

Lineare Algebra I

Wintersemester 2025/2026

Universität Bayreuth

MICHAEL STOLL

INHALTSVERZEICHNIS

1. Einige allgemeine Vorbemerkungen	2
2. Die Sprache der Mathematik: Logik und Mengenlehre	4
3. Algebraische Strukturen: Gruppen, Ringe, Körper	27
4. Der Körper der komplexen Zahlen	33
5. Vektorräume: Definition und Beispiele	38
6. Untervektorräume	43
7. Erzeugendensysteme	47
8. Lineare Unabhängigkeit	54
9. Basis und Dimension	59
10. Lineare Abbildungen	72
11. Matrizen	87
12. Der Normalformalgorithmus und Lineare Gleichungssysteme	92
13. Matrizen und lineare Abbildungen	105
14. Die Determinante	110
15. Eigenwerte und Eigenvektoren	122
16. Diagonalisierbarkeit	133

1. EINIGE ALLGEMEINE VORBEMERKUNGEN

Die meisten von Ihnen kommen mehr oder weniger direkt von der Schule (die prominenteste Ausnahme sind die zukünftigen Realschullehrer/innen, die die Lineare Algebra I erst im dritten Semester hören). Das Erste, das Sie sich zu Beginn Ihres Mathematik-Studiums klar machen müssen, ist, dass das, was Sie in der Schule unter der Bezeichnung „Mathematik“ kennen gelernt haben, nicht wirklich Mathematik ist. Das bedeutet, dass Sie hier an der Universität im Grunde auf völlig andere Art und Weise gefordert sein werden als an der Schule. Das heißt jetzt nicht, dass Sie die ganze Schulmathematik vergessen können — manches kann als Beispielmateriale noch nützlich sein, und es schadet auch nicht, wenn man eine gewisse Fertigkeit im Rechnen hat, wie man sie an der Schule lernt.

Was folgt aus diesem in Deutschland leider traditionellen Bruch zwischen Schule und Universität?

- Die meisten von Ihnen werden sich erst einmal eher schwer tun. Das ist völlig normal und kein Grund zur Beunruhigung.
- Wenn Sie in der Schule in Mathe sehr gut waren, heißt das nicht, dass Ihnen die Mathematik an der Universität auch leicht fällt. Umgekehrt kann es sein, dass Ihnen die Mathematik an der Schule langweilig war und Sie dann hier auf den Geschmack kommen.
- Sie sollten nicht erwarten, den Stoff sofort während der Vorlesung zu verstehen. Das Nacharbeiten der Vorlesung ist sehr wichtig, da man Mathematik nur verstehen kann, wenn man darüber nachdenkt. (Das Modulhandbuch sieht drei Stunden pro Woche dafür vor.) Ganz wichtig ist auch, dass Sie die Übungsaufgaben bearbeiten, denn richtig versteht man den Stoff erst, wenn man ihn anwendet. (Das Modulhandbuch sieht dafür fünf Stunden pro Woche vor.) Dabei hilft es, gemeinsam in kleinen Gruppen zu arbeiten, denn für das Verständnis ist es ungemein förderlich, wenn man versucht, jemand anderem etwas zu erklären. Eine sehr gut dafür geeignete Umgebung ist das *Lernzentrum* der Mathematik (jeden Nachmittag im S79).
- Für diejenigen von Ihnen, die Lehrerinnen oder Lehrer werden wollen, heißt das umgekehrt auch, dass Sie den größten Teil von dem, was Sie hier lernen, in der Schule nicht direkt verwenden können. Es ist zu hoffen, dass sich das bald einmal ändert und Sie die Möglichkeit haben werden, die „richtige“ Mathematik Ihren Schülern nahezubringen. In jedem Fall sollte Sie die Ausbildung, die Sie an der Universität erhalten, in die Lage versetzen, Ihren Unterricht innerhalb der Mathematik einzuordnen und weiter gehende Fragen Ihrer Schülerinnen und Schüler souverän zu beantworten.
- Das Lernziel im Mathematikstudium allgemein und in dieser Vorlesung im Besonderen ist, dass Sie die Mathematik *verstehen*. Das bedeutet, dass Sie selbständig mit den Begriffen und Aussagen umgehen können und in der Lage sind, sie in geeigneter Weise zu kombinieren, um mathematische Probleme zu lösen. Dabei handelt es sich meistens nicht um Rechenaufgaben, sondern häufig darum, eine stichhaltige Begründung (einen mathematischen Beweis) für die Richtigkeit einer gegebenen Aussage zu finden. Wie oben schon betont, lernt man das nur, wenn man es *selbst versucht*. Die Prüfung zu dieser Vorlesung wird Ihr Verständnis dann testen. Die wöchentlichen Übungsaufgaben sind von ähnlicher Art wie die Prüfungsaufgaben; sie selbständig zu bearbeiten ist die beste Vorbereitung. Es kann verführerisch sein, Lösungen abzuschreiben (egal,

ob sie von einer willigen Kommilitonin, einem willigen Kommilitonen oder einer KI zur Verfügung gestellt werden), aber damit betrügen Sie sich am Ende nur selbst. Wenn Sie Schwierigkeiten haben, in der vorgesehenen Zeit mit den Übungsaufgaben fertig zu werden, dann nehmen Sie die vorhandenen Angebote wahr: Gehen Sie ins Lernzentrum, in die Fragestunde und natürlich in die Übungen.

Lassen Sie sich von den Schwierigkeiten am Anfang nicht zu sehr frustrieren! Bei den meisten von Ihnen wird in den ersten beiden Semestern der Groschen fallen. Falls Sie aber nach zwei Semestern immer noch das Gefühl haben, nichts zu verstehen, dann kann es auch sein, dass das Mathematikstudium doch nicht das Richtige für Sie ist.

Ich habe in dieses Skript an manchen Stellen Links zu Webseiten eingebaut, die so aussehen (dieser Link führt auf meine Homepage). Die meisten davon verweisen auf die Wikipedia, die für den Zweck einer ersten Orientierung meistens gut geeignet ist. (Als Hauptquelle für Zitate in einer wissenschaftlichen Arbeit wie z.B. einer Bachelor- oder Masterarbeit ist die Wikipedia aber nicht geeignet. Da müssen Sie Lehrbücher oder Fachartikel zitieren.)

2. DIE SPRACHE DER MATHEMATIK: LOGIK UND MENGENLEHRE

Worum geht es nun in der Mathematik?

Die Wikipedia schreibt (Stand Oktober 2025):

Für *Mathematik* gibt es keine allgemein anerkannte Definition; heute wird sie üblicherweise als eine Wissenschaft beschrieben, die durch logische Definitionen selbst geschaffene abstrakte Strukturen mittels der Logik auf ihre Eigenschaften und Muster untersucht.

Es geht also unter anderem um *Abstraktion*. Man abstrahiert von den speziellen Eigenschaften, die man in verschiedenen Situationen vorliegen hat, und zieht das Gemeinsame heraus. Dann versucht man auf der Grundlage nur dieser wesentlichen Merkmale möglichst viele Aussagen abzuleiten, die dann auf *alle* Situationen zutreffen, die diese Merkmale aufweisen. Dies geschieht durch den zentralen Vorgang aller mathematischen Tätigkeit, nämlich durch das Führen eines *mathematischen Beweises*: Die englische Version der Wikipedia-Seite zur Mathematik enthält den Satz

Mathematics uses pure reason to prove properties of objects, a *proof* consisting of a succession of applications of deductive rules to already established results.

Zugespitzt kann man sagen, dass eine *Mathematikerin* bzw. ein *Mathematiker* sich dadurch auszeichnet, einen solchen Beweis führen zu können:

- Das wichtigste „Lernziel“ in den Grundvorlesungen besteht darin, dass Sie lernen, *wie man mathematische Beweise führt*.

Sie sollen hier natürlich auch und nicht zuletzt Ergebnisse und Methoden der Linearen Algebra kennen lernen, aber ohne die mathematische Grundfertigkeit des Beweisens würde Ihnen das kaum etwas nützen.

Bevor wir damit beginnen können, müssen wir die Vokabeln und Grammatik der Sprache der Mathematik lernen. Mathematische Aussagen und Beweise werden in der Sprache der *Logik* formuliert; die Objekte, von denen die Rede ist, in der Sprache der *Mengenlehre*. Beide werden wir hier kurz einführen (oder wiederholen, je nachdem wie viel Sie davon schon aus der Schule kennen). Es handelt sich um das „Handwerkszeug“, mit dem Sie täglich zu tun haben werden, also passen Sie gut auf!

2.A Aussagenlogik.

Die Aussagenlogik verknüpft mathematische *Aussagen* (die entweder wahr oder falsch sind) miteinander und untersucht, wie das Wahr- oder Falschsein einer zusammengesetzten Aussage von den beteiligten Aussagen abhängt. Aussagen in diesem Sinn sind zum Beispiel „37 ist eine Primzahl“ (eine wahre Aussage) oder „ $6 \cdot 9 = 42$ “ (eine falsche Aussage).

(Die Farbe grün, in der der folgende Text erscheint, kennzeichnet Definitionen in diesem Skript.)

2.1. **Definition.** Seien A und B Aussagen. Die logischen Verknüpfungen sind:

DEF

- (1) Die *Negation*: wir schreiben „nicht A “ oder „ $\neg A$ “ für die Verneinung der Aussage A . $\neg A$ ist genau dann wahr, wenn A falsch ist, und umgekehrt.
- (2) Die *Konjunktion*: wir schreiben „ A und B “ oder „ $A \wedge B$ “; diese Aussage ist genau dann wahr, wenn sowohl A als auch B wahr sind.
- (3) Die *Disjunktion*: wir schreiben „ A oder B “ oder „ $A \vee B$ “; diese Aussage ist genau dann wahr, wenn wenigstens eine der Aussagen A und B wahr ist.
- (4) Die *Implikation*: wir schreiben „aus A folgt B “, „ A impliziert B “ oder „ $A \Rightarrow B$ “; diese Aussage ist genau dann wahr, wenn A falsch oder B wahr ist (oder beides).
- (5) Die *Äquivalenz*: wir schreiben „ A genau dann, wenn B “, „ A und B sind äquivalent“ oder „ $A \Leftrightarrow B$ “; diese Aussage ist genau dann wahr, wenn entweder A und B beide wahr oder A und B beide falsch sind. $\diamond$

$\neg A$
 $A \wedge B$
 $A \vee B$
 $A \Rightarrow B$
 $A \Leftrightarrow B$

Alle hier aufgeführten Schreibweisen sind möglich und erlaubt; die Schreibweise „ $A \wedge B$ “ ist zum Beispiel nicht besser oder schlechter als „ A und B “ (nur kürzer). Bei verschachtelten Verknüpfungen werden Klammern gesetzt, um die Bedeutung klar zu machen: Bei „ A und B oder C “ ist sonst nicht klar, ob $(A \wedge B) \vee C$ oder $A \wedge (B \vee C)$ gemeint ist.

Die wichtigste (und gleichzeitig die am schwersten zu verstehende) dieser Verknüpfungen ist die Implikation. Sie ist wichtig, weil die große Mehrzahl aller mathematischen Sätze die Form einer Implikation haben: Wenn gewisse Voraussetzungen A gelten, dann folgt eine Aussage B . Sie ist ein wenig schwierig, weil mit ihr im täglichen Leben oft ungenau bis falsch umgegangen wird. Vor allem neigen viele Menschen dazu, zwischen „aus A folgt B “ und „aus B folgt A “ nicht sorgfältig zu unterscheiden. Das mag damit zusammenhängen, dass man gerne die Voraussetzung A in $A \Rightarrow B$ als „Ursache“ für die Folgerung B missversteht und man im täglichen Leben häufig von der Wirkung auf die Ursache schließt. Den Unterschied zwischen $A \Rightarrow B$ und $B \Rightarrow A$ zu begreifen, ist die erste wichtige Hürde für Sie als zukünftige Mathematiker. Machen Sie sich Folgendes klar:

- $A \Rightarrow B$ ist jedenfalls immer dann wahr, wenn A falsch ist.
- $A \Rightarrow B$ ist auch immer dann wahr, wenn B wahr ist.
- $A \Rightarrow B$ kann *nur dann* falsch sein, wenn A wahr, aber B falsch ist.



Wir verwenden manchmal die Schreibweise „ $\perp$ “ für das *Falsum*, also eine stets falsche Aussage oder einen Widerspruch. Analog gibt es die stets wahre Aussage „ $\top$ “. Dann können wir also schreiben

$$\perp \Rightarrow B \quad \text{und} \quad A \Rightarrow \top \quad \text{gelten stets.}$$

Für die Lateiner unter Ihnen: Die erste dieser Tatsachen ist auch unter dem schönen Namen *Ex falso quodlibet* bekannt.

Folgende Schlussweise ist *nicht* erlaubt:

Wir wissen, dass $A \Rightarrow B$ gilt. Wir wollen A zeigen. Also nehmen wir einmal an, dass A stimmt. Dann müsste auch B gelten. B ist aber richtig, also muss auch A gelten.

(Die Farbe blau wird in diesem Skript für Beispiele verwendet.)

2.2. Beispiel. Wir wollen $0 = 1$ zeigen. Dazu formen wir um: Aus $0 = 1$ folgt durch Verdoppeln $0 = 2$, dann durch Subtraktion von 1 auf beiden Seiten $-1 = 1$, schließlich durch Quadrieren $1 = 1$, was offensichtlich stimmt. Also gilt auch die ursprüngliche Gleichung $0 = 1$.

BSP
falscher
Schluss

Hier ist alles korrekt bis auf das „Also“ im letzten Satz, denn der Schluss von $A \Rightarrow B$ und B auf A ist nicht möglich.

Häufig sieht man solche „Beweise“ in der Form

$$0 = 1 \xrightarrow{\cdot 2} 0 = 2 \xrightarrow{-1} -1 = 1 \xrightarrow{\text{quadrieren}} 1 = 1 \quad \checkmark$$

Das macht das Problem etwas klarer: Die Pfeile zeigen in die falsche Richtung! Für ein korrektes Argument müssen die Pfeile in Richtung der zu beweisenden Aussage zeigen. Tatsächlich sind die beiden ersten Implikationen sogar Äquivalenzen:

$$0 = 1 \iff 0 = 2 \iff -1 = 1 \implies 1 = 1$$

Man kann also von $-1 = 1$ auf $0 = 1$ schließen, aber der letzte Implikationspfeil lässt sich nicht umdrehen. ♣

Der Schluss von $A \Rightarrow B$ und A auf B ist aber natürlich sehr wohl möglich und stellt eine der grundlegenden Schlussweisen in Beweisen dar. Häufig ist „ $A \Rightarrow B$ “ ein mathematischer Satz, der angewendet werden soll. Wir weisen nach, dass die Voraussetzung A gilt, und können dann auf B schließen.

Einen mathematischen Beweis zu führen, kann man sich wie ein Spiel vorstellen: Auf dem „Spielbrett“ befindet sich zum Einen die Aussage, die wir gerade beweisen wollen, und zum Anderen befinden sich dort auch die Objekte (zum Beispiel gewisse Mengen oder Zahlen) und gewisse Aussagen, deren Gültigkeit wir angenommen haben. Zu Beginn des Spiels befindet sich lediglich die zu beweisende Aussage auf dem Brett. *Schlussregeln* wie die eben erwähnte beschreiben die erlaubten Spielzüge. Diese können dazu führen, dass es mehrere Spielbretter gibt. Das Spiel ist gewonnen, wenn alle Spielbretter leer geworden sind, also keine zu beweisende Aussage mehr vorhanden ist. Ein Spielbrett kann abgeräumt werden, wenn sich die zu beweisende Aussage unter den Annahmen befindet.

Die oben beschriebene Schlussregel sagt, wie man eine Implikation, die sich unter den Annahmen befindet, *verwenden* kann. Solche Regeln heißen Eliminationsregeln:

2.3. Definition. Die *Eliminationsregeln* für die verschiedenen aussagenlogischen Verknüpfungen sind wie folgt.

DEF
Eliminations-
regeln

- **Implikation:** Wenn sich sowohl $A \Rightarrow B$ als auch A unter den Annahmen befinden, dann kann man B zu den Annahmen hinzufügen.
- **Konjunktion:** Wenn sich $A \wedge B$ unter den Annahmen befindet, dann kann man A und B zu den Annahmen hinzufügen.
- **Falsum:** Da $\perp \Rightarrow A$ für jede Aussage A gilt, erhalten wir eine Eliminationsregel für das Falsum $\perp$: Wenn sich $\perp$ unter den Annahmen befindet, kann man jede beliebige Aussage zu den Annahmen hinzufügen. Man kann $\perp$ als „Widerspruch“ interpretieren. Die Regel besagt dann, dass aus widersprüchlichen Annahmen alles folgt.
- **Negation:** Die Negation $\neg A$ ist gleichwertig mit $A \Rightarrow \perp$. Die Eliminationsregel für die Implikation besagt dann, dass wir $\perp$ (und damit auch jede beliebige andere Aussage) zu den Annahmen hinzufügen können, wenn sowohl A als auch $\neg A$

unter den Annahmen sind: Eine Aussage zusammen mit ihrer Negation führt auf einen Widerspruch.

- **Äquivalenz:** Eine Äquivalenz $A \Leftrightarrow B$ ist gleichwertig mit der Konjunktion von zwei Implikationen: $(A \Rightarrow B) \wedge (B \Rightarrow A)$. Befindet sich $A \Leftrightarrow B$ unter den Annahmen, können wir also die beiden Implikationen $A \Rightarrow B$ und $B \Rightarrow A$ zu den Annahmen hinzufügen.
- **Disjunktion:** Die Eliminationsregel für die Disjunktion $A \vee B$ ist etwas komplizierter; sie verkörpert das Prinzip der *Fallunterscheidung*: Um zu zeigen, dass eine Aussage unter der Voraussetzung $A \vee B$ gilt, genügt es, die Aussage einmal unter der Voraussetzung A und einmal unter der Voraussetzung B zu beweisen. Das bedeutet, dass wir das betreffende Spielbrett durch *zwei* Spielbretter ersetzen, bei denen wir $A \vee B$ in den Annahmen einmal durch A und einmal durch B ersetzen. $\diamond$

Als Gegenstück zu den Eliminationsregeln gibt es die Einführungsregeln, die uns sagen, wie man eine Aussage *beweisen* kann, die eine Verknüpfung anderer Aussagen ist.

2.4. Definition. Wir haben folgende *Einführungsregeln* für die aussagenlogischen Verknüpfungen.

DEF
Einführungs-
regeln

- **Implikation:** Um die Implikation $A \Rightarrow B$ zu beweisen, genügt es, die Aussage B unter der Annahme A zu zeigen. Auf dem Spielbrett ersetzen wir also die zu beweisende Aussage $A \Rightarrow B$ durch B und fügen A zu den Annahmen hinzu.
- **Konjunktion:** Um eine Konjunktion zu beweisen, zeigt man die beiden Aussagen einzeln. Man ersetzt also das Spielbrett mit der zu beweisenden Aussage $A \wedge B$ durch zwei Spielbretter, auf denen die zu beweisende Aussage A bzw. B ist (mit denselben Annahmen wie vorher).
- **Äquivalenz:** Daraus ergibt sich auch die Einführungsregel für die Äquivalenz: Statt $A \Leftrightarrow B$ zu beweisen, zeigen wir sowohl $A \Rightarrow B$ als auch $B \Rightarrow A$ unter den vorliegenden Annahmen. Das ist ein Beweisschritt, dem wir in der Vorlesung recht oft begegnen werden.
- **Disjunktion:** Um eine Disjunktion $A \vee B$ zu beweisen, kann man entweder A zeigen oder B . Man ersetzt also die zu beweisende Aussage $A \vee B$ auf dem Spielbrett durch A oder durch B . Man beachte: Es ist hier möglich, einen „falschen Zug“ zu machen. Es kann sein, dass man zum Beispiel A nicht unter den gegebenen Annahmen beweisen kann. In diesem Fall führt das Ersetzen von $A \vee B$ durch A in eine Sackgasse.
- **Negation:** Eine Negation $\neg A$ ist äquivalent zu $A \Rightarrow \perp$. Man zeigt also die Negation von A , indem man unter der Annahme von A einen Widerspruch ($\perp$) herbeiführt. $\diamond$

Trotz ihres Namens bewirken also die Einführungsregeln, dass man die entsprechende Verknüpfung aus der zu beweisenden Aussage *entfernt*. Das kommt daher, dass man diese Regeln gewissermaßen „rückwärts“ anwendet: Statt erst A und B separat zu beweisen und dann auf $A \wedge B$ zu schließen, führt man den Beweis von $A \wedge B$ auf die Beweise von A und B zurück.



Man kann die Einführungsregeln aber auch dazu verwenden, aus vorhandenen Annahmen neue zu generieren. Hat man zum Beispiel A und B unter den Annahmen, dann kann man $A \wedge B$ hinzufügen. Hat man A , dann kann man $A \vee B$ (oder auch $B \vee A$) hinzufügen.

2.5. Beispiele. Hier sind zwei Beispiele für Beweise mit den oben aufgezählten Regeln. Wir schreiben in Klammern „elim($\wedge$)“ oder „intro($\Rightarrow$)“, wenn wir die Eliminationsregel für $\wedge$ bzw. die Einführungsregel für $\Rightarrow$ verwenden (und analog für die anderen Verknüpfungen). Im Folgenden nennen wir die jeweils zu beweisende Aussage das *Ziel*.

BSP
Beweise

- (1) Wir beweisen die Implikation $(A \wedge B) \Rightarrow (A \vee B)$ für beliebige Aussagen A und B („wenn beide Aussagen gelten, dann gilt wenigstens eine“).
 - Zu Beginn ist das Ziel $(A \wedge B) \Rightarrow (A \vee B)$.
 - Um das zu zeigen, nehmen wir $A \wedge B$ an (fügen diese Aussage also zu den bisher noch nicht vorhandenen Annahmen hinzu) und ersetzen das Ziel durch $A \vee B$ (intro($\Rightarrow$)).
 - Mit der Annahme $A \wedge B$ haben wir auch A (und B) (elim($\wedge$)).
 - Wir können das Ziel ersetzen durch A (intro($\vee$)).
 - Da A auch unter den Annahmen ist, ist der Beweis vollständig (das Spielbrett kann abgeräumt werden).
 (Wir hätten auch mit B statt A argumentieren können; das ist in diesem Fall egal.)
- (2) Als zweites Beispiel beweisen wir $((A \Rightarrow B) \wedge (B \Rightarrow C)) \Rightarrow (A \Rightarrow C)$ (wenn B aus A folgt und C aus B folgt, dann folgt C aus A).
 - Wir beginnen mit dem Ziel $((A \Rightarrow B) \wedge (B \Rightarrow C)) \Rightarrow (A \Rightarrow C)$.
 - Um diese Implikation zu zeigen, nehmen wir $(A \Rightarrow B) \wedge (B \Rightarrow C)$ an und ersetzen das Ziel durch $A \Rightarrow C$ (intro($\Rightarrow$)).
 - Wir können dann auch $A \Rightarrow B$ und $B \Rightarrow C$ zu den Annahmen hinzufügen (elim($\wedge$)).
 - Um $A \Rightarrow C$ zu zeigen, nehmen wir A an und zeigen C (intro($\Rightarrow$)).
 - Da wir jetzt sowohl $A \Rightarrow B$ als auch A unter den Annahmen haben, können wir B hinzufügen (elim($\Rightarrow$)).
 - Jetzt haben wir sowohl $B \Rightarrow C$ als auch B unter den Annahmen, also können wir C hinzufügen (elim($\Rightarrow$)).
 - Damit befindet sich das Ziel unter den Annahmen, und der Beweis ist vollständig. ♣

Zusammengesetzte Aussagen, die stets wahr sind, wie in den Beispielen oben, heißen *Tautologien*. Solche Aussagen (oder allgemeiner, Aussagen, die wir bereits bewiesen haben) können wir jederzeit zu den aktuellen Annahmen hinzufügen. Eine solche Aussage ist die Tautologie $A \vee \neg A$ (für eine beliebige Aussage A), auch als die *Regel vom ausgeschlossenen Dritten* (eine Aussage ist wahr oder falsch; eine dritte Möglichkeit gibt es nicht) bekannt. Wir können das dann mit der Fallunterscheidung (also der Eliminationsregel für die Disjunktion) kombinieren.

DEF
Tautologie

Wir geben weitere Beispiele, in denen auch die Negation vorkommt. Dabei werden wir diese Art der Fallunterscheidung verwenden.

2.6. Beispiele.

BSP
Regeln von
de Morgan

- (1) Die Konjunktion $A \wedge B$ ist genau dann falsch, wenn wenigstens eine der Aussagen A und B falsch ist: $\neg(A \wedge B) \Leftrightarrow (\neg A \vee \neg B)$.

Wir zerlegen diese Äquivalenz in zwei Implikationen (intro($\Leftrightarrow$)). Wir beginnen mit der Richtung „von rechts nach links“:

$$(\neg A \vee \neg B) \Rightarrow \neg(A \wedge B)$$

Dazu nehmen wir $\neg A \vee \neg B$ an; das neue Ziel ist $\neg(A \wedge B)$ (intro($\Rightarrow$)). Dafür nehmen wir $A \wedge B$ an und müssen einen Widerspruch ($\perp$) herbeiführen (intro($\neg$)). Aus der Annahme $A \wedge B$ bekommen wir A und B (elim($\wedge$)). Die Eliminationsregel der Disjunktion (Fallunterscheidung) für $\neg A \vee \neg B$ führt dann auf *zwei* Spielbretter jeweils mit dem Ziel $\perp$ und einmal mit den Annahmen $\neg A$, A und B , zum anderen mit den Annahmen $\neg B$, A und B . In beiden Fällen haben wir eine Aussage und ihre Negation in den Annahmen, sodass wir $\perp$ zu den Annahmen hinzufügen können (elim($\neg$)) und damit beide Spiele „gewinnen“.

Jetzt betrachten wir die Richtung „von links nach rechts“:

$$\neg(A \wedge B) \Rightarrow (\neg A \vee \neg B)$$

Wir nehmen also $\neg(A \wedge B)$ an; das neue Ziel ist $\neg A \vee \neg B$ (intro($\Rightarrow$)). Aber jetzt stecken wir fest, denn wir haben keine Regel, mit der wir hier sinnvoll weitermachen können: Wir müssten ja entweder $\neg A$ oder $\neg B$ zeigen. Um weiterzukommen, verwenden wir die oben beschriebene Regel: Wir nehmen die Aussage $A \vee \neg A$ zu den Annahmen hinzu und wenden anschließend die Eliminationsregel der Disjunktion an. Das ergibt zwei Spielbretter mit Ziel $\neg A \vee \neg B$ und den Annahmen $\neg(A \wedge B)$ und A bzw. den Annahmen $\neg(A \wedge B)$ und $\neg A$. Das zweite Spielbrett wird gelöst durch die Einführungsregel der Disjunktion: Wir ersetzen das Ziel durch $\neg A$, was eine unserer Annahmen ist.

Im ersten Spielbrett nehmen wir die Aussage $B \vee \neg B$ zu den Annahmen hinzu und machen dann wieder eine Fallunterscheidung. Das liefert wieder zwei Spielbretter mit dem Ziel $\neg A \vee \neg B$ und den Annahmen $\neg(A \wedge B)$, $\neg A$ und B bzw. den Annahmen $\neg(A \wedge B)$, $\neg A$ und $\neg B$. Das zweite Brett erledigt sich wie vorher, indem wir das Ziel durch $\neg B$ ersetzen. Im ersten Brett wenden wir eine Version der Einführungsregel der Konjunktion an: Da wir $A \wedge B$ aus A und B beweisen können, können wir die beiden Annahmen A und B durch $A \wedge B$ ersetzen. Damit haben wir eine Aussage $(A \wedge B)$ und ihre Negation unter den Annahmen, woraus wir (zunächst $\perp$ (elim($\neg$)) und dann) alles beweisen können.

- (2) Analog gilt, dass die Disjunktion $A \vee B$ genau dann falsch ist, wenn A und B beide falsch sind: $\neg(A \vee B) \Leftrightarrow (\neg A \wedge \neg B)$.

Die Richtung „von links nach rechts“: wir nehmen $\neg(A \vee B)$ an und wollen $\neg A \wedge \neg B$ zeigen (intro($\Rightarrow$)). Dafür müssen wir (immer unter der Annahme $\neg(A \vee B)$) sowohl $\neg A$ als auch $\neg B$ beweisen (intro($\wedge$)). Dazu nehmen wir A (bzw. B) an; das Ziel wird zu $\perp$ (intro($\neg$)). Aus A folgt $A \vee B$, ebenso mit B ; daher können wir jeweils $A \vee B$ zu den Annahmen hinzufügen, womit wir widersprüchliche Annahmen erhalten. Damit ist der Beweis erfolgreich abgeschlossen.

Der Beweis der Gegenrichtung ist eine Übungsaufgabe. ♣

Sobald wir eine Äquivalenz von zwei Aussagen allgemein bewiesen haben, können wir jederzeit in einem Beweis eine der beiden durch die andere ersetzen: Aus der Äquivalenz $A \Leftrightarrow B$ folgt (z.B.) die Implikation $A \Rightarrow B$. Wenn also $A \Leftrightarrow B$

allgemein gilt, dann können wir $A \Rightarrow B$ zu unseren Annahmen hinzufügen; wenn A bereits unter den Annahmen ist, dann bekommen wir auch B ($\text{elim}(\Rightarrow)$).

Die Regel des ausgeschlossenen Dritten erlaubt auch einen Beweis der Äquivalenz $A \Leftrightarrow \neg\neg A$ (Übung). Das ergibt das Prinzip des *indirekten Beweises*: Um eine Aussage A zu zeigen, genügt es, aus ihrer Verneinung $\neg A$ einen Widerspruch abzuleiten.

DEF
indirekter
Beweis

2.7. Beispiel. Eine weitere sehr wichtige solche Äquivalenz ist

$$(A \Rightarrow B) \Leftrightarrow (\neg B \Rightarrow \neg A).$$

BSP
Kontraposition

Man sagt, $\neg B \Rightarrow \neg A$ sei die *Kontraposition* von $A \Rightarrow B$. Man kann also diese Implikation auch beweisen, indem man $\neg A$ unter der Annahme von $\neg B$ zeigt. (Oder indem man einen Widerspruch aus den Annahmen A und $\neg B$ ableitet.)

DEF
Kontraposition

Wir müssen aber erst einmal die obige Äquivalenz beweisen. Für die Richtung „von links nach rechts“ nehmen wir erst $A \Rightarrow B$, dann zusätzlich $\neg B$ und schließlich A an; das Ziel ist dann $\perp$ (zweimal $\text{intro}(\Rightarrow)$, dann $\text{intro}(\neg)$). Aus A und $A \Rightarrow B$ ergibt sich B ($\text{elim}(\Rightarrow)$) und aus B und $\neg B$ dann der gewünschte Widerspruch ($\text{elim}(\neg)$).

Für den Beweis der Gegenrichtung nehmen wir erst $\neg B \Rightarrow \neg A$ und dann A an; das Ziel ist dann B (zweimal $\text{intro}(\Rightarrow)$). An dieser Stelle verwenden wir das Prinzip des indirekten Beweises: Wir fügen $\neg B$ zu den Annahmen hinzu und ersetzen das Ziel durch $\perp$. Das „Endspiel“ des Beweises ist dann sehr ähnlich wie für die andere Richtung. ♣

In der Praxis (und sehr bald auch in dieser Vorlesung) schreibt man Beweise natürlich nicht in der Form auf, dass man genau sagt, welche Regeln man verwendet, sondern man gibt an, was man annimmt und was daraus jeweils folgt. Den einzelnen Regeln entsprechen dabei gewissen Formulierungen. Der Beweis für die Äquivalenz im letzten Beispiel könnte wie folgt formuliert sein:

Wir wollen die Äquivalenz „ $(A \Rightarrow B) \Leftrightarrow (\neg B \Rightarrow \neg A)$ “ zeigen.

„ $\Rightarrow$ “ (Richtung von links nach rechts): Wir nehmen $A \Rightarrow B$ und $\neg B$ an; zu zeigen ist dann $\neg A$. Dafür nehmen wir auch A an; wir müssen dann einen Widerspruch ableiten. Aus $A \Rightarrow B$ und A folgt B , und aus B und $\neg B$ ergibt sich der gewünschte Widerspruch.

„ $\Leftarrow$ “ (Richtung von rechts nach links): Wir nehmen $\neg B \Rightarrow \neg A$ und A an; zu zeigen ist dann B . Um B zu zeigen, nehmen wir zusätzlich $\neg B$ an und leiten einen Widerspruch ab. Aus $\neg B \Rightarrow \neg A$ und $\neg B$ folgt dann $\neg A$, und aus $\neg A$ und A ergibt sich der gewünschte Widerspruch.

Im späteren Verlauf dieser Vorlesung und in fortgeschritteneren Vorlesungen werden dann auch zunehmend als „offensichtlich“ betrachtete Beweisschritte nicht mehr explizit erwähnt. Letzten Endes muss aber jeder Beweis sich als ein gewonnenes Spiel nach den oben beschriebenen Regeln formulieren lassen, um gültig zu sein.

Manche Teile dieses Skripts sind — wie hier — kleiner gedruckt. Dort wird Material behandelt, das über den eigentlichen Stoff der Vorlesung hinausgeht, aber vielleicht für den Einen oder die Andere von Ihnen interessant ist.

Hier geht es zunächst um die Frage, ob es schwierig ist, für eine gegebene aussagenlogische Formel zu entscheiden, ob sie eine Tautologie ist. Im Prinzip kann man das mit Hilfe

einer sogenannten „Wahrheitstafel“ herausfinden: Man betrachtet alle möglichen Kombinationen von „wahr“ und „falsch“ für die beteiligten Elementaraussagen $A, B, C, \dots$ und ermittelt daraus schrittweise den Wahrheitswert der zusammengesetzten Aussage. Ist dieser in allen Fällen „wahr“, dann hat man eine Tautologie, sonst nicht. Allerdings gibt es, wenn n verschiedene elementare Aussagen (wie A und B oben) beteiligt sind, 2^n mögliche Kombinationen von Wahrheitswerten, die überprüft werden müssen. Diese Zahl wächst sehr schnell mit n : $2^{100} = 1\,267\,650\,600\,228\,229\,401\,496\,703\,205\,376$ („exponentielles Wachstum“), sodass es praktisch unmöglich ist, alles durchzuprobieren. Man kann sich aber fragen, ob es auch eine einigermaßen effiziente (also mit vertretbarem Aufwand durchführbare) Methode gibt. Solche Fragen werden von der *Komplexitätstheorie* studiert, die im Bereich zwischen mathematischer Logik und theoretischer Informatik angesiedelt ist. Im vorliegenden Fall ist die Antwort „wahrscheinlich Nein“: Das eng verwandte Erfüllbarkeitsproblem ist *NP-vollständig* (eine Aussage ist genau dann nicht erfüllbar, wenn ihre Negation eine Tautologie ist), und für solche Probleme sind keine effizienten Lösungsverfahren (Algorithmen) *bekannt*. Die Frage danach, ob es tatsächlich keine *gibt*, ist der Inhalt des „ $P = NP?$ “-Problems, für dessen Lösung man eine Million Dollar bekommen würde.

Die zweite Bemerkung in diesem klein gedruckten Abschnitt ist, dass man das Beweisen tatsächlich als Spiel betreiben kann: Es gibt Computerprogramme, mit deren Hilfe man mathematische Beweise konstruieren kann, die das System auf Korrektheit überprüft, und das lässt sich auch als ein Computerspiel umsetzen. Ein solches System ist Lean; auf dem Lean Game Server ($\leftarrow$ das ist ein Link!) lassen sich einige solche Spiele spielen. Probieren Sie es aus!

2.B Mengen.

Ich setze voraus, dass Sie in der Schule gelernt haben, mit *Mengen* umzugehen. Daher werde ich mich auf eine kurze Wiederholung bzw. Einführung von Schreibweisen und grundlegenden Operationen und Rechenregeln beschränken.

Endliche Mengen können durch Aufzählung ihrer Elemente angegeben werden:

$$\{1, 2, 4, 8\}, \quad \{\{1\}, \{1, 2\}\}.$$

Beachte: Die Elemente einer Menge können selbst wieder Mengen sein. Sehr wichtig ist die *leere Menge*, die $\{\}$ geschrieben werden kann. Es ist aber die Schreibweise $\emptyset$ allgemein gebräuchlich; wir werden uns ebenfalls daran halten. Wir schreiben „ $x \in M$ “ für die Aussage „ x ist Element der Menge M “ und „ $x \notin M$ “ für ihre Negation. Zum Beispiel ist $x \in \emptyset$ stets falsch, da die leere Menge keine Elemente hat. Zwei Mengen sind *gleich*, wenn sie dieselben Elemente haben. Insbesondere kommt es nicht darauf an, wie oft man ein Element aufführt:

$$\{1, 3, 2, 1, 2, 2\} = \{1, 2, 3\}.$$

Man kann Mengen auch durch Angabe der Eigenschaften beschreiben, die ihre Elemente haben, wie zum Beispiel

$$\{n \mid n \text{ ist Primzahl}\}.$$

(Statt des senkrechten Strichs „ $|$ “ ist auch ein Doppelpunkt „ $:$ “ gebräuchlich.) Es gibt Symbole für gewisse häufig benötigte Mengen, wie

- die Menge $\mathbb{N} = \{0, 1, 2, 3, \dots\}$ der *natürlichen Zahlen*,
- die Menge $\mathbb{Z} = \{\dots, -3, -2, -1, 0, 1, 2, 3, \dots\}$ der *ganzen Zahlen*,
- die Menge $\mathbb{Q} = \{\frac{a}{b} \mid a, b \in \mathbb{Z}, b \neq 0\}$ der *rationalen Zahlen* und
- die Menge $\mathbb{R}$ der *reellen Zahlen*.

DEF
 $\emptyset, x \in M$

DEF
 $M = N$

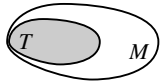
DEF
 $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$

Warnung. Die Definition von $\mathbb{N}$ ist in der Literatur nicht einheitlich; sehr häufig wird auch $\mathbb{N} = \{1, 2, 3, \dots\}$ (also ohne die Null) gesetzt. Hier gibt es kein Richtig oder Falsch; letzten Endes ist das eine Geschmacksfrage. Für mich sind die natürlichen Zahlen gerade die Mächtigkeiten (die *Mächtigkeit* einer Menge ist die Anzahl ihrer Elemente) von endlichen Mengen, und da die leere Menge endlich ist, sollte auch die Null eine natürliche Zahl sein.



2.8. Definition. Eine Menge T heißt *Teilmenge* der Menge M , kurz $T \subset M$, wenn jedes Element von T auch ein Element von M ist. Man beachte, dass der Fall $T = M$ hier erlaubt ist. Um auszudrücken, dass T eine *echte* Teilmenge von M ist (also Teilmenge von M , aber nicht ganz M), schreiben wir $T \subsetneq M$. Statt $M \subset N$ kann man auch $N \supset M$ schreiben. Die Teilmengenbeziehung heißt auch *Inklusion*. $\diamond$

DEF
Teilmenge



Warnung. Die Schreibweise wird in der Literatur nicht einheitlich verwendet; oft findet man $T \subset M$ für echte Teilmengen und $T \subseteq M$ für beliebige Teilmengen. Machen Sie sich solche Unterschiede bewusst, wenn Sie Lehrbücher benutzen!



2.9. Beispiele. Einfache Beispiele von Teilmengen sind die leere Menge, die Teilmenge *jeder* Menge ist: $\emptyset \subset M$, und natürlich ist jede Menge Teilmenge von sich selbst: $M \subset M$. Für die oben eingeführten Zahlenmengen haben wir die Beziehungen $\mathbb{N} \subsetneq \mathbb{Z} \subsetneq \mathbb{Q} \subsetneq \mathbb{R}$. $\clubsuit$

BSP
Teilmengen

2.10. Definition. Die Menge aller Teilmengen von M heißt die *Potenzmenge* von M ; wir schreiben

$$\mathcal{P}(M) = \{T \mid T \subset M\}$$

dafür. $\diamond$

DEF
Potenzmenge
 $\mathcal{P}(M)$

2.11. Beispiel. Zum Beispiel gilt

$$\mathcal{P}(\emptyset) = \{\emptyset\}, \quad \mathcal{P}(\{\emptyset\}) = \{\emptyset, \{\emptyset\}\} \quad \text{und} \quad \mathcal{P}(\{1, 2\}) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}. \quad \clubsuit$$

BSP
Potenzmenge

An dieser Stelle gleich noch ein wichtiger Hinweis:

- Man muss sorgfältig zwischen Mengen und ihren Elementen unterscheiden. Zum Beispiel haben „ $a \in M$ “ und „ $a \subset M$ “ völlig verschiedene Bedeutungen.
- Besonders schwer fällt die Unterscheidung zwischen dem Element a und der Einermenge $\{a\}$. Es ist sehr wichtig, sich diese Unterschiede gleich zu Beginn klar zu machen!



Auf der anderen Seite ist $a \in M$ äquivalent zu $\{a\} \subset M$ — wenn man also beides falsch macht, wird es wieder richtig.

Mengen können miteinander verknüpft werden:

2.12. Definition. Seien M und N zwei Mengen.

- (1) Die *Vereinigung* von M und N ist

$$M \cup N = \{x \mid x \in M \vee x \in N\}.$$

- (2) Der *Durchschnitt* oder die *Schnittmenge* von M und N ist

$$M \cap N = \{x \mid x \in M \wedge x \in N\}.$$

Zwei Mengen, deren Durchschnitt leer ist ($M \cap N = \emptyset$), heißen *disjunkt*.

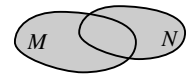
- (3) Die *Mengendifferenz* von M und N ist

$$M \setminus N = \{x \mid x \in M \wedge x \notin N\}.$$

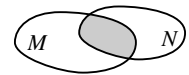
Sie besteht aus den Elementen von M , die keine Elemente von N sind. $\diamond$

DEF

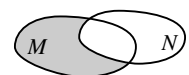
$M \cup N$



$M \cap N$



$M \setminus N$



Für diese Verknüpfungen gelten gewisse Rechenregeln und es gibt Beziehungen zum Begriff der Teilmenge. Ich werde hier einige davon angeben und auch beweisen; andere können Sie sich selbst herleiten, was eine gute erste Übung darstellt, einfache Beweise zu führen. Einige solche Aufgaben finden Sie auf dem Übungsblatt. (Die Farbe rot wird in diesem Skript für mathematische Sätze oder Lemmata (Hilfssätze) verwendet.)

2.13. Satz.

- (1) Es gilt für alle Mengen M und N , dass M genau dann eine Teilmenge von N ist, wenn die Mengen $M \cup N$ und N übereinstimmen:

$$M \subset N \iff M \cup N = N.$$

- (2) Für je zwei Mengen X und Y gilt das „Absorptionsgesetz“

$$(X \cap Y) \cup Y = Y.$$

- (3) Für alle Mengen A, B, C gilt das „Distributivgesetz“

$$(A \cap B) \cup C = (A \cup C) \cap (B \cup C).$$

Es gelten auch die analogen folgenden Aussagen, die wir hier aber nicht beweisen werden:

- (1) Für alle Mengen M und N gilt $M \subset N \iff M \cap N = M$.

- (2) Für je zwei Mengen X und Y gilt $(X \cup Y) \cap Y = Y$.

- (3) Für alle Mengen A, B, C gilt $(A \cup B) \cap C = (A \cap C) \cup (B \cap C)$.

(Die Farbe dunkelrot kennzeichnet Beweise.)

Beweis.

- (1) Zu zeigen ist eine Äquivalenz $A \Leftrightarrow B$. In den meisten Fällen ist es am besten, den Beweis in zwei Teile, nämlich den Beweis von $A \Rightarrow B$ und den Beweis von $B \Rightarrow A$, zu zerlegen.

„ $\Rightarrow$ “: (D.h., wir beweisen die Richtung „von links nach rechts“, also die Aussage $M \subset N \Rightarrow M \cup N = N$.) Wir setzen voraus, dass M eine Teilmenge von N ist, und wir müssen zeigen, dass $M \cup N = N$ ist. Dies ist wiederum eine Äquivalenz, nämlich die Aussage $x \in M \cup N \Leftrightarrow x \in N$. Wir zerlegen den Beweis wieder in zwei Schritte:

SATZ

Eigensch.
Mengen

„ $M \cup N \subset N$ “: Sei $x \in M \cup N$. (Das ist die übliche Formulierung dafür, dass man annimmt, dass $x \in M \cup N$ richtig ist.) Das bedeutet nach Definition $x \in M$ oder $x \in N$. Im ersten Fall ($x \in M$) folgt $x \in N$, da nach Voraussetzung M Teilmenge von N ist. Im zweiten Fall gilt $x \in N$ bereits.

„ $N \subset M \cup N$ “: Das gilt immer, denn aus $x \in N$ folgt $x \in M \vee x \in N$.

Damit ist die Gleichheit $M \cup N = N$ gezeigt und der Beweis der einen Richtung beendet.

„ $\Leftarrow$ “: (Jetzt beweisen wir die Richtung „von rechts nach links“, also die Aussage $M \cup N = N \Rightarrow M \subset N$.) Es gelte $M \cup N = N$. Zu zeigen ist $M \subset N$, also die Implikation $x \in M \Rightarrow x \in N$. Sei also $x \in M$. Dann ist auch $x \in M \cup N$, aber $M \cup N = N$, also folgt $x \in N$. Damit ist gezeigt, dass M eine Teilmenge von N ist.

