nur Ursprungsgeraden) auf Geraden abbildet. Parallelität von Geraden kann man auch durch Abstände charakterisieren. Deshalb muss f auch Parallelogramme auf Parallelogramme abbilden. Da wir die Eigenschaft $f(0) = 0$ in unsere Definition von f eingebaut haben, folgt, dass f mit der Addition verträglich ist. Dann ergibt sich auch $f(nv) = f(v + \dots + v) = nf(v)$ für alle $n \in \mathbb{N}$ und mit etwas Zusatzarbeit auch $f(av) = af(v)$ für alle $a \in \mathbb{Q}$. Da man jede reelle Zahl zwischen rationalen Zahlen »einschachteln« kann, kann man daraus auch $f(av) = af(v)$ für alle $a \in \mathbb{R}$ ableiten. ◇

Wir können die Eigenschaft, eine Isometrie zu sein, auch in Termen der Norm, oder des Skalarprodukts umformulieren:

LEMMA II.23. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung. Dann sind äquivalent:

- (i) Die Abbildung f ist abstandserhaltend.
- (ii) Für alle $v \in \mathbb{R}^n$ gilt $\|f(v)\| = \|v\|$.
- (iii) Für alle $v, w \in \mathbb{R}^n$ gilt $f(v) \cdot f(w) = v \cdot w$.

BEWEIS. Sei f abstandserhaltend. Dann gilt insbesondere $d(o, f(v)) = d(o, v)$ für alle v , also (ii). Sei nun f eine lineare Abbildung mit Eigenschaft (ii). Weil wir das Skalarprodukt (wie in unserer Definition 11.14) durch die Norm von Vektoren ausdrücken können, folgt (iii). Ist schließlich Eigenschaft (iii) erfüllt, so folgt aus $d(v, w) = \sqrt{(w-v) \cdot (w-v)}$ (und entsprechend für $d(f(v), f(w))$), dass f auch die Bedingung (i) erfüllt. $\square$

In Termen der darstellenden Matrix $M(f)$ von f erhalten wir die folgende Charakterisierung:

LEMMA 11.24. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung. Dann sind äquivalent:

- (i) Die Abbildung f ist abstandserhaltend.
- (ii) Die Matrix $M(f)$ ist invertierbar und es gilt $M(f)^{-1} = M(f)^t$.

BEWEIS. Wir haben schon gesehen, dass eine abstandserhaltende Abbildung f ein Isomorphismus ist, folglich ist in diesem Fall $M(f)$ invertierbar. Wir müssen noch zeigen, dass $M(f)^t M(f) = E_n$ gilt. Wir bezeichnen dazu mit $e_1, \dots, e_n$ die Standardbasis von $\mathbb{R}^n$. Der Eintrag in Zeile i und Spalte j irgendeiner Matrix $A \in M_n(\mathbb{R})$ ist dann $e_i^t A e_j$. Für den entsprechenden Eintrag der Matrix $M(f)^t M(f)$ erhalten wir

$$e_i^t M(f)^t M(f) e_j = f(e_i)^t f(e_j) = f(e_i) \cdot f(e_j) = e_i \cdot e_j,$$

und dies ist 1 oder 0, je nachdem, ob $i = j$ oder $i \neq j$ – genau wie bei der Einheitsmatrix.

Ist andererseits $M(f)$ invertierbar mit inverser Matrix $M(f)^t$, so erhalten wir für $v, w \in \mathbb{R}^n$:

$$f(v) \cdot f(w) = f(v)^t f(w) = (M(f)v)^t (M(f)w) = v^t M(f)^t M(f)w = v \cdot w.$$

$\square$

SATZ 11.25. Die Teilmenge

$$O_n(\mathbb{R}) := \{A \in GL_n(\mathbb{R}); A^{-1} = A^t\} \subseteq GL_n(\mathbb{R})$$

bildet bezüglich des Matrizenprodukts eine Gruppe, die sogenannte orthogonale Gruppe. Ihre Elemente heißen orthogonale Matrizen.

BEWEIS. Es ist klar, dass $E_n \in O_n(\mathbb{R})$. Dass das Produkt orthogonaler Matrizen wieder die definierende Eigenschaft hat, ebenso wie die inverse Matrix einer orthogonalen Matrix, folgt aus den Formeln $(AB)^t = B^t A^t$ und $(A^t)^{-1} = (A^{-1})^t$. $\square$

Für die Ebene erhalten wir die folgende Beschreibung:

SATZ 11.26. Sei $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ eine Isometrie. Dann existieren $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$, so dass

$$M(f) = \begin{pmatrix} a & -b \\ b & a \end{pmatrix} \quad \text{oder} \quad M(f) = \begin{pmatrix} a & b \\ b & -a \end{pmatrix}.$$

BEWEIS. Wir setzen $\begin{pmatrix} a \\ b \end{pmatrix} = f(e_1)$. Dass $\|f(e_1)\| = 1$ gilt, sagt genau, dass $a^2 + b^2 = 1$.

Sei $A = M(f) = \begin{pmatrix} a & a' \\ b & b' \end{pmatrix}$. Aus

$$\det(A)^2 = \det(A) \det(A^t) = \det(AA^t) = \det(E_2) = 1$$

folgt $\det(A) \in \{1, -1\}$. (Wir benutzen hier grundlegende Eigenschaften der Determinante $\det(A) = \delta(A)$, wie sie in Kapitel 9 erklärt werden. Siehe auch Bemerkung 5.56; im (2×2) -Fall lassen sich die oben behaupteten Gleichheiten auch leicht direkt überprüfen.)

Damit erhalten wir für das Inverse von A :

$$A^{-1} = \begin{pmatrix} b' & -a' \\ -b & a \end{pmatrix} \quad \text{bzw.} \quad A^{-1} = \begin{pmatrix} -b' & a' \\ b & -a \end{pmatrix},$$

je nachdem, ob $\det(A) = 1$ oder $\det(A) = -1$ gilt. Weil $A^{-1} = A^t$ ist, folgt die Behauptung. $\square$

SATZ II.27. Für eine Isometrie $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ sind die folgenden Bedingungen äquivalent. Sind sie erfüllt, so nennen wir f eine Drehung.

- (i) $\det(f) = 1$
- (ii) Es existiert eine Isometrie $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit $f = g \circ g$.
- (iii) $M(f)$ hat die Form $\begin{pmatrix} a & -b \\ b & a \end{pmatrix}$ für $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$.
- (iv) $f = \text{id}$ oder $f = -\text{id}$ oder f besitzt keinen Eigenvektor, d.h. es existiert kein $v \in \mathbb{R}^2 \setminus \{0\}$, so dass $f(v) \in \langle v \rangle$ ist.

BEWEIS. Die Folgerung (ii) $\Rightarrow$ (i) folgt aus der Multiplikativität der Determinante. Die Implikationen (i) $\Leftrightarrow$ (iv) $\Leftrightarrow$ (iii) ergeben sich leicht aus Satz II.26.

Dass (ii) aus (iii) folgt, kann man durch eine direkte Rechnung überprüfen. Nachdem wir den Begriff des Winkels eingeführt haben, können wir das aber wesentlich leichter sehen, weil wir dann f als Drehung um einen Winkel α verstehen und g als die Drehung um den Winkel $\alpha / 2$ definieren können. $\square$

Die erste im Satz genannte Eigenschaft ist geeignet, um in beliebiger Dimension den Begriff der Drehung zu definieren. Siehe Ergänzung 7.59, Ergänzung 9.24.

- DEFINITION II.28.** (1) Eine Isometrie $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt *Drehung*, wenn $\det(f) = 1$ gilt.
- (2) Die Untergruppe von $O_n(\mathbb{R})$, die aus allen Drehungen besteht, mit anderen Worten der Kern der Gruppenhomomorphismus $O_n(\mathbb{R}) \rightarrow \mathbb{R}^\times, A \mapsto \det(A)$, heißt die *spezielle orthogonale Gruppe* und wird mit $SO_n(\mathbb{R})$ bezeichnet.

⊥

Ein anderer wichtiger Typ von Isometrien ist die Spiegelung an einem Untervektorraum der Dimension $n - 1$ in $\mathbb{R}^n$ (wir sprechen auch von einem Unterraum »der Kodimension 1« oder von einer *Hyperebene*).

DEFINITION II.29. Eine Isometrie $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt eine *Spiegelung*, wenn eine Hyperebene $U \subseteq \mathbb{R}^n$, d.h. ein Untervektorraum der Dimension $n - 1$, und ein Vektor $v \in \mathbb{R}^n$, der zu allen Elementen von U senkrecht ist, existieren, so dass $f(u) = u$ für alle $u \in U$ und $f(v) = -v$. ⊥

Für $\mathbb{R}^2$ erhalten wir, dass eine Isometrie genau dann eine Spiegelung ist, wenn ihre darstellende Matrix die Form $\begin{pmatrix} a & b \\ b & -a \end{pmatrix}$ für $a, b \in \mathbb{R}$ mit $a^2 + b^2 = 1$ hat.

Auf diesen Webseiten können Sie mit Drehungen und Spiegelungen der Ebene $\mathbb{R}^2$ »herumspielen«:

<https://math.ug/applets/drehungen.html>^a,
<https://math.ug/applets/spiegelungen.html>^b.

^a Drehungen

^b Spiegelungen

11.3. Weglängen und die trigonometrischen Funktionen

Wir beginnen mit einer heuristischen Diskussion, um den weiteren Fahrplan zu erläutern. Wir möchten einer Drehung der Ebene um den Ursprung einen *Winkel* α zuordnen, so dass wir eine Bijektion zwischen der Menge der Winkel und der Menge der Drehungen erhalten. Dabei möchten wir gerne der Drehung ρ als Winkel die »Länge des Kreisbogens vom Punkt e_1 zum Punkt $\rho(e_1)$ auf dem Einheitskreis« (gegen den Uhrzeigersinn) zuordnen. Als Einheitskreis bezeichnen wir die Menge S^1 aller Punkte in $\mathbb{R}^2$ mit Abstand 1 zum Ursprung $(0, 0)^t$.

Wenn wir mit π die Hälfte des Umfangs des Einheitskreises bezeichnen, bedeutet das, dass wir zu jeder Zahl im halboffenen Intervall $[0, 2\pi)$ den in e_1 beginnenden Kreisbogen der entsprechenden Länge und den zugehörigen Winkel betrachten können.

Die Verkettung der Drehungen um den Winkel α und den Winkel β ist dann die Drehung um den Winkel $\alpha + \beta$ (bzw. um $\alpha + \beta - 2\pi$, falls $\alpha + \beta \geq 2\pi$ ist). Diese Additivitätseigenschaft ist neben der einfachen Möglichkeit, Winkel zu »messen« und dadurch zu vergleichen, ein wichtiger Grund, den Winkelbegriff einzuführen.

Wir definieren dann die Kosinus- und Sinusfunktion dadurch, dass für $\alpha \in [0, 2\pi)$ der Punkt $(\cos(\alpha), \sin(\alpha))^t$ der eindeutig bestimmte Punkt P ist, so dass die Länge des Kreisbogens von e_1 bis P auf dem Einheitskreis gegen den Uhrzeigersinn gleich α ist. Indem wir die Funktionen durch die Regeln $\cos(\alpha + 2k\pi) = \cos(\alpha)$ und $\sin(\alpha + 2k\pi) = \sin(\alpha)$ für $k \in \mathbb{Z}$ auf ganz $\mathbb{R}$ »periodisch« fortsetzen, erhalten wir die üblichen trigonometrischen Funktionen.

Die Drehung ρ_α um den Winkel α ist demzufolge durch die Matrix

$$\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}$$

gegeben, und dass die Verkettung $\rho_\alpha \circ \rho_\beta$ die Drehung um den Winkel $\alpha + \beta$ ist, entspricht der Gleichheit

$$\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix} \begin{pmatrix} \cos(\beta) & -\sin(\beta) \\ \sin(\beta) & \cos(\beta) \end{pmatrix} = \begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix}$$

von Matrizen, oder explizit ausgeschrieben den Gleichheiten

$$\cos(\alpha) \cos(\beta) - \sin(\alpha) \sin(\beta) = \cos(\alpha + \beta)$$

$$\sin(\alpha) \cos(\beta) + \cos(\alpha) \sin(\beta) = \sin(\alpha + \beta),$$

die man als die *Additionstheoreme* für $\cos$ und $\sin$ bezeichnet.

Wenn Sie bereit sind zu glauben, dass man jedem Kreisbogen »in vernünftiger Weise« eine Länge zuordnen kann, so dass – mit 2π als Umfang des Einheitskreises – jede Länge zwischen $[0, 2\pi)$ als Länge eines eindeutig bestimmten Kreisbogens auftritt, und man auf diese Art und Weise wie oben beschrieben stetige Funktionen $\cos$ und $\sin$ erhält, dann können Sie die folgenden Abschnitte, die der formalen Herleitung dieser Tatsachen dienen, erstmal überspringen.

II.3.1. Länge einer Kurve. Wir müssen, um den Begriff des Winkels so definieren zu können, wie wir es vorhaben, nicht nur die Länge einer Strecke, sondern auch die Länge eines »Bogens«, also entlang einer »Kurve«, messen können – zum Beispiel den Umfang des Einheitskreises $S^1 := \{x \in \mathbb{R}^2; \|x\| = 1\}$ oder eines Teils davon.

Wie man die Länge solcher Wege messen könnte, war lange Zeit ein wichtiges ungelöstes mathematisches Problem.

Es würde zu weit führen, die entsprechende Theorie hier vollständig auszuarbeiten, aber wir skizzieren die Kernpunkte. Weitere Details finden Sie in der Literatur, siehe die Hinweise am Ende dieses Abschnitts.

Wie üblich bezeichnen wir für reelle Zahlen $a < b$ mit $[a, b] := \{x \in \mathbb{R}; a \leq x \leq b\}$ das abgeschlossene Intervall mit den Endpunkten a und b , mit $(a, b) := \{x \in \mathbb{R}; a < x < b\}$ das offene Intervall und mit $[a, b)$ bzw. $(a, b]$ die entsprechenden halboffenen Intervalle.

Unser erstes Ziel ist es, einen vernünftigen Begriff für die Länge eines »Weges« in $\mathbb{R}^n$ anzugeben. Unter einem Weg verstehen wir dabei Folgendes:

DEFINITION II.30. Sei $n \in \mathbb{N}$. Ein *Weg* ist eine stetige Abbildung $\gamma: [a, b] \rightarrow \mathbb{R}^n$ für $a < b \in \mathbb{R}$. $\dashv$

Dass eine Abbildung $\gamma: [a, b] \rightarrow \mathbb{R}^n$ stetig ist, ist damit gleichbedeutend (und Sie können das dementsprechend als Definition nehmen), dass alle Verkettungen $p_i \circ \gamma: [a, b] \rightarrow \mathbb{R}$ stetig sind, wobei $p_i: \mathbb{R}^n \rightarrow \mathbb{R}$, $(x_1, \dots, x_n)^t \mapsto x_i$, die Projektion auf den i -ten Eintrag bezeichnet.

Für uns ist nur der Fall, dass γ injektiv ist, von Bedeutung, und es schadet nicht, wenn Sie das von vorneherein annehmen. Formal ist es im Moment aber noch nicht erforderlich.

DEFINITION II.31. Wir nennen einen Weg $\gamma: [a, b] \rightarrow \mathbb{R}^n$ *rektifizierbar*, wenn das Supremum

$$L(\gamma) = \sup_Z \sum_{i=1}^{l_Z} d(\gamma(z_i), \gamma(z_{i-1}))$$

über alle Zerlegungen $Z = (a = z_0 < z_1 < \dots < z_{l_Z} = b)$ von $[a, b]$ in $\mathbb{R}$ existiert.

In diesem Fall nennen wir $L(\gamma)$ die *Länge* des Weges γ (und schreiben $L_a^b(\gamma)$, wenn es vorteilhaft ist, den betrachteten Anfangs- und Endpunkt explizit anzugeben). $\dashv$

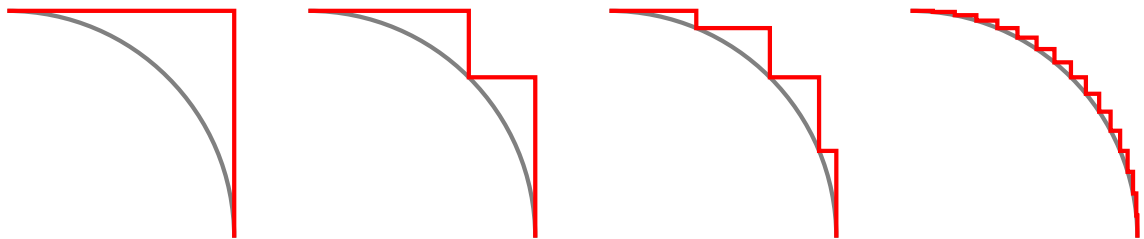


ABBILDUNG 2. Wir wollen die Länge einer »Kurve« durch Approximation durch Strecken definieren. Die Länge einer Strecke ist der Abstand zwischen den beiden Endpunkten der Strecke, den wir bereits in zufriedenstellender Weise definiert haben. Die Abbildungen hier zeigen, dass man etwas vorsichtig sein muss, um eine geeignete Form der Approximation auszuwählen: In allen vier Abbildungen ist die Länge des roten Streckenzuges gleich 2. Auch wenn man den Kreisbogen so »immer besser« approximiert, ist dennoch der Umfang des Einheitskreises nicht gleich 8 ...

Wir nennen eine Funktion $\gamma: [a, b] \rightarrow \mathbb{R}^n$ (komponentenweise) stetig differenzierbar auf $[a, b]$, wenn für alle Projektionen p_i die Verkettung $\gamma_i := p_i \circ \gamma$ auf dem $[a, b]$ differenzierbar ist, und die Ableitungsfunktion stetig ist. In den Randpunkten a und b betrachten wir dazu die einseitige Ableitung und fordern die einseitige Stetigkeit.

Wir bezeichnen dann mit $\gamma': [a, b] \rightarrow \mathbb{R}^n$ die Funktion $x \mapsto (\gamma'_1(x), \dots, \gamma'_n(x))^t$.

SATZ 11.32. Sei $\gamma: [a, b] \rightarrow \mathbb{R}^n$ (komponentenweise) stetig differenzierbar. Dann ist γ rektifizierbar und für die Länge $L_a^b(\gamma)$ gilt

$$L_a^b(\gamma) = \int_a^b \|\gamma'(x)\| dx.$$

BEWEIS. Für den Beweis verweisen wir auf die Literatur (konkrete Referenzen befinden sich weiter unten), hier nur ein paar Worte dazu. Der Hauptsatz der Differential- und Integralrechnung und die Dreiecksungleichung für Integrale liefern für Punkte $z_{i-1}, z_i \in [a, b]$, $z_{i-1} < z_i$, dass

$$\|\gamma(z_i) - \gamma(z_{i-1})\| = \left\| \int_{z_{i-1}}^{z_i} \gamma'(t) dt \right\| \leq \int_{z_{i-1}}^{z_i} \|\gamma'(t)\| dt,$$

wobei wir den Ausdruck $\int_{z_{i-1}}^{z_i} \gamma'(t) dt$ als den Vektor $\left(\int_{z_{i-1}}^{z_i} \gamma'_1(t) dt, \dots, \int_{z_{i-1}}^{z_i} \gamma'_n(t) dt \right)^t$ verstehen. Dementsprechend benötigt man im zweiten Schritt eine »mehrdimensionale« Dreiecksungleichung für das Integral. Dies liefert durch Aufsummieren über i die Abschätzung $L_a^b(\gamma) \leq \int_a^b \|\gamma'(x)\| dx$.

Für die andere Abschätzung zeigt man, dass für jedes $\varepsilon > 0$ eine Zerlegung $Z = (a = z_0 < z_1 < \dots < z_l = b)$ von $[a, b]$ existiert, so dass $\int_a^b \|\gamma'(x)\| dx \leq \sum_{i=1}^l d(\gamma(z_i), \gamma(z_{i-1})) + \varepsilon$ gilt. Dazu kann man ausnutzen, dass die stetige Funktion γ auf dem kompakten Intervall sogar gleichmäßig stetig ist und sich daher durch Streckenzüge mit hinreichend vielen Stützpunkten beliebig genau annähern lässt. $\square$

Sind $\gamma_1: [a, b] \rightarrow \mathbb{R}^n$ und $\gamma_2: [b, c] \rightarrow \mathbb{R}^n$ rektifizierbare Wege mit $\gamma_1(b) = \gamma_2(b)$, so ist der Weg $\gamma: [a, c] \rightarrow \mathbb{R}^n$, der durch »Zusammensetzen« der Wege γ_1 und γ_2 entsteht (d.h. $\gamma(t) = \gamma_1(t)$ für $t \in [a, b]$, $\gamma(t) = \gamma_2(t)$ für $t \in [b, c]$) ebenfalls rektifizierbar, und es gilt $L_a^c(\gamma) = L_a^b(\gamma_1) + L_b^c(\gamma_2)$. Wenn man das ausnutzt, kann man den vorherigen Satz auch anwenden, um die Länge von Wegen durch Integrale auszudrücken, die durch Zusammensetzen von stetig differenzierbaren Wegen entstehen, die aber an den »Verbindungsstellen« nicht unbedingt stetig differenzierbar sein müssen.

Nun wollen wir von der Weglänge übergehen zur Länge des Bildes $\text{Im}(\gamma)$ eines Weges γ . Es ist klar, dass nur dann vernünftig ein Zusammenhang zur Weglänge hergestellt werden kann, wenn der betrachtete Weg (jedenfalls »im wesentlichen«, siehe unten) eine injektive Abbildung ist (also wenn wir nicht »unterwegs umkehren, ein Stück zurücklaufen, und dann auf demselben Weg weiterlaufen«). Daher definieren wir

DEFINITION 11.33. Eine Teilmenge $\Gamma \subseteq \mathbb{R}^n$ heißt *Jordanbogen in $\mathbb{R}^n$* , wenn ein injektiver Weg $\gamma: [a, b] \rightarrow \mathbb{R}^n$ mit $\text{Im}(\gamma) = \Gamma$ existiert. $\dashv$

Um die Länge eines Jordanbogens Γ definieren zu können, ist entscheidend, dass alle injektiven Wege mit Bild Γ dieselbe Weglänge haben. Das werden wir aus dem folgenden Satz folgern können:

SATZ 11.34. Seien $\gamma_1: [a_1, b_1] \rightarrow \mathbb{R}^n$, $\gamma_2: [a_2, b_2] \rightarrow \mathbb{R}^n$ injektive Wege mit $\Gamma = \text{Im}(\gamma_1) = \text{Im}(\gamma_2)$. Dann existiert eine stetige und streng monotone Abbildung $\varphi: [a_1, b_1] \rightarrow [a_2, b_2]$ mit $\gamma_1 = \gamma_2 \circ \varphi$.

Zum Beweis definiert man $\varphi(t) = \gamma_2^{-1}(\gamma_1(t))$, was möglich ist, weil γ_2 eine Bijektion $[a_2, b_2] \rightarrow \Gamma$ ist und daher eine Umkehrfunktion besitzt. Es ist nicht schwer nachzuprüfen, dass φ stetig und streng monoton ist.

KOROLLAR II.35. Seien $\gamma_1: [a_1, b_1] \rightarrow \mathbb{R}^n$, $\gamma_2: [a_2, b_2] \rightarrow \mathbb{R}^n$ injektive Wege mit $\Gamma = \text{Im}(\gamma_1) = \text{Im}(\gamma_2)$. Dann gilt: γ_1 ist genau dann rektifizierbar, wenn γ_2 rektifizierbar ist, und in diesem Fall gilt $L_{a_1}^{b_1}(\gamma_1) = L_{a_2}^{b_2}(\gamma_2)$.

In diesem Fall nennen wir den Jordanbogen Γ rektifizierbar und definieren seine Länge als die Zahl $L(\Gamma) = L_{a_1}^{b_1}(\gamma_1) = L_{a_2}^{b_2}(\gamma_2)$.

Mit etwas zusätzlicher Arbeit kann man dieses Korollar, insbesondere also die Definition der Länge eines Jordanbogens verallgemeinern auf den Fall, dass γ injektiv auf $[a, b]$ ist und dass $\gamma(a) = \gamma(b)$ gilt. Man spricht dann von einem *geschlossenen* Jordanbogen.

Ist f eine Isometrie oder eine Verschiebung, dann ist die Verkettung eines rektifizierbaren Weges mit f ein rektifizierbarer Weg derselben Länge (denn f erhält die Länge von Strecken), und daher bildet f Jordanbögen auf Jordanbögen derselben Länge ab.

Der Begriff des rektifizierbaren Weges und die Bogenlänge werden in vielen Lehrbüchern zur Analysis besprochen, zum Beispiel

O. Deiser, [Analysis 2](#)^a, Kapitel 3.2.

H. Heuser, *Lehrbuch der Analysis, Teil 2*, Teubner 1993, Abschnitte 177, 178

K. Königsberger, *Analysis I*, Springer 1992, Kapitel 13

L. Rogge, [Vorlesungsskript Analysis](#)^b, Kapitel IX.

^a<https://aleph1.info/?call=Puc&permalink=analysis2>

^bhttps://www.uni-due.de/%7Ehn213me/sk/rogge/Ana_1-2.pdf

II.3.2. Der Einheitskreis als Jordanbogen und die Zahl π . Wir können nun zeigen, dass der Einheitskreis ein Jordanbogen ist. Wenn man die trigonometrischen Funktionen $\sin$ und $\cos$ schon kennt, und beispielsweise die Zahl π so definiert, dass $\frac{\pi}{2}$ die kleinste positive Nullstelle der Kosinusfunktion ist, dann kann man den Weg $\gamma: [0, 2\pi] \rightarrow \mathbb{R}^2$, $t \mapsto (\cos(t), \sin(t))^t$ benutzen, dessen Bild der Einheitskreis ist und der auf $[0, 2\pi]$ injektiv ist, um dies zu zeigen, und sieht dann auch direkt, dass der Einheitskreis Umfang 2π hat.

Wir wollen das aber hier noch nicht benutzen. Wir können dann zwar nicht einen stetig differenzierbaren Weg mit Bild S^1 angeben, aber es genügt ja, den Einheitskreis in einige Stücke zu zerlegen und für die einzelnen Stücke nachzuweisen, dass es sich um Jordanbögen handelt.

Die Funktion $\gamma: [-1, 1] \rightarrow \mathbb{R}^2$, $t \mapsto (t, \sqrt{1-t^2})^t$, hat als Bild die »obere Hälfte« des Einheitskreises, also die Menge $\{(x, y)^t \in \mathbb{R}^2; x^2 + y^2 = 1, y \geq 0\}$. Mit der Vorschrift $t \mapsto -\sqrt{1-t^2}$ erhalten wir die untere Hälfte, also die Punkte $(x, y)^t$ des Einheitskreises mit $y \leq 0$.

Allerdings sind diese beiden Wege nur auf $(-1, 1)$, nicht jedoch auch $[-1, 1]$ stetig differenzierbar. Wir können aber auch die Wege $t \mapsto (\sqrt{1-t^2}, t)^t$ und analog $t \mapsto (-\sqrt{1-t^2}, t)^t$ betrachten, die für $t \in [-1, 1]$ die rechte bzw. linke Hälfte des Einheitskreises (also die Kreisbögen mit $x \geq 0$, bzw. mit $x \leq 0$) als Bild haben und auch auf dem offenen Intervall $(-1, 1)$ stetig differenzierbar sind. Man kann dann diese Wege geeignet zusammensetzen, beispielsweise indem man jeden von ihnen für ein Viertel des Einheitskreises »benutzt«.

Insgesamt sehen wir

SATZ II.36. Der Einheitskreis $S^1 \subset \mathbb{R}^2$ ist ein geschlossener Jordanbogen.

Als ersten wichtigen Schritt zum Winkelbegriff können wir nun die Zahl π definieren.

DEFINITION 11.37. Die Zahl $\pi \in \mathbb{R}$ ist definiert als die Hälfte der Länge des Jordanbogens S^1 . $\dashv$

Anhand unserer Definition ist klar, dass der Umfang eines regelmäßigen n -Ecks, dessen Eckpunkte alle auf dem Einheitskreis liegen, eine untere Schranke für 2π ist. Es ist nicht schwer zu zeigen, dass die Folge dieser Umfänge für $n \rightarrow \infty$ gegen 2π konvergiert.

11.4. Die trigonometrischen Funktionen

Als nächstes möchten wir zeigen, dass es für jede Zahl $\alpha \in [0, 2\pi)$ genau einen Punkt P auf dem Einheitskreis gibt, so dass die Länge des Kreisbogens von 0 nach P (gegen den Uhrzeigersinn) gleich α ist. Dann können wir Funktionen $\cos: [0, 2\pi) \rightarrow \mathbb{R}$, $\sin: [0, 2\pi) \rightarrow \mathbb{R}$ definieren durch

$$P =: \begin{pmatrix} \cos(\alpha) \\ \sin(\alpha) \end{pmatrix},$$

d.h. $\cos(\alpha)$ und $\sin(\alpha)$ sind die Koordinaten dieses Punktes P , der einen Kreisbogen der Länge α vom Punkt $(1, 0)^t$ aus markiert. Wir nennen die Funktionen $\cos: [0, 2\pi) \rightarrow \mathbb{R}$ und $\sin: [0, 2\pi) \rightarrow \mathbb{R}$ die Kosinus- und Sinusfunktion.

Dann werden wir sagen, dass der positive Teil der x -Achse und die Halbgerade durch 0 und P , also die Menge aller nicht-negativen Vielfachen von P , den Winkel α einschließen.

Wir können die Funktionen $\cos$ und $\sin$ dann durch $\cos(x + 2k\pi) := \cos(x)$, $\sin(x + 2k\pi) := \sin(x)$ für alle $k \in \mathbb{Z}$ periodisch fortsetzen und erhalten so die übliche Kosinus- und Sinusfunktion.

Es gibt grundsätzlich zwei mögliche Ansätze, um dieses Programm durchzuführen:

- (1) Man definiert zunächst die Sinus- und Kosinusfunktionen auf »analytischem« Wege und zeigt, dass es sich um differenzierbare Abbildungen handelt, berechnet ihre Ableitungen und daraus weitere Eigenschaften, und zeigt dann, mithilfe von Satz 11.32, dass die Länge des Kreisbogens von 0 bis zum Punkt $(\cos(\alpha), \sin(\alpha))^t$ gleich α ist (für $\alpha \in [0, 2\pi)$).

Je nach Geschmack ist ein gefühlter Nachteil dieses Ansatzes, dass die Definition von Sinus und Kosinus nicht geometrisch motiviert ist sondern ein bisschen »vom Himmel fällt«. Mit der unten vorgestellten Definition mit der komplexen Exponentialfunktion bleibt man aber ziemlich nah an der geometrischen Anschauung. Mit einem gewissen Aufwand ist es verbunden zu zeigen, dass die Sinus- und Kosinusfunktionen überhaupt periodisch sind. Alle weiteren gewünschten Eigenschaften ergeben sich ohne größere Schwierigkeiten.

- (2) Alternativ kann man den Weg $[-1, 1] \rightarrow \mathbb{R}^2, t \mapsto (t, \sqrt{1-t^2})^t$, der den Halbkreis des Einheitskreises oberhalb der x -Achse parametrisiert, und seine Varianten benutzen (mit denen wir schon gezeigt haben, dass der Einheitskreis ein Jordanbogen ist), um die Sinus- und Kosinusfunktionen zu definieren. Während dieser Ansatz auf den ersten Blick natürlicher erscheint, ist er in der Durchführung an mehreren Stellen lästiger. Unter anderem deshalb, weil der dieser Weg bei 1 und -1 nicht differenzierbar ist.

11.4.1. Geometrie der komplexen Zahlenebene. Es wird im folgenden nützlich sein, die reelle Ebene $\mathbb{R}^2$ mit der komplexen Zahlenebene zu identifizieren. Wir betrachten die komplexen Zahlen $\mathbb{C}$ als zweidimensionalen $\mathbb{R}$ -Vektorraum. Im folgenden fixieren wir die Basis $1, i$ von $\mathbb{C}$, mit anderen Worten den $\mathbb{R}$ -Vektorraum-Isomorphismus $\iota: \mathbb{C} \rightarrow \mathbb{R}^2$, der 1 auf e_1 und i auf e_2 abbildet. Eine komplexe Zahl $a + bi$ wird also auf den Vektor $(a, b)^t$ abgebildet. (Das entspricht ja auch genau unserer ursprünglichen Konstruktion der komplexen Zahlen.)

Die Multiplikation mit $z = a + bi$, also die Abbildung $\mathbb{C} \rightarrow \mathbb{C}, x \mapsto zx$, ist ein Homomorphismus von $\mathbb{R}$ -Vektorräumen, der bezüglich der Basis $(1, i)$ durch die Matrix

$$\begin{pmatrix} a & -b \\ b & a \end{pmatrix}$$

beschrieben wird. Vergleiche Bemerkung 5.4.2.

Die Spiegelung an der Gerade $\mathbb{R} \subset \mathbb{C}$ ist dann die Abbildung $a + bi \mapsto a - bi$; wir bezeichnen sie als die *komplexe Konjugation* und nennen $a - bi$ die zu $a + bi$ *komplex konjugierte Zahl*. Üblicherweise wird die zu $z \in \mathbb{C}$ komplex konjugierte Zahl mit $\bar{z}$ bezeichnet.

Für eine komplexe Zahl z heißt die Länge des Vektors $\iota(z) \in \mathbb{R}^2$, also der Abstand $d(0, \iota(z))$ von $\iota(z)$ zum Ursprung, der *komplexe Absolutbetrag* von z . Explizit gilt für $z = a + bi$, dass $\|z\| = \sqrt{a^2 + b^2}$, was wir auch als $\sqrt{z\bar{z}}$ schreiben können. Meist schreibt man $|z|$ statt $\|z\|$ für den Betrag von z .

Den Einheitskreis können wir dann als die Menge aller komplexen Zahlen vom Betrag 1 betrachten.

Ist $z = a + bi$ eine komplexe Zahl mit $|z| = 1$, dann ist $a^2 + b^2 = 1$. Die Multiplikation mit z ist dann eine Drehung. Im allgemeinen können wir jede komplexe Zahl $z \neq 0$ schreiben als $|z| \cdot \frac{z}{|z|}$, und $\frac{z}{|z|}$ hat Betrag 1. Die Multiplikation mit z ist also die Verkettung der zu $\frac{z}{|z|}$ gehörigen Drehung und der Streckung um den (reellen) Faktor $|z|$.