- (2) Es ist eine Gleichheit von Mengen zu beweisen. Wir haben eben schon gesehen, dass das eine Äquivalenz $a \in (X \cap Y) \cup Y \Leftrightarrow a \in Y$ ist, die man in der Regel am besten in zwei Implikationen aufspaltet. In diesem Fall verwenden wir häufig folgende Kurzschreibweise:

„ $\subset$ “: (Wir beweisen die Inklusion $(X \cap Y) \cup Y \subset Y$.) Sei $a \in (X \cap Y) \cup Y$. Das bedeutet $a \in X \cap Y$ oder $a \in Y$. Der zweite Fall ($a \in Y$) ist gerade, was wir zeigen wollen. Im ersten Fall gilt $a \in X$ und $a \in Y$, also insbesondere (ein Wort, das Mathematiker gerne verwenden) wieder $a \in Y$.

„ $\supset$ “: (Wir beweisen die Inklusion $(X \cap Y) \cup Y \supset Y$, also $Y \subset (X \cap Y) \cup Y$.) Sei $a \in Y$. Dann gilt die schwächere Aussage „ $a \in X \cap Y$ oder $a \in Y$ “, und das bedeutet gerade $a \in (X \cap Y) \cup Y$.

Alternativ kann man auch argumentieren, dass Aussage (2) aus Aussage (1) folgt, wenn man dort $M = X \cap Y$ und $N = Y$ setzt, denn die Inklusion $X \cap Y \subset Y$, die dann auf der linken Seite der Äquivalenz steht, ist immer richtig.

- (3) Wie eben teilen wir den Beweis der Gleichheit zweier Mengen auf in die Teile „ $\subset$ “ und „ $\supset$ “.

„ $\subset$ “: Sei $s \in (A \cap B) \cup C$. Das bedeutet $s \in A \cap B$ oder $s \in C$. Im ersten Fall gilt $s \in A$ und $s \in B$, daraus folgt $s \in A \cup C$ und $s \in B \cup C$ und damit $s \in (A \cup C) \cap (B \cup C)$. Im zweiten Fall (also $s \in C$) gilt ebenfalls $s \in A \cup C$ und $s \in B \cup C$, also folgt auch in diesem Fall $s \in (A \cup C) \cap (B \cup C)$. Damit ist die Inklusion $(A \cap B) \cup C \subset (A \cup C) \cap (B \cup C)$ gezeigt.

„ $\supset$ “: Sei $s \in (A \cup C) \cap (B \cup C)$. Das bedeutet $s \in A \cup C$ und $s \in B \cup C$. Die erste dieser Aussagen heißt $s \in A$ oder $s \in C$. Wenn $s \in C$ ist, dann ist auch $s \in (A \cap B) \cup C$. Wenn $s \notin C$ ist, dann muss $s \in A$ und $s \in B$ sein, also ist $s \in A \cap B$ und damit auch $s \in (A \cap B) \cup C$. $\square$

2.C Prädikatenlogik.

Nun hat man es in der Mathematik nicht nur mit einfachen Aussagen zu tun, die man irgendwie verknüpft, sondern in aller Regel hängen die Aussagen noch von gewissen Parametern oder Variablen ab. Ein typisches Beispiel ist die Aussage „ $x \in M$ “, deren Wahrheitswert davon abhängt, wofür x und M stehen. (Man nennt solche parameterabhängigen Aussagen manchmal *Aussageformen*, weil sie erst dadurch zu einer Aussage mit festgelegtem Wahrheitswert werden, dass man den Parametern Werte zuweist. Auch der Begriff *Prädikat* ist gebräuchlich, was die Bezeichnung „Prädikatenlogik“ erklärt.) Um aus solchen von Variablen abhängigen

DEF
Quantoren
 $\forall, \exists$

Aussagen wiederum Aussagen zu machen, die nicht mehr von (einigen oder allen) Variablen abhängen, gibt es im Wesentlichen zwei Möglichkeiten. Sei dafür $A(x)$ eine (möglicherweise) von der Variablen x abhängige Aussage.

- Wir machen die Aussage „für alle x gilt $A(x)$ “ oder kurz „ $\forall x: A(x)$ “.
- Wir machen die Aussage „es gibt ein x , sodass $A(x)$ gilt“, kurz „ $\exists x: A(x)$ “.

Im Fachjargon spricht man von *Quantifizierung*, da man eine Aussage darüber macht, für wie viele x (alle oder wenigstens eines) $A(x)$ stimmt. In diesem Zusammenhang heißen die Symbole $\forall$ und $\exists$ auch *Quantoren*, genauer *Allquantor* ($\forall$) und *Existenzquantor* ($\exists$).

In der Praxis kommen fast nur die Kombinationen

$$\forall x: x \in M \Rightarrow A(x) \quad \text{und} \quad \exists x: x \in M \wedge A(x)$$

vor, die man dann zu

$$\forall x \in M: A(x) \quad \text{„für alle } x \in M \text{ gilt } A(x)\text{“}$$

und

$$\exists x \in M: A(x) \quad \text{„es gibt ein } x \in M \text{ mit } A(x)\text{“}$$

abkürzt. An der ausführlicheren Form oben erkennt man, dass

$$\forall x \in \emptyset: A(x)$$

immer wahr ist, denn die Voraussetzung $x \in \emptyset$ in der Implikation „ $x \in \emptyset \Rightarrow A(x)$ “ ist falsch. Entsprechend ist

$$\exists x \in \emptyset: A(x)$$

immer falsch, denn es gibt ja kein Element der leeren Menge, also erst recht keines mit zusätzlichen Eigenschaften.

Für den Umgang mit den Quantoren sind folgende Regeln wichtig:

$$\neg \forall x \in M: A(x) \quad \text{ist gleichbedeutend mit} \quad \exists x \in M: \neg A(x)$$

und

$$\neg \exists x \in M: A(x) \quad \text{ist gleichbedeutend mit} \quad \forall x \in M: \neg A(x).$$

Die erste zeigt, wie man eine „Allaussage“ widerlegt: *Man gibt ein Gegenbeispiel an*. Das macht auch verständlich, warum $\forall x \in \emptyset: A(x)$ wahr sein muss: Es gibt kein Gegenbeispiel! Das klingt jetzt vielleicht wie esoterische Spielerei, das ist es aber keineswegs: *Es ist sehr wichtig, Grenzfälle zu verstehen*. Die leere Menge ist ein typischer Grenzfall in vielen Situationen, und nur wenn Sie diesen Grenzfall verstehen, haben Sie die Situation wirklich verstanden! Zum Beispiel gilt die auf den ersten Blick offensichtliche Implikation

$$(\forall x \in M: A(x)) \Rightarrow (\exists x \in M: A(x))$$

(„wenn alle Elemente von M die Eigenschaft A haben, dann hat wenigstens eines diese Eigenschaft“) *nur dann, wenn M nicht die leere Menge ist*. Das zu wissen kann einem manche Falltüren im Beweisgeschäft ersparen.

Eine wichtige Regel ist, dass es auf den Namen der Variablen nicht ankommt: Wenn y in $A(x)$ nicht vorkommt, dann sind

$$\forall x \in M: A(x) \quad \text{und} \quad \forall y \in M: A(y)$$



äquivalent („gebundene Umbenennung“); analog für Existenzaussagen. Sie kennen das vielleicht in ähnlicher Form von Integralen:

$$\int_a^b f(x) dx = \int_a^b f(y) dy.$$

Sehr wichtig ist auch, dass es auf die *Reihenfolge* der Quantoren ankommt: Die Aussagen

$$\forall x \in M \exists y \in N: A(x, y) \quad \text{und} \quad \exists y \in N \forall x \in M: A(x, y)$$

haben unterschiedliche Bedeutung — in der ersten Aussage kann das y , dessen Existenz behauptet wird, von x abhängen, in der zweiten Aussage gibt es *ein festes* y , das für alle x funktionieren muss. Deswegen sollte man sich auch angewöhnen, alle Quantoren systematisch *vor* die Aussage zu stellen, auf die sie sich beziehen. Etwas in der Art von

$$\exists a \in M: A(a, x), \forall x \in N$$

ist unbedingt zu vermeiden! Das gilt entsprechend auch für die Formulierung von Aussagen in Textform. (Allerdings ist es durchaus üblich, in Textform das „für alle x “ auch ans Ende zu stellen; das sollte man aber nur dann tun, wenn es keine anderen Quantoren gibt.)



2.14. Beispiel.

$$\forall x \in \mathbb{N} \exists y \in \mathbb{N}: y > x$$

BSP
 $\forall \exists$ und $\exists \forall$

(„für jede natürliche Zahl x gibt es eine größere natürliche Zahl y “) ist sicher richtig, während

$$\exists y \in \mathbb{N} \forall x \in \mathbb{N}: y > x$$

(„es gibt eine natürliche Zahl y , die größer ist als alle natürlichen Zahlen x “) falsch ist. Die Variante

$$\exists x \in \mathbb{N} \forall y \in \mathbb{N}: y > x$$

(„es gibt eine natürliche Zahl x , die kleiner ist als alle natürlichen Zahlen y “) ist ebenfalls falsch, aber sozusagen nur knapp, denn mit „ $\geq$ “ statt „ $>$ “ wäre sie richtig (mit $x = 0$). ♣

Auf der anderen Seite sind

$$\forall x \in M \forall y \in N: A(x, y) \quad \text{und} \quad \forall y \in N \forall x \in M: A(x, y)$$

äquivalent; deswegen kürzt man das auch gerne in der Form

$$\forall x \in M, y \in N: A(x, y)$$

oder, falls $M = N$,

$$\forall x, y \in M: A(x, y)$$

ab. Das gilt dann analog auch für zwei (oder mehr) Existenzquantoren.

Wie geht man jetzt in Beweisen mit diesen Quantoren um? Auch dafür gibt es Eliminations- und Einführungsregeln.

2.15. **Definition.** Die Eliminationsregeln für die Quantoren sind wie folgt.

DEF
Eliminations-
regeln $\forall/\exists$

- **Allquantor:** Haben wir $\forall x \in M: A(x)$ als Annahme und ist $a \in M$ ein konkretes Element von M , dann können wir die Aussage $A(a)$ zu den Annahmen hinzufügen (man sagt auch, dass man die Allaussage für a *spezialisiert*).
- **Existenzquantor:** Haben wir $\exists x \in M: A(x)$ als Annahme, dann können wir ein unspezifiziertes Element $a \in M$ zusammen mit der Aussage $A(a)$ zu den Annahmen hinzufügen. $\diamond$

2.16. **Definition.** Die Einführungsregeln für die Quantoren sind wie folgt.

DEF
Einführungs-
regeln $\forall/\exists$

- **Allquantor:** Um die Aussage $\forall x \in M: A(x)$ zu beweisen, fügen wir ein unspezifiziertes $a \in M$ zu den Annahmen hinzu und ersetzen die zu beweisende Aussage durch $A(a)$. („Sei $a \in M$ beliebig. ...“)
- **Existenzquantor:** Um die Aussage $\exists x \in M: A(x)$ zu beweisen, geben wir ein konkretes Element $a \in M$ an und zeigen die Aussage $A(a)$. (Das Element a ist ein „Zeuge“ für die Korrektheit der Existenzaussage.) $\diamond$

Bei einem Existenzbeweis muss man sich also überlegen, *welchen* Zeugen man wählt, damit der Beweis funktioniert.

2.17. **Beispiel.** Als Beispiel beweisen wir das „Drinkers’ Paradox“: In einer Bar befindet sich mindestens ein Gast, und jeder Gast trinkt entweder Alkohol oder nicht. Dann gibt es einen Gast in der Bar mit der Eigenschaft, dass genau dann *alle* Gäste Alkohol trinken, wenn dieser eine Gast Alkohol trinkt.

BSP
Drinkers’
Paradox

Um das in Formeln zu fassen, schreiben wir B für die Menge der Gäste der Bar und $A(x)$ für die Aussage „(der Gast) x trinkt Alkohol“. Dann lautet die Behauptung:

$$B \neq \emptyset \Rightarrow \exists g \in B: \left(A(g) \Leftrightarrow \forall x \in B: A(x) \right)$$

Wir nehmen an, dass B nicht leer ist (intro($\Rightarrow$)). Das ist gleichbedeutend mit $\exists b \in B$. Sei also $b \in B$ ein Gast (elim($\exists$)). Wir müssen die Existenzaussage

$$\exists g \in B: \left(A(g) \Leftrightarrow \forall x \in B: A(x) \right)$$

zeigen. Wir brauchen also einen passenden Zeugen $g \in B$ (intro($\exists$)). Dafür machen wir eine Fallunterscheidung: Entweder es trinken alle Gäste Alkohol, oder das ist nicht der Fall. Im ersten Fall wählen wir b als Zeugen. Dann gelten sowohl $A(b)$ als auch $\forall x \in B: A(x)$, also ist die Äquivalenz wahr. Im zweiten Fall verwenden wir die Äquivalenz $\neg \forall x \in B: A(x) \iff \exists x \in B: \neg A(x)$. Sei daher $g \in B$ mit $\neg A(g)$ (elim($\exists$)). Mit diesem g als Zeugen gilt dann weder $A(g)$ noch $\forall x \in B: A(x)$, also ist die Äquivalenz auch in diesem Fall wahr. $\clubsuit$

Einen Beweis einer Existenzaussage mithilfe der Einführungsregel nennt man auch einen *konstruktiven* Existenzbeweis, weil man einen Zeugen explizit angibt oder konstruiert. Alternativ kann man die äquivalente Aussage $\neg \forall x \in M: \neg A(x)$ beweisen, indem man die Annahme, dass kein $x \in M$ die Eigenschaft A hat, zum Widerspruch führt. Dabei muss kein Element von M angegeben werden, das tatsächlich die Eigenschaft A hat.

Zum besseren Verständnis hier ein Beispiel: Sei $N > 1$ eine gegebene natürliche Zahl. Sie wollen beweisen, dass N einen echten Teiler $d > 1$ hat, also die Aussage

$$\exists d, m \in \mathbb{N}: d > 1 \wedge m > 1 \wedge N = d \cdot m.$$

Das kann man natürlich tun, indem man einen echten Teiler d findet. Man kann aber auch versuchen, die Negation der Aussage, nämlich „ N ist Primzahl“ zum Widerspruch

zu führen. Dazu kann man etwa den „kleinen Satz von Fermat“ verwenden, der aussagt, dass für jede Primzahl p und jede ganze Zahl a die Zahl $a^p - a$ durch p teilbar ist. Wenn Sie also eine ganze Zahl a finden, sodass $a^N - a$ *nicht* durch N teilbar ist, dann folgt daraus, dass N keine Primzahl ist, also muss N einen echten Teiler haben, ohne dass Sie einen angeben können. Dieser Unterschied ist durchaus auch praktisch relevant. Es gibt nämlich effiziente Algorithmen, die feststellen, ob N eine Primzahl ist oder nicht, aber es sind bisher keine effizienten Algorithmen bekannt, die eine zusammengesetzte Zahl faktorisieren können.

Weitere Beispiele für Beweise werden in großer Zahl im Lauf der Vorlesung folgen.

An den bisherigen Beispielen von Beweisen können Sie jedenfalls schon sehen, dass *die Struktur der zu beweisenden Aussage* (über die entsprechenden Regeln) *die Struktur des Beweises vorgibt*: Bis zu einem gewissen Grad gibt es auch für das Beweisen Rezepte, die man anwenden kann! Allerdings braucht man an manchen Stellen dann doch eine Idee (welche Fallunterscheidung ist hilfreich? welcher Zeuge funktioniert?); dort wird es dann schwierig, aber auch interessant, und wenn man die richtige Idee hat, kann das sehr befriedigend sein.

2.D Geordnete Paare und kartesische Produkte.

Häufig möchte man mit zwei (oder vielleicht auch mehr, siehe unten) Elementen von verschiedenen Mengen gemeinsam arbeiten, wobei es auf die Reihenfolge ankommt. (Wenn die Reihenfolge keine Rolle spielt, also bei *ungeordneten* Paaren, kann man Zweiermengen $\{a, b\}$ verwenden.) Dazu führt man geordnete Paare ein:

2.18. Definition. Sind a und b Elemente irgendwelcher Mengen, dann steht (a, b) für das daraus gebildete *geordnete Paar*. Die wesentliche Eigenschaft dieser geordneten Paare ist, dass zwei solche Paare genau dann gleich sind, wenn sie in beiden Komponenten übereinstimmen:

DEF
geordnetes
Paar (a, b)

$$(a, b) = (x, y) \iff (a = x \text{ und } b = y). \quad \diamond$$

Man kann geordnete Paare innerhalb der Mengenlehre definieren, indem man

$$(a, b) = \{\{a\}, \{a, b\}\}$$

setzt. (Man beachte den Sonderfall $(a, a) = \{\{a\}\}$.) Man muss dann zeigen, dass die so definierten Paare die obige Eigenschaft haben. Das sollten Sie als Aufforderung begreifen!

2.19. Definition. Sind M und N zwei Mengen, dann schreibt man

DEF
 $M \times N$

$$M \times N = \{(m, n) \mid m \in M, n \in N\}$$

(lies „ M kreuz N “) für die Menge der geordneten Paare, deren erste Komponente aus M und deren zweite Komponente aus N kommt; die Menge $M \times N$ heißt das *kartesische Produkt* der Mengen M und N . $\diamond$

(„Kartesisch“ leitet sich vom latinisierten Namen *Cartesius* des Mathematikers und Philosophen („cogito ergo sum“: „ich denke, also bin ich“) René Descartes ab.)



R. Descartes
(1596–1650)

2.20. Definition. Analog zu Paaren kann man (geordnete) *Tripel* (a, b, c) , *Quadrupel* (a, b, c, d) , *Quintupel* (a, b, c, d, e) , *Sextupel* (a, b, c, d, e, f) und ganz allgemein *n-Tupel* $(a_1, a_2, \dots, a_n)$ einführen und kartesische Produkte mit mehr als zwei Faktoren definieren, zum Beispiel

$$A \times B \times C \times D = \{(a, b, c, d) \mid a \in A, b \in B, c \in C, d \in D\}.$$

Die Gleichheit zweier *n-Tupel* $(a_1, a_2, \dots, a_n)$ und $(b_1, b_2, \dots, b_n)$ definiert man analog zur Gleichheit von Paaren:

$$(a_1, a_2, \dots, a_n) = (b_1, b_2, \dots, b_n) \iff a_1 = b_1 \wedge a_2 = b_2 \wedge \dots \wedge a_n = b_n.$$

Einen wichtigen Spezialfall des allgemeinen kartesischen Produkts erhalten wir, wenn alle beteiligten Mengen übereinstimmen. Dann schreibt man kurz M^2 für $M \times M$, M^3 für $M \times M \times M$ und allgemein

$$M^n = \{(m_1, m_2, \dots, m_n) \mid m_1, m_2, \dots, m_n \in M\}$$

für die Menge der *n-Tupel*, deren Komponenten aus der Menge M kommen. Für $n = 1$ haben wir „1-Tupel“; wir setzen $(a) = a$, damit ist dann $M^1 = M$. Für $n = 0$ gibt es das (eindeutig bestimmte) Nulltupel $()$, und $M^0 = \{()\}$ ist eine einelementige Menge. $\diamond$

2.21. Beispiele. Zum Beispiel ist $\mathbb{R}^2$ die Menge der Paare reeller Zahlen. Wenn man die Komponenten als *x*- und *y*-Koordinate interpretiert, dann kann man $\mathbb{R}^2$ als die Menge der Punkte der Ebene auffassen, und entsprechend $\mathbb{R}^3$ als die Menge der Punkte des (dreidimensionalen) Raumes. Diese Mengen und ihre allgemeinere Form $\mathbb{R}^n$ werden uns bald wieder als Standardbeispiele von „Vektorräumen“ begegnen. $\clubsuit$

2.E Abbildungen.

Der (vorläufig) letzte wichtige Begriff, den wir einführen müssen, ist der der Abbildung zwischen zwei Mengen.

2.22. Definition. Seien M und N zwei Mengen. Dann ist eine *Abbildung* f von M nach N eine Vorschrift, die jedem $x \in M$ ein eindeutig bestimmtes $y \in N$ zuordnet; für dieses y schreiben wir dann $f(x)$ („ f von x “). Wir schreiben

$$f: M \longrightarrow N$$

oder, wenn wir die Abbildungsvorschrift angeben wollen,

$$f: M \longrightarrow N, \quad x \longmapsto f(x),$$

wobei statt „ $f(x)$ “ meistens eine konkrete Formel oder Ähnliches steht. Beachten Sie die beiden unterschiedlichen Pfeile „ $\rightarrow$ “ und „ $\mapsto$ “! Der erste steht zwischen den *Mengen* M und N , der zweite zwischen den *Elementen* $x \in M$ und $f(x) \in N$. $f(x) \in N$ heißt dann das *Bild* von $x \in M$ unter f . Gilt $f(x) = y$ für ein $y \in N$, dann heißt x ein *Urbild* von y unter f . Man beachte: Es ist durchaus möglich, dass ein $y \in N$ kein Urbild oder viele verschiedene Urbilder unter f hat.

Häufig (vor allem in der Analysis) verwendet man auch das Wort *Funktion* für *Abbildung* (was die häufig verwendete Bezeichnung „ f “ für Abbildungen erklärt). $\diamond$

Mit „Vorschrift“ ist hier nicht gemeint, dass das Bild von x unter f durch einen Rechenausdruck oder so etwas gegeben sein muss. Es kommt *nur* darauf an, dass *jedem* $x \in M$ *genau ein* $f(x) \in N$ zugeordnet ist. Man kann sich f als eine

DEF
Tripel, ...
n-Tupel
 $M_1 \times \dots \times M_n$
 M^n

BSP
kartesische
Produkte

DEF
Abbildung

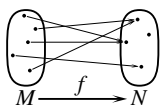


Bild
Urbild



„Black Box“ vorstellen, die einem, wenn man ein $x \in M$ hineinsteckt, ein Element $f(x) \in N$ herausgibt (und zwar für dasselbe x immer dasselbe $f(x)$):

$$M \ni x \longrightarrow \boxed{f} \longrightarrow f(x) \in N$$

(Die Pfeile hier dienen nur der Illustration und sind nicht Teil der oben eingeführten Notation $f: M \rightarrow N, x \mapsto f(x)$.)

2.23. Definition. M heißt die *Definitions- oder Definitionsmenge* oder die *Quelle* der Abbildung $f: M \rightarrow N$, N dementsprechend der *Wertebereich* oder das *Ziel* von f . Wichtig ist dabei, dass zur Angabe einer Abbildung immer auch Quelle und Ziel gehören; die Abbildungsvorschrift alleine genügt nicht.

Zwei Abbildungen f und g sind genau dann gleich (und man schreibt $f = g$), wenn ihre Definitions- und Wertebereiche übereinstimmen und für alle Elemente x des Definitionsbereichs gilt $f(x) = g(x)$: Abbildungen sind (bei gegebenem Definitions- und Wertebereich) durch ihre Werte festgelegt.

Wir schreiben $\text{Abb}(M, N)$ für die Menge aller Abbildungen mit Definitionsbereich M und Wertebereich N . $\diamond$

Der Wertebereich N von f ist von der *Wertemenge* oder *Bildmenge*

$$\{f(x) \mid x \in M\} \subset N$$

von f zu unterscheiden. Die Bildmenge kann eine echte Teilmenge des Wertebereichs sein.

2.24. Beispiele. Beispiele von Abbildungen sind

$$n: \mathbb{R} \longrightarrow \mathbb{R}, \quad x \longmapsto 0$$

(die Nullfunktion; es gilt $n(x) = 0 \in \mathbb{R}$ für alle $x \in \mathbb{R}$),

$$p: \mathbb{R} \longrightarrow \mathbb{R}, \quad x \longmapsto x^3 - 2x^2 + x - 5$$

(eine Polynomfunktion; es gilt zum Beispiel $p(1) = p(0) = -5$),

$$s: \mathbb{R} \longrightarrow \{-1, 0, 1\}, \quad x \longmapsto \begin{cases} 1 & \text{falls } x > 0, \\ 0 & \text{falls } x = 0, \\ -1 & \text{falls } x < 0 \end{cases}$$

(die Vorzeichenfunktion). Für eine beliebige Menge M gibt es die „Einermengenabbildung“

$$e: M \longrightarrow \mathcal{P}(M), \quad x \longmapsto \{x\}.$$

Zum kartesischen Produkt $M \times N$ gehören die *Projektionsabbildungen*

$$\text{pr}_1: M \times N \longrightarrow M, \quad (a, b) \longmapsto a \quad \text{und} \quad \text{pr}_2: M \times N \longrightarrow N, \quad (a, b) \longmapsto b.$$

Ist T eine Teilmenge von M , dann hat man die *Inklusionsabbildung*

$$i: T \longrightarrow M, \quad x \longmapsto x.$$

Für jede Menge X gibt es (als Spezialfall der Inklusionsabbildung) die *identische Abbildung* oder kurz *Identität*

$$\text{id}_X: X \longrightarrow X, \quad x \longmapsto x,$$

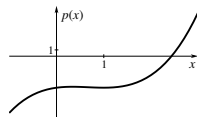
die jedes Element von X auf sich selbst abbildet. Als Grenzfälle haben wir für jede Menge X genau eine Abbildung $\emptyset \rightarrow X$; eine Abbildung $X \rightarrow \emptyset$ gibt es jedoch nur dann, wenn X selbst die leere Menge ist, denn wenn X ein Element x hat,

DEF
Definitions-,
Wertebereich
 $f = g$
 $\text{Abb}(M, N)$



DEF
Bildmenge

BSP
Abbildungen



DEF
Projektion

DEF
Inklusionsabb.

DEF
Identität id_X

könnte es auf kein Element abgebildet werden (denn die leere Menge hat keine Elemente). ♣

Wenn Sie den Begriff „Vorschrift“, den wir oben verwendet haben, zu schwammig finden, dann erfahren Sie hier, wie man den Abbildungsbegriff auf eine solide Grundlage stellen kann. Man greift dazu auf die Mengenlehre zurück und identifiziert eine Abbildung $f: M \rightarrow N$ mit ihrem *Graphen*

$$\Gamma(f) = \{(x, f(x)) \mid x \in M\} \subset M \times N.$$

(Das verallgemeinert die Funktionsgraphen von Funktionen $\mathbb{R} \rightarrow \mathbb{R}$, die Sie aus der Schule kennen.) Dann kann man sagen, dass eine Teilmenge $F \subset M \times N$ genau dann einer Abbildung $f: M \rightarrow N$ entspricht, wenn die Bedingungen

$$\forall x \in M \exists y \in N: (x, y) \in F$$

und

$$\forall x \in M \forall y_1, y_2 \in N: ((x, y_1) \in F \wedge (x, y_2) \in F) \Rightarrow y_1 = y_2$$

erfüllt sind. Die erste Bedingung drückt aus, dass *jedes* $x \in M$ auf ein Element von N abgebildet werden muss, und die zweite Bedingung sagt, dass es *höchstens ein* solches Element von N gibt.

Es gibt gewisse wichtige Eigenschaften, die eine Abbildung haben kann oder nicht.

2.25. Definition. Sei $f: M \rightarrow N$ eine Abbildung.

DEF
injektiv
surjektiv
bijektiv

- (1) f heißt *injektiv* oder eine *Injektion*, wenn f keine zwei verschiedenen Elemente von M auf dasselbe Element von N abbildet:

$$\forall x_1, x_2 \in M: f(x_1) = f(x_2) \Rightarrow x_1 = x_2$$

(So weist man auch nach, dass f injektiv ist: Man nimmt an, zwei Elemente hätten dasselbe Bild unter f und zeigt dann, dass diese beiden Elemente gleich sein müssen.)

- (2) f heißt *surjektiv* oder eine *Surjektion* (das „sur“ ist französisch für „auf“, daher ist die korrekte Aussprache „ßür“), wenn jedes Element von N als Bild unter f eines Elements von M auftritt:

$$\forall y \in N \exists x \in M: f(x) = y$$

- (3) f heißt *bijektiv* oder eine *Bijektion*, wenn f sowohl injektiv als auch surjektiv ist. ◇

Man kann das auch so ausdrücken:

- f ist genau dann injektiv, wenn jedes Element von N *höchstens* ein Urbild unter f hat.
- f ist genau dann surjektiv, wenn jedes Element von N *mindestens* ein Urbild unter f hat.
- f ist genau dann bijektiv, wenn jedes Element von N *genau* ein Urbild unter f hat.

2.26. Definition. Sei $f: M \rightarrow N$ bijektiv. Wir definieren dann eine Abbildung $f^{-1}: N \rightarrow M$ dadurch, dass wir für $f^{-1}(y)$ das eindeutig bestimmte $x \in M$ mit $f(x) = y$ nehmen. Diese Abbildung f^{-1} heißt die *Umkehrabbildung* oder *inverse Abbildung* von f . Eine bijektive Abbildung $f: X \rightarrow X$ heißt auch eine *Permutation* von X . ◇

DEF
Umkehrabb.
Permutation

2.27. Beispiele. Wir schreiben $\mathbb{R}_{\geq 0}$ für die Menge $\{x \in \mathbb{R} \mid x \geq 0\}$ der nichtnegativen reellen Zahlen. Dann gilt:

BSP
injektiv
surjektiv

- (1) $f_1: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$, ist weder injektiv noch surjektiv, denn es gilt zum Beispiel $f_1(1) = f_1(-1) = 1$ und $-1 \in \mathbb{R}$ hat kein Urbild.
- (2) $f_2: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}, x \mapsto x^2$, ist injektiv, aber nicht surjektiv.
- (3) $f_3: \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}, x \mapsto x^2$, ist surjektiv, aber nicht injektiv.
- (4) $f_4: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}, x \mapsto x^2$, ist bijektiv.
Die Umkehrabbildung ist $f_4^{-1}: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}, x \mapsto \sqrt{x}$.

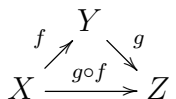
Daran sieht man auch sehr schön, dass Definitions- und Wertebereich wesentlich für eine Abbildung sind. Weitere allgemeine Beispiele sind:

- (5) Für jede Menge M ist die identische Abbildung id_M bijektiv.
- (6) Für jede Menge M ist die „leere Abbildung“ $\emptyset \rightarrow M$ injektiv.
- (7) Jede Abbildung $\{a\} \rightarrow M$ ist injektiv.
- (8) Eine Abbildung $M \rightarrow \{a\}$ ist genau dann surjektiv, wenn M nicht leer ist.
- (9) Die Einermengenabbildung $e: M \rightarrow \mathcal{P}(M)$ ist injektiv, aber nicht surjektiv (Letzteres, weil zum Beispiel die leere Menge kein Urbild hat).
- (10) Die Projektionsabbildung $\text{pr}_1: M \times N \rightarrow M$ ist surjektiv, wenn $N \neq \emptyset$ ist (sei $b \in N$; dann hat jedes $a \in M$ das Urbild (a, b)). Entsprechend gilt, dass $\text{pr}_2: M \times N \rightarrow N$ surjektiv ist, wenn $M \neq \emptyset$ ist. ♣

Abbildungen können verknüpft werden, indem man sie „hintereinanderschaltet“:

2.28. Definition. Sind $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ Abbildungen, sodass der Wertebereich von f mit dem Definitionsbereich von g übereinstimmt, dann kann man die zusammengesetzte Abbildung $g \circ f: X \rightarrow Z$ bilden, die $x \in X$ auf $g(f(x)) \in Z$ abbildet:

DEF
Komposition



$$x \longrightarrow \boxed{f} \longrightarrow f(x) \longrightarrow \boxed{g} \longrightarrow g(f(x)) \quad \diamond$$

Man muss sich merken, dass in $g \circ f$ die Abbildung f *zuerst* ausgeführt wird, obwohl sie hinter g steht. Die Sprechweise „ g nach f “ für $g \circ f$ hilft dabei.



Diese Verknüpfung oder *Komposition* von Abbildungen hat einige wichtige Eigenschaften:

2.29. Satz.

SATZ
Eigensch.
Abbildungen

- (1) Sind $f: W \rightarrow X$, $g: X \rightarrow Y$ und $h: Y \rightarrow Z$ Abbildungen, dann gilt $(h \circ g) \circ f = h \circ (g \circ f)$. Man lässt deswegen meistens die Klammern weg und schreibt $h \circ g \circ f$.
- (2) Ist $f: X \rightarrow Y$ eine Abbildung, dann gilt
$$f \circ \text{id}_X = f \quad \text{und} \quad \text{id}_Y \circ f = f.$$
- (3) Sind $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ injektive Abbildungen, dann ist auch $g \circ f: X \rightarrow Z$ injektiv.
- (4) Sind $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ surjektive Abbildungen, dann ist auch $g \circ f: X \rightarrow Z$ surjektiv.

- (5) Ist $f: X \rightarrow Y$ bijektiv mit Umkehrabbildung $f^{-1}: Y \rightarrow X$, dann gilt
- $$f^{-1} \circ f = \text{id}_X \quad \text{und} \quad f \circ f^{-1} = \text{id}_Y .$$
- (6) Sind $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ zwei Abbildungen, dann gilt
- $$g \circ f \text{ injektiv} \implies f \text{ injektiv} \quad \text{und} \quad g \circ f \text{ surjektiv} \implies g \text{ surjektiv}.$$
- (7) Ist $f: X \rightarrow Y$ eine Abbildung, dann ist f genau dann injektiv, wenn X leer ist oder es eine Abbildung $g: Y \rightarrow X$ mit $g \circ f = \text{id}_X$ gibt.
- (8) Ist $f: X \rightarrow Y$ eine Abbildung, dann ist f genau dann surjektiv, wenn es eine Abbildung $g: Y \rightarrow X$ gibt mit $f \circ g = \text{id}_Y$.
- (9) Ist $f: X \rightarrow Y$ eine Abbildung, dann ist f genau dann bijektiv, wenn es eine Abbildung $g: Y \rightarrow X$ gibt mit $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$.

Man kann also auf (mindestens) zwei verschiedene Arten beweisen, dass eine Abbildung $f: X \rightarrow Y$ bijektiv ist:



- Man weist nach, dass f injektiv und surjektiv ist, oder
- man findet einen Kandidaten g für die Umkehrabbildung und rechnet nach, dass $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$ ist.

In vielen Fällen ist die zweite Methode einfacher durchzuführen.

Beweis.

- (1) Erst einmal ist klar, dass die Abbildungen den gemeinsamen Definitionsbereich W und den gemeinsamen Wertebereich Z haben. Die Aussage „ $(h \circ g) \circ f = h \circ (g \circ f)$ “ bedeutet dann

$$\forall w \in W: ((h \circ g) \circ f)(w) = (h \circ (g \circ f))(w).$$

Sei also $w \in W$. Dann ist

$$((h \circ g) \circ f)(w) = (h \circ g)(f(w)) = h(g(f(w)))$$

und ebenso

$$(h \circ (g \circ f))(w) = h((g \circ f)(w)) = h(g(f(w))),$$

also gilt die behauptete Gleichheit für w . Da $w \in W$ beliebig war, gilt die Gleichheit für alle $w \in W$.

$$w \longrightarrow \boxed{f} \longrightarrow f(w) \longrightarrow \boxed{g} \longrightarrow g(f(w)) \longrightarrow \boxed{h} \longrightarrow h(g(f(w)))$$

- (2) In beiden Fällen haben alle beteiligten Abbildungen denselben Definitionsbereich X und denselben Wertebereich Y . Für $x \in X$ gilt

$$(f \circ \text{id}_X)(x) = f(\text{id}_X(x)) = f(x),$$

also ist $f \circ \text{id}_X = f$, und

$$(\text{id}_Y \circ f)(x) = \text{id}_Y(f(x)) = f(x),$$

also ist auch $\text{id}_Y \circ f = f$.

- (3) Übung.

- (4) Übung.

- (5) Die Definitions- und Wertebereiche stimmen jeweils überein. Für $x \in X$ gilt $f^{-1}(f(x)) = x = \text{id}_X(x)$ nach Definition der Umkehrabbildung, also ist $f^{-1} \circ f = \text{id}_X$. Für $y \in Y$ gilt $f(f^{-1}(y)) = y = \text{id}_Y(y)$ ebenfalls nach Definition der Umkehrabbildung, also ist $f \circ f^{-1} = \text{id}_Y$.
- (6) Wir nehmen an, dass $g \circ f$ injektiv ist; wir müssen zeigen, dass auch f injektiv ist. Seien dazu $x_1, x_2 \in X$ mit $f(x_1) = f(x_2)$. Dann folgt

$$(g \circ f)(x_1) = g(f(x_1)) = g(f(x_2)) = (g \circ f)(x_2),$$

und weil $g \circ f$ injektiv ist, muss $x_1 = x_2$ sein. Damit ist gezeigt, dass f injektiv ist.

Jetzt nehmen wir an, dass $g \circ f$ surjektiv ist; wir müssen zeigen, dass auch g surjektiv ist. Sei dazu $z \in Z$. Da nach Voraussetzung $g \circ f$ surjektiv ist, gibt es $x \in X$ mit $(g \circ f)(x) = z$. Das heißt aber $g(f(x)) = z$, also gilt mit $y = f(x) \in Y$ auch $g(y) = z$. Das zeigt, dass g surjektiv ist.

- (7) Zu zeigen ist die Äquivalenz

$$f: X \rightarrow Y \text{ injektiv} \iff (X = \emptyset \vee \exists g \in \text{Abb}(Y, X): g \circ f = \text{id}_X).$$

„ $\Rightarrow$ “: Wir nehmen an, f sei injektiv. Wenn X leer ist, dann gilt die rechte Seite. Wenn X nicht leer ist, dann sei $x_0 \in X$ irgendein Element. Wir konstruieren eine passende Abbildung $g: Y \rightarrow X$ wie folgt: Sei $y \in Y$. Wenn es ein $x \in X$ gibt mit $f(x) = y$, dann setzen wir $g(y) = x$. Da es (weil f injektiv ist) dann genau ein solches x gibt, ist $g(y)$ eindeutig bestimmt. Wenn es kein $x \in X$ gibt mit $f(x) = y$, dann setzen wir $g(y) = x_0$. Jetzt müssen wir nachprüfen, dass g die geforderte Eigenschaft $g \circ f = \text{id}_X$ hat. Definitions- und Wertebereich beider Seiten stimmen überein, und für $x \in X$ gilt nach Definition von g , dass $(g \circ f)(x) = g(f(x)) = x = \text{id}_X(x)$ ist. Damit ist die Gleichheit der Abbildungen gezeigt.

„ $\Leftarrow$ “: Wenn $X = \emptyset$ ist, dann ist f injektiv. Wenn es $g: Y \rightarrow X$ gibt mit $g \circ f = \text{id}_X$, dann ist f ebenfalls injektiv nach Teil (6), denn id_X ist injektiv.

- (8) „ $\Rightarrow$ “: Ist f surjektiv, dann können wir zu jedem $y \in Y$ ein $x_y \in X$ auswählen mit $f(x_y) = y$ (denn es gibt ja immer mindestens ein Urbild). Wir setzen dann $g(y) = x_y$ und es folgt $f \circ g = \text{id}_Y$.

„ $\Leftarrow$ “: Das folgt aus Teil (6), denn id_Y ist surjektiv.

- (9) „ $\Rightarrow$ “: Ist f bijektiv, dann hat $g = f^{-1}$ die verlangte Eigenschaft nach Teil (5).

„ $\Leftarrow$ “: Nach Teil (7) ist f injektiv und nach Teil (8) auch surjektiv, also bijektiv. $\square$

Wenn man Abbildungen definieren möchte, die von zwei (oder mehr) Elementen möglicherweise verschiedener Mengen abhängen, dann kann man dies unter Zuhilfenahme von kartesischen Produkten tun: Möchte man einem Element von M_1 und einem Element von M_2 ein Element von N zuordnen, so entspricht das einer Abbildung $M_1 \times M_2 \rightarrow N$. Zum Beispiel kann man die Addition reeller Zahlen als eine Abbildung $\mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$, $(x, y) \mapsto x + y$, auffassen. Ist $f: M_1 \times M_2 \rightarrow N$ eine Abbildung, dann schreibt man auch $f(m_1, m_2)$ für $f((m_1, m_2))$. Eine Abbildung der Form $M \times M \rightarrow M$ heißt auch eine *Verknüpfung* auf M .

Schließlich führen wir noch eine weitere Interpretation und Schreibweise für Abbildungen ein, die häufig vorkommt.

DEF
Verknüpfung

2.30. Definition. Wenn $a: I \rightarrow M$ eine Abbildung ist, dann schreibt man dafür auch $(a_i)_{i \in I}$ und nennt das eine *Familie* von Elementen von M mit der *Indexmenge* I . Dabei ist $a_i = a(i)$ der Wert der Abbildung a an der Stelle $i \in I$. Sie kennen das von Folgen $(a_n)_{n \in \mathbb{N}}$. Die n -Tupel, die wir vor einer Weile eingeführt haben, kann man als den Spezialfall $I = \{1, 2, \dots, n\}$ einer solchen Familie betrachten. In Analogie zur Schreibweise M^n für die Menge der n -Tupel mit Komponenten aus M schreibt man auch M^I für die Menge der Familien von Elementen von M mit Indexmenge I . Das ist (bis auf die Schreibweise) nichts anderes als die Menge $\text{Abb}(I, M)$ der Abbildungen von I nach M . $\diamond$

DEF
Familie
 M^I

Der Unterschied zwischen einer *Abbildung* $a: I \rightarrow M$ und einer *Familie* $(a_i)_{i \in I}$ von Elementen von M ist lediglich einer der Perspektive:

- Bei der Abbildung interessiert man sich mehr für den Vorgang der *Zuordnung*.
- Bei der Familie interessiert man sich mehr für die vorkommenden *Werte*. (Man kann eine Familie als „nummerierte Menge“ betrachten; allerdings kann dabei ein Element auch mehrfach (mit verschiedenen Indizes = „Nummern“) vorkommen.)

An dieser Stelle bietet es sich an, etwas mehr zur Mengenlehre zu sagen. Was wir hier betreiben, ist „naive“ Mengenlehre; wir machen uns hier also keine Gedanken darüber, welche Konstruktionen mit Mengen tatsächlich möglich oder erlaubt sind. Das führt normalerweise auch nicht zu Problemen. Sie sollten aber wissen, dass die Mengenlehre durchaus nicht so harmlos ist, wie sie einem zunächst erscheinen mag. Wenn man bei der Bildung von Mengen zu viel erlaubt, kommt man in Schwierigkeiten, wie die berühmte *Russellsche Antinomie* zeigt. Denn dann könnte man die „Menge aller Mengen, die sich nicht selbst als Element enthalten“, also $M = \{x \mid x \notin x\}$ konstruieren. Die Frage, ob M ein Element von M ist, führt auf einen unauflösbaren Widerspruch. (In der Unterhaltungsmathematik gibt es die Variante mit dem Dorfbarbier, der genau die Männer im Dorf rasiert, die sich nicht selbst rasieren. Rasiert sich nun der Barbier oder nicht?) Um diesen Widerspruch zu vermeiden, muss man genaue Regeln formulieren, wie man Mengen konstruieren darf. Das führt zur *axiomatischen Mengenlehre*.



B. Russell
(1872–1970)

Die meisten der Axiome sind recht „harmlos“; sie besagen etwa, dass die leere Menge existiert, dass man Einer- und Zweiermengen bilden kann, dass man immer Teilmengen bilden kann, und dass Vereinigungsmengen und Potenzmengen existieren. Es gibt aber ein Axiom, das *Auswahlaxiom*, das von einigen Mathematikern abgelehnt wurde. Es besagt, dass „es zu jeder Familie nichtleerer Mengen eine Auswahlfunktion gibt“. Genauer: Ist $(X_i)_{i \in I}$ eine Familie von Mengen mit $X_i \neq \emptyset$ für alle $i \in I$, dann gibt es eine *Auswahlfunktion* $f: I \rightarrow X$, wobei $X = \{x \mid \exists i \in I: x \in X_i\}$ die Vereinigung aller Mengen X_i ist (die nach einem der harmlosen Axiome existiert), sodass für jedes $i \in I$ das Bild $f(i)$ ein Element von X_i ist. Die Auswahlfunktion wählt also aus jeder Menge X_i ein Element aus. Wir haben dieses Auswahlaxiom im Beweis von Teil (8) von Satz 2.29 benutzt, als wir für jedes $y \in Y$ ein Urbild x_y ausgewählt haben. Der Grund für die Ablehnung des Auswahlaxioms liegt darin, dass es nicht „konstruktiv“ ist: Es macht eine Existenzaussage („es gibt eine Auswahlfunktion“), sagt aber nicht, *wie* man eine Auswahlfunktion bekommt. Heutzutage vertreten die meisten Mathematiker den pragmatischen Standpunkt, dass das Auswahlaxiom nützlich ist und es deswegen seine Berechtigung hat. Vor allem in der Analysis käme man ohne das Auswahlaxiom nicht weit. Es ist bekannt, dass die Hinzunahme des Auswahlaxioms nicht zu einem Widerspruch in der Mengenlehre führt, wenn die Mengenlehre ohne das Auswahlaxiom widerspruchsfrei ist (allerdings gilt das auch für seine Verneinung).

Zum Abschluss dieses Abschnitts über Grundlagen gibt es hier noch eine Tabelle mit griechischen Buchstaben. Als Mathematiker gehen einem schnell die Buchstaben aus, um die verschiedenen Objekte zu bezeichnen, mit denen man es zu tun hat. Darum wird gerne auf das griechische Alphabet zurückgegriffen.

klein	groß	Name
α	A	Alpha
β	B	Beta
γ	Γ	Gamma
δ	Δ	Delta
ε, ϵ	E	Epsilon
ζ	Z	Zeta
η	H	Eta
θ, ϑ	Θ	Theta

klein	groß	Name
ι	I	Iota
κ	K	Kappa
λ	Λ	Lambda
μ	M	My
ν	N	Ny
ξ	Ξ	Xi
o	O	Omikron
π	Π	Pi

klein	groß	Name
ρ, ϱ	P	Rho
σ	Σ	Sigma
τ	T	Tau
υ	Υ	Ypsilon
ϕ, φ	Φ	Phi
χ	X	Chi
ψ	Ψ	Psi
ω	Ω	Omega

3. ALGEBRAISCHE STRUKTUREN: GRUPPEN, RINGE, KÖRPER

In diesem Abschnitt werden wir die wichtigsten algebraischen Strukturen einführen. Gruppen treten in vielen Zusammenhängen in der Mathematik auf, allerdings wird das hier in der Linearen Algebra noch nicht so deutlich werden. Für uns wichtig sind Körper (das sind Strukturen, in denen man die vier Grundrechenarten zusammen mit den üblichen Rechenregeln zur Verfügung hat), denn zu einem Vektorraum (das ist die Struktur, die in der Linearen Algebra hauptsächlich betrachtet wird) gehört immer ein Körper, aus dem die „Skalare“ kommen. Ringe sind gewissermaßen Körper ohne Division; sie sind als Zwischenschritt bei der Definition von Körpern praktisch und auch wichtig in der Algebra. Sie werden ausführlicher in der Vorlesung „Einführung in die Zahlentheorie und algebraische Strukturen“ untersucht.

Wir beginnen mit dem Minimum, das man für eine halbwegs interessante algebraische Struktur braucht.

3.1. Definition. Eine *Halbgruppe* ist ein Paar $(H, *)$, bestehend aus einer Menge H und einer Verknüpfung $*$: $H \times H \rightarrow H$, $(a, b) \mapsto a * b$, die das *Assoziativgesetz* erfüllt:

DEF
Halbgruppe

$$\forall a, b, c \in H: (a * b) * c = a * (b * c).$$

Die Halbgruppe heißt *kommutativ*, wenn zusätzlich das *Kommutativgesetz* gilt:

$$\forall a, b \in H: a * b = b * a.$$

Wenn die Verknüpfung $*$ aus dem Kontext klar ist, spricht man der Einfachheit halber meist von „der Halbgruppe H “.

◇

Eine Bemerkung zur Notation: Verknüpfungen in algebraischen Strukturen wie $*$ in obiger Definition werden gerne in „Infix-Notation“ geschrieben, also $a * b$ statt $*(a, b)$.

Das Assoziativgesetz bewirkt, dass es nicht darauf ankommt, wie Ausdrücke, die drei oder mehr Elemente miteinander verknüpfen, geklammert sind. Zum Beispiel gilt für beliebige Elemente a, b, c, d, e von H :

$$\begin{aligned} a * ((b * c) * d) &= a * (b * (c * d)) = (a * b) * (c * d) \\ &= ((a * b) * c) * d = (a * (b * c)) * d \quad \text{und} \\ a * (b * (c * (d * e))) &= (a * b) * (c * (d * e)) = ((a * b) * (c * d)) * e = \dots \end{aligned}$$

Man kann deswegen einfach $a * b * c * d$ bzw. $a * b * c * d * e$ schreiben.

Hier ergibt sich die interessante kombinatorische Frage, *wie viele* verschiedene Klammierungen es für eine Verknüpfung von n Elementen gibt. Wir schreiben C_n für diese Zahl. Dann gilt offenbar $C_1 = C_2 = 1$, $C_3 = 2$ und $C_4 = 5$. Wenn man sich überlegt, dass man n Elemente dadurch verknüpfen kann, dass man eine Verknüpfung von k Elementen (mit $1 \leq k < n$) mit einer Verknüpfung von $n - k$ Elementen verknüpft, dann sieht man die folgende Rekursion für die Zahlen C_n :

$$C_n = \sum_{k=1}^{n-1} C_k C_{n-k} = C_1 C_{n-1} + C_2 C_{n-2} + \dots + C_{n-2} C_2 + C_{n-1} C_1 \quad \text{für alle } n \geq 2.$$

Damit kann man dann zum Beispiel $C_5 = 1 \cdot 5 + 1 \cdot 2 + 2 \cdot 1 + 5 \cdot 1 = 14$, $C_6 = 42$, $C_7 = 132$ usw. berechnen. Es gibt auch eine Formel für C_n , nämlich

$$C_n = \frac{1}{n} \binom{2n-2}{n-1} = \frac{1}{2n-1} \binom{2n-1}{n-1} = \frac{(2n-2)!}{(n-1)!n!},$$

die aber direkt nicht so einfach zu beweisen ist (was Sie natürlich nicht von einem Versuch abhalten soll!). Die Zahlen C_n heißen *Catalan-Zahlen* (was die Bezeichnung erklärt; oft ist der Index verschoben und man fängt mit $C_0 = C_1 = 1$ an) und treten in der Kombinatorik in vielen verschiedenen Zusammenhängen auf.

Wenn die Halbgruppe kommutativ ist, dann kommt es auch nicht auf die Reihenfolge an:

$$a * b * c = b * a * c = b * c * a = c * b * a = c * a * b = a * c * b.$$

Im Folgenden schreiben wir $\mathbb{N}_{>0} = \{n \in \mathbb{N} \mid n > 0\} = \{1, 2, 3, \dots\}$ für die Menge der echt positiven natürlichen Zahlen.

DEF
 $\mathbb{N}_{>0}$

3.2. Beispiele. Das Trivialbeispiel einer Halbgruppe ist $(\emptyset, *)$, wobei $*$: $\emptyset \times \emptyset \rightarrow \emptyset$ die leere Abbildung ist (beachte: $\emptyset \times \emptyset = \emptyset$).

BSP
Halbgruppen

Beispiele von kommutativen Halbgruppen sind $(\mathbb{N}_{>0}, +)$, $(\mathbb{N}, +)$, $(\mathbb{Z}, +)$, $(\mathbb{N}_{>0}, \cdot)$, $(\mathbb{N}, \cdot)$, $(\mathbb{Z}, \cdot)$. Die Halbgruppe $(\text{Abb}(X, X), \circ)$ für eine beliebige Menge X , mit der Komposition von Abbildungen als Verknüpfung, ist im Allgemeinen nicht kommutativ. (Diese Halbgruppe ist genau dann kommutativ, wenn X höchstens ein Element hat — Übung!)



Mit Halbgruppen kann man allerdings noch nicht allzu viel anfangen. Deshalb fordern wir zusätzliche Eigenschaften.

3.3. Definition. Ein *Monoid* ist ein Tripel $(M, *, e)$, bestehend aus einer Menge M , einer Verknüpfung $*$: $M \times M \rightarrow M$ und einem Element $e \in M$, sodass $(M, *)$ eine Halbgruppe mit *neutralem Element* e ist:

DEF
Monoid

$$\forall a \in M: e * a = a = a * e.$$

Das Monoid heißt *kommutativ*, wenn die Halbgruppe $(M, *)$ kommutativ ist. $\diamond$

Wenn es ein neutrales Element gibt, dann ist es eindeutig bestimmt, wie das folgende Lemma zeigt. (Ein *Lemma* ist eine Hilfsaussage oder ein weniger wichtiger mathematischer Satz. Plural: Lemmas oder Lemmata (wenn man seine klassische Bildung heraushängen lassen will).)

3.4. Lemma. Sei $(H, *)$ eine Halbgruppe. Ist e ein links- und e' ein rechtsneutrales Element in dieser Halbgruppe, also

LEMMA
Eindeutigkeit
des neutralen
Elements

$$\forall a \in H: e * a = a \text{ und } a * e' = a,$$

dann gilt $e = e'$.

Beweis. Da e linksneutral ist, gilt $e * e' = e'$. Da e' rechtsneutral ist, gilt $e * e' = e$. Es folgt $e = e'$. $\square$

Aus diesem Grund lässt man meistens die Angabe des neutralen Elements weg und spricht vom „Monoid $(M, *)$ “ oder auch nur vom „Monoid M “, wenn die Verknüpfung aus dem Kontext klar ist.

Es ist allerdings möglich, dass es in einer Halbgruppe zum Beispiel mehrere linksneutrale Elemente (und dann natürlich kein rechtsneutrales Element) gibt. Wenn etwa M beliebig ist und man als Verknüpfung pr_2 wählt (also $a * b = b$), dann hat man eine Halbgruppe, in der *alle* Elemente linksneutral sind.

3.5. Beispiele. Da die Definition von „Monoid“ ein neutrales Element fordert, kann die leere Menge kein Monoid sein. Das triviale Monoid ist dann $(\{e\}, *, e)$, wobei $*$ die einzige Abbildung $\{e\} \times \{e\} \rightarrow \{e\}$ ist (es ist also $e * e = e$).

Bis auf $(\mathbb{N}_{>0}, +)$, wo es kein neutrales Element gibt, lassen sich alle Beispiele von Halbgruppen aus 3.2 als Monoide $(\mathbb{N}, +, 0)$, $(\mathbb{Z}, +, 0)$, $(\mathbb{N}_{>0}, \cdot, 1)$, $(\mathbb{N}, \cdot, 1)$, $(\mathbb{Z}, \cdot, 1)$ und $(\text{Abb}(X, X), \circ, \text{id}_X)$ betrachten. ♣

Noch schöner ist es, wenn sich die Verknüpfung mit einem Element durch die Verknüpfung mit einem (in der Regel) anderen Element wieder rückgängig machen lässt. Das führt auf den Begriff der Gruppe.

3.6. Definition. Eine *Gruppe* ist ein Quadrupel $(G, *, e, i)$, bestehend aus einer Menge G , einer Verknüpfung $*$: $G \times G \rightarrow G$, einem Element $e \in G$ und einer Abbildung $i: G \rightarrow G$, sodass $(G, *, e)$ ein Monoid ist und für jedes $g \in G$ das Element $i(g) \in G$ ein *Inverses* von g ist:

$$\forall g \in G: i(g) * g = e = g * i(g).$$

Die Gruppe heißt *kommutativ* oder *abelsch*, wenn das Monoid $(G, *, e)$ kommutativ ist. ◇

Die Bezeichnung „abelsch“ ehrt den norwegischen Mathematiker Niels Henrik Abel, nach dem auch der *Abelpreis* benannt ist, ein dem Nobelpreis vergleichbarer Preis für Mathematik, der seit 2003 jährlich verliehen wird.

Auch Inverse sind eindeutig bestimmt:

3.7. Lemma. Sei $(M, *, e)$ ein Monoid und sei $a \in M$. Ist $b \in M$ ein Linksinverses und $c \in M$ ein Rechtsinverses von a , also

$$b * a = e = a * c,$$

dann gilt $b = c$.

Beweis. Wir haben

$$b = b * e = b * (a * c) = (b * a) * c = e * c = c.$$

Analog zu Monoiden spricht man deshalb auch einfach von „der Gruppe $(G, *)$ “ oder auch von „der Gruppe G “, wenn die Verknüpfung aus dem Kontext klar ist.

Gruppen schreibt man gerne „multiplikativ“, dann ist die Verknüpfung $a \cdot b$ oder kurz ab , das neutrale Element heißt 1 und das Inverse von a wird a^{-1} geschrieben. Kommutative Gruppen schreibt man auch häufig „additiv“, dann ist die Verknüpfung $a + b$, das neutrale Element heißt 0 und das Inverse von a wird als das Negative von a geschrieben: $-a$. Dann schreibt man auch kurz $a - b$ für $a + (-b)$.

3.8. Beispiele. Das triviale Monoid lässt sich auch als Gruppe betrachten, denn das einzige Element e ist sein eigenes Inverses.

Von den übrigen Beispielen von Monoiden in 3.5 kann nur $(\mathbb{Z}, +, 0, -)$ auch als Gruppe betrachtet werden (und im letzten Beispiel $\text{Abb}(X, X)$, wenn X höchstens ein Element hat; dann hat man eine triviale Gruppe). Ein weiteres Beispiel einer kommutativen Gruppe ist $(\mathbb{R}_{>0}, \cdot, 1, x \mapsto 1/x)$, wobei $\mathbb{R}_{>0}$ die Menge der positiven reellen Zahlen ist.

BSP
Monoide

DEF
Gruppe



N.H. Abel
1802–1829

LEMMA
Eindeutigkeit
des Inversen

BSP
Gruppen

Wenn man sich bei den Abbildungen $X \rightarrow X$ auf die bijektiven Abbildungen beschränkt, dann erhält man eine Gruppe $(S(X), \circ, \text{id}_X, f \mapsto f^{-1})$, die auch die *symmetrische Gruppe* von X heißt. Dabei ist

$$S(X) = \{f: X \rightarrow X \mid f \text{ bijektiv}\}.$$

Diese Gruppe ist genau dann kommutativ, wenn X höchstens zwei Elemente enthält (Übung).

Gruppen treten häufig in der Mathematik als „Symmetriegruppen“ von irgendwelchen Objekten auf. Zum Beispiel bilden die Drehungen und Spiegelungen der Ebene, die ein regelmäßiges n -Eck auf sich abbilden, eine Gruppe, oder die Drehungen des dreidimensionalen Raumes, die ein reguläres Tetraeder, einen Würfel (oder ein reguläres Oktaeder) oder ein reguläres Dodekaeder (oder Ikosaeder) in sich abbilden, bilden jeweils eine Gruppe, die Tetraeder-, Oktaeder- und Ikosaedergruppe. In einem recht allgemeinen Sinn ist die symmetrische Gruppe $S(X)$ die Symmetriegruppe der Menge X ohne weitere Struktur. In der Algebra treten Symmetriegruppen als „Automorphismengruppen“ auf. Zum Beispiel bildet für eine Gruppe $(G, *)$ die Menge

$$\text{Aut}(G) = \{f: G \rightarrow G \mid f \text{ bijektiv und } \forall g, g' \in G: f(g * g') = f(g) * f(g')\}$$

mit der Komposition von Abbildungen eine Gruppe, die *Automorphismengruppe* von G . Sie besteht aus den bijektiven Abbildungen $G \rightarrow G$, die mit der Struktur von G als Gruppe verträglich sind. ♣

Damit eine Halbgruppe sogar eine Gruppe ist, genügt es, die Existenz eines linksneutralen Elements e und für jedes Element x die Existenz eines Linksinversen $i(x)$ (also mit $i(x) * x = e$) zu fordern. Dann folgt zunächst, dass e auch rechtsneutral ist, denn es gilt

$$x * e = e * x * e = i(i(x)) * i(x) * x * e = i(i(x)) * e * e = i(i(x)) * e = i(i(x)) * i(x) * x = e * x = x.$$

Daraus ergibt sich auch $i(i(x)) = x$. Damit kann man dann zeigen, dass $i(x)$ auch ein Rechtsinverses von x ist:

$$x * i(x) = i(i(x)) * i(x) = e.$$

Ganz analog funktioniert das natürlich auch, wenn man „links“ jeweils durch „rechts“ ersetzt. Auf der anderen Seite gibt es aber Halbgruppen mit linksneutralen und rechtsinversen Elementen, die keine Gruppen sind. Finden Sie ein Beispiel!

Eine wichtige Eigenschaft von Gruppen ist, dass sich gewisse Gleichungen stets eindeutig lösen lassen. Zuerst beweisen wir aber eine Kürzungsregel.

3.9. Lemma. Sei $(G, *, e, i)$ eine Gruppe und seien $a, b, c \in G$. Dann gilt

$$a * c = b * c \iff a = b \iff c * a = c * b.$$

LEMMA
Kürzungsregel
in Gruppen

Beweis. Wir beweisen die erste Äquivalenz; der Beweis der zweiten ist analog.

„ $\Leftarrow$ “ ist klar. Für „ $\Rightarrow$ “ haben wir

$$\begin{aligned} a * c = b * c &\implies (a * c) * i(c) = (b * c) * i(c) \implies a * (c * i(c)) = b * (c * i(c)) \\ &\implies a * e = b * e \implies a = b. \end{aligned}$$

□

3.10. Lemma. Sei $(G, *, e, i)$ eine Gruppe und seien $a, b \in G$. Dann haben die Gleichungen

$$a * x = b \quad \text{und} \quad x * a = b$$

jeweils eine eindeutige Lösung $x \in G$, nämlich $x = i(a) * b$ bzw. $x = b * i(a)$.

LEMMA
Gleichungen
in Gruppen

Beweis. Wir führen den Beweis exemplarisch für die erste Gleichung:

$$a * x = b \iff i(a) * a * x = i(a) * b \iff e * x = i(a) * b \iff x = i(a) * b.$$

Für die erste Äquivalenz haben wir Lemma 3.9 benutzt. $\square$

Als Nächstes betrachten wir Strukturen mit zwei Verknüpfungen.

3.11. Definition. Ein *Ring* ist ein Sextupel $(R, +, 0, -, \cdot, 1)$, bestehend aus einer Menge R , Verknüpfungen $+, \cdot: R \times R \rightarrow R$, Elementen $0, 1 \in R$ und einer Abbildung $-: R \rightarrow R$, sodass $(R, +, 0, -)$ eine kommutative Gruppe und $(R, \cdot, 1)$ ein Monoid ist und die *Distributivgesetze*

DEF
Ring

$$\forall a, b, c \in R: a \cdot (b + c) = a \cdot b + a \cdot c \quad \text{und} \quad (a + b) \cdot c = a \cdot c + b \cdot c$$

gelten. Der Ring heißt *kommutativ*, wenn das Monoid $(R, \cdot, 1)$ kommutativ ist. $\diamond$

Da die neutralen und inversen Elemente eindeutig bestimmt sind, spricht man oft nur vom „Ring $(R, +, \cdot)$ “ oder sogar vom „Ring R “, wenn die Verknüpfungen aus dem Kontext klar sind. Ist der Ring kommutativ, dann genügt es, eines der beiden Distributivgesetze zu fordern (das andere folgt dann). Für das Produkt $a \cdot b$ zweier Elemente schreibt man auch kurz ab .

In einem Ring kann man also addieren, subtrahieren und multiplizieren, und die üblichen Rechenregeln gelten, wie zum Beispiel $0 \cdot a = a \cdot 0 = 0$, $-(a + b) = -a - b$, $(-a) \cdot (-b) = a \cdot b$ (Übung). Was aber im Allgemeinen *nicht* gelten muss, ist die Implikation $a \cdot b = 0 \Rightarrow a = 0 \vee b = 0$.



Wenn sowohl Summen als auch Produkte in einem Term vorkommen, gilt „Punkt vor Strich“, d.h., die Multiplikation bindet stärker als die Addition:

$$a \cdot b + c \cdot d = (a \cdot b) + (c \cdot d) \quad \text{und nicht} \quad a \cdot (b + c) \cdot d.$$

3.12. Beispiele. Das Trivialbeispiel für einen Ring ist der sogenannte *Nullring* $(\{0\}, +, 0, -, \cdot, 0)$, in dem $0 = 1$ und $0 + 0 = -0 = 0 \cdot 0 = 0$ gelten. Jeder Ring R , in dem $0_R = 1_R$ gilt, ist so ein Nullring, denn für alle $r \in R$ gilt dann $r = 1_R \cdot r = 0_R \cdot r = 0_R$.

BSP
Ringe

Das Standardbeispiel für einen (kommutativen) Ring ist der Ring $\mathbb{Z}$ der ganzen Zahlen mit der üblichen Addition und Multiplikation als Verknüpfungen. Ein etwas anders geartetes Beispiel ist $(\mathcal{P}(X), \Delta, \emptyset, \text{id}_{\mathcal{P}(X)}, \cap, X)$ für eine beliebige Menge X ; dabei ist $T_1 \Delta T_2 = (T_1 \setminus T_2) \cup (T_2 \setminus T_1)$ die „symmetrische Differenz“ der Mengen T_1 und T_2 (Übung).

Falls Sie aus der Schule Matrizen kennen und wissen, wie man sie addiert und multipliziert, dann können Sie nachprüfen, dass die Menge der 2×2 -Matrizen mit Einträgen aus $\mathbb{R}$ zusammen mit der Addition und Multiplikation von Matrizen einen nicht-kommutativen Ring bildet. $\clubsuit$

Schließlich kommen wir zu den Körpern. (Der Stern im linken Rand neben der Definition bedeutet, dass dies eine der Definitionen ist, die möglicherweise in der Klausur explizit abgefragt werden. Ähnliches gilt für Sätze „mit Stern“. Das heißt natürlich nicht, dass Sie die übrigen Definitionen und Sätze nicht kennen müssten!) 

*

3.13. Definition. Ein *Körper* ist ein Septupel $(K, +, 0, -, \cdot, 1, i)$, bestehend aus einer Menge K , Verknüpfungen $+, \cdot: K \times K \rightarrow K$, Elementen $0, 1 \in K$ und Abbildungen $-: K \rightarrow K$, $i: K \setminus \{0\} \rightarrow K \setminus \{0\}$, sodass $(K, +, 0, -, \cdot, 1)$ ein kommutativer Ring und $(K \setminus \{0\}, \cdot, 1, i)$ eine (kommutative) Gruppe ist. Für $i(a)$ schreibt man a^{-1} . $\diamond$

DEF
Körper

Wie üblich spricht man meistens einfach von dem „Körper $(K, +, \cdot)$ “ oder von dem „Körper K “. Aus der Definition folgt, dass 0 und 1 in einem Körper verschieden sein müssen, denn 1 soll das neutrale Element der Gruppe $K \setminus \{0\}$ sein. Für diese Gruppe $(K \setminus \{0\}, \cdot)$ schreibt man auch $K^\times$ und nennt sie die *multiplikative Gruppe* von K . (Häufig findet man auch die Schreibweise K^* dafür.)

DEF
multiplikative
Gruppe $K^\times$

Man kann natürlich auch ohne Rückgriff auf Ringe und Gruppen definieren, was ein Körper ist. Dann hat man für alle $a, b, c \in K$ die folgenden Axiome:

$$\begin{array}{ll} (a+b)+c = a+(b+c), & a+b = b+a \\ a+0 = a, & a+(-a) = 0 \\ (a \cdot b) \cdot c = a \cdot (b \cdot c) & a \cdot b = b \cdot a \\ a \cdot 1 = a, & a \neq 0 \Rightarrow a \cdot a^{-1} = 1 \\ 0 \neq 1, & a \cdot (b+c) = a \cdot b + a \cdot c \end{array}$$

Für $a, b \in K$, $b \neq 0$, kann man die Division definieren durch $a/b = a \cdot b^{-1}$. Dann hat man die vier Grundrechenarten zur Verfügung und die üblichen Rechenregeln dafür gelten, denn man kann sie aus den Körperaxiomen herleiten. Zum Beispiel gilt in einem Körper stets, dass aus $a \cdot b = 0$ folgt, dass $a = 0$ oder $b = 0$ ist. (Denn ist $a \neq 0$, dann folgt $0 = a^{-1} \cdot 0 = a^{-1} \cdot a \cdot b = 1 \cdot b = b$.)

3.14. Beispiele. Das kleinste Beispiel für einen Körper hat nur die beiden Elemente 0 und 1, die in der Definition gefordert werden. Für die Addition und Multiplikation folgen $0+0=0$, $0+1=1+0=1$, $0 \cdot 0=0 \cdot 1=1 \cdot 0=0$ und $1 \cdot 1=1$ direkt aus der Definition; für die verbleibende Summe $1+1$ bleibt nur der Wert 0, da die Gleichung $a+1=0$ lösbar sein muss. Man kann (einfach, aber länglich) nachprüfen, dass dieser Körper, der mit $\mathbb{F}_2$ bezeichnet wird, die Axiome erfüllt.

BSP
Körper

Es gibt noch weitere endliche Körper: Zu jeder Potenz p^e einer Primzahl p (mit $e \geq 1$) gibt es im Wesentlichen („bis auf Isomorphie“) genau einen Körper $\mathbb{F}_{p^e}$ mit p^e Elementen, und es gibt keine anderen endlichen Körper. Das wird in der „Einführung in die Algebra“ genauer besprochen.

Standardbeispiele für Körper sind der Körper $\mathbb{Q}$ der rationalen Zahlen und der Körper $\mathbb{R}$ der reellen Zahlen, jeweils mit der bekannten Addition und Multiplikation. Im nächsten Abschnitt werden wir einen weiteren Körper konstruieren, den Körper $\mathbb{C}$ der komplexen Zahlen. $\clubsuit$

4. DER KÖRPER DER KOMPLEXEN ZAHLEN

Der Körper $\mathbb{R}$ der reellen Zahlen hat, wie Sie in der Analysis lernen (werden), viele schöne Eigenschaften. Eine Eigenschaft allerdings fehlt ihm: Es sind nicht alle Gleichungen der Form

$$x^n + a_{n-1}x^{n-1} + \dots + a_1x + a_0 = 0$$

(mit $n \geq 1$ und $a_0, a_1, \dots, a_{n-1} \in \mathbb{R}$) in $\mathbb{R}$ lösbar.

Für ungerades n folgt aus dem *Zwischenwertsatz*, dass es stets eine Lösung geben muss; das lernen Sie in der Analysis.

Die einfachste Gleichung dieser Art ohne Lösung ist $x^2 + 1 = 0$: Die linke Seite ist stets ≥ 1 , kann also niemals null werden. Wir werden jetzt einen $\mathbb{R}$ umfassenden Körper konstruieren, in dem diese Gleichung eine Lösung hat.

Um zu sehen, wie man dabei vorgehen kann, stellen wir uns einfach einmal vor, dass wir schon so einen Körper hätten; wir nennen ihn $\mathbb{C}$. Dann haben wir eine Lösung i obiger Gleichung, also ein Element $i \in \mathbb{C}$ mit $i^2 = -1$. Wir haben natürlich auch die reellen Zahlen in $\mathbb{C}$. Mit $a, b \in \mathbb{R}$ können wir dann das Element $a + bi \in \mathbb{C}$ erzeugen. Muss es noch weitere Elemente geben? Dazu müssen wir überprüfen, ob die vier Grundrechenarten aus der Menge der Elemente der Form $a + bi$ herausführen. Seien $a, b, a', b' \in \mathbb{R}$. Dann gilt, wenn $\mathbb{C}$ ein Körper ist,

$$\begin{aligned}(a + bi) + (a' + b'i) &= (a + a') + (b + b')i \quad \text{und} \\ (a + bi) \cdot (a' + b'i) &= aa' + ab'i + ba'i + bb'i^2 = (aa' - bb') + (ab' + ba')i.\end{aligned}$$

Dabei haben wir $i^2 = -1$ benutzt. Offensichtlich ist das additive Inverse (also das Negative) von $a + bi$ gerade $(-a) + (-b)i$. Wie sieht es mit dem multiplikativen Inversen aus (also dem Kehrwert)? Dazu überlegen wir uns erst, dass genau dann $a + bi = 0$ ist, wenn $a = b = 0$ gilt. Eine Richtung („ $\Leftarrow$ “) ist klar. Umgekehrt sei $a + bi = 0$. Dann folgt

$$0 = (a - bi) \cdot 0 = (a - bi) \cdot (a + bi) = a^2 + b^2.$$

Da a und b reelle Zahlen sind, ist das nur möglich, wenn $a = b = 0$ gilt. Seien also a und b nicht beide null. Dann sollte gelten (das ist der alte Trick, wie man „Quadratwurzeln aus dem Nenner entfernt“; man beachte, dass $i = „\sqrt{-1}“$ ist):

$$\frac{1}{a + bi} = \frac{a - bi}{(a - bi)(a + bi)} = \frac{a - bi}{a^2 + b^2} = \frac{a}{a^2 + b^2} + \frac{-b}{a^2 + b^2}i.$$

Offenbar brauchen wir also keine zusätzlichen Elemente.

Um das Ganze formal auf eine solide Grundlage zu stellen, ersetzen wir einen Ausdruck der Form $a + bi$ durch das Paar $(a, b) \in \mathbb{R} \times \mathbb{R}$. Wir schreiben $\mathbb{C}$ für $\mathbb{R} \times \mathbb{R}$ und definieren die folgenden Abbildungen:

$$\begin{aligned}+_C: \mathbb{C} \times \mathbb{C} &\longrightarrow \mathbb{C}, & ((a, b), (a', b')) &\longmapsto (a + a', b + b') \\ \cdot_C: \mathbb{C} \times \mathbb{C} &\longrightarrow \mathbb{C}, & ((a, b), (a', b')) &\longmapsto (aa' - bb', ab' + ba') \\ -_C: \mathbb{C} &\longrightarrow \mathbb{C}, & (a, b) &\longmapsto (-a, -b) \\ i_C: \mathbb{C} \setminus \{(0, 0)\} &\longrightarrow \mathbb{C} \setminus \{(0, 0)\}, & (a, b) &\longmapsto \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right)\end{aligned}$$

Außerdem schreiben wir 0_C und 1_C für $(0, 0)$ und $(1, 0)$.

4.1. Satz. Die Menge $\mathbb{C} = \mathbb{R} \times \mathbb{R}$ zusammen mit den oben definierten Abbildungen und Elementen bildet einen Körper.

SATZ
Körper $\mathbb{C}$

Beweis. Es sind die verschiedenen Axiome nachzuprüfen. Für die additive Gruppe $(\mathbb{C}, +_{\mathbb{C}}, 0_{\mathbb{C}}, -_{\mathbb{C}})$ ist das sehr leicht; darum lassen wir das hier weg (es sei Ihnen aber als Übung empfohlen). Wir prüfen Assoziativität und Kommutativität der Multiplikation. Dabei benutzen wir, dass $\mathbb{R}$ ein Körper ist, dass also dort die bekannten Rechenregeln gelten.

$$\begin{aligned} ((a, b) \cdot_{\mathbb{C}} (a', b')) \cdot_{\mathbb{C}} (a'', b'') &= (aa' - bb', ab' + a'b) \cdot_{\mathbb{C}} (a'', b'') \\ &= ((aa' - bb')a'' - (ab' + a'b)b'', (aa' - bb')b'' + (ab' + a'b)a'') \\ &= (aa'a'' - ab'b'' - ba'b'' - bb'a'', aa'b'' + ab'a'' + ba'a'' - bb'b'') \end{aligned}$$

und dasselbe Resultat erhalten wir aus $(a, b) \cdot_{\mathbb{C}} ((a', b') \cdot_{\mathbb{C}} (a'', b''))$. Ebenso gilt

$$(a, b) \cdot_{\mathbb{C}} (a', b') = (aa' - bb', ab' + ba') = (a'a - b'b, ba' + ab') = (a', b') \cdot_{\mathbb{C}} (a, b).$$

Dass $1_{\mathbb{C}} = (1, 0)$ neutrales Element der Multiplikation ist, folgt aus

$$(1, 0) \cdot_{\mathbb{C}} (a, b) = (1 \cdot a - 0 \cdot b, 1 \cdot b + 0 \cdot a) = (a, b).$$

Wir rechnen nach, dass $i_{\mathbb{C}}((a, b))$ das multiplikative Inverse von $(a, b) \neq (0, 0)$ ist:

$$\begin{aligned} (a, b) \cdot_{\mathbb{C}} i_{\mathbb{C}}((a, b)) &= (a, b) \cdot_{\mathbb{C}} \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right) \\ &= \left(\frac{a^2}{a^2 + b^2} - \frac{-b^2}{a^2 + b^2}, \frac{-ab}{a^2 + b^2} + \frac{ba}{a^2 + b^2} \right) \\ &= (1, 0) = 1_{\mathbb{C}}. \end{aligned}$$

$0_{\mathbb{C}} \neq 1_{\mathbb{C}}$ ist klar. Es bleibt das Distributivgesetz nachzuprüfen:

$$\begin{aligned} (a, b) \cdot_{\mathbb{C}} ((a', b') +_{\mathbb{C}} (a'', b'')) &= (a, b) \cdot_{\mathbb{C}} (a' + a'', b' + b'') \\ &= (a(a' + a'') - b(b' + b''), a(b' + b'') + b(a' + a'')) \\ &= (aa' + aa'' - bb' - bb'', ab' + ab'' + ba' + ba'') \\ &= (aa' - bb' + aa'' - bb'', ab' + ba' + ab'' + ba'') \\ &= (aa' - bb', ab' + ba') +_{\mathbb{C}} (aa'' - bb'', ab'' + ba'') \\ &= (a, b) \cdot_{\mathbb{C}} (a', b') +_{\mathbb{C}} (a, b) \cdot_{\mathbb{C}} (a'', b''). \quad \square \end{aligned}$$

Ist a eine reelle Zahl, dann haben wir das Element $a_{\mathbb{C}} = (a, 0) \in \mathbb{C}$. Für $a, b \in \mathbb{R}$ gilt

$$a = b \iff a_{\mathbb{C}} = b_{\mathbb{C}}, \quad (a + b)_{\mathbb{C}} = a_{\mathbb{C}} +_{\mathbb{C}} b_{\mathbb{C}} \quad \text{und} \quad (ab)_{\mathbb{C}} = a_{\mathbb{C}} \cdot_{\mathbb{C}} b_{\mathbb{C}}.$$

Mit den Elementen $a_{\mathbb{C}}$ rechnet man also genauso wie mit den zugehörigen reellen Zahlen a . Deswegen macht man keinen Unterschied zwischen a und $a_{\mathbb{C}}$ und betrachtet $\mathbb{R}$ als eine Teilmenge von $\mathbb{C}$. Wir schreiben also einfach a für das Element $a_{\mathbb{C}} = (a, 0)$ von $\mathbb{C}$. Außerdem schreiben wir ab jetzt der Einfachheit halber meistens $+$, $\cdot$ und so weiter statt $+_{\mathbb{C}}$, $\cdot_{\mathbb{C}}$ etc.

4.2. Definition. Der in Satz 4.1 eingeführte Körper $\mathbb{C}$ heißt der *Körper der komplexen Zahlen*. Wir schreiben i für das Element $(0, 1) \in \mathbb{C}$. Dann gilt $i^2 = -1$, und jedes Element $z = (a, b) \in \mathbb{C}$ kann geschrieben werden als $z = a + bi$ (oder $a + ib$) mit $a, b \in \mathbb{R}$. Dann heißt a der *Realteil* $\operatorname{Re} z$ und b der *Imaginärteil* $\operatorname{Im} z$ von z . Gilt $\operatorname{Re} z = 0$, dann heißt z *rein imaginär*. $\diamond$

DEF
Körper der
komplexen
Zahlen

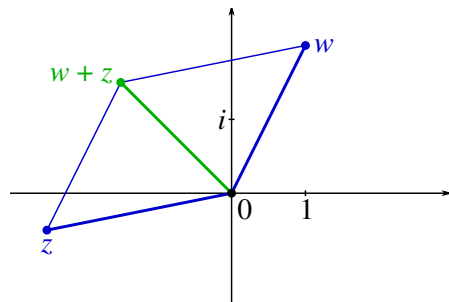
Die in der Definition gemachten Behauptungen sollten wir nachprüfen:

$$i^2 = (0, 1) \cdot_{\mathbb{C}} (0, 1) = (0 \cdot 0 - 1 \cdot 1, 0 \cdot 1 + 1 \cdot 0) = (-1, 0) = (-1)_{\mathbb{C}} = -1$$

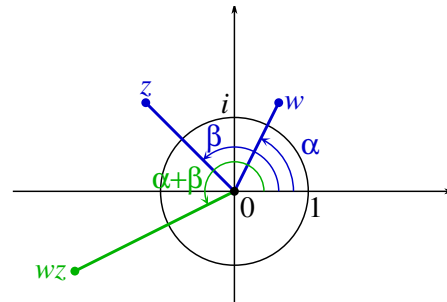
und

$$a + bi = (a, 0) +_{\mathbb{C}} (b, 0) \cdot_{\mathbb{C}} (0, 1) = (a, 0) +_{\mathbb{C}} (0, b) = (a, b).$$

Man kann sich die komplexen Zahlen ganz gut veranschaulichen, wenn man sich daran erinnert, dass $\mathbb{C} = \mathbb{R} \times \mathbb{R} = \mathbb{R}^2$ der Menge der Punkte der Ebene entspricht. Wenn man die Ebene so interpretiert, spricht man auch von der *komplexen (Zahlen-)Ebene*. Die Addition entspricht dann dem, was Sie aus der Physik als „Kräfteparallelogramm“ kennen.



Addition $w + z$



Multiplikation $w \cdot z$

Auch die Multiplikation lässt sich geometrisch interpretieren. Wir betrachten dazu $z = a + bi \in \mathbb{C}$. Dann ist $a^2 + b^2 \geq 0$; man setzt $|z| = \sqrt{a^2 + b^2}$ und nennt das den *Absolutbetrag* von z . Das entspricht dem Abstand des Punktes z in der komplexen Ebene vom Ursprung $0 \in \mathbb{C}$. Für $z \in \mathbb{R}$ (also $b = 0$) bekommt man den bekannten Absolutbetrag auf $\mathbb{R}$. Ist $z \neq 0$, dann hat $w = z/|z|$ den Absolutbetrag 1. Wenn wir $w = u + vi$ schreiben, dann gilt $u^2 + v^2 = 1$, also liegt der Punkt (u, v) auf dem Einheitskreis. Es gibt dann $\alpha \in \mathbb{R}$ mit $u = \cos \alpha$, $v = \sin \alpha$. Dieser Winkel α heißt auch das *Argument* von w und von z (der Winkel ist aber nur bis auf Vielfache von $2\pi \hat{=} 360^\circ$ eindeutig bestimmt). Es gilt die Beziehung

DEF
Absolut-
betrag $|z|$

DEF
Argument

$$(\cos \alpha + i \sin \alpha) \cdot (\cos \beta + i \sin \beta) = \cos(\alpha + \beta) + i \sin(\alpha + \beta)$$

— das ist äquivalent zu den Additionstheoremen für Sinus und Cosinus:

$$\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta, \quad \sin(\alpha + \beta) = \cos \alpha \sin \beta + \sin \alpha \cos \beta.$$

Daher addieren sich die Winkel bei Multiplikation. Man kann das dann so formulieren: Multiplikation mit $z \neq 0$ bewirkt eine *Drehstreckung* der komplexen Ebene mit dem Drehwinkel α und dem Streckfaktor $|z|$.

In der Analysis werden Sie lernen, dass

$$\cos \alpha + i \sin \alpha = e^{i\alpha}$$

ist. Die Relation oben folgt dann aus $e^{x+y} = e^x \cdot e^y$.

Wir können jetzt immerhin zeigen, dass man quadratische Gleichungen in $\mathbb{C}$ stets lösen kann.

4.3. Satz. Seien $a, b, c \in \mathbb{C}$ mit $a \neq 0$. Dann hat die Gleichung

$$az^2 + bz + c = 0$$

mindestens eine Lösung $z \in \mathbb{C}$.

SATZ
quadratische
Gleichungen
in $\mathbb{C}$

Beweis. Die Gleichung ist äquivalent zu $(2az + b)^2 = b^2 - 4ac$. Es genügt also zu zeigen, dass jede komplexe Zahl eine Quadratwurzel in $\mathbb{C}$ hat. (Ist $z' \in \mathbb{C}$ mit $z'^2 = b^2 - 4ac$, dann ist $z = (-b + z')/(2a)$ eine Lösung der Gleichung. Das ist die bekannte Lösungsformel für quadratische Gleichungen.) Sei also $w \in \mathbb{C}$. Wir wollen $z \in \mathbb{C}$ finden mit $z^2 = w$. Ist $w = 0$, dann ist $z = 0$ eine Lösung. Sonst können wir wie oben w als $w = |w|(\cos \alpha + i \sin \alpha)$ schreiben. Dann ist $z = \sqrt{|w|}(\cos \frac{\alpha}{2} + i \sin \frac{\alpha}{2})$ eine Lösung. $\square$

Analog zeigt man, dass es für jede komplexe Zahl w und jede natürliche Zahl $n \geq 1$ eine n te Wurzel $z \in \mathbb{C}$ von w gibt, also eine Lösung der Gleichung $z^n = w$. Es gilt aber sogar noch viel mehr.

4.4. Satz. Jede Gleichung

$$z^n + a_{n-1}z^{n-1} + \dots + a_1z + a_0 = 0$$

mit $n \geq 1$ und $a_0, a_1, \dots, a_{n-1} \in \mathbb{C}$ hat mindestens eine Lösung $z \in \mathbb{C}$.

SATZ
Fundamental-
satz der
Algebra

Beweisen können wir diesen Satz hier nicht. Es gibt verschiedene Beweise; der wohl einfachste verwendet den *Satz von Liouville* aus der Funktionentheorie. Sie werden ihn in der Vorlesung „Funktionentheorie“ kennenlernen.

Zwar lassen sich Gleichungen der obigen Form mit $n \leq 4$ durch Ziehen von Quadrat- und Kubikwurzeln lösen (das wurde bereits im 16. Jahrhundert von del Ferro, Tartaglia und Ferrari entdeckt); Gleichungen mit $n \geq 5$ lassen sich jedoch im Allgemeinen nicht mehr mit Hilfe der vier Grundrechenarten und durch das Ziehen von beliebigen m ten Wurzeln lösen (Satz von Abel-Ruffini; erster vollständiger Beweis 1824 von Abel). Die Aussage von Satz 4.4 ist also viel stärker als die Existenz von n ten Wurzeln in $\mathbb{C}$.

Ein Körper K mit der Eigenschaft, dass jede Gleichung

$$x^n + a_{n-1}x^{n-1} + \dots + a_1x + a_0 = 0$$

mit $n \geq 1$ und $a_0, a_1, \dots, a_{n-1} \in K$ eine Lösung $x \in K$ hat, heißt *algebraisch abgeschlossen*. Der „Fundamentalsatz der Algebra“ lässt sich also auch so formulieren:

Der Körper der komplexen Zahlen ist algebraisch abgeschlossen.

DEF
algebraisch
abgeschlossen

Demgegenüber ist der Körper der reellen Zahlen *nicht* algebraisch abgeschlossen, wie wir gesehen haben. In dieser Hinsicht ist $\mathbb{C}$ also „besser“ als $\mathbb{R}$. Auf der anderen Seite ist $\mathbb{C}$ kein angeordneter Körper mehr; man verliert also auch etwas beim Übergang von $\mathbb{R}$ zu $\mathbb{C}$. (In einem angeordneten Körper K gilt $x^2 \geq 0$ für $x \in K$. Damit müsste in $\mathbb{C}$ gelten, dass $-1 = i^2 \geq 0$ ist, aber -1 ist in einem angeordneten Körper immer negativ, und wir haben einen Widerspruch.)

Da mit i auch $-i$ eine Lösung von $x^2 + 1 = 0$ ist, könnte man überall i durch $-i$ ersetzen und alles würde genauso funktionieren. Für $z = a + bi \in \mathbb{C}$ setzen wir daher $\bar{z} = a - bi$; die Abbildung $\mathbb{C} \rightarrow \mathbb{C}$, $z \mapsto \bar{z}$, heißt die *komplexe Konjugation*. Es gilt $\overline{w + z} = \bar{w} + \bar{z}$ und $\overline{wz} = \bar{w} \cdot \bar{z}$ (leichte Übung); außerdem $z\bar{z} = a^2 + b^2 = |z|^2$ (das haben wir schon benutzt). Daraus bekommt man die Formel $z^{-1} = \bar{z}/|z|^2$ für den Kehrwert einer komplexen Zahl $z \neq 0$; das ist derselbe Ausdruck, den wir bereits hergeleitet

hatten, in einer etwas abgekürzten Form. Außerdem hat die komplexe Konjugation noch die folgenden Eigenschaften:

$$z \in \mathbb{R} \iff z = \bar{z}, \quad \operatorname{Re} z = \frac{z + \bar{z}}{2}, \quad \operatorname{Im} z = \frac{z - \bar{z}}{2i}.$$

5. VEKTORRÄUME: DEFINITION UND BEISPIELE

In diesem Abschnitt beginnen wir mit dem Studium der Linearen Algebra. Was ist „Lineare Algebra“? Die Lineare Algebra befasst sich mit „linearen Strukturen“, genauer mit *Vektorräumen* und *linearen Abbildungen* zwischen ihnen. Diese Begriffe sind zunächst einmal sehr abstrakt, aber darin liegt gerade die Stärke der Linearen Algebra: Vektorräume und lineare Abbildungen treten sehr häufig in der Mathematik in den unterschiedlichsten Zusammenhängen auf. Gerade weil man von den jeweils konkreten individuellen Umständen abstrahiert und sich auf die wesentlichen gemeinsamen Eigenschaften beschränkt, lassen sich die Ergebnisse der Linearen Algebra in all diesen unterschiedlichen Situationen anwenden. Es war, historisch gesehen, ein langwieriger Prozess, zu dieser Abstraktion zu gelangen, aber am Endpunkt dieser Entwicklung steht eine sehr leistungsfähige, allgemein anwendbare und erfolgreiche Theorie. Das hat dazu geführt, dass *lineare* Probleme als *einfach* gelten, während *nichtlineare* Probleme sehr häufig besonders *schwierig* sind. In Ausschreibungen für Mathematik-Professuren findet man zum Beispiel häufiger das Wort „nichtlinear“ (etwa im Kontext von „nichtlinearen partiellen Differentialgleichungen“), aber so gut wie niemals das Wort „linear“. Zwei Beispiele mit physikalischem Hintergrund: Die *Wärmeleitungsgleichung*, die die zeitliche Entwicklung der Temperaturverteilung in einem Körper beschreibt, ist eine *lineare* partielle Differentialgleichung. Die zugehörige Lösungstheorie wurde bereits von Jean-Baptiste-Joseph Fourier entwickelt („Théorie analytique de la chaleur“, 1822). Im Gegensatz dazu sind die *Navier-Stokes-Gleichungen*, die die Bewegung von Flüssigkeiten beschreiben, *nichtlineare* partielle Differentialgleichungen, und die Frage, ob sie für vernünftige Anfangsbedingungen im dreidimensionalen Raum immer eindeutig für beliebig lange Zeiträume lösbar sind, ist eines der sieben Millenniumprobleme der Clay Foundation; für die Lösung bekommt man eine Million US-Dollar.