II.4.2. Die trigonometrischen Funktionen und die komplexe Exponentialfunktion. Am elegantesten (und am schnellsten, wenn man hinreichende Vorkenntnisse über die komplexen Zahlen hat) lassen sich die trigonometrischen Funktionen mithilfe der komplexen Exponentialfunktion einführen. Dies soll hier skizziert werden, und wir beginnen mit der Definition der Exponentialfunktion als

$$\exp: \mathbb{C} \rightarrow \mathbb{C}, \quad x \mapsto \sum_{n=0}^{\infty} \frac{x^n}{n!}.$$

Die hier auftretende Reihe ist für jedes $x \in \mathbb{C}$ absolut konvergent.

SATZ II.38 (Eigenschaften der komplexen Exponentialfunktion). (1) Es ist $\exp(0) = 1$.

(2) Die Funktion $\exp$ ist auf ganz $\mathbb{C}$ komplex differenzierbar, d.h. für alle $x \in \mathbb{C}$ existiert der Grenzwert $\lim_{h \rightarrow 0} \frac{\exp(x+h) - \exp(x)}{h}$, wobei $h \in \mathbb{C}$ ist. Als Ableitungsfunktion erhalten wir

$$\exp'(x) = \exp(x),$$

die Exponentialfunktion ist also gleich ihrer Ableitung.

(3) (Funktionalgleichung) Für alle $x, y \in \mathbb{C}$ gilt

$$\exp(x+y) = \exp(x) \exp(y).$$

Insbesondere gilt $\exp(x) \in \mathbb{C}^\times$ für alle $x \in \mathbb{C}$ und $\exp: \mathbb{C} \rightarrow \mathbb{C}^\times$ ist ein Gruppenhomomorphismus.

(4) Für $t \in \mathbb{R}$ gilt $|\exp(it)| = 1$.

BEWEIS. zu (1). Das folgt direkt aus der Definition von $\exp$ als Reihe.

zu (2). Um die Differenzierbarkeit zu zeigen, kann man, ähnlich wie im Reellen, das Restglied der Reihe geeignet abschätzen, um zu zeigen, dass man die Ableitung der »Potenzreihe« $\sum_{n \geq 0} \frac{x^n}{n!}$ summandenweise bilden darf, obwohl es sich um eine unendliche Reihe handelt. Damit ist dann auch klar, dass $\exp' = \exp$ gilt.

zu (3). Als erstes leiten wir die Funktion $x \mapsto \exp(x) \exp(-x)$ ab. Als Ableitung erhalten wir (wegen (2), der Produktregel und der Kettenregel, die auch für komplex differenzierbare Funktionen gelten) die Nullfunktion. Das impliziert, dass die betrachtete Funktion konstant

ist, und indem wir $x = 0$ einsetzen, sehen wir, dass der Wert für alle $x \in \mathbb{C}$ gleich 1 ist. Wir haben damit bewiesen, dass $\exp(-x) = \exp(x)^{-1}$ und insbesondere $\exp(x) \neq 0$ für alle $x \in \mathbb{C}$ gilt.

Wir fixieren nun $y \in \mathbb{C}$ und betrachten die Funktion

$$\mathbb{C} \rightarrow \mathbb{C}, \quad x \mapsto \frac{\exp(x+y)}{\exp(x)},$$

von der wir gerne zeigen möchte, dass sie konstant ist und den Wert $\exp(y)$ hat. Die Funktion ist komplex differenzierbar und hat bei $x = 0$ wegen Teil (1) den gewünschten Wert $\exp(y)$. Es genügt daher zu zeigen, dass die Ableitung dieser Funktion die Nullfunktion ist. Für die Ableitung erhalten wir nach (2) und unter Verwendung von Quotienten- und Kettenregel

$$\frac{\exp(x+y) \exp(x) - \exp(x+y) \exp(x)}{\exp(x)^2} = 0.$$

Alternativ kann man die Funktionalgleichung auch direkt anhand der Reihendarstellung beweisen, indem man ein geeignetes Ergebnis über das Produkt von unendlichen Reihen heranzieht.

Teil (4) folgt direkt aus Teil (3), denn

$$|\exp(it)| = \sqrt{\exp(it)\exp(it)} = \sqrt{\exp(it)\exp(i\bar{t})} = \sqrt{\exp(it)\exp(-it)} = 1,$$

wobei wir auch noch benutzt haben, dass $\overline{\exp(z)} = \exp(\bar{z})$ für alle $z \in \mathbb{C}$ gilt. Das liegt daran, dass die komplexe Konjugation eine stetige Abbildung $\mathbb{C} \rightarrow \mathbb{C}$ ist und daher mit der Grenzwertbildung der unendlichen Reihe, die die Exponentialfunktion definiert, vertauscht werden darf. $\square$

DEFINITION II.39. (1) Die Kosinusfunktion ist definiert als

$$\cos: \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto \cos(t) := \text{Realteil von } \exp(it).$$

(2) Die Sinusfunktion ist definiert als

$$\sin: \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto \sin(t) := \text{Imaginärteil von } \exp(it).$$

Für alle $t \in \mathbb{R}$ gilt also $\exp(it) = \cos(t) + i \sin(t)$. $\dashv$

Indem man die entsprechenden Eigenschaften der komplexen Exponentialfunktion, eingeschränkt auf Argumente der Form it , $t \in \mathbb{R}$, in ihren Real- und Imaginärteil »zerlegt«, erhält man:

KOROLLAR II.40 (Eigenschaften von sin und cos). (1) Für alle $x \in \mathbb{R}$ gilt

$$\cos(x)^2 + \sin(x)^2 = 1.$$

(2) Die Funktionen sin und cos sind auf ganz $\mathbb{R}$ differenzierbar und es gilt

$$\sin'(x) = \cos(x), \quad \cos'(x) = -\sin(x).$$

(3) Additionstheoreme. Für alle $s, t \in \mathbb{R}$ gilt

$$\begin{aligned} \cos(s+t) &= \cos(s)\cos(t) - \sin(s)\sin(t), \\ \sin(s+t) &= \cos(s)\sin(t) + \sin(s)\cos(t). \end{aligned}$$

BEMERKUNG II.41. Wir erhalten hier eine weitere Rechtfertigung dafür, Winkel als Zahl in $[0, 2\pi)$ zu messen (und nicht etwa als Zahl zwischen 0 und 360). Natürlich kann man zwischen beiden Systemen einfach umrechnen, indem man mit $2\pi / 360$ bzw. mit $360 / (2\pi)$ multipliziert. Für die Kosinus- und Sinusfunktion würde eine derartige Umskalierung aber bedeuten, dass auch in den Formeln für die Ableitung von $\sin$ und $\cos$ der entsprechende Faktor auftreten würde. $\diamond$

SATZ II.42. (I) Es gilt $\exp(i\pi / 2) = i$.

- (2) Der Kern des Gruppenhomomorphismus $\mathbb{R} \rightarrow \mathbb{C}^\times, t \mapsto \exp(it)$, ist die Untergruppe $2\pi\mathbb{Z} \subseteq \mathbb{R}$ aller Vielfachen von 2π .
- (3) Die Abbildung $[0, 2\pi) \rightarrow S^1, t \mapsto \exp(it)$, ist eine Bijektion vom Intervall $[0, 2\pi)$ auf den Einheitskreis S^1 .

BEWEIS. Wir zeigen zuerst, dass es eine reelle Zahl c mit $\exp(ic) = i$ gibt. Das ist der Schlüssel dazu, zu zeigen, dass die Abbildung $\mathbb{R} \rightarrow S^1, t \mapsto \exp(it)$, surjektiv ist. Die Existenz von c , die dazu äquivalent ist, dass die Kosinusfunktion eine Nullstelle hat, kann man rein analytisch zeigen, etwa durch geeignete Abschätzung der Reihendarstellung oder durch Argumente mit den Ableitungen von $\cos$ und $\sin$. Man könnte auf diesem Wege die Zahl π erst definieren, und nachträglich zeigen, dass sie die Länge des halben Einheitskreisbogens misst (oder für einen großen Teil der Theorie einfach darauf verzichten, diese Interpretation als Länge zu beweisen und die Zahl π sowie die trigonometrischen Funktionen einfach benutzen). Siehe zum Beispiel [Fo] §14, [Hi] Abschnitt 3.5 oder [Gr] 13.3.

Wir wollen hier, unserer Motivation entsprechend, ein geometrisches Argument skizzieren, das direkter mit der Bogenlänge arbeitet, auch wenn es etwas länger ist. (Es hat schon gute Gründe, dass meist der rein analytische Zugang bevorzugt wird.)

Angenommen, die Kosinusfunktion hätte keine Nullstelle. Dann wäre der Kern des Homomorphismus $t \mapsto \exp(it)$ trivial, diese Abbildung mithin injektiv. Denn für $C \neq 0$ mit $\exp(Ci) = 1$ folgt $\exp(Ci / 2) \in \{1, -1\}$. Ist $C > 0$ minimal mit $\exp(Ci) = 1$, so muss $\exp(Ci / 2) = -1$ gelten und es folgt $\exp(Ci / 4) \in \{i, -i\}$ und damit $\cos(C / 4) = 0$.

Weil $|\exp'(it)| = |i \exp(it)| = 1$ für alle t gilt, erhalten wir für den Weg $[0, x] \rightarrow \mathbb{C}, t \mapsto \exp(it)$ nach Satz II.32 die Länge

$$\int_0^x |\exp'(it)| dt = x,$$

die wir wegen der Injektivität auch als Bogenlänge interpretieren können. Unter der Annahme, dass $\cos(t)$ niemals Null ist, müsste diese Länge, also x , für alle x kleiner als ein Viertel des Umfangs des Einheitskreises sein, was absurd ist.

Wir haben damit gezeigt, dass $\cos$ eine Nullstelle in $\mathbb{R}_{>0}$ besitzt und definieren c als das Infimum aller Nullstellen von $\cos$ in $\mathbb{R}_{>0}$. Aus der Stetigkeit von $\cos$ und weil $\cos(0) = 1$ gilt, folgt $c > 0$.

Dass eine reelle Zahl x eine Nullstelle des Kosinus ist, ist dazu äquivalent, dass $|\sin(x)| = 1$ gilt. Wir wollen begründen, dass $\sin(c) = 1$ ist. Wegen der Minimalität von c ist die Kosinusfunktion auf $[0, c)$ positiv, und wegen $\sin' = \cos$ ergibt sich, dass die Sinusfunktion auf diesem Intervall streng monoton steigend ist, so dass nicht $\sin(c) = -1$ gelten kann.

Damit ist gezeigt, dass $\exp(ic) = i$ gilt. Um den Beweis von Teil (I) abzuschließen, müssen wir noch zeigen, dass $c = \pi / 2$ gelten muss. Die Minimalität von c impliziert, dass $\exp(it)$ auf $[0, c]$ injektiv ist. Denn sonst gäbe es in $[0, c]$ auch ein Element x mit $\exp(ix) = 1$, und für das minimale solche x in $(0, c)$ müsste $\exp(ix / 2) = -1$ und deshalb $\exp(ix / 4) \in \{i, -i\}$, also $\cos(ix / 4) = 0$, gelten – ein Widerspruch zur Definition von c .

Wir sehen mit dem Zwischenwertsatz, dass $t \mapsto \exp(it)$ eine Bijektion $[0, c] \rightarrow \{(a, b)^t \in S_1; 0 \leq a \leq 1, b \geq 0\}$ zwischen dem Intervall $[0, c]$ und dem »Viertel des Einheitskreises«, das im rechten oberen Quadranten liegt, definiert. Mit der oben durchgeführten Berechnung der Bogenlänge sehen wir, dass tatsächlich $c = \pi / 2$ gilt. Damit ist Teil (1) bewiesen.

Aus der Funktionalgleichung folgt nun direkt, dass $\exp(2\pi i) = 1$ ist, und damit ist klar, dass auch alle Vielfachen von 2π im Kern des Homomorphismus $t \mapsto \exp(it)$ liegen. Wir haben uns außerdem im ersten Teil des Beweises schon überlegt, dass es keine Zahlen x mit $0 < x < 2\pi$ und $\exp(ix) = 1$ geben kann. Daraus ergibt sich die Aussage aus Teil (2).

Für Teil (3) bemerken wir zunächst, dass aus dem Zwischenwertsatz und Teil (1) folgt, dass der Kosinus jedenfalls alle Zahlen in $[0, 1]$ als Werte annimmt. Daraus und aus der Funktionalgleichung lässt sich nun leicht schließen, dass die Abbildung $\mathbb{R} \rightarrow S^1, t \mapsto \exp(it)$, surjektiv ist. Dass die Einschränkung $[0, 2\pi) \rightarrow S^1$ dann eine Bijektion ist, folgt aus Teil (2). $\square$

BEMERKUNG II.43 (Polarkoordinaten). Teil (3) des Satzes, also die Tatsache, dass wir jede komplexe Zahl mit Betrag 1 in der Form $\exp(it)$, $t \in [0, 2\pi)$, schreiben können, erlaubt es uns, komplexe Zahlen durch *Polarkoordinaten* zu beschreiben. Damit ist gemeint, dass jede komplexe Zahl z in der Form $z = r \exp(it)$ für $r \in \mathbb{R}_{\geq 0}$ und $t \in [0, 2\pi)$ geschrieben werden kann. Wir setzen dafür einfach $r := |z|$. Ist $z \neq 0$, dann sind r und t eindeutig bestimmt.

Die Multiplikation komplexer Zahlen lässt sich in Polarkoordinaten sehr einfach beschreiben. Es ist $(r \exp(it)) \cdot (r' \exp(it')) = (rr') \exp(i(t + t'))$. Ist $t + t' > 2\pi$, so kann man das Ergebnis umschreiben als $(rr') \exp(i(t + t' - 2\pi))$, um zu den Polarkoordinaten des Produkts zu kommen. $\diamond$

BEMERKUNG II.44. Wenn wir die Ableitung im physikalischen Sinne als Geschwindigkeit deuten, heißt die Gleichheit $|\exp'(it)| = 1$ für alle t , dass der Weg $t \mapsto \exp(it)$ den Einheitskreis mit konstanter Geschwindigkeit 1 entlangläuft. Im Sinne dieser Anschauung ist klar, dass wir tatsächlich zum Zeitpunkt x genau Weglänge x zurückgelegt haben. $\diamond$

BEMERKUNG II.45. Alternativ kann man die Funktionen $\sin, \cos$ direkt durch ihre Darstellung als (reelle) Potenzreihen definieren,

$$\cos(x) = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!}, \quad \sin(x) = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!}.$$

Das vermeidet die Benutzung der komplexen Zahlen, ist aber etwas weniger »durchsichtig«, insbesondere, was die Herleitung der Additionstheoreme angeht. $\diamond$

BEMERKUNG II.46 (Alternativer Ansatz zur Definition von Sinus und Kosinus). Wir wollen kurz einen weiteren möglichen Ansatz skizzieren, um die trigonometrischen Funktionen einzuführen. (Es gibt natürlich auch noch andere Möglichkeiten, siehe zum Beispiel [Hi] 3.8 (insbes. S. 293) und 4.1.)

Für einen Punkt P des Einheitskreises bezeichnen wir mit $\ell(P)$ die Länge des Kreisbogens von $e_1 = (1, 0)^t$ nach P (wie üblich entgegen dem Uhrzeigersinn). Dass dieser Kreisbogen ein Jordanbogen ist, überlegt man sich ähnlich wie in Abschnitt II.3.2. Wir erhalten eine Funktion (den »Arkuskosinus«)

$$\arccos: [-1, 1] \rightarrow \mathbb{R}, \quad x \mapsto \ell((x, \sqrt{1-x^2})^t),$$

denn für alle $x \in [-1, 1]$ ist ja $(x, \sqrt{1-x^2})^t$ ein Punkt des Einheitskreises. Aus den Eigenschaften der Länge von Jordanbögen folgt, dass die Funktion $\arccos$ streng monoton wachsend (und insbesondere injektiv) mit Bild $[0, \pi]$ ist.

Wir würden gerne (vergleiche Satz II.3.2) für $c: [-1, 1] \rightarrow \mathbb{R}^2, x \mapsto (x, \sqrt{1-x^2})^t$ die Kreisbogenlänge als

$$\arccos(x) = \int_x^1 \|c'(t)\| dt$$

schreiben. Der genannte Satz gibt das aber nicht her, weil die Funktion c' , die auf dem offenen Intervall $(-1, 1)$ definiert ist, dort unbeschränkt ist und sich nicht stetig nach $[-1, 1]$ fortsetzen lässt. Dieselbe Problematik hat uns schon in Abschnitt II.3.2

Man muss (und kann) die Integraldarstellung der Funktion $\arccos$ daher als uneigentliches Integral auffassen. Wir wollen diese Schwierigkeit, um die Darstellung nicht noch weiter in die Länge zu ziehen, im folgenden stillschweigend ignorieren. Man kann dann zeigen, dass $\arccos$ auf dem offenen Intervall $(-1, 1)$ differenzierbar ist mit

$$\arccos'(x) = \|c'(x)\| = -\frac{1}{\sqrt{1-x^2}}.$$

Gäbe es nicht die (kleine) Schwierigkeit an den Integrationsgrenzen, so würde dies ja direkt aus dem Hauptsatz der Differential- und Integralrechnung folgen.

Weil $\arccos$ auf $[-1, 1]$ streng monoton wachsend ist, ist die Funktion dort injektiv, also eine Bijektion $[-1, 1] \rightarrow [0, \pi]$. Wir definieren die Abbildung $\cos: [0, \pi] \rightarrow [-1, 1]$ als ihre Umkehrfunktion. Weiterhin definieren wir $\sin(x) = \sqrt{1 - \cos(x)^2}$ für $x \in [0, \pi]$, so dass für alle diese x der Punkt $(\cos(x), \sin(x))^t$ auf dem Einheitskreis liegt. Aus dem Satz über die Ableitung der Umkehrfunktion einer differenzierbaren Funktion erhält man ohne Schwierigkeiten, dass $\cos' = -\sin$ gilt. Es folgt dann auch $\sin' = \cos$.

Man kann nun die Funktionen $\cos$ und $\sin$ von $[0, \pi]$ auf ganz $\mathbb{R}$ fortsetzen und mit ein bisschen zusätzlicher Arbeit zeigen, dass man überall differenzierbare Funktionen mit den gewünschten Eigenschaften erhält. Die Additionstheoreme erhält man aus der Additivität der Länge von Jordanbögen für zusammengesetzte Wege. $\diamond$

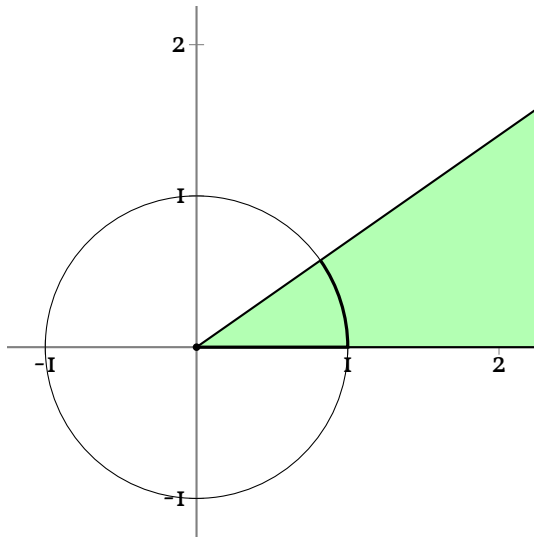
Letztlich erhalten wir mit beiden Ansätzen zur Definition von $\cos$ und $\sin$ die gewünschte Darstellung des Einheitskreises als Bild des Weges $t \mapsto (\cos(t), \sin(t))^t$. Wir können damit zum Beispiel auch den Umfang eines Kreises vom Radius $r \in \mathbb{R}_{>0}$ ausrechnen, denn dieser ist das Bild des Weges $\gamma: [0, 2\pi] \rightarrow \mathbb{R}^2, t \mapsto (r \cos(t), r \sin(t))^t$, wir erhalten als Länge also

$$\int_0^{2\pi} \|\gamma'(t)\| dt = \int_0^{2\pi} \sqrt{r^2 \cos'(t)^2 + r^2 \sin'(t)^2} dt = 2\pi r.$$

II.5. Der Begriff des Winkels

Wir beginnen damit, den Winkel zwischen zwei Ursprungsgeraden in der Ebene $\mathbb{R}^2$ zu definieren. Später definieren wir dann auch den Winkel zwischen Vektoren (oder Geraden) in $\mathbb{R}^n$ für beliebiges n .

Der Winkel (oder genauer: das *Winkelmaß*) soll die Größe des »Sektors« messen, der von zwei Geraden eingeschlossen wird. Man könnte nun der Größe des »Vollwinkels« (gewissermaßen der ganzen Ebene) willkürlich eine Zahl zuweisen (360 »Grad« ist die übliche Wahl) und dann andere Winkel entsprechend ihrer Größe mit Zahlen zwischen 0 und 360 Grad messen. Dann wäre zum Beispiel ein rechter Winkel, also ein Viertel des Vollwinkels, gerade 90 Grad groß.



Aus mathematischer Sicht ist es aber eleganter, diese willkürliche Wahl zu eliminieren und stattdessen die Größe des Winkels zwischen zwei Ursprungsgeraden durch die Länge des Kreisbogens zu messen, den die beiden Geraden aus dem Einheitskreis ausschneiden. Das bedeutet, dass wir dem Vollwinkel die Größe 2π , den Umfang des Einheitskreises, zuordnen. So hatten wir es auch in der heuristischen Diskussion am Anfang von Abschnitt II.3 angelegt. Siehe auch Bemerkung II.41. Der rechte Winkel entspricht dann $\frac{\pi}{2}$.

Wir hatten bereits den Begriff der Drehung, der mit dem Winkelbegriff eng zusammenhängt, definiert. Jeder Winkel α bestimmt eine Drehung (die *Drehung um den Winkel α*) und umgekehrt. Wir erinnern daran, dass $SO_2(\mathbb{R}) \subset GL_2(\mathbb{R})$ (die »spezielle orthogonale Gruppe«) die Untergruppe aller Drehungen bezeichnet.

SATZ II.47. *Die Abbildung*

$$\rho: \mathbb{R} \rightarrow SO_2(\mathbb{R}), \quad \alpha \mapsto \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix},$$

ist ein surjektiver Gruppenhomomorphismus mit Kern $\{2\pi k; k \in \mathbb{Z}\}$.

Die Isometrie $\rho_\alpha := f_{\rho(\alpha)}$ von $\mathbb{R}^2$ bezeichnen wir als die Drehung um den Winkel α .

BEWEIS. Dass ρ ein Gruppenhomomorphismus ist, bedeutet $\rho(\alpha + \beta) = \rho(\alpha)\rho(\beta)$, oder in Termen von Matrizen ausgeschrieben

$$\begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix} = \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix} \begin{pmatrix} \cos(\beta) & -\sin(\beta) \\ \sin(\beta) & \cos(\beta) \end{pmatrix}.$$

Dies entspricht genau den Additionstheoremen.

Die Surjektivität folgt daraus, dass sich jeder Punkt $(a, b)^t$ auf dem Einheitskreis in der Form $(\cos(\alpha), \sin(\alpha))^t$ und aus unserer Beschreibung der Drehungen von $\mathbb{R}^2$.

Die Aussage über den Kern folgt aus Satz II.42. □

Insbesondere gilt:

- (1) Wir erhalten eine Bijektion $[0, 2\pi) \rightarrow SO_2(\mathbb{R})$,
- (2) für alle $k \in \mathbb{Z}$ und alle $\alpha \in \mathbb{R}$ gilt $\rho_\alpha = \rho_{\alpha+2k\pi}$ und
- (3) für $\alpha, \beta \in \mathbb{R}$ gilt

$$\rho_\alpha \circ \rho_\beta = \rho_{\alpha+\beta}.$$

Das führt uns auf die folgende Definition des Winkels.

DEFINITION II.48. Wir sagen, der *Winkel* zwischen Punkten v und w auf dem Einheitskreis (bzw. zwischen den von ihnen aufgespannten »Strahlen« $\mathbb{R}_{\geq 0}v$ und $\mathbb{R}_{\geq 0}w$) sei die eindeutig bestimmte Zahl $\alpha \in [0, 2\pi)$, so dass $w = \rho_\alpha(v)$. ⊥

Äquivalent können wir die Zahl α aus der Definition als die Länge des Kreisbogens von v nach w (gegen den Uhrzeigersinn) definieren.

Wir erhalten damit für die Sinus- und Kosinusfunktion die »naive« Interpretation als das Verhältnis von Gegenkathete und Hypotenuse bzw. von Ankathete und Hypotenuse im rechtwinkligen Dreieck.

SATZ II.49. Seien $o, v, w \in \mathbb{R}^2$ Eckpunkte eines rechtwinkligen Dreiecks mit dem rechten Winkel an der Ecke v (d.h. $-v \cdot (w - v) = 0$). Sei α der Winkel zwischen v und w bzw. zwischen w und v (wir wählen unter diesen beiden den Winkel, der kleiner als der rechte Winkel $\pi / 2$ ist).

Dann gilt

$$\cos(\alpha) = \frac{\|v\|}{\|w\|}, \quad \sin(\alpha) = \frac{\|w - v\|}{\|w\|}.$$

BEWEIS. Wenn v auf der waagerechten Koordinatenachse liegt und w ein Punkt auf dem Einheitskreis im rechten oberen Quadranten ist, dann gilt nach unsere Definition von $\cos$ und $\sin$, dass $w = (\cos(\alpha), \sin(\alpha))^t$ und aus $v \cdot (w - v) = 0$ folgt $v = (\cos(\alpha), 0)^t$. Es ist dann $\|w\| = 1$, $\|v\| = \cos(\alpha)$ und $\|w - v\| = \sin(\alpha)$, wie gewünscht.

Den allgemeinen Fall führen wir durch eine geeignete Verschiebung, eine Drehung und eine Skalierung auf diese Situation zurück. $\square$

Für Vektoren $v, w \in \mathbb{R}^2$ könnten wir den Winkel $\angle(v, w)$ als den Winkel zwischen den Halbgeraden $\mathbb{R}_{\geq 0}v$ und $\mathbb{R}_{\geq 0}w$ definieren. Wir wollen aber eine Definition geben, die dann auch für höher-dimensionale Räume $\mathbb{R}^n$ vernünftig ist. Dort ist es dann wieder nicht möglich, der Richtung »entgegen des Uhrzeigersinns« eine Bedeutung zu geben. Wir werden daher von vorneherein das Ziel verfolgen, »den kleineren« der beiden Winkel zwischen v und w zu definieren.

Die Interpretation von $\cos$ und $\sin$ aus Satz II.49 erlaubt es uns, den Winkel zwischen Vektoren auf den Abstandsbegriff (der uns schon in allen $\mathbb{R}^n$ zur Verfügung steht) zurückzuführen. Mit der Beschreibung des Skalarprodukts aus Lemma II.17 kommen wir auf die folgende Definition.

DEFINITION II.50. Seien $v, w \in \mathbb{R}^n \setminus \{0\}$. Der *Innenwinkel* $\angle(v, w)$ zwischen v und w (oft spricht man auch einfach vom *Winkel* zwischen diesen Vektoren) ist die eindeutig bestimmte Zahl $\alpha \in [0, \pi)$, so dass

$$\cos(\alpha) = \frac{vw}{\|v\| \|w\|}.$$

—

Man beachte, dass – wie schon gesagt – der Winkel zwischen zwei Vektoren im Sinne dieser Definition immer zwischen 0 und dem gestreckten Winkel π liegt – um das zu betonen, spricht man vom Innenwinkel. Werte zwischen π und 2π (oder negative Werte für den Winkel) treten in der obigen Definition nicht auf. Mit anderen Worten: Der Innenwinkel ist der kleinere der beiden Winkel, die man erhält, wenn man einerseits von v nach w »geht« oder andererseits von w nach v »geht«. Es gilt $\angle(v, w) = \angle(w, v)$.

Diese Definition des Winkels hat die erwarteten Eigenschaften, beispielsweise sind die gegenüberliegenden Winkel, die durch zwei sich schneidende Geraden gebildet werden, gleich groß, und zwei nebeneinander liegende addieren sich zu π , dem halben Vollwinkel. Der Winkelbegriff ist invariant unter Translationen. Die Innenwinkelsumme im Dreieck ist π .

11.5.1. Der Kosinussatz. Als eine kleine »Anwendung« zeigen wir, dass der Kosinussatz sich aus unseren Definitionen nun leicht ergibt.

SATZ 11.51 (Kosinussatz). Seien $a, b, c > 0$ die Seitenlängen eines Dreiecks, und sei γ der der Seite c gegenüberliegende Winkel. Dann gilt

$$\cos(\gamma) = \frac{a^2 + b^2 - c^2}{2ab}.$$

BEWEIS. Wir wählen Punkte $v, w \in \mathbb{R}^2$, so dass das Dreieck mit den Eckpunkten $0, v, w$ die vorgegebenen Seitenlängen hat, so dass $a = \|v\|, b = \|w\|, c = d(v, w) = \|w - v\|$. Dann ist γ der Winkel zwischen den Vektoren v und w .

Es folgt

$$a^2 + b^2 - c^2 = (v, v) + (w, w) - (w - v, w - v) = 2(v, w).$$

Andererseits haben wir

$$\cos(\gamma) = \frac{(v, w)}{\|v\| \|w\|} = \frac{2(v, w)}{2ab}.$$

Setzen wir hier den obigen Ausdruck für $2(v, w)$ ein, so erhalten wir die Behauptung des Satzes. $\square$

BEMERKUNG 11.52. In Frankreich wird der Kosinussatz als *Théorème d'Al-Kashi* bezeichnet, das bezieht sich auf den persischen Mathematiker (und Arzt) [Dschamschid Mas'ud al-Kaschi](https://de.wikipedia.org/wiki/Dschamschid_Mas%CA%BFud_al-Kaschi)³ (ca. 1380 – 1429), der den Kosinussatz in einem seiner Bücher angegeben hat. Er führte auch weitere trigonometrische Berechnungen durch und berechnete die Zahl 2π auf 16 Dezimalstellen genau. $\diamond$

11.6. Volumina und die Determinante

11.6.1. Das Volumen von Teilmengen von $\mathbb{R}^n$. Es entspricht zwar der Anschauung, dass man Teilmengen von $\mathbb{R}^2$ einen Flächeninhalt in $\mathbb{R}_{\geq 0} \cup \{\infty\}$, also eine nicht-negative Zahl oder »unendlich«, zuordnen kann. Allgemeiner würden wir gerne jeder Teilmenge von $\mathbb{R}^n$ ein Volumen zuordnen. *Dies ist aber nicht in vernünftiger Weise möglich!* Siehe unten für eine genauere Aussage (und Literaturverweise). Für den Moment halten wir fest, dass wir nicht jeder Teilmenge von $\mathbb{R}^n$ ein Volumen zuordnen werden, sondern nur »geeigneten« Teilmengen, darunter alle solche, über die wir in der Praxis sprechen möchten.

Wir beginnen damit, einige offensichtliche Anforderungen zu formulieren. Wir fixieren $n \in \mathbb{N}_{>0}$ und bezeichnen mit $P(\mathbb{R}^n)$ die Potenzmenge von $\mathbb{R}^n$, also die Menge aller Teilmengen von $\mathbb{R}^n$. Wir suchen dann nach einer Teilmenge $\mathcal{M}_n \subseteq P(\mathbb{R}^n)$ (die aus allen Teilmengen von $\mathbb{R}^n$ besteht, denen wir überhaupt ein Volumen zuordnen werden) und einer *Volumenfunktion* (oder einem *Maß*)

$$\mu_n: \mathcal{M}_n \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}.$$

Der Wert $\mu_n(X)$ für eine Teilmenge $X \subseteq \mathbb{R}^n$, die in $\mathcal{M}_n$ liegt, ist also entweder eine nicht-negative reelle Zahl, oder unendlich (∞). Beispielsweise sollte sicherlich $\mu_n(\emptyset) = 0$ und $\mu_n(\mathbb{R}^n) = \infty$ gelten.

Etwas systematischer wollen wir jedenfalls die folgenden Eigenschaften verlangen:

Normierung. Es gilt $[0, 1]^n \in \mathcal{M}_n$ und $\mu_n([0, 1]^n) = 1$.

Translationsinvarianz. Für $X \in \mathcal{M}_n$ und $v \in \mathbb{R}^n$ gilt $v + X \in \mathcal{M}_n$ und $\mu_n(v + X) = \mu_n(X)$.

Additivität für endliche disjunkte Vereinigungen. Sind $X, Y \in \mathcal{M}_n$, und ist $X \cap Y = \emptyset$, so gilt

$$\mu_n(X \cup Y) = \mu_n(X) + \mu_n(Y).$$

³https://de.wikipedia.org/wiki/Dschamschid_Mas%CA%BFud_al-Kaschi

Hier ist $[0, 1] = \{x \in \mathbb{R}; 0 \leq x \leq 1\}$ das abgeschlossene Intervall mit den Endpunkten 0, 1. Die Teilmenge $[0, 1]^n \subseteq \mathbb{R}^n$ nennt man den *Einheitswürfel*. Wir schreiben $v + X = \{v + x; x \in X\}$, mit anderen Worten: $v + X = t_v(X)$ ist das Bild von X unter der Translation $t_v: x \mapsto v + x$. Ist in der dritten Bedingung eines (oder beide) der Volumen $\mu_n(X), \mu_n(Y)$ gleich ∞ , so soll auch die Summe den Wert ∞ haben.

BEMERKUNG II.53. Es gibt für $n \geq 3$ keine Funktion $\mu_n: P(\mathbb{R}^n) \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}$ mit den obigen drei Eigenschaften. Das hat Hausdorff 1914 bewiesen. Banach zeigte 1923, dass es für $n = 1$ und $n = 2$ zwar Abbildungen $\mu_1: \mathbb{R} \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}, \mu_2: \mathbb{R}^2 \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}$ mit diesen Eigenschaften gibt, dass diese aber nicht eindeutig bestimmt sind. Siehe [E1] I.1. $\diamond$

Die Bemerkung zeigt, dass wir nur für eine geeignete Teilmenge $\mathcal{M}_n \subset P(\mathbb{R}^n)$ eine Volumenfunktion μ_n werden definieren können. Bevor wir darüber sprechen, welchen Teilmengen wir ein Volumen zuordnen können, ziehen wir einige Folgerungen aus den oben geforderten Eigenschaften.