J.-B.-J. Fourier
(1768–1830)

Was bedeutet nun „linear“? Dazu als Beispiele drei lineare Gleichungen (oder Gleichungssysteme):

5.1. Beispiele.

- (1) Wir suchen $w, x, y, z \in \mathbb{R}$ mit

$$w + x + y + z = 0 \quad \text{und} \quad x + 2y + 3z = 0.$$

Wahrscheinlich haben Sie in der Schule gelernt, wie man solche Gleichungssysteme löst (und in jedem Fall werden wir das auch in dieser Vorlesung besprechen). Als Lösungen erhält man

$$(w, x, y, z) = (a, -2a + b, a - 2b, b) \quad \text{mit } a, b \in \mathbb{R}.$$

- (2) Wir suchen Folgen $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen, für die gilt

$$a_{n+2} = a_{n+1} + a_n \quad \text{für alle } n \in \mathbb{N}.$$

Die Folge $(0, 1, 1, 2, 3, 5, 8, \dots)$ der Fibonacci-Zahlen ist eine Lösung, aber es gibt noch mehr. Alle Lösungen lassen sich darstellen in der Form

$$a_n = a \left(\frac{1 + \sqrt{5}}{2} \right)^n + b \left(\frac{1 - \sqrt{5}}{2} \right)^n \quad \text{mit } a, b \in \mathbb{R}.$$

- (3) Wir suchen (zweimal differenzierbare) Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$, für die gilt

$$f''(x) + f(x) = 0 \quad \text{für alle } x \in \mathbb{R}.$$

BSP
lineare
Gleichungen

Hier sind die Lösungen gegeben durch

$$f(x) = a \cos x + b \sin x \quad \text{mit } a, b \in \mathbb{R}.$$



Obwohl die betrachteten Objekte ganz unterschiedlich sind (Quadrupel von reellen Zahlen, Folgen reeller Zahlen, zweimal differenzierbare reelle Funktionen), ist die Struktur der Lösungsmenge in allen drei Fällen sehr ähnlich. Dass dies so sein muss, ist ein allgemeines Resultat über lineare Gleichungen. Etwas konkreter äußert sich die Linearität darin, dass die *Summe* zweier Lösungen wieder eine Lösung ist, und dass *Vielfache* einer Lösung wieder Lösungen sind. Diese beiden Operationen, also Addition und Vervielfachung, d.h. Multiplikation mit einem „Skalar“ (in den Beispielen ist das jeweils eine reelle Zahl), ergeben die lineare Struktur, die in der folgenden Definition formalisiert ist.

✱

5.2. Definition. Sei K ein Körper. Ein K -Vektorraum oder Vektorraum über K oder linearer Raum über K ist ein Quintupel $(V, +, \mathbf{0}, -, \cdot)$, bestehend aus einer Menge V , einer Verknüpfung $+: V \times V \rightarrow V$ (genannt Addition), einem Element $\mathbf{0} \in V$, einer Abbildung $-: V \rightarrow V$ und einer Abbildung $\cdot: K \times V \rightarrow V$ (Skalarmultiplikation), sodass $(V, +, \mathbf{0}, -)$ eine kommutative Gruppe ist und die folgenden weiteren Bedingungen („Axiome“) erfüllt sind:

DEF
Vektorraum

(1) $\forall v \in V: 1 \cdot v = v$ (hier ist $1 \in K$ das Einselement des Körpers K).

(2) (Assoziativität der Skalarmultiplikation)
 $\forall \lambda, \mu \in K \forall v \in V: \lambda \cdot (\mu \cdot v) = (\lambda \mu) \cdot v.$

(3) (Distributivgesetze)
 $\forall \lambda, \mu \in K \forall v \in V: (\lambda + \mu) \cdot v = \lambda \cdot v + \mu \cdot v$ und
 $\forall \lambda \in K \forall v, w \in V: \lambda \cdot (v + w) = \lambda \cdot v + \lambda \cdot w.$

Statt $\lambda \cdot v$ schreibt man oft kurz λv . Die Elemente eines Vektorraums werden auch *Vektoren* genannt. $\mathbf{0} \in V$ heißt der *Nullvektor* des Vektorraums V .

Ein $\mathbb{R}$ -Vektorraum heißt auch ein *reeller Vektorraum*, ein $\mathbb{C}$ -Vektorraum ein *komplexer Vektorraum*. $\diamond$

Machen Sie sich klar, dass „+“ in diesen Axiomen zwei verschiedene Bedeutungen hat: Es kann die Addition im Körper K gemeint sein oder die Addition im Vektorraum V !

Der Vollständigkeit halber und zur Erinnerung sind hier noch einmal die vier Axiome für eine kommutative Gruppe $(V, +, \mathbf{0}, -)$ angegeben:

(1) (Assoziativität der Addition) $\forall v_1, v_2, v_3 \in V: (v_1 + v_2) + v_3 = v_1 + (v_2 + v_3).$

(2) (Kommutativität der Addition) $\forall v, w \in V: v + w = w + v.$

(3) (Nullelement) $\forall v \in V: v + \mathbf{0} = v.$

(4) (Negative Elemente) $\forall v \in V: v + (-v) = \mathbf{0}.$

Wir kürzen $v + (-w)$ wie gewohnt zu $v - w$ ab.

Wie üblich kann man sich auf die Angabe von Addition und Skalarmultiplikation beschränken, da das Nullelement und die Negation eindeutig bestimmt sind. Wenn die Verknüpfungen aus dem Kontext klar sind, spricht man einfach nur vom „ K -Vektorraum V “; wenn auch der Körper K aus dem Kontext klar ist, vom „Vektorraum V “.

Wir kommen zu einigen einfachen Eigenschaften.

5.3. Lemma. Sei $(V, +, \mathbf{0}, -, \cdot)$ ein K -Vektorraum. Dann gilt:

LEMMA
Rechenregeln
Vektorraum

- (1) $\forall v \in V: 0 \cdot v = \mathbf{0}$.
- (2) $\forall \lambda \in K: \lambda \cdot \mathbf{0} = \mathbf{0}$.
- (3) $\forall v \in V: (-1) \cdot v = -v$.
- (4) $\forall \lambda \in K \forall v \in V: \lambda \cdot v = \mathbf{0} \iff \lambda = 0 \text{ oder } v = \mathbf{0}$.

Beweis.

- (1) Es ist (mit einem der beiden Distributivgesetze)

$$0 \cdot v + 0 \cdot v = (0 + 0) \cdot v = 0 \cdot v$$

und Addition von $-(0 \cdot v)$ auf beiden Seiten liefert

$$0 \cdot v = 0 \cdot v + 0 \cdot v - 0 \cdot v = 0 \cdot v - 0 \cdot v = \mathbf{0}.$$

- (2) Das geht analog unter Verwendung des anderen Distributivgesetzes:

$$\mathbf{0} = \lambda \cdot \mathbf{0} - \lambda \cdot \mathbf{0} = \lambda \cdot (\mathbf{0} + \mathbf{0}) - \lambda \cdot \mathbf{0} = \lambda \cdot \mathbf{0} + \lambda \cdot \mathbf{0} - \lambda \cdot \mathbf{0} = \lambda \cdot \mathbf{0}.$$

- (3) Es gilt

$$v + (-1) \cdot v = 1 \cdot v + (-1) \cdot v = (1 + (-1)) \cdot v = 0 \cdot v = \mathbf{0},$$

also muss $(-1) \cdot v$ das eindeutig bestimmte Negative $-v$ von v sein.

- (4) Seien $\lambda \in K$ und $v \in V$. Die Implikation „ $\Leftarrow$ “ wurde bereits in den ersten beiden Teilen des Lemmas bewiesen. Es gelte also $\lambda \cdot v = \mathbf{0}$. Ist $\lambda = 0$, dann gilt die rechte Seite. Anderenfalls gibt es $\lambda^{-1} \in K$ und es folgt (mit Teil (2) und der Assoziativität der Skalarmultiplikation)

$$\mathbf{0} = \lambda^{-1} \cdot \mathbf{0} = \lambda^{-1} \cdot (\lambda \cdot v) = (\lambda^{-1} \lambda) \cdot v = 1 \cdot v = v. \quad \square$$

Hier sind einige Beispiele von Vektorräumen:

5.4. Beispiele. Sei K ein Körper.

BSP
Vektorräume

- (1) Der kleinste K -Vektorraum besteht nur aus dem Nullvektor: $V = \{\mathbf{0}\}$ und es gilt $\mathbf{0} + \mathbf{0} = \mathbf{0}$ und $\lambda \cdot \mathbf{0} = \mathbf{0}$ für alle $\lambda \in K$. Dieser Vektorraum heißt der *Null-Vektorraum*. Er ist als Vektorraum nicht besonders interessant, spielt aber in der Linearen Algebra eine ähnliche Rolle wie die leere Menge in der Mengenlehre.
- (2) Das nächste Beispiel ist der Körper K selbst mit seiner Addition und Multiplikation. Die Vektorraum-Axiome entsprechen einem Teil der Körper-Axiome.
- (3) Sehr wichtig ist die folgende Klasse von Beispielen, denn es sind die Standardbeispiele für K -Vektorräume. Als Menge nimmt man K^n , die Menge der n -Tupel von Elementen von K , und die Verknüpfungen definiert man „komponentenweise“:

$$(x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n) = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \quad \text{und} \\ \lambda \cdot (x_1, x_2, \dots, x_n) = (\lambda x_1, \lambda x_2, \dots, \lambda x_n).$$

Dann kann man die Axiome leicht nachprüfen. Wir führen das hier exemplarisch für eines der Distributivgesetze durch:

$$\begin{aligned}
 \lambda \cdot ((x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n)) &= \lambda \cdot (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \\
 &= (\lambda(x_1 + y_1), \lambda(x_2 + y_2), \dots, \lambda(x_n + y_n)) \\
 &= (\lambda x_1 + \lambda y_1, \lambda x_2 + \lambda y_2, \dots, \lambda x_n + \lambda y_n) \\
 &= (\lambda x_1, \lambda x_2, \dots, \lambda x_n) + (\lambda y_1, \lambda y_2, \dots, \lambda y_n) \\
 &= \lambda \cdot (x_1, x_2, \dots, x_n) + \lambda \cdot (y_1, y_2, \dots, y_n).
 \end{aligned}$$

Man sieht, dass das direkt aus dem Distributivgesetz $\lambda(x + y) = \lambda x + \lambda y$ von K folgt. Für die übrigen Axiome geht das ganz analog. In diesem Beispiel sind die beiden vorigen Beispiele als Grenzfälle enthalten: Für $n = 0$ hat die Menge K^0 nur ein Element (das Nulltupel, das keine Komponenten hat) und ist somit ein Null-Vektorraum. Für $n = 1$ ist $K^1 = K$ und man bekommt K als Vektorraum über K . Für $K = \mathbb{R}$ und $K = \mathbb{C}$ haben wir den reellen Vektorraum $\mathbb{R}^n$ und den komplexen Vektorraum $\mathbb{C}^n$ für jedes $n \in \mathbb{N}$.

- (4) Man kann das vorige Beispiel noch verallgemeinern: K^n kann als der Spezialfall $I = \{1, 2, \dots, n\}$ der Menge K^I der Familien von Elementen von K mit Indexmenge I aufgefasst werden. (Zur Erinnerung: Familien $(x_i)_{i \in I}$ mit $x_i \in K$ sind nur eine andere Schreibweise für Abbildungen $I \rightarrow K$.) Man macht K^I zu einem K -Vektorraum, indem man Addition und Skalarmultiplikation wieder komponentenweise definiert:

$$\begin{aligned}
 (x_i)_{i \in I} + (y_i)_{i \in I} &= (x_i + y_i)_{i \in I} \quad \text{und} \\
 \lambda \cdot (x_i)_{i \in I} &= (\lambda x_i)_{i \in I}.
 \end{aligned}$$

Wenn man statt mit Familien mit Abbildungen $I \rightarrow K$ arbeitet, dann sieht das wie folgt aus (in diesem Fall sagt man auch „punktweise“ statt „komponentenweise“):

$$\begin{aligned}
 f + g: I &\longrightarrow K, \quad i \longmapsto f(i) + g(i), & \text{d.h.} \quad (f + g)(i) &= f(i) + g(i) & \text{und} \\
 \lambda \cdot f: I &\longrightarrow K, \quad i \longmapsto \lambda f(i), & \text{d.h.} \quad (\lambda \cdot f)(i) &= \lambda f(i).
 \end{aligned}$$

Das Nachprüfen der Axiome funktioniert im Wesentlichen genauso wie für die n -Tupel. Als Beispiel hier das andere Distributivgesetz (in der Abbildungs-Schreibweise): Seien $\lambda, \mu \in K$ und $f: I \rightarrow K$ eine Abbildung. Dann gilt für $i \in I$:

$$\begin{aligned}
 ((\lambda + \mu) \cdot f)(i) &= (\lambda + \mu)f(i) = \lambda f(i) + \mu f(i) \\
 &= (\lambda \cdot f)(i) + (\mu \cdot f)(i) = (\lambda \cdot f + \mu \cdot f)(i),
 \end{aligned}$$

also folgt $(\lambda + \mu) \cdot f = \lambda \cdot f + \mu \cdot f$. Zum Beispiel können wir den reellen Vektorraum $\mathbb{R}^{\mathbb{R}} = \text{Abb}(\mathbb{R}, \mathbb{R})$ aller reellen Funktionen betrachten oder den Vektorraum $\mathbb{R}^{\mathbb{N}}$ aller Folgen reeller Zahlen.

- (5) Der Körper $\mathbb{C}$ der komplexen Zahlen ist ein reeller Vektorraum: Die Addition ist die von $\mathbb{C}$, die Skalarmultiplikation ist die Multiplikation von $\mathbb{C}$, eingeschränkt auf $\mathbb{R} \times \mathbb{C}$. (Die *Einschränkung* einer Abbildung $f: X \rightarrow Y$ auf eine Teilmenge $T \subset X$ ist die Abbildung $f|_T: T \rightarrow Y$, $x \mapsto f(x)$, bei der der Definitionsbereich verkleinert (also eingeschränkt) wird, die Abbildungsvorschrift aber unverändert bleibt.) Wenn wir $\mathbb{C}$ als $\mathbb{R}^2$ betrachten, dann ist das derselbe reelle Vektorraum wie in (3) oben mit $K = \mathbb{R}$ und $n = 2$.

DEF
Einschränkung

- (6) Ein auf den ersten Blick ganz anders gearteter Vektorraum ist der folgende: Sei X eine Menge. Dann definieren wir eine Addition auf der Potenzmenge $\mathcal{P}(X)$ durch

$$A + B = A \triangle B = (A \setminus B) \cup (B \setminus A)$$

(symmetrische Differenz, siehe Beispiel 3.12) und eine Skalarmultiplikation mit Elementen des Körpers $\mathbb{F}_2 = \{0, 1\}$ in der einzig möglichen Form, nämlich durch $0 \cdot A = \mathbf{0} = \emptyset$ und $1 \cdot A = A$. Dann erhält man einen $\mathbb{F}_2$ -Vektorraum. Man kann die Axiome wieder nachrechnen, aber man tut sich etwas leichter, wenn man sich klar macht, dass die Potenzmenge $\mathcal{P}(X)$ und die Menge $\mathbb{F}_2^X$ der Abbildungen $X \rightarrow \mathbb{F}_2$ einander bijektiv entsprechen durch

$$\begin{aligned} \mathcal{P}(X) &\longrightarrow \mathbb{F}_2^X, & A &\longmapsto \left(x \mapsto \begin{cases} 0 & \text{falls } x \notin A, \\ 1 & \text{falls } x \in A, \end{cases} \right) \\ \mathbb{F}_2^X &\longrightarrow \mathcal{P}(X), & f &\longmapsto \{x \in X \mid f(x) = 1\}. \end{aligned}$$

Dann entsprechen sich auch Addition und Skalarmultiplikation auf beiden Seiten, also folgt die Gültigkeit der Axiome für $\mathcal{P}(X)$ aus ihrer Gültigkeit für $\mathbb{F}_2^X$.



Weitere Beispiele von Vektorräumen erhalten wir als *Untervektorräume* von anderen Vektorräumen; das werden wir im nächsten Abschnitt genauer betrachten.

In den Beispielen 5.1 für lineare Gleichungen vom Beginn dieses Abschnitts sind Lösungen in gewissen reellen Vektorräumen gesucht: Im ersten Beispiel in $\mathbb{R}^4$, im zweiten Beispiel in $\mathbb{R}^{\mathbb{N}}$ und im dritten Beispiel in einem Untervektorraum von $\text{Abb}(\mathbb{R}, \mathbb{R})$.

6. UNTERVEKTORRÄUME

Häufig möchte man, wenn man einen Vektorraum V gegeben hat, nicht mit dem ganzen Vektorraum arbeiten, sondern mit einer Teilmenge. Damit stellt sich die Frage, wann so eine Teilmenge (wenn man die Addition und Skalarmultiplikation darauf einschränkt) selbst wieder ein Vektorraum ist. Damit diese Frage sinnvoll ist, müssen die Addition und Skalarmultiplikation auf der Teilmenge *wohldefiniert* sein, das heißt, dass Summen und Vielfache von Elementen der Teilmenge wieder in der Teilmenge liegen müssen. Außerdem brauchen wir natürlich das Nullelement. Das führt auf folgende Definition:

* **6.1. Definition.** Seien K ein Körper, V ein K -Vektorraum und $U \subset V$ eine Teilmenge von V . Dann heißt U ein *Untervektorraum* oder *linearer Unterraum* von V , wenn U die folgenden Bedingungen erfüllt: **DEF**
Unter-
vektorraum

- (1) $\mathbf{0} \in U$,
- (2) $\forall u_1, u_2 \in U: u_1 + u_2 \in U$
(„ U ist abgeschlossen unter der Addition“),
- (3) $\forall \lambda \in K \forall u \in U: \lambda \cdot u \in U$
(„ U ist abgeschlossen unter der Skalarmultiplikation“). ◇

Wir zeigen gleich, dass diese Definition sinnvoll ist.

6.2. Lemma. Seien K ein Körper, V ein K -Vektorraum und $U \subset V$ ein Untervektorraum. Dann gilt für alle $u \in U$, dass auch $-u$ ein Element von U ist. **LEMMA**
Unter-VR ist
Vektorraum

Wir schreiben $+_U$ für die auf U (auch im Wertebereich) eingeschränkte Addition $+_U: U \times U \rightarrow U$, $(u_1, u_2) \mapsto u_1 + u_2$, $-_U$ für die auf U eingeschränkte Negationsabbildung $-_U: U \rightarrow U$, $u \mapsto -u$, und $\cdot_U$ für die auf U eingeschränkte Skalarmultiplikation $\cdot_U: K \times U \rightarrow U$, $(\lambda, u) \mapsto \lambda \cdot u$. Dann ist $(U, +_U, \mathbf{0}, -_U, \cdot_U)$ ein K -Vektorraum.

Beweis. Die erste Behauptung ist $\forall u \in U: -u \in U$. Das folgt aber aus der Definition, denn $-u = (-1) \cdot u$, vgl. Lemma 5.3. Deshalb und nach der Definition können wir $+_U$, $-_U$ und $\cdot_U$ wie angegeben definieren (denn die Bilder liegen jeweils in U). Es bleiben die Vektorraum-Axiome für U nachzuprüfen. Diese haben aber alle die Form von „Allaussagen“, es wird also verlangt, dass eine Aussage für alle Elemente $u_1, u_2, \dots$ von U gilt. Da V ein Vektorraum ist, gelten diese Aussagen aber sogar für alle Elemente von V , also erst recht für alle Elemente von U . □

In der Literatur finden Sie meistens eine Definition von „Vektorraum“ (und analog für Gruppen, Ringe, Körper, ...), die von dem Tripel $(V, +, \cdot)$ ausgeht und dann die *Existenz* eines Nullelements und von Inversen bezüglich der Addition fordert. Im Gegensatz dazu haben wir hier das Nullelement und die Negationsabbildung mit in die „Daten“ des Vektorraums aufgenommen. Der Vorteil ist, dass die Axiome dann alle zu Allaussagen werden, die man leichter nachprüfen kann, wie im obigen Beweis. Auf der anderen Seite muss man sich aber vorher überlegen, was das Nullelement ist und wie die Negationsabbildung aussieht. Im gerade bewiesenen Lemma geschieht dies dadurch, dass wir zeigen, dass U auch unter der Negation abgeschlossen ist, sodass wir die Negationsabbildung $-_U$ definieren können. Wenn man die andere Formulierung der Axiome benutzt, dann muss man diesen Beweisschritt ausführen, wenn man die Existenz des zu u negativen Elements zeigt. Im Endeffekt muss man also das Gleiche tun, nur die Reihenfolge ist etwas anders. Man beachte allerdings, dass die logische Struktur der Axiome bei der anderen

Variante der Definitionen komplizierter ist. Für eine Gruppe $(G, *)$ hat man neben dem Assoziativgesetz dann eine geschachtelte Existenzaussage

$$\exists e \in G: \left(\forall a \in G: (e * a = a = a * e \wedge \exists a' \in G: a' * a = e = a * a') \right)$$

statt zwei reinen Allaussagen.

Die Schreibweise $+_U$ usw. für die auf U eingeschränkten Abbildungen diene nur der Verdeutlichung für die Formulierung des Lemmas. Wir schreiben normalerweise einfach $+$ usw. für die Addition usw. auf U .

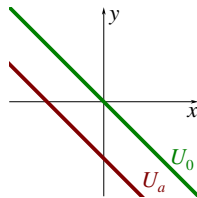
6.3. Beispiele. Jeder Vektorraum V hat die Untervektorräume $U = \{0\} \subset V$ (ein Null-Vektorraum) und $U = V$.

BSP
♣ triviale
Unter-VR

6.4. Beispiel. Sei $a \in \mathbb{R}$. Wir betrachten den reellen Vektorraum $V = \mathbb{R}^2$ und setzen $U_a = \{(x, y) \in \mathbb{R}^2 \mid x + y = a\}$. Für welche a ist U_a ein Untervektorraum von $\mathbb{R}^2$?

BSP
Unter-VR
von $\mathbb{R}^2$

Dazu müssen wir die Bedingungen in der Definition nachprüfen. Die erste davon sagt, dass der Nullvektor $0 = (0, 0)$ ein Element von U_a sein muss. Das bedeutet $0 + 0 = a$, also ist das nur für $a = 0$ möglich. Wir prüfen die beiden anderen Bedingungen:



- U_0 ist abgeschlossen unter der Addition, denn für Elemente $u_1 = (x_1, y_1)$ und $u_2 = (x_2, y_2)$ von U_0 gilt $u_1 + u_2 = (x_1 + x_2, y_1 + y_2)$ und

$$(x_1 + x_2) + (y_1 + y_2) = (x_1 + y_1) + (x_2 + y_2) = 0 + 0 = 0,$$

also ist $u_1 + u_2 \in U_0$.

- U_0 ist abgeschlossen unter der Skalarmultiplikation, denn für $u = (x, y) \in U_0$ und $\lambda \in \mathbb{R}$ ist $\lambda \cdot u = (\lambda x, \lambda y)$ und es gilt

$$\lambda x + \lambda y = \lambda(x + y) = \lambda \cdot 0 = 0,$$

also ist $\lambda \cdot u \in U_0$.

♣

Weitere interessante Beispiele sind „Folgenräume“ und „Funktionsräume“, die als Untervektorräume des Vektorraums $\mathbb{R}^{\mathbb{N}}$ der Folgen reeller Zahlen oder des Vektorraums $\text{Abb}(\mathbb{R}, \mathbb{R})$ der reellen Funktionen auftreten.

6.5. Beispiele. Sei $V = \mathbb{R}^{\mathbb{N}}$ der reelle Vektorraum, dessen Elemente alle Folgen reeller Zahlen sind.

BSP
Folgenräume

- (1) Sei $U_b = \{(a_n)_{n \in \mathbb{N}} \in \mathbb{R}^{\mathbb{N}} \mid (a_n)_{n \in \mathbb{N}} \text{ ist beschränkt}\}$. Dann ist U_b ein Untervektorraum von $\mathbb{R}^{\mathbb{N}}$.

Beweis. Wir prüfen die Bedingungen nach. Die konstante Nullfolge (mit $a_n = 0$ für alle $n \in \mathbb{N}$) ist beschränkt, also gilt $0 \in U_b$. Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ zwei beschränkte Folgen. Dann gibt es $A, B \in \mathbb{R}$ mit $|a_n| \leq A$ und $|b_n| \leq B$ für alle $n \in \mathbb{N}$. Es folgt $|a_n + b_n| \leq A + B$, also ist auch die Summenfolge $(a_n)_{n \in \mathbb{N}} + (b_n)_{n \in \mathbb{N}} = (a_n + b_n)_{n \in \mathbb{N}}$ beschränkt. Ist zusätzlich $\lambda \in \mathbb{R}$, dann gilt $|\lambda a_n| \leq |\lambda|A$, also ist auch die Folge $\lambda \cdot (a_n)_{n \in \mathbb{N}} = (\lambda a_n)_{n \in \mathbb{N}}$ beschränkt. $\square$

- (2) Sei $U_n = \{(a_n)_{n \in \mathbb{N}} \in \mathbb{R}^{\mathbb{N}} \mid (a_n)_{n \in \mathbb{N}} \text{ ist eine Nullfolge}\}$. Dann ist U_n ein Untervektorraum von $\mathbb{R}^{\mathbb{N}}$ (oder auch von U_b). (Übung.)

- (3) Sei $U_k = \{(a_n)_{n \in \mathbb{N}} \in \mathbb{R}^{\mathbb{N}} \mid (a_n)_{n \in \mathbb{N}} \text{ konvergiert}\}$. Dann ist U_k ein Untervektorraum von $\mathbb{R}^{\mathbb{N}}$ (oder auch von U_b).

Beweis. Die konstante Nullfolge konvergiert (gegen 0), also ist sie in U_k . In der Analysis lernen Sie, dass die Summe zweier konvergenter Folgen wieder konvergiert und dass jedes Vielfache einer konvergenten Folge konvergiert. Damit sind die drei Bedingungen erfüllt. $\square$

Für diese drei Untervektorräume gilt $U_n \subset U_k \subset U_b$ (denn jede Nullfolge konvergiert gegen 0 und jede konvergente Folge ist beschränkt, vgl. Analysis). $\clubsuit$

6.6. Beispiele. Sei $V = \text{Abb}(\mathbb{R}, \mathbb{R})$ der reelle Vektorraum, dessen Elemente alle Funktionen $\mathbb{R} \rightarrow \mathbb{R}$ sind.

BSP
Funktions-
räume

- (1) Sei $\mathcal{C}(\mathbb{R}) = \{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) \mid f \text{ ist stetig}\}$. Dann ist $\mathcal{C}(\mathbb{R})$ ein Untervektorraum von V .

Beweis. Die Nullfunktion $x \mapsto 0$ ist stetig. In der Analysis lernen Sie, dass Summen und Vielfache stetiger Funktionen wieder stetig sind. $\square$

- (2) Sei $n \in \mathbb{N}$ und $\mathcal{C}^n(\mathbb{R}) = \{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) \mid f \text{ ist } n\text{-mal differenzierbar und } f^{(n)} \text{ ist stetig}\}$ der Raum der n -mal stetig differenzierbaren Funktionen. Aus Ergebnissen der Analysis folgt, dass $\mathcal{C}^n(\mathbb{R})$ ein Untervektorraum von V ist.

- (3) Sei $a > 0$ und $\mathcal{P}(a) = \{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) \mid \forall x \in \mathbb{R}: f(x+a) = f(x)\}$ die Menge der periodischen Funktionen mit Periode a (zum Beispiel sind $\sin$ und $\cos$ Elemente von $\mathcal{P}(2\pi)$). Dann ist $\mathcal{P}(a)$ ein Untervektorraum von V .

Beweis. Die Nullfunktion ist periodisch, also ein Element von $\mathcal{P}(a)$. Seien $f, g \in \mathcal{P}(a)$ und $\lambda \in \mathbb{R}$. Wir zeigen $f+g, \lambda f \in \mathcal{P}(a)$: Für alle $x \in \mathbb{R}$ gilt

$$(f+g)(x+a) = f(x+a) + g(x+a) \stackrel{f, g \in \mathcal{P}(a)}{=} f(x) + g(x) = (f+g)(x) \quad \text{und}$$

$$(\lambda f)(x+a) = \lambda f(x+a) \stackrel{f \in \mathcal{P}(a)}{=} \lambda f(x) = (\lambda f)(x).$$

Damit sind alle drei Bedingungen erfüllt. $\square$



Auch in der *Codierungstheorie* spielt der Begriff des Untervektorraums eine sehr wichtige Rolle.

6.7. Beispiel. Seien F ein endlicher Körper (zum Beispiel $F = \mathbb{F}_2$) und $n \in \mathbb{N}$. Dann heißt ein Untervektorraum von F^n ein *linearer Code* der Länge n über F . Ein Beispiel ist der *Hamming-Code* der Länge 7 über $\mathbb{F}_2$, der gegeben ist durch

BSP
Lineare
Codes

$$H = \{(x_1, x_2, x_3, x_4, x_1+x_2+x_4, x_1+x_3+x_4, x_2+x_3+x_4) \in \mathbb{F}_2^7 \mid x_1, x_2, x_3, x_4 \in \mathbb{F}_2\}.$$

In der Codierungstheorie interessiert man sich dann für die „Größe“ (genauer: die *Dimension*, die wir bald einführen werden) des Codes und dafür, wie viele Fehler er korrigieren kann. Dafür ist wichtig, dass je zwei verschiedene Codewörter (also Elemente des Codes) sich an möglichst vielen Stellen unterscheiden. Wegen der linearen Struktur kann man Differenzen bilden und daher annehmen, dass eines der Codewörter null ist. Dann ist die Frage, an mindestens wie vielen Stellen ein von $\mathbf{0}$ verschiedenes Codewort eine von 0 verschiedene Komponente hat. Für den Hamming-Code H ist diese „Minimaldistanz“ 3, was bedeutet, dass er „einen

Fehler korrigieren“ kann. (Wenn ein Codewort an einer Stelle verändert wird, kann man es rekonstruieren, da sich jedes andere Codewort von dem veränderten Wort an mindestens zwei Stellen unterscheidet.) ♣

7. ERZEUGENDENSYSTEME

Wir erinnern uns an die Beispiele von Funktionenräumen im letzten Abschnitt. Dort hatten wir gesehen, dass der Raum $\mathcal{C}(\mathbb{R})$ der stetigen reellen Funktionen und der Raum $\mathcal{P}(a)$ der a -periodischen reellen Funktionen beides Untervektorräume von $\text{Abb}(\mathbb{R}, \mathbb{R})$ sind. Wie sieht es mit stetigen periodischen Funktionen aus? Muss $\mathcal{C}(\mathbb{R}) \cap \mathcal{P}(a)$ auch ein Untervektorraum sein?

7.1. Lemma. *Sei V ein K -Vektorraum mit zwei Untervektorräumen U_1 und U_2 . Dann ist $U_1 \cap U_2$ ebenfalls ein Untervektorraum von V .*

LEMMA
Schnitt von
zwei UVR

Beweis. Wir müssen die drei Bedingungen aus der Definition von „Untervektorraum“ nachprüfen.

- (1) Da U_1 und U_2 Untervektorräume sind, ist $\mathbf{0} \in U_1$ und $\mathbf{0} \in U_2$, also auch $\mathbf{0} \in U_1 \cap U_2$.
- (2) Seien $u, u' \in U_1 \cap U_2$. Dann gilt $u, u' \in U_1$ und $u, u' \in U_2$. Da U_1 und U_2 Untervektorräume sind, folgt $u + u' \in U_1$ und $u + u' \in U_2$, also auch $u + u' \in U_1 \cap U_2$.
- (3) Seien $\lambda \in K$ und $u \in U_1 \cap U_2$. Dann ist $u \in U_1$ und $u \in U_2$. Da U_1 und U_2 Untervektorräume sind, folgt $\lambda u \in U_1$ und $\lambda u \in U_2$, also auch $\lambda u \in U_1 \cap U_2$. $\square$

Wir wollen diese Aussage jetzt auf Durchschnitte von beliebig (auch unendlich) vielen Untervektorräumen verallgemeinern. Dazu führen wir erst eine Schreibweise für Vereinigungen und Durchschnitte von vielen Mengen ein.

7.2. Definition. Ist $(A_i)_{i \in I}$ eine Familie von Mengen, dann schreiben wir

$$\bigcup_{i \in I} A_i = \{x \mid \exists i \in I: x \in A_i\}$$

für die Vereinigung aller Mengen A_i . (Ist I die leere Menge, dann ist diese Vereinigung ebenfalls leer.) Ist $I \neq \emptyset$, dann schreiben wir analog

$$\bigcap_{i \in I} A_i = \{x \mid \forall i \in I: x \in A_i\}$$

für den Durchschnitt aller Mengen A_i .

Ist $\mathcal{M}$ eine Menge, deren Elemente selbst Mengen sind, dann schreiben wir

$$\bigcup \mathcal{M} = \bigcup_{A \in \mathcal{M}} A = \{x \mid \exists A \in \mathcal{M}: x \in A\}$$

für die Vereinigung all dieser Mengen und, falls $\mathcal{M}$ nicht leer ist,

$$\bigcap \mathcal{M} = \bigcap_{A \in \mathcal{M}} A = \{x \mid \forall A \in \mathcal{M}: x \in A\}$$

für ihren Durchschnitt. $\diamond$

Im Fall $I = \emptyset$ wäre die Bedingung $\forall i \in I: x \in A_i$ für alle x erfüllt und man bekäme die Menge, die alles enthält. Diese Menge kann es aber nicht geben, denn sie würde die Menge enthalten, die zur Russellschen Antinomie führt, siehe die Bemerkungen zur Mengenlehre am Ende von Abschnitt 2.

Damit können wir jetzt die Verallgemeinerung formulieren:

DEF
 $\bigcup_{i \in I} A_i$
 $\bigcap_{i \in I} A_i$
 $\bigcup \mathcal{M}$
 $\bigcap \mathcal{M}$

7.3. Lemma. Sei V ein K -Vektorraum und sei $(U_i)_{i \in I}$ eine Familie von Untervektorräumen von V mit $I \neq \emptyset$. Dann ist

$$U = \bigcap_{i \in I} U_i$$

ebenfalls ein Untervektorraum von V .

Für $I = \{1, 2, 3, \dots, n\}$ (mit $n \geq 1$) haben wir den Spezialfall

$$\begin{aligned} U_1, U_2, \dots, U_n \subset V \text{ Untervektorräume} \\ \implies U_1 \cap U_2 \cap \dots \cap U_n \subset V \text{ Untervektorraum.} \end{aligned}$$

Beweis. Wir müssen die Bedingungen aus Definition 6.1 für U nachprüfen.

- (1) Da jede Teilmenge U_i ein Untervektorraum von V ist, gilt $\forall i \in I: \mathbf{0} \in U_i$. Das bedeutet gerade $\mathbf{0} \in U$.
- (2) Seien $u_1, u_2 \in U$. Nach Definition von U bedeutet das $\forall i \in I: u_1, u_2 \in U_i$. Da alle U_i Untervektorräume von V sind, folgt $\forall i \in I: u_1 + u_2 \in U_i$, also $u_1 + u_2 \in U$.
- (3) Sei $\lambda \in K$ und $u \in U$. Dann gilt $\forall i \in I: u \in U_i$. Da alle U_i Untervektorräume von V sind, folgt $\forall i \in I: \lambda u \in U_i$, also $\lambda u \in U$. $\square$

7.4. Beispiel. Der Raum

$$\mathcal{C}(\mathbb{R}) \cap \mathcal{P}(a) = \{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) \mid f \text{ ist stetig und } a\text{-periodisch}\}$$

ist ein Untervektorraum von $\text{Abb}(\mathbb{R}, \mathbb{R})$. 

Wie sieht es mit Vereinigungen von Untervektorräumen aus? Im Allgemeinen erhält man daraus *keinen* Untervektorraum. Die Vereinigung von zwei Untervektorräumen U_1 und U_2 zum Beispiel ist nur dann wieder ein Untervektorraum, wenn einer der beiden im anderen enthalten ist (Übung). Man hat aber immerhin das folgende Resultat. 

7.5. Lemma. Sei V ein K -Vektorraum und sei $(U_n)_{n \in \mathbb{N}}$ eine aufsteigende Folge von Untervektorräumen von V (d.h. $U_n \subset U_{n+1}$ für alle $n \in \mathbb{N}$). Dann ist

$$U = \bigcup_{n \in \mathbb{N}} U_n$$

ebenfalls ein Untervektorraum von V .

Beweis. Wir prüfen die Bedingungen für U .

- (1) $\mathbf{0} \in U_0$, also ist auch $\mathbf{0} \in U$.
- (2) Seien $u_1, u_2 \in U$. Dann gibt es $n_1, n_2 \in \mathbb{N}$ mit $u_1 \in U_{n_1}$ und $u_2 \in U_{n_2}$. Sei n die größere der beiden Zahlen n_1 und n_2 . Da wir eine aufsteigende Folge von Untervektorräumen haben, gilt dann $U_{n_1} \subset U_n$ und $U_{n_2} \subset U_n$ und damit $u_1, u_2 \in U_n$. Da U_n ein Untervektorraum ist, folgt $u_1 + u_2 \in U_n \subset U$.
- (3) Sei $\lambda \in K$ und $u \in U$. Dann gibt es $n \in \mathbb{N}$, sodass $u \in U_n$ ist. Da U_n ein Untervektorraum ist, folgt $\lambda u \in U_n \subset U$. $\square$

Lemma 7.3 erlaubt es uns nun, den kleinsten Untervektorraum zu konstruieren, der eine gegebene Teilmenge eines Vektorraums V enthält.

LEMMA
Durchschnitt
von Unter-VR

BSP

LEMMA
aufsteigende
Vereinigung
von Unter-VR

* **7.6. Definition.** Sei V ein K -Vektorraum und $A \subset V$ eine beliebige Teilmenge von V . Dann heißt der Untervektorraum

DEF
Lineare
Hülle

$$\langle A \rangle = \langle A \rangle_K = \bigcap \{U \subset V \mid U \text{ Untervektorraum von } V \text{ und } A \subset U\}$$

(also der Durchschnitt aller A enthaltenden Untervektorräume von V) der von A erzeugte oder aufgespannte Untervektorraum von V , die $(K\text{-})$ lineare Hülle von A oder der $(K\text{-})$ Spann von A . Ist $A = \{v_1, v_2, \dots, v_n\}$ endlich, dann schreiben wir auch

$$\langle v_1, v_2, \dots, v_n \rangle \quad \text{oder} \quad \langle v_1, v_2, \dots, v_n \rangle_K$$

statt $\langle A \rangle$ oder $\langle A \rangle_K$.

Wir werden statt Mengen häufiger auch Familien $(v_i)_{i \in I}$ von Elementen von V betrachten. Wir schreiben dann

$$\langle (v_i)_{i \in I} \rangle \quad \text{oder auch} \quad \langle v_i \mid i \in I \rangle$$

für den von allen v_i erzeugten Untervektorraum $\langle \{v_i \mid i \in I\} \rangle$. $\diamond$

Lemma 7.3 garantiert uns, dass $\langle A \rangle$ tatsächlich ein Untervektorraum von V ist, denn $\langle A \rangle$ ist definitionsgemäß der Durchschnitt einer nichtleeren Menge von Untervektorräumen (nichtleer, weil V selbst immer ein A enthaltender Untervektorraum von V ist).

Wir benutzen die Schreibweise $\langle A \rangle_K$, um zu verdeutlichen, welcher Körper zugrunde gelegt wird. Zum Beispiel gilt im $\mathbb{R}$ -Vektorraum $\mathbb{C}$, dass $\langle 1 \rangle = \langle 1 \rangle_{\mathbb{R}} = \mathbb{R}$ ist. Wird $\mathbb{C}$ aber als $\mathbb{C}$ -Vektorraum betrachtet, dann haben wir $\langle 1 \rangle = \langle 1 \rangle_{\mathbb{C}} = \mathbb{C}$.

7.7. Beispiel. In Definition 7.6 können wir für A die leere Menge wählen. Was ist der von A erzeugte Untervektorraum?

BSP
 $\langle \emptyset \rangle = \{0\}$

Da jeder Untervektorraum von V die leere Menge enthält, müssen wir den Durchschnitt über alle Untervektorräume von V bilden. Da jeder Untervektorraum den Nullvektor enthält und $\{0\}$ ein Untervektorraum ist, folgt $\langle \emptyset \rangle = \{0\}$. $\clubsuit$

* **7.8. Definition.** Sei V ein K -Vektorraum und $E \subset V$ eine Teilmenge von V . Dann heißt E ein $(K\text{-})$ Erzeugendensystem von V , wenn $V = \langle E \rangle$ gilt. Analog heißt eine Familie $(v_i)_{i \in I}$ von Elementen von V ein $(K\text{-})$ Erzeugendensystem von V , wenn $V = \langle (v_i)_{i \in I} \rangle$ ist. $\diamond$

DEF
Erzeugenden-
system

Zum Beispiel ist die leere Menge ein Erzeugendensystem des Null-Vektorraums.

Definition 7.6 ist sehr elegant, aber nicht besonders praktisch, weil sie uns nicht sagt, „wie die lineare Hülle von A aussieht“, also was ihre Elemente sind. In gewisser Weise ist es eine Definition „von oben“ — wir betrachten alle Untervektorräume, die mindestens so groß sind wie gewünscht, und wählen dann den kleinsten (im Sinne der Inklusion von Mengen) aus. (Das ist übrigens völlig analog zur Definition des Abschlusses $\bar{A}$ einer Menge A als Durchschnitt aller A enthaltenden abgeschlossenen Mengen oder auch zur Definition des Supremums einer Menge reeller Zahlen als kleinste obere Schranke.) Was wir aber gerne hätten, ist eine Definition „von unten“, die die Elemente von $\langle A \rangle$ aus den Elementen von A konstruiert.

Dafür betrachten wir als Beispiel ein Paar (v_1, v_2) von Vektoren in V . Welche Elemente muss $\langle v_1, v_2 \rangle$ mindestens enthalten?

Wir wissen, dass v_1 und v_2 Elemente von $\langle v_1, v_2 \rangle$ sind; außerdem ist $\langle v_1, v_2 \rangle$ ein Untervektorraum, also unter Addition und Skalarmultiplikation abgeschlossen. Es

müssen also insbesondere Summen von Vielfachen von v_1 und v_2 in $\langle v_1, v_2 \rangle$ enthalten sein:

$$\{\lambda_1 v_1 + \lambda_2 v_2 \mid \lambda_1, \lambda_2 \in K\} \subset \langle v_1, v_2 \rangle.$$

Auf der anderen Seite überlegt man sich leicht, dass diese Menge selbst schon ein Untervektorraum von V ist. Da dieser Untervektorraum v_1 und v_2 enthält und gleichzeitig in allen v_1 und v_2 enthaltenden Untervektorräumen enthalten ist, muss er gleich $\langle v_1, v_2 \rangle$ sein. (Das ist analog zu unserer Konstruktion des Körpers $\mathbb{C}$: Wir haben erst überlegt, welche Elemente er enthalten muss, und dann gezeigt, dass diese bereits ausreichen, da sie einen Körper bilden.) Diese Beobachtung lässt sich verallgemeinern.

7.9. Satz. Sei V ein K -Vektorraum.

SATZ
Beschreibung
von $\langle A \rangle$

(1) Sind $v_1, v_2, \dots, v_n \in V$, dann gilt

$$\langle v_1, v_2, \dots, v_n \rangle = \{\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n \mid \lambda_1, \lambda_2, \dots, \lambda_n \in K\}.$$

(2) Ist $A \subset V$ beliebig, dann gilt

$$\langle A \rangle = \{\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n \mid n \in \mathbb{N}, v_1, v_2, \dots, v_n \in A, \lambda_1, \lambda_2, \dots, \lambda_n \in K\}.$$

(3) Ist $(v_i)_{i \in I}$ eine Familie von Elementen von V , dann gilt

$$\langle v_i \mid i \in I \rangle = \{\lambda_1 v_{i_1} + \lambda_2 v_{i_2} + \dots + \lambda_n v_{i_n} \mid n \in \mathbb{N}, i_1, \dots, i_n \in I, \lambda_1, \dots, \lambda_n \in K\}.$$

Für $n = 0$ setzen wir dabei $\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mathbf{0}$ („leere Summe“).

Beweis.