LEMMA II.54. Sei $\mathcal{M}_n \subset P(\mathbb{R}^n)$ und $\mu_n: \mathcal{M}_n \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}$ eine Funktion, die die oben aufgeführten Eigenschaften (Normierung, Translationsinvarianz, Additivität für endliche disjunkte Vereinigungen) hat. Dann gilt

- (1) (Monotonie) Sind $X, Y \in \mathcal{M}$, und ist $X \subseteq Y$, so gilt $\mu_n(X) \leq \mu_n(Y)$.
 (2) (Volumen von Quadern) Für $a_i, b_i \in \mathbb{R}, a_i \leq b_i, i = 1, \dots, n$, gilt $[a_1, b_1] \times \dots \times [a_n, b_n] \in \mathcal{M}_n$ und

$$\mu_n([a_1, b_1] \times \dots \times [a_n, b_n]) = \prod_{i=1}^n (b_i - a_i)$$

In (1) setzen wir $a \leq \infty$ für alle $a \in \mathbb{R}_{\geq 0}$.

Wir schreiben hier wie üblich $[a, b] = \{x \in \mathbb{R}; a \leq x \leq b\}$ für das abgeschlossene Intervall mit den Endpunkten $a \leq b, a, b \in \mathbb{R}$. Das Produkt, das wir hier betrachten, ist also ein » n -dimensionaler Quader«.

BEWEIS. Teil (1) folgt, indem wir Y als die disjunkte Vereinigung $Y \cup (X \setminus Y)$ schreiben, weil $\mu_n(X \setminus Y) \geq 0$ ist. Für Teil (2) bemerken wir zunächst, dass es für das Volumen eines Quaders keinen Unterschied macht, ob wir abgeschlossene oder (halb-)offene Intervalle betrachten, denn wir können beliebig viele Mengen der Form $\{a\} \times [a_2, b_2] \times \dots \times [a_n, b_n]$ in $[a_1, b_1] \times \dots \times [a_n, b_n]$ unterbringen. Das ist wegen der Monotonie nur möglich, wenn $\mu_n(\{a\} \times [a_2, b_2] \times \dots \times [a_n, b_n]) = 0$ gilt.

Durch Zusammensetzen von Translaten $x + [0, 1]^n$ des halboffenen Einheitswürfels, verschoben um einen Vektor x mit Koordinaten in $\mathbb{Z}$, kann man alle Quader, die als Produkte von halboffenen Intervallen mit ganzzahligen Koordinaten gebildet werden, erhalten, und sieht, dass diese das im Lemma gegebene Maß haben. Andersherum kann man Translate von Quadern, die Produkte von Intervallen mit rationaler Länge sind, zusammensetzen zu einem Quader mit ganzzahligen Seitenlängen. Damit erhält man das Ergebnis unter der Voraussetzung, dass $b_i - a_i \in \mathbb{Q}$ für alle i . Den allgemeinen Fall erhält man dann schließlich, indem man ausnutzt, dass sich reelle Zahlen beliebig genau durch rationale Zahlen annähern lassen, und wieder die Monotonie verwendet. $\square$

An dieser Stelle können Sie zu Abschnitt II.6.4 springen, wenn Sie bereit sind zu glauben, dass $\mu_n: \mathcal{M}_n \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}$ existiert mit einer Menge $\mathcal{M}_n$ von Teilmengen von $\mathbb{R}^n$, die alle Teilmengen enthält, die »uns interessieren«, und dass für diese Teilmengen der Wert μ_n eindeutig durch die obigen Eigenschaften bestimmt ist.

Etwas mehr Details zu zwei möglichen Ansätzen geben wir in den folgenden beiden Abschnitten:

- Für *Jordan-messbare* Teilmengen X von $\mathbb{R}^n$ ist der Wert μ_n schon durch die obigen Eigenschaften eindeutig bestimmt, und man kann auf der Teilmenge $\mathcal{M}_n^J \subset P(\mathbb{R}^n)$ aller Jordan-messbaren Mengen die Funktion $\mu_n: \mathcal{M}_n^J \rightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}$ relativ leicht definieren und ihre Eigenschaften beweisen. Der Nachteil dieser Definition ist, dass viele Mengen, über die man »in der Praxis« (Analysis, Wahrscheinlichkeitstheorie, ...) sprechen möchte, nicht Jordan-messbar sind, siehe Bemerkung II.6.1. Für unsere Zwecke reicht diese Menge von »messbaren Mengen« aber aus.
- Wegen der genannten Nachteile der Jordan-Messbarkeit ist der weitaus häufiger genutzte Ansatz die Definition von messbaren Mengen, die auf Lebesgue (und Borel) zurückgeht. Siehe Abschnitt II.6.3.

II.6.2. Jordan-messbare Mengen. Wir wollen nun schrittweise eine Zahl $\mu_n(M)$ für viele Teilmengen von $\mathbb{R}^n$ definieren. Wir beginnen mit dem Produkt von n Intervallen; diesen Fall hatten wir schon in Lemma II.5.4 betrachtet.

DEFINITION II.55. (I) Ein Quader (in $\mathbb{R}^n$) ist ein Produkt von n Intervallen mit Endpunkten $a_i \leq b_i, i = 1, \dots, n$, zum Beispiel

$$[a_1, b_1] \times \cdots \times [a_n, b_n],$$

aber wir lassen alternativ auch offene und halboffene Intervalle zu.

(2) Das *Volumen* (oder: *Maß*) eines Quaders wie in (I) ist definiert als

$$\mu_n([a_1, b_1] \times \cdots \times [a_n, b_n]) = \prod_{i=1}^n (b_i - a_i).$$

(Dieselbe Definition verwenden wir, wenn der Quader mit offenen/halboffenen Intervallen gebildet wurde.)

—

Statt von Quadern spricht man präziser auch von *n-dimensionalen Quadern* oder manchmal auch von *Hyperquadern* oder *Hyperrechtecken*. Als nächsten Schritt betrachten wir Vereinigungen von Quadern.

DEFINITION II.56. Eine *elementare Teilmenge* von $\mathbb{R}^n$ ist eine Teilmenge, die sich als endliche Vereinigung von Quadern schreiben lässt.

—

Es ist dann nicht schwer zu beweisen, dass man auch jeder elementaren Teilmenge in sinnvoller Weise ein Volumen zuordnen kann.

LEMMA II.57. Jede elementare Teilmenge $M \subset \mathbb{R}^n$ lässt sich als disjunkte Vereinigung $Q_1 \sqcup \cdots \sqcup Q_r$ von Quadern Q_i schreiben, und die Zahl $\mu_n(M) := \sum_{i=1}^r \mu_n(Q_i)$ ist unabhängig von der Wahl dieser Zerlegung. Wir nennen $\mu_n(M)$ das Volumen oder das Maß von M .

BEWEIS. Man kann das auf verschiedene Weisen beweisen. Für einen geschickten Beweis siehe [Ta], Lemma I.1.2. □

Nur elementare Mengen zu betrachten, reicht uns noch nicht aus. Wir dehnen die Theorie weiter aus, indem wir Teilmengen von $\mathbb{R}^n$ hinzufügen, die sich »mit demselben Ergebnis von innen und von außen durch elementare Teilmengen approximieren lassen«. Die Theorie des Jordan-Maßes ist nicht gut geeignet, um Mengen mit »unendlicher Ausdehnung« zu betrachten, wir beschränken uns daher auf beschränkte Mengen im Sinne der folgenden Definition.

DEFINITION II.58. Eine Teilmenge $X \subseteq \mathbb{R}^n$ heißt beschränkt, wenn $N \geq 0$ existiert mit $X \subseteq \{x \in \mathbb{R}^n; \|x\| \leq N\}$, oder äquivalent: wenn es einen Quader in $\mathbb{R}^n$ gibt, in dem X enthalten ist. $\dashv$

DEFINITION II.59. Sei $X \subseteq \mathbb{R}^n$ eine beschränkte Teilmenge.

- (1) Das *innere Jordan-Maß* $\mu_{n,*}(X)$ von X ist definiert als

$$\mu_{n,*}(X) = \sup_{E \subseteq X} \mu_n(E),$$

wobei das Supremum über alle elementaren Teilmengen $E \subset \mathbb{R}^n$ gebildet wird, die in X enthalten sind.

- (2) Das *äußere Jordan-Maß* $\mu_n^*(X)$ von X ist definiert als

$$\mu_n^*(X) = \inf_{X \subseteq E} \mu_n(E),$$

wobei das Infimum über alle elementaren Teilmengen $E \subset \mathbb{R}^n$ gebildet wird, die X enthalten.

- (3) Die Teilmenge X heißt *Jordan-messbar*, wenn $\mu_{n,*}(X) = \mu_n^*(X)$ gilt. In diesem Fall nennen wir $\mu_n(X) := \mu_{n,*}(X) = \mu_n^*(X)$ das (*Jordan-*)*Maß* oder (*Jordan-*)*Volumen* von X . $\dashv$

Jede elementare Teilmenge E ist Jordan-messbar, und ihr Jordan-Maß im Sinne der obigen Definition stimmt mit der vorherigen Definition von $\mu_n(E)$ überein. Für die leere Menge verstehen wir die Definition so, dass das Supremum über die leere Indexmenge gleich 0 ist, so dass $\emptyset$ messbar mit Maß $\mu_n(\emptyset) = 0$ ist.

Man könnte für unbeschränkte Mengen X als Jordan-Maß ∞ festsetzen, aber das ist nicht befriedigend, weil auch unbeschränkten Mengen unter Umständen ein endliches Maß zugeordnet werden sollte, ähnlich wie es konvergente uneigentliche Integrale gibt.

Für die Experten die Bemerkung dass das, was wir hier (und man üblicherweise) als Jordan-Maß bezeichnen, kein »Maß« (sondern nur ein »Inhalt«) im Sinne der Maßtheorie ist.

Wir fassen einige wichtige Eigenschaften des so definierten Begriffs zusammen; darunter sind insbesondere die Eigenschaften, mit denen wir in Abschnitt II.6.1 die Diskussion begonnen hatten.

SATZ II.60 (Eigenschaften des Jordan-Maßes). *Das Jordan-Maß ist eine Abbildung μ_n von der Menge $\mathcal{M}_n^J$ aller Jordan-messbaren Teilmengen nach $\mathbb{R}_{\geq 0}$ mit den folgenden Eigenschaften.*

- (1) (Normierung) *Der Einheitsquader $[0, 1]^n$ ist Jordan-messbar und hat Maß 1. Die leere Menge ist Jordan-messbar und hat Maß 0.*
- (2) (Translationsinvarianz) *Ist X Jordan-messbar und $v \in \mathbb{R}^n$, dann ist $v + X := \{v + x; x \in X\}$ Jordan-messbar und $\mu_n(v + X) = \mu_n(X)$.*
- (3) *Sind X und Y Jordan-messbar, so auch $X \cap Y$, $X \cup Y$ und $X \setminus Y$.*
- (4) (Endliche Vereinigungen) *Sind $X, Y \subseteq \mathbb{R}^n$ Jordan-messbare Teilmengen, so sind $X \cup Y$ und $X \cap Y$ messbar und es gilt*

$$\mu_n(X \cup Y) = \mu_n(X) + \mu_n(Y) - \mu_n(X \cap Y).$$

- (5) (Monotonie) *Sind $X, Y \subseteq \mathbb{R}^n$ messbar, und ist $X \subseteq Y$, so gilt $\mu_n(X) \leq \mu_n(Y)$.*
- (6) *Jede beschränkte konvexe Teilmenge von $\mathbb{R}^n$ ist Jordan-messbar.*

BEWEIS. Wir geben nur einige skizzenhafte Hinweise. Teile (1), (2) und (5) ergeben sich leicht aus der Definition.

zu (3). Das kann man daraus folgern, dass für elementare Teilmengen $X, Y \subseteq \mathbb{R}^n$ auch die Mengen $X \cup Y$, $X \cap Y$ und $X \setminus Y$ elementare Teilmengen sind.

zu (4). Wir zerlegen $X \cup Y$ als die disjunkte Vereinigung $(X \setminus Y) \sqcup (Y \setminus X) \sqcup (X \cap Y)$. Das reduziert die Aufgabe darauf, die Aussage für disjunkte Vereinigungen zu zeigen, und das ist nicht schwierig.

zu (6). Eine Teilmenge $X \subseteq \mathbb{R}^n$ heißt *konvex*, wenn für alle $x, x' \in X$ auch die gesamte Strecke zwischen x und x' , also alle Punkte der Form $ax + (1-a)x'$ für $a \in [0, 1]$, in X liegen. Diese Eigenschaft des Jordan-Maßes ist nicht so leicht zu zeigen. Siehe [El] II.7.4. $\square$

Es ist leicht, Teilmengen von $\mathbb{R}^n$ anzugeben, die nicht Jordan-messbar sind, zum Beispiel $[0, 1]^n \cap \mathbb{Q}^n$ und $[0, 1]^n \setminus \mathbb{Q}^n$ für $n \geq 1$. Für die Belange der Vorlesung Lineare Algebra kommen wir aber mit den Jordan-messbaren Teilmengen aus. Es sei aber bemerkt, dass selbst kompakte Teilmengen von $\mathbb{R}$ nicht notwendig Jordan-messbar sind; ebenso gibt es beschränkte offene Teilmengen, die nicht Jordan-messbar sind. Die folgende Bemerkung enthält einen Verweis auf ein konkretes Beispiel.

BEMERKUNG 11.61. Die sogenannten **fetten Cantor-Mengen**⁴ sind kompakte Teilmengen von $[0, 1]$, die kein Intervall $[a, b]$ enthalten, deren Abschluss aber das gesamte Einheitsintervall $[0, 1]$ ist. Sie haben folglich inneres Jordan-Maß 0 und äußeres Jordan-Maß 1 und sind damit ein Beispiel für nicht-Jordan-messbare kompakte Teilmengen von $\mathbb{R}$. Das Komplement $[0, 1] \setminus C$ einer solchen Teilmenge $C \subset [0, 1]$ ist ein Beispiel für eine beschränkte offene Teilmenge von $\mathbb{R}$, die nicht Jordan-messbar ist. $\diamond$

Die Definition des Riemann-Integrals (oder präziser: des Darboux-Integrals) mit Ober- und Untersumme passt gut zum Begriff der Jordan-Messbarkeit, und man kann zeigen:

SATZ 11.62 (Jordan-Maß und Riemann-Integral). *Seien $a \leq b$ reelle Zahlen und $f: [a, b] \rightarrow \mathbb{R}$ eine Abbildung. Dann sind äquivalent:*

- (i) *Die Funktion f ist auf dem Intervall $[a, b]$ Riemann-integrierbar.*
- (ii) *Die Mengen $E_+ := \{(x, y); x \in [a, b], 0 \leq y \leq f(x)\}$ und $E_- := \{(x, y); x \in [a, b], 0 \geq y \geq f(x)\}$ sind Jordan-messbar.*

In diesem Fall gilt

$$\int_a^b f(x) dx = \mu_2(E_+) - \mu_2(E_-).$$

Siehe [Ta] I.1.

11.6.3. Lebesgue-messbare Mengen. Wir wollen den Begriff des Lebesgue-Maßes und der Lebesgue-messbaren Teilmenge hier nicht definieren, sondern nur darauf hinweisen, dass man eine Theorie entwickeln kann, so dass für die Menge $\mathcal{M}_n^L$ aller messbaren Teilmengen gilt, dass $\mathcal{M}_n^J \subset \mathcal{M}_n^L \subset P(\mathbb{R}^n)$ – jede Jordan-messbare Teilmenge ist Lebesgue-messbar, und so dass alle offenen Teilmengen von $\mathbb{R}^n$ sowie alle abgeschlossenen Teilmengen von $\mathbb{R}^n$ Lebesgue-messbar sind. Außerdem sind endliche Durchschnitte und abzählbare Vereinigungen von Lebesgue-messbaren Teilmengen wieder Lebesgue-messbar. Vereinigungen mit abzählbar unendlicher Indexmenge in die Theorie miteinzubeziehen, ist ein wichtiger Baustein der »Maßtheorie«.

⁴ <https://de.wikipedia.org/wiki/Smith-Volterra-Cantor-Menge>

Das Lebesgue-Maß μ_n^L ist dann eine Abbildung

$$\mu_n^L: \mathcal{M}_n^L \longrightarrow \mathbb{R}_{\geq 0} \cup \{\infty\},$$

die jeder Lebesgue-messbaren Teilmenge als Maß eine nicht-negative Zahl in $\mathbb{R}$ oder das Maß ∞ – »unendlich« – zuordnet. Mengen mit Maß unendlich können in dieser Theorie in natürlicher Weise berücksichtigt werden. Für Jordan-messbare Teilmengen stimmen das Jordan-Maß und das Lebesgue-Maß überein.

Es gibt »viel mehr« Lebesgue-messbare Mengen als Jordan-messbare Mengen, und das Lebesgue-Maß ist ausreichend für (fast) alle mathematischen Zwecke, aber die Theorie ist dementsprechend komplizierter »aufzusetzen«.

Eindeutigkeit (erst) dann, wenn man Additivität für *abzählbare* disjunkte Vereinigungen verlangt, also für disjunkte Vereinigungen mit abzählbarer Indexmenge (Abschnitt 3.15).

Mit dem Begriff des Lebesgue-Maßes lässt sich auch der Integralbegriff verallgemeinern. Das »Lebesgue-Integrals« liefert für Riemann-integrierbare Funktionen dasselbe Ergebnis, wie das »gewöhnliche« Integral aus der Analysis I, es gibt aber Funktionen, die Lebesgue-integrierbar, jedoch nicht Riemann-integrierbar sind. Ein Beispiel ist die Funktion $[0, 1] \rightarrow \mathbb{R}$, die rationale Zahlen auf 1 und irrationale Zahlen auf 0 abbildet. (Ihr Lebesgue-Integral über das Intervall $[0, 1]$ ist gleich 0.)

Die Theorie des Lebesgue-Maßes und -Integrals werden in den bereits genannten Büchern [Ta] von Tao (insbesondere Kapitel I.2), und [El] von Elstrodt (Kapitel I, II) entwickelt. Eine alternative Quelle ist das Buch von H. Bauer, *Maß- und Integrations-theorie*, DeGruyter 2002.

Die Maßtheorie ist sowohl in der Analysis, als auch ganz besonders in der Wahrscheinlichkeitstheorie von großer Bedeutung.

II.6.4. Der Flächeninhalt der Einheitskreisscheibe. Wir wollen die Theorie des Jordan-Maßes, die wir entwickelt haben, benutzen, um den Flächeninhalt der Einheitskreisscheibe $\{x \in \mathbb{R}^2; \|x\| \leq 1\}$ zu berechnen. Eine Möglichkeit dazu, die schon auf [Archimedes](#)⁵ und damit mehr als 2000 Jahre zurückgeht und die genau dem Prinzip des Jordan-Maßes entspricht, ist, die Kreisscheibe durch Polygone, genauer durch regelmäßige n -Ecke zu approximieren.

Dabei verstehen wir unter einem regelmäßigen n -Eck eine Teilmenge von $\mathbb{R}^2$, die aus n kongruenten gleichseitigen Dreiecken zusammengesetzt ist, deren kleinster Winkel jeweils $2\pi / n$ ist, und die alle an der Ecke mit diesem kleinsten Winkel zusammentreffen und sich nur mit den Seitenlinien überschneiden. Der Kreis, auf dem alle Eckpunkte der Dreiecke außer des Mittelpunkts des n -Ecks liegen, heißt der Umkreis des n -Ecks. Sein Radius ist die Länge der beiden gleichen Seiten der gleichschenkligen Dreiecke, an denen je zwei Dreiecke aneinander angrenzen.

Es ist nicht schwer zu sehen, dass jedes Dreieck Jordan-messbar ist und dass die übliche Formel für den Flächeninhalt gilt: Er berechnet sich als die Hälfte des Produkts der Länge der Grundseite mit der Höhe des Dreiecks. Man kann das leicht auf die Flächeninhalte von Rechtecken zurückführen.

Wir betrachten nun ein regelmäßiges n -Eck, dessen Umkreis Radius r hat. Wir können jedes der n gleichschenkligen Dreiecke in 2 rechtwinklige Dreiecke unterteilen, deren Hypotenuse die Länge r und deren Katheten die Längen $r \sin(\pi / n)$ und $r \cos(\pi / n)$ haben. Vergleiche Satz II.49.

⁵<https://de.wikipedia.org/wiki/Archimedes>

Der Flächeninhalt jedes dieser $2n$ rechtwinkligen Dreiecke ist mithin

$$\frac{1}{2} r^2 \sin\left(\frac{\pi}{n}\right) \cos\left(\frac{\pi}{n}\right)$$

und für den Flächeninhalt des n -Ecks erhalten wir das $2n$ -fache, also

$$2n \cdot \frac{1}{2} r^2 \sin\left(\frac{\pi}{n}\right) \cos\left(\frac{\pi}{n}\right) = \pi r^2 \cdot \frac{n}{\pi} \sin\left(\frac{\pi}{n}\right) \cos\left(\frac{\pi}{n}\right).$$

Dieser Ausdruck konvergiert für $n \rightarrow \infty$ wegen $\sin'(0) = 1$ und der Stetigkeit von $\cos$ gegen πr^2 . Als inneres Jordan-Maß der Kreisscheibe vom Radius r erhalten wir damit $\geq \pi r^2$.

Um das äußere Jordan-Maß zu berechnen, betrachten wir ein regelmäßiges n -Eck, das eine Kreisscheibe vom Radius r enthält. Ein solches können wir zusammengesetzen aus $2n$ rechtwinkligen Dreiecken mit einem Winkel π/n , der an diesem Winkel anliegenden Kathete von Länge r , folglich Hypotenuse von Länge $r / \cos(\pi/n)$ und zweiter Kathete der Länge $r \sin(\pi/n) / \cos(\pi/n)$. Damit erhalten wir für den Flächeninhalt des n -Ecks

$$2n \cdot \frac{1}{2} r^2 \sin\left(\frac{\pi}{n}\right) / \cos\left(\frac{\pi}{n}\right) = \pi r^2 \cdot \cos\left(\frac{\pi}{n}\right)^{-1} \cdot \frac{n}{\pi} \sin\left(\frac{\pi}{n}\right),$$

der für $n \rightarrow \infty$ ebenfalls gegen πr^2 konvergiert. Damit ist das äußere Jordan-Maß gleich dem inneren Jordan-Maß, die Einheitskreisscheibe also meßbar mit Jordan-Maß πr^2 .

Alternativ kann man den Zusammenhang zwischen Jordan-Maß und (Riemann-)Integral ausnutzen. Dass der Flächeninhalt der Einheitskreisscheibe gleich π ist, entspricht dem folgenden Ergebnis.

SATZ 11.63. *Es gilt*

$$\int_{-1}^1 \sqrt{1-x^2} dx = \frac{\pi}{2}.$$

BEWEIS. Nach der Substitutionsregel, angewandt auf $x = \cos(t)$, können wir das Integral umschreiben als

$$\int_{-1}^1 \sqrt{1-x^2} dx = \int_{\pi}^0 \sqrt{1-\cos(t)^2} \cos'(t) dt = \int_0^{\pi} \sin(t)^2 dt.$$

(Die Vorzeichenänderung wegen der Vertauschung der Integrationsgrenzen im zweiten Schritt hebt sich damit auf, dass wir mit $\sin$ das Negative der Ableitung $\cos'$ stehen lassen.)

Um das so erhaltene Integral auszurechnen, gibt es verschiedene Möglichkeiten. Wir verwenden die Gleichheit $\sin(t)^2 = \frac{1}{2}(1 - \cos(2t))$, die leicht aus dem Additionstheorem für $\cos(2x) = \cos(x+x)$ folgt. Damit ergibt sich

$$\dots = \int_0^{\pi} \left(\frac{1}{2}(1 - \cos(2t)) \right) dt = \frac{1}{2} \left[t - \frac{1}{2} \sin(2t) \right]_0^{\pi} = \frac{\pi}{2},$$

wobei wir im letzten Schritt noch benutzen, dass $\sin(0) = \sin(2\pi) = 0$. □

11.6.5. Volumenänderung unter linearen Abbildungen. Wir können nach all diesen Vorbereitungen nun die lineare Algebra ins Spiel bringen und wollen untersuchen, was wir über das Volumen $f(M)$ des Bildes einer (messbaren) Menge M unter einer *linearen* Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ sagen können.

Wir beginnen mit linearen Abbildungen von besonders einfacher Form.

LEMMA 11.64. (I) *Sei $A = \text{diag}(a_1, \dots, a_n) \in M_n(\mathbb{R})$ eine Diagonalmatrix und sei f die durch $x \mapsto Ax$ gegebene lineare Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$. Ist $M \subseteq \mathbb{R}^n$ Jordan-messbar, so gilt dies auch für $f(M)$, und es gilt $\mu_n(f(M)) = |a_1 \cdots a_n| \mu_n(M)$.*

- (2) Sei $A = E_{ij}(a) \in M_n(\mathbb{R})$ eine Elementarmatrix ($i \neq j, a \in \mathbb{R}$) und sei f die durch $x \mapsto Ax$ gegebene lineare Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$. Ist $M \subseteq \mathbb{R}^n$ Jordan-messbar, so gilt dies auch für $f(M)$, und es gilt $\mu_n(f(M)) = \mu_n(M)$.

BEWEIS. In beiden Fällen genügt es zu beweisen, dass Quader unter f auf Jordan-messbare Teilmengen abgebildet werden und dass in diesen Fällen die gegebenen Formeln für das Maß des Bildes unter f richtig sind. Dann folgt dieselbe Aussage auch zunächst für elementare Teilmengen, indem wir sie als disjunkte Vereinigungen von Quadern schreiben, und sodann für beliebige Jordan-messbare Teilmengen von $\mathbb{R}^n$, indem wir sie geeignet durch elementare Teilmengen »approximieren« (und beachten, dass die Änderung des Maßes sich in einem konstanten, d.h. von M unabhängigen, Faktor widerspiegelt).

Teil (1) ist damit klar, denn f bildet Quader auf Quader ab, und die Seitenlängen des Bildquaders entstehen aus denen des ursprünglich gegebenen Quaders durch Multiplikation mit den Absolutbeträgen der Diagonaleinträge der Diagonalmatrix A .

Für Teil (2) argumentieren wir wie folgt. Wir betrachten der Einfachheit halber den Fall $n = 2$. Der allgemeine Fall lässt sich darauf zurückführen, weil f auf dem Unterraum der Dimension $n - 2$, der von den Standardbasisvektoren e_k mit $k \neq i, j$ erzeugt wird, als Identität operiert, und nur auf dem Komplement $\langle e_i, e_j \rangle$ »etwas passiert«. Ferner sei ohne Einschränkung $i = 1, j = 2$.

Das Bild des Rechtecks $[a, b] \times [a', b']$ unter f ist ein Parallelogramm mit derselben Grundseite $[a, b]$ und der Höhe $b' - a'$. Es ist leicht zu sehen, dass dieses Jordan-messbar ist und Flächeninhalt $(b - a)(b' - a')$ hat. Entweder, indem man es durch Rechtecke geeignet ausschöpft bzw. überdeckt und so das innere und äußere Jordan-Maß abschätzt, oder indem man den Zusammenhang zwischen Jordan-Maß und Riemann-Integral ausnutzt. $\square$

Eine weitreichende Verallgemeinerung von Teil (2) des Lemmas ist der folgende Satz, dessen Beweis wir aber an dieser Stelle auslassen. (Es handelt sich um einen Spezialfall des Satzes von Fubini für das Riemann-Integral.)

SATZ II.65 (Prinzip von Cavalieri). Seien $X, Y \subseteq \mathbb{R}^n$ Jordan-messbare Teilmengen. Es sei $n = r + s$ und für $a \in \mathbb{R}^r$ sei

$$X_a = \{y \in \mathbb{R}^s; (a, y) \in X \subset \mathbb{R}^n = \mathbb{R}^r \times \mathbb{R}^s\}.$$

Entsprechend bilden wir $Y_a \subset \mathbb{R}^s$ für $a \in \mathbb{R}^r$.

Wir setzen voraus, dass alle Teilmengen X_a, Y_a von $\mathbb{R}^s$ Jordan-messbar seien und dass $\mu_s(X_a) = \mu_s(Y_a)$ für alle $a \in \mathbb{R}^r$ gilt. Dann folgt $\mu_n(X) = \mu_n(Y)$.

Aus Lemma II.64 erhalten wir nun leicht den folgenden allgemeinen Satz, der noch einmal die Bedeutung der Determinante einer linearen Abbildung unterstreicht und eine geometrische Interpretation dafür liefert.

SATZ II.66. Sei $M \subseteq \mathbb{R}^n$ eine Jordan-messbare Teilmenge und sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung. Dann ist auch $f(M)$ Jordan-messbar und es gilt

$$\mu_n(f(M)) = |\det(f)| \mu_n(M).$$

BEWEIS. Ist die Abbildung f nicht invertierbar, so ist ihr Bild (und damit auch $f(M)$) enthalten in einem echten Untervektorraum von $\mathbb{R}^n$, und es folgt $\mu_n(f(M)) = 0$.

Andernfalls können wir f schreiben als eine Verkettung einer linearen Abbildung mit Determinante 1 und einer linearen Abbildung, die durch eine Diagonalmatrix gegeben ist, also letztlich als Verkettung von Abbildungen der Form, wie sie in Lemma II.64 betrachtet wurden. Aus dem Lemma folgt dann direkt die Behauptung. $\square$

Diese Interpretation der Determinante ist unter anderem wichtig in höherdimensionalen Analysis, zum Beispiel für die Transformationsformel, die höherdimensionale Verallgemeinerung der Substitutionsregel für Integrale.

Als konkrete Anwendung können wir den Flächeninhalt der Ellipse

$$E = \left\{ (x, y)^t \in \mathbb{R}^2; \frac{x^2}{a^2} + \frac{y^2}{b^2} \leq 1 \right\}$$

(für $a, b \in \mathbb{R}_{>0}$) berechnen. Denn E ist das Bild der Einheitskreisscheibe unter der Abbildung $x \mapsto \text{diag}(a, b)x$ und es folgt $\mu_2(E) = \pi ab$.

Kodierungstheorie *

12.1. Einführung und Definitionen

In diesem Abschnitt erklären wir einige grundlegende Konzepte der *Kodierungstheorie*, um zu illustrieren, wie der Begriff des endlichen Körpers und speziell Methoden der linearen Algebra über endlichen Körpern dabei helfen, ein »real world problem« einer mathematischen Behandlung zugänglich zu machen.

Das Grundproblem der Kodierungstheorie ist, eine Nachricht effizient über einen Kommunikationskanal (ein Netzkabel, eine WLAN-Verbindung, Speichern und Auslesen von Daten auf einer Festplatte, CD etc.) zu übermitteln, der die übertragenen Nachrichten mit einer gewissen Wahrscheinlichkeit verfälscht. Mit anderen Worten: Auch wenn ein Kratzer auf der CD ist, soll es möglich sein, die gewünschten Informationen zu extrahieren. Eine naive Möglichkeit wäre, zum Beispiel die Zahl 321 654 987 in der Form

321 654 987 321 654 987 321 654 987

zu speichern, d.h. die Information wird dreimal wiederholt. Wenn eine dieser 27 Ziffern geändert wird, kann man durch Vergleich der drei Kopien immer noch herausfinden, was die ursprüngliche Zahl war. Allerdings braucht man mit dieser Methode erheblich mehr Speicherkapazität (oder: eine Übertragung der Nachricht würde entsprechend länger dauern), nämlich dreimal soviel, wie für die eigentliche Nachricht.

In der Kodierungstheorie werden Möglichkeiten gesucht (und gefunden und untersucht), dasselbe Ziel zu erreichen und die Ursprungsnachricht möglichst wenig zu verlängern. Dieses Ziel ist nicht zu verwechseln mit der Verschlüsselungstheorie oder Kryptographie, in der man versucht, Nachrichten so umzuschreiben (zu *verschlüsseln*), dass sie einem Außenstehenden nicht verständlich sind, aber vom Adressaten wieder lesbar gemacht (*entschlüsselt*) werden können. Auch dabei gibt es interessante mathematische Fragen, die jedoch hier nicht das Thema sind; siehe aber Abschnitt 12.5 für eine Verbindung zwischen diesen beiden Themen.

Wir machen zwei Grundannahmen: Erstens: Alle Codewörter werden mit der gleichen Wahrscheinlichkeit als Nachricht verschickt. (Sonst wäre es vielleicht besser, den Code auf die häufiger verschickten Nachrichten hin zu optimieren.) Zweitens setzen wir voraus, dass für jedes gesendete Zeichen auch genau ein Zeichen beim Empfänger ankommt (möglicherweise jedoch ein anderes) und dass es für $m < n$ wahrscheinlicher ist, dass m (also weniger) Fehler bei der Übertragung passieren, als dass n Fehler auftreten. (Es wären auch Kommunikationskanäle denkbar, die immer direkt mehrere Zeichen verändern, wo aber nur sehr selten einzelne Fehler auftreten.) In sehr vielen Praxisanwendungen sind das realistische Voraussetzungen. Man kann die Theorie erweitern auf Situationen, wo diese Annahmen verletzt sind, das wollen wir aber an dieser Stelle nicht tun.

BEISPIEL 12.1. Ein einfaches Beispiel, das das grundlegende Prinzip illustriert, ist die Verwendung eines Paritätsbits: Die zu sendende Nachricht wird in Pakete von »Wörtern« einer festen Länge (zum Beispiel sieben Zeichen), die jeweils 0 oder 1 sein können (man spricht von 7 *Bits*). Zusätzlich wird immer ein weiteres Zeichen hinzugefügt, und zwar eine 0 oder 1, so dass von den acht Bits eine gerade Anzahl gleich 1 ist.