(1) Sei $U = \{\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n \mid \lambda_1, \lambda_2, \dots, \lambda_n \in K\}$ die Menge auf der rechten Seite der Gleichung. Da $v_1, v_2, \dots, v_n \in \langle v_1, v_2, \dots, v_n \rangle$ und $\langle v_1, v_2, \dots, v_n \rangle$ unter Skalarmultiplikation und Addition abgeschlossen ist, muss jedes Element der Form $\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n$ ebenfalls in $\langle v_1, v_2, \dots, v_n \rangle$ liegen. Es gilt also $U \subset \langle v_1, v_2, \dots, v_n \rangle$.

Auf der anderen Seite gilt $v_1, v_2, \dots, v_n \in U$ (wähle $\lambda_j = 1$ und $\lambda_i = 0$ für alle $i \in \{1, 2, \dots, n\} \setminus \{j\}$, um zu sehen, dass $v_j \in U$ ist) und U ist ein Untervektorraum von V :

- $\mathbf{0} \in U$ (setze $\lambda_i = 0$ für alle i).
- U ist abgeschlossen unter der Addition, denn

$$\begin{aligned} (\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n) + (\mu_1 v_1 + \mu_2 v_2 + \dots + \mu_n v_n) \\ = (\lambda_1 + \mu_1) v_1 + (\lambda_2 + \mu_2) v_2 + \dots + (\lambda_n + \mu_n) v_n. \end{aligned}$$

- U ist abgeschlossen unter der Skalarmultiplikation, denn

$$\lambda(\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n) = (\lambda \lambda_1) v_1 + (\lambda \lambda_2) v_2 + \dots + (\lambda \lambda_n) v_n.$$

Da U ein $v_1, v_2, \dots, v_n$ enthaltender Untervektorraum von V ist, folgt (nach Definition 7.6) $\langle v_1, v_2, \dots, v_n \rangle \subset U$; insgesamt erhalten wir die behauptete Gleichheit.

(2) Sei wieder U die Menge auf der rechten Seite der Gleichung. Wie in Teil (1) ist klar, dass $U \subset \langle A \rangle$ ist. Es gilt wieder, dass U ein A enthaltender Untervektorraum ist. Die einzige Schwierigkeit tritt beim Nachweis der Abgeschlossenheit unter der Addition auf, denn in den beiden zu addierenden Summen können verschiedene Elemente von A auftreten. Da aber nicht

vorausgesetzt ist, dass die auftretenden Elemente paarweise verschieden¹ sein müssen, können wir die beiden Summen einfach „formal“ addieren:

$$\begin{aligned} (\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n) + (\mu_1 w_1 + \mu_2 w_2 + \dots + \mu_m w_m) \\ = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n + \lambda_{n+1} v_{n+1} + \dots + \lambda_{n+m} v_{n+m}, \end{aligned}$$

wenn wir $\lambda_{n+j} = \mu_j$ und $v_{n+j} = w_j$ setzen für $j \in \{1, 2, \dots, m\}$.

(3) Das folgt wegen $\langle v_i \mid i \in I \rangle = \langle \{v_i \mid i \in I\} \rangle$ aus der vorigen Aussage. $\square$

Es ist eine gute Übung, sich zu überlegen, an welcher Stelle in diesem Beweis welche der Vektorraum-Axiome verwendet werden.

Weil die Ausdrücke der Form $\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n$ so wichtig sind, haben sie einen eigenen Namen.

7.10. Definition. Seien K ein Körper und V ein K -Vektorraum.

DEF
Linear-
kombination

(1) Sind $v_1, v_2, \dots, v_n \in V$ und $\lambda_1, \lambda_2, \dots, \lambda_n \in K$, dann heißt

$$\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n$$

eine *(K-)Linearkombination* von $v_1, v_2, \dots, v_n$.

Dabei heißt λ_j der *Koeffizient* von v_j in der Linearkombination.

(2) Ist $A \subset V$ eine beliebige Teilmenge von V , dann heißt jede K -Linearkombination von Elementen $v_1, v_2, \dots, v_n \in A$ eine *(K-)Linearkombination von Elementen von A*.

(3) Ist $(v_i)_{i \in I}$ eine Familie von Elementen von V , dann heißt jede K -Linearkombination von Vektoren $v_{i_1}, v_{i_2}, \dots, v_{i_n}$ mit $i_1, i_2, \dots, i_n \in I$ eine *(K-)Linearkombination der Familie $(v_i)_{i \in I}$* . $\diamond$

Satz 7.9 kann dann so formuliert werden:

Die lineare Hülle von $A \subset V$ besteht genau aus allen Linearkombinationen von Elementen von A. Die entsprechende Aussage gilt für Familien.

Eine Teilmenge $E \subset V$ ist genau dann ein Erzeugendensystem von V , wenn jedes Element von V eine Linearkombination von Elementen von E ist. Die analoge Aussage gilt für Familien als Erzeugendensysteme.

Warnung. In einer Linearkombination kommen immer nur **endlich viele** Elemente vor! In der Linearen Algebra gibt es (im Gegensatz zur Analysis) keine unendlichen Summen!



7.11. Definition. Um Summen von beliebig vielen Elementen präzise und kurz hinschreiben zu können, führen wir folgende Summenschreibweise ein: Wir schreiben

DEF
 $\sum_{i \in I} a_i$

$$\sum_{i \in I} a_i \quad \text{bzw.} \quad \sum_{i=1}^n a_i$$

für die Summe der Glieder der Familie $(a_i)_{i \in I}$ bzw. für die Summe der Komponenten des n -Tupels $(a_1, a_2, \dots, a_n)$. Dabei sind die a_i aus einer kommutativen Gruppe (bei uns fast immer Elemente eines Vektorraums) und die Menge I ist

¹„ $v_1, v_2, \dots, v_n$ sind paarweise verschieden“ bedeutet „ $\forall i, j \in \{1, 2, \dots, n\}: i \neq j \Rightarrow v_i \neq v_j$ “.

endlich. Ist I leer (bzw. $n = 0$), dann ist der Wert dieser „leeren Summe“ das Nullelement der Gruppe. Eine Linearkombination kann dann in der Form

$$\sum_{i=1}^n \lambda_i v_i$$

geschrieben werden.

Für unendliche Indexmengen I verlangen wir, dass alle bis auf endlich viele Summanden a_i null sind, und setzen dann

$$\sum_{i \in I} a_i = \sum_{i \in J} a_i,$$

wobei $J \subset I$ endlich ist mit $a_i = 0$ für alle $i \in I \setminus J$; die rechts stehende Summe ist wieder endlich. In dieser Schreibweise sind dann die Linearkombinationen der Familie $(v_i)_{i \in I}$ gegeben durch

$$\sum_{i \in I} \lambda_i v_i,$$

wobei $(\lambda_i)_{i \in I}$ eine Familie von Skalaren (Elementen von K) ist, sodass nur endlich viele $\lambda_i \neq 0$ sind. Wir schreiben $K^{(I)}$ für die Menge aller dieser Familien von Skalaren. $\diamond$

7.12. Beispiel. Sei K ein Körper und $n \in \mathbb{N}$. Im Standard-Vektorraum K^n haben wir die Elemente

$$\mathbf{e}_1 = (1, 0, 0, \dots, 0), \quad \mathbf{e}_2 = (0, 1, 0, \dots, 0), \quad \dots, \quad \mathbf{e}_n = (0, 0, 0, \dots, 0, 1).$$

Dabei sind alle Komponenten von $\mathbf{e}_j$ null mit Ausnahme der j -ten, die den Wert 1 hat. $(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$ ist ein Erzeugendensystem von K^n , denn jedes Element von K^n ist eine Linearkombination dieser Elemente:

$$(x_1, x_2, \dots, x_n) = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n.$$



BSP
Erzeugenden-
system
von K^n

7.13. Beispiel. Ein Vektorraum hat im Allgemeinen viele Erzeugendensysteme. Zum Beispiel sind

$$(\mathbf{e}_1, \mathbf{e}_2), \quad ((1, 1), (1, -1)), \quad \{(1, 2), (2, 3), (3, 4)\}, \quad \mathbb{Z} \times \mathbb{Z} \quad \text{und} \quad \mathbb{R} \times \mathbb{R}$$

alles Erzeugendensysteme des reellen Vektorraums $V = \mathbb{R}^2$.



BSP
Viele
Erzeugenden-
systeme

7.14. Beispiel. Im reellen Vektorraum $V = \text{Abb}(\mathbb{R}, \mathbb{R})$ betrachten wir die *Potenzfunktionen*

$$f_0: x \mapsto 1, \quad f_1: x \mapsto x, \quad f_2: x \mapsto x^2, \quad \dots, \quad f_n: x \mapsto x^n, \quad \dots$$

Wie sieht der von $(f_0, f_1, f_2, \dots) = (f_n)_{n \in \mathbb{N}}$ erzeugte Untervektorraum P von V aus?

Seine Elemente sind gerade die Linearkombinationen von endlich vielen der Potenzfunktionen. Indem wir eventuell Potenzfunktionen mit Koeffizient 0 hinzufügen (was am Wert der Linearkombination nichts ändert) und gleichartige Terme zusammenfassen, können wir annehmen, dass die Linearkombination die Form

$$f = a_0 f_0 + a_1 f_1 + \dots + a_n f_n$$

hat mit $n \in \mathbb{N}$ und $a_0, a_1, \dots, a_n \in \mathbb{R}$. Dann gilt

$$f(x) = a_0 f_0(x) + a_1 f_1(x) + \dots + a_n f_n(x) = a_0 + a_1 x + a_2 x^2 + \dots + a_n x^n.$$

BSP
Vektorraum
der Polynom-
funktionen

Die Elemente von P sind also gerade die *Polynomfunktionen*. ♣

Wir notieren noch einige einfache Eigenschaften der linearen Hülle.

7.15. **Lemma.** *Sei V ein K -Vektorraum.*

LEMMA

(1) *Für Teilmengen $A \subset B \subset V$ gilt $\langle A \rangle \subset \langle B \rangle$.*

Eigensch.

(2) *Sei E ein Erzeugendensystem von V . Eine Teilmenge $A \subset V$ ist genau dann ein Erzeugendensystem von V , wenn $E \subset \langle A \rangle$ gilt.*

lineare

Hülle

Die analogen Aussagen gelten natürlich auch für Familien statt Teilmengen.

Beweis. Übung. □

8. LINEARE UNABHÄNGIGKEIT

Wir haben gesehen, dass ein K -Vektorraum V sehr viele Erzeugendensysteme haben kann; eines davon ist zum Beispiel die Menge V selbst. Das erscheint aber ein wenig verschwenderisch, sodass sich die Frage stellt, ob es auch minimale Erzeugendensysteme gibt und wie sie gegebenenfalls charakterisiert werden können. Dazu überlegen wir Folgendes: Sei E ein Erzeugendensystem von V , das nicht minimal ist in dem Sinn, dass es ein Element $v_0 \in E$ gibt, sodass $E_0 = E \setminus \{v_0\}$ auch schon ein Erzeugendensystem von V ist. Dann können wir v_0 als Linearkombination von Elementen von E_0 schreiben:

$$v_0 = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n$$

mit $v_1, v_2, \dots, v_n \in E_0$ und $\lambda_1, \lambda_2, \dots, \lambda_n \in K$. Dabei können wir annehmen, dass $v_1, v_2, \dots, v_n$ paarweise verschieden sind (sonst fassen wir die Terme entsprechend zusammen). Wenn wir $\lambda_0 = -1$ setzen, dann können wir das auch in symmetrischer Form schreiben als

$$\lambda_0 v_0 + \lambda_1 v_1 + \dots + \lambda_n v_n = \mathbf{0}.$$

Es gibt also eine *nichttriviale* Linearkombination (das ist eine, in der nicht alle Koeffizienten null sind; hier ist $\lambda_0 = -1 \neq 0$) von paarweise verschiedenen Elementen von E , die den Nullvektor ergibt.

DEF
nichttriviale
Linearkomb.

Umgekehrt gilt: Gibt es eine solche nichttriviale Linearkombination von Elementen von E , deren Wert der Nullvektor ist, etwa

$$\lambda_0 v_0 + \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mathbf{0}$$

mit $v_0, v_1, v_2, \dots, v_n \in E$ paarweise verschieden, dann ist $\lambda_j \neq 0$ für wenigstens ein $j \in \{0, 1, 2, \dots, n\}$. Wir können (falls nötig) die Nummerierung so ändern, dass $\lambda_0 \neq 0$ ist. Dann ist die Gleichung äquivalent zu

$$v_0 = (-\lambda_0^{-1} \lambda_1) v_1 + (-\lambda_0^{-1} \lambda_2) v_2 + \dots + (-\lambda_0^{-1} \lambda_n) v_n.$$

Wir können also ein Element v_0 von E als Linearkombination von Elementen von $E \setminus \{v_0\}$ schreiben. Daraus folgt, dass $E_0 = E \setminus \{v_0\}$ immer noch ein Erzeugendensystem von V ist. Das folgt aus Lemma 7.15: Wir müssen zeigen, dass jedes Element von E eine Linearkombination von Elementen von $E \setminus \{v_0\}$ ist. Für $v \in E$ mit $v \neq v_0$ können wir aber einfach $v = 1 \cdot v$ schreiben (denn $v \in E \setminus \{v_0\}$). Für v_0 haben wir die obige Darstellung.

E ist also genau dann ein minimales Erzeugendensystem, wenn der Nullvektor *nicht* als nichttriviale Linearkombination von (paarweise verschiedenen) Elementen von E geschrieben werden kann. Diese Eigenschaft ist sehr wichtig und hat einen eigenen Namen.

* **8.1. Definition.** Sei V ein K -Vektorraum.

DEF
Linear
unabhängig

- (1) Sei $n \in \mathbb{N}$. Das n -Tupel $(v_1, v_2, \dots, v_n)$ von Vektoren in V heißt *(K-)linear unabhängig*, wenn gilt:

$$\forall \lambda_1, \lambda_2, \dots, \lambda_n \in K: (\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mathbf{0} \Rightarrow \lambda_1 = \lambda_2 = \dots = \lambda_n = 0).$$

Anderenfalls heißt es *(K-)linear abhängig*. Wir sagen auch etwas informeller, die Vektoren $v_1, v_2, \dots, v_n$ seien linear (un)abhängig. Wir definieren „linear (un)abhängig“ analog für beliebige endliche Familien $(v_i)_{i \in I}$ (also mit I endlich) von Vektoren von V .

- (2) Sei I eine Menge. Eine Familie $(v_i)_{i \in I}$ von Elementen von V heißt *(K-)linear unabhängig*, wenn für jede endliche Teilmenge $J \subset I$ die endliche Familie $(v_i)_{i \in J}$ linear unabhängig ist. Mit der verallgemeinerten Summenschreibweise bedeutet das

$$\forall (\lambda_i)_{i \in I} \in K^{(I)}: \sum_{i \in I} \lambda_i \cdot v_i = \mathbf{0} \Rightarrow \forall i \in I: \lambda_i = 0.$$

Anderenfalls heißt $(v_i)_{i \in I}$ *(K-)linear abhängig*.

- (3) Eine Teilmenge $A \subset V$ heißt *(K-)linear unabhängig*, wenn die Familie $(v)_{v \in A}$ linear unabhängig ist, sonst heißt A *(K-)linear abhängig*. $\diamond$

Eine Familie oder Menge von Vektoren ist also genau dann linear abhängig, wenn man den Nullvektor als nichttriviale Linearkombination von Vektoren aus der Familie oder von verschiedenen Vektoren aus der Menge schreiben kann.

Der Unterschied zwischen Familien und Mengen ist, dass die Elemente in einer Familie gewissermaßen durch die Indexmenge nummeriert sind und sich wiederholen können, während die Elemente einer Menge keine weitere Ordnung haben und nicht mehrfach vorkommen. Zum Beispiel ist im K -Vektorraum K die Familie $(1)_{i \in \{1,2\}}$ linear abhängig, weil $1 \cdot 1 + (-1) \cdot 1 = 0$ eine nichttriviale Linearkombination ist, die den Nullvektor $0 \in K$ darstellt. (Allgemeiner ist jede Familie, in der ein Element mehrfach vorkommt, linear abhängig, wie man auf die gleiche Weise sehen kann.) Dagegen ist die Menge $\{1 \mid i \in \{1,2\}\} = \{1\}$ linear unabhängig, vergleiche Beispiel 8.3 unten.

Eine Menge A von Vektoren ist genau dann linear unabhängig, wenn jede *endliche* Teilmenge von A linear unabhängig ist.

Wie wir oben gesehen haben, ist ein Erzeugendensystem genau dann minimal, wenn es linear unabhängig ist. Aus unseren Überlegungen hat sich auch Folgendes ergeben:

Das Tupel $(v_1, v_2, \dots, v_n) \in V^n$ ist genau dann linear abhängig, wenn sich einer der Vektoren als Linearkombination der übrigen schreiben lässt.

Wichtig: Die Definition der Linearen Unabhängigkeit ist zentral für die Lineare Algebra. Es ist äußerst wichtig, dass Sie sie verstehen!



8.2. Beispiel. Wir betrachten den Grenzfall: Ist die leere Menge (oder die leere Familie) linear unabhängig oder linear abhängig?

Die einzige Linearkombination der leeren Menge ist die leere Summe mit dem Wert $\mathbf{0}$. Ist diese Linearkombination trivial oder nicht? Da „trivial“ bedeutet, dass alle Koeffizienten null sind, muss die leere Linearkombination trivial sein, denn da es keine Koeffizienten gibt, ist jede Allaussage über die Koeffizienten wahr. Die leere Menge ist also linear unabhängig.

Das passt auch mit der obigen Beobachtung zusammen, dass ein Erzeugendensystem genau dann minimal ist, wenn es linear unabhängig ist, denn die leere Menge kann man ja nicht verkleinern. $\clubsuit$

BSP
 $\emptyset$ ist linear
unabhängig

8.3. Beispiel. Wann ist die Einermenge $\{v\}$ linear unabhängig?

Die Linear„kombinationen“ haben die Form λv mit λ aus dem jeweiligen Körper. Aus $\lambda v = \mathbf{0}$ folgt $\lambda = 0$ oder $v = \mathbf{0}$ (vgl. Lemma 5.3). Das zeigt, dass v linear unabhängig ist, wenn v nicht der Nullvektor ist. Auf der anderen Seite ist $1 \cdot \mathbf{0} = \mathbf{0}$ eine nichttriviale Linearkombination, die den Nullvektor ergibt, also ist $\mathbf{0}$ linear abhängig. ♣

BSP
Wann ist
 $\{v\}$ linear
unabhängig?

8.4. Beispiel. Nach unseren Überlegungen vom Anfang dieses Abschnitts ist ein Paar (v_1, v_2) von Vektoren in V genau dann linear abhängig, wenn einer der beiden ein Vielfaches des anderen ist: $v_2 = \lambda v_1$ oder $v_1 = \lambda v_2$ für ein $\lambda \in K$. (Ist $v_1 = \mathbf{0}$, $v_2 \neq \mathbf{0}$, dann ist v_1 ein Vielfaches von v_2 , aber nicht umgekehrt.) ♣

BSP
Lineare Un-
abhängigkeit
von
zwei Vektoren

8.5. Beispiel. Hier ist ein sehr konkretes (und typisches) Beispiel. Sind die Vektoren $v_1 = (1, 1, 1, 1)$, $v_2 = (1, 2, 3, 4)$ und $v_3 = (1, 3, 5, 7)$ in $V = \mathbb{R}^4$ linear unabhängig oder nicht?

BSP
3 Vektoren
im $\mathbb{R}^4$

Wir müssen die Bedingung überprüfen. Seien also $\lambda_1, \lambda_2, \lambda_3 \in \mathbb{R}$ mit

$$\lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = \mathbf{0} = (0, 0, 0, 0).$$

Die Frage ist, ob daraus zwingend $\lambda_1 = \lambda_2 = \lambda_3 = 0$ folgt. Ausgeschrieben lautet die Gleichung

$$(\lambda_1 + \lambda_2 + \lambda_3, \lambda_1 + 2\lambda_2 + 3\lambda_3, \lambda_1 + 3\lambda_2 + 5\lambda_3, \lambda_1 + 4\lambda_2 + 7\lambda_3) = (0, 0, 0, 0);$$

das ist äquivalent zu den vier Gleichungen

$$\begin{aligned}\lambda_1 + \lambda_2 + \lambda_3 &= 0 \\ \lambda_1 + 2\lambda_2 + 3\lambda_3 &= 0 \\ \lambda_1 + 3\lambda_2 + 5\lambda_3 &= 0 \\ \lambda_1 + 4\lambda_2 + 7\lambda_3 &= 0\end{aligned}$$

Dieses Gleichungssystem hat $(\lambda_1, \lambda_2, \lambda_3) = (1, -2, 1)$ als eine nichttriviale Lösung. Das bedeutet, dass die Vektoren linear abhängig sind. ♣

8.6. Beispiel. Das Erzeugendensystem $(e_1, e_2, \dots, e_n)$ von K^n ist linear unabhängig, denn

$$\lambda_1 e_1 + \lambda_2 e_2 + \dots + \lambda_n e_n = (\lambda_1, \lambda_2, \dots, \lambda_n)$$

ist genau dann der Nullvektor, wenn alle Koeffizienten null sind. ♣

BSP
 $(e_1, \dots, e_n)$
ist linear
unabhängig

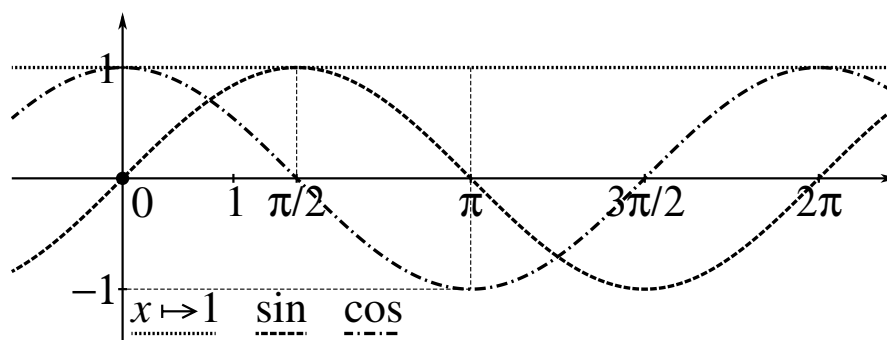


ABBILDUNG 1. Zu Beispiel 8.7

8.7. Beispiel. Die Funktionen $x \mapsto 1, \sin, \cos, \sin^2, \cos^2$ aus dem Raum $\mathcal{C}(\mathbb{R})$ der stetigen reellen Funktionen sind linear abhängig, denn es gilt

$$\forall x \in \mathbb{R}: \sin^2(x) + \cos^2(x) - 1 = 0,$$

also haben wir eine nichttriviale Linearkombination, die die Nullfunktion darstellt:

$$(-1) \cdot (x \mapsto 1) + 0 \cdot \sin + 0 \cdot \cos + 1 \cdot \sin^2 + 1 \cdot \cos^2 = \mathbf{0}.$$

Andererseits sind $x \mapsto 1, \sin$ und $\cos$ linear unabhängig:

Aus $\lambda_1 + \lambda_2 \sin(x) + \lambda_3 \cos(x) = 0$ für alle $x \in \mathbb{R}$ folgt durch Einsetzen von $x = 0, \pi, \pi/2$

$$\lambda_1 + \lambda_3 = \lambda_1 - \lambda_3 = \lambda_1 + \lambda_2 = 0$$

und damit $\lambda_1 = \lambda_2 = \lambda_3 = 0$.

BSP
Lineare Un-
abhängigkeit
in $\mathcal{C}(\mathbb{R})$



8.8. Beispiel. Die Familie $(f_n)_{n \in \mathbb{N}}$ der Potenzfunktionen $f_n: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^n$, ist linear unabhängig. (Das ist ein Beispiel für eine *unendliche* linear unabhängige Familie.) Das bedeutet

$$\forall n \in \mathbb{N} \forall a_0, a_1, \dots, a_n \in \mathbb{R}:$$

$$(\forall x \in \mathbb{R}: a_0 + a_1 x + \dots + a_n x^n = 0) \Rightarrow a_0 = a_1 = \dots = a_n = 0.$$

Für den Beweis stützen wir uns auf die (aus der Schule bekannte?) Tatsache, dass eine Polynomfunktion vom Grad n (also eine Funktion wie oben mit $a_n \neq 0$) höchstens n verschiedene Nullstellen (in $\mathbb{R}$) haben kann. Da es (für jedes $n \in \mathbb{N}$) mehr als n reelle Zahlen gibt, muss es dann, wenn wenigstens einer der Koeffizienten $\neq 0$ ist (dann hat die Polynomfunktion einen Grad $m \in \mathbb{N}$; m ist der größte Index, für den $a_m \neq 0$ gilt), ein $x \in \mathbb{R}$ geben, sodass

$$a_0 + a_1 x + \dots + a_n x^n \neq 0$$

ist. Das zeigt die Kontraposition der Implikation und somit auch die Implikation in der Aussage oben.

Man kann das auch durch vollständige Induktion beweisen. Für $n = 0$ reduziert sich die Behauptung auf die triviale Aussage $a_0 = 0 \Rightarrow a_0 = 0$. Für den Induktionsschritt nehmen wir an, dass die Aussage für n gilt. Um die Aussage für $n + 1$ zu beweisen, seien $a_0, a_1, \dots, a_{n+1} \in \mathbb{R}$ mit

$$\forall x \in \mathbb{R}: a_0 + a_1 x + a_2 x^2 + \dots + a_{n+1} x^{n+1} = 0.$$

Einsetzen von $x = 0$ liefert $a_0 = 0$, also haben wir

$$\forall x \in \mathbb{R}: x(a_1 + a_2 x + \dots + a_{n+1} x^n) = 0,$$

was bedeutet

$$\forall x \in \mathbb{R} \setminus \{0\}: a_1 + a_2 x + \dots + a_{n+1} x^n = 0.$$

Weil Polynomfunktionen stetig sind (das lernen Sie in der Analysis), gilt dies dann auch für $x = 0$, also

$$\forall x \in \mathbb{R}: a_1 + a_2 x + \dots + a_{n+1} x^n = 0.$$

Aus der Induktionsvoraussetzung folgt dann $a_1 = a_2 = \dots = a_{n+1} = 0$ wie gewünscht.



Wir schreiben noch eine einfache, aber nützliche Beobachtung auf, die unsere Überlegungen vom Beginn dieses Abschnitts formalisiert.

8.9. Lemma. Sei V ein K -Vektorraum.

LEMMA
Erzeugnis
einer linear
unabhängigen
Menge

- (1) Sei $(v_1, v_2, \dots, v_n)$ ein linear unabhängiges Tupel von Vektoren in V . Dann gilt für alle $v \in V$:

$$v \in \langle v_1, v_2, \dots, v_n \rangle \iff (v, v_1, v_2, \dots, v_n) \text{ linear abhängig.}$$

- (2) Sei $(v_i)_{i \in I}$ eine linear unabhängige Familie von Vektoren in V . Sei $v \in V$ beliebig. Sei weiter $i_0 \notin I$ und $I' = I \cup \{i_0\}$; wir setzen $v_{i_0} = v$. Dann gilt:

$$v \in \langle v_i \mid i \in I \rangle \iff (v_i)_{i \in I'} \text{ linear abhängig.}$$

- (3) Für linear unabhängige Teilmengen $A \subset V$ und Vektoren $v \in V$ gilt entsprechend:

$$v \in \langle A \rangle \iff v \in A \quad \text{oder} \quad A \cup \{v\} \text{ linear abhängig.}$$

Beweis.

- (1) „ $\Rightarrow$ “: $v \in \langle v_1, v_2, \dots, v_n \rangle$ bedeutet nach Satz 7.9, dass

$$v = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n$$

eine K -Linearkombination der Vektoren v_j ist. Dann ist

$$(-1)v + \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mathbf{0}$$

und diese Linearkombination ist nichttrivial (denn der Koeffizient -1 von v ist $\neq 0$). Also ist $(v, v_1, \dots, v_n)$ linear abhängig.

„ $\Leftarrow$ “: Da $(v, v_1, \dots, v_n)$ linear abhängig ist, gibt es eine nichttriviale K -Linearkombination

$$\lambda v + \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mathbf{0}.$$

Dabei kann λ nicht null sein, denn sonst hätten wir eine nichttriviale Linearkombination von $(v_1, v_2, \dots, v_n)$, die den Nullvektor darstellt, im Widerspruch zur linearen Unabhängigkeit dieser Vektoren. Dann können wir die Gleichung aber nach v auflösen:

$$v = (-\lambda^{-1}\lambda_1)v_1 + (-\lambda^{-1}\lambda_2)v_2 + \dots + (-\lambda^{-1}\lambda_n)v_n,$$

was (wiederum nach Satz 7.9) $v \in \langle v_1, v_2, \dots, v_n \rangle$ zeigt.

- (2) Das folgt aus dem ersten Teil, da in den jeweils zu betrachtenden Linearkombinationen nur endlich viele Vektoren v_i vorkommen.
- (3) Die Aussage für Teilmengen folgt aus der Aussage für Familien. □

9. BASIS UND DIMENSION

Linear unabhängige Erzeugendensysteme spielen eine fundamentale Rolle in der Linearen Algebra.

* **9.1. Definition.** Sei V ein K -Vektorraum. Eine Familie $(v_i)_{i \in I}$ von Elementen von V heißt $(K\text{-})Basis(familie)$ von V , wenn sie linear unabhängig und gleichzeitig ein Erzeugendensystem von V ist. Analog heißt eine Teilmenge $B \subset V$ eine $(K\text{-})Basis(menge)$ von V , wenn sie ein linear unabhängiges Erzeugendensystem von V ist. $\diamond$

DEF
Basis

Manchmal ist es praktischer, mit Familien (also „nummerierten Mengen“) zu arbeiten, und manchmal ist es praktischer, mit Mengen zu arbeiten, darum haben wir den Begriff der Basis in beiden Versionen definiert. Der Unterschied ist gering, denn in einer linear unabhängigen Familie kann kein Element mehrfach auftreten.

9.2. Beispiele.

BSP
Basen

- (1) Ist V ein Vektorraum und $A \subset V$ linear unabhängig, dann ist A eine Basis von $\langle A \rangle$ (denn A ist ein linear unabhängiges Erzeugendensystem von $\langle A \rangle$).
- (2) Die leere Menge ist Basis des Null-Vektorraums $\{0\}$.
- (3) Das Tupel $(e_1, e_2, \dots, e_n)$ ist eine K -Basis von K^n , die sogenannte *Standardbasis* von K^n .
- (4) Die Folge $(f_n)_{n \in \mathbb{N}}$ der Potenzfunktionen ist eine Basis des Vektorraums P der Polynomfunktionen. (Das ist ein Spezialfall von Teil (1)) $\clubsuit$

DEF
Standardbasis
von K^n

Wir hatten zu Beginn von §8 gesehen, dass ein Erzeugendensystem genau dann minimal ist, wenn es linear unabhängig (also eine Basis) ist. Wir formulieren das hier noch einmal und ergänzen es um eine ähnliche Aussage über linear unabhängige Familien.

* **9.3. Lemma.** Seien V ein Vektorraum und $B = (b_i)_{i \in I}$ eine Familie von Vektoren in V . Dann sind die folgenden Aussagen äquivalent:

LEMMA
Charakterisierung
von Basen

- (1) B ist eine Basis von V .
- (2) B ist ein minimales Erzeugendensystem von V .
- (3) B ist eine maximale linear unabhängige Familie in V .

„Maximal“ heißt dabei, dass für jedes $v \in V$ die um v erweiterte Familie nicht mehr linear unabhängig ist.

Beweis. Nach Definition 9.1 ist eine Basis ein linear unabhängiges Erzeugendensystem.

„(1) $\Leftrightarrow$ (2)“: Das hatten wir uns bereits zu Beginn von Abschnitt 8 überlegt.

„(1) $\Rightarrow$ (3)“: Es ist zu zeigen, dass jede echt größere Familie linear abhängig ist. Das folgt aber unmittelbar aus Lemma 8.9, (2), da für alle $v \in V$ ja $v \in \langle B \rangle$ gilt.

„(3) $\Rightarrow$ (1)“: B ist bereits linear unabhängig; es bleibt zu zeigen, dass B ein Erzeugendensystem ist. Dazu sei $v \in V$. Dann ist die um v erweiterte Familie B' linear abhängig; nach Lemma 8.9. (2) folgt $v \in \langle B \rangle$. Da $v \in V$ beliebig war, folgt $\langle B \rangle = V$, also ist B ein Erzeugendensystem von V . $\square$

Wir können die Eigenschaften, ein Erzeugendensystem, linear unabhängig oder eine Basis zu sein, auch durch die Anzahl der Linearkombinationen ausdrücken, die ein gegebenes Element von V darstellen. Wir formulieren das hier für endlich viele Vektoren.

9.4. Lemma. Sei V ein K -Vektorraum und seien $v_1, v_2, \dots, v_n \in V$. Wir definieren die zugehörige Linearkombinationenabbildung

$$\phi_{(v_1, v_2, \dots, v_n)}: K^n \longrightarrow V, \quad (\lambda_1, \lambda_2, \dots, \lambda_n) \longmapsto \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n.$$

Dann gilt:

- (1) $(v_1, v_2, \dots, v_n)$ ist genau dann ein **Erzeugendensystem** von V , wenn jeder Vektor $v \in V$ auf **mindestens** eine Weise als Linearkombination von $(v_1, v_2, \dots, v_n)$ geschrieben werden kann, also genau dann, wenn $\phi_{(v_1, v_2, \dots, v_n)}$ **surjektiv** ist.
- (2) $(v_1, v_2, \dots, v_n)$ ist genau dann **linear unabhängig**, wenn jeder Vektor $v \in V$ auf **höchstens** eine Weise als Linearkombination von $(v_1, v_2, \dots, v_n)$ geschrieben werden kann, also genau dann, wenn $\phi_{(v_1, v_2, \dots, v_n)}$ **injektiv** ist.
- (3) $(v_1, v_2, \dots, v_n)$ ist genau dann eine **Basis** von V , wenn jeder Vektor $v \in V$ auf **genau** eine Weise als Linearkombination von $(v_1, v_2, \dots, v_n)$ geschrieben werden kann, also genau dann, wenn $\phi_{(v_1, v_2, \dots, v_n)}$ **bijektiv** ist.

LEMMA
EZS/LU/Basis
über Anzahl
Lin.komb.
DEF
Linearkomb.-
abbildung

Beweis. Teil (1) folgt direkt aus Definition 7.8 und Satz 7.9.

Wir beweisen Teil (2). „ $\Rightarrow$ “: Wir nehmen an, dass $v_1, v_2, \dots, v_n$ linear unabhängig sind. Sei $v \in V$. Wenn wir zwei Linearkombinationen haben, also

$$v = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = \mu_1 v_1 + \mu_2 v_2 + \dots + \mu_n v_n$$

mit $\lambda_1, \lambda_2, \dots, \lambda_n, \mu_1, \mu_2, \dots, \mu_n \in K$, dann bilden wir die Differenz:

$$(\lambda_1 - \mu_1)v_1 + (\lambda_2 - \mu_2)v_2 + \dots + (\lambda_n - \mu_n)v_n = \mathbf{0}.$$

Weil $v_1, v_2, \dots, v_n$ linear unabhängig sind, muss das die triviale Linearkombination sein, also folgt $\lambda_1 = \mu_1, \lambda_2 = \mu_2, \dots, \lambda_n = \mu_n$.

„ $\Leftarrow$ “: Wir nehmen an, dass jedes $v \in V$ höchstens auf eine Weise als Linearkombination von $v_1, v_2, \dots, v_n$ darstellbar ist. Das gilt dann auch für $v = \mathbf{0}$. Da die triviale Linearkombination $\mathbf{0}$ darstellt, muss es die einzige sein. Damit ist gezeigt, dass $v_1, v_2, \dots, v_n$ linear unabhängig sind.

Teil (3) folgt dann aus (1) und (2). □

Wir können das vorstehende Lemma auch für beliebige Familien $(v_i)_{i \in I}$ formulieren, indem wir die in Definition 7.11 eingeführte Notation

$$K^{(I)} = \{(\lambda_i)_{i \in I} \in K^I \mid \{i \in I \mid \lambda_i \neq 0\} \text{ ist endlich}\}.$$

und die Summenschreibweise verwenden. Damit können wir die Linearkombinationenabbildung analog definieren als

$$\phi_{(v_i)_{i \in I}}: K^{(I)} \longrightarrow V, \quad (\lambda_i)_{i \in I} \longmapsto \sum_{i \in I} \lambda_i v_i.$$

Dann gilt wieder:

- (1) $(v_i)_{i \in I}$ Erzeugendensystem $\iff \phi_{(v_i)_{i \in I}}$ surjektiv.
- (2) $(v_i)_{i \in I}$ linear unabhängig $\iff \phi_{(v_i)_{i \in I}}$ injektiv.
- (3) $(v_i)_{i \in I}$ Basis $\iff \phi_{(v_i)_{i \in I}}$ bijektiv.

$K^{(I)}$ ist genau der K -Untervektorraum von K^I , der durch die Familien $\mathbf{e}_i = (\delta_{ij})_{j \in I}$ für $i \in I$ erzeugt wird. Dabei ist $\delta_{ij} = 1$ für $i = j$ und $\delta_{ij} = 0$ für $i \neq j$ (das sogenannte *Kronecker-Delta*); die Familie $\mathbf{e}_i$ hat also als i -te Komponente eine Eins, alle anderen Komponenten sind null. Das verallgemeinert die Standardbasis von K^n auf den Vektorraum $K^{(I)}$.

Eine Basis $(v_1, v_2, \dots, v_n)$ von V verhilft uns also zu einer bijektiven Abbildung $K^n \rightarrow V$. Damit können wir die Elemente von V durch ihr Koeffiziententupel $(\lambda_1, \lambda_2, \dots, \lambda_n) \in K^n$ beschreiben (und Addition und Skalarmultiplikation von V verhalten sich genauso wie die von K^n). Das ist natürlich eine schöne Sache. Es stellt sich dann die Frage, ob jeder Vektorraum eine Basis hat. Wir werden das hier für endlich erzeugte Vektorräume positiv beantworten. (Ein Vektorraum ist *endlich erzeugt*, wenn er ein endliches Erzeugendensystem hat.) Dafür beweisen wir sogar eine stärkere Aussage, die viele nützliche Anwendungen haben wird.

DEF
endl. erzeugter
Vektorraum

* **9.5. Satz.** *Sei V ein Vektorraum und seien $v_1, v_2, \dots, v_n$ und $w_1, w_2, \dots, w_m$ Elemente von V , sodass*

SATZ
Basis-
ergänzungs-
satz

- (1) $(v_1, v_2, \dots, v_n)$ linear unabhängig und
- (2) $(v_1, v_2, \dots, v_n, w_1, w_2, \dots, w_m)$ ein Erzeugendensystem von V ist.

Dann kann man $(v_1, v_2, \dots, v_n)$ durch Hinzunahme geeigneter Vektoren w_j zu einer Basis von V ergänzen.

Genauer bedeutet das: Es gibt $k \in \mathbb{N}$ und Indizes $j_1, j_2, \dots, j_k \in \{1, 2, \dots, m\}$, sodass

$$(v_1, v_2, \dots, v_n, w_{j_1}, w_{j_2}, \dots, w_{j_k})$$

eine Basis von V ist.

Die natürlichen Zahlen n und m dürfen und k kann auch null sein. Wenn $m = 0$ ist, dann ist $(v_1, v_2, \dots, v_n)$ schon eine Basis, und es ist nichts zu tun (dann ist auch $k = 0$). Das werden wir im Beweis als Induktionsanfang benutzen.

Wenn $n = 0$ ist, dann sagt der Satz, dass jedes endliche Erzeugendensystem eine Basis enthält. Das ist plausibel, denn man kann ja immer Elemente entfernen, solange das Erzeugendensystem nicht minimal ist. Irgendwann (nach spätestens m -maligem Entfernen eines Elements) muss man bei einem minimalen Erzeugendensystem ankommen; das ist dann eine Basis.

Wenn sich $k = 0$ ergibt, dann bedeutet das, dass $(v_1, v_2, \dots, v_n)$ bereits eine Basis ist.

Im folgenden Beweis wird vollständige Induktion verwendet. Die vollständige Induktion wurde in der Analysis behandelt. Für diejenigen, die die Analysis I erst nach der Linearen Algebra hören, ist hier ein kleiner Crashkurs.

Das Beweisprinzip der vollständigen Induktion beruht auf der rekursiven Definition der natürlichen Zahlen: 0 ist eine natürliche Zahl (oder 1, je nach Vorliebe...), mit n ist auch $n + 1$ eine natürliche Zahl, und alle natürlichen Zahlen entstehen sukzessive auf diese Weise:

$$0 \rightarrow 1 \rightarrow 2 \rightarrow 3 \rightarrow \dots \rightarrow n \rightarrow n + 1 \rightarrow \dots$$

Sei nun $A(n)$ eine Aussage über die natürliche Zahl n . Wenn wir zeigen können, dass $A(0)$ gilt und dass außerdem die Implikationen

$$A(0) \Rightarrow A(1), \quad A(1) \Rightarrow A(2), \quad \dots, \quad A(n) \Rightarrow A(n+1), \quad \dots$$

gelten, dann folgt $A(n)$ für alle $n \in \mathbb{N}$: $A(0)$ haben wir gezeigt, aus $A(0)$ und $A(0) \Rightarrow A(1)$ folgt $A(1)$, aus $A(1)$ und $A(1) \Rightarrow A(2)$ folgt $A(2)$, und so weiter. Daraus ergibt sich das folgende **Induktionsprinzip**:

Aus $A(0)$ und $\forall n \in \mathbb{N}: A(n) \Rightarrow A(n+1)$ folgt $\forall n \in \mathbb{N}: A(n)$.

Dabei heißt der Beweis von $A(0)$ der *Induktionsanfang* und der Beweis der Implikation $A(n) \Rightarrow A(n+1)$ der *Induktionsschritt*. Im Induktionsschritt heißt $A(n)$ meistens die *Induktionsannahme* oder *Induktionsvoraussetzung*.

Das folgende Standardbeispiel sollte das etwas klarer machen. Wir wollen zeigen, dass

$$\forall n \in \mathbb{N}: \sum_{j=0}^n j = \frac{n(n+1)}{2}$$

gilt (links in der Gleichung steht die Summe $0 + 1 + 2 + \dots + (n-1) + n$).

Induktionsanfang: Wir setzen $n = 0$; die Aussage ist dann

$$\sum_{j=0}^0 j = \frac{0(0+1)}{2};$$

beide Seiten sind null, also gilt die Aussage.

Induktionsschritt: Jetzt ist $n \in \mathbb{N}$ beliebig. Wir nehmen an (Induktionsvoraussetzung), dass die Aussage für n stimmt. Unter dieser Annahme müssen wir zeigen, dass sie auch für $n+1$ richtig ist. Das geht in diesem Fall zum Beispiel so:

$$\sum_{j=0}^{n+1} j = \sum_{j=0}^n j + (n+1) \stackrel{\text{IV}}{=} \frac{n(n+1)}{2} + n+1 = \frac{n(n+1) + 2(n+1)}{2} = \frac{(n+1)(n+2)}{2}$$

An der mit „IV“ markierten Stelle haben wir die Induktionsvoraussetzung benutzt.

Varianten des Induktionsprinzips ergeben sich, wenn man nicht bei 0 anfängt, sondern mit einer anderen ganzen Zahl n_0 ; dann folgt die Aussage für alle $n \geq n_0$. Manchmal braucht man die Voraussetzung für zwei Werte von n (dann zeigt man $A(0)$ und $A(1)$ als Induktionsanfang und $(A(n) \wedge A(n+1)) \Rightarrow A(n+2)$ als Induktionsschritt), oder man stützt sich gleich auf *alle* kleineren Fälle:

Aus $\forall n \in \mathbb{N}: (\forall k \in \mathbb{N}, k < n: A(k)) \Rightarrow A(n)$ folgt $\forall n \in \mathbb{N}: A(n)$.

Die Induktionsvoraussetzung ist dann also, dass $A(k)$ für *alle* $k < n$ gilt; daraus muss $A(n)$ hergeleitet werden. Der Induktionsanfang ist dabei implizit eingeschlossen, denn für $n = 0$ ist die Voraussetzung trivialerweise erfüllt (die Menge der kleineren k ist leer), also zeigt man hier $A(0)$ ohne Zusatzannahmen.

Alternativ dazu kann man das **Prinzip des kleinsten Verbrechers** (*minimal criminal*) verwenden: Ist die Aussage $\forall n \in \mathbb{N}: A(n)$ falsch, dann gibt es ein *kleinstes* n , für das $A(n)$ nicht gilt. Damit muss aber $A(k)$ für alle $k < n$ gelten, woraus man dann einen Widerspruch ableitet. Die Annahme, dass es ein Gegenbeispiel gibt, ist also nicht haltbar; damit muss $A(n)$ für alle $n \in \mathbb{N}$ gelten.

Beweis. Der Beweis benutzt vollständige Induktion nach m . Er basiert auf der anschaulichen Idee, dass man ausgehend von $(v_1, \dots, v_n, w_1, \dots, w_m)$ sukzessive Elemente w_j entfernt, bis man linear unabhängige Vektoren hat. Dazu fixieren wir das n -Tupel $(v_1, \dots, v_n)$ und nehmen an (Voraussetzung (1)), dass es linear unabhängig ist. Die Aussage $A(m)$, die wir durch Induktion beweisen wollen, ist dann, dass die Behauptung für das fest gewählte Tupel $(v_1, \dots, v_n)$ und jedes m -Tupel $(w_1, \dots, w_m)$ von Vektoren von V gilt.