Wenn bei der Übertragung genau ein Bit falsch übertragen wird, hat das empfangene Wort eine ungerade Anzahl von Einsen. Die Empfänger*in kann den Übertragungsfehler also feststellen (allerdings nicht herausfinden, was die ursprüngliche Nachricht war). $\diamond$

Wenn über den Kommunikationskanal die Symbole 0 oder 1 übertragen werden können, und die Wahrscheinlichkeit der korrekten Übertragung für jedes Zeichen durch dieselbe Zahl p , $0 \leq p \leq 1$ gegeben ist (zum Beispiel bedeutet $p = 0,85$, dass in 85% der Fälle das gesendete Zeichen richtig übertragen wird), dann gilt der Satz von Shannon ([vL] Theorem 2.2.3), der umgangssprachlich ausgedrückt Folgendes besagt (die präzise Formulierung dort beinhaltet auch eine Aussage zur Übertragungsrate):

Sei $p \neq \frac{1}{2}$. Für jede vorgegebene positive (kleine) Wahrscheinlichkeit ε gibt es Codes, bei denen die Wahrscheinlichkeit, ein gesendetes Codewort falsch zu dekodieren, kleiner als ε ist. (Wichtig ist dabei, dass man sich erlaubt, ausreichend lange Codewörter zu benutzen. Je kleiner ε ist, desto längere Codewörter wird man in der Regel benötigen.)

Nachrichten zu betrachten, die als eine Folge von Nullen und Einsen geschrieben werden, ist oft naheliegend. Unter anderem, weil auch ein Computer seine Daten so abspeichert. Natürliche Zahlen haben die Binärdarstellung, die aus Nullen und Einsen besteht; Buchstaben können in Zahlen umgeschrieben und dann ebenso in Binärdarstellung geschrieben werden, usw. Zerlegt man die Nachricht in Abschnitte der Länge m , so kann man jeden solchen Abschnitt als ein Element von $\mathbb{F}_2^m$ betrachten.

Vom mathematischen Aufwand können wir aber an dieser Stelle statt $\mathbb{F}_2$ irgendeinen endlichen Körper betrachten und wollen das dementsprechend tun. (Sie können aber das Kapitel auch einfach mit $q = 2$ weiterlesen.)

Sei K ein endlicher Körper, und sei q die Anzahl der Elemente von K . Man kann zeigen (Ergänzung 6.57), dass q eine Primzahlpotenz sein muss, etwa $q = p^d$ mit einer Primzahl p . Dann ist p die Charakteristik von K , also die kleinste positive Zahl, so dass $1 + \dots + 1 = 0$ (mit p Summanden auf der linken Seite), vergleiche Abschnitt 4.2.2. Wir setzen wie oben voraus, dass alle Nachrichten in Abschnitte von k -Tupeln in K zerteilt werden, also als eine Folge von Vektoren in K^k . Sei $N \subseteq K^k$ die Menge aller »Wörter«, also aller Elemente von K^k , die tatsächlich als Nachrichten(-teile) verwendet werden sollen.

Das Grundprinzip der Kodierungstheorie ist nun, die Nachrichtenwörter vor der Übertragung durch andere Wörter zu ersetzen und damit auf geschickte Weise eine Redundanz hinzuzufügen, die es der Empfänger*in ermöglicht, Übertragungsfehler festzustellen und bestenfalls automatisch zu korrigieren. Dazu suchen wir eine injektive Abbildung $c: N \rightarrow K^n$ von der Menge aller Nachrichtenwörter nach K^n (für ein geeignetes n , das in der Regel größer sein wird als m). Statt einer Nachricht $v \in K^k$ wird dann $c(v)$ übertragen. Entscheidend ist dabei, eine geeignete Wahl für die Abbildung c zu treffen, d.h. zu entscheiden, welche der q^n Elemente von K^n in ihrem Bild liegen und daher tatsächlich verwendet werden, damit man auch bei (wenigen) Fehlern in der Übertragung noch auf den gesendeten Vektor zurückschließen kann. Das Bild der Abbildung c nennt man den verwendeten *Code*; c ist die *Kodierungsfunktion*. Für unsere weiteren Betrachtungen werden N und c keine große Rolle spielen. Wir konzentrieren uns auf das Bild der Abbildung c und machen die folgende Definition:

DEFINITION 12.2. Sei K ein endlicher Körper, sei $n \geq 1$. Ein *Code* der Länge n über K ist eine Teilmenge $C \subseteq K^n$. $\dashv$

Wie in der Kodierungstheorie üblich, wollen wir in diesem Kapitel **die Elemente von K^n als Zeilenvektoren verstehen**, wir identifizieren also $K^n = M_{1 \times n}(K)$.

Um der Empfänger*in einer Nachricht ein Verfahren an die Hand zu geben, um Übertragungsfehler festzustellen und sie gegebenenfalls korrigieren zu können, definiert man

DEFINITION 12.3. Seien K ein endlicher Körper und C ein Code der Länge n über K .

- (1) Für $v = (v_1, \dots, v_n), w = (w_1, \dots, w_n) \in K^n$ ist die *Hamming-Distanz* zwischen v und w definiert als

$$d(v, w) = \#\{i; v_i \neq w_i\},$$

also die Anzahl der Einträge der Vektoren v und w , die sich unterscheiden.

- (2) Die (minimale) Hamming-Distanz von C ist

$$d(C) = \min\{d(v, w); v, w \in C, v \neq w\}.$$

—

Es ist klar, dass man keine vernünftigen Aussagen treffen kann, wenn der Kommunikationskanal zu viele Einträge der gesendeten Nachricht verfälscht. Der Begriff der Hamming-Distanz ermöglicht uns eine präzise Aussage (Satz 12.5), mit wie vielen Fehler in einer Nachricht es noch möglich ist festzustellen, dass ein Übertragungsfehler aufgetreten sein muss, und bei wie vielen Fehlern sogar die ursprüngliche Nachricht gefunden werden kann.

Der Code aus Beispiel 12.1 hat die minimale Hamming-Distanz 2: Es gibt zwar Codewörter, die sich an nur zwei Stellen unterscheiden. Die Paritätsbedingung, dass die Anzahl der Einsen in jedem Codewort gerade sein muss, bewirkt aber, dass es keine zwei Codewörter mit Hamming-Distanz 1 geben kann.

BEISPIEL 12.4. Der folgende Code $C \subset \mathbb{F}_2^7$ hat Hamming-Distanz 3, wie man leicht überprüft. Wir schreiben die Zeilenvektoren hier ohne Kommata, um es etwas übersichtlicher zu machen.

$$C = \{(1000111), \\ (0100110), \\ (0010101), \\ (0001011)\}$$

Wenn bei der Übertragung eines der Wörter aus C ein oder zwei Einträge falsch übertragen werden, kommt bei der Empfänger*in ein Wort an, das kein Element von C ist. Dieser Übertragungsfehler lässt sich also erkennen. Wenn genau ein Eintrag falsch übertragen wird, dann lässt sich durch Vergleich mit den Elementen von C auch feststellen, welches Wort übertragen wurde. Der Code kann also 1 Fehler »korrigieren«. Siehe Satz 12.5.

Es gibt Teilmengen $C' \subset \mathbb{F}_2^7$, die C als echte Teilmenge enthalten, und auch Hamming-Distanz 3 haben (und daher »bessere« Codes sind). Versuchen Sie, ein möglichst großes C' zu finden! Siehe auch Abschnitt 12.4. $\diamond$

Die Funktion $d: K^n \times K^n \rightarrow \mathbb{Z}_{\geq 0}$ hat die Eigenschaften einer Metrik (oder Distanzfunktion), d.h.

- (a) Es gilt $d(v, w) = 0$ genau dann, wenn $v = w$,
 (b) Es gilt $d(v, w) = d(w, v)$ für alle $v, w \in K^n$,
 (c) Es gilt die *Dreiecksungleichung*

$$d(u, v) + d(v, w) \geq d(u, w) \quad \text{für alle } u, v, w \in K^n.$$

Alle drei Eigenschaften sind leicht einzusehen.

Im folgenden Satz bezeichnet $\lceil - \rceil$ die *Aufrundungsfunktion*, manchmal auch die *obere Gaußklammer* genannt: Für eine reelle Zahl x ist $\lceil x \rceil$ die kleinste ganze Zahl, die $\geq x$ ist. Zum Beispiel ist $\lceil x \rceil = x$ für alle $x \in \mathbb{Z}$ und $\lceil \frac{1}{2} \rceil = 1$. Wie man leicht sieht, gilt $2\lceil \frac{d}{2} \rceil \leq d + 1$ für jede ganze Zahl d . Das werden wir im Beweis benutzen.

SATZ 12.5. Sei C ein Code der Länge n über K . Wir betrachten das Szenario, dass die Absender*in eine Nachricht $v \in K^n$ verschickt. Den Vektor, den die Empfänger*in erhält, nennen wir w .

- (1) Wenn bei der Übertragung höchstens $d(C) - 1$ Einträge des Vektors v falsch übertragen werden, dann gilt: Ist $w \in C$, so ist $w = v$. Ist $w \notin C$, so muss es Übertragungsfehler gegeben haben. Wir sagen, dass der Code C bis zu $d(C) - 1$ Übertragungsfehler erkennen kann.
- (2) Wenn bei der Übertragung höchstens $\lceil d(C) / 2 \rceil - 1$ Einträge des Vektors v falsch übertragen werden, dann ist v das eindeutig bestimmte Element von C mit $d(v, w) \leq \lceil d(C) / 2 \rceil - 1$. Wir können also die Originalnachricht aus w rekonstruieren. Wir sagen, dass der Code C bis zu $\lceil d(C) / 2 \rceil - 1$ Übertragungsfehler erkennen kann.

BEWEIS. zu (1). Da höchstens $d(C) - 1$ Einträge falsch übertragen wurden, gilt $d(v, w) < d(C)$. Ist $w \in C$, so folgt daraus $v = w$ nach Definition von $d(C)$ als Minimum der Distanzen von verschiedenen Elementen aus C .

Ist $w \notin C$, dann müssen natürlich Übertragungsfehler aufgetreten sein, da *alle* Nachrichten Elemente aus C sind.

zu (2). Es ist klar, dass die Voraussetzung $d(v, w) \leq \lceil d(C) / 2 \rceil - 1$ impliziert. Es ist nur noch zu zeigen, dass es kein Element $v' \in C$, $v \neq v'$, geben kann mit $d(v', w) \leq \lceil d(C) / 2 \rceil - 1$. In diesem Fall wäre aber

$$d(C) \leq d(v, v') \leq d(v, w) + d(v', w) \leq d(C) - 1,$$

ein Widerspruch. Hier haben wir benutzt, dass die Hamming-Distanz die Dreiecksungleichung erfüllt und symmetrisch ist. $\square$

Bei der Wahl eines Codes würden wir also prinzipiell C gerne so wählen, dass $\#C$ groß, $d(C)$ groß und n klein sind. Außerdem ist es für die praktische Anwendung wichtig, dass die Funktionen *Kodieren einer Nachricht* und *Dekodieren einer Nachricht* möglichst einfach sind. Zum Beispiel wäre es praktisch, das De-/Kodieren einer Nachricht durch eine einfache Rechnung durchzuführen (im Vergleich dazu, dass man ein riesiges »Wörterbuch« speichern müsste, in dem jede Übersetzung nachgeschlagen werden muss).

12.2. Lineare Codes

Bisher haben wir noch keine lineare Algebra gesehen. Das wollen wir jetzt ändern. Nach dem vorherigen Abschnitt stellt sich die Frage, wie man geeignete Codes C findet. Dabei ist es naheliegend, die Vektorraumstruktur von K^n auszunutzen.

DEFINITION 12.6. Ein Code $C \subseteq K^n$ heißt *linearer Code*, wenn C ein K -Untervektorraum von K^n ist.

Wenn $k = \dim C$ ist, dann sagt man auch, C sei ein $[n, k]$ -Code (oder ein $[n, k, d]$ -Code wobei $d = d(C)$ die minimale Hamming-Distanz von C ist). $\dashv$

Wenn C ein Untervektorraum ist, können wir die minimale Hamming-Distanz berechnen, indem wir alle Vektoren mit 0 vergleichen:

LEMMA 12.7. Ist C ein linearer Code, so gilt

$$d(C) = \min\{d(0, v); v \in C \setminus \{0\}\}.$$

BEWEIS. Für $v, w \in K^n$ gilt $d(v, w) = d(v - w, 0)$, wie man unmittelbar nachprüft. Weil für $v, w \in C$ auch $v - w$ in C liegt, folgt daraus die Behauptung. $\square$

Wir suchen dann lineare Codes, so dass $\dim C$ und $d(C)$ möglichst groß sind, aber dabei n möglichst klein ist.

Wie oben gesagt, wollen wir die Elemente von K^n als *Zeilenvektoren* betrachten.

DEFINITION 12.8. Sei $C \subseteq K^n$ ein linearer Code. Eine *Erzeugermatrix* (oder *Generatormatrix*) ist eine Matrix, deren Zeilen eine Basis von C bilden.

Wir sagen, eine Erzeugermatrix sei in *Standardform*, wenn sie eine Blockmatrix der Form $\begin{pmatrix} E_k & A \end{pmatrix}$ mit $A \in M_{k \times (n-k)}(K)$ ist. $\dashv$

Wenn G eine Erzeugermatrix des linearen Codes C ist, dann ist die Abbildung $w \mapsto wG$ ein Isomorphismus $K^k \rightarrow C$. Wenn wir die Ursprungsnachrichten als Wörter in K^k schreiben, ist die Kodierungsfunktion also einfach die durch G gegebene lineare Abbildung (im Zeilenvektor-Sinn).

BEISPIEL 12.9. Wir hatten zu Beginn den Paritätscheck-Code betrachtet, der einem Wort aus sieben Bits (Nullen und Einsen) eine achte Null oder Eins so hinzufügt, dass die Gesamtzahl der Einsen gerade ist.

Die Menge der Nachrichten ist also $\mathbb{F}_2^7$, der Code C ist gegeben durch

$$C = \{(x_1, \dots, x_8) \in \mathbb{F}_2^8; \sum_{i=1}^8 x_i = 0\},$$

wobei die Summe in $\mathbb{F}_2$ zu bilden ist. Der Code C besitzt eine Erzeugermatrix in Standardform, und zwar

$$G = \begin{pmatrix} I & & I \\ & \ddots & \vdots \\ & & I & I \end{pmatrix}.$$

$\diamond$

Ist C ein linearer Code, so hat C höchstens eine Erzeugermatrix in Standardform.

Für die Qualität eines Codes tut es offenbar nichts zur Sache, wenn wir eine Permutation der Standardbasisvektoren von K^n durchführen (also C durch sein Bild unter einem Automorphismus von K^n ersetzen, der durch eine Permutationsmatrix gegeben ist). Wir nennen zwei Codes, die durch eine Operation dieser Art ineinander übergehen, *äquivalent*. Da wir jede Matrix durch elementare Zeilenumformungen auf reduzierte Zeilenstufenform bringen können und diese den von den Zeilen aufgespannten Untervektorraum nicht verändern, ist klar, dass es zu jedem linearen Code einen äquivalenten Code gibt, der eine Erzeugermatrix in Standardform besitzt.

DEFINITION 12.10. Sei C ein linearer $[n, k]$ -Code über K . Eine *Paritätsprüfmatrix* ist eine Matrix $H \in M_{(n-k) \times n}(K)$, so dass für alle $c \in K^n$ gilt:

$$c \in C \iff Hc^t = 0,$$

d.h. C ist die Lösungsmenge des durch $Hx = 0$ gegebenen linearen Gleichungssystems (wenn wir die Elemente von C zu Spaltenvektoren transponieren). $\dashv$

Ist $G = (E_k \mid A)$ eine Erzeugermatrix von C in Standardform, dann ist $H = (-A^t \mid E_{n-k})$ eine Paritätsprüfmatrix, denn es gilt offenbar $HG^t = 0$, also $\{c^t; c \in C\} \subseteq \text{Ker}(H)$, und $\text{rg}(H) = n - k$, so dass sogar $\{c^t; c \in C\} = \text{Ker}(H)$ folgt.

Die minimale Hamming-Distanz eines Codes hat die folgende Interpretation in Termen einer Paritätsprüfmatrix:

LEMMA 12.II. Sei C ein linearer $[n, k]$ -Code mit minimaler Hamming-Distanz $d(C)$. Sei H eine Paritätsprüfmatrix von C , und sei d die kleinste Zahl, so dass je $d - 1$ Spalten von H stets linear unabhängig sind, aber so dass es d Spalten in H gibt, die linear abhängig sind. Dann gilt $d = d(C)$.

BEWEIS. Dass es d Spalten in H gibt, die linear abhängig sind, ist äquivalent dazu, dass die Lösungsmenge des linearen Gleichungssystems $Hx = 0$ einen Vektor x enthält, in dem höchstens d Einträge $\neq 0$ sind. Aus diesem Argument und Lemma 12.7 folgt sofort die Behauptung. $\square$

12.3. Die Gilbert-Varshamov-Schranke

12.4. Hamming-Codes

12.5. Das McEliece-Verfahren

Ein wichtiger Aspekt, den wir bisher außen vor gelassen haben (und den wir auch hier nur indirekt streifen), ist die Frage nach einer effizienten Kodierungs- bzw. Dekodierungsfunktion für einen Code. Je nach Anzahl der Codewörter ist es sicherlich unpraktikabel, ein empfangenes Wort mit jedem Codewort zu vergleichen und dasjenige mit minimalem Hammingabstand herauszusuchen. Wenn man nur irgendeine Erzeugermatrix G eines Codes kennt, sind die Möglichkeiten der schnellen Dekodierung (bei großen Codes) eingeschränkt. Das Kodieren geht aber relativ schnell, weil man hierzu nur ein Matrizenprodukt der Form wG (für ein zu versendendes Wort w) berechnen muss.

Dass es Codes gibt, bei denen es auch für große Codes schnelle Dekodierfunktionen gibt (zum Beispiel die sogenannten Goppa-Codes) macht sich das Verschlüsselungsverfahren von Robert McEliece zu Nutze, das wir hier skizzieren wollen. Es handelt sich um ein Public-Key-Verfahren, das auf Kodierungstheorie und Linearer Algebra basiert. Eine Besonderheit ist, dass es nach jetzigem Kenntnisstand nicht anfällig ist für Angriffe mit Quantencomputern (anders als die meisten Verfahren, die momentan gängig sind). Ein Nachteil ist, dass die Schlüssel, die man dafür speichern und versenden muss, relativ groß sind.

Unter einem Public-Key-Verfahren versteht man ein Verschlüsselungsverfahren, bei dem die spätere Empfänger*in einer verschlüsselten Nachricht alle Informationen, die die Absender*in der Nachricht zum Verschlüsseln benötigt, öffentlich zur Verfügung stellen kann (insbesondere könnte man diese Informationen über Kanäle verschicken, die von anderen abgehört werden können). Dieser »Public Key« muss natürlich die Eigenschaft haben, dass er zwar das Verschlüsseln, jedoch nicht das Entschlüsseln einer Nachricht erlaubt.

Die Empfängerin der verschlüsselten Nachricht, die wir wie in der Kryptographie üblich Alice nennen wollen, wählt die folgenden Informationen. Diese bilden den *privaten Schlüssel*, den sie geheim hält.

Einen (geeigneten) linearen $[n, k]$ -Code C über $\mathbb{F}_2$, der bis zu t Fehler korrigieren kann. (Damit das Verfahren sicher ist, werden Werte von $n = 2048$, $k = 1751$, $t = 27$ empfohlen; für Sicherheit gegen Quantencomputerangriffe noch höhere Werte.) Sie wählt den Code so, dass sie eine schnelle Dekodierfunktion kennt. Sei G eine Erzeugermatrix dieses Codes in Standardform.

Alice wählt außerdem eine invertierbare Matrix $S \in GL_k(\mathbb{F}_2)$ und eine Permutationsmatrix $P \in GL_n(\mathbb{F}_2)$. Sie berechnet das Matrizenprodukt $G' := SGP$ und veröffentlicht dann den folgenden

Öffentlichen Schlüssel: (G', t) .

Während die Matrix G (in Standardform) in der Regel Rückschlüsse darauf zulassen wird, aus welcher Familie von Codes Alice ihren Code gewählt hat, und damit auf eine effektive Dekodierfunktion, ist das für die Matrix G' nicht der Fall.

Bob, der Alice eine Nachricht schicken möchte, schreibt seine Nachricht im Klartext als Folge von Wörtern $m \in \mathbb{F}_2^k$. Für jedes m bildet er $c' := mG'$ (die »regulär« mit dem Code mit Erzeugermatrix G' kodierte Nachricht), wählt zufällig ein Element $z \in \mathbb{F}_2^n$ mit t Einsen und $n - t$ Nullen, und schickt Alice die Nachricht $c := c' + z$. Es werden durch die Addition von z also künstlich t Fehler eingefügt.

Alice berechnet aus c das Wort $cP^{-1} (= c'P^{-1} + zP^{-1})$. Da P eine Permutationsmatrix ist, unterscheidet sich dieses Wort an genau t Stellen von mSG . Mit der Dekodierfunktion für den Code mit Erzeugermatrix G kann Alice daraus mS berechnen, und findet damit die ursprüngliche Nachricht als $m = (mS)S^{-1}$.

Die Sicherheit des Verfahrens beruht darauf, dass die Werte n , k und t so gewählt werden, dass die Berechnung von m aus $mG' + z$ nicht in akzeptabler Zeit möglich ist. (Die Durchführbarkeit beruht darauf, dass es eben doch möglich ist, wenn man die zusätzlichen strukturellen Informationen kennt, die Alice zur Hand hat.)

Weitere Informationen: [Wikipedia \(English\)](#)¹

12.6. Quellen/Weiterführende Literatur

J. H. van Lint, *Introduction to Coding Theory*, Springer Graduate Texts in Mathematics **86**, 3rd ed., Springer 1999.

J. H. van Lint, G. van der Geer, *Introduction to Coding Theory and Algebraic Geometry*, DMV Seminar **12**, Birkhäuser 1988.

Siehe auch [Ma] Kap. 28, Kap. 29 für weiterführende Fragen in der Kodierungstheorie, die man mit linearer Algebra beantworten kann.

¹https://en.wikipedia.org/wiki/McEliece_cryptosystem

Graphentheorie *

13.1. Definition

DEFINITION 13.1. Ein (endlicher) Graph $G = (\mathbb{E}, \mathbb{K})$ ist ein Paar, das besteht aus

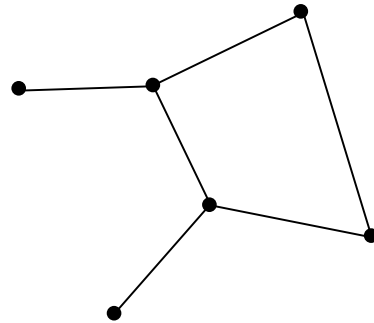
- (a) einer Menge $\mathbb{E}$, deren Elemente die *Ecken* oder *Knoten* des Graphen genannt werden,
- (b) einer Menge $\mathbb{K}$ von zweielementigen Teilmengen von $\mathbb{E}$, deren Elemente die *Kanten* des Graphen G genannt werden.

⊢

Wir können einen Graphen durch eine Zeichnung visualisieren, indem wir die Ecken des Graphen als »dicke Punkte« zeichnen, und zwei Ecken P_1, P_2 genau dann verbinden, wenn die Teilmenge $\{P_1, P_2\}$ ein Element der Menge $\mathbb{K}$ ist.

Die Abbildung zeigt einen Graphen mit 6 Ecken.

Man beachte aber: Nicht jeder Graph lässt sich »überschneidungsfrei« in der Ebene zeichnen. Siehe Abschnitt 13.6.



Der hier definierte Begriff des Graphen hat nichts mit Funktionsgraphen zu tun.

BEMERKUNG 13.2. In unserer Definition

- kann keine Kante einen Knoten mit sich selbst verbinden,
- gibt es zwischen zwei Knoten entweder keine oder eine Kante (aber nicht mehrere).

Je nachdem, wozu man die Theorie benutzen möchte, könnte man die Definition abändern und allgemeinere Formen von Graphen erlauben, als wir es hier tun.

Eine andere nützliche Variante ist der Begriff des *gerichteten Graphen*, in dem jede Kante des Graphen mit einer Richtung versehen wird (und dann auch erlaubt ist, dass zwei Knoten durch je eine Kante in die beiden möglichen Richtungen miteinander verbunden sind). Die Netzwerke in Frage 2.7 kann man als gerichtete Graphen betrachten. ◇

Wie genau die Ecken eines Graphen benannt werden (ob mit 1, 2, 3, ... oder A, B, C, ... , oder noch anders) spielt für uns keine Rolle. Wir definieren deshalb:

DEFINITION 13.3. Seien $G = (\mathbb{E}, \mathbb{K})$ und $G' = (\mathbb{E}', \mathbb{K}')$ Graphen. Ein *Isomorphismus* ist eine Bijektion $\varphi: \mathbb{E} \rightarrow \mathbb{E}'$, so dass für alle $i, j \in \mathbb{E}$ gilt: $\{i, j\} \in \mathbb{K} \Leftrightarrow \{\varphi(i), \varphi(j)\} \in \mathbb{K}'$. Zwei Graphen heißen *isomorph*, wenn es einen Isomorphismus zwischen ihnen gibt. ⊢

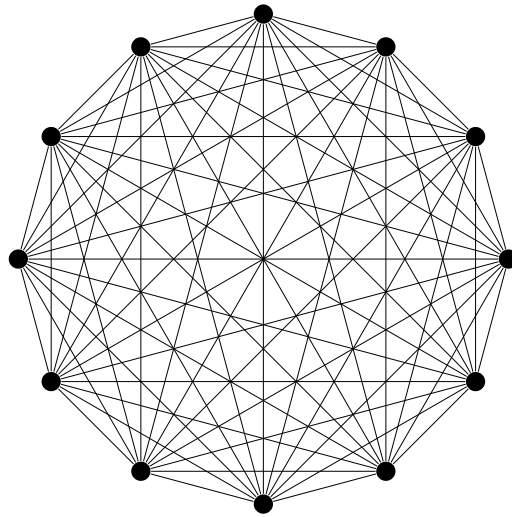


ABBILDUNG 1. Der Graph K_{12} : Jede der 12 Ecken ist mit jeder anderen Ecke durch eine Kante verbunden.

Alle Eigenschaften von Graphen, die wir im folgenden betrachten, werden durch Isomorphismen erhalten.

Es ist allerdings mitunter nicht leicht festzustellen, ob zwei gegebene Graphen zueinander isomorph sind oder nicht. (Natürlich gibt es prinzipiell die Möglichkeit, alle Bijektionen zwischen den Eckenmengen durchzuprobieren und zu testen, ob es sich um Isomorphismen von Graphen handelt. Wenn die Anzahl der Ecken nicht sehr klein ist, ist das aber nicht praktikabel.)

DEFINITION 13.4. Sei $G = (\mathbb{E}, \mathbb{K})$ ein Graph. Ein Graph $G' = (\mathbb{E}', \mathbb{K}')$ heißt ein *Teilgraph* von G , wenn $\mathbb{E}' \subseteq \mathbb{E}$ und $\mathbb{K}' \subseteq \mathbb{K}$. $\dashv$

13.2. Ramsey-Zahlen

BEISPIEL 13.5. Sei $n \geq 1$ eine natürliche Zahl. Der *vollständige Graph* mit n Ecken ist der Graph K_n mit Eckenmenge $\mathbb{E} = \{1, \dots, n\}$ und Kantenmenge $\{\{i, j\}; i, j = 1, \dots, n, i \neq j\}$, d.h. je zwei verschiedene Punkte sind durch eine Kante verbunden. $\diamond$

Zum Beispiel ist der Graph K_3 einfach ein »Dreieck«, d.h. er besteht aus drei Ecken, die jeweils durch eine Kante miteinander verbunden sind. Mit wachsendem n steigt die Anzahl der Kanten in K_n allerdings sehr schnell.

Sei $n \geq 1$ eine fixierte natürliche Zahl. Wir wollen nun als zusätzliche Information in dem Graphen K_n jede Kante entweder blau oder rot einfärben, wie im Beispiel $n = 5$ in Abbildung 2 gezeigt. Die Frage, die wir dann betrachten, ist, ob der eingefärbte Graph K_n einen Teilgraphen der Form K_m hat, dessen Kanten alle blau oder alle rot sind. In dieser Sache gilt der folgende Satz.

SATZ 13.6 (Ramsey). Seien $b, r \geq 1$ natürliche Zahlen. Dann existiert eine Zahl N , so dass für alle $n \geq N$ und jede blau/rote Kantenfärbung des Graphen K_n ein roter r -Teilgraph oder ein blauer b -Teilgraph existiert.

Die kleinste solche Zahl N nennen wir die Ramsey-Zahl zu b und r und bezeichnen sie mit $R(b, r)$.

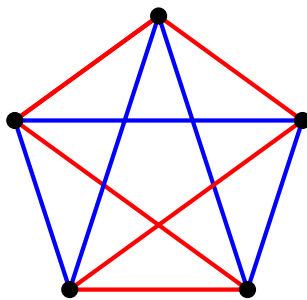


ABBILDUNG 2. Der Graph K_5 mit blau/rot eingefärbten Kanten. Es gibt weder ein blaues noch ein rotes Dreieck.

REFERENZ ZUM BEWEIS. Man kann dies durch Induktion nach $b + r$ beweisen, siehe zum Beispiel [Ramsey's theorem \(Wikipedia\)](https://en.wikipedia.org/wiki/Ramsey%27s_theorem)¹. $\square$

BEISPIEL 13.7. Es ist $R(3, 3) = 6$. In der Tat zeigt Abbildung 2, dass $R(3, 3) > 5$ sein muss. Es ist daher nur noch zu zeigen, dass es in jedem rot/blau gefärbten Graphen K_6 ein rotes Dreieck oder ein blaues Dreieck gibt.

Wir fixieren eine Färbung der Kanten von K_6 . Sei i eine fixierte Ecke. Wie jede Ecke in K_6 ist sie mit 5 anderen Ecken verbunden, und von diesen Kanten können wir mindestens 3 finden, die dieselbe Farbe haben. Indem wir rot und blau vertauschen, falls erforderlich, können wir für das Argument ohne Einschränkung annehmen, dass von i drei rote Kanten ausgehen. Wenn von den Verbindungen zwischen den drei Endpunkten dieser Kanten eine rot ist, dann haben wir ein rotes Dreieck gefunden. Sind die Verbindungen alle drei blau, dann bilden sie ein blaues Dreieck.

Wir können der Aussage folgendermaßen einen Alltagsbezug geben (auch wenn der Nutzen fragwürdig ist): Stellen Sie sich vor, auf einer Party sind mindestens 6 Personen anwesend, von denen je zwei entweder miteinander bekannt, oder einander unbekannt sind. Dann gibt es eine Dreiergruppe von Personen, die sich alle gegenseitig kennen, oder eine Dreiergruppe von Personen, die sich gegenseitig nicht kennen. $\diamond$

BEISPIEL 13.8. Die Ramsey-Zahl $R(5, 5)$ ist nicht bekannt!

Man kann zeigen, dass

$$43 \leq R(5, 5) \leq 48,$$

es handelt sich also nicht um eine besonders große Zahl. Dennoch ist die Berechnung so komplex, dass sie bisher nicht gelungen ist. (Das Problem ist, dass zum Beispiel K_{43} insgesamt 946 Kanten hat und damit die Zahl der möglichen Kantenfärbungen astronomisch hoch ist.) $\diamond$

13.3. Die Adjazenzmatrix eines Graphen

Problemstellung Was ist eine gute Methode, um in einem gegebenen Graphen alle Wege einer fixierten Länge zu zählen? Wir werden dieses Problem mithilfe des Matrizenprodukts »lösen«. Je nachdem, wie groß der Graph ist, ist die Rechnung eventuell nicht ohne Weiteres durchführbar. Den Matrizenkalkül ins Spiel zu bringen ist aber jedenfalls eine Verbesserung.

Wir stellen durch die folgende Definition eine Verbindung zur linearen Algebra her, die zwar recht banal aussieht, aber dennoch sehr nützlich ist.

¹https://en.wikipedia.org/wiki/Ramsey%27s_theorem

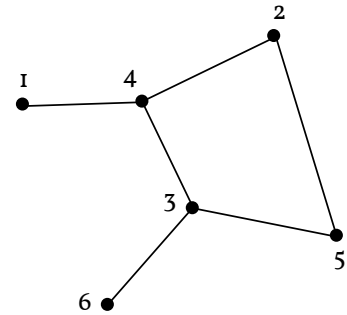
DEFINITION 13.9. Sei $G = (\mathbb{E}, \mathbb{K})$ ein Graph. Wir fixieren eine Identifizierung $\mathbb{E} = \{1, \dots, n\}$. Die Adjazenzmatrix (oder Nachbarschaftsmatrix) von G ist die Matrix $A = (a_{ij}) \in M_{n \times n}(\mathbb{Q})$ mit

$$a_{ij} = \begin{cases} 1 & i, j \text{ sind durch eine Kante verbunden,} \\ 0 & \text{sonst.} \end{cases}$$

—

BEISPIEL 13.10. (1) Die Adjazenzmatrix des Graphen in der nebenstehenden Abbildung ist

$$\begin{pmatrix} 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 \\ 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \end{pmatrix}.$$



(2) Die Adjazenzmatrix von K_n ist die Matrix, deren Diagonaleinträge 0 und deren andere Einträge alle 1 sind.

◇

Interessanterweise haben auch die Potenzen der Adjazenzmatrix eine graphentheoretische Interpretation. Sei G ein Graph mit Eckenmenge $\{1, \dots, n\}$ und Adjazenzmatrix A . Ein Produkt der Form $a_{ij}a_{jk}$ ist genau dann 1, wenn in G ein »Weg« $i-j-k$ existiert, d.h. wenn sowohl i und j als auch j und k miteinander verbunden sind. Wir definieren

DEFINITION 13.11. Sei G ein Graph. Ein Weg (oder Pfad) in G ist ein Tupel $(e_0, \dots, e_l)$ von Ecken in G , so dass für alle i die Ecken e_i und e_{i+1} durch eine Kante verbunden sind. Wir sprechen auch von einem Weg von e_0 nach e_l .

Die Zahl l heißt die Länge des Wegs.

—

Dann ist $\sum_{j=1}^n a_{ij}a_{jk}$ gerade die Anzahl der Wege der Länge 2 von i nach k .

SATZ 13.12. Sei G ein Graph mit Eckenmenge $\{1, \dots, n\}$ und Adjazenzmatrix A . Sei l eine natürliche Zahl. Sei $B = (b_{ij})_{i,j} = A^l$. Dann ist b_{ij} die Anzahl der Wege der Länge l von i nach j in G .

BEWEIS. Wir führen Induktion nach l . Für $l = 1$ ist die Sache klar. (Und auch für $l = 0$, wenn man möchte: Ein Weg der Länge 0 ist nichts anderes als eine Ecke in G .)

Sei nun $l > 1$, und sei $B = (b_{ij})_{i,j} = A^{l-1}$. Ein Weg der Länge l von i nach k lässt sich zerlegen als ein Weg von i zu einer Ecke j der Länge $l-1$ und dem letzten Schritt des Weges von j nach k , d.h. j muss mit k benachbart sein. Die Anzahl der Wege der Länge l ist daher die Summe aller b_{ij} für diejenigen j , die mit k benachbart sind – also die Summe $\sum_{j=1}^n b_{ij}a_{jk}$, und das ist gerade der Eintrag von A^l an der Stelle (i, k) . □

Wir können also Matrizenrechnung benutzen, um die Anzahl der Wege einer vorgegebenen Länge zwischen zwei Knoten in einem Graphen zu berechnen. Um zum Beispiel alle Wege der Länge 16 zu zählen, muss man für die Adjazenzmatrix A die 16-te Potenz ausrechnen. Ähnlich wie in Beispiel 5.60 kann man das mit 4 Matrixmultiplikationen machen ($A^2 = AA$, $A^4 = A^2A^2$, ...). Das ist wesentlich effizienter, als direkt im Graphen »herumzuprobieren«.

BEMERKUNG 13.13. Offenbar ist ein Graph durch seine Adjazenzmatrix (bis auf Isomorphie, also bis auf die Benennung der Ecken) eindeutig bestimmt. Überlegen Sie sich zur Übung, welche Matrizen als Adjazenzmatrix eines Graphen auftreten.

Andererseits ist die Adjazenzmatrix eines Graphen erst dann eindeutig festgelegt, wenn wir eine Nummerierung der Ecken fixieren, denn wir müssen ja wissen, auf welche Ecke sich die Einträge der ersten Zeile beziehen, usw. Wenn wir die Ecken anders nummerieren, erhalten wir in aller Regel eine andere Matrix. Es ist aber leicht zu sehen, wie sich die beiden Matrizen voneinander unterscheiden: Wir müssen die Zeilen und die Spalten der ersten Matrix so vertauschen, wie es der Umnummerierung der Ecken entspricht. Das bedeutet, dass die Matrizen durch Konjugation mit einer Permutationsmatrix auseinander hervorgehen.

Das gleiche Argument können wir benutzen, um zu sehen, dass die Adjazenzmatrizen von zwei isomorphen Graphen zueinander konjugiert sind. Daraus ergibt sich eine Möglichkeit, um zu zeigen, dass zwei Graphen nicht isomorph sind: Nämlich zu zeigen, dass die zugehörigen Adjazenzmatrizen nicht zueinander konjugiert sind (zum Beispiel, weil sie nicht dieselben Eigenwerte haben). $\diamond$

13.4. Teildreiecke suchen

In diesem Abschnitt untersuchen wir die Frage, wie wir in einem Graphen G »Dreiecke«, das heißt Teilgraphen der Form K_3 , finden können. Dazu arbeiten wir wieder mit der Adjazenzmatrix des Graphen.

Sei also A die Adjazenzmatrix (für eine fixierte Eckenummerierung) und sei $B = (b_{ij})_{i,j} = A^2$.

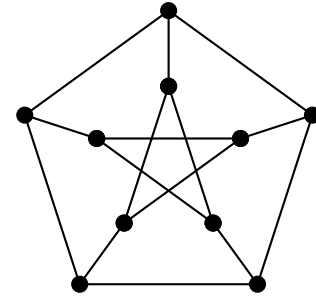
Dann können wir die Frage nach einem Dreieck in G auch so umformulieren: Wir suchen Paare (i, j) von Ecken in G mit $a_{ij} \neq 0, b_{ij} \neq 0$ (denn das bedeutet erstens, dass i und j benachbart sind, und zweitens, dass zwischen i und j ein Weg der Länge 2 existiert; bezeichnet k den Zwischenschritt auf diesem Weg, dann sind i, j und k paarweise miteinander verbunden und bilden daher ein Dreieck).

Abgesehen davon, dass das ursprüngliche Problem damit (vielleicht) eleganter ausgedrückt ist, kann man auch (für sehr große Graphen) einen Vorteil daraus ziehen, was die konkrete Berechnung angeht. Wenn man naiv in dem Graphen G nach einem Dreieck sucht, müsste man alle drei-elementigen Teilmengen der Eckenmenge $\mathbb{E}$ untersuchen; das wären größenordnungsmäßig $(\#\mathbb{E})^3$ Rechenoperationen. Für die Berechnung des Quadrats A^2 der Adjazenzmatrix braucht man mit dem naiven Verfahren zur Berechnung des Matrizenprodukts zwar auch größenordnungsmäßig $(\#\mathbb{E})^3$ Rechenoperationen. Es gibt aber bessere Algorithmen (vgl. Ergänzung 5.33).

Quelle: [Ma] Kap. 10.

13.5. Eigenwerte der Adjazenzmatrix

Der Petersen-Graph ist der hier abgebildete Graph. Er hat 10 Ecken und an jeder Ecke gehen 3 Kanten ab. Es ist klar, dass wir ihn als Teilgraph des vollständigen Graphen K_{10} mit 10 Ecken einbetten können. Da in K_{10} an jeder Ecke 9 Kanten abgehen, ist es denkbar, dass man den Graphen K_{10} vollständig überdecken kann, indem man den Petersen-Graph auf drei verschiedene Weisen als Teilgraph einbettet (also so, dass insgesamt jede Kante von K_{10} in einem der drei Teilgraphen vorkommt).



Wir werden unten sehen, dass dies nicht möglich ist. Es ist aber klar, dass es eine sehr langwierige (und langweilige) Aufgabe wäre, das »per Hand« durch direktes Überprüfen aller Möglichkeiten zu beweisen. Stattdessen gibt es eine elegante Beweismethode mit linearer Algebra.

Zur Lösung dieses Problems untersuchen wir die Eigenwerte der Adjazenzmatrix.

Die Adjazenzmatrix A von K_{10} ist sehr einfach: Es ist die (10×10) -Matrix mit Nullen auf der Diagonalen und Einsen in allen anderen Feldern. Wir schreiben J_{10} für die Matrix, deren Einträge sämtlich $= 1$ sind. Dann gilt also $A = J_{10} - E_{10}$.

Wenn wir den Petersen-Graph als Teilgraph in K_{10} einbetten, so ergibt sich eine Bijektion auf den Eckenmengen. Wir können die Situation also so beschreiben, dass wir die Ecken $\{1, \dots, 10\}$ von K_{10} so mit Kanten verbinden, dass sich ein zum Petersen-Graph isomorpher Graph ergibt. Je nachdem, wie wir das machen, ergibt sich eine unterschiedliche Adjazenzmatrix (bezogen auf die Eckenmenge $\{1, \dots, 10\}$, die wir festhalten).

Angenommen, es gäbe drei Einbettungen des Petersen-Graphs als Teilgraphen von K_{10} , die insgesamt alle Kanten überdecken. Wegen der zahlenmäßigen Übereinstimmung wird dann jede Kante genau einmal getroffen. Sind M_1, M_2, M_3 die drei Adjazenzmatrizen dieser Teilgraphen, so gilt also

$$A = M_1 + M_2 + M_3.$$

Außerdem sind die Matrizen M_i alle zueinander und zur Adjazenzmatrix des Petersen-Graphen (bezüglich irgendeiner fixierten Eckenummerierung) konjugiert (Bemerkung 13.13). Das bedeutet, dass M_1, M_2 und M_3 dieselben Eigenwerte haben, und dass die zugehörigen Eigenräume dieselben Dimensionen haben, Lemma 10.7.

LEMMA 13.14. Sei M eine der drei Matrizen M_1, M_2, M_3 . Dann gilt:

- (1) Der Eigenraum von M zum Eigenwert 1 hat Dimension 5.
- (2) Der Eigenraum von M zum Eigenwert 1 ist enthalten in dem Untervektorraum

$$U_0 := \{(x_1, \dots, x_{10}); \sum_{i=1}^{10} x_i = 0\}$$

von $\mathbb{Q}^{10}$.

BEWEIS. zu (1). Berechne mit dem Gauß-Algorithmus die Dimension von $\text{Ker}(M - E_{10})$. Wir lassen die einfache Rechnung hier aus.

zu (2). In jede Spalte von $M - E_{10}$ sind genau drei Einträge gleich 1, genau einer gleich -1 (der Diagonaleintrag) und alle anderen gleich 0. Die Summe aller Zeilen von $M - E_{10}$ ist

also $(2, \dots, 2)$. Das bedeutet, dass alle Lösungen $(x_1, \dots, x_{10})$ des linearen Gleichungssystems $(M - E_{10})x = 0$ auch die Gleichung

$$2x_1 + \dots + 2x_{10} = 0$$

erfüllen, also in U_0 liegen. □

Die Eigenräume zum Eigenwert 1 von M_1 und M_2 liegen in dem 9-dimensionalen Vektorraum U_0 , müssen sich also nicht-trivial schneiden. Sei $x \neq 0$ ein Element des Durchschnitts. Es gilt dann $J_{10}x = 0$, weil $x \in U_0$. Wir rechnen nun

$$M_3x = (A - M_1 - M_2)x = (J_{10} - I_{10} - M_1 - M_2)x = -3x.$$

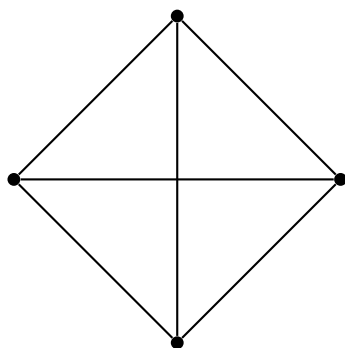
Den gewünschten Widerspruch erhält man nun, indem man (wieder mit dem Gauß-Algorithmus) überprüft, dass der Kern von $M_3 + 3E_{10}$ trivial ist, also dass -3 kein Eigenwert der Adjazenzmatrix des Petersen-Graphen ist.

Quelle: [Ma] Kap. 13; Matoušek verweist auf eine Arbeit von Lossers und Schwenk. Für ein ähnliches aber etwas komplizierteres Problem siehe auch loc. cit. Kap. 14.

13.6. Ausblick: Planare Graphen

Zum Schluss noch ein kleiner Ausblick auf ein weiteres Thema der Graphentheorie, auch wenn wir es hier nicht mit Methoden der Linearen Algebra in Verbindung bringen.

DEFINITION 13.15. Sei G ein Graph. Wir sagen, G sei ein *planarer Graph*, wenn sich G ohne Überschneidung von Kanten in der Ebene zeichnen lässt. ⇐



Die Definition ist ein bisschen salopp, soll an dieser Stelle aber genügen. Der subtile Punkt dabei ist, dass man es einem Graphen nicht unbedingt auf den ersten Blick ansieht, ob er planar ist. Zum Beispiel ist der Graph K_4 ein planarer Graph, auch wenn er hier mit einer Kantenüberschneidung gezeichnet ist. Man kann denselben Graphen nämlich auch so zeichnen, dass sich keine zwei Kanten schneiden (zum Beispiel, indem man die vertikale Kante außen um das Quadrat herumführt). Entsprechend ist es auch nicht trivial zu zeigen, dass ein gegebener Graph *nicht* planar ist; auch wenn es nicht gelingt, den Graph überschneidungsfrei in der Ebene zu zeichnen, könnte es ja sein, dass man eine geschickte Möglichkeit übersehen hat.

Ein berühmter Satz über planare Graphen ist der

SATZ 13.16 (Vierfarbensatz). *In jedem planaren Graphen lassen sich die Ecken mit höchstens 4 Farben so einfärben, dass je zwei benachbarte Ecken unterschiedliche Farben haben.*

Drei Farben reichen im allgemeinen nicht aus, um einen planaren Graphen in dieser Weise einzufärben, wie das Beispiel des Graphen K_4 zeigt.

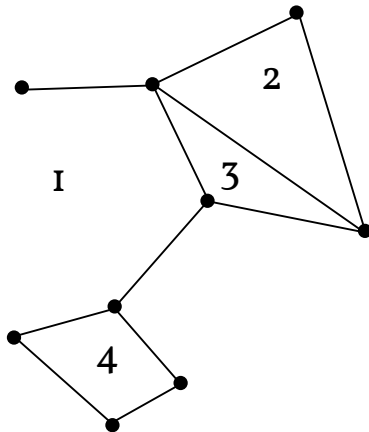
Der Vierfarbensatz lässt sich auch so uminterpretieren, dass man jede Landkarte so einfärben kann, dass je zwei entlang einer Grenze benachbarte Länder unterschiedlich gefärbt werden. Der Satz wurde 1976 von Appel und Haken mit Hilfe von Computern bewiesen. Es ist der erste berühmte mathematische Satz, der mit einem computergestützten Beweis gezeigt wurde. Bis heute ist kein Beweis bekannt, der ohne Computerberechnungen auskommt.

Gleichzeitig ist dieses Problem eine Quelle von Beispielen für Fehler und falsche Beweise. Siehe [Vierfarbensatz \(Wikipedia\)](#)².

Für die folgende Diskussion definieren wir:

DEFINITION 13.17. Ein Graph heißt *zusammenhängend*, wenn je zwei Ecken durch einen Weg verbunden sind (und die Eckenmenge nicht leer ist). $\dashv$

Für zusammenhängende planare Graphen gilt die Eulersche Formel, die zurückgeht auf L. Euler (um 1750). Euler hat mit seiner Lösung des [Königsberger Brückenproblems](#)³ sozusagen die Graphentheorie begründet.



Um die Euler-Formel anzugeben, beobachten wir, dass ein (in der Ebene ohne Kantenüberschneidungen gezeichneter) planarer Graph durch seine Kanten die Ebene in endlich viele »Gebiete« unterteilt. Das »äußere«, unbegrenzte Gebiet zählen wir dabei mit. Die Form und Lage der Gebiete hängt natürlich davon ab, wie der Graph gezeichnet wird. A priori hängt auch die Anzahl der Gebiete davon ab, die Euler-Formel zeigt allerdings, dass das nicht der Fall ist. Der hier gezeichnete Graph unterteilt die Ebene in 4 Gebiete.

SATZ 13.18 (Eulersche Formel für planare Graphen). Sei G ein zusammenhängender planarer Graph, den wir uns in der Ebene gezeichnet vorstellen, sei E die Anzahl seiner Ecken, K die Anzahl der Kanten, und F die Anzahl der Gebiete (»Flächen«), in die G die Ebene unterteilt. Dann gilt

$$E - K + F = 2.$$

Der Beweis des Satzes ist nicht sehr schwierig. Eine naheliegende Möglichkeit ist es, induktiv vorzugehen. Dazu überlegt man sich, dass man jeden zusammenhängenden planaren Graphen, ausgehend vom Graphen mit einer einzigen Ecke, schrittweise aufbauen kann, indem man entweder eine Kante zwischen zwei Ecken hinzufügt, oder eine Ecke und eine Kante, die die neue Ecke mit einer anderen Ecke verbindet, hinzufügt. Da die Formel für den Graphen mit einem einzigen Punkt gilt, muss man sich dann nur noch überlegen, dass diese beiden Schritte die Gültigkeit der Formel erhalten. Versuchen Sie, die Details einzufüllen! Oder schauen Sie zum Beispiel bei Grieser [Gr] nach.

Eng mit dieser Formel verwandt ist die [Eulersche Polyederformel](#)⁴, die besagt, dass dieselbe Relation $E - K + F = 2$ gilt, wenn E die Anzahl der Ecken, K die Anzahl der Kanten und F die Anzahl der Außenflächen eines konvexen Polyeders in $\mathbb{R}^3$ ist. Ein Beispiel ist der Würfel mit 8 Ecken, 12 Kanten und 6 Flächen; in der Tat gilt $8 - 12 + 6 = 2$.

KOROLLAR 13.19. Sei G ein zusammenhängender planarer Graph mit $E > 2$ Ecken und K Kanten.

- (1) Es gilt $3E - 6 \geq K$.
- (2) Es gibt eine Ecke, von der weniger als 6 Kanten abgehen.

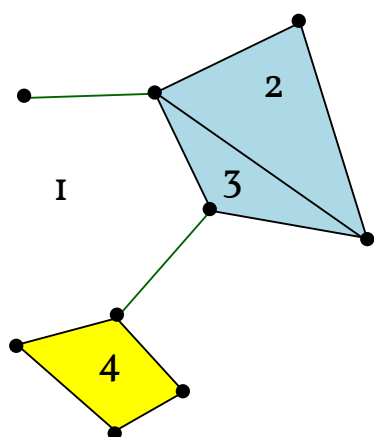
BEWEIS. zu (1). Wir stellen uns G in der Ebene gezeichnet vor und betrachten wie bei der Eulerschen Formel die Gebiete, in die die Ebene durch G zerteilt wird. Sei F_i die Anzahl

² <https://de.wikipedia.org/wiki/Vier-Farben-Satz>

³ https://de.wikipedia.org/wiki/K%C3%B6nigsberger_Br%C3%BCckenproblem

⁴ https://de.wikipedia.org/wiki/Eulerscher_Polyedersatz

der Gebiete, die durch i Kanten von G begrenzt werden. (Dabei zählen wir Kanten doppelt, wenn das Gebiet auf beiden Seiten der Kante liegt.)



Die beiden hellblau gefärbten Gebiete haben jeweils 3 Begrenzungskanten (also $F_3 = 2$), das gelbe Gebiet hat 4 Begrenzungskanten (also $F_4 = 1$) und das weiße Außengebiet hat 12 Begrenzungskanten, weil wir die beiden dunkelgrün gezeichneten Kanten jeweils doppelt zählen. Also ist $F_{12} = 1$. Für $i \notin \{3, 4, 12\}$ gilt $F_i = 0$.

Dann gilt für die Gesamtzahl F der Gebiete:

$$F = F_3 + F_4 + F_5 + \dots,$$

denn jedes Gebiet muss von mindestens 3 Kanten begrenzt werden.

Andererseits ist jede Kante eine »Begrenzungskante« für genau zwei Gebiete, es gilt also auch

$$2K = 3F_3 + 4F_4 + 5F_5 + \dots$$

Wenn wir das Dreifache der ersten Gleichung von der zweiten abziehen, dann erhalten wir

$$K - 3F = F_4 + 2F_5 + \dots \geq 0,$$

folglich unter Ausnutzung der Euler-Formel

$$2K \geq 3F = 3(2 + K - E) = 6 + 3K - 3E,$$

also $3E - 6 \geq K$, und das wollten wir zeigen.

zu (2). Angenommen, von jeder Ecke würden 6 oder mehr Kanten abgehen. Dann hätten wir

$$K \geq 6E / 2 = 3E > 3E - 6,$$

also einen Widerspruch zu Teil (1). □

KOROLLAR 13.20. Der Graph K_5 ist kein planarer Graph.

BEWEIS. Dies folgt aus Teil (1) des vorherigen Korollars. Der Graph K_5 hat 5 Ecken und 10 Kanten. Es gilt aber nicht $3 \cdot 5 - 6 \geq 10$. □

Außerdem können wir den »Sechsfarbensatz« beweisen, eine (wesentlich) abgeschwächte Version des Vierfarbensatzes:

KOROLLAR 13.21 (Sechsfarbensatz). Sei G ein planarer Graph. Dann lässt sich jeder Ecke eine von sechs Farben zuordnen, derart dass niemals zwei benachbarte Ecken dieselbe Farbe haben.

Dies ist natürlich eine wesentlich schwächere Aussage als der oben genannte Vierfarbensatz, aber weil der Beweis so einfach ist, soll er hier trotzdem skizziert werden. (Auch der analoge Fünffarbensatz ist wesentlich einfacher zu beweisen als der Vierfarbensatz.)

BEWEIS. Wir führen Induktion nach Anzahl der Ecken. Für alle Graphen mit 6 oder weniger Ecken ist die Aussage klar. Für den Induktionsschritt wählen wir eine Ecke, von der höchstens 5 Ecken abgehen. Eine solche existiert nach Korollar 13.19. Den Graphen G' , der aus G entsteht indem wir diese Ecke und alle Kanten, die von ihr ausgehen, entfernen, können

wir nach Induktionsvoraussetzung mit höchstens 6 Farben einfärben, so dass benachbarte Ecken stets unterschiedliche Farben zugewiesen bekommen.

Nun färben wir G , indem wir für alle Ecken in G' die Färbung übernehmen. Für die eine verbleibende Ecke können wir auch noch eine gültige Farbe finden, weil ja von dieser Ecke höchstens 5 Kanten ausgehen, also höchstens 5 der 6 Farben ausgeschlossen sind. $\square$

13.7. Weitere Literatur

R. Diestel, *Graphentheorie*, Springer Spektrum 2017

Im Buch [Ma] befinden sich mehrere Beispiele (über die oben genannten hinaus), in denen Methoden der linearen Algebra in der Graphentheorie zur Anwendung kommen.

Zusammenfassung *

Grundlagen über Mengen und Abbildungen, wie sie in Kapitel 3 dargestellt werden, setzen wir als bekannt voraus.

A.1. Gruppen und Körper

Referenz: Kapitel 8, Kapitel 4.

DEFINITION A.1. Eine *Gruppe* ist eine Menge G zusammen mit einer Verknüpfung $G \times G \rightarrow G$, die assoziativ ist und ein neutrales Element besitzt, und so dass jedes Element ein inverses Element hat. ⊢

Das neutrale Element ist eindeutig bestimmt, ebenso wie das inverse Element zu $g \in G$. Wird die Verknüpfung als Multiplikation geschrieben, so schreiben wir 1 für das neutrale Element von G und g^{-1} für das inverse Element von $g \in G$. Eine Gruppe, deren Verknüpfung kommutativ ist, heißt *kommutative* oder *abelsche* Gruppe. In diesem Fall wird die Verknüpfung oft als Addition $+$ geschrieben, und entsprechend das neutrale Element als 0 und das Inverse eines Elements g als $-g$. Man schreibt dann $g - h := g + (-h)$.

Eine nicht-leere Teilmenge $H \subseteq G$ einer Gruppe G heißt *Untergruppe*, wenn $hh' \in H$ und $h^{-1} \in H$ für alle $h, h' \in H$.

Ist $H \subseteq G$ eine Untergruppe, so ist H mit der Einschränkung der Gruppenverknüpfung selbst eine Gruppe. Der Durchschnitt von Untergruppen einer Gruppe ist eine Untergruppe.

Ist $M \subseteq G$ eine Teilmenge der Gruppe G , so ist der Durchschnitt aller Untergruppen von G , die M enthalten, die kleinste Untergruppe von G , die M enthält. Wir bezeichnen diese mit $\langle M \rangle$ und nennen sie die *von M erzeugte Untergruppe*.

DEFINITION A.2. Seien G, H (multiplikativ geschriebene) Gruppen. Ein *Gruppenhomomorphismus* ist eine Abbildung $f: G \rightarrow H$, so dass $f(gg') = f(g)f(g')$ für alle $g, g' \in G$. Ein *Isomorphismus* von Gruppen ist ein Gruppenhomomorphismus $f: G \rightarrow H$, so dass ein Gruppenhomomorphismus $g: H \rightarrow G$ existiert, derart dass f und g Umkehrabbildungen voneinander sind. ⊢

Die Verkettung von Gruppenhomomorphismen ist ein Gruppenhomomorphismus. Die Identitätsabbildung einer Gruppe G ist ein Gruppenhomomorphismus. Jeder bijektive Gruppenhomomorphismus ist ein Isomorphismus von Gruppen.

Ist $f: G \rightarrow H$ ein Gruppenhomomorphismus, so heißt $\text{Ker}(f) := f^{-1}(\{1\})$ der *Kern* und $\text{Im}(f)$ das *Bild* des Homomorphismus f . Es handelt sich dabei um Untergruppen von G bzw. H . Ein Homomorphismus ist genau dann injektiv, wenn sein Kern trivial (d.h. die triviale Gruppe $\{1\}$) ist.

BEISPIEL A.3. Sei $n \in \mathbb{N}$. Die Gruppe S_n der Bijektionen $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$ nennen wir die *symmetrische Gruppe*. Ihre Elemente heißen *Permutationen*. Eine Transposition ist eine Permutation σ , so dass $i \neq j$ existieren mit $\sigma(i) = j$, $\sigma(j) = i$ und $\sigma(k) = k$ für alle $k \neq i, j$. Es gibt einen eindeutig bestimmten Gruppenhomomorphismus $\text{sgn}: S_n \rightarrow \{1, -1\}$, der alle Transpositionen auf -1 abbildet. Wir nennen $\text{sgn}(\sigma)$ das *Signum* der Permutation σ . $\diamond$

DEFINITION A.4. Ein *Körper* ist eine Menge K mit Verknüpfungen $+: K \times K \rightarrow K$ (Addition) und $\cdot: K \times K \rightarrow K$ (Multiplikation), so dass $(K, +)$ eine kommutative Gruppe mit neutralem Element $0 \in K$ und $(K \setminus \{0\}, \cdot)$ eine kommutative Gruppe (mit neutralem Element 1) ist, und so dass für Addition und Multiplikation das Distributivgesetz gilt. $\dashv$

Ist K ein Körper, so schreiben wir $K^\times = K \setminus \{0\}$ für die *multiplikative Gruppe* von K .

Sei L ein Körper. Unter einem Teilkörper K von L verstehen wir eine Teilmenge $K \subseteq L$, so dass für die Inklusionsabbildung $\iota: K \rightarrow L$ und alle $x, y \in K$ gilt, dass $\iota(x + y) = \iota(x) + \iota(y)$, $\iota(xy) = \iota(x)\iota(y)$.

BEISPIEL A.5. (1) Die Körper der rationalen, reellen und komplexen Zahlen $\mathbb{Q} \subseteq \mathbb{R} \subseteq \mathbb{C}$ sind Teilkörper voneinander.

(2) Ist p eine Primzahl, so ist die Menge $\mathbb{Z}/p =: \mathbb{F}_p$ der Restklassen von $\mathbb{Z}$ modulo p mit der Addition und Multiplikation von Restklassen ein Körper. (Abschnitt 4.2.1) $\diamond$

A.2. Vektorräume

Referenz: Kapitel 6. Sei K ein Körper.

A.2.1. Definitionen.

DEFINITION A.6. Ein *Vektorraum* (über K) ist eine Menge V zusammen mit Verknüpfungen $+: V \times V \rightarrow V$ (Addition) und $\cdot: K \times V \rightarrow V$ (Skalarmultiplikation), so dass gilt:

- (a) $(V, +)$ ist eine kommutative Gruppe,
- (b) $(ab)v = a(bv)$ für alle $a, b \in K, v \in V$,
- (c) $1v = v$ für alle $v \in V$,
- (d) $(a + b)v = av + bv, a(v + w) = av + aw$ für alle $a, b \in K, v, w \in V$.

Die Elemente eines Vektorraums heißen *Vektoren*. $\dashv$

DEFINITION A.7. Sei V ein K -Vektorraum. Eine nicht-leere Teilmenge $U \subseteq V$ heißt *Untervektorraum* von V , wenn $u + u' \in U$ und $au \in U$ für alle $u, u' \in U$ und $a \in K$ gilt. $\dashv$

Ist $U \subseteq V$ ein Untervektorraum, so ist U mit der Einschränkung der Addition und Skalarmultiplikation selbst ein Vektorraum. Der Durchschnitt von Untervektorräumen eines Vektorraums ist ein Untervektorraum.

Ist $M \subseteq V$ eine Teilmenge, so ist der Durchschnitt aller Untervektorräume von V , die M enthalten, der kleinste Untervektorraum von V , der M enthält. Wir bezeichnen diesen mit $\langle M \rangle$ und nennen ihn den *von M erzeugten Untervektorraum*.

BEISPIEL A.8. (1) Der Vektorraum $\{0\}$ heißt der *Nullvektorraum*. Der Körper K ist mit der Addition und Multiplikation auf K ein K -Vektorraum.

- (2) Für Vektorräume $V_i, i \in I$ ist das kartesische Produkt $\prod_{i \in I} V_i$ mit der komponentenweisen Addition und Multiplikation ein K -Vektorraum, den wir als das *Produkt* der Vektorräume V_i bezeichnen. Die Teilmenge $\bigoplus_{i \in I} V_i \subseteq \prod_{i \in I} V_i$ der Elemente des Produkts, in denen nur endlich viele Einträge $\neq 0$ ist, ist ein Untervektorraum, die sogenannte *direkte Summe* der V_i . Wir schreiben $V^I := \prod_{i \in I} V, V^{(I)} := \bigoplus_{i \in I} V, V^n := \prod_{i=1}^n V$.
- (3) Für $n \in \mathbb{N}$ nennen wir den Vektorraum K^n den Standardvektorraum »der Dimension n « (siehe Definition A.13 unten). Allgemeiner können wir für eine Menge I den Vektorraum $K^{(I)}$ als Standardvektorraum bezeichnen.
- (4) Sind V, W Vektorräume über K , so ist die Menge $\text{Abb}(V, W)$ aller Abbildungen $V \rightarrow W$ mit der Addition und Skalarmultiplikation von Abbildungen:

$$f + g: v \mapsto f(v) + g(v), \quad af: v \mapsto af(v),$$

ein K -Vektorraum.

◇

DEFINITION A.9 (Definition 7.1). Seien V und W Vektorräume über dem Körper K . Eine *lineare Abbildung* (oder: ein *Vektorraum-Homomorphismus*) ist ein Homomorphismus $f: V \rightarrow W$ von additiven Gruppen, so dass zusätzlich für alle $v \in V, a \in K$ gilt: $f(av) = af(v)$.

Ein *Isomorphismus* von Vektorräumen ist ein Homomorphismus $f: V \rightarrow W$, so dass ein Homomorphismus $g: W \rightarrow V$ existiert, derart dass f und g Umkehrabbildungen voneinander sind. $\dashv$

Wir bezeichnen mit $\text{Hom}_K(V, W) \subset \text{Abb}(V, W)$ die Teilmenge aller linearen Abbildungen. Dies ist ein Untervektorraum. Wir schreiben auch $\text{End}_K(V) := \text{Hom}_K(V, V)$ und nennen die Elemente dieser Menge *Endomorphismen* von V . Ein *Automorphismus* von V ist ein Isomorphismus $V \rightarrow V$.

Jeder bijektive Vektorraumhomomorphismus ist ein Isomorphismus.

Ist $f: V \rightarrow W$ eine lineare Abbildung, so heißt $\text{Ker}(f) := f^{-1}(\{0\})$ der *Kern* und $\text{Im}(f)$ das *Bild* der linearen Abbildung f . Es handelt sich dabei um Untervektorräume von V bzw. W . Ein Homomorphismus ist genau dann injektiv, wenn sein Kern trivial (d.h. der Nullvektorraum) ist.

BEMERKUNG A.10. Ist V ein K -Vektorraum und $(v_i)_{i \in I}$ eine Familie von Vektoren, so erhalten wir eine lineare Abbildung

$$K^{(I)} \rightarrow V, \quad (a_i)_{i \in I} \mapsto \sum_i a_i v_i.$$

Das Bild dieser Abbildung ist $\langle \{v_i; i \in I\} \rangle$.

◇

A.2.2. Basen.

DEFINITION A.11. Sei V ein K -Vektorraum. Für eine Familie $(v_i)_{i \in I}$ von Vektoren aus V nennen wir die zugehörige Abbildung aus Bemerkung A.10 die natürliche Abbildung $K^{(I)} \rightarrow V$ (zu der gegebenen Familie von Vektoren).

- (1) Eine Familie $(v_i)_{i \in I}$ heißt *Erzeugendensystem* von V , wenn die folgenden äquivalenten Bedingungen erfüllt sind:
- (a) Jedes Element von V lässt sich als Linearkombination der Familie $(v_i)_i$ schreiben.
 - (b) Die natürliche Abbildung $K^{(I)} \rightarrow V$ ist surjektiv.

- (2) Eine Familie $(v_i)_{i \in I}$ heißt *linear unabhängig*, wenn die folgenden äquivalenten Bedingungen erfüllt sind:
- (a) Die einzige Möglichkeit, den Nullvektor $0 \in V$ als Linearkombination der Familie $(v_i)_i$ auszudrücken, ist die triviale Linearkombination.
 - (b) Jeder Vektor $v \in \langle \{v_i; i \in I\} \rangle$ lässt sich in eindeutiger Weise als Linearkombination der Familie $(v_i)_{i \in I}$ schreiben.
 - (c) Die natürliche Abbildung $K^{(I)} \rightarrow V$ ist injektiv.
- (3) Eine Familie $(v_i)_{i \in I}$ heißt eine *Basis von V* , wenn die folgenden äquivalenten Bedingungen erfüllt sind. (Siehe Satz 6.36.)
- (a) Jedes Element von V lässt sich in eindeutiger Weise als Linearkombination der Familie $(v_i)_i$ schreiben.
 - (b) Die Familie $(v_i)_i$ ist ein linear unabhängiges Erzeugendensystem.
 - (c) Die Familie $(v_i)_i$ ist ein minimales Erzeugendensystem.
 - (d) Die Familie $(v_i)_i$ ist ein maximales linear unabhängiges System von Vektoren.
 - (e) Die natürliche Abbildung $K^{(I)} \rightarrow V$ ist ein Isomorphismus.

⊥

Sei $V = K^{(I)}$ für eine Menge I . Für $i \in I$ sei $e_i \in K^{(I)}$ der Vektor, der an der Stelle i den Eintrag 1, und an allen anderen Stellen den Eintrag 0 hat. Dann bildet die Familie $(e_i)_{i \in I}$ eine Basis von V , die sogenannte *Standardbasis*.

Ist V ein Vektorraum mit Basis $\mathcal{B} = (v_1, \dots, v_n)$, so bezeichnen wir die Umkehrabbildung $c_{\mathcal{B}}: V \rightarrow K^n$ des Isomorphismus $K^n \rightarrow V, e_i \mapsto v_i$, als die Koordinatenabbildung bezüglich der Basis $\mathcal{B}$.

Wir nennen einen Vektorraum V *endlich erzeugt*, wenn ein Erzeugendensystem von V existiert, das aus nur endlich vielen Elementen besteht.