Im Induktionsanfang ist $m = 0$. Dann ist nach Voraussetzung (2) $(v_1, \dots, v_n)$ ein Erzeugendensystem von V und nach Voraussetzung (1) linear unabhängig, also bereits eine Basis: Die Behauptung gilt mit $k = 0$.

Für den Induktionsschritt nehmen wir an, die Behauptung gelte für m -Tupel $(w_1, \dots, w_m)$. Um sie für $m + 1$ zu zeigen, betrachten wir ein $(m + 1)$ -Tupel $(w_1, \dots, w_{m+1})$ von Vektoren von V , sodass $(v_1, \dots, v_n, w_1, \dots, w_{m+1})$ ein Erzeugendensystem von V ist. Es gibt nun zwei Möglichkeiten:

- $(v_1, \dots, v_n, w_1, \dots, w_{m+1})$ ist linear unabhängig. Dann haben wir eine Basis; die Behauptung gilt mit $k = m + 1$ und $j_1 = 1, \dots, j_{m+1} = m + 1$.
- Anderenfalls ist $(v_1, \dots, v_n, w_1, \dots, w_{m+1})$ linear abhängig. Dann gibt es eine nichttriviale Linearkombination

$$\lambda_1 v_1 + \dots + \lambda_n v_n + \mu_1 w_1 + \dots + \mu_m w_m + \mu_{m+1} w_{m+1} = \mathbf{0}.$$

Hier können nicht alle $\mu_j = 0$ sein, denn dann hätten wir eine nichttriviale Linearkombination nur der v_i , die den Nullvektor darstellt, im Widerspruch zur linearen Unabhängigkeit von $(v_1, \dots, v_n)$. Also gibt es ein $j_0 \in \{1, 2, \dots, m + 1\}$ mit $\mu_{j_0} \neq 0$. Falls nötig, vertauschen wir w_{j_0} und w_{m+1} ; dann können wir annehmen, dass $j_0 = m + 1$ ist. Wir können obige Gleichung dann nach w_{m+1} auflösen. Das zeigt, dass

$$w_{m+1} \in \langle v_1, \dots, v_n, w_1, \dots, w_m \rangle$$

ist, also ist

$$\langle v_1, \dots, v_n, w_1, \dots, w_m \rangle = \langle v_1, \dots, v_n, w_1, \dots, w_m, w_{m+1} \rangle = V.$$

Voraussetzung (2) ist also für $(w_1, \dots, w_m)$ erfüllt, sodass wir die Induktionsannahme verwenden können, was uns die gewünschte Aussage liefert.

Wenn man ganz genau sein will, dann setzt man $w'_i = w_i$ für $i \in \{1, 2, \dots, m\} \setminus \{j_0\}$ und $w'_{j_0} = w_{m+1}$ (falls $j_0 \neq m + 1$), und wendet die Induktionsannahme auf $(w'_1, \dots, w'_m)$ an. Dann bekommt man $k \leq m$ und $j'_1, \dots, j'_k$, sodass $(v_1, \dots, v_n, w'_{j'_1}, \dots, w'_{j'_k})$ eine Basis ist. Mit $j_i = j'_i$ für $j'_i \neq j_0$ und $j_i = m + 1$, falls $j'_i = j_0 \neq m + 1$, ist dann $(v_1, \dots, v_n, w_{j_1}, \dots, w_{j_k})$ eine Basis. $\square$

Man kann alternativ den Beweis auch so formulieren, dass man nacheinander Vektoren w_j zu den $v_1, v_2, \dots, v_n$ hinzunimmt, solange das entstehende Tupel linear unabhängig ist. Ist das nicht mehr möglich, dann muss man eine Basis haben.

Man beachte, dass in diesem Beweis die Zahl n nicht fixiert ist. (Er ist deswegen etwas weniger leicht zu verstehen, weswegen ich zunächst den anderen Beweis formuliert habe.)

Der Induktionsanfang, also der Fall $m = 0$, ist klar, denn dann ist $(v_1, v_2, \dots, v_n)$ bereits ein linear unabhängiges Erzeugendensystem, also eine Basis. Die Behauptung gilt also mit $k = 0$.

Für den Induktionsschritt nehmen wir an, dass die Aussage für ein gegebenes m stimmt, und beweisen sie für $m + 1$. Sei also $(v_1, v_2, \dots, v_n)$ linear unabhängig und seien außerdem $w_1, w_2, \dots, w_m, w_{m+1} \in V$, sodass $(v_1, v_2, \dots, v_n, w_1, w_2, \dots, w_m, w_{m+1})$ ein Erzeugendensystem von V ist. Wir unterscheiden zwei Fälle:

- (1) $w_{m+1} \in \langle v_1, v_2, \dots, v_n \rangle$.
Dann ist $(v_1, v_2, \dots, v_n, w_1, w_2, \dots, w_m)$ auch schon ein Erzeugendensystem; die Behauptung folgt direkt aus der Induktionsannahme.
- (2) $w_{m+1} \notin \langle v_1, v_2, \dots, v_n \rangle$.
Wir schreiben v_{n+1} für w_{m+1} . Dann ist $(v_1, v_2, \dots, v_n, v_{n+1})$ linear unabhängig (wir benutzen hier wieder Lemma 8.9) und

$$(v_1, v_2, \dots, v_n, v_{n+1}, w_1, w_2, \dots, w_m)$$

ist ein Erzeugendensystem (dasselbe wie vorher). Nach der Induktionsannahme gibt es $j'_1, \dots, j'_{k'} \in \{1, 2, \dots, m\}$, sodass

$$(v_1, v_2, \dots, v_n, v_{n+1}, w_{j'_1}, w_{j'_2}, \dots, w_{j'_{k'}}) = (v_1, \dots, v_n, w_{m+1}, w_{j'_1}, \dots, w_{j'_{k'}})$$

eine Basis von V ist. Wir setzen

$$k = k' + 1, \quad j_1 = m + 1 \quad \text{und} \quad j_2 = j'_1, j_3 = j'_2, \dots, j_k = j'_{k'}$$

und erhalten die Behauptung.

9.6. Folgerung. *Jeder Vektorraum, der ein endliches Erzeugendensystem besitzt, hat eine Basis.*

FOLG
Existenz
einer Basis

Beweis. Das folgt aus Satz 9.5, wenn man $n = 0$ nimmt. Genauer erhalten wir die Aussage, dass man eine Basis finden kann, die aus Elementen eines gegebenen endlichen Erzeugendensystems besteht. (Der Beweis reduziert sich in diesem Fall darauf, dass man aus dem gegebenen Erzeugendensystem solange Elemente entfernen kann, bis es minimal, also linear unabhängig, ist.) $\square$

Was passiert, wenn es kein endliches Erzeugendensystem gibt? Dann gibt es auch noch einen Basisergänzungssatz, den wir hier für Mengen formulieren:

Satz. *Sei V ein Vektorraum und seien A und E Teilmengen von V , sodass A linear unabhängig und $A \cup E$ ein Erzeugendensystem von V ist. Dann gibt es eine Teilmenge $B \subset E$, sodass $A \cup B$ eine Basismenge von V ist.*

SATZ
Basis-
ergänzungs-
satz

Den Beweis kann man jetzt natürlich nicht mehr durch vollständige Induktion führen. Man braucht ein anderes Werkzeug dafür, zum Beispiel das sogenannte *Zornsche Lemma*. Es besagt Folgendes.

Satz. *Seien X eine Menge und $\mathcal{M} \subset \mathcal{P}(X)$ eine Menge von Teilmengen von X . Eine **Kette** in $\mathcal{M}$ ist eine Teilmenge $\mathcal{K} \subset \mathcal{M}$, sodass je zwei Elemente von $\mathcal{K}$ vergleichbar sind, das heißt*

$$\forall T_1, T_2 \in \mathcal{K}: T_1 \subset T_2 \quad \text{oder} \quad T_2 \subset T_1.$$

SATZ
Zornsches
Lemma

*Wenn jede solche Kette $\mathcal{K}$ eine **obere Schranke** in $\mathcal{M}$ hat, wenn es also zu $\mathcal{K}$ ein Element $S \in \mathcal{M}$ gibt, so dass gilt*

$$\forall T \in \mathcal{K}: T \subset S,$$

*dann hat $\mathcal{M}$ **maximale** Elemente. Es gibt dann also (mindestens) ein $M \in \mathcal{M}$, sodass gilt*

$$\forall T \in \mathcal{M}: M \subset T \Rightarrow M = T$$

(d.h., es gibt keine echt größere Menge in $\mathcal{M}$).

Man kann zeigen, dass das Zornsche Lemma (wenn man die „harmlosen“ Axiome der Mengenlehre als gegeben annimmt) zum Auswahlaxiom (siehe die Diskussion im Kleingedruckten auf Seite 25) äquivalent ist.

Der Beweis des Basisergänzungssatzes geht dann so: E ist die Menge X im Zornschen Lemma und $\mathcal{M} = \{B \subset E \mid A \cup B \text{ linear unabhängig}\}$. Wir müssen die Voraussetzung des Zornschen Lemmas nachprüfen. Sei dazu $\mathcal{K} \subset \mathcal{M}$ eine Kette. Wir setzen $S = \bigcup \mathcal{K}$ (das ist also die Vereinigung all der Teilmengen von E , die Elemente der Kette $\mathcal{K}$ sind). Es ist dann klar, dass $T \subset S$ für alle $T \in \mathcal{K}$ gilt. Wir müssen noch zeigen, dass $S \in \mathcal{M}$ ist, dass also $A \cup S$ linear unabhängig ist. Angenommen, das wäre falsch, dann gäbe es eine nichttriviale Linearkombination von Elementen von $A \cup S$, die den Nullvektor darstellt. In dieser Linearkombination kommen nur endlich viele Elemente $v_1, v_2, \dots, v_n$ von S vor. Da $S = \bigcup \mathcal{K}$, gibt es für jedes v_j ein $K_j \in \mathcal{K}$ mit $v_j \in K_j$. Nach eventueller Umnummerierung können wir annehmen, dass $K_1 \subset K_2 \subset \dots \subset K_n$ ist (hier wird verwendet, dass $\mathcal{K}$

eine Kette ist). Dann sind aber $v_1, v_2, \dots, v_n \in K_n$, und es würde folgen, dass $A \cup K_n$ linear abhängig ist. Weil $K_n \in \mathcal{M}$ ist, ist das ein Widerspruch, also muss $A \cup S$ linear unabhängig sein. (Für dieses Argument ist die Endlichkeit von Linearkombinationen entscheidend!) Damit ist S eine obere Schranke von $\mathcal{K}$ in $\mathcal{M}$ und die Voraussetzung im Zornschen Lemma ist erfüllt. Es folgt, dass $\mathcal{M}$ ein maximales Element B hat. Da $B \in \mathcal{M}$ ist, ist $A \cup B$ linear unabhängig. Wäre $A \cup B$ kein Erzeugendensystem, dann gäbe es $v \in E$ mit $v \notin \langle A \cup B \rangle$. Dann wäre aber $A \cup (B \cup \{v\})$ ebenfalls linear unabhängig. Das würde $B \cup \{v\} \in \mathcal{M}$ bedeuten, aber das kann nicht sein, da B maximal ist (v kann kein Element von B sein, sonst wäre $v \in \langle A \cup B \rangle$). Also ist $A \cup B$ auch ein Erzeugendensystem und somit eine Basis.

Wir erhalten daraus sofort (mit $A = \emptyset$ und $E = V$):

Folgerung. *Jeder Vektorraum hat eine Basis.*

FOLG
Existenz
von Basen

Aus dem Auswahlaxiom folgt also zum Beispiel, dass $\mathbb{R}$ als $\mathbb{Q}$ -Vektorraum (als den man $\mathbb{R}$ mit seiner Addition und der auf $\mathbb{Q} \times \mathbb{R}$ eingeschränkten Multiplikation betrachten kann) eine Basis hat. Gesehen hat so eine Basis aber noch niemand. Wie schon früher erwähnt ist das Auswahlaxiom (und damit auch das Zornsche Lemma) inhärent inkonstruktiv, sodass unser Beweis oben (im Gegensatz zum endlichen Fall) keinerlei Hinweis darauf gibt, wie die gesuchte Teilmenge B zu finden wäre.

Eine weitere wichtige Folgerung besagt, dass man (in einem endlich erzeugten Vektorraum) beliebige linear unabhängige Vektoren stets zu einer Basis ergänzen kann.

9.7. Folgerung. *Sei V ein Vektorraum mit endlichem Erzeugendensystem und sei $(v_1, v_2, \dots, v_n) \in V^n$ linear unabhängig. Dann gibt es $k \in \mathbb{N}$ und Vektoren $v_{n+1}, v_{n+2}, \dots, v_{n+k} \in V$, sodass $(v_1, v_2, \dots, v_{n+k})$ eine Basis von V ist.*

FOLG
Erweiterung
zu Basis

Beweis. Sei $(w_1, w_2, \dots, w_m)$ ein endliches Erzeugendensystem von V . Dann sind für $v_1, \dots, v_n$ und $w_1, \dots, w_m$ die Voraussetzungen von Satz 9.5 erfüllt. Die Aussage des Satzes liefert dann die Behauptung, wenn man $v_{n+1} = w_{j_1}, \dots, v_{n+k} = w_{j_k}$ setzt. $\square$

9.8. Beispiel. Wir finden eine Basis des Untervektorraums

BSP
Basis

$$U = \{(x, y, z) \in \mathbb{R}^3 \mid z = x + y\} \subset \mathbb{R}^3.$$

Dazu finden wir möglichst viele linear unabhängige Vektoren und prüfen dann, ob wir ein Erzeugendensystem haben. Zum Beispiel sind $(1, 0, 1)$ und $(0, 1, 1)$ linear unabhängige Elemente von U , denn

$$\lambda(1, 0, 1) + \mu(0, 1, 1) = \mathbf{0} \iff (\lambda, \mu, \lambda + \mu) = (0, 0, 0) \iff \lambda = \mu = 0.$$

Diese beiden Vektoren bilden auch ein Erzeugendensystem, denn für $(x, y, z) \in U$ gilt $z = x + y$, also

$$(x, y, z) = (x, y, x + y) = x(1, 0, 1) + y(0, 1, 1) \in \langle (1, 0, 1), (0, 1, 1) \rangle.$$

Damit ist $((1, 0, 1), (0, 1, 1))$ eine Basis von U . $\clubsuit$

Eine weitere wichtige Konsequenz des Basisergänzungssatzes ist der *Basisaustauschsatz*.

* **9.9. Satz.** Sei V ein Vektorraum und seien $(v_1, v_2, \dots, v_n)$ und $(w_1, w_2, \dots, w_m)$ zwei Basen von V . Für jedes $i \in \{1, 2, \dots, n\}$ gibt es ein $j \in \{1, 2, \dots, m\}$, sodass $(v_1, \dots, v_{i-1}, w_j, v_{i+1}, \dots, v_n)$ ebenfalls eine Basis von V ist.

SATZ
Basis-
austausch-
satz

Man tauscht also das Basiselement v_i der ersten Basis durch ein Element der zweiten Basis aus.

Beweis. Wir können ohne Einschränkung $i = n$ annehmen (sonst ändere man die Nummerierung entsprechend). Wir wenden den Basisergänzungssatz 9.5 an mit $v_1, v_2, \dots, v_{n-1}$ und $w_1, w_2, \dots, w_m$. Die Voraussetzungen sind erfüllt, da Teilfamilien von linear unabhängigen Familien immer linear unabhängig sind und die w_j schon alleine ein Erzeugendensystem bilden. Es gibt also $k \in \mathbb{N}$ und Indizes $j_1, \dots, j_k \in \{1, 2, \dots, m\}$, sodass $(v_1, \dots, v_{n-1}, w_{j_1}, w_{j_2}, \dots, w_{j_k})$ eine Basis von V ist. Die Behauptung bedeutet gerade, dass man $k = 1$ wählen kann; wir setzen dann $j = j_1$. Es ist klar, dass $k > 0$ sein muss, denn $(v_1, v_2, \dots, v_{n-1})$ ist kein Erzeugendensystem mehr ($(v_1, v_2, \dots, v_{n-1}, v_n)$ ist ein minimales Erzeugendensystem, aus dem wir ein Element entfernt haben). Wir zeigen, dass $(v_1, v_2, \dots, v_{n-1}, w_{j_1})$ ein Erzeugendensystem ist; daraus folgt die Behauptung, denn dieses Tupel ist als Teilfamilie der Basis $(v_1, \dots, v_{n-1}, w_{j_1}, w_{j_2}, \dots, w_{j_k})$ linear unabhängig.

Wir haben $w_{j_1} \in V = \langle v_1, v_2, \dots, v_n \rangle$. Nach Lemma 8.9 bedeutet das, dass $(v_1, v_2, \dots, v_{n-1}, v_n, w_{j_1})$ linear abhängig ist. Da $(v_1, v_2, \dots, v_{n-1}, w_{j_1})$ als Teil der Basis $(v_1, \dots, v_{n-1}, w_{j_1}, w_{j_2}, \dots, w_{j_k})$ linear unabhängig ist, folgt dann wieder mit Lemma 8.9, dass $v_n \in \langle v_1, v_2, \dots, v_{n-1}, w_{j_1} \rangle$ ist. Da natürlich auch $v_1, v_2, \dots, v_{n-1}$ in diesem Untervektorraum enthalten sind, enthält er ein Erzeugendensystem von V ; es folgt (mit Lemma 7.15(2)) $\langle v_1, v_2, \dots, v_{n-1}, w_{j_1} \rangle = V$ wie behauptet. $\square$

9.10. Folgerung. Sei V ein Vektorraum und seien $(v_1, \dots, v_n)$ und $(w_1, \dots, w_m)$ zwei (endliche) Basen von V . Dann ist $n = m$.

FOLG
Größe
von Basen

Je zwei Basen haben also gleich viele Elemente.

Beweis. Wir nehmen $n > m$ an und leiten einen Widerspruch her (der Fall $n < m$ geht genauso). Durch n -malige Anwendung von Satz 9.9 (mit $i = 1, 2, \dots, n$) erhalten wir Indizes $j_1, j_2, \dots, j_n \in \{1, 2, \dots, m\}$, sodass $(w_{j_1}, w_{j_2}, \dots, w_{j_n})$ eine Basis von V ist. Da m kleiner als n ist, müssen sich in diesem Tupel Vektoren wiederholen. Dann sind $w_{j_1}, w_{j_2}, \dots, w_{j_n}$ aber nicht linear unabhängig. Dies ist der gewünschte Widerspruch. $\square$

Wir führen eine Schreibweise für die Anzahl der Elemente einer Menge ein.

9.11. Definition. Sei M eine Menge. Wir schreiben $\#M$ für die Anzahl der Elemente von M . Wenn M unendlich ist, dann setzen wir $\#M = \infty$. $\diamond$

DEF
 $\#M$

Eine andere häufig anzutreffende Schreibweise ist $|M|$. Ich bevorzuge $\#M$, weil es dabei keine Verwechslungsgefahr gibt.



Wir setzen hier ein intuitives Verständnis davon voraus, was die „Anzahl der Elemente“ einer (endlichen) Menge ist. Wenn man das formal sauber definieren will, ist es aber gar nicht so einfach. Man kann zum Beispiel für eine Menge M und $n \in \mathbb{N}$ definieren:

$$\#M = n \iff \exists f: M \rightarrow \mathbb{N}_{<n} \text{ mit } f \text{ bijektiv.}$$

(Beachte: $\mathbb{N}_{<n} = \{m \in \mathbb{N} \mid m < n\} = \{0, 1, 2, \dots, n-1\}$ hat gerade n Elemente.) Das formalisiert die Vorstellung, dass man die Elemente von 0 bis $n-1$ (oder analog von 1

bis n) durchnummerieren kann. Man muss dann noch zeigen, dass eine solche Bijektion nicht mit zwei Mengen $\mathbb{N}_{<n}$ zu verschiedenen n möglich ist und dass eine Menge genau dann unendlich ist, wenn es für kein $n \in \mathbb{N}$ eine solche Bijektion gibt.

Das Erste lässt sich durch vollständige Induktion erledigen: Angenommen, es gäbe eine Bijektion $f: \mathbb{N}_{<n+m} \rightarrow \mathbb{N}_{<n}$ mit $m > 0$. Für $n = 0$ ist das unmöglich (es gibt keine Abbildung von der nichtleeren Menge $\mathbb{N}_{<m}$ in die leere Menge $\mathbb{N}_{<0}$). Für $n > 0$ konstruiert man aus f eine neue Bijektion $f': \mathbb{N}_{<n-1+m} \rightarrow \mathbb{N}_{<n-1}$, die es nach Induktionsannahme aber nicht geben kann. Dabei ist $f'(k) = f(k)$, außer $f(k) = n - 1$; dann setzt man $f'(k) = f(n - 1 + m)$.

Für das Zweite braucht man eine Definition, was eine „unendliche Menge“ ist. Eine Möglichkeit ist

$$M \text{ ist unendlich} \iff \exists f: \mathbb{N} \rightarrow M \quad \text{mit } f \text{ injektiv.}$$

Eine alternative Definition ist

$$M \text{ ist unendlich} \iff \exists f: M \rightarrow M \quad \text{mit } f \text{ injektiv, aber nicht surjektiv.}$$

Es ist eine interessante Übung, die Äquivalenz dieser beiden Definitionen zu zeigen. Man zeigt wieder durch vollständige Induktion mit der zweiten Definition, dass die Mengen $\mathbb{N}_{<n}$ nicht unendlich sind (eine injektive Abbildung $\mathbb{N}_{<n} \rightarrow \mathbb{N}_{<n}$ ist immer auch bijektiv). Es ist dann noch zu zeigen, dass eine nicht unendliche Menge M bijektiv zu einer der Mengen $\mathbb{N}_{<n}$ ist. Dazu definiert man eine injektive Abbildung f von einem Anfangsstück von $\mathbb{N}$ nach M : Falls M leer ist, dann ist die Abbildung leer; das Anfangsstück ist $\mathbb{N}_{<0}$ und $\#M = 0$. Sonst gibt es ein $m_0 \in M$; wir setzen $f(0) = m_0$ und $M_1 = M \setminus \{m_0\}$. Wenn f schon auf $\mathbb{N}_{<n}$ definiert und parallel M_n konstruiert ist, dann ist entweder M_n leer; in diesem Fall haben wir eine Bijektion $f: \mathbb{N}_{<n} \rightarrow M$ und $\#M = n$; oder es gibt ein $m_n \in M_n$, dann setzen wir $f(n) = m_n$ und $M_{n+1} = M_n \setminus \{m_n\}$. Damit ist f auf $\mathbb{N}_{<n+1}$ definiert. Wenn diese Konstruktion nicht abbricht, dann erhalten wir eine injektive Abbildung $f: \mathbb{N} \rightarrow M$, also muss M (nach der ersten Definition) unendlich sein.

Allgemeiner kann man für beliebige (auch unendliche) Mengen M_1 und M_2 definieren

$$\#M_1 \leq \#M_2 \iff \exists f: M_1 \rightarrow M_2 \quad \text{mit } f \text{ injektiv.}$$

Nach dem Satz von Cantor-Bernstein-Schröder folgt aus $\#M_1 \leq \#M_2$ und $\#M_2 \leq \#M_1$, dass es eine Bijektion zwischen M_1 und M_2 gibt; deshalb ist es sinnvoll zu definieren

$$\#M_1 = \#M_2 \iff \exists f: M_1 \rightarrow M_2 \quad \text{mit } f \text{ bijektiv}$$

(man sagt dann auch, M_1 und M_2 seien *gleichmächtig*). Damit kann man dann auch bei unendlichen Mengen zwischen verschiedenen Mächtigkeiten oder *Kardinalitäten* unterscheiden. Zum Beispiel gilt für jede unendliche Menge M , dass $\#\mathbb{N} \leq \#M$ ist (das war eine der beiden Definitionen oben); es gibt aber auch „überabzählbare“ Mengen M mit $\#\mathbb{N} < \#M$ (also $\#\mathbb{N} \leq \#M$ und $\#\mathbb{N} \neq \#M$). In der Analysis lernen Sie, dass $M = \mathbb{R}$ dafür ein Beispiel ist. Allgemein gilt für jede Menge M , dass $\#M < \#\mathcal{P}(M)$ ist; es gibt zu jeder Menge also eine noch „größere“.

Wir können jetzt die Dimension eines Vektorraums einführen.

✱

9.12. Definition. Sei V ein Vektorraum. Wenn V eine endliche Basis $(v_1, \dots, v_n)$ hat, dann sagen wir, dass V *Dimension n* hat oder *n -dimensional* ist und schreiben $\dim V = n$. Hat V keine endliche Basis, dann sagen wir, dass V *unendlich-dimensional* ist und schreiben $\dim V = \infty$. Hat V Dimension n für ein $n \in \mathbb{N}$, dann heißt V *endlich-dimensional* und wir schreiben $\dim V < \infty$.

DEF
Dimension

Wenn wir betonen wollen, dass es um die Dimension von V als K -Vektorraum geht, dann schreiben wir genauer $\dim_K V$. $\diamond$

Zum Beispiel ist $\dim_{\mathbb{C}} \mathbb{C} = 1$ ($\mathbb{C}$ -Basis (1)), aber $\dim_{\mathbb{R}} \mathbb{C} = 2$ ($\mathbb{R}$ -Basis $(1, i)$).

Folgerung 9.10 sagt uns, dass diese Definition sinnvoll ist, weil alle endlichen Basen von V (wenn es sie gibt) dieselbe Anzahl von Elementen haben.

9.13. Beispiele.

BSP
Dimension

- (1) Die leere Menge ist Basis des Null-Vektorraums, also ist $\dim\{\mathbf{0}\} = 0$. Ist umgekehrt V ein Vektorraum mit $\dim V = 0$, dann hat V eine Basis aus null Vektoren, also ist $V = \{\mathbf{0}\}$.
- (2) Für $n \in \mathbb{N}$ gilt $\dim K^n = n$, denn K^n hat die n -elementige Standardbasis $(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$.
- (3) Für den Vektorraum der Polynomfunktionen gilt $\dim P = \infty$, denn er hat eine unendliche Basis und kann deswegen nicht endlich-dimensional sein (siehe Folgerung 9.15 unten). ♣

Die Dimension eines Vektorraums ist eine wichtige Größe, wie die folgenden Aussagen zeigen.

9.14. **Satz.** Seien $m, n \in \mathbb{N}$, sei V ein n -dimensionaler Vektorraum und seien $v_1, v_2, \dots, v_m \in V$.

SATZ
Eigensch.
Dimension

- (1) Wenn $(v_1, v_2, \dots, v_m)$ linear unabhängig ist, dann ist $m \leq n$. Im Fall $m = n$ ist $(v_1, v_2, \dots, v_m)$ eine Basis von V .
- (2) Wenn $(v_1, v_2, \dots, v_m)$ ein Erzeugendensystem von V ist, dann ist $m \geq n$. Im Fall $m = n$ ist $(v_1, v_2, \dots, v_m)$ eine Basis von V .

Beweis.

- (1) Nach Folgerung 9.7 können wir $(v_1, v_2, \dots, v_m)$ durch Hinzunehmen von geeigneten Vektoren von V zu einer Basis von V ergänzen. Diese Basis hat n Elemente, also gilt $m \leq n$. Wenn $m = n$ ist, dann werden keine Elemente hinzugefügt, also liegt bereits eine Basis vor.
- (2) Nach dem Basisergänzungssatz 9.5 (mit $n = 0$ in der dortigen Notation) gibt es eine Basis, die durch Weglassen von geeigneten Vektoren v_j aus $(v_1, v_2, \dots, v_m)$ entsteht. Diese Basis hat Länge n , also gilt $m \geq n$. Wenn $m = n$ ist, dann kann nichts weggelassen werden, also liegt bereits eine Basis vor. □

Weil dieser Satz so wichtig ist, gebe ich eine weitere Formulierung.

Man kann den ersten Teil der beiden Aussagen auch so ausdrücken:

- (1) In einem **n -dimensionalen** Vektorraum sind **mehr als n Vektoren immer linear abhängig**.
- (2) Die **lineare Hülle** von **m Vektoren** hat **Dimension höchstens m** :

$$\dim\langle v_1, v_2, \dots, v_m \rangle \leq m.$$

Die erste dieser Aussagen ist eine starke *Existenzaussage*. Sie besagt nämlich Folgendes: Sind $v_1, v_2, \dots, v_m \in V$ mit $m > \dim V$, dann gibt es eine nichttriviale Linearkombination

$$\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m = \mathbf{0}.$$

Der zweite Teil der beiden Aussagen im Satz oben bedeutet:

- (1) *In einem n -dimensionalen Vektorraum sind n linear unabhängige Vektoren immer schon eine Basis.*
- (2) *In einem n -dimensionalen Vektorraum ist ein Erzeugendensystem mit n Elementen immer schon eine Basis.*

Linear unabhängige Familien geben also untere Schranken und Erzeugendensysteme geben obere Schranken für die Dimension. Es ist daher plausibel, dass wir unendlich-dimensionale Vektorräume wie folgt charakterisieren können.

9.15. Folgerung. *Sei V ein Vektorraum. Die folgenden Aussagen sind äquivalent:*

FOLG
 $\dim = \infty$

- (1) *Es gibt in V eine (unendliche) Folge $(v_n)_{n \in \mathbb{N}}$ linear unabhängiger Vektoren.*
- (2) $\dim V = \infty$.

Beweis. „(1) $\Rightarrow$ (2)“: Sei $(v_n)_{n \in \mathbb{N}}$ linear unabhängig. Wenn $\dim V = m < \infty$ wäre, dann müssten nach Satz 9.14 die $m + 1$ Vektoren $v_0, v_1, \dots, v_m$ linear abhängig sein, was aber der linearen Unabhängigkeit von $(v_n)_{n \in \mathbb{N}}$ widerspricht. Also muss V unendlich-dimensional sein.

„(2) $\Rightarrow$ (1)“: Sei V unendlich-dimensional. Das bedeutet, dass V keine endliche Basis hat; damit hat V (nach Folgerung 9.6) auch kein endliches Erzeugendensystem. Wir konstruieren rekursiv eine linear unabhängige Folge $(v_n)_{n \in \mathbb{N}}$ in V . Sei dazu $(v_0, v_1, \dots, v_{n-1})$ bereits konstruiert und linear unabhängig. (Für $n = 0$ ist das das leere Tupel, das immer linear unabhängig ist.) Da $(v_0, v_1, \dots, v_{n-1})$ kein Erzeugendensystem ist, gibt es $v_n \in V \setminus \langle v_0, v_1, \dots, v_{n-1} \rangle$. Nach Lemma 8.9 ist dann $(v_0, \dots, v_{n-1}, v_n)$ linear unabhängig. Die so konstruierte Folge $(v_n)_{n \in \mathbb{N}}$ ist linear unabhängig, weil das für alle endlichen Anfangsstücke gilt (in einer Linearkombination kommen nur endlich viele Vektoren v_n vor). $\square$

Wir können also sagen:

- Die Dimension von V ist die **maximale** Anzahl **linear unabhängiger** Vektoren in V .
- Die Dimension von V ist die **minimale** Anzahl von **Erzeugern** von V .

Es folgt eine Anwendung der Aussage, dass $n+1$ Vektoren in einem n -dimensionalen Vektorraum linear abhängig sein müssen. Dafür brauchen wir noch eine Definition und ein Lemma.

9.16. Definition. Wir sagen, eine Polynomfunktion $f \in P$ habe $\text{Grad} < n$ (und wir schreiben $\deg(f) < n$), wenn sie in der Form

DEF
Grad einer
Polynomfkt.

$$f(x) = a_0 + a_1x + \dots + a_{n-1}x^{n-1}$$

(mit $a_0, a_1, \dots, a_{n-1} \in \mathbb{R}$) geschrieben werden kann. f hat $\text{Grad } n$ ($\deg(f) = n$), wenn gilt $\deg(f) < n + 1$, aber nicht $\deg(f) < n$ (dann hat f die Form

$$f(x) = a_0 + a_1x + \dots + a_nx^n$$

mit $a_n \neq 0$). Man beachte, dass die Nullfunktion die einzige Polynomfunktion vom $\text{Grad} < 0$ ist. $\diamond$

Sie wissen aus der Schule, dass eine Polynomfunktion vom Grad n höchstens n reelle Nullstellen haben kann. Das kann man auch wie folgt ausdrücken (siehe auch Beispiel 8.8):

9.17. Lemma. *Ist f eine Polynomfunktion mit $\deg(f) < n$, die mindestens n reelle Nullstellen hat, dann ist f die Nullfunktion.*

LEMMA
Polynom = 0

Beweisen kann man das zum Beispiel mit Mitteln der Analysis. Dazu verwendet man einerseits, dass die Ableitung einer Polynomfunktion vom Grad $< n$ eine Polynomfunktion vom Grad $< n - 1$ ist, und andererseits den *Satz von Rolle*: Zwischen zwei Nullstellen einer differenzierbaren Funktion f liegt eine Nullstelle der Ableitung f' .

Der Beweis geht dann durch Induktion über n . Ist $n = 0$, dann ist f schon die Nullfunktion (denn „ $\deg(f) < 0$ “ bedeutet, dass in der Darstellung von f als $f(x) = a_0 + a_1x + \dots$ kein von null verschiedener Koeffizient a_j auftreten kann). Für den Induktionsschritt nehmen wir an, dass die Aussage für Polynomfunktionen vom Grad $< n$ gilt. Sei f eine Polynomfunktion vom Grad $< n + 1$ mit mindestens $n + 1$ reellen Nullstellen. Dann ist die Ableitung f' eine Polynomfunktion vom Grad $< n$ mit mindestens n reellen Nullstellen (nach dem Satz von Rolle liegt zwischen je zwei aufeinanderfolgenden Nullstellen von f eine Nullstelle von f'). Die Induktionsannahme liefert jetzt $f' = 0$. Dann muss f konstant sein; weil f aber mindestens eine Nullstelle hat ($n + 1 > 0$), muss $f = 0$ sein.

9.18. Beispiel. *Seien $x_1, \dots, x_n \in \mathbb{R}$ paarweise verschieden und $y_1, \dots, y_n \in \mathbb{R}$. Dann gibt es eine Polynomfunktion f mit $\deg(f) < n$, sodass $f(x_j) = y_j$ ist für alle $j \in \{1, 2, \dots, n\}$.*

BSP
Interpolation

Beweis. Wir betrachten die folgenden $n + 1$ Vektoren in $\mathbb{R}^n$:

$$\begin{aligned} v_0 &= (1, 1, 1, \dots, 1) \\ v_1 &= (x_1, x_2, x_3, \dots, x_n) \\ v_2 &= (x_1^2, x_2^2, x_3^2, \dots, x_n^2) \\ &\vdots \\ v_{n-1} &= (x_1^{n-1}, x_2^{n-1}, x_3^{n-1}, \dots, x_n^{n-1}) \\ v_n &= (y_1, y_2, y_3, \dots, y_n) \end{aligned}$$

Dann wissen wir, dass $v_0, v_1, \dots, v_n$ linear abhängig sein müssen, denn es ist $\dim \mathbb{R}^n = n < n + 1$. Es gibt also $\lambda_0, \lambda_1, \dots, \lambda_n \in \mathbb{R}$, nicht alle null, mit

$$\lambda_0 + \lambda_1 x_j + \lambda_2 x_j^2 + \dots + \lambda_{n-1} x_j^{n-1} + \lambda_n y_j = 0$$

für alle $j \in \{1, 2, \dots, n\}$. Ich behaupte, dass λ_n nicht null sein kann. Denn sonst hätte die Polynomfunktion

$$x \mapsto \lambda_0 + \lambda_1 x + \dots + \lambda_{n-1} x^{n-1}$$

vom Grad $< n$ mindestens die n Nullstellen $x_1, x_2, \dots, x_n$, müsste also nach Lemma 9.17 die Nullfunktion sein, was $\lambda_0 = \lambda_1 = \dots = \lambda_{n-1} = 0$ bedeuten würde. (Diese Implikation folgt aus der linearen Unabhängigkeit der Potenzfunktionen; vgl. Beispiel 8.8.) Dann wäre die obige Linearkombination aber trivial, ein Widerspruch. Also ist $\lambda_n \neq 0$. Wir setzen

$$a_0 = -\frac{\lambda_0}{\lambda_n}, \quad a_1 = -\frac{\lambda_1}{\lambda_n}, \quad \dots, \quad a_{n-1} = -\frac{\lambda_{n-1}}{\lambda_n}$$

und

$$f(x) = a_0 + a_1 x + \dots + a_{n-1} x^{n-1};$$

dann gilt

$$f(x_j) = a_0 + a_1 x_j + \dots + a_{n-1} x_j^{n-1} = y_j$$

wie gewünscht. $\square$

Als Nebenprodukt unserer Überlegungen ergab sich, dass die Vektoren $v_j = (x_1^j, \dots, x_n^j)$ für $j \in \{0, 1, \dots, n-1\}$ linear unabhängig sind.

Wir sehen also, dass man die *Existenz* einer interpolierenden Polynomfunktion im Wesentlichen direkt aus der linearen Abhängigkeit von „zu vielen“ Vektoren bekommt, ohne dafür rechnen zu müssen. Wenn man allerdings wissen möchte, wie diese Funktion konkret aussieht, braucht man andere Methoden. Wir werden darauf zurückkommen. $\clubsuit$

Die Dimension ist ein Maß für die „Größe“ eines Vektorraums. Das wird deutlich, wenn man die Dimension eines Untervektorraums betrachtet.

9.19. Satz. *Seien V ein Vektorraum und $U \subset V$ ein Untervektorraum. Dann gilt $\dim U \leq \dim V$. Ist V endlich-dimensional und gilt $\dim U = \dim V$, dann ist $U = V$.*

SATZ
Dimension
von Unter-VR

Dabei gelte $n \leq \infty$ für alle $n \in \mathbb{N}$ und $\infty \leq \infty$.

Beweis. Im Fall $\dim V = \infty$ ist die Aussage trivialerweise richtig. Sei jetzt also $\dim V = n \in \mathbb{N}$. Wäre $\dim U = \infty$, dann gäbe es nach Folgerung 9.15 unendlich viele linear unabhängige Elemente in U und damit auch in V , ein Widerspruch. Also ist U endlich-dimensional mit $\dim U = m \in \mathbb{N}$. Eine Basis von U besteht aus m linear unabhängigen Vektoren von V . Nach Satz 9.14 folgt $\dim U = m \leq n = \dim V$. Gilt $m = n$, dann ist die Basis von U bereits eine Basis von V und es folgt $U = V$. $\square$

9.20. Beispiel. Ein unendlich-dimensionaler Vektorraum kann durchaus echte Untervektorräume haben, die ihrerseits unendlich-dimensional sind. Zum Beispiel können wir im Vektorraum P der Polynomfunktionen den Untervektorraum P_g der geraden Polynomfunktionen betrachten:

$$P_g = \{f \in P \mid \forall x \in \mathbb{R}: f(-x) = f(x)\}.$$

(Prüfen Sie nach, dass P_g tatsächlich ein Untervektorraum von P ist!) Da die Funktion $x \mapsto x$, die ein Element von P ist, nicht in P_g liegt, gilt $P_g \neq P$. Auf der anderen Seite sind die geraden Potenzfunktionen $x \mapsto x^{2n}$ für $n \in \mathbb{N}$ alle linear unabhängig, also ist $\dim P_g = \infty$. $\clubsuit$

BSP
 $\dim U =$
 $\dim V = \infty$
 $U \subsetneq V$

10. LINEARE ABBILDUNGEN

Sei V ein K -Vektorraum und seien $v_1, v_2, \dots, v_n \in V$. Sei weiter $\phi = \phi_{(v_1, \dots, v_n)}$ die zugehörige Linearkombinationenabbildung

$$\phi: K^n \longrightarrow V, \quad (x_1, x_2, \dots, x_n) \longmapsto x_1 v_1 + x_2 v_2 + \dots + x_n v_n.$$

Dann gilt für $\mathbf{x} = (x_1, x_2, \dots, x_n) \in K^n$, $\mathbf{y} = (y_1, y_2, \dots, y_n) \in K^n$ und $\lambda \in K$:

$$\begin{aligned} \phi(\mathbf{x} + \mathbf{y}) &= \phi((x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n)) \\ &= \phi(x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \\ &= (x_1 + y_1)v_1 + (x_2 + y_2)v_2 + \dots + (x_n + y_n)v_n \\ &= (x_1 v_1 + x_2 v_2 + \dots + x_n v_n) + (y_1 v_1 + y_2 v_2 + \dots + y_n v_n) \\ &= \phi(x_1, x_2, \dots, x_n) + \phi(y_1, y_2, \dots, y_n) \\ &= \phi(\mathbf{x}) + \phi(\mathbf{y}) \end{aligned}$$

und

$$\begin{aligned} \phi(\lambda \mathbf{x}) &= \phi(\lambda(x_1, x_2, \dots, x_n)) \\ &= \phi(\lambda x_1, \lambda x_2, \dots, \lambda x_n) \\ &= (\lambda x_1)v_1 + (\lambda x_2)v_2 + \dots + (\lambda x_n)v_n \\ &= \lambda(x_1 v_1 + x_2 v_2 + \dots + x_n v_n) \\ &= \lambda \phi(x_1, x_2, \dots, x_n) \\ &= \lambda \phi(\mathbf{x}). \end{aligned}$$

(Man beachte, dass Addition und Skalarmultiplikation hier einmal in K^n und einmal in V stattfinden.)

Die Abbildung ϕ ist also mit Addition und Skalarmultiplikation verträglich: Das Bild einer Summe ist die Summe der Bilder und das Bild eines skalaren Vielfachen ist das entsprechende Vielfache des Bildes. Solche mit der linearen Struktur verträgliche Abbildungen heißen *lineare Abbildungen*.

*

10.1. Definition. Sei K ein Körper und seien V_1 und V_2 zwei K -Vektorräume. Eine Abbildung $\phi: V_1 \rightarrow V_2$ heißt *(K-)linear* oder ein *Homomorphismus* (von K -Vektorräumen), wenn sie die folgenden beiden Bedingungen erfüllt:

- (1) $\forall v, w \in V_1: \phi(v + w) = \phi(v) + \phi(w)$.
- (2) $\forall \lambda \in K \forall v \in V_1: \phi(\lambda v) = \lambda \phi(v)$.

DEF
Lineare
Abbildung
Homo-
morphismus

Eine lineare Abbildung heißt ein *Monomorphismus*, wenn sie injektiv ist, ein *Epi-morphismus*, wenn sie surjektiv ist, und ein *Isomorphismus*, wenn sie bijektiv ist. Eine lineare Abbildung $\phi: V \rightarrow V$ heißt ein *Endomorphismus* von V ; ϕ heißt ein *Automorphismus* von V , wenn ϕ außerdem bijektiv ist. Zwei Vektorräume V_1 und V_2 heißen (zueinander) *isomorph*, $V_1 \cong V_2$, wenn es einen Isomorphismus $\phi: V_1 \rightarrow V_2$ gibt. $\diamond$

Mono-, Epi-,
Iso-, Endo-,
Automorph.
isomorph

10.2. Beispiele.

- (1) Für beliebige K -Vektorräume V_1 und V_2 ist die *Nullabbildung* $V_1 \rightarrow V_2$, $v \mapsto \mathbf{0}$, eine lineare Abbildung.
- (2) Für jeden K -Vektorraum V ist die identische Abbildung $\text{id}_V: V \rightarrow V$ ein Automorphismus von V .

BSP
lineare
Abbildungen

- (3) Ist V ein Vektorraum und $U \subset V$ ein Untervektorraum, dann ist die Inklusionsabbildung $U \rightarrow V$ linear. ♣

10.3. Beispiel. Seien $V, v_1, v_2, \dots, v_n$ und ϕ wie zum Beginn dieses Abschnitts. Dann ist ϕ ein Homomorphismus. Nach Lemma 9.4 gilt außerdem:

- $(v_1, v_2, \dots, v_n)$ ist linear unabhängig $\iff \phi$ ist ein Monomorphismus.
- $(v_1, v_2, \dots, v_n)$ ist ein Erzeugendensystem von $V \iff \phi$ ist ein Epimorphismus.
- $(v_1, v_2, \dots, v_n)$ ist eine Basis von $V \iff \phi$ ist ein Isomorphismus.

BSP
Linearkomb.-
Abbildung
ist linear

Aus dem letzten Punkt ergibt sich die Aussage

$$\dim V < \infty \implies V \cong K^{\dim V}.$$



Wir überzeugen uns noch davon, dass eine lineare Abbildung wirklich mit der gesamten Struktur verträglich ist und dass sich lineare Abbildungen bezüglich Komposition und Inversion gut verhalten.

10.4. Lemma. V_1, V_2 und V_3 seien K -Vektorräume.

LEMMA
Lin. Abb.:
Eigensch.

- (1) Sei $\phi: V_1 \rightarrow V_2$ eine lineare Abbildung. Dann gilt

$$\phi(\mathbf{0}) = \mathbf{0} \quad \text{und} \quad \forall v \in V_1: \phi(-v) = -\phi(v).$$

- (2) Seien $\phi_1: V_1 \rightarrow V_2$ und $\phi_2: V_2 \rightarrow V_3$ lineare Abbildungen. Dann ist auch $\phi_2 \circ \phi_1: V_1 \rightarrow V_3$ linear.

- (3) Sei $\phi: V_1 \rightarrow V_2$ ein Isomorphismus. Dann ist die Umkehrabbildung $\phi^{-1}: V_2 \rightarrow V_1$ ebenfalls ein Isomorphismus.

Teil (3) zeigt, dass es in der Definition von „isomorph“ nicht darauf ankommt, ob man einen Isomorphismus $V_1 \rightarrow V_2$ oder einen Isomorphismus $V_2 \rightarrow V_1$ fordert.

Beweis.

- (1) Es gilt $\phi(\mathbf{0}) = \phi(\mathbf{0} + \mathbf{0}) = \phi(\mathbf{0}) + \phi(\mathbf{0})$. Durch Addition von $-\phi(\mathbf{0})$ folgt $\phi(\mathbf{0}) = \mathbf{0}$.

Außerdem gilt für $v \in V_1$: $\phi(-v) = \phi((-1)v) = (-1)\phi(v) = -\phi(v)$.

- (2) Wir müssen die beiden Eigenschaften aus Definition 10.1 für $\phi_2 \circ \phi_1$ nachweisen. Seien dazu $v, w \in V_1$ und $\lambda \in K$. Dann gilt

$$\begin{aligned} (\phi_2 \circ \phi_1)(v + w) &= \phi_2(\phi_1(v + w)) = \phi_2(\phi_1(v) + \phi_1(w)) \\ &= \phi_2(\phi_1(v)) + \phi_2(\phi_1(w)) = (\phi_2 \circ \phi_1)(v) + (\phi_2 \circ \phi_1)(w) \end{aligned}$$

und

$$(\phi_2 \circ \phi_1)(\lambda v) = \phi_2(\phi_1(\lambda v)) = \phi_2(\lambda \phi_1(v)) = \lambda \phi_2(\phi_1(v)) = \lambda (\phi_2 \circ \phi_1)(v).$$

- (3) Wir weisen die Eigenschaften aus Definition 10.1 für ϕ^{-1} nach. Seien dazu $v, w \in V_2$ und $\lambda \in K$. Wir setzen $v' = \phi^{-1}(v)$ und $w' = \phi^{-1}(w)$, sodass $v = \phi(v')$ und $w = \phi(w')$. Dann gilt

$$\phi^{-1}(v + w) = \phi^{-1}(\phi(v') + \phi(w')) = \phi^{-1}(\phi(v' + w')) = v' + w' = \phi^{-1}(v) + \phi^{-1}(w)$$

und

$$\phi^{-1}(\lambda v) = \phi^{-1}(\lambda \phi(v')) = \phi^{-1}(\phi(\lambda v')) = \lambda v' = \lambda \phi^{-1}(v).$$

Außerdem ist ϕ^{-1} als Umkehrabbildung einer bijektiven Abbildung bijektiv. $\square$

Bevor wir weitere Eigenschaften untersuchen, führen wir noch eine Schreibweise ein.

10.5. Definition. Sei $f: X \rightarrow Y$ eine Abbildung zwischen beliebigen Mengen X und Y . Ist T eine Teilmenge von X , dann schreiben wir

$$f(T) = \{f(x) \mid x \in T\} \subset Y$$

für die Menge der Bilder der Elemente von T und nennen $f(T)$ das *Bild von T unter f* . Im Spezialfall $T = X$ schreiben wir auch $\text{im}(f)$ für $f(X)$; $\text{im}(f)$ heißt das *Bild* oder die *Bildmenge* von f . Ist U eine Teilmenge von Y , dann schreiben wir

$$f^{-1}(U) = \{x \in X \mid f(x) \in U\} \subset X$$

für die Menge der Urbilder der Elemente von U und nennen $f^{-1}(U)$ das *Urbild von U unter f* . $\diamond$

Aus $T \subset T' \subset X$ folgt $f(T) \subset f(T')$, und aus $U \subset U' \subset Y$ folgt $f^{-1}(U) \subset f^{-1}(U')$.

Häufig wird auch $f^{-1}(y) = \{x \in X \mid f(x) = y\}$ für die Menge der Urbilder eines Elements $y \in Y$ geschrieben. Das kann zu Verwirrung führen, denn wenn f bijektiv ist, dann bedeutet $f^{-1}(y)$ auch das Bild von y unter der Umkehrabbildung, also das Urbild von y und nicht die Urbildmenge von y . Wir werden die „Datentypen“ (Elemente bzw. Teilmengen) hier aber sorgfältig auseinanderhalten und immer $f^{-1}(\{y\})$ für diese Menge schreiben.

Wenn f bijektiv ist, dann hat $f^{-1}(U)$ zwei mögliche Bedeutungen: einerseits ausgehend von f wie oben definiert und andererseits ausgehend von der Umkehrfunktion f^{-1} . Zum Glück stimmen beide Versionen überein.

Noch eine **Warnung**: Die hier eingeführte Schreibweise kann einen dazu verführen zu denken, dass $f^{-1}(f(T)) = T$ und $f(f^{-1}(U)) = U$ sein muss. Das ist aber im Allgemeinen **falsch**! Es gilt immer $f^{-1}(f(T)) \supset T$ und $f(f^{-1}(U)) \subset U$; die Inklusionen können jedoch echt sein. Finden Sie Beispiele dafür!

Da der Nullvektor eine ausgezeichnete Rolle in einem Vektorraum spielt, ist die Menge seiner Urbilder unter einer linearen Abbildung ein wichtiges Datum.

***** **10.6. Definition.** Sei $\phi: V_1 \rightarrow V_2$ eine lineare Abbildung. Der *Kern* von ϕ ist die Menge der Urbilder von $\mathbf{0} \in V_2$:

$$\ker(\phi) = \phi^{-1}(\{\mathbf{0}\}) = \{v \in V_1 \mid \phi(v) = \mathbf{0}\} \subset V_1.$$

$\diamond$

Nach Lemma 10.4 gilt stets $\mathbf{0} \in \ker(\phi)$.

10.7. Beispiel. Sei $V \subset \mathbb{R}^{\mathbb{N}}$ der Vektorraum der konvergenten Folgen. Dann ist

$$\lim: V \longrightarrow \mathbb{R}, \quad (a_n)_{n \in \mathbb{N}} \longmapsto \lim_{n \rightarrow \infty} a_n$$

eine $\mathbb{R}$ -lineare Abbildung. Das folgt aus den Rechenregeln für Grenzwerte.

Der Kern $\ker(\lim)$ ist gerade die Menge der Nullfolgen, denn das sind definitionsgemäß die Folgen mit Limes null. $\clubsuit$

Eine wichtige Eigenschaft des Kerns ist die folgende:

DEF
Bilder und
Urbilder von
Teilmengen



DEF
Kern

BSP
Limes ist
linear

10.8. **Lemma.** Sei $\phi: V_1 \rightarrow V_2$ eine lineare Abbildung. Dann gilt:

$$\phi \text{ ist injektiv} \iff \ker(\phi) = \{\mathbf{0}\}.$$

LEMMA

injektiv

$\iff$

$\ker = \{\mathbf{0}\}$

Beweis. „ $\Rightarrow$ “: Sei ϕ injektiv und $v \in \ker(\phi)$. Dann ist $\phi(v) = \mathbf{0} = \phi(\mathbf{0})$, also $v = \mathbf{0}$.

„ $\Leftarrow$ “: Es gelte $\ker(\phi) = \{\mathbf{0}\}$. Seien weiter $v, w \in V_1$ mit $\phi(v) = \phi(w)$. Dann folgt $\mathbf{0} = \phi(v) - \phi(w) = \phi(v - w)$, also ist $v - w \in \ker(\phi) = \{\mathbf{0}\}$ und damit $v - w = \mathbf{0}$; das bedeutet $v = w$. $\square$

Wenn $\ker(\phi) = \{\mathbf{0}\}$ ist, sagt man auch, der Kern sei *trivial*.

DEF

Kern trivial

Wie zu erwarten, vertragen sich lineare Abbildungen sehr gut mit Untervektorräumen.

10.9. **Satz.** Sei $\phi: V_1 \rightarrow V_2$ eine K -lineare Abbildung.

SATZ

lin. Abb.

und UVR

- (1) Ist $U_1 \subset V_1$ ein Untervektorraum, dann ist $\phi(U_1) \subset V_2$ wieder ein Untervektorraum. Insbesondere ist $\text{im}(\phi) = \phi(V_1) \subset V_2$ ein Untervektorraum von V_2 . Außerdem ist die auf U_1 eingeschränkte Abbildung $\phi|_{U_1}: U_1 \rightarrow V_2$ ebenfalls linear.
- (2) Ist $U_2 \subset V_2$ ein Untervektorraum, dann ist $\phi^{-1}(U_2) \subset V_1$ wieder ein Untervektorraum. Insbesondere ist $\ker(\phi) = \phi^{-1}(\{\mathbf{0}\}) \subset V_1$ ein Untervektorraum von V_1 .

(3) Seien

$$M_1 = \{U_1 \mid U_1 \subset V_1 \text{ Untervektorraum mit } \ker(\phi) \subset U_1\} \quad \text{und}$$

$$M_2 = \{U_2 \mid U_2 \subset V_2 \text{ Untervektorraum mit } U_2 \subset \text{im}(\phi)\}.$$

Die Abbildungen $M_1 \rightarrow M_2$, $U_1 \mapsto \phi(U_1)$, und $M_2 \rightarrow M_1$, $U_2 \mapsto \phi^{-1}(U_2)$, sind zueinander inverse Bijektionen.

Die Aussage, dass der Kern einer linearen Abbildung ein Untervektorraum ist, ist oft nützlich, wenn man zeigen möchte, dass eine Teilmenge eines Vektorraums ein Untervektorraum ist. Oft kann man nämlich Untervektorräume in natürlicher Weise als Kerne schreiben.

Beweis.

- (1) Wir müssen die Bedingungen für einen Untervektorraum für $\phi(U_1)$ nachprüfen:

- $\mathbf{0} = \phi(\mathbf{0}) \in \phi(U_1)$, da $\mathbf{0} \in U_1$.
- Seien $v, w \in \phi(U_1)$. Dann gibt es $v', w' \in U_1$ mit $\phi(v') = v$, $\phi(w') = w$. Es folgt $v + w = \phi(v') + \phi(w') = \phi(v' + w') \in \phi(U_1)$, denn $v' + w' \in U_1$.
- Sei $v \in \phi(U_1)$ und $\lambda \in K$. Dann gibt es $v' \in U_1$ mit $\phi(v') = v$. Es folgt $\lambda v = \lambda \phi(v') = \phi(\lambda v') \in \phi(U_1)$, denn $\lambda v' \in U_1$.

Da V_1 ein Untervektorraum von V_1 ist, folgt, dass $\text{im}(\phi)$ ein Untervektorraum von V_2 ist.

Dass $\phi|_{U_1}$ linear ist, folgt daraus, dass man $\phi|_{U_1}$ als Verkettung der (linearen) Inklusionsabbildung $U_1 \rightarrow V_1$ und ϕ schreiben kann.

- (2) Wir prüfen die Bedingungen für $\phi^{-1}(U_2)$:

- $\mathbf{0} \in \phi^{-1}(U_2)$, da $\phi(\mathbf{0}) = \mathbf{0} \in U_2$.

- Seien $v, w \in \phi^{-1}(U_2)$. Dann sind $\phi(v), \phi(w) \in U_2$. Es folgt $\phi(v + w) = \phi(v) + \phi(w) \in U_2$ und damit $v + w \in \phi^{-1}(U_2)$.
- Seien $v \in \phi^{-1}(U_2)$ und $\lambda \in K$. Dann ist $\phi(v) \in U_2$, also auch $\phi(\lambda v) = \lambda\phi(v) \in U_2$, und damit $\lambda v \in \phi^{-1}(U_2)$.

Da $\{0\}$ ein Untervektorraum von V_2 ist, folgt, dass $\ker(\phi)$ ein Untervektorraum von V_1 ist.

- (3) Wir überlegen zunächst, dass die Abbildungen wohldefiniert sind: Für $U_1 \subset V_1$ gilt $\phi(U_1) \subset \phi(V_1) = \text{im}(\phi)$ und für $U_2 \subset V_2$ gilt $\phi^{-1}(U_2) \supset \phi^{-1}(\{0\}) = \ker(\phi)$. Nach den Teilen (1) und (2) werden Untervektorräume auf Untervektorräume abgebildet. Damit haben wir tatsächlich Abbildungen zwischen den beiden angegebenen Mengen.

Wir zeigen jetzt, dass die Abbildungen zueinander invers sind. Daraus folgt dann auch, dass sie bijektiv sind. Sei also $U_1 \subset V_1$ ein Untervektorraum mit $\ker(\phi) \subset U_1$. Dann gilt

$$\begin{aligned}
 v \in \phi^{-1}(\phi(U_1)) &\iff \phi(v) \in \phi(U_1) \\
 &\iff \exists v' \in U_1: \phi(v) = \phi(v') \\
 &\iff \exists v' \in U_1: \phi(v - v') = 0 \\
 &\iff \exists v' \in U_1: v - v' \in \ker(\phi) \\
 &\iff v \in U_1.
 \end{aligned}$$

(Die letzte Äquivalenz sieht man so: „ $\Leftarrow$ “: wähle $v' = v$. „ $\Rightarrow$ “: Sei $v'' = v - v' \in \ker(\phi) \subset U_1$, dann ist $v = v' + v'' \in U_1$.) Das zeigt $\phi^{-1}(\phi(U_1)) = U_1$. Sei jetzt $U_2 \subset \text{im}(\phi)$ ein Untervektorraum von V_2 . Dann gilt

$$\begin{aligned}
 v \in \phi(\phi^{-1}(U_2)) &\iff \exists v' \in \phi^{-1}(U_2): \phi(v') = v \\
 &\iff \exists v' \in V_1: \phi(v') \in U_2 \text{ und } \phi(v') = v \\
 &\iff v \in U_2 \quad \text{und} \quad v \in \text{im}(\phi) \\
 &\iff v \in U_2 \cap \text{im}(\phi) \\
 &\iff v \in U_2.
 \end{aligned}$$

Das zeigt $\phi(\phi^{-1}(U_2)) = U_2$. □

Den Zusammenhang, der im dritten Teil dieses Satzes formuliert wird, kann man sich etwa (sehr schematisch) wie in Abbildung 2 veranschaulichen.

10.10. Beispiel. Seien K ein Körper, X eine Menge und V ein Untervektorraum von $K^X = \text{Abb}(X, K)$ (zum Beispiel können wir $X = K = \mathbb{R}$ setzen und für V den Vektorraum der stetigen reellen Funktionen nehmen). Sei weiter $x \in X$. Dann ist die *Auswertungsabbildung*

$$\text{ev}_x: V \longrightarrow K, \quad f \longmapsto f(x)$$

linear. Das ergibt sich direkt aus der Definition der Addition und der skalaren Multiplikation von Funktionen:

$$\begin{aligned}
 \text{ev}_x(f + g) &= (f + g)(x) = f(x) + g(x) = \text{ev}_x(f) + \text{ev}_x(g) \quad \text{und} \\
 \text{ev}_x(\lambda f) &= (\lambda f)(x) = \lambda f(x) = \lambda \text{ev}_x(f).
 \end{aligned}$$

(Man kann sagen, dass die Addition und Skalarmultiplikation in K^X gerade so definiert sind, *damit* die Auswertungsabbildungen linear werden!)

BSP
Auswertungs-
abbildung
DEF
Auswertungs-
abbildung

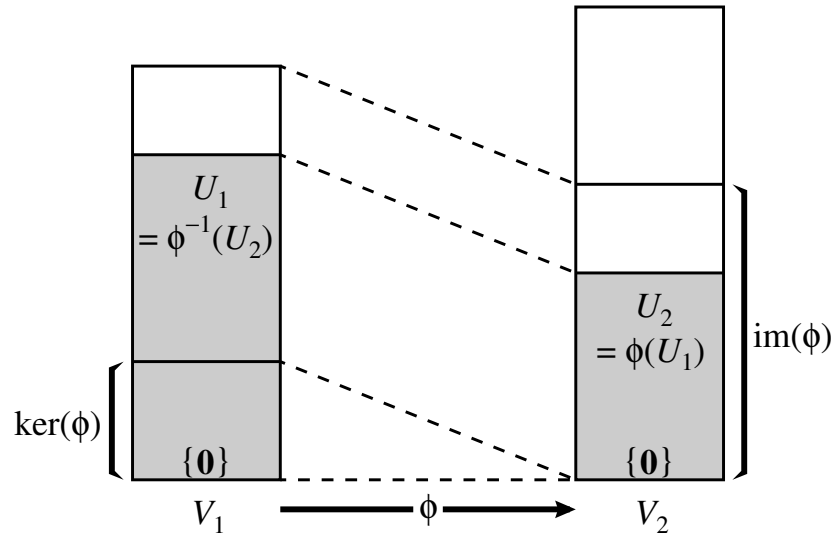


ABBILDUNG 2. Skizze zu Satz 10.9

Sei T eine nichtleere Teilmenge von X . Dann ist

$$\{f \in V \mid \forall x \in T: f(x) = 0\} = \bigcap_{x \in T} \ker(\text{ev}_x)$$

ein Untervektorraum von V .

Im Spezialfall $X = \{1, 2, \dots, n\}$ haben wir $K^X = K^n$; dann heißen die Abbildungen ev_j (für $j \in \{1, 2, \dots, n\}$) *Projektionen* und werden pr_j geschrieben:

$$\text{pr}_j: K^n \longrightarrow K, \quad (a_1, a_2, \dots, a_n) \longmapsto a_j$$

DEF
Projektion

Sie sind also ebenfalls linear. ♣

Wir zeigen jetzt, dass eine lineare Abbildung dadurch festgelegt ist, was sie auf einer Basis macht.

* **10.11. Satz.** Sei V ein K -Vektorraum mit Basis $(b_1, b_2, \dots, b_n)$ und sei W ein weiterer K -Vektorraum. Seien weiter $w_1, w_2, \dots, w_n \in W$. Dann gibt es genau eine K -lineare Abbildung $\phi: V \rightarrow W$ mit $\phi(b_j) = w_j$ für alle $j \in \{1, 2, \dots, n\}$.

SATZ
Basen und
lin. Abb.

Beweis. Wir beweisen zuerst die Eindeutigkeit. Seien also $\phi_1, \phi_2: V \rightarrow W$ lineare Abbildungen mit $\phi_1(b_j) = w_j = \phi_2(b_j)$ für alle $j \in \{1, 2, \dots, n\}$. Sei $v \in V$ beliebig. Dann ist v eine Linearkombination der Basisvektoren:

$$v = \lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n.$$

Es folgt

$$\begin{aligned} \phi_1(v) &= \phi_1(\lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n) \\ &= \lambda_1 \phi_1(b_1) + \lambda_2 \phi_1(b_2) + \dots + \lambda_n \phi_1(b_n) \\ &= \lambda_1 w_1 + \lambda_2 w_2 + \dots + \lambda_n w_n \\ &= \lambda_1 \phi_2(b_1) + \lambda_2 \phi_2(b_2) + \dots + \lambda_n \phi_2(b_n) \\ &= \phi_2(\lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n) \\ &= \phi_2(v), \end{aligned}$$

also ist $\phi_1 = \phi_2$.

Dieser Eindeutigkeitsbeweis zeigt uns, wie wir die Existenz beweisen können: Wenn es eine lineare Abbildung $\phi: V \rightarrow W$ gibt mit $\phi(b_j) = w_j$ für alle $j \in \{1, 2, \dots, n\}$, dann muss für $v \in V$ wie oben gelten

$$\phi(v) = \lambda_1 w_1 + \lambda_2 w_2 + \dots + \lambda_n w_n.$$

Wir müssen prüfen,

- (1) dass ϕ wohldefiniert ist, dass also $\phi(v)$ nicht davon abhängt, wie v als Linearkombination der b_j geschrieben wurde, und
- (2) dass die so definierte Abbildung ϕ linear ist.

Die Wohldefiniertheit folgt daraus, dass v nur auf genau eine Weise als Linearkombination der Basisvektoren geschrieben werden kann (vgl. Lemma 9.4). Die Linearität rechnet man nach. Etwas eleganter ist es, wenn man bemerkt, dass $\phi = \phi_{(w_1, w_2, \dots, w_n)} \circ \phi_{(b_1, b_2, \dots, b_n)}^{-1}$ ist (mit den zu $(b_1, b_2, \dots, b_n)$ und zu $(w_1, w_2, \dots, w_n)$ gehörigen Linearkombinationenabbildungen $K^n \rightarrow V$ bzw. $K^n \rightarrow W$ — man beachte, dass $\phi_{(b_1, b_2, \dots, b_n)}$ hier ein Isomorphismus ist). Die Linearität von ϕ folgt dann daraus, dass die Linearkombinationenabbildungen linear sind (Beispiel 10.3) und aus Lemma 10.4. $\square$

Das analoge Resultat gilt auch für (nicht unbedingt endliche) Basismengen:

Sind V und W K -Vektorräume und ist $B \subset V$ eine Basis, dann gibt es zu jeder Abbildung $f: B \rightarrow W$ genau eine lineare Abbildung $\phi: V \rightarrow W$ mit $\phi(b) = f(b)$ für alle $b \in B$ (oder kurz: $\phi|_B = f$).

Der Beweis geht im Wesentlichen genauso unter Verwendung der allgemeinen Linearkombinationenabbildungen $K^{(B)} \rightarrow V$ bzw. $K^{(B)} \rightarrow W$.

Für praktische Zwecke lässt sich der Inhalt von Satz 10.11 wie folgt zusammenfassen. Seien dazu V ein K -Vektorraum mit Basis $(b_1, b_2, \dots, b_n)$ und W ein weiterer K -Vektorraum.

- Wir können eine lineare Abbildung $V \rightarrow W$ definieren, indem wir die Bilder der Basisvektoren b_j beliebig festlegen.
- Zwei lineare Abbildungen $V \rightarrow W$ sind schon dann gleich, wenn die Bilder der Basisvektoren b_j unter beiden Abbildungen übereinstimmen.

Da eine lineare Abbildung also durch das Bild einer Basis eindeutig bestimmt ist, sollten sich auch Eigenschaften wie injektiv oder surjektiv zu sein durch das Bild der Basis ausdrücken lassen.

10.12. Satz. *Seien V und W K -Vektorräume und sei $\phi: V \rightarrow W$ linear. Sei weiter $(b_1, b_2, \dots, b_n)$ eine Basis von V .*

SATZ
inj./surj.
lin. Abb.

- (1) ϕ ist genau dann injektiv, wenn $(\phi(b_1), \phi(b_2), \dots, \phi(b_n))$ linear unabhängig ist.
- (2) ϕ ist genau dann surjektiv, wenn $(\phi(b_1), \phi(b_2), \dots, \phi(b_n))$ ein Erzeugendensystem von W ist.
- (3) ϕ ist genau dann ein Isomorphismus, wenn $(\phi(b_1), \phi(b_2), \dots, \phi(b_n))$ eine Basis von W ist.

Beweis. Seien $w_1 = \phi(b_1)$, $w_2 = \phi(b_2)$, $\dots$, $w_n = \phi(b_n)$. Wie im Beweis von Satz 10.11 ist dann $\phi = \phi_{(w_1, w_2, \dots, w_n)} \circ \phi_{(b_1, b_2, \dots, b_n)}^{-1}$. Da $\phi_{(b_1, b_2, \dots, b_n)}$ bijektiv ist, ist

ϕ injektiv bzw. surjektiv genau dann, wenn $\phi_{(w_1, w_2, \dots, w_n)}$ die entsprechende Eigenschaft hat (beachte dafür $\phi \circ \phi_{(b_1, b_2, \dots, b_n)} = \phi_{(w_1, w_2, \dots, w_n)}$). Die Behauptungen folgen dann sofort aus den Aussagen von Lemma 9.4 (siehe auch Beispiel 10.3). $\square$

Daraus können wir gleich zwei wichtige Folgerungen ziehen.

10.13. Folgerung. *Sind V und W zwei K -Vektorräume derselben endlichen Dimension n , dann sind V und W isomorph.*

FOLG
endl.-dim.
VR gleicher
Dimension
sind
isomorph

Beweis. Sei $(b_1, b_2, \dots, b_n)$ eine Basis von V und sei $(b'_1, b'_2, \dots, b'_n)$ eine Basis von W . Dann gibt es nach Satz 10.11 eine lineare Abbildung $\phi: V \rightarrow W$ mit $\phi(b_j) = b'_j$ für alle $j \in \{1, 2, \dots, n\}$. Nach Satz 10.12 ist ϕ ein Isomorphismus. $\square$

10.14. Folgerung. *Seien V und W zwei K -Vektorräume derselben endlichen Dimension n und sei $\phi: V \rightarrow W$ eine lineare Abbildung. Dann sind die folgenden Aussagen äquivalent:*

FOLG
lin. Abb.
bei gleicher
Dimension

- (1) ϕ ist ein Isomorphismus.
- (2) ϕ ist injektiv.
- (3) ϕ ist surjektiv.

Beweis. Es ist klar, dass aus (1) die beiden Aussagen (2) und (3) folgen. Sei $(b_1, \dots, b_n)$ eine Basis von V . Nach Satz 10.12 ist ϕ genau dann injektiv, wenn $(\phi(b_1), \dots, \phi(b_n))$ linear unabhängig ist. n linear unabhängige Vektoren bilden aber eine Basis (wegen $\dim W = n$, siehe Satz 9.14); nach Satz 10.12 ist ϕ dann ein Isomorphismus. Analog ist ϕ genau dann surjektiv, wenn $(\phi(b_1), \dots, \phi(b_n))$ den Vektorraum W erzeugt. Ein Erzeugendensystem aus n Elementen ist aber wieder eine Basis, also ist ϕ dann ein Isomorphismus. $\square$

Als Merkhilfe für diese wichtige Aussage kann die Analogie zu Abbildungen zwischen endlichen Mengen dienen. Es gilt nämlich (wie leicht einzusehen ist):

Seien X und Y zwei endliche Mengen mit $\#X = \#Y = n$ und sei $f: X \rightarrow Y$ eine Abbildung. Dann sind die folgenden Aussagen äquivalent:

- (1) f ist bijektiv.
- (2) f ist injektiv.
- (3) f ist surjektiv.

Die im folgenden Beispiel verwendete Produktschreibweise

$$\prod_{i \in I} a_i$$

ist analog definiert wie die Schreibweise mit dem Summenzeichen $\sum$, nur dass die angegebenen Elemente multipliziert werden statt addiert. Ist $I = \emptyset$, dann hat das („leere“) Produkt den Wert 1.

10.15. Beispiel. Der Vektorraum $P_{<n}$ der Polynomfunktionen vom Grad $< n$ wird von den n Potenzfunktionen $x \mapsto x^j$ für $j \in \{0, 1, \dots, n-1\}$ erzeugt. Diese Funktionen sind linear unabhängig, also hat $P_{<n}$ Dimension n .

BSP
Interpolation

Seien $x_1, x_2, \dots, x_n \in \mathbb{R}$ paarweise verschieden. Wir definieren für $j \in \{1, 2, \dots, n\}$ die Polynomfunktion $p_j \in P_{<n}$ durch

$$p_j(x) = \prod_{i \in \{1, \dots, n\} \setminus \{j\}} \frac{x - x_i}{x_j - x_i} = \frac{x - x_1}{x_j - x_1} \cdots \frac{x - x_{j-1}}{x_j - x_{j-1}} \cdot \frac{x - x_{j+1}}{x_j - x_{j+1}} \cdots \frac{x - x_n}{x_j - x_n}.$$

Für $n = 3$ und $(x_1, x_2, x_3) = (-1, 0, 1)$ bedeutet das zum Beispiel:

$$\begin{aligned} p_1(x) &= \frac{(x - x_2)(x - x_3)}{(x_1 - x_2)(x_1 - x_3)} = \frac{x(x - 1)}{(-1)(-2)} = \frac{1}{2}x^2 - \frac{1}{2}x \\ p_2(x) &= \frac{(x - x_1)(x - x_3)}{(x_2 - x_1)(x_2 - x_3)} = \frac{(x + 1)(x - 1)}{1 \cdot (-1)} = -x^2 + 1 \\ p_3(x) &= \frac{(x - x_1)(x - x_2)}{(x_3 - x_1)(x_3 - x_2)} = \frac{(x + 1)x}{2 \cdot 1} = \frac{1}{2}x^2 + \frac{1}{2}x \end{aligned}$$

Es gilt dann für $i, j \in \{1, 2, \dots, n\}$:

$$p_j(x_i) = \begin{cases} 1 & \text{falls } i = j; \\ 0 & \text{falls } i \neq j. \end{cases}$$

Wir definieren

$$\phi: P_{<n} \longrightarrow \mathbb{R}^n, \quad f \longmapsto (f(x_1), f(x_2), \dots, f(x_n))$$

(ϕ ist aus Auswertungsabbildungen zusammengesetzt und daher linear) und eine lineare Abbildung $\psi: \mathbb{R}^n \rightarrow P_{<n}$ durch Festlegung der Bilder der Standardbasis:

$$\psi(\mathbf{e}_j) = p_j \quad \text{für alle } j \in \{1, 2, \dots, n\}.$$

(ψ ist gerade die Linearkombinationenabbildung $\phi_{(p_1, \dots, p_n)} \cdot$.) Dann gilt $\phi \circ \psi = \text{id}_{\mathbb{R}^n}$:

$$\phi(\psi(\mathbf{e}_j)) = \phi(p_j) = (p_j(x_1), \dots, p_j(x_n)) = \mathbf{e}_j.$$

$\phi \circ \psi$ und die identische Abbildung stimmen auf einer Basis überein, also sind sie gleich. Dann muss ψ injektiv sein und ϕ surjektiv. Nach Folgerung 10.14 sind wegen $\dim \mathbb{R}^n = n = \dim P_{<n}$ beide Abbildungen (zueinander inverse) Isomorphismen. Das bedeutet:

Seien $x_1, x_2, \dots, x_n \in \mathbb{R}$ paarweise verschieden und $y_1, y_2, \dots, y_n \in \mathbb{R}$. Dann gibt es **genau eine** Polynomfunktion $f \in P_{<n}$ mit $f(x_j) = y_j$ für alle $j \in \{1, 2, \dots, n\}$, nämlich $f = y_1 p_1 + y_2 p_2 + \dots + y_n p_n$.

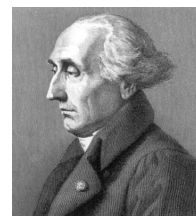
Das sieht man so: Die Bedingung an f bedeutet $f \in P_{<n}$ und $\phi(f) = (y_1, y_2, \dots, y_n)$. Letzteres ist aber äquivalent zu

$$f = \psi(y_1, y_2, \dots, y_n) = y_1 p_1 + y_2 p_2 + \dots + y_n p_n.$$

Diese Formel für das Interpolationspolynom heißt Lagrangesche Interpolationsformel.

(Dass ϕ bijektiv sein muss, kann man alternativ auch so sehen: Aus Beispiel 9.18 wissen wir bereits, dass ϕ surjektiv ist. Da $\dim P_{<n} = n = \dim \mathbb{R}^n$ gilt, sagt Folgerung 10.14, dass ϕ schon bijektiv sein muss. Das sagt uns aber noch nicht, wie die Umkehrabbildung aussieht.) ♣

Wir bleiben bei den Polynomfunktionen und geben weitere Beispiele für lineare Abbildungen.



J.-L. Lagrange
1736–1813

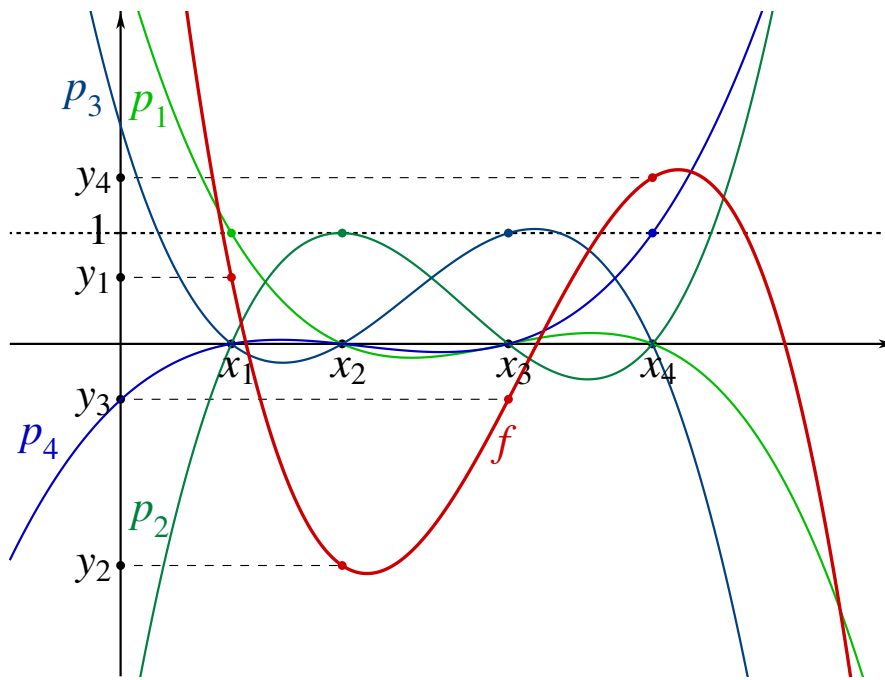


ABBILDUNG 3. Interpolationspolynom (zu Beispiel 10.15); hier ist der Fall mit vier Stützstellen $x_1, \dots, x_4$ dargestellt.

Beispiele.

BSP

lin. Abb. auf
Polynomfkt.

- (1) Wir haben schon in Beispiel 10.10 gesehen, dass für $a \in \mathbb{R}$ die Auswertungsabbildung

$$\text{ev}_a: P \longrightarrow \mathbb{R}, \quad f \longmapsto f(a)$$

linear ist.

- (2) Die *Differentiation* von Polynomfunktionen ist linear:

$$D: P \longrightarrow P, \quad f \longmapsto f'.$$

Für $f(x) = a_0 + a_1x + \dots + a_nx^n$ gilt dabei $f'(x) = a_1 + 2a_2x + \dots + na_nx^{n-1}$. Man könnte D also definieren als diejenige lineare Abbildung, die für $n > 0$ die Potenzfunktion $f_n: x \mapsto x^n$ auf nf_{n-1} abbildet und f_0 auf die Nullfunktion.

D ist surjektiv (also ein Epimorphismus) und der Kern von D besteht genau aus den konstanten Funktionen. (An diesem Beispiel kann man sehen, dass die Aussage von Satz 10.12 für unendlich-dimensionale Vektorräume nicht gelten muss.)

- (3) Die Berechnung des *bestimmten Integrals* von a bis b ist linear:

$$I_{a,b}: P \longrightarrow \mathbb{R}, \quad f \longmapsto \int_a^b f(x) dx.$$

Für die Potenzfunktionen f_n gilt $I_{a,b}(f_n) = \frac{b^{n+1} - a^{n+1}}{n+1}$.

- (4) Auch die *unbestimmte Integration* mit Anfangspunkt $a \in \mathbb{R}$ ist linear:

$$I_a: P \longrightarrow P, \quad f \longmapsto \left(x \mapsto \int_a^x f(t) dt \right).$$

Das ist die lineare Abbildung mit $I_a(f_n) = \left(x \mapsto \frac{1}{n+1}(x^{n+1} - a^{n+1}) \right)$. Die Abbildung I_a ist injektiv, aber nicht surjektiv: ihr Bild ist gerade der Kern von ev_a (die Integralfunktionen verschwinden alle an der Stelle a).

(5) Die *Translation* (also Verschiebung) um $a \in \mathbb{R}$ ist linear:

$$T_a: P \longrightarrow P, \quad f \longmapsto (x \mapsto f(x-a)).$$

T_a ist ein Automorphismus von P , der inverse Automorphismus ist T_{-a} .

Zwischen diesen Abbildungen bestehen eine Reihe von Relationen, wie zum Beispiel

$$\begin{aligned} \text{ev}_b \circ I_a &= I_{a,b}, & D \circ I_a &= \text{id}_P, & (I_a \circ D)(f) &= f - \text{ev}_a(f)f_0, \\ T_a \circ D &= D \circ T_a, & T_a \circ T_b &= T_{a+b}, & I_a \circ T_b &= T_b \circ I_{a-b}, \\ I_{a,b} \circ T_c &= I_{a-c,b-c}, & \text{ev}_a \circ T_b &= \text{ev}_{a-b}. \end{aligned}$$

Man kann sie leicht auf den Potenzfunktionen nachprüfen (das genügt, weil die Potenzfunktionen eine Basis von P sind).

In der Analysis lernen Sie, dass Differentiation und Integration ganz allgemein lineare Abbildungen sind. ♣

Kern und Bild einer linearen Abbildung sind wichtige Daten. Für die Dimension des Bildes gibt es sogar einen eigenen Namen.

* 10.16. **Definition.** Ist $\phi: V \rightarrow W$ eine lineare Abbildung, dann heißt

$$\text{rk}(\phi) = \dim \text{im}(\phi)$$

der *Rang* von ϕ . ◇

DEF
Rang einer
linearen Abb.

Zwischen dem Rang und der Dimension des Kerns besteht ein einfacher Zusammenhang. (Siehe Abbildung 2 für eine Veranschaulichung.)

* 10.17. **Satz.** Sei $\phi: V \rightarrow W$ eine lineare Abbildung. Dann gilt

$$\dim \ker(\phi) + \text{rk}(\phi) = \dim V.$$

SATZ
Rangsatz

Dabei sei $n + \infty = \infty + n = \infty + \infty = \infty$ für $n \in \mathbb{N}$.

Beweis. Ist $\dim \ker(\phi) = \infty$, dann muss auch $\dim V = \infty$ sein, denn $\ker(\phi)$ ist ein Untervektorraum von V (siehe Satz 9.19). Also ist die Behauptung in diesem Fall richtig. Ist $\text{rk}(\phi) = \infty$, dann können wir unendlich viele linear unabhängige Vektoren $w_j \in \text{im}(\phi)$ finden ($j \in \mathbb{N}$). Sei $v_j \in V$ ein Urbild von w_j ; dann sind auch die v_j linear unabhängig. Denn sei $\lambda_0 v_0 + \lambda_1 v_1 + \dots + \lambda_n v_n = \mathbf{0}$, dann folgt durch Anwenden von ϕ auch

$$\begin{aligned} \lambda_0 w_0 + \lambda_1 w_1 + \dots + \lambda_n w_n &= \lambda_0 \phi(v_0) + \lambda_1 \phi(v_1) + \dots + \lambda_n \phi(v_n) \\ &= \phi(\lambda_0 v_0 + \lambda_1 v_1 + \dots + \lambda_n v_n) = \mathbf{0}. \end{aligned}$$

Weil $w_0, w_1, \dots, w_n$ linear unabhängig sind, müssen alle Koeffizienten λ_j null sein, was zeigt, dass $v_0, v_1, \dots, v_n$ linear unabhängig sind. Es gibt also unendlich viele linear unabhängige Vektoren in V ; damit ist $\dim V = \infty$ und die Behauptung des Satzes stimmt. Wir können also jetzt annehmen, dass $\dim \ker(\phi)$ und $\text{rk}(\phi)$ beide endlich sind.

Seien $k = \dim \ker(\phi)$, $r = \text{rk}(\phi)$. Wir können eine Basis $(b_1, \dots, b_k)$ von $\ker(\phi)$ und eine Basis $(b'_1, \dots, b'_r)$ vom $\text{im}(\phi)$ wählen. Es gibt dann Vektoren $b_{k+1}, \dots, b_{k+r}$ mit $\phi(b_{k+1}) = b'_1$, $\phi(b_{k+2}) = b'_2$, $\dots$, $\phi(b_{k+r}) = b'_r$. Ich behaupte jetzt, dass $(b_1, \dots, b_k, b_{k+1}, \dots, b_{k+r})$ eine Basis von V ist. Daraus folgt $k + r = \dim V$, also die Behauptung des Satzes.

- Erzeugendensystem:

Sei $v \in V$. Da $(b'_1, \dots, b'_r)$ eine Basis von $\text{im}(\phi)$ ist, gibt es Skalare $\mu_1, \dots, \mu_r$ mit $\phi(v) = \mu_1 b'_1 + \dots + \mu_r b'_r$. Dann ist

$$\begin{aligned}\phi(v - (\mu_1 b_{k+1} + \dots + \mu_r b_{k+r})) &= \phi(v) - (\mu_1 \phi(b_{k+1}) + \dots + \mu_r \phi(b_{k+r})) \\ &= \phi(v) - (\mu_1 b'_1 + \dots + \mu_r b'_r) = \mathbf{0},\end{aligned}$$

also ist $v - (\mu_1 b_{k+1} + \dots + \mu_r b_{k+r}) \in \ker(\phi)$. Es gibt also Skalare $\lambda_1, \dots, \lambda_k$ mit

$$v - (\mu_1 b_{k+1} + \dots + \mu_r b_{k+r}) = \lambda_1 b_1 + \dots + \lambda_k b_k.$$

Damit ist

$$v = \lambda_1 b_1 + \dots + \lambda_k b_k + \mu_1 b_{k+1} + \dots + \mu_r b_{k+r}$$

eine Linearkombination von $(b_1, \dots, b_{k+r})$.

- linear unabhängig:

Seien $\lambda_1, \dots, \lambda_{k+r}$ Skalare mit

$$\lambda_1 b_1 + \dots + \lambda_k b_k + \lambda_{k+1} b_{k+1} + \dots + \lambda_{k+r} b_{k+r} = \mathbf{0}.$$

Dann ist

$$\begin{aligned}\mathbf{0} &= \phi(\lambda_1 b_1 + \dots + \lambda_k b_k + \lambda_{k+1} b_{k+1} + \dots + \lambda_{k+r} b_{k+r}) \\ &= \lambda_1 \phi(b_1) + \dots + \lambda_k \phi(b_k) + \lambda_{k+1} \phi(b_{k+1}) + \dots + \lambda_{k+r} \phi(b_{k+r}) \\ &= \lambda_{k+1} \phi(b_{k+1}) + \dots + \lambda_{k+r} \phi(b_{k+r}) \\ &= \lambda_{k+1} b'_1 + \dots + \lambda_{k+r} b'_r,\end{aligned}$$

denn $\phi(b_1) = \dots = \phi(b_k) = \mathbf{0}$. Da $(b'_1, \dots, b'_r)$ linear unabhängig ist, muss $\lambda_{k+1} = \dots = \lambda_{k+r} = 0$ gelten. Eingesetzt in die ursprüngliche Relation liefert das

$$\lambda_1 b_1 + \dots + \lambda_k b_k = \mathbf{0}.$$

Weil auch $(b_1, \dots, b_k)$ linear unabhängig ist, folgt daraus $\lambda_1 = \dots = \lambda_k = 0$, also war auch unsere ursprüngliche Linearkombination trivial. $\square$

10.18. Beispiel. Seien $V = K^n$ mit $n \geq 1$ und

$$\phi: V \longrightarrow K, \quad (x_1, x_2, \dots, x_n) \longmapsto x_1 + x_2 + \dots + x_n.$$

BSP
Dimension
eines Kerns

Dann ist ϕ linear (leicht nachzurechnen). Damit ist

$$U = \{(x_1, x_2, \dots, x_n) \in V \mid x_1 + x_2 + \dots + x_n = 0\} = \ker(\phi) \subset V$$

ein Untervektorraum von V . Es gilt $\text{im}(\phi) = K$:

$$\text{für } \lambda \in K \text{ ist } \phi((\lambda, 0, \dots, 0)) = \lambda.$$

Damit ist $\text{rk}(\phi) = \dim_K K = 1$. Es folgt

$$\dim U = \dim \ker(\phi) = \dim V - \text{rk}(\phi) = n - 1.$$



In vielen Fällen ist es einfacher, den Kern und seine Dimension direkt zu bestimmen als den Rang. Mit Satz 10.17 kann man daraus dann den Rang berechnen.

Die Konstruktion des Vektorraums $K^X = \text{Abb}(X, K)$ lässt sich verallgemeinern.

10.19. Definition. Seien K ein Körper und V ein K -Vektorraum. Sei weiter X eine Menge. Dann können wir auf $V^X = \text{Abb}(X, V)$ eine Struktur als K -Vektorraum definieren durch

DEF
Vektorraum
 V^X

$$f + g: x \mapsto f(x) + g(x) \quad \text{und} \quad \lambda f: x \mapsto \lambda f(x).$$

Der Beweis ist analog zu dem für K^X .

Für $X = \{1, 2, \dots, n\}$ identifizieren wir V^X mit V^n . ◇

Man kann diese Konstruktion noch weiter verallgemeinern. Ist $(V_i)_{i \in I}$ eine beliebige Familie von K -Vektorräumen, dann wird ihr kartesisches Produkt

$$\prod_{i \in I} V_i = \{(v_i)_{i \in I} \mid \forall i \in I: v_i \in V_i\}$$

(das verallgemeinert das kartesische Produkt von endlich vielen Mengen) in der gleichen Weise (also komponentenweise) zu einem K -Vektorraum.

Wir können also insbesondere zwei *lineare* Abbildungen $V \rightarrow W$ addieren oder eine solche Abbildung mit einem Skalar multiplizieren (da wir das sogar für beliebige Abbildungen $V \rightarrow W$ können).

10.20. Satz. Seien V und W zwei K -Vektorräume. Die Menge der linearen Abbildungen $V \rightarrow W$ bildet einen K -Untervektorraum von $\text{Abb}(V, W)$.