THEOREM A.12. (1) Sei V ein [endlich erzeugter] K -Vektorraum. Ist E ein Erzeugendensystem von V und M ein linear unabhängiges System von Vektoren aus V , dessen Elemente sämtlich in E enthalten sind, dann lässt sich M durch Hinzunahme von Elementen aus E zu einer Basis von V ergänzen.

- (2) Je zwei Basen eines [endlich erzeugten] Vektorraums haben dieselbe Mächtigkeit (d.h. haben gleich viele Elemente).

Dieses Theorem ist ein wichtiges Ergebnis der linearen Algebra, dessen Beweis nicht offensichtlich ist. Für unseren Ansatz spielt das Austauschprinzip eine wichtige Rolle, das in seiner einfachsten Form sagt, dass jeder Vektor $v \neq 0$ eines [endlich erzeugten] K -Vektorraums V durch Hinzufügen von Elementen aus einem vorgegebenen Erzeugendensystem zu einer Basis ergänzt werden kann. Teil (2) des Theorems erlaubt es uns, die folgende Definition zu machen:

DEFINITION A.13. Sei V ein [endlich erzeugter] Vektorraum. Die Anzahl $\dim(V)$ der Elemente in einer Basis nennen wir die *Dimension* von V .

⊥

Dann gilt: [Endlich erzeugte] K -Vektorräume V und W sind genau dann isomorph, wenn sie dieselbe Dimension haben.

SATZ A.14. Sei V ein K -Vektorraum mit Basis $(v_i)_{i \in I}$, sei W ein K -Vektorraum und seien $w_i \in W, i \in I$. Dann existiert genau eine lineare Abbildung $f: V \rightarrow W$ mit $f(v_i) = w_i$ für alle $i \in I$.

Mit anderen Worten: Sei V ein K -Vektorraum mit Basis $(v_i)_{i \in I}$. Dann ist für jeden K -Vektorraum die Abbildung

$$\operatorname{Hom}_K(V, W) \rightarrow W^I, \quad f \mapsto (f(v_i))_{i \in I},$$

ein Isomorphismus.

Es ist klar, dass die Abbildung f im ersten Teil des Satzes wegen der Linearität durch $f(\sum_i a_i v_i) = \sum_i a_i f(v_i)$ eindeutig bestimmt ist; man prüft leicht nach, dass die so definierte Abbildung tatsächlich linear ist.

SATZ A.15. Sei V ein endlich-dimensionaler K -Vektorraum.

- (1) Sei $U \subseteq V$ ein Untervektorraum. Dann ist U endlich erzeugt und $\dim(U) \leq \dim(V)$. Der Fall der Gleichheit tritt nur für $U = V$ ein.
- (2) Zu jedem Untervektorraum $U \subseteq V$ existiert ein Komplement, d.h. ein Untervektorraum $W \subseteq V$, so dass die Abbildung $U \oplus W \rightarrow V, (u, w) \mapsto u + w$, ein Isomorphismus ist.
- (3) Seien $U, W \subseteq V$ Untervektorräume. Dann gilt

$$\dim(U + W) + \dim(U \cap W) = \dim(U) + \dim(W).$$

Teil (1) folgt aus den obigen beiden Theoremen. Für Teil (2) wähle man eine Basis von U und setze diese zu einer Basis von V fort; die hinzugenommenen Basisvektoren erzeugen ein Komplement von U . Um Teil (3) zu beweisen, kann man ebenfalls den Basisergänzungssatz benutzen; oder die Dimensionsformel für lineare Abbildungen, siehe unten.

Für eine lineare Abbildung $f: V \rightarrow W$ nennen wir $\operatorname{rg}(f) := \dim(\operatorname{Im}(f))$ den *Rang* von f . Es gilt dann

THEOREM A.16 (Dimensionsformel für lineare Abbildungen, Theorem 7.23). Seien V, W Vektorräume über K , und sei V endlich-dimensional. Dann gilt

$$\dim(V) = \dim(\operatorname{Ker}(f)) + \operatorname{rg}(f).$$

Insbesondere ist ein Endomorphismus f eines endlich-dimensionalen K -Vektorraums genau dann injektiv, wenn er surjektiv ist (Korollar 7.24, vergleiche auch Satz A.27 weiter unten).

Sind $U, W \subseteq V$ Untervektorräume, so liefert die Dimensionsformel für lineare Abbildungen, angewandt auf die Abbildung $U \oplus W \rightarrow V, (u, w) \mapsto u - w$, einen Beweis der Dimensionsformel für Schnitt und Summe von Untervektorräumen aus Satz A.15 (2).

A.2.3. Der Dualraum eines Vektorraums. (Abschnitt 7.5) Sei V ein Vektorraum. Der Vektorraum $V^\vee := \operatorname{Hom}_K(V, K)$ heißt der *Dualraum* von V . Ist $f: V \rightarrow W$ eine lineare Abbildung, so heißt $f^\vee: W^\vee \rightarrow V^\vee, \lambda \mapsto \lambda \circ f$, die *duale Abbildung* von f ; dies ist ein Vektorraumhomomorphismus. Für eine Verkettung $f \circ g$ gilt $(f \circ g)^\vee = g^\vee \circ f^\vee$.

SATZ A.17. Sei $f: V \rightarrow W$ ein Homomorphismus.

- (1) Die Abbildung $f^\vee$ ist genau dann surjektiv, wenn f injektiv ist.
- (2) Die Abbildung $f^\vee$ ist genau dann injektiv, wenn f surjektiv ist.
- (3) Sind V und W endlich-dimensional, so gilt $\operatorname{rg}(f) = \operatorname{rg}(f^\vee)$.

Ist V ein K -Vektorraum mit Basis $v_1, \dots, v_n$, so bilden die Elemente $\lambda_1, \dots, \lambda_n \in V^\vee$ mit $\lambda_i(v_j) = 1$ für $i = j$, $\lambda_i(v_j) = 0$ für $i \neq j$, eine Basis von $V^\vee$, die sogenannte *duale Basis* zur Basis $v_1, \dots, v_n$. Insbesondere gilt: Ist V endlich-dimensional, dann ist $\dim(V) = \dim(V^\vee)$. Die Koordinatenabbildung für die duale Basis ist der Isomorphismus aus Satz A.14, angewandt auf $W = K$ und die gewählte Basis von V .

Die kanonische Abbildung $V \rightarrow V^{\vee\vee} := (V^\vee)^\vee$ von V in seinen Doppeldualraum, die gegeben ist durch $v \mapsto (V^\vee \rightarrow K, \lambda \mapsto \lambda(v))$, ist immer injektiv. Ist V endlich-dimensional, so ist sie ein Isomorphismus, weil V und $V^{\vee\vee}$ nach dem oben Gesagten dieselbe Dimension haben.

A.3. Matrizen

Referenz: Abschnitte 5.1.3, 5.3

A.3.1. Definition und das Produkt von Matrizen.

DEFINITION A.18. Eine *Matrix* über K der Größe $m \times n$ (mit $m, n \in \mathbb{N}$) ist eine Familie $(a_{ij})_{i,j}$ von Elementen aus K mit Indexmenge $\{1, \dots, m\} \times \{1, \dots, n\}$. Die Menge $M_{m \times n}(K)$ der Matrizen der Größe $m \times n$ ist mit der eintragsweisen Addition und Skalarmultiplikation ein K -Vektorraum. $\dashv$

Eine *Diagonalmatrix* ist eine Matrix $A = (a_{ij})_{i,j} \in M_n(K)$, so dass $a_{ij} = 0$ für alle $i \neq j$. Für $a_1, \dots, a_n \in K$ bezeichnen wir mit $\text{diag}(a_1, \dots, a_n)$ die Diagonalmatrix $(a_{ij})_{i,j}$ mit $a_{ii} = a_i$ für alle i . Für $A = (a_{ij})_{i,j} \in M_{m \times n}(K)$ nennen wir $A^t := (a_{ji})_{i,j} \in M_{n \times m}(K)$ die zu A *transponierte Matrix*.

DEFINITION A.19 (Matrizenprodukt). Sind $A \in M_{l \times m}(K)$, $B \in M_{m \times n}(K)$, so definieren wir das Produkt $AB = (c_{ij})_{i,j} \in M_{l \times n}(K)$ der Matrizen A und B durch

$$c_{ij} = \sum_{k=1}^m a_{ik} b_{kj}.$$

$\dashv$

Für das Matrizenprodukt gelten das Assoziativ- und das Distributivgesetz. Die Diagonalmatrix $E_n := \text{diag}(1, \dots, 1) \in M_n(K)$ heißt die *Einheitsmatrix*. Es gilt $E_m A = A$, $A E_n = A$ für alle $A \in M_{m \times n}(K)$. Wir schreiben $M_n(K) := M_{n \times n}(K)$.

Eine Matrix $A \in M_n(K)$ heißt *invertierbar*, wenn eine Matrix $B \in M_n(K)$ mit $AB = BA = E_n$ existiert. Die *allgemeine lineare Gruppe* ist die Gruppe $GL_n(K)$ der invertierbaren Matrizen in $M_n(K)$ bezüglich der Multiplikation von Matrizen.

BEISPIEL A.20 (Permutationsmatrizen, siehe Abschnitt 8.1.5). Für eine Permutation $\sigma \in S_n$ sei $P_\sigma \in M_n(K)$ die Matrix, deren j -te Spalte der Vektor $e_{\sigma(j)}$ ist. Die Abbildung $S_n \rightarrow GL_n(K)$, $\sigma \mapsto P_\sigma$, ist ein injektiver Gruppenhomomorphismus. Die Matrizen P_σ bezeichnen wir als *Permutationsmatrizen*. Es gilt $P_\sigma^{-1} = P_{\sigma^{-1}} = (P_\sigma)^t$. $\diamond$

A.3.2. Matrizen und lineare Abbildungen. (Abschnitt 7.3)

Für eine Matrix $A \in M_{m \times n}(K)$ ist die Abbildung $f_A: K^n \rightarrow K^m, x \mapsto Ax$, eine lineare Abbildung. Wir schreiben auch $\text{Ker}(A)$ statt $\text{Ker}(f_A)$ und $\text{Im}(A)$ statt $\text{Im}(f_A)$.

SATZ A.21. Die Abbildung $M_{m \times n}(K) \rightarrow \text{Hom}_K(K^n, K^m), A \mapsto f_A$, ist ein Isomorphismus von K -Vektorräumen. Wir bezeichnen mit $f \mapsto M(f)$ seine Umkehrabbildung. Dann ist die j -te Spalte von $M(f)$ der Vektor $f(e_j)$.

Für endlich-dimensionale Vektorräume V und W mit Basen $\mathcal{B}$ und $\mathcal{C}$ induzieren die Koordinatenisomorphismen $c_{\mathcal{B}}: V \rightarrow K^n$ und $c_{\mathcal{C}}: W \rightarrow K^m$ (mit $n := \dim(V)$, $m := \dim(W)$) einen Isomorphismus $\text{Hom}_K(V, W) \rightarrow \text{Hom}_K(K^n, K^m), f \mapsto c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1}$. Durch Kombination mit dem vorherigen Satz bekommen wir:

SATZ A.22. Seien V und W endlich-dimensionale K -Vektorräume mit Basen $\mathcal{B}$ und $\mathcal{C}$. Die Abbildung $\text{Hom}_K(V, W) \rightarrow M_{m \times n}(K), f \mapsto M_{\mathcal{C}}^{\mathcal{B}}(f) := M(c_{\mathcal{C}} \circ f \circ c_{\mathcal{B}}^{-1})$ ist ein Isomorphismus.

Schreiben wir $\mathcal{B} = (v_1, \dots, v_n)$, so ist die j -te Spalte von $M_{\mathcal{C}}^{\mathcal{B}}(f)$ der Vektor $c_{\mathcal{C}}(f(v_j))$, der aus den Koeffizienten besteht, mit denen $f(v_j)$ als Linearkombination der Basis $\mathcal{C}$ geschrieben wird.

Für die Verkettung von Homomorphismen $f: U \rightarrow V$ und $g: V \rightarrow W$ zwischen endlich-dimensionalen Vektorräumen mit Basen $\mathcal{B}, \mathcal{C}, \mathcal{D}$ erhalten wir

$$M_{\mathcal{D}}^{\mathcal{B}}(g \circ f) = M_{\mathcal{D}}^{\mathcal{C}}(g) M_{\mathcal{C}}^{\mathcal{B}}(f).$$

Sind $\mathcal{B}, \mathcal{B}'$ Basen desselben Vektorraums V , so nennen wir $M_{\mathcal{B}}^{\mathcal{B}'} := M_{\mathcal{B}}^{\mathcal{B}'}(\text{id}_V)$ die Basiswechselmatrix zwischen den Basen $\mathcal{B}'$ und $\mathcal{B}$. Sind $f: V \rightarrow W$ ein Homomorphismus und $\mathcal{B}, \mathcal{B}'$ Basen von V und $\mathcal{C}, \mathcal{C}'$ Basen von W , so erhalten wir aus der obigen Formel die Basiswechselformel (Korollar 7.33)

$$M_{\mathcal{C}'}^{\mathcal{B}'}(f) = M_{\mathcal{C}'}^{\mathcal{C}} M_{\mathcal{C}}^{\mathcal{B}}(f) M_{\mathcal{B}}^{\mathcal{B}'}.$$

Ein Homomorphismus f ist genau dann ein Isomorphismus, wenn die Matrix $M_{\mathcal{C}}^{\mathcal{B}}(f)$ invertierbar ist. Insbesondere ist die Basiswechselmatrix $M_{\mathcal{B}}^{\mathcal{B}'}$ invertierbar; ihre inverse Matrix ist $M_{\mathcal{B}'}^{\mathcal{B}}$.

SATZ A.23 (Die darstellende Matrix der dualen Abbildung, Satz 7.54). *Sind V und W endlich-dimensionale Abbildungen mit Basen $\mathcal{B}$ und $\mathcal{C}$, und sind $\mathcal{B}^\vee$ bzw. $\mathcal{C}^\vee$ die dualen Basen von $\mathcal{B}$ und $\mathcal{C}$, so gilt*

$$M_{\mathcal{B}^\vee}^{\mathcal{C}^\vee}(f^\vee) = M_{\mathcal{C}}^{\mathcal{B}}(f)^t.$$

Für eine Matrix $A \in M_{m \times n}(K)$ heißt $\text{rg}(A) := \dim(\text{Im}(f_A))$ der Rang der Matrix A .

THEOREM A.24 (Zeilenrang = Spaltenrang, Theorem 7.41). *Sei $A \in M_{m \times n}(K)$. Dann gilt $\text{rg}(A) = \text{rg}(A^t)$.*

SATZ A.25 (Smith-Normalform, Satz 7.36). (1) *Sei $f: V \rightarrow W$ ein Homomorphismus zwischen endlich-dimensionalen K -Vektorräumen. Dann existieren Basen $\mathcal{B}$ von V und $\mathcal{C}$ von W , so dass*

$$M_{\mathcal{C}}^{\mathcal{B}}(f) = \begin{pmatrix} E_r & 0 \\ 0 & 0 \end{pmatrix} \text{ (als Blockmatrix geschrieben). Hier ist } r = \text{rg}(f).$$

(2) *Sei $A \in M_{m \times n}(K)$. Dann existieren invertierbare Matrizen $S \in M_m(K)$ und $T \in M_n(K)$, so dass*

$$SAT = \begin{pmatrix} E_r & 0 \\ 0 & 0 \end{pmatrix} \text{ (als Blockmatrix geschrieben). Hier ist } r = \text{rg}(A).$$

A.3.3. Lineare Gleichungssysteme. Eine lineare Gleichung in n Unbestimmten über K ist eine Gleichung der Form

$$a_1 X_1 + \dots + a_n X_n = b$$

mit $a_i, b \in K$. Sind m Gleichungen dieser Form gegeben, so sprechen wir von einem *linearen Gleichungssystem*. Wir schreiben die Koeffizienten a_{ij} dieser Gleichungen zusammen mit den Werten b_i auf der rechten Seiten der Gleichungen in die *erweiterte Koeffizientenmatrix* $(A \mid b)$ mit $A \in M_{m \times n}(K)$, $b \in K^m = M_{m \times 1}(K)$. Ist $b = 0$, so sprechen wir auch von einem *homogenen* linearen Gleichungssystem. Die Lösungsmenge des linearen Gleichungssystems ist $f_A^{-1}(b) = \{x \in K^n; Ax = b\}$. Ist diese nicht leer, d.h. ist $b \in \text{Im}(A)$, und ist t ein Element dieser Lösungsmenge, so gilt $f_A^{-1}(b) = t + \text{Ker}(A) := \{t + v; v \in \text{Ker}(A)\}$. Hier ist $\text{Ker}(A) = \{x \in K^n; Ax = 0\}$ die Lösungsmenge des zugehörigen *homogenen linearen Gleichungssystems*.

THEOREM A.26 (Gauß-Algorithmus). (1) *Sei $m \in \mathbb{N}$. Die Matrizen $E_{ij}(a)$, $1 \leq i \neq j \leq m$, $a \in K$ und $\text{diag}(a_1, \dots, a_m)$, $a_i \in K^\times$, erzeugen die Gruppe $GL_m(K)$.*

(2) *Ist $A \in M_{m \times n}$, so existieren $S \in GL_m(K)$ sowie eine Permutationsmatrix $T \in GL_n(K)$, so dass SAT die Form*

$$\begin{pmatrix} E_r & A' \\ 0 & 0 \end{pmatrix} \text{ (als Blockmatrix geschrieben)}$$

hat. Dabei ist r der Rang von A und $A' \in M_{r \times (n-r)}(K)$ eine Matrix, die durch A bis auf Vertauschung der Spalten eindeutig bestimmt ist.

- (3) Mit der Notation aus Teil (2) ist die Lösungsmenge $\text{Ker}(A)$ des durch A gegebenen homogenen linearen Gleichungssystems der von den Spalten der Matrix $T \begin{pmatrix} -A' \\ E_{n-r} \end{pmatrix}$ erzeugte Untervektorraum von K^n .

A.3.4. Invertierbare Matrizen.

SATZ A.27. Sei $A \in M_n(K)$. Dann sind äquivalent:

- (i) A ist invertierbar.
- (ii) Es existiert $B \in M_n(K)$ mit $AB = E_n$.
- (iii) Es existiert $B \in M_n(K)$ mit $BA = E_n$.
- (iv) Der Homomorphismus f_A ist ein Isomorphismus.
- (v) Der Homomorphismus f_A ist surjektiv.
- (vi) Der Homomorphismus f_A ist injektiv.
- (vii) Für jedes $b \in K^n$ ist das lineare Gleichungssystem $Ax = b$ eindeutig lösbar.
- (viii) Für jedes $b \in K^n$ ist das lineare Gleichungssystem $Ax = b$ lösbar.
- (ix) Das homogene lineare Gleichungssystem $Ax = 0$ ist eindeutig lösbar.
- (x) Die reduzierte Zeilenstufenform der Matrix A ist E_n . (Definition 5.13)
- (xi) Es gilt $\det(A) \neq 0$ (siehe den Abschnitt A.4.2 weiter unten).

A.4. Determinante

Referenz: Kapitel 9.

A.4.1. Definition und Existenz. Seien K ein Körper und V ein K -Vektorraum.

- DEFINITION A.28. (1) Wir nennen eine Abbildung $f: V^n \rightarrow W$ in einen K -Vektorraum W *multilinear*, wenn für alle i und $v_1, \dots, v_{i-1}, v_{i+1}, \dots, v_n \in V$ die Abbildung $V \rightarrow W$, $v \mapsto f(v_1, \dots, v_{i-1}, v, v_{i+1}, \dots, v_n)$ linear ist.
- (2) Wir nennen eine multilineare Abbildung $f: V^n \rightarrow W$ in einen K -Vektorraum W *alternierend*, wenn für alle $v_1, \dots, v_n \in V$, so dass $i \neq j$ mit $v_i = v_j$ existieren, gilt, dass $f(v_1, \dots, v_n) = 0$.
- (3) Sei nun V endlich-dimensional. Eine alternierende multilineare Abbildung $V^n \rightarrow K$ heißt *Determinantenfunktion*. Die Menge der Determinantenfunktionen bildet einen Untervektorraum von $\text{Abb}(V^n, K)$, den wir mit $\mathcal{D}_V$ bezeichnen.

⊥

Die Funktion $\det: M_n(K) \rightarrow K$,

$$A = (a_{ij})_{i,j} \mapsto \sum_{\sigma \in S_n} \text{sgn}(\sigma) \prod_{i=1}^n a_{i,\sigma(i)} \quad (\text{Leibniz-Formel})$$

ist eine Determinantenfunktion auf $M_n(K) = (K^n)^n$. Wir nennen $\det(A) \in K$ die *Determinante* der Matrix A . Für $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in M_2(K)$ ist $\det(A) = ad - bc$.

SATZ A.29. Der Vektorraum $\mathcal{D}_V$ hat Dimension 1.

A.4.2. Eigenschaften der Determinante.

SATZ A.30. Seien $A, B \in M_n(K)$. Dann gilt $\det(AB) = \det(A) \det(B)$.

Für $A \in M_n(K)$ vom Rang $< n$ gilt $\det(A) = 0$. Für invertierbare Matrizen ist die Determinante $\neq 0$, genauer gilt:

THEOREM A.31. Die Determinante induziert einen Gruppenhomomorphismus $GL_n(K) \rightarrow K^\times$. Sein Kern ist die spezielle lineare Gruppe $SL_n(K)$, das ist die Untergruppe von $GL_n(K)$, die erzeugt wird von den Matrizen der Form $E_{ij}(a)$, $a \in K$.

Es gilt $\det(A) = \det(A^t)$. Für eine Permutation σ mit zugehöriger Permutationsmatrix P_σ ist $\det(P_\sigma) = \operatorname{sgn}(\sigma)$ (wobei hier $\operatorname{sgn}(\sigma)$ als Element des Grundkörpers K aufzufassen ist).

DEFINITION A.32 (Determinante eines Endomorphismus). Sei V ein endlich-dimensionaler K -Vektorraum und $f \in \operatorname{End}_K(V)$. Ist $\mathcal{B}$ eine Basis von V , so nennen wir $\det(f) := \det(M_{\mathcal{B}}^{\mathcal{B}}(f)) \in K$ die Determinante des Endomorphismus f . Diese ist unabhängig von der Wahl von $\mathcal{B}$. $\dashv$

SATZ A.33. Für eine quadratische Matrix $A \in M_n(K)$ bezeichnen wir mit $A'_{ij} \in M_{n-1}(K)$ die Matrix, die aus A durch Streichen der i -ten Zeile und j -ten Spalte hervorgeht. Die Komplementärmatrix von A ist $A^{ad} \in M_n(K)$ mit Eintrag $(-1)^{i+j} \det(A'_{ji})$ in Zeile i und Spalte j .

(1) Laplacescher Entwicklungssatz, Satz 9.29. Für alle i gilt die »Entwicklung von $\det A$ nach der i -ten Zeile«:

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A'_{ij}$$

und für alle j die »Entwicklung von $\det A$ nach der j -ten Spalte«:

$$\det A = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det A'_{ij}.$$

(2) Cramersche Regel, Satz 9.32.

$$AA^{ad} = A^{ad}A = \det(A)E_n.$$

A.4.3. Die Spur einer Matrix. Für eine quadratische Matrix $A = (a_{ij})_{i,j} \in M_n(K)$ definieren wir die *Spur* $\operatorname{Spur}(A) \in K$ von A als die Summe $\sum_{i=1}^n a_{ii}$ der Diagonaleinträge von A . Wir erhalten einen Gruppenhomomorphismus $M_n(K) \rightarrow K$ von additiven Gruppen. Es gilt $\operatorname{Spur}(AB) = \operatorname{Spur}(BA)$ für alle $A, B \in M_n(K)$ und insbesondere $\operatorname{Spur}(SAS^{-1}) = \operatorname{Spur}(A)$ für alle $A \in M_n(K)$, $S \in GL_n(K)$.

A.5. Eigenwerte

DEFINITION A.34. Sei V ein K -Vektorraum und $f \in \operatorname{End}_K(V)$. Ein *Eigenvektor* von f zum *Eigenwert* $\lambda \in K$ ist ein Element $v \in V \setminus \{0\}$, so dass $f(v) = \lambda v$. Ein Element $\lambda \in K$ wird als Eigenwert von f bezeichnet, wenn ein Eigenvektor zum Eigenwert λ existiert.

Ist $\lambda \in K$ ein Eigenwert von f , so nennen wir

$$V_\lambda(f) = \{v \in V; f(v) = \lambda v\}$$

den zugehörigen *Eigenraum*. $\dashv$

Ist $A \in M_n(K)$ eine Matrix, so nennt man die Eigenwerte und Eigenvektoren der Abbildung f_A auch die Eigenwerte bzw. Eigenvektoren von A . Für jede Basis $\mathcal{B}$ von V stimmen dann die Eigenwerte von $f \in \text{End}_K(V)$ und von $M_{\mathcal{B}}^{\mathcal{B}}(f)$ überein (aber nicht notwendigerweise die jeweiligen Eigenvektoren, die ja sogar in unterschiedlichen Vektorräumen liegen, wenn V nicht der Standardvektorraum K^n ist).

SATZ A.35. Seien V ein K -Vektorraum und $f \in \text{End}_K(V)$. Sind $v_1, \dots, v_m \in V$ Eigenvektoren von f zu paarweise verschiedenen Eigenwerten, so ist die Familie $v_1, \dots, v_m$ linear unabhängig.

DEFINITION A.36. (1) Sei V ein endlich-dimensionaler K -Vektorraum und f ein Endomorphismus von V . Wir nennen f *diagonalisierbar*, wenn die folgenden äquivalenten Bedingungen erfüllt sind:

- (a) Es existiert eine Basis $\mathcal{B}$ von V , so dass $M_{\mathcal{B}}^{\mathcal{B}}(f)$ eine Diagonalmatrix ist.
 - (b) Es existiert eine Basis von V , die aus Eigenvektoren von f besteht.
 - (c) Der Vektorraum V ist die (direkte) Summe der Eigenräume von f .
- (2) Eine Matrix $A \in M_n(K)$ heißt *diagonalisierbar*, wenn die folgenden äquivalenten Bedingungen erfüllt sind:
- (a) Es existiert eine invertierbare Matrix S , so dass SAS^{-1} eine Diagonalmatrix ist.
 - (b) Der Endomorphismus f_A von K^n ist diagonalisierbar im Sinn von Teil (1).

—

Mathematische Ergänzungen *

B.1. Das Lemma von Zorn

Wir haben in Lemma 3.61 gesehen, dass endliche Produkte von nicht-leeren Mengen stets nicht leer sind – eine Aussage, die sehr einleuchtend ist. Interessanterweise hat die entsprechende Aussage für unendliche Produkte eine besondere Stellung unter den Axiomen der Mengenlehre:

AXIOM B.1 (Auswahlaxiom). *Seien I eine Menge und $X_i, i \in I$, nicht-leere Mengen. Dann ist das Produkt $\prod_{i \in I} X_i$ nicht leer.*

Dieses Axiom wurde erstmals 1904 von Zermelo formuliert, gehört aber nicht zu dem Axiomensystem von Zermelo und Fraenkel, das heutzutage mit dem Kürzel ZF bezeichnet wird (diese Axiome wurden nach Vorarbeiten von Zermelo 1921 von Fraenkel formuliert). Es wurde 1938 von Gödel bewiesen, dass das Auswahlaxiom nicht im Widerspruch zu den Axiomen ZF steht, und erst 1963 konnte Cohen zeigen, dass es andererseits auch nicht aus diesen Axiomen folgt.

Das Axiomensystem »ZF + Auswahlaxiom« bezeichnet man mit ZFC (»C« für »axiom of choice«), und es ist das in überwiegenden Teilen der heutigen Mathematik gängige Axiomensystem für die Mengenlehre und damit die Grundlage der allermeisten mathematischen Theorien. Dass das Auswahlaxiom eine Sonderstellung hat, liegt daran, dass man einerseits einen großen Teil der Mathematik auch ohne Auswahlaxiom aus dem Axiomensystem ZF herleiten kann, und andererseits, dass das Auswahlaxiom auch einige Konsequenzen hat, die der Intuition widersprechen.

Ein bekanntes Beispiel dafür ist das [Banach-Tarski-Paradoxon](#)¹. Dieses hängt eng damit zusammen, dass es *nicht* möglich ist, jeder beschränkten Teilmenge von $\mathbb{R}^3$ ein »Volumen«, also eine Zahl in $\mathbb{R}_{\geq 0}$ zuzuordnen, so dass das Volumen nicht konstant 0, additiv für endliche disjunkte Vereinigungen und invariant unter Verschiebungen, Drehungen und Spiegelungen ist.

[Video](#)^a zum Banach-Tarski-Paradoxon von Vsauce.

^a <https://www.youtube.com/watch?v=s86-Z-CbaHA>

”

The Axiom of Choice is obviously true, the well-ordering principle obviously false, and who can tell about Zorn's lemma?

Jerry Bona

Fundort: <https://mathoverflow.net/a/7194>

¹ <https://de.wikipedia.org/wiki/Banach-Tarski-Paradoxon>

Auch wenn das Auswahlaxiom einige Herausforderungen an unsere Intuition stellt, ist die oben formulierte Aussage doch so überzeugend, dass es üblicherweise (und auch in diesem Skript) in die Liste der verwendeten Axiome aufgenommen wird. Wie erwähnt, ist bewiesen worden, dass es nicht im Widerspruch zu den Axiomen in ZF steht, so dass es aus mathematischer Sicht unkritisch ist, dieses Axiom zu verwenden. Und es ist oft praktisch, weil es so mächtig ist. In der linearen Algebra erlaubt es uns zu beweisen, dass jeder Vektorraum eine Basis besitzt (siehe Ergänzung 6.48).

Man kann zeigen, dass das Auswahlaxiom zu der Aussage, dass jeder Vektorraum eine Basis besitzt, äquivalent ist.

Leicht zu sehen ist, dass das Auswahlaxiom äquivalent ist zu

SATZ B.2. *Sei $f: X \rightarrow Y$ eine surjektive Abbildung. Dann existiert ein Schnitt von f , d.h. eine Abbildung $g: Y \rightarrow X$ mit $f \circ g = \text{id}_Y$.*

BEWEIS. Die Existenz von g entspricht der Existenz eines Elements von $\prod_{y \in Y} f^{-1}(\{y\})$. Wegen der Surjektivität von f sind die Mengen $f^{-1}(\{y\})$ alle nicht-leer. $\square$

Um aus dem Satz das Auswahlaxiom abzuleiten, betrachtet man für eine Familie X_i die surjektive Abbildung $\coprod_{i \in I} X_i \rightarrow I$, wobei $\coprod_{i \in I} X_i$ die »disjunkte Vereinigung« der X_i ist. Dies ist eine Menge, die alle X_i als Teilmengen enthält, die gleich der Vereinigung aller dieser Teilmengen X_i ist, und die $X_i \cap X_j = \emptyset$ für alle $i \neq j$ erfüllt. (Formal definiert man die disjunkte Vereinigung der X_i , indem man X_i identifiziert mit $\{i\} \times X_i$, $i \in I$, und dann die Vereinigung aller dieser Mengen bildet. Durch das Hinzufügen der Komponente i erreicht man, dass die Kopien der verschiedenen X_i sich nicht schneiden. Hier wird eine Vereinigung von Mengen gebildet, die nicht Teilmenge einer vorgegebenen Menge sind. Diese Konstruktion ist im Rahmen von ZF durchführbar.)

Eine andere bekannte Aussage, die äquivalent ist zum Auswahlaxiom, ist der

SATZ B.3 (Wohlordnungssatz). *Sei X eine Menge. Dann existiert eine Wohlordnung auf X , d.h. eine totale Ordnung $\preceq$, derart dass jede nicht-leere Teilmenge von X ein kleinstes Element bezüglich $\preceq$ hat.*

Siehe Abschnitt 3.14.3 für die hier verwendeten Begriffe. So plausibel das Auswahlaxiom ist, so schwierig ist es, sich beispielsweise eine Wohlordnung der Menge der reellen Zahlen vorzustellen.

Um zu beweisen, dass jeder Vektorraum eine Basis hat, benötigen wir eine andere zum Auswahlaxiom äquivalente Aussage, und zwar das sogenannte Lemma von Zorn.

Dafür benötigen wir zuerst einige Sprechweisen: Zunächst benutzen wir die Begriffe der *partiellen* und der *totalen Ordnung*, wie sie in Abschnitt 3.14.3 eingeführt wurden. Ein besonders passendes Beispiel in unserem Kontext ist die Teilmengenbeziehung: Ist M eine Menge, so bezeichnen wir mit $P(M)$ die Potenzmenge von M , also die Menge aller Teilmengen von M . Diese ist mittels der Inklusion $\subseteq$ partiell geordnet, das bedeutet: Für Elemente $X, Y, Z \in P(M)$ (also Teilmengen $X, Y \subseteq M$) gilt

- (1) $X \subseteq X$,
- (2) aus $X \subseteq Y$ und $Y \subseteq Z$ folgt $X \subseteq Z$,
- (3) aus $X \subseteq Y$ und $Y \subseteq X$ folgt $X = Y$.

Bei einer totalen Ordnung müssen zusätzlich je zwei Elemente vergleichbar sein (wie bei der üblichen Anordnung der ganzen Zahlen: Für $x, y \in \mathbb{Z}$ gilt $x \leq y$ oder $y \leq x$). Bei der Inklusion von Teilmengen ist das offensichtlich nicht der Fall: In vielen Situationen gilt weder $X \subseteq Y$ noch $Y \subseteq X$. Man spricht daher von einer *partiellen* Ordnung. Jede Teilmenge $U \subseteq P(M)$, also

eine Menge von Teilmengen von M , der nicht alle Teilmengen angehören müssen, ist dann auch durch die Inklusion $\subseteq$ partiell geordnet. Während die folgende Diskussion für $P(M)$ mit der partiellen Ordnung $\subseteq$ nicht so interessant ist, genügt es, wenn Sie sich alles im Fall einer Teilmenge $U \subseteq P(M)$ der Potenzmenge irgendeiner Menge M vorstellen (mit $\subseteq$ als partieller Ordnung).

Sei nun U eine Menge mit einer partiellen Ordnung $\preceq$. Wir hatten in Definition 3.80 definiert, dass ein Element $x \in U$ *maximal* (bezüglich $\preceq$) heißt, wenn für alle $y \in U$ aus $x \preceq y$ schon $x = y$ folgt.

BEISPIEL B.4. (1) Sei M eine Menge und $U = P(M)$ mit $\subseteq$ als partieller Ordnung. Dann ist M ein maximales Element, und zwar das einzige.

(2) Sei M eine nicht-leere Menge und $U \subset P(M)$ die Teilmenge von $P(M)$, die aus allen echten Teilmengen $X \subsetneq M$ von M besteht. Dann ist für jedes $m \in M$ das Element $M \setminus \{m\}$ ein maximales Element von U bezüglich $\subseteq$.

(3) Seien K ein Körper und V ein K -Vektorraum. Sei $U \subset P(V)$ die Menge aller linear unabhängigen Teilmengen von V . Wir haben in Satz 6.36 gesehen, dass die maximalen Elemente von U genau die Basen von V sind. Das Lemma von Zorn wird uns ein Kriterium für die Existenz eines maximalen Elements in irgendeiner partiell geordneten Menge zur Verfügung stellen. Angewandt auf das Beispiel hier erhalten wir damit das Ergebnis, dass jeder Vektorraum eine Basis besitzt. Siehe Ergänzung 6.48.

◇

DEFINITION B.5. Sei U eine durch $\preceq$ partiell geordnete Menge. Sei U' eine Teilmenge von U . Ein Element $x \in U$ heißt *obere Schranke* von U' , wenn für alle $y \in U'$ gilt, dass $y \preceq x$. $\dashv$

In der vorherigen Definition ist es wichtig, dass wir nicht verlangen, dass die obere Schranke x ein Element von U' ist. Im Zornschen Lemma wird die Existenz oberer Schranken für gewisse Teilmengen von U verlangt. (Eine obere Schranke von U selbst wäre gerade ein größtes Element; dessen Existenz wollen wir sicherlich nicht zur Voraussetzung machen.) Dafür bemerken wir, dass jede Teilmenge $U' \subseteq U$ einer partiell geordneten Menge selbst wieder partiell geordnet ist. Je nachdem kann es natürlich passieren, dass die partielle Ordnung von U auf der Teilmenge U' sogar eine *totale* Ordnung induziert; in diesem Fällen fordert man im Lemma von Zorn die Existenz einer oberen Schranke.

SATZ B.6 (Lemma von Zorn). Sei U eine nicht-leere Menge mit einer partiellen Ordnung $\preceq$, die die folgende Eigenschaft hat: Jede Teilmenge $U' \subseteq U$, die durch $\preceq$ total geordnet ist, besitzt eine obere Schranke in U .

Dann existiert in U ein maximales Element bezüglich $\preceq$.

Formal gesehen kann man darauf verzichten, die Voraussetzung $U \neq \emptyset$ in die Aussage des Lemmas aufzunehmen, weil die folgende Bedingung, angewandt auf die total geordnete Teilmenge $U' := \emptyset$, garantiert, dass U ein Element enthält, nämlich eine obere Schranke der leeren Teilmenge.

Man nennt total geordnete Teilmengen einer partiell geordneten Menge manchmal auch *Ketten*, und nennt eine partiell geordnete Menge *induktiv geordnet*, wenn jede Kette eine obere Schranke besitzt. Dann kann man das Lemma von Zorn formulieren als: Jede induktiv geordnete Menge besitzt ein maximales Element.

Wir verzichten darauf, an dieser Stelle zu beweisen, dass das Lemma von Zorn aus dem Auswahlaxiom folgt. Siehe [So] I.9 für einen gut lesbaren Beweis. Man kann auch zeigen, dass das Auswahlaxiom aus dem Zornschen Lemma folgt (das ist eher noch einfacher zu

beweisen als die andere Implikation). Die beiden Aussagen sind also (unter den Axiomen in ZF) äquivalent. Sie könnten daher auch einfach das Zornsche Lemma als eines der Axiome betrachten, die der Mengenlehre zugrunde liegen sollen.

BEMERKUNG B.7. Wenn in der obigen Situation eine Kette U' nur endlich viele Elemente hat, so können wir sie uns in der Form

$$u_0 \preceq \cdots \preceq u_n$$

vorstellen. Dann ist natürlich u_n eine obere Schranke. Sobald es unendlich viele Elemente in U' gibt, wird die Sache offenbar komplizierter. Es genügt nicht, den Fall von Ketten der Form

$$u_0 \preceq u_1 \preceq u_1 \preceq \dots,$$

die durch $\mathbb{N}$ indiziert sind, zu betrachten! Um das Zornsche Lemma anwenden zu können, muss man zeigen, dass *jede* Kette eine obere Schranke besitzt, und eine Kette muss nicht abzählbar sein.

Sonst könnte man zum Beispiel zeigen, dass die Menge der reellen Zahlen eine maximale abzählbare Teilmenge besitzt. Das ist offenbar nicht möglich, da eine abzählbare Teilmenge von $\mathbb{R}$ eine echte Teilmenge sein muss, und sie nach Hinzufügen eines weiteren Elements abzählbar bleibt. $\diamond$

Verschiedenes *

C.1. Mathematische Sprechweisen, Code words

C.1.1. Begriffe. In der mathematischen Fachsprache gibt es eine ganze Menge von Wörtern, die man kennen muss, um einen mathematischen Text zu lesen, die aber nicht ein mathematisches Objekt im engeren Sinne bezeichnen und daher nicht Bestandteil einer Definition sind. Einige davon werden in diesem Abschnitt erklärt. Wenn Sie andere Wörter sehen, deren Bedeutung Sie nicht verstehen, fragen Sie nach! Einige der Wörter in der folgenden Liste haben eine mathematische Bedeutung (und werden also an geeigneter Stelle definiert). Sie befinden sich hier, weil sie auch Bestandteil der Alltagssprache sind, so dass man die präzise mathematische Bedeutung womöglich leicht vergisst (zum Beispiel: es ist ein wichtiger Unterschied, ob ein Element das *kleinste* Element einer partiell geordneten Menge, oder ein *minimales* Element ist).

Absolutterm, absolutes Glied – In einem polynomiellen Ausdruck $a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$ heißt a_0 (also der Summand, der nicht mit einer positiven Potenz von x multipliziert wird) der Absolutterm oder das absolute Glied.

a fortiori – Dies ist Latein und bedeutet *erst recht*. Wenn man schon eine stärkere Aussage bewiesen hat, aber nur ein schwächeres Ergebnis im weiteren Verlauf verwenden möchte, kann man diese Sprechweise benutzen. (»Wir haben gesehen, dass $n > 5$ ist, a fortiori ist n positiv.«)

a priori – Dies ist Latein und bedeutet *von vorneherein*. Der Ausdruck wird (auch) in mathematischen Texten benutzt, um etwas zu kennzeichnen, was man erwarten könnte, aber dann doch nicht eintritt. (»A priori könnte n negativ sein, aber weil ... gilt, kann das nicht sein.«)

Annahme, angenommen – Auch wenn das nicht immer so strikt gesehen wird, finde ich es sinnvoll, die Wörter *Annahme* und *angenommen* (nur) dann zu verwenden, wenn man in einem Widerspruchsbeweis (siehe Abschnitt 3.7.3) die Annahme angibt, die dann zum Widerspruch geführt werden soll. Die Tatsachen, die in einem Satz vorausgesetzt werden, sollte man auch so nennen: *Voraussetzungen*.

im allgemeinen/in der Regel – Diese Ergänzung schiebt die Mathematiker*in gerne ein, wenn eine Aussage zwar so gut wie immer wahr ist, es aber ein paar wenige Ausnahmefälle gibt: Ist n eine natürliche Zahl, dann gibt es im allgemeinen mehr als eine Bijektion $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$. (Aber für die Ausnahmefälle $n = 0$ und $n = 1$ gibt es eben nur eine einzige. Die obige Formulierung ist klarer, als zu sagen, dass es mehr als eine Bijektion $\{1, \dots, n\} \rightarrow \{1, \dots, n\}$ geben kann; denn das ließe im Dunklen, dass es nur in ganz wenigen Fällen nicht so ist. Manchmal ist es natürlich noch besser, die Ausnahmen genau anzugeben.)

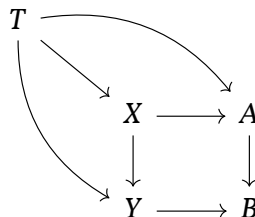
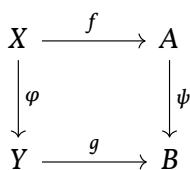
Implikation – Dies ist ein anderes Wort für Folgerung, oder genauer für den Sachverhalt, dass eine Aussage aus einer anderen folgt (»Beweisen Sie die Implikation $A \Rightarrow B$.«).

Inklusion – Die Eigenschaft, eine Teilmenge zu sein. (»Die Inklusion $X \subseteq Y$ ist leicht zu zeigen.«)

kleinstes, größtes, minimales, maximales Element bezüglich einer (partiellen) Ordnung – siehe Abschnitt 3.14.3.

kanonisch – Man bezeichnet ein mathematisches Objekt (oft eine Abbildung) als kanonisch, wenn es sich um die »offensichtliche« Wahl handelt. Dieser Begriff hat also keine mathematische Definition. Man könnte ihn auch umschreiben mit »so wie wir es immer machen«. Zum Beispiel nennt man die Standardbasis von K^n manchmal die *kanonische Basis*. Ähnlich gelagert ist die Sache mit dem Begriff *natürlich*, wobei man hiermit je nach Kontext auch eine mathematische Bedeutung verbindet (die Eigenschaft der Funktorialität, die wir hier nicht genauer erläutern).

Kommutatives Diagramm – Eine Abbildung zwischen Mengen X und Y schreibt man üblicherweise mit einem Pfeil: $X \rightarrow Y$. Hat man mehrere Abbildungen, dann ist es oft übersichtlicher, alle Abbildungen zusammen mit den entsprechenden Pfeilen darzustellen, auch wenn das mehr Platz benötigt:



wobei man die Namen der Abbildungen wie im linken Beispiel an die Pfeile schreiben kann. Eine solche Darstellung nennt man ein *Diagramm* (egal, wie viele Abbildungen darin kombiniert werden). Ein Diagramm heißt *kommutativ*, wenn alle Möglichkeiten, eine Abbildung von einer Menge zu irgendeiner anderen Menge in dem Diagramm als Verkettung von Abbildungen des Diagramms zu definieren, dieselbe Abbildung liefern. Im linken Beispiel heißt das also einfach, dass $\psi \circ f = g \circ \varphi$ gelten muss. Im rechten Beispiel sind die beiden Verkettungen $X \rightarrow B$ und die Verkettungen $T \rightarrow A$ und $T \rightarrow Y$ zu vergleichen. Es folgt dann, dass auch die verschiedenen Abbildungen $T \rightarrow B$ alle übereinstimmen.

Diese Sprechweise kann man ebenso für Vektorräume und Homomorphismen von Vektorräumen, oder für Gruppen und Homomorphismen von Gruppen usw. verwenden.

Lemma – Ein Lemma ist eine mathematische Aussage, die untergeordnete Bedeutung hat, aber doch im Text ausdrücklich (und in der Regel mit einer Nummer versehen) festgehalten werden soll, damit man später darauf verweisen kann. Für den großen Überblick braucht man sich die Lemmata (so der Plural dieses griechischstämmigen Worts; Lemmas ist aber auch gebräuchlich) eher nicht zu merken. Statt von einem Lemma spricht man auch manchmal von einem Hilfssatz. Demgegenüber ist ein Satz (oder eine Proposition) und erst recht ein Theorem (manchmal auch: Lehrsatz) ein Ergebnis von großer (oder sogar überragender) Bedeutung. Ein Korollar ist ein Ergebnis, dass sich ohne große zusätzliche Arbeit aus einem bereits bewiesenen Satz (oder Theorem ...) ableiten lässt.

loc.cit. – Dies ist die Abkürzung für den lateinischen Ausdruck *loco citato*, am angegebenen Orte, und bedeutet, dass sich die Literaturangabe auf dieselbe Quelle bezieht, die davor angegeben wurde.

Ohne Einschränkung / ohne Beschränkung der Allgemeinheit – (Abkürzung: OE beziehungsweise OBdA) Diese Floskel benutzt man um zu begründen, dass es (in einem Beweis) erlaubt ist, sich auf einen speziellen Fall zurückzuziehen. Zum Beispiel, weil die ausgeschlossenen Fälle ohnehin leicht zu erledigen sind, oder weil eine symmetrische Situation vorliegt.

paarweise – Wenn man über eine Eigenschaft spricht, die eine Menge (oder eine Familie) haben kann, meint man manchmal, dass die Eigenschaft für je zwei Elemente gelten soll, und ergänzt deshalb das Wort *paarweise*. Zum Beispiel sind die Zahlen 3, 5, 18 teilerfremd, denn es gibt keine natürliche Zahl > 1 , die alle drei Zahlen teilt. Aber 3 und 18 sind nicht

teilerfremd, und deshalb sind 3, 5, 18 nicht paarweise teilerfremd. In ähnlicher Weise ist es ein Unterschied, ob Teilmengen einer Menge disjunkt (der Durchschnitt ist leer) oder paarweise disjunkt (der Durchschnitt von je zweien ist leer) sind. Oder ob eine Familie von Vektoren linear unabhängig, oder paarweise linear unabhängig ist. (Wenn man sagt, dass Zahlen/Objekte $x_1, \dots, x_n$ verschieden sein sollen, dann meint man aber in der Regel, dass sie paarweise verschieden sind, ohne das explizit dazuzusagen.)

Schubfachprinzip – Unter dem Schubfachprinzip (das manchmal noch pompöser als das Dirichletsche Schubfachprinzip bezeichnet wird) versteht man die folgende banale (aber manchmal sehr nützliche) Feststellung: Wenn man n Dinge in weniger als n Fächer legt, dann sind in mindestens einem Fach mindestens zwei Dinge. (Oder mathematischer ausgedrückt: Ist $f: X \rightarrow Y$ eine Abbildung zwischen endlichen Mengen und ist $\#X > \#Y$, dann ist f nicht injektiv. Siehe Lemma 3.63.)

trivial – Das Wort *trivial* wird in zwei Kontexten verwendet. Einerseits sagt man gelegentlich, ein Beweis (oder ein Ergebnis) sei trivial, wenn der Beweis sehr einfach ist. (Man sollte sich allerdings zwingen, mindestens ein zweites Mal nachzudenken, bevor man das schreibt, denn viele Fehler in mathematischen Texten verstecken sich gerade an den Stellen, wo etwas als trivial oder offensichtlich bezeichnet wird. Dann war es eben nicht als *Das ist einfach*. zu verstehen, sondern eher als *Das sollte kein Problem sein, aber ich bin zu faul, es mir genau zu überlegen und erst recht es aufzuschreiben*, und da kommt es dann leicht zu Fehlern.)

Zweitens spricht man manchmal von trivialen Objekten oder Abbildungen, meist solche, die besonders uninteressant sind: Die Gruppe $\{1\}$ mit einem einzigen Element wird als die triviale Gruppe bezeichnet. Andere Beispiele: Die triviale Lösung eines homogenen linearen Gleichungssystems, die triviale Linearkombination einer Familie von Vektoren.

wohldefiniert – Man sagt, ein Begriff (oder ein mathematisches Objekt; sehr oft eine Abbildung) sei *wohldefiniert*, wenn die Definition zwar a priori von gewissen zusätzlichen Wahlen abhängt, diese aber letztlich keine Rolle spielen. Zum Beispiel: Wir definieren die Summe von zwei Bruchzahlen $\frac{a}{b}$ und $\frac{c}{d}$ als $\frac{ad+cb}{bd}$. Um zu überprüfen, dass dies eine *wohldefinierte* Abbildung $+: \mathbb{Q} \times \mathbb{Q} \rightarrow \mathbb{Q}$ liefert, müssen wir sicherstellen, dass das Ergebnis nicht davon abhängt, wie wir den Bruch $\frac{a}{b}$ (und entsprechend $\frac{c}{d}$) schreiben. Da man Brüche kürzen und erweitern kann, gilt ja $\frac{a}{b} = \frac{a'}{b'}$ für viele a', b' , und die obige Vorschrift ergibt nur dann Sinn, wenn dann auch $\frac{ad+cb}{bd} = \frac{a'd+cb'}{b'd}$ gilt (was in der Tat der Fall ist). Die Vorschrift $(\frac{a}{b}, \frac{c}{d}) \mapsto \frac{a+c}{b+d}$ ist dahingegen nicht wohldefiniert – wenn man einen der Brüche $\frac{a}{b}, \frac{c}{d}$ erweitert, erhält man in der Regel ein ganz anderes Ergebnis. Die Frage nach Wohldefiniertheit tritt immer dann auf, wenn sich eine Definition auf eine Äquivalenzklasse (Abschnitt 3.14.2) bezieht, aber für die Definition ein Repräsentant der Äquivalenzklasse herangezogen wird. Ein anderes Beispiel ist Definition 3.58, in der wir die Mächtigkeit einer endlichen Menge X definiert haben; dort wählt man eine Bijektion $X \rightarrow \{1, \dots, n\}$ und muss wissen, dass die Zahl n , die dabei auftritt, eindeutig bestimmt ist (auch wenn es natürlich in der Regel viele solche Bijektionen gibt). Sie sollten den Begriff wohldefiniert *nicht* verwenden, um auf andere Punkte hinzuweisen, die für die Korrektheit einer Definition erforderlich sind. (Beispiel: Wollen Sie eine Abbildung $f: X \rightarrow Y$ durch eine Formel für $f(x)$ angeben, so ist es erforderlich, dass $f(x)$ für alle x ein Element von Y ist; dass das so ist, beschreibt man aber nicht mit dem Wort wohldefiniert.)

C.1.2. Abkürzungen. Vor allem an der Tafel benutzt man manchmal die folgenden Abkürzungen, um etwas Zeit und Platz zu sparen:

- gdw. – genau dann, wenn (auf Englisch: iff, if and only if)
- IA, IS, IV – Induktionsanfang, Induktionsschritt, Induktionsvoraussetzung.
- Kp. – Körper

- l. u., l. a. – linear unabhängig, linear abhängig
- OE, OBdA (s.o.), auf Englisch: WLOG – without loss of generality
- TFAE (Englisch) – the following are equivalent, die folgenden Aussagen sind äquivalent
- VR – Vektorraum
- zz, gzz – zu zeigen, genügt zu zeigen (auf Englisch: ETS – enough to show)

C.2. Mathematische Texte am Computer schreiben.

Zu Beginn Ihres Studiums besteht keine Notwendigkeit, mathematische Texte am Computer zu schreiben. Was die Hausaufgaben betrifft, ist meine Empfehlung, diese handschriftlich anzufertigen.

Irgendwann, spätestens, wenn Sie Ihre Bachelor-Arbeit schreiben, wird sich diese Anforderung aber ergeben.

Das Standardprogramm für mathematischen Textsatz ist $\text{\LaTeX}$ (gesprochen Latech, das X am Ende ist der große griechische Buchstabe Chi; auf Englisch wird er üblicherweise als K gesprochen). Mathematik mit einem herkömmlichen Textverarbeitungsprogramm wie LibreOffice oder Word zu schreiben, ist einerseits viel mehr Arbeit, andererseits sieht das Ergebnis deutlich schlechter aus (speziell die mathematischen Formeln, aber auch bei gewöhnlichem Text sind die Algorithmen von $\text{\LaTeX}$ in der Regel besser, um optimale Zeilenumbrüche zu finden).

Das zugrundeliegende $\text{\TeX}$ und $\text{\LaTeX}$ wurden von D. Knuth und L. Lamport Ende der 1970'er, Anfang der 1980'er entwickelt und haben das Verfassen und Veröffentlichen von mathematischen Texten revolutioniert. Die Programme werden seitdem behutsam weiterentwickelt. Sie werden von praktisch allen Mathematikern benutzt, um ihre Ergebnisse in Artikeln und Büchern festzuhalten (und auch in vielen anderen Fachgebieten). Dieses Skript wurde natürlich auch mit $\text{\LaTeX}$ geschrieben, genauer mit LuaLaTeX. Das ist eine Variante, die unter anderem besser mit verschiedenen Schriftarten umgehen kann und modernere Schnittstellen zur Programmierung bietet. Eine andere Variante ist XeLaTeX.

Daher sollten Sie bei Gelegenheit, zum Beispiel in den Semesterferien, wenn Sie etwas Zeit haben, $\text{\LaTeX}$ lernen. Die $\text{\LaTeX}$ -Umgebung lässt sich auf allen gängigen Plattformen (Windows, Mac OS, Linux, ...) installieren. Es gibt im Netz dazu jede Menge Anleitungen, und es gibt auch entsprechende Lehrbücher, wenn Sie es gerne etwas ausführlicher hätten.

Etwas mehr Informationen und weitere Referenzen: [Wikipedia^a](https://de.wikipedia.org/wiki/LaTeX).

Eine Suche nach »Latex Einführung« liefert auch jede Menge brauchbarer Treffer.

Ein Buch über $\text{\LaTeX}$ benötigen Sie erstmal wahrscheinlich nicht.

Mein Tipp wäre: Suchen Sie sich jemanden, der $\text{\LaTeX}$ auf einem Rechner mit demselben Betriebssystem benutzt, wie Sie es haben, und lassen Sie sich die ersten Schritte direkt am Computer erklären. Danach können Sie vieles mit »learning by doing« machen.

^a<https://de.wikipedia.org/wiki/LaTeX>

Einige wichtige Grundregeln des mathematischen Schriftsatzes, die nicht offensichtlich sind:

- Es sollte nie ein mathematisches Symbol direkt am Satzanfang stehen.
- Es sollten nie zwei Symbole direkt aufeinander folgen (mindestens ein Komma oder anderes Satzzeichen muss dazwischen sein, um Missverständnisse zu vermeiden).

- Operatoren sind nicht dasselbe wie eine Folge von Variablen: $\det(A)$ und nicht $\det(A)$. In $\text{\LaTeX}$ muss man also `\det(A)` schreiben, und für Namen von Operatoren, die noch nicht in $\text{\LaTeX}$ definiert sind, ein entsprechendes Makro selber schreiben (keine Angst, das ist nur eine Zeile, zum Beispiel `\newcommand\Spur{\mathop{\rm Spur}}`).

Wesentlich ausführlicher hat das K. Conrad als [Advice on mathematical writing](https://kconrad.math.uconn.edu/blurbs/proofs/writingtips.pdf)¹ aufgeschrieben, zwar auf Englisch, aber das meiste lässt sich direkt auf deutsche Texte übertragen.

C.2.1. Mathematik im Forum auf der Moodle-Seite. Auch auf der Moodle-Seite können Sie die gängigen mathematischen Notationen verwenden. Auch dort wird die $\text{\LaTeX}$ -Notation verwendet. Mathematische Ausdrücke müssen in *doppelte Dollar-Zeichen* eingeschlossen werden, um erkannt zu werden. Zum Beispiel:

<code>\$\$a^n\$\$</code>	a^n
<code>\$\$U\oplus W\$\$</code>	$U \oplus W$
<code>\$\$m\times n\$\$</code>	$m \times n$
<code>\$\$A=(a_{ij})_{i,j}\$\$</code>	$A = (a_{ij})_{i,j}$
<code>\$\$\sum_{i=1}^n i = \frac{(n+1)n}{2}\$\$</code>	$\sum_{i=1}^n i = \frac{(n+1)n}{2}$
<code>\$\$A=\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}\$\$</code>	$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}$
<code>\$\$e_2=\begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}\$\$</code>	$e_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$
<code>\$\$\bigoplus_{i \in I} V_i \subseteq \prod_{i \in I} V_i\$\$</code>	$\bigoplus_{i \in I} V_i \subseteq \prod_{i \in I} V_i$

Schließt man Text in `\(` (und `\)` ein, so wird der dazwischen angegebene Text als $\text{\LaTeX}$ interpretiert und das Ergebnis als abgesetzte Zeile dargestellt.

¹<https://kconrad.math.uconn.edu/blurbs/proofs/writingtips.pdf>

Bemerkungen zur Literatur *

Vorbemerkungen:

- Sie brauchen neben dem Material, das Ihnen im Rahmen der Vorlesung zur Verfügung gestellt wird, nicht unbedingt weitere Quellen. Die unten genannten Verweise dienen nur der Ergänzung – für den Fall, dass Sie das Material, das wir in der Vorlesung behandeln, noch einmal in etwas anderer Form anschauen möchten.
- Auch wenn die Liste unten schon recht lang ist, ist sie bei weitem nicht vollständig. Es gibt sehr viele Lehrbücher und Skripte zur Linearen Algebra, und Sie können und sollten ruhig auch selbst einmal schauen, was Sie sonst noch an Material finden – sei es im Internet, sei es in der Bibliothek. Da der Stoff der Linearen Algebra sehr weitgehend standardisiert ist, ist das Risiko klein, dass Sie woanders »etwas Falsches« lernen.
- Als Anknüpfung an den letzten Punkt, aber als eigener Punkt, weil er so wichtig ist und sich auf alle Quellen bezieht, egal, ob Sie sie selbst gefunden haben oder ob sie unten erwähnt werden (und genauso gilt der Punkt für dieses Skript): Da in der Mathematik ohnehin alle Aussagen zu beweisen sind, können (und **müssen**) Sie selbst überprüfen, ob das, was Sie lesen, überhaupt richtig ist. *Glauben Sie niemandem!*

Mit den Definitionen verhält es sich natürlich ein bisschen anders, denn eine Definition kann ja nicht falsch sein (höchstens kann sie ungeschickt oder nutzlos sein). Hier muss man tatsächlich manchmal ein bisschen aufpassen. Was den Begriff des Vektorraums oder des Körpers angeht, gibt es wohl keine Unterschiede, aber in manchen Büchern werden nur Vektorräume über den reellen und komplexen Zahlen betrachtet. Die Null wird manchmal nicht als Element der Menge der natürlichen Zahlen betrachtet.

Viele der Bücher können Sie als Studierende der Universität Duisburg-Essen kostenfrei herunterladen. Dazu müssen Sie aus dem Adressbereich der Universität auf die Verlagsseiten zugreifen; entweder von einem Rechner aus, der sich an der Uni befindet (oder über Eduroam das WLAN der Uni nutzt), oder über eine [VPN-Verbindung](#)¹.

Alle Quellen, zu denen unten ein Link angegeben ist, sollten entweder allgemein zugänglich, oder aus dem Netz der Universität Duisburg-Essen zugänglich sein.

D.1. Lehrbücher

(Alphabetisch nach den Namen der Autoren.)

C. Baer, *Lineare Algebra und Analytische Geometrie*, Springer Spektrum 2018,
<https://doi.org/10.1007/978-3-658-22620-6>

S. Bosch, *Lineare Algebra*, Springer 2014,
<https://doi.org/10.1007/978-3-642-55260-1>

E. Brieskorn, *Lineare Algebra und Analytische Geometrie I*, Vieweg+Teubner 1983

¹ <https://www.uni-due.de/zim/services/internetzugang/vpn>

Das Buch von Brieskorn ist sehr ausführlich und enthält auch Abschnitte über die Geschichte der linearen Algebra (und analytischen Geometrie) und Bemerkungen zur Motivation vieler Begriffe. (Es gibt auch noch einen zweiten und dritten Band.)

G. Fischer, *Lineare Algebra*, Springer Spektrum 2014,
<https://doi.org/10.1007/978-3-658-03945-5>

Vielleicht *das* Standard-Lehrbuch der linearen Algebra im deutschsprachigen Raum. Der »Fischer« wurde auch schon empfohlen, als ich im ersten Semester war.

G. Fischer, *Lernbuch Lineare Algebra und Analytische Geometrie*, Springer 2019,
<https://doi.org/10.1007/978-3-658-27343-9>

K. Jänich, *Lineare Algebra*, Springer 2008, <https://doi.org/10.1007/978-3-662-08382-6>
Mit ausführlichen Bemerkungen zu den Querverbindungen zur Physik.

W. Klingenberg, *Lineare Algebra und Geometrie*, Springer, 3. Aufl., 1992.

Neben dem Standardstoff der linearen Algebra erhält auch die (analytische) Geometrie ausführliche Beachtung.

H.-J. Kowalsky, *Lineare Algebra*, de Gruyter 1963.

Manchmal lohnt auch der Blick in ein älteres Buch, und sei es nur um zu sehen, dass sich der Stoff der Grundvorlesung Lineare Algebra in den vergangenen 50 Jahren nur wenig geändert hat. (Allerdings hoffe ich, dass aus den Ergänzungen im Skript hervorgeht, dass genau diese lineare Algebra auch die Grundlage für viele aktuelle(re) Ergebnisse und Anwendungen ist.) Zum Beispiel in dieses, von dem es danach noch mehrere neue Auflagen gab, die dann auch die Basis für das Buch von Kowalsky und Michler zur linearen Algebra waren.

J. Liesen, V. Mehrmann, *Lineare Algebra*, Vieweg+Teubner 2012,
<https://doi.org/10.1007/978-3-8348-8290-5>

Der Untertitel »Ein Lehrbuch über die Theorie mit Blick auf die Praxis« wird unter anderem anhand mehrerer konkreter Fragestellungen umgesetzt, die im ersten Kapitel erklärt und dann in den späteren Kapiteln wieder aufgegriffen werden. Enthält auch viele Übungsaufgaben.

F. Lorenz, *Lineare Algebra I*, Spektrum Akad. Verlag 2004.

S. Waldmann, *Lineare Algebra I*, Springer Spektrum 2017,
<https://doi.org/10.1007/978-3-662-49914-6>,
Lineare Algebra 2, Springer Spektrum 2017,
<https://doi.org/10.1007/978-3-662-53348-2>

Dieses Buch enthält viele Bezüge zur Physik, besonders auch im zweiten Band, und außerdem viele schöne Übungsaufgaben.

D.2. Vorlesungsskripte zur Linearen Algebra

C. Löh, *Lineare Algebra I*, Wintersemester 2016/17, Univ. Regensburg.

http://www.mathematik.uni-regensburg.de/loeh/teaching/linalg1_ws1617/lecture_notes.pdf

Ein aufwändig gemachtes Skript mit vielen Randbemerkungen, sowohl zu »angrenzenden« mathematischen Themen als auch zu (inner- und außermathematischen) Anwendungen der linearen Algebra. Im Anhang sind die Übungsaufgaben der damaligen Vorlesung enthalten, darunter viele schöne »Nicht-Standard-Aufgaben«.

R. Pink, *Lineare Algebra I und II*, Zusammenfassung, 2016, ETH Zürich.

<https://people.math.ethz.ch/~pink/ftp/Lineare-Algebra-Zusammenfassung-20161006.pdf>

Der Stoff der Vorlesungen Lineare Algebra I und 2 auf circa 100 Seiten (allerdings größtenteils ohne Beweise), dazu ungefähr 20 Seiten an lesenswerten »Vorbemerkungen«.

W. Soergel, *Lineare Algebra I*, Univ. Freiburg

<http://home.mathematik.uni-freiburg.de/soergel/Skripten/XXLA1.pdf>

Ein ausführliches Skript zur Linearen Algebra I mit vielen Bemerkungen, die den Zugang zum Stoff erleichtern sollen.

J. Stix, *Lineare Algebra*, Univ. Frankfurt

https://www.uni-frankfurt.de/74414804/Stix_LineareAlgebra_Skript.pdf

Dieses Skript ist in einem relativ knappen Stil geschrieben, enthält aber eine Menge an Material (das wir teilweise erst in der Linearen Algebra 2 behandeln werden).

A. Werner, *Lineare Algebra I*, Univ. Frankfurt

Basiskurs: https://www.uni-frankfurt.de/50581195/basiskurs_LA1.pdf,

Aufbaukurs: https://www.uni-frankfurt.de/50581204/Aufbaukurs_LA1.pdf

Das Skript von A. Werner ist etwas anders aufgebaut als die meisten hier genannten Skripte und Bücher: Es ist unterteilt in einen Basiskurs und einen Aufbaukurs. Im Basiskurs wird ein »rechnerischer« Zugang gewählt, und der theoretische Aufbau der Theorie hintangestellt bzw. teilweise in den Aufbaukurs verlagert.

D.3. Bücher übers Problemlösen und Beweisen

C. Ableitinger, A. Herrmann, *Lernen aus Musterlösungen zur Analysis und Linearen Algebra*, Vieweg+Teubner 2011, <https://doi.org/10.1007/978-3-8348-8335-3>

Ausführlich diskutierte Musterlösungen zu Aufgaben zu den Vorlesungen Analysis I/2 und Lineare Algebra I/2.

D. Grieser, *Mathematisches Problemlösen und Beweisen*, Springer Spektrum 2017,

<https://doi.org/10.1007/978-3-658-14765-5>

Der Untertitel *Eine Entdeckungsreise in die Mathematik* beschreibt das Buch gut: Anhand von gut zugänglichen Problemstellungen werden verschiedene Beweismethoden und Strategien, um mathematische Probleme zu lösen, vorgestellt. Das Buch richtet sich explizit an Studienanfänger*innen und ist mit Schulmathematikkenntnissen (und Neugier auf Mathematik) lesbar.

C. Alsina, R. Nelsen, *Bezaubernde Beweise*, Springer Spektrum 2013
 C. Alsina, R. Nelsen, *Perlen der Mathematik*, Springer Spektrum 2015,
<https://doi.org/10.1007/978-3-662-45461-9>

Zwei Bücher derselben Autoren, die jeweils eine große und reich illustrierte Sammlung von mathematischen Problemen, Bemerkungen und (teilweise) Kuriositäten enthalten, deren Lösung überwiegend mit einfachen Mitteln möglich ist. Leider ist nur das Neuere im E-Book-Paket enthalten, das Uni-Mitgliedern der Universität Duisburg-Essen zur Verfügung steht.

In seiner Rezension über das Buch *Bezaubernde Beweise* schreibt der Rezensent M. Drmota: »Es ist ein schönes Buch zum Schmökern und zur mathematischen Entspannung und kann allen mathematisch Interessierten wärmstens empfohlen werden.« (Math. Semesterberichte **60** (2013), 274–275.)

K. Houston, *Wie man mathematisch denkt*, Springer Spektrum 2012

Eine Einführung in das »Beweisen« und allgemeiner die mathematische Arbeitstechnik an der Universität.