SATZ
Vektorraum
der lin. Abb.

Beweis. Wir müssen die Bedingungen für einen Untervektorraum nachprüfen.

- Die Nullabbildung ist linear.
- Seien $\phi, \psi: V \rightarrow W$ linear. Wir müssen zeigen, dass $\phi + \psi$ ebenfalls linear ist. Seien dazu $v, v' \in V$, $\lambda \in K$. Dann haben wir

$$\begin{aligned} (\phi + \psi)(v + v') &= \phi(v + v') + \psi(v + v') = \phi(v) + \phi(v') + \psi(v) + \psi(v') \\ &= \phi(v) + \psi(v) + \phi(v') + \psi(v') = (\phi + \psi)(v) + (\phi + \psi)(v') \end{aligned}$$

und

$$\begin{aligned} (\phi + \psi)(\lambda v) &= \phi(\lambda v) + \psi(\lambda v) = \lambda \phi(v) + \lambda \psi(v) \\ &= \lambda(\phi(v) + \psi(v)) = \lambda(\phi + \psi)(v). \end{aligned}$$

- Sei $\phi: V \rightarrow W$ linear und $\lambda \in K$. Wir müssen zeigen, dass $\lambda \phi$ ebenfalls linear ist. Seien dazu $v, v' \in V$, $\mu \in K$. Dann haben wir

$$\begin{aligned} (\lambda \phi)(v + v') &= \lambda \phi(v + v') = \lambda(\phi(v) + \phi(v')) \\ &= \lambda \phi(v) + \lambda \phi(v') = (\lambda \phi)(v) + (\lambda \phi)(v') \end{aligned}$$

und

$$\begin{aligned} (\lambda \phi)(\mu v) &= \lambda \phi(\mu v) = \lambda \cdot \mu \phi(v) \\ &= \mu \cdot \lambda \phi(v) = \mu(\lambda \phi)(v). \end{aligned}$$

□

10.21. Definition. Der Vektorraum der linearen Abbildungen $V \rightarrow W$ wird mit $\text{Hom}(V, W)$ (oder $\text{Hom}_K(V, W)$) bezeichnet. Im Fall $V = W$ schreiben wir auch $\text{End}(V) = \text{Hom}(V, V)$ (oder $\text{End}_K(V)$) für den Vektorraum der Endomorphismen von V . ◇

DEF
 $\text{Hom}(V, W)$
 $\text{End}(V)$

10.22. Satz. Seien V und W zwei K -Vektorräume mit $\dim V = n < \infty$. Sei weiter $(b_1, b_2, \dots, b_n)$ eine Basis von V . Dann ist

$$\Phi: \operatorname{Hom}(V, W) \longrightarrow W^n, \quad \phi \longmapsto (\phi(b_1), \phi(b_2), \dots, \phi(b_n))$$

ein Isomorphismus. Insbesondere ist im Fall von $\dim W = m < \infty$

$$\dim \operatorname{Hom}(V, W) = \dim W^n = n \dim W = mn = (\dim V)(\dim W).$$

SATZ
 $\operatorname{Hom}(V, W)$
 $\cong W^{\dim V}$

Beweis. Es ist klar, dass Φ linear ist (denn Φ setzt sich aus Auswertungsabbildungen zusammen; die Auswertungsabbildungen $\operatorname{Abb}(V, W) \rightarrow W$, $\phi \mapsto \phi(v)$, sind auch in diesem allgemeineren Kontext linear; man sieht das wie in Beispiel 10.10). Nach Satz 10.11 gibt es zu jeder Wahl der Bilder von $b_1, \dots, b_n$ in W genau eine lineare Abbildung; das bedeutet, dass Φ bijektiv ist. Isomorphe Vektorräume haben dieselbe Dimension; der Beweis von $\dim W^n = n \dim W$ ist eine Übungsaufgabe. $\square$

10.23. Folgerung. Seien V und W zwei K -Vektorräume; sei $(b_1, b_2, \dots, b_n)$ eine Basis von V und $(b'_1, b'_2, \dots, b'_m)$ eine Basis von W .

Für $(i, j) \in \{1, 2, \dots, m\} \times \{1, 2, \dots, n\}$ sei $\phi_{ij}: V \rightarrow W$ die lineare Abbildung mit $\phi_{ij}(b_k) = \mathbf{0}$ für $k \neq j$ und $\phi_{ij}(b_j) = b'_i$. Dann ist $(\phi_{ij})_{(i,j) \in \{1, 2, \dots, m\} \times \{1, 2, \dots, n\}}$ eine Basis von $\operatorname{Hom}(V, W)$.

FOLG
 Basis von
 $\operatorname{Hom}(V, W)$

Beweis. Nach Satz 10.11 existieren eindeutig bestimmte ϕ_{ij} wie angegeben. Wir zeigen, dass die $\phi_{ij} \in \operatorname{Hom}(V, W)$ linear unabhängig sind. Seien dazu λ_{ij} Skalare mit

$$\sum_{i=1}^m \sum_{j=1}^n \lambda_{ij} \phi_{ij} = \mathbf{0}.$$

Sei $k \in \{1, 2, \dots, n\}$. Einsetzen von b_k liefert dann

$$\mathbf{0} = \left(\sum_{i=1}^m \sum_{j=1}^n \lambda_{ij} \phi_{ij} \right) (b_k) = \sum_{i=1}^m \sum_{j=1}^n \lambda_{ij} \phi_{ij}(b_k) = \sum_{i=1}^m \lambda_{ik} b'_i$$

Im letzten Schritt haben wir verwendet, dass $\phi_{ij}(b_k) = \mathbf{0}$ ist für $k \neq j$; in der inneren Summe bleibt dann nur $\lambda_{ik} \phi_{ik}(b_k) = \lambda_{ik} b'_i$ stehen. Da die b'_i linear unabhängig sind, folgt $\lambda_{ik} = 0$ für alle i . Da k beliebig war, sind also alle $\lambda_{ij} = 0$, was zu zeigen war. Nach Satz 10.22 ist $\dim \operatorname{Hom}(V, W) = nm$ gleich der Anzahl der linear unabhängigen Elemente $\phi_{ij} \in \operatorname{Hom}(V, W)$, nach Satz 9.14 sind die ϕ_{ij} dann bereits eine Basis von $\operatorname{Hom}(V, W)$. $\square$

Im Fall $V = K^n$, $W = K^m$ mit den Standardbasen kann man das, was ϕ_{ij} bewirkt, so beschreiben: Man nimmt die j -te Komponente von $(x_1, x_2, \dots, x_n) \in K^n$ und steckt sie in die i -te Komponente des Resultats in K^m ; die übrigen Komponenten sind null.

10.24. Beispiel. Als einfaches Beispiel betrachten wir $V = \mathbb{R}^3$, $W = \mathbb{R}^2$, jeweils mit der Standardbasis $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$ bzw. $(\mathbf{e}'_1, \mathbf{e}'_2)$. Die Basis von $\text{Hom}(\mathbb{R}^3, \mathbb{R}^2)$ aus Folgerung 10.23 sieht in diesem Fall so aus:

BSP
Basis von
 $\text{Hom}(\mathbb{R}^3, \mathbb{R}^2)$

$$\begin{aligned}\phi_{11}: (x, y, z) &\longmapsto (x, 0) \\ \phi_{12}: (x, y, z) &\longmapsto (y, 0) \\ \phi_{13}: (x, y, z) &\longmapsto (z, 0) \\ \phi_{21}: (x, y, z) &\longmapsto (0, x) \\ \phi_{22}: (x, y, z) &\longmapsto (0, y) \\ \phi_{23}: (x, y, z) &\longmapsto (0, z)\end{aligned}$$

Jede lineare Abbildung $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ lässt sich als Linearkombination dieser sechs Abbildungen schreiben; es gibt also $a, b, c, d, e, f \in \mathbb{R}$, sodass

$$\phi = a\phi_{11} + b\phi_{12} + c\phi_{13} + d\phi_{21} + e\phi_{22} + f\phi_{23}$$

ist; also gilt

$$\phi(x, y, z) = (ax + by + cz, dx + ey + fz).$$



Die Endomorphismen eines Vektorraums V bilden sogar einen Ring, den *Endomorphismenring* von V :

DEF
Endomorphi-
menring
SATZ
 $\text{End}(V)$ ist
ein Ring

10.25. Satz. Sei V ein K -Vektorraum. Dann ist $\text{End}(V)$ ein Ring mit der Addition des K -Vektorraums $\text{End}(V) = \text{Hom}(V, V)$ und der Komposition von Abbildungen als Multiplikation; das Einselement ist die identische Abbildung id_V .

Beweis. Die Vektorraum-Axiome, die in $\text{End}(V)$ gelten, liefern uns die Ring-Axiome für die Addition. Es bleibt noch zu zeigen, dass die Multiplikation assoziativ ist mit Einselement id_V und dass die beiden Ring-Distributivgesetze gelten. Seien also $f, g, h \in \text{End}(V)$. Die Assoziativität $(f \circ g) \circ h = f \circ (g \circ h)$ gilt für Abbildungen ganz allgemein, ebenso wie $\text{id}_V \circ f = f = f \circ \text{id}_V$. Zum Nachweis der Distributivgesetze rechnen wir für $v \in V$:

$$\begin{aligned}((f + g) \circ h)(v) &= (f + g)(h(v)) = f(h(v)) + g(h(v)) \\ &= (f \circ h)(v) + (g \circ h)(v) = (f \circ h + g \circ h)(v),\end{aligned}$$

also ist $(f + g) \circ h = f \circ h + g \circ h$, und

$$\begin{aligned}(f \circ (g + h))(v) &= f((g + h)(v)) = f(g(v) + h(v)) \\ &= f(g(v)) + f(h(v)) = (f \circ g)(v) + (f \circ h)(v) \\ &= (f \circ g + f \circ h)(v),\end{aligned}$$

also ist $f \circ (g + h) = f \circ g + f \circ h$ (dabei haben wir verwendet, dass f linear ist). $\square$

Der Endomorphismenring ist nicht kommutativ, wenn $\dim V \geq 2$ ist (Übung!). Für $\dim V = 1$ ist $\text{End}(V) = K$, da alle Endomorphismen durch Multiplikation mit Skalaren gegeben sind; für $\dim V = 0$ ist $\text{End}(V)$ der Nullring.

Die Automorphismen von V bilden eine Gruppe, die *Automorphismengruppe* $\text{Aut}(V)$ von V (das ist auch die Gruppe der invertierbaren Elemente des Rings $\text{End}(V)$).

11. MATRIZEN

Die Ergebnisse des letzten Abschnitts zeigen uns, dass wir lineare Abbildungen zwischen zwei endlich-dimensionalen K -Vektorräumen V und W der Dimensionen n und m durch mn Koeffizienten aus K beschreiben können. Dazu müssen wir Basen von V und W wählen; daraus bekommen wir eine Basis von $\text{Hom}(V, W)$ wie in Folgerung 10.23 und die gesuchten Koeffizienten sind dann die Koeffizienten in der Darstellung der gegebenen linearen Abbildung als Linearkombination bezüglich dieser Basis. Für diese Koeffizienten führt man eine spezielle Form der Darstellung ein.

11.1. Definition. Sei K ein Körper und seien $m, n \in \mathbb{N}$. Eine $m \times n$ -Matrix mit Einträgen aus K (oder kurz über K) ist ein rechteckiges Schema aus mn Elementen von K , das wie folgt notiert wird:

DEF
Matrix

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

Zur Abkürzung schreiben wir auch $(a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ (oder auch $(a_{ij})_{i,j}$, falls die Zahlen m und n aus dem Kontext klar sind) für diese Matrix. Im Fall $m = n$ heißt die Matrix *quadratisch*. Für $i \in \{1, 2, \dots, m\}$ heißt das n -Tupel $(a_{i1}, a_{i2}, \dots, a_{in})$ die i -te Zeile der Matrix, für $j \in \{1, 2, \dots, n\}$ heißt das m -Tupel $(a_{1j}, a_{2j}, \dots, a_{mj})$ die j -te Spalte der Matrix.

Wir schreiben $\text{Mat}(m \times n, K)$ für die Menge aller $m \times n$ -Matrizen mit Einträgen aus K ; im Fall $m = n$ auch kürzer $\text{Mat}(n, K)$ für $\text{Mat}(n \times n, K)$. $\diamond$

Ebenso gebräuchlich ist die Notation $K^{m \times n}$. Matrizen werden auch (insbesondere in der englischsprachigen Literatur) mit eckigen Klammern notiert:

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}.$$

Im Grunde ist eine $m \times n$ -Matrix über K nichts anderes als eine Familie von Elementen von K mit der Indexmenge $\{1, 2, \dots, m\} \times \{1, 2, \dots, n\}$, also

$$\text{Mat}(m \times n, K) = K^{\{1, 2, \dots, m\} \times \{1, 2, \dots, n\}}.$$

Da wir auf beliebigen Mengen der Form K^I eine Struktur als K -Vektorraum definiert haben, folgt sofort:

11.2. Lemma. Seien K ein Körper und $m, n \in \mathbb{N}$. Die Menge $\text{Mat}(m \times n, K)$ mit komponentenweise definierter Addition und Skalarmultiplikation ist ein K -Vektorraum der Dimension mn .

LEMMA
Vektorraum
der $m \times n$ -
Matrizen

Ist $m = 0$ oder $n = 0$ (oder beides), dann ist $\text{Mat}(m \times n, K)$ ein Null-Vektorraum; sein einziges Element ist eine leere Matrix (mit null Zeilen und n Spalten oder mit m Zeilen und null Spalten).

Matrizen (mit der gleichen Anzahl an Zeilen und Spalten) werden also wie folgt addiert und mit Skalaren multipliziert:

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} + \begin{pmatrix} b_{11} & b_{12} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2n} \\ \vdots & \vdots & & \vdots \\ b_{m1} & b_{m2} & \cdots & b_{mn} \end{pmatrix} = \begin{pmatrix} a_{11} + b_{11} & a_{12} + b_{12} & \cdots & a_{1n} + b_{1n} \\ a_{21} + b_{21} & a_{22} + b_{22} & \cdots & a_{2n} + b_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} + b_{m1} & a_{m2} + b_{m2} & \cdots & a_{mn} + b_{mn} \end{pmatrix}$$

und

$$\lambda \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} = \begin{pmatrix} \lambda a_{11} & \lambda a_{12} & \cdots & \lambda a_{1n} \\ \lambda a_{21} & \lambda a_{22} & \cdots & \lambda a_{2n} \\ \vdots & \vdots & & \vdots \\ \lambda a_{m1} & \lambda a_{m2} & \cdots & \lambda a_{mn} \end{pmatrix}.$$

Wie zu Beginn dieses Abschnitts beschrieben, können wir linearen Abbildungen Matrizen zuordnen. Wir betrachten zunächst $V = K^n$ und $W = K^m$ mit den Standardbasen $B = (\mathbf{e}_1, \dots, \mathbf{e}_n)$ von V und $B' = (\mathbf{e}'_1, \dots, \mathbf{e}'_m)$ von W (wir schreiben hier $\mathbf{e}'_i$ für den i -ten Standard-Basisvektor in K^m zur Unterscheidung von den Basisvektoren $\mathbf{e}_j$ in K^n). Wir haben dann die Basis $(\phi_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ von $\text{Hom}(K^n, K^m)$ wie in Folgerung 10.23 mit $\phi_{ij}(\mathbf{e}_k) = \mathbf{0}$ für $k \neq j$ und $\phi_{ij}(\mathbf{e}_j) = \mathbf{e}'_i$. Ist $\phi: K^n \rightarrow K^m$ eine lineare Abbildung, dann schreiben wir ϕ als Linearkombination

$$\phi = \sum_{i=1}^m \sum_{j=1}^n a_{ij} \phi_{ij} \quad \text{mit } a_{ij} \in K.$$

Die zugehörige Matrix ist dann $A = (a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$.

11.3. Beispiel. Wie wir gesehen haben, hat eine lineare Abbildung $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ die Form $\phi(x, y, z) = (ax + by + cz, dx + ey + fz)$ mit geeigneten $a, b, c, d, e, f \in \mathbb{R}$. Dann ist $\phi = a\phi_{11} + b\phi_{12} + c\phi_{13} + d\phi_{21} + e\phi_{22} + f\phi_{23}$ (vergleiche Beispiel 10.24), also ist die zugehörige Matrix

$$A = \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix}.$$



Die j -te Spalte der zu $\phi: K^n \rightarrow K^m$ gehörigen $m \times n$ -Matrix enthält gerade die Koeffizienten des Bildes des j -ten Standard-Basisvektors $\mathbf{e}_j$, denn

$$\phi(\mathbf{e}_j) = \sum_{i=1}^m \sum_{k=1}^n a_{ik} \phi_{ik}(\mathbf{e}_j) = \sum_{i=1}^m a_{ij} \mathbf{e}'_i = (a_{1j}, a_{2j}, \dots, a_{mj})$$

ähnlich wie im Beweis von Folgerung 10.23.

11.4. Lemma. Die oben beschriebene Zuordnung definiert einen Isomorphismus $\text{Hom}(K^n, K^m) \rightarrow \text{Mat}(m \times n, K)$. Wenn man $\text{Mat}(m \times n, K)$ mit $K^{\{1, \dots, m\} \times \{1, \dots, n\}}$ identifiziert, dann ist dieser Isomorphismus invers zu der Linearkombinationenabbildung $K^{\{1, 2, \dots, m\} \times \{1, 2, \dots, n\}} \rightarrow \text{Hom}(K^n, K^m)$, die zur Basis $(\phi_{ij})_{(i,j) \in \{1, \dots, m\} \times \{1, \dots, n\}}$ von $\text{Hom}(K^n, K^m)$ gehört.

BSP
Matrix für
 $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^2$

LEMMA
 $\text{Mat}(m \times n, K)$
 $\cong$
 $\text{Hom}(K^n, K^m)$

Beweis. Die erwähnte Linearkombinationenabbildung

$$\Phi: \text{Mat}(m \times n, K) = K^{\{1,2,\dots,m\} \times \{1,2,\dots,n\}} \rightarrow \text{Hom}(K^n, K^m)$$

bildet eine Matrix $(a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ auf die Linearkombination $\sum_{i,j} a_{ij} \phi_{ij}$ ab; sie ist ein Isomorphismus, da $(\phi_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ eine Basis von $\text{Hom}(K^n, K^m)$ ist (siehe Beispiel 10.3). Die Abbildung, die einer linearen Abbildung $\phi: K^n \rightarrow K^m$ ihre Matrix zuordnet, ist offenbar die Umkehrabbildung von Φ , insbesondere also ebenfalls ein Isomorphismus. $\square$

Wie stellt sich die Anwendung der linearen Abbildung $\phi: K^n \rightarrow K^m$ dar, wenn wir die zugehörige Matrix $A = (a_{ij})_{i,j}$ verwenden? Es gilt

$$\phi(x_1, x_2, \dots, x_n) = \sum_{i=1}^m \sum_{j=1}^n a_{ij} \phi_{ij}(x_1, x_2, \dots, x_n) = \sum_{i=1}^m \sum_{j=1}^n a_{ij} x_j \mathbf{e}'_i,$$

also ist die i -te Komponente von $\phi(x_1, x_2, \dots, x_n)$ gegeben durch

$$\sum_{j=1}^n a_{ij} x_j = a_{i1} x_1 + a_{i2} x_2 + \dots + a_{in} x_n.$$

Man schreibt das dann gerne als Multiplikation der Matrix A mit dem $(x_1, \dots, x_n)$ entsprechenden *Spaltenvektor*: Man identifiziert also K^n mit $\text{Mat}(n \times 1, K)$ und K^m mit $\text{Mat}(m \times 1, K)$. Dann haben wir für das Resultat der Anwendung von ϕ :

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix}.$$

Das Ergebnis ist wieder ein Spaltenvektor, diesmal der Länge m . Seine i -te Komponente ergibt sich aus der i -ten Zeile der Matrix und dem Spaltenvektor zu $(x_1, \dots, x_n)$ als das *Skalarprodukt*

$$a_{i1}x_1 + a_{i2}x_2 + \dots + a_{in}x_n.$$

DEF
Skalarprodukt

(Das Skalarprodukt heißt so, weil sein Wert ein Skalar ist:

„Vektor mal Vektor = Skalar“.

Man beachte den Unterschied zur Skalarmultiplikation

„Skalar mal Vektor = Vektor“!)



11.5. Beispiele. 2×3 -Matrizen mit Einträgen in $\mathbb{R}$ entsprechen linearen Abbildungen $\mathbb{R}^3 \rightarrow \mathbb{R}^2$. In diesem Fall sieht obige Formel so aus:

BSP

$$\begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} ax + by + cz \\ dx + ey + fz \end{pmatrix}$$

3×2 -Matrizen über $\mathbb{R}$ entsprechen linearen Abbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^3$. Dann haben wir:

$$\begin{pmatrix} a & b \\ c & d \\ e & f \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} ax + by \\ cx + dy \\ ex + fy \end{pmatrix}$$



Die Verknüpfung von linearen Abbildungen entspricht der Multiplikation von Matrizen.

11.6. Definition. Sei K ein Körper, seien weiter $l, m, n \in \mathbb{N}$. Für Matrizen $A \in \text{Mat}(l \times m, K)$, $B \in \text{Mat}(m \times n, K)$ ist das *Produkt* $A \cdot B \in \text{Mat}(l \times n, K)$ definiert als die zu $f \circ g$ gehörende Matrix, wobei $f: K^m \rightarrow K^l$ und $g: K^n \rightarrow K^m$ die den Matrizen A und B entsprechenden linearen Abbildungen sind. Wie üblich schreibt man auch AB für $A \cdot B$. $\diamond$

DEF
Matrix-
multiplikation

So wie man Abbildungen nur dann miteinander verknüpfen kann, wenn der Wertebereich der einen Abbildung mit dem Definitionsbereich der anderen übereinstimmt, kann man Matrizen nur dann miteinander multiplizieren, wenn sie in der Größe „zueinander passen“, wenn also die Spaltenanzahl des linken Faktors gleich der Zeilenanzahl des rechten Faktors ist.

Wie sieht diese Matrixmultiplikation konkret aus? Seien $A = (a_{ij})_{1 \leq i \leq l, 1 \leq j \leq m}$, $B = (b_{jk})_{1 \leq j \leq m, 1 \leq k \leq n}$ und $C = (c_{ik})_{1 \leq i \leq l, 1 \leq k \leq n} = AB$. Dann sollte c_{ik} die i -te Komponente von $f(g(\mathbf{e}_k))$ sein. Es ist

$f(g(\mathbf{e}_k)) = f(b_{1k}\mathbf{e}'_1 + b_{2k}\mathbf{e}'_2 + \dots + b_{mk}\mathbf{e}'_m) = b_{1k}f(\mathbf{e}'_1) + b_{2k}f(\mathbf{e}'_2) + \dots + b_{mk}f(\mathbf{e}'_m)$
und die i -te Komponente von $f(\mathbf{e}'_j)$ ist a_{ij} . Also ist

$$c_{ik} = a_{i1}b_{1k} + a_{i2}b_{2k} + \dots + a_{im}b_{mk} = \sum_{j=1}^m a_{ij}b_{jk}$$

das **Skalarprodukt der i -ten Zeile von A mit der k -ten Spalte von B :**

Die (i, k) -Komponente von AB ist „ i -te Zeile von A mal k -te Spalte von B “.

Die oben eingeführte Multiplikation „Matrix mal Spaltenvektor“ ist dann also ein Spezialfall dieser allgemeinen Matrixmultiplikation.

11.7. Beispiel. Wir berechnen das Produkt zweier Matrizen über $\mathbb{R}$:

BSP

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix} = \begin{pmatrix} 1 \cdot 1 + 2 \cdot 3 + 3 \cdot 5 & 1 \cdot 2 + 2 \cdot 4 + 3 \cdot 6 \\ 4 \cdot 1 + 5 \cdot 3 + 6 \cdot 5 & 4 \cdot 2 + 5 \cdot 4 + 6 \cdot 6 \end{pmatrix} = \begin{pmatrix} 22 & 28 \\ 49 & 64 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} = \begin{pmatrix} 1 \cdot 1 + 2 \cdot 4 & 1 \cdot 2 + 2 \cdot 5 & 1 \cdot 3 + 2 \cdot 6 \\ 3 \cdot 1 + 4 \cdot 4 & 3 \cdot 2 + 4 \cdot 5 & 3 \cdot 3 + 4 \cdot 6 \\ 5 \cdot 1 + 6 \cdot 4 & 5 \cdot 2 + 6 \cdot 5 & 5 \cdot 3 + 6 \cdot 6 \end{pmatrix} = \begin{pmatrix} 9 & 12 & 15 \\ 19 & 26 & 33 \\ 29 & 40 & 51 \end{pmatrix} \clubsuit$$

Zur identischen Abbildung gehört eine spezielle Matrix.

11.8. Definition. Seien K ein Körper und $n \in \mathbb{N}$. Die Matrix $I_n \in \text{Mat}(n, K)$, die der identischen Abbildung id_{K^n} entspricht, heißt die *Einheitsmatrix* (der Größe n über K). $\diamond$

DEF
Einheits-
matrix

In der j -ten Spalte von I_n muss der j -te Standard-Basisvektor stehen, also sieht I_n so aus:

$$I_n = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Man schreibt das auch $I_n = (\delta_{ij})_{1 \leq i, j \leq n}$ mit dem *Kronecker-Delta*

DEF
Kronecker-
Delta

$$\delta_{ij} = \begin{cases} 1 & \text{falls } i = j, \\ 0 & \text{falls } i \neq j. \end{cases}$$

11.9. Lemma. Sei K ein Körper. Die Matrixmultiplikation ist assoziativ und hat die Einheitsmatrix als neutrales Element; sie erfüllt die Distributivgesetze bezüglich der Matrixaddition:

LEMMA
Eigensch.
Matrixmult.

- (1) Für alle $A \in \text{Mat}(k \times l, K)$, $B \in \text{Mat}(l \times m, K)$, $C \in \text{Mat}(m \times n, K)$ gilt $(AB)C = A(BC)$.
- (2) Für alle $A \in \text{Mat}(m \times n, K)$ gilt $I_m A = A = A I_n$.
- (3) Für alle $A \in \text{Mat}(l \times m, K)$ und $B, C \in \text{Mat}(m \times n, K)$ gilt $A(B + C) = AB + AC$.
- (4) Für alle $A, B \in \text{Mat}(l \times m, K)$ und $C \in \text{Mat}(m \times n, K)$ gilt $(A + B)C = AC + BC$.
- (5) Außerdem gilt für alle $\lambda \in K$ und $A \in \text{Mat}(l \times m, K)$, $B \in \text{Mat}(m \times n, K)$: $\lambda(AB) = (\lambda A)B = A(\lambda B)$.

Insbesondere ist $\text{Mat}(n, K)$ mit der Matrixaddition und Matrixmultiplikation als Verknüpfungen ein Ring.

Beweis. Das ist eine unmittelbare Übersetzung der entsprechenden Aussagen für lineare Abbildungen, vergleiche den Beweis von Satz 10.25 (die Beweise etwa für die Distributivgesetze funktionieren auch in der etwas allgemeineren Situation, die hier vorliegt). Alternativ kann man das auch leicht direkt nachrechnen. $\square$

11.10. Definition. Der Ring $\text{Mat}(n, K)$ heißt der *Matrizenring* (der Größe n über K). Eine Matrix $A \in \text{Mat}(n, K)$ heißt *invertierbar*, wenn es eine Matrix $B \in \text{Mat}(n, K)$ gibt mit $AB = I_n$. Dann gilt auch $BA = I_n$; wir schreiben A^{-1} für B und nennen B die *Inverse* von A . $\diamond$

DEF
Matrizen-
ring
invertierbare
Matrix

Für die zu A und B gehörenden linearen Abbildungen $f, g: K^n \rightarrow K^n$ bedeutet $AB = I_n$, dass $f \circ g = \text{id}_{K^n}$ ist. Dann ist f surjektiv, also ein Isomorphismus (siehe Folgerung 10.14) und $g = f^{-1}$, also ist auch $g \circ f = \text{id}_{K^n}$, d.h., $BA = I_n$. Die Matrix $B = A^{-1}$ ist also die zu f^{-1} gehörende Matrix.

11.11. Beispiel. Die Matrix $A = \begin{pmatrix} 1 & \lambda \\ 0 & 1 \end{pmatrix} \in \text{Mat}(2, K)$ (mit $\lambda \in K$ beliebig) ist invertierbar, denn

BSP

$$\begin{pmatrix} 1 & \lambda \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & -\lambda \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

$\clubsuit$

Im nächsten Abschnitt werden wir lernen, wie wir Basen von Kern und Bild einer linearen Abbildung $f: K^n \rightarrow K^m$ anhand der zugehörigen Matrix berechnen können. Wir werden auch sehen, wie man feststellt, ob eine Matrix invertierbar ist, und wie man gegebenenfalls ihre Inverse findet.

12. DER NORMALFORMALGORITHMUS UND LINEARE GLEICHUNGSSYSTEME

Wie können wir den Rang einer durch eine Matrix $A \in \text{Mat}(m \times n, K)$ gegebenen linearen Abbildung $f: K^n \rightarrow K^m$ bestimmen und eine Basis ihres Kerns finden? Dazu überlegen wir uns, wie man die Matrix verändern kann, ohne dass sich der Kern ändert. Dann können wir versuchen, die Matrix in eine Form zu bringen, aus der sich zum Beispiel eine Basis des Kerns leicht ablesen lässt. Eine solche Form ist die *Zeilenstufenform*:

12.1. Definition. Seien K ein Körper, $m, n \in \mathbb{N}$ und $A = (a_{ij}) \in \text{Mat}(m \times n, K)$. Die Matrix A ist in *Zeilenstufenform*, wenn sie folgende Form hat (ein Stern steht für einen beliebigen Eintrag):

DEF
Zeilenstufenform

$$A = \begin{pmatrix} 0 \dots 0 & \mathbf{1} & * \dots * & * & * \dots * & * & * \dots * \\ 0 \dots 0 & 0 & 0 \dots 0 & \mathbf{1} & * \dots * & * & * \dots * \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & \mathbf{1} & * \dots * \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 \end{pmatrix}$$

Formal bedeutet das, dass es $0 \leq r \leq m$ und Indizes $1 \leq j_1 < j_2 < \dots < j_r \leq n$ gibt mit $a_{ij} = 0$, wenn $i > r$ oder $i \leq r$ und $j < j_i$, und $a_{ij_i} = 1$ für alle $i \in \{1, 2, \dots, r\}$. (Dabei ist r die Anzahl der Zeilen, die nicht nur aus Nullen bestehen, und $j_1, j_2, \dots, j_r$ sind die Nummern der Spalten, in denen die führenden Einsen der ersten r Zeilen stehen.)

A ist in *reduzierter Zeilenstufenform*, wenn zusätzlich $a_{ijk} = 0$ ist für alle $1 \leq i < k$ und alle $k \in \{1, 2, \dots, r\}$:

$$A = \begin{pmatrix} 0 \dots 0 & \mathbf{1} & * \dots * & 0 & * \dots * & 0 & * \dots * \\ 0 \dots 0 & 0 & 0 \dots 0 & \mathbf{1} & * \dots * & 0 & * \dots * \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & \mathbf{1} & * \dots * \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 & 0 & 0 \dots 0 \end{pmatrix}$$

Die „führenden Einsen“ der ersten r Zeilen (in den Spalten $j_1, j_2, \dots, j_r$) sind durch Fettdruck hervorgehoben, die zugehörigen Spalten in der Bildschirmversion auch farblich abgesetzt. $\diamond$

Zur Vereinfachung führen wir folgende Sprechweise ein:

12.2. Definition. Sei $A \in \text{Mat}(m \times n, K)$. Dann ist das *Bild* von A , $\text{im}(A)$, das Bild der zugehörigen linearen Abbildung $f: K^n \rightarrow K^m$, der *Rang* von A , $\text{rk}(A)$, ist der Rang von f und der *Kern* von A , $\text{ker}(A)$, ist der Kern von f . $\diamond$

DEF
Bild, Rang,
Kern
einer Matrix

Da die Bilder unter f der Standard-Basisvektoren von K^n gerade die Spalten von A sind, ist $\text{im}(A)$ das Erzeugnis der Spalten von A . Deswegen heißt $\text{im}(A)$ auch der *Spaltenraum* von A . Analog definiert man den *Zeilenraum* von A als den von den Zeilen von A erzeugten Untervektorraum von K^n .

DEF
Spaltenraum
Zeilenraum

12.3. Lemma. Sei $A \in \text{Mat}(m \times n, K)$ in Zeilenstufenform mit r und $j_1, j_2, \dots, j_r$ wie in Definition 12.1. Dann ist $\text{rk}(A) = r$. Hat die Matrix sogar reduzierte Zeilenstufenform, dann erhalten wir eine Basis von $\ker(A)$ wie folgt:

Sei $J = \{1, 2, \dots, n\} \setminus \{j_1, j_2, \dots, j_r\}$ die Menge der Indizes von Spalten ohne „führende Eins“ und sei für $j \in J$ der Vektor $b_j \in K^n$ definiert als

$$b_j = \mathbf{e}_j - \sum_{i=1}^r a_{ij} \mathbf{e}_{j_i}.$$

Dann ist $(b_j)_{j \in J}$ eine Basis von $\ker(A)$.

Etwas anschaulicher bekommen wir die Basis des Kerns so: Die Indizes in J , die den Spalten ohne führende Eins einer Zeile entsprechen, sind Positionen, für die wir die Komponenten frei wählen können. Wir setzen eine (die Position $j \in J$) davon auf 1, die anderen auf 0 und lösen die aus $Ab_j = \mathbf{0}$ entstehenden Gleichungen nach den übrigen Komponenten auf.

Beweis. Das Bild der zu A gehörenden linearen Abbildung f wird von den Spalten der Matrix erzeugt. Die Spalten mit den Nummern $j_1, j_2, \dots, j_r$ sind linear unabhängig: Aus

$$\lambda_1(1, 0, \dots, 0) + \lambda_2(*, 1, 0, \dots, 0) + \dots + \lambda_r(*, \dots, *, 1, 0, \dots, 0) = (0, \dots, 0)$$

folgt sukzessive $\lambda_r = 0, \dots, \lambda_2 = 0, \lambda_1 = 0$. Außerdem sind alle Spalten im Untervektorraum $\langle \mathbf{e}'_1, \mathbf{e}'_2, \dots, \mathbf{e}'_r \rangle$ der Dimension r von K^m enthalten. Also ist $\text{im}(f) = \langle \mathbf{e}'_1, \mathbf{e}'_2, \dots, \mathbf{e}'_r \rangle$; es folgt $\text{rk}(A) = r$.

Wir nehmen jetzt an, dass die Matrix reduzierte Zeilenstufenform hat. Aus der Dimensionsformel in Satz 10.17 folgt, dass der Kern von A Dimension $n - r = \#J$ hat. Es genügt also zu zeigen, dass die b_j im Kern liegen und linear unabhängig sind. Wir schreiben A_j für die j -te Spalte von A . Dann ist $A_j = \sum_{i=1}^r a_{ij} \mathbf{e}'_i$ (dabei verwenden wir, dass $a_{ij} = 0$ ist für $i > r$), und $A_{j_i} = \mathbf{e}'_i$ für $1 \leq i \leq r$ (das ist Teil der Aussage, dass A reduzierte Zeilenstufenform hat). Es folgt

$$f(b_j) = A_j - \sum_{i=1}^r a_{ij} A_{j_i} = \sum_{i=1}^r a_{ij} \mathbf{e}'_i - \sum_{i=1}^r a_{ij} \mathbf{e}'_i = \mathbf{0},$$

also ist b_j im Kern. Um zu zeigen, dass die b_j linear unabhängig sind, betrachten wir eine Linearkombination:

$$\mathbf{0} = \sum_{j \in J} \lambda_j b_j = \sum_{j \in J} \lambda_j \mathbf{e}_j - \sum_{i=1}^r \left(\sum_{j \in J} a_{ij} \lambda_j \right) \mathbf{e}_{j_i}$$

Da $(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$ eine Basis von K^n und $\{1, 2, \dots, n\}$ die disjunkte Vereinigung von J und $\{j_1, j_2, \dots, j_r\}$ ist, folgt $\lambda_j = 0$ für alle $j \in J$. $\square$

12.4. Beispiel. Sei $K = \mathbb{R}$ und A die folgende Matrix über $\mathbb{R}$:

$$A = \begin{pmatrix} 0 & 1 & 2 & 0 & 0 & -2 \\ 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 5 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Dann ist A in reduzierter Zeilenstufenform mit $r = 3$ (das ist die Anzahl der Zeilen, die keine Null-Zeilen sind) und $j_1 = 2, j_2 = 4, j_3 = 5$. Der Rang ist also 3, $J = \{1, 3, 6\}$ und eine Basis des Kerns ist gegeben durch

$$b_1 = (1, 0, 0, 0, 0, 0), \quad b_3 = (0, -2, 1, 0, 0, 0), \quad b_6 = (0, 2, 0, -1, -5, 1).$$

LEMMA
Rang und
Kern einer
Matrix
in ZSF

BSP

Die frei wählbaren Komponenten (Positionen 1, 3, 6) sind durch Fettdruck hervorgehoben. Die restlichen Komponenten von b_j ergeben sich aus den Negativen der ersten r Einträge der j -ten Spalte von A .

Was hier passiert, wird vielleicht klarer, wenn man die Bedingung „ $(x_1, x_2, \dots, x_6) \in \ker(A)$ “ explizit formuliert. Das führt auf das Gleichungssystem

$$\begin{array}{rcccccccl} x_2 & + & 2x_3 & & & - & 2x_6 & = & 0 \\ & & & x_4 & & + & x_6 & = & 0 \\ & & & & x_5 & + & 5x_6 & = & 0 \end{array}$$

Dass A in reduzierter Zeilenstufenform ist, bedeutet, dass man dieses Gleichungssystem nach x_2, x_4, x_5 auflösen kann:

$$\begin{aligned} x_2 &= -2x_3 + 2x_6 \\ x_4 &= -x_6 \\ x_5 &= -5x_6 \end{aligned}$$

Man kann also x_1, x_3, x_6 beliebig vorgeben und daraus x_2, x_4, x_5 bestimmen. Wenn man (x_1, x_3, x_6) die Standardbasis von $K^3 = K^{n-j}$ durchlaufen lässt, dann bekommt man so eine Basis des Kerns. ♣

Wie bekommen wir nun eine Matrix in diese Zeilenstufenform, ohne ihren Kern zu ändern? Dazu gehen wir schrittweise vor und führen kleine Veränderungen durch, von denen man leicht einsehen kann, dass sie diese Eigenschaft haben.

12.5. Definition. Seien K ein Körper, $m, n \in \mathbb{N}$ und $A \in \text{Mat}(m \times n, K)$.

DEF
(elementare)
Zeilen-
umformungen

- (1) Eine *elementare Zeilenumformung vom Typ I* an der Matrix A besteht darin, die i -te Zeile von A mit λ zu multiplizieren. Dabei ist $i \in \{1, 2, \dots, m\}$ und $\lambda \in K \setminus \{0\}$. Wir schreiben $\mathbf{I}_i(\lambda)$ für diese Umformung.
- (2) Eine *elementare Zeilenumformung vom Typ II* an der Matrix A besteht darin, das λ -fache der j -ten Zeile von A zur i -ten Zeile zu addieren. Dabei sind $i, j \in \{1, 2, \dots, m\}$ mit $i \neq j$ und $\lambda \in K$. Wir schreiben $\mathbf{II}_{i,j}(\lambda)$ für diese Umformung.
- (3) Eine *elementare Zeilenumformung vom Typ III* an der Matrix A besteht darin, in A die i -te und die j -te Zeile miteinander zu vertauschen. Dabei sind $i, j \in \{1, 2, \dots, m\}$ mit $i \neq j$. Wir schreiben $\mathbf{III}_{i,j}$ für diese Umformung.

Eine *Zeilenumformung* an der Matrix A ist eine Abfolge von sukzessiven elementaren Zeilenumformungen, beginnend mit der Matrix A . ◇

Eine elementare Zeilenumformung vom Typ III kann durch eine Abfolge geeigneter Umformungen der Typen I und II erreicht werden (Übung). Diese Art der Umformung ist also eigentlich nicht nötig, stellt aber häufig eine praktische Abkürzung dar.

12.6. Lemma. Seien K ein Körper, $m, n \in \mathbb{N}$ und $A \in \text{Mat}(m \times n, K)$. Sei weiter A' eine Matrix, die aus A durch eine elementare Zeilenumformung hervorgeht. Dann ist $\ker(A') = \ker(A)$ und daher auch $\text{rk}(A') = \text{rk}(A)$. Außerdem haben A' und A denselben Zeilenraum.

LEMMA
Zeilenumf.
erhalten
Kern, Rang,
Zeilenraum

Beweis. Ein Vektor $v = (x_1, x_2, \dots, x_n) \in K^n$ ist genau dann im Kern von $A = (a_{ij})$, wenn für alle $i \in \{1, 2, \dots, m\}$ gilt $\sum_{j=1}^n a_{ij}x_j = 0$. Eine elementare Zeilenumformung vom Typ I ersetzt eine dieser Gleichungen durch ihr λ -faches mit $\lambda \neq 0$, was ihre Gültigkeit nicht ändert. Bei einer elementaren Zeilenumformung vom Typ II wird zu einer der Gleichungen das λ -fache einer anderen Gleichung addiert, die neuen Gleichungen sind also gültig, wenn die alten es sind. Da man die Umformung rückgängig machen kann (durch Subtraktion des λ -fachen der j -ten Zeile von der i -ten), gelten die neuen Gleichungen genau dann, wenn die alten gelten. (Umformungen vom Typ III brauchen nicht extra betrachtet zu werden; da sie aber nur die Reihenfolge der Gleichungen ändern, ist klar, dass der Kern dabei erhalten bleibt.) Das zeigt, dass v genau dann im Kern von A ist, wenn v im Kern von A' ist. Die Gleichheit der Ränge folgt aus dem Rangsatz 10.17.

Elementare Zeilenumformungen ersetzen eine Zeile durch eine Linearkombination der Zeilen der Matrix; da diese Linearkombination im Zeilenraum der ursprünglichen Matrix liegt, ist der Zeilenraum der neuen Matrix jedenfalls im Zeilenraum der alten Matrix enthalten. Da sich elementare Zeilenumformungen rückgängig machen lassen, gilt auch die umgekehrte Inklusion. $\square$

Durch Induktion (über die Anzahl der elementaren Zeilenumformungen) folgt dann sofort, dass Kern, Rang und Zeilenraum auch unter beliebigen Zeilenumformungen erhalten bleiben.

Der Rang bleibt unter Zeilenumformungen zwar erhalten, das Bild der Matrix kann sich jedoch ändern!



Wir zeigen jetzt, dass man jede Matrix durch Zeilenumformungen in Zeilenstufenform überführen kann.

* **12.7. Satz.** Seien K ein Körper, $m, n \in \mathbb{N}$ und $A \in \text{Mat}(m \times n, K)$. Dann lässt sich A durch sukzessive elementare Zeilenumformungen in eine Matrix A' in reduzierter Zeilenstufenform überführen.

SATZ
Normal-
form von
Matrizen

Beweis. Wir zeigen zuerst, dass sich A in (nicht notwendig reduzierte) Zeilenstufenform bringen lässt. Der Beweis dafür geht durch Induktion nach der Zeilenanzahl m . Im Fall $m = 0$ ist die Matrix (trivialerweise) bereits in Zeilenstufenform. Sei also $m > 0$ und die Behauptung für alle Matrizen mit weniger als m Zeilen schon gezeigt. Ist A die Nullmatrix, dann ist A in Zeilenstufenform und es ist nichts zu zeigen. Wir können also annehmen, dass A einen von null verschiedenen Eintrag hat. Sei j_1 der kleinste Index einer Spalte mit einem solchen Eintrag. Ist $a_{1j_1} = 0$, dann können wir durch eine Typ-III-Umformung erreichen, dass $a_{1j_1} \neq 0$ ist. Eine Umformung vom Typ I mit $\lambda = a_{1j_1}^{-1}$, angewandt auf die erste Zeile, ergibt $a_{1j_1} = 1$. Die Matrix hat jetzt die Form

$$\begin{pmatrix} 0 & \cdots & 0 & 1 & * & \cdots & * \\ 0 & \cdots & 0 & a_{2j_1} & * & \cdots & * \\ & & \vdots & \vdots & & \vdots & \\ 0 & \cdots & 0 & a_{mj_1} & * & \cdots & * \end{pmatrix}$$

Durch die Umformungen $\mathbf{II}_{2,1}(-a_{2j_1}), \mathbf{II}_{3,1}(-a_{3j_1}), \dots, \mathbf{II}_{m,1}(-a_{mj_1})$ können wir die j_1 -te Spalte unterhalb der ersten Zeile „ausräumen“, sodass wir nun die Form

$$\left(\begin{array}{cccc|ccc} 0 & \cdots & 0 & 1 & * & \cdots & * \\ 0 & \cdots & 0 & 0 & & & \\ & & \vdots & \vdots & & & \\ 0 & \cdots & 0 & 0 & & \tilde{A} & \end{array} \right)$$

haben mit einer $(m-1) \times (n-j_1)$ -Matrix $\tilde{A}$. Zeilenumformungen an $\tilde{A}$ können auch als Zeilenumformungen an dieser