D.4. Englische Lehrbücher zur linearen Algebra

Schauen Sie ruhig auch einmal in ein Buch, das auf Englisch geschrieben ist, und sei es nur um zu sehen, dass »mathematisches Englisch« in aller Regel sehr leicht zu verstehen ist. Später im Studium werden Sie mit hoher Wahrscheinlichkeit Quellen auf Englisch lesen müssen.

D. Austin, *Understanding Linear Algebra*,
<http://merganser.math.gvsu.edu/david/linear.algebra/ula/ula/ula.html>

Ein Skript/Buch zur linearen Algebra mit vielen Anwendungsbeispielen und eingebetteten Berechnungen/Berechnungsaufgaben mit dem Computeralgebra-System *Sage*.

J. Hefferon, *Linear Algebra*, <http://joshua.smcvt.edu/linearalgebra/>

Dieses Buch enthält viele *Topics* (Ausblicke, Ergänzungen) zu Anwendungen und Verbindungen mit anderen Themen.

S. Axler, *Linear Algebra Done Right*, Springer 2015,
<https://doi.org/10.1007/978-3-319-11080-6>

S. Treil, *Linear Algebra Done Wrong*,
<https://www.math.brown.edu/~treil/papers/LADW/LADW.html>.

Zum Schluss noch zwei »fortgeschrittenere« Bücher, in denen die lineare Algebra als eines von mehreren Themen behandelt wird.

T. Hungerford, *Algebra*, Springer Graduate Texts in Math. **73**, 1974.

Ein Standard-Lehrbuch im US-amerikanischen Raum, das allerdings wesentlich mehr Stoff als die lineare Algebra enthält (die Algebra und ein Teil der kommutativen Algebra sind auch dabei, und noch ein paar andere Themen). Dafür finden Sie hier das meiste, was wir in der Vorlesung machen, kurz und knapp zusammengefasst.

E. Vinberg, *A Course in Algebra*, Graduate Studies in Math. **56**, AMS 2003.

Dieses Buch enthält ebenfalls wesentlich mehr Stoff als die lineare Algebra. Es ist vielleicht insgesamt etwas einfacher zu lesen als das Buch von Hungerford.

D.5. Verschiedene Buchempfehlungen

D.5.1. Was ist Mathematik? Diese Frage ist natürlich nicht so leicht zu beantworten. Was ist Ihre Antwort? Die folgenden Bücher versuchen in sehr unterschiedlicher Weise, eine Antwort zu geben.

T. Gowers, *Mathematik*, Reclam-Sachbuch 2011 (oder das englische Original: *Mathematics. A very short introduction*, Oxford Univ. Press 2002).

Eine kurze Abhandlung zur Fragestellung *Was ist Mathematik* (aber vielleicht doch nicht *very short*, auch wenn es dem Autor Timothy Gowers, der 1998 für seine Forschung mit der Fields-Medaille ausgezeichnet wurde, womöglich so vorkam). Gut lesbar, auch vor dem Beginn des Mathematikstudiums, und mit einem Preis von aktuell 5,60 € mehr als preiswert. Besondere Empfehlungen: Die Kapitel 3 (Beweise), 5 (Dimension), 6 (Geometrie).

T. Gowers, *The Princeton Companion to Mathematics*, Princeton Univ. Press 2008.

Tausend Seiten über Mathematik, von einem Kapitel über die Geschichte der Mathematik über mathematische Konzepte, Theoreme und Vermutungen, Kurzbiographien von Mathematikern zum Einfluss der Mathematik in Biologie, Chemie, Informatik, den Wirtschaftswissenschaften sowie Kunst und Musik. Wenn Sie anfangen, in diesem Buch zu blättern, finden Sie sicher etwas, was Sie interessiert.

Behrends, Fünf Minuten Mathematik <https://doi.org/10.1007/978-3-8348-9529-5>

100 Beiträge aus der Mathematik-Kolumne, die E. Behrends für die Tageszeitung *Die Welt* geschrieben hat. Die *5 Minuten* sind durchaus wörtlich zu nehmen, so dass die einzelnen Texte nicht in die Tiefe gehen. Dafür lassen sie sich aber eben mal schnell zwischendurch lesen, geben einen Eindruck von der Vielseitigkeit der Mathematik und machen oft Appetit auf mehr.

R. Courant, H. Robbins, *Was ist Mathematik*, Springer 1992

Ein Klassiker (der inzwischen vielleicht ein bisschen angestaubt ist), der schon deshalb nicht in dieser Liste fehlen darf, weil der Abschnitt mit dem Titel dieses Buchs überschrieben ist. Das Buch richtet sich (auch) an Studienanfänger*innen (und teils an Schüler*innen), ist aber nicht immer leicht zu lesen.

D.5.2. Dies und das. Verschiedene Buchtipps, die nicht unbedingt etwas mit linearer Algebra zu tun haben – für den Fall, dass Sie sich langweilen oder noch etwas suchen, was Sie sich zu Weihnachten wünschen können (und dabei einen intellektuellen Eindruck hinterlassen).

T. Körner, *The pleasures of counting*, Cambridge Univ. Press 1996 (Es gibt auch eine deutsche Übersetzung: *Mathematisches Denken – vom Vergnügen am Umgang mit Zahlen*, Birkhäuser 1998)

Ein dickes Buch (über 500 Seiten), in dem alle möglichen Anwendungen von Mathematik beschrieben werden, beginnend mit der Bekämpfung der Cholera um 1850, bei der man durch eine simple Statistik feststellen konnte, dass die Krankheit über das Trinkwasser übertragen wird – damals ein Durchbruch, über Verschlüsselungstheorie bis zu verschiedenen Anwendungen von Methoden der Analysis.

T. Körner, *Calculus for the ambitious*, Cambridge Univ. Press 2014.

Auch wenn dieses Buch gar nichts mit linearer Algebra zu tun hat, nehme ich es in die Liste auf, weil es mir so gut gefallen hat. Eine Einführung in den Stoff der Vorlesung Analysis I, in der der formale Anspruch, wie er in einer Vorlesung (aus guten Gründen) besteht, zurückgestellt wird und die Theorie von einem intuitiven Standpunkt aus entwickelt wird. Das macht das Buch nicht unbedingt leichter zu lesen als ein herkömmliches Lehrbuch, ich denke aber, dass es sehr lehrreich ist, sich auf diese Lektüre einzulassen.

J. Matoušek, *Thirty-three Miniatures*, Mathematical and Algorithmic Applications of Linear Algebra, Student Math. Library **35**, AMS 2010.

Siehe auch <https://kam.mff.cuni.cz/%7Ematousek/stml-53-matousek-1.pdf>

Ein Buch mit direktem Bezug zur linearen Algebra (das auch im Text mehrfach zitiert wird), aber kein Lehrbuch. Stattdessen hat der Autor 33 Problemstellungen ausgewählt und beschrieben, bei deren Lösung die Methoden der linearen Algebra nützlich sind und oft auf überraschende Art und Weise eingesetzt werden.

M. Aigner, G. Ziegler, *Das BUCH der Beweise*, Springer 2018

<https://doi.org/10.1007/978-3-662-57767-7>

Der Mathematiker Paul Erdős hat die Redensart vom *Buch der Beweise* geprägt, in dem Gott für jedes mathematische Ergebnis den perfekten Beweis aufgeschrieben hat. Das Buch von Aigner und Ziegler ist sozusagen ein Auszug davon. Einige der Beweise sind auch mit wenig Grundkenntnissen zugänglich.

C. Rousseau, Y. Saint-Aubin, *Mathematik und Technologie*, Springer 2012,

<https://doi.org/10.1007/978-3-642-30092-9>

In fünfzehn Kapiteln werden verschiedene Anwendungen von mathematischen Methoden auf »Alltagsprobleme« erklärt. Es wird dadurch hervorragend sichtbar, wie viel Mathematik in der Technologisierung der Gesellschaft steckt. Zum Beispiel: Navigation mit GPS, Radiochirurgie/Computertomographie, Robotik, das RSA-Verfahren aus der Kryptographie, Zufallszahlengeneratoren. Auch einige der im Skript angesprochenen Themen werden diskutiert (zum Beispiel der Page-rank-Algorithmus und Codes).

D.5.3. Mathematik (auf-)schreiben. Mathematik verständlich aufzuschreiben, ist nicht immer einfach. Im Grunde sollte es ausreichen, wenn Sie das mit den Übungsaufgaben trainieren, aber es gibt auch Texte, in denen das explizit thematisiert wird. (Für den technischen Aspekt des mathematischen Schriftsatzes am Computer siehe C.2.)

Beutelspacher, *Das ist o.B.d.A. trivial!*, Vieweg 2006

<https://doi.org/10.1007/978-3-8348-9075-7>

Vielleicht ist es übertrieben, ein ganzes Buch zu dem Thema zu lesen, wie man Mathematik klar verständlich aufschreibt, aber einmal in dieses Buch hereinzuschauen, kann sich lohnen.

K. Conrad, [Advice on mathematical writing](#)²

(auf Englisch, aber das meiste lässt sich direkt auf deutsche Texte übertragen)

J. P. Serre, How to write mathematics badly,

<https://www.youtube.com/watch?v=ECQyFzzBH1o>.

Der berühmte Mathematiker J. P. Serre erklärt in diesem unterhaltsamen Vortrag, wie man es nicht machen sollte ...

²<https://kconrad.math.uconn.edu/blurbs/proofs/writingtips.pdf>

D.6. Webseiten, Videos

Unter dem Namen [3blue1brown](https://www.3blue1brown.com/)³ hat G. Sanderson viele verschiedene kurze Videos zu mathematischen Themen produziert und veröffentlicht, die ich größtenteils sehr gelungen finde, darunter eine Serie *Essence of linear algebra*. Von einigen dieser Videos gibt es auch eine deutsche Version. Direkt zum [Youtube-Kanal](https://www.youtube.com/c/3blue1brown)⁴.

[mathe-vital.de](http://www.mathe-vital.de/)⁵ ist ein Projekt des Lehrstuhls »Geometrie und Visualisierung« der TU München, das von J. Richter-Gebert gegründet wurde. Sie finden dort viele Visualisierungen und Applets zu Themen aus der linearen Algebra und aus anderen Gebieten.

C. Baer hat eine [Seite mit Übungsaufgaben](https://www.cbaer.eu/joomla/index.php/en/mathematics/practice)⁶ zu verschiedenen Themen.

Auf der Webseite www.mathematik.de⁷ der Deutschen Mathematiker-Vereinigung findet sich eine Menge Material, zum Beispiel viele Buchempfehlungen in der [Lesecke](https://www.mathematik.de/mathematik/lesecke)⁸.

Es gibt natürlich noch viel, viel mehr Videos und Webseiten zur Linearen Algebra und allgemeiner zur Mathematik, auch zu Themen des Mathematikstudiums – von unterschiedlich hoher Qualität...Wenn Sie einen Vorschlag haben, was an dieser Stelle noch ergänzt werden sollte (das gilt natürlich auch für Bücher), teilen Sie es mir mit!

³<https://www.3blue1brown.com/>

⁴<https://www.youtube.com/c/3blue1brown>

⁵<http://www.mathe-vital.de/>

⁶<https://www.cbaer.eu/joomla/index.php/en/mathematics/practice>

⁷<https://www.mathematik.de/>

⁸<https://www.mathematik.de/mathematik/lesecke>

Literaturverzeichnis

- [Ba] C. Baer, *Lineare Algebra und Analytische Geometrie*, Springer Spektrum 2018,
<https://doi.org/10.1007/978-3-658-22620-6>
- [Bo] S. Bosch, *Lineare Algebra*, Springer Spektrum 2014,
<https://doi.org/10.1007/978-3-642-55260-1>
- [El] J. Elstrodt, *Maß- und Integrationstheorie*, Springer 2011,
<https://doi.org/10.1007/978-3-642-17905-1>
- [Fi] G. Fischer, *Lineare Algebra*, Springer Spektrum 2014,
<https://doi.org/10.1007/978-3-658-03945-5>
- [Fi-AG] G. Fischer, *Analytische Geometrie*, Vieweg+Teubner, 7. Aufl., 2001
- [Fo] O. Forster, *Analysis I*, 12. Aufl., Springer Spektrum 2016,
<https://doi.org/10.1007/978-3-658-11545-6>
- [GP] I. M. Gelfand, V. A. Ponomarev, *Problems of linear algebra and classification of quadruples of subspaces in a finite dimensional vector space*, Colloq. Math. Soc. János Bolyai, vol. 5, Hilbert Space Operators, Tihany, 1970, 163–237.
- [Go] T. Gowers, *Mathematik*, Reclam-Sachbuch 2011 (oder das englische Original: *Mathematics. A very short introduction*, Oxford Univ. Press 2002)
- [Gr] D. Grieser, *Mathematisches Problemlösen und Beweisen*, Springer Spektrum 2017,
<https://doi.org/10.1007/978-3-658-14765-5>
- [He] H. Heuser, *Lehrbuch der Analysis, Teil I*, Teubner-Verlag, 9. Aufl., 1991.
- [Hi] S. Hildebrandt, *Analysis I*, Springer, 2. Aufl., 2006,
<https://doi.org/10.1007/3-540-29285-3>
- [Ho] K. Houston, *Wie man mathematisch denkt*, Springer Spektrum 2012
- [Hu] T. Hungerford, *Algebra*, Springer Graduate Texts in Math. **73**, 1974.
- [Jä] K. Jähnich, *Lineare Algebra*, Springer Hochschultext, 2. Aufl., 1981.
- [Jo] D. Joyner, *Adventures in Group Theory: Rubik's Cube, Merlin's Machine, and Other Mathematical Toys*, The Johns Hopkins University Press, 2008.
- [Kl] W. Klingenberg, *Lineare Algebra und Geometrie*, Springer, 3. Aufl., 1992.
- [LM] J. Liesen, V. Mehrmann, *Lineare Algebra*, Springer 2015,
<https://doi.org/10.1007/978-3-658-06610-9>
- [Lo] F. Lorenz, *Lineare Algebra I*, Spektrum Akad. Verlag 2004.
- [Ma] J. Matoušek, *Thirty-three Miniatures*, Mathematical and Algorithmic Applications of Linear Algebra, Student Math. Library **35**, AMS 2010. Siehe auch [preliminary version](#)⁹.
- [Nar] L. A. Nazarova, *Representations of a tetrad*, Izv. AN SSSR Ser. Math. **31** (6) (1967), 1361–1378 (Russisch); engl. Übersetzung in Math. USSR Izvestija **1** (1967), 1305–1321.
- [Naz] L. A. Nazarova, *Representations of quivers of infinite type*, Izv. AN SSSR, Ser. Mat. **37** (1973), 752–791 (Russisch); engl. Übersetzung in Math. USSR Izvestija **7** (1973), 749–792.
- [Pi] R. Pink, *Lineare Algebra I und II*, Zusammenfassung, 2016.
<https://people.math.ethz.ch/~pink/ftp/Lineare-Algebra-Zusammenfassung-20161006.pdf>
- [RS] C. Rousseau, Y. Saint-Aubin, *Mathematik und Technologie*, Springer 2012,
<https://doi.org/10.1007/978-3-642-30092-9>
- [So] W. Soergel, *Lineare Algebra I*,
<http://home.mathematik.uni-freiburg.de/soergel/Skripten/XXLA1.pdf>
- [So-AZT] W. Soergel, *Algebra und Zahlentheorie*,
<http://home.mathematik.uni-freiburg.de/soergel/Skripten/XXAL.pdf>
- [Ta] T. Tao, *An introduction to measure theory*, Graduate Studies in Mathematics **126**, Amer. Math. Soc. 2011, siehe auch <https://terrytao.files.wordpress.com/2011/01/measure-book1.pdf>
- [VK] S. Verma, J. P. Krishna, *Image compression and linear algebra*¹⁰, 2013.
- [vL] J. van Lint, *Introduction to Coding Theory*, Springer Graduate Texts in Mathematics **86**, 3rd ed., 1999.

⁹<https://kam.mff.cuni.cz/%7Ematousek/stml-53-matousek-1.pdf>

¹⁰<https://www.cmi.ac.in/%7Eksutar/NLA2013/imagecompression.pdf>

Index

- $-^I$, 53
- $-^c$, 52
- $-^n$, 53
- $:=$, 48
- $\cap$, 52
- $\cup$, 52
- $\in$, 49
- $\mapsto$, 54
- $\setminus$, 52
- $\rightarrow$, 54
- $f \circ g$, 57
- f^{-1} , 59
- $f^{-1}(Z)$, 56
- $f|_U$, 57
- $-^t$, 97, 117
- $\oplus$, 139
- $V^\vee$, 182
- $K^\times$, 81
- $f^\vee$, 182
- $\cong$, 166, 205
- $\xrightarrow{\sim}$, 166
- $\langle M \rangle$, 142, 206
- $-^{\text{ad}}$, 251
- Abb, 56
- Abbildung, 54
 - abstandserhaltend, 192, 282
 - alternierend, 238
 - bijektiv, 57
 - Bild, 56
 - duale, 182, 329
 - Einschränkung, 57
 - identische, 56
 - injektiv, 57
 - Komposition, 57
 - konstant, 56
 - linear, 163, 327
 - multilinear, 238
 - surjektiv, 57
 - Urbild, 56
 - Verkettung, 57
- abc-Vermutung, 41
- Abelsche Gruppe, 202
- Absolutbetrag, 63
- Absolutes Glied, 94, 339
- Absolutterm, 94, 339
- Abstand, 192, 276
- abstandserhaltend, 192, 282
- abzählbar, 76
- Adjazenzmatrix, 318
- ähnlich, 206
- Äquivalenzklasse, 71
- Äquivalenzrelation, 70
- Affiner Raum, 274
- Affiner Teilraum, 275
- Affiner Unterraum, 275
- Allgemeine lineare Gruppe, 203
- Allquantor, 47
- alternierend, 238
- Ananas, 21
- antisymmetrisch, 74
- Assoziativgesetz, 79, 80, 201
- Aufrundungsfunktion, 309
- Aussage, 46
- Automorphismus, 166
- Basis, 141, 328
 - duale, 185, 329
 - geordnete, 173
- Basisaustauschsatz, 147
- Basisreproduktionszahl, 269
- Basiswechsel
 - für lineare Abbildung, 176
- Basiswechselmatrix, 176
- Betrag, 63
- Beweis
 - direkt, 47
 - durch Widerspruch, 47
 - indirekt, 47
- Bijektion, 57
- bijektiv, 57
- Bild, 56
 - einer Matrix, 120
 - eines Gruppenhomomorphismus, 208
- Bild (einer Abbildung), 56
- Bild (unter einer Abbildung), 54
- Bildbereich, 56
- Bilinearform, 279
- Blockmatrix, 118
- Bruhat-Zerlegung, 219
- $\mathbb{C}$, 82
- Cauchy-Schwarzsche Ungleichung, 280
- Charakteristik, 91
- Chinesischer Restsatz, 223
- Code, 308
 - Erzeugermatrix, 311
 - Generatormatrix, 311

- linear, 310
- Cramersche Regel, 251, 333
- Darstellende Matrix, 174
- Definitionsbereich, 55
- $\det(A)$, 240
- Determinante, 125, 240, 332
 - eines Endomorphismus, 246
 - Produktsatz, 243
- Determinantenfunktion, 239
- $\det(f)$, 246
- diag, 118
- diagonalisierbar, 257, 334
- Diagonalmatrix, 118
- Diagramm
 - kommutativ, 340
- Diedergruppe, 210, 211
- Differentialgleichungen, 269
- Differenz (von Mengen), 52
- Dimension, 328
- Dimensionsformel
 - Durchschnitt von UVR, 150
 - für lineare Abbildungen, 171
- Direkte Summe
 - (von UVR), 139
 - (von Vektorräumen), 155
 - innere, 155
 - äußere, 155
- disjunkt, 53
- Distributivgesetz, 80
- Division mit Rest, 64
- Divisionsalgebra, 27
 - endlich-dimensional, 160
- Doppeldualraum, 186
- Drehmatrix, 194
- Drehung, 121, 191, 194, 247, 284
 - $\mathbb{R}^2$, 194
- Dreiecksungleichung, 277, 278
- Duale Basis, 185
- Duale Abbildung, 182, 329
- Duale Basis, 329
- Dualer Vektorraum, 182
- Dualraum, 182, 329
- Durchschnitt, 52, 53
- $\mathcal{D}_V$, 239
- Echte Teilmenge, 51
- Ecke (eines Graphen), 315
- Eigenraum, 258, 333
- Eigenvektor, 258, 333
- Eigenwert, 258, 333
- Einheiten, 81
- Einheitsmatrix, 102, 118
- Einheitsvektor
 - kanonisch, 141
- Einschränkung, 57
- Elementare Spaltenumformungen, 103
- Elementare Zeilenumformungen
 - eines LGS, 103
 - von Matrizen, 103
- Elementarmatrix, 212
- E_n , 102, 118
- endlich erzeugt, 144
- endlich-dimensional, 148
- Endomorphismus, 165
 - diagonalisierbar, 257
- Entwicklungssatz
 - von Laplace, 250, 333
- Epimorphismus, 171
- Erweiterungskörper, 82
- Erzeugendensystem, 143, 327
- Euler-Formel für ebene Graphen, 322
- Eulersche φ -Funktion, 223
- Eulersche Polyederformel, 322
- Existenzquantor, 47
- $\mathbb{F}_2$, 83
- $\mathbb{F}_3$, 84
- f_A , 119
- Familie, 57
- Fehlstand, 217, 218
- Fibonacci-Zahlen, 21, 126, 158, 229, 245, 265
- $\mathbb{F}_p$, 90
- Funktion, 55
- Ganze Zahlen, 50
- Gauß-Algorithmus, 102, 105
- Gaußklammer, 309
- ggT, 64
- gleichmächtig, 75
- $GL_n(K)$, 203
- Goldbachsche Vermutung, 41
- Goldener Schnitt, 21
- Grad, 96
- Graph, 315
 - isomorph, 315
 - planar, 321
 - zusammenhängend, 322
- Griechische Buchstaben, 48
- Grundkörper, 135
- Gruppe, 201, 325
 - abelsch, 202, 325
 - isomorph, 205
 - kommutativ, 202, 325
 - spezielle lineare, 212
 - symmetrische, 214
 - trivial, 202
 - zyklisch, 207
- Gruppenhomomorphismus, 204, 325
- Gruppenisomorphismus, 205
- Größter gemeinsamer Teiler, 64
- Größtes Element, 75
- Halbordnung, 74
- Hamming-Distanz, 309
- $\text{Hom}(G, H)$, 204
- $\text{Hom}_K(V, W)$, 165
- homogen
 - (LGS), 97
- Homomorphisms
 - von Gruppen, 204
- Homomorphismus
 - von Vektorräumen, 163, 327
- Hyperebene, 284

- Identische Abbildung, 56
- Identität(-sabbildung), 56
- $\text{Im}(A)$, 120
- Imaginärteil, 83
- $\text{Im}(f)$, 56, 170, 208
- Indexmenge, 53
- Induktion, 60
- Induktionsanfang, 60
- Induktionsschritt, 60
- inhomogen
 - (LGS), 97
- Injektion, 57
- injektiv, 57
- Innenwinkel, 296
- Inverse Matrix, 122
- Inverses Element, 80, 201
- invertierbar, 122
- Isometrie, 282
- isomorph, 166, 205
 - (Graphen), 315
- Isomorphismus
 - von Gruppen, 205
 - von Vektorräumen, 166
 - von Graphen, 315
- Jordanbogen, 287
 - geschlossen, 288
 - Länge, 288
 - rektifizierbar, 288
- K -Vektorraum, 135
- kanonisch, 340
- Kante (eines Graphen), 315
- Kardinalität, 68
- Kartesisches Produkt, 53
- $\text{Ker}(A)$, 120
- $\text{Ker}(f)$, 170, 208
- Kern
 - einer linearen Abbildung, 170
 - einer Matrix, 120
 - eines Gruppenhomomorphismus, 208
- Kleiner Fermatscher Satz, 92
- Kleinstes Element, 75
- Knoten (eines Graphen), 315
- Koeffizient
 - (Matrix), 100, 113
- Koeffizienten, 97
- Koeffizientenmatrix, 101
 - erweiterte, 101
- Körper, 79, 203, 326
 - Charakteristik, 91
 - Einheiten, 81
 - endlich, 87, 157
 - multiplikative Gruppe, 81
- Kommutative Gruppe, 202
- Kommutatives Diagramm, 340
- Kommutativgesetz, 80, 202
- Komplement, 52
 - eines Untervektorraums, 139
- Komplementärmatrix, 251
- Komplementärraum, 139
- Komplexe Zahlen
 - Polarkoordinaten, 293
- Komplexe Zahlen, 82, 119, 289
 - Imaginärteil, 83
 - Realteil, 83
- Komposition, 57
- kongruent, 91
- Konjugation, 205
- konjugiert, 206
- Konstante Abbildung, 56
- Kontinuumschypothese, 77
 - verallgemeinerte, 77
- Kontraposition, 47
- Koordinatenabbildung, 169
- Koordinatenvektor, 169
- Koprodukt, 155
- Koprodukt (von Vektorräumen), 155
- Kosinussatz, 297
- Laplacescher Entwicklungssatz, 250, 333
- $\text{\LaTeX}$, 342
- Leere Menge, 49
- Leere Summe, 63
- Leeres Produkt, 53, 63
- Legendre-Symbol, 226
- Leibniz-Formel, 240, 254
- Lemma, 340
- LGS, 97
- linear abhängig, 145
- linear unabhängig, 145, 328
- Lineare Abbildung, 163, 327
 - Kern, 170
- Lineare Hülle, 143
- Lineare Ordnung, 74
- Linearer Code, 310
- Lineares Gleichungssystem, 28, 97
 - eindeutig lösbar, 98
 - homogen, 97
 - inhomogen, 97
 - lösbar, 98
 - unterbestimmt, 110
 - zugehöriges homogenes, 98
- Linearform, 182
- Linearkombination, 140
 - triviale, 144
- Länge
 - einer Permutation, 217
 - eines Weges, 286
 - eines Vektors, 278
- Lösung
 - trivial, 98
- Lösungsmenge, 97
- Lösungsvektor, 97
- Matrix, 100, 113, 330
 - Abbildung f_A , 119
 - ähnlich, 206
 - äquivalent, 181
 - Bild, 120
 - Blockmatrix, 118
 - Diagonale, 102
 - diagonalisierbar, 257
 - Diagonalmatrix, 118

- Einheitsmatrix, 118
- Eintrag, 100, 113
- inverse, 122
- invertierbar, 122
- Kern, 120
- Koeffizient, 100, 113
- konjugiert, 206
- obere Dreiecksmatrix, 118
- orthogonale, 283
- quadratisch, 100
- Rang, 179
- Spaltenrang, 179
- Spur, 253
- transponierte, 117
- Zeilenrang, 179
- Matrizenprodukt, 113
- max, 48
- Maximales Element, 75
- Maximum, 48
- McEliece-Verfahren, 312
- Menge, 49
 - abzählbar, 76
 - endlich, 68
 - leere, 49
- min, 48
- Minimales Element, 75
- Minimum, 48
- $M_{m \times n}(K)$, 100, 113
- $M_n(K)$, 100, 113
- modulo, 87
- Monomorphismus, 171
- multilinear, 238
- Multiplikative Gruppe
 - (eines Körpers), 81
- Mächtigkeit, 68
- $\mathbb{N}$, 50
- n -tupel, 53
- Nachbarschaftsmatrix, 318
- Natürliche Zahlen, 50
- Negation, 46
- Neutrales Element, 80, 201
- Norm (eines Vektors), 278
- Nullabbildung, 163
- Nullmatrix, 100, 113
- Nullraum, 111, 136
- Nullstelle, 94
- Nullvektor, 99, 135
- Nullvektorraum, 136, 167
- OBdA, 340
- Obere Dreiecksmatrix, 118
- Obere Schranke, 337
- Obermenge, 52
- OE, 340
- ord, 222
- Ordnung, 74
 - einer Gruppe, 222
 - eines Gruppenelements, 222
 - linear, 74
 - partiell, 74, 336
 - total, 74, 336
- Orientierung, 248
- orientierungserhaltend, 248
- orthogonal, 279
- Orthogonale Matrix, 283
- Orthogonale Gruppe, 283
- Page-rank-Algorithmus, 23, 127, 197, 267
- Partielle Ordnung, 74, 336
- Permutation, 215
 - gerade, 217
 - Länge, 217
 - Signum, 217
 - ungerade, 217
- Permutationsmatrix, 209
- Pfad, 318
- $\prod$, 63
- Planarer Graph, 321
- Polarkoordinaten, 293
- Polynomfunktion, 94
 - Grad, 96
- positiv definit, 279
- Potenz, 81
- Potenzmenge, 74
- Primeigenschaft, 65, 89
- Primfaktorzerlegung in $\mathbb{Z}$, 66
- Primitivwurzel, 225
- Primzahl, 38, 63
- Prinzip vom kleinsten Element, 60
- Produkt
 - leeres, 53, 63
 - von Gruppen, 203
 - von Matrizen, 113
- Produkt (von Mengen), 53
- Produkt (von Vektorräumen), 155
- Produktsymbol, 63
- P_σ , 209
- Public-Key-Verfahren, 93
- 15-Puzzle, 230
- $\mathbb{Q}$, 50
- Quadratisches Reziprozitätsgesetz, 226
- Quadratisches Sieb, 22, 159
- Quadrupel, 53
- Quantor, 47
- Quaternionen, 27, 132, 203, 256
- Quintupel, 53
- $\mathbb{R}$, 50
- Ramsey-Zahl, 316
- Rang, 180, 329
 - der dualen Abbildung, 185
 - einer linearen Abbildung, 171
 - einer Matrix, 179
- Rationale Zahlen, 50
- Realteil, 83
- Reduzierte Zeilenstufenform, 105
- Reelle Zahlen, 50
- reflexiv, 70
- Regel von Sarrus, 241
- rektifizierbar, 286, 288
- Relation, 70
 - Äquivalenzrelation, 70
 - antisymmetrisch, 74

- reflexiv, 70
- symmetrisch, 70
- total, 74
- transitiv, 70
- Repräsentant, 71
- Restklasse, 71, 87
- Restklassenring, 203
- $\text{rg}(A)$, 179
- Ring, 203
- Ringschluss, 47
- RSA-Verfahren, 22, 93, 224
- Rubiks Würfel, 231
- Russells Paradox, 50
- Satz
 - (kleiner), von Fermat, 92
 - von Euler, 223
 - von Lagrange, 222
 - von Ramsey, 316
 - von Schröder-Bernstein, 76
- Scherung, 122
- Schiefkörper, 86, 203, 266
- Schnitt, 336
- Schnittmenge, 52
- Schubfachprinzip, 341
- Sechsfarbensatz, 323
- SEIR-Modell, 269
- senkrecht, 278, 279
- Σ , 62
- Signum, 217
- Skalar, 99
- Skalarmultiplikation
 - auf K^n , 99
- Skalarprodukt, 278
- $SL_n(K)$, 212, 244
- Smith-sche Normalform, 177
- S_n , 214
- Spaltenrang, 180
- Spaltenrank, 179
- Spaltenumformungen, 103, 154
- Spaltenvektor, 97
- Spann, 143
- Spezielle lineare Gruppe, 212, 244
- Spiegelung, 121, 193, 284
- Spur, 253, 333
 - eines Endomorphismus, 253
- Standard-Skalarprodukt, 278
- Standardbasis, 141, 156
- Standardbasisvektor, 141
- Standardvektorraum, 136
- Strassen-Algorithmus, 115
- Streckung, 122
- Summe
 - auf K^n , 99
 - leer, 63
 - von Matrizen, 113
- Summe von UVR, 139
- Summensymbol, 62
- Surjektion, 57
- surjektiv, 57
- Symmetriegruppe, 209
- symmetrisch, 70
- Symmetrische Gruppe, 214
- System, 57
- $-^t$, 117
- Teilbarkeit, 63
- Teiler, 63, 75
- Teilgraph, 316
- Teilkörper, 82
- Teilmenge, 51
 - echte, 51
- Teilraum
 - eines Vektorraums, 138
 - von K^n , 111
- Theorema aureum, 226
- Totale Ordnung, 74, 336
- transitiv, 70
- Translation, 165, 275
- Transponierte Matrix, 117
- Transposition, 217
 - elementare, 217
- Tripel, 53
- trivial, 341
- Triviale Linearkombination, 144
- Triviale Lösung, 98
- Träger, 215
- Tupel, 53
- überabzählbar, 76
- Umkehrabbildung, 59
- Ungleichung
 - von Cauchy-Schwarz, 280
- Untergruppe, 206, 325
 - von Teilmenge erzeugte, 206
- Untermenge, 51
- Untervektorraum, 138, 326
 - trivialer, 138
 - von Teilmenge aufgespannter, 142
 - von Teilmenge erzeugter, 142
- Urbild, 56
- Vandermonde-Matrix, 245, 263
- Vektor, 99, 137
- Vektorraum, 135, 326
 - direkte Summe, 155
 - endlich erzeugt, 144
 - endlich-dimensional, 148
 - isomorph, 166
 - Produkt, 155
- Vektorraumisomorphismus, 166
- Vereinigung, 52, 53
- Verkettung, 57
- Verknüpfung
 - (auf einer Menge), 79
 - (von Abbildungen), 57
- Verschiebung, 165, 275
- Vier-Unterraum-Problem, 199
- Vierfarbensatz, 321
- V_λ , 258
- $V_\lambda(f)$, 258
- Vollständige Induktion, 60
- Weg
 - in $\mathbb{R}^n$, 286

- Länge, 286
 - rektifizierbar, 286
- in einem Graph, 318
 - Länge, 318
- Wert (einer Abbildung), 54
- Wertebereich, 55
- Widerspruchsbeweis, 47
- Winkel, 295, 296
- W_n , 209
- wohldefiniert, 72, 341
- Wohlordnungssatz, 336

- $\mathbb{Z}$, 50
- Zauberwürfel, 231
- Zeilenrang, 179, 180
- Zeilenstufenform, 104
 - reduzierte, 105
- Zeilenumformungen
 - eines LGS, 103
 - von Matrizen, 103
- Ziel (einer Abbildung), 55
- $\mathbb{Z}/n$, 87, 203
- zusammenhängend
 - (Graph), 322
- Zykel, 215
- Zyklische Gruppe, 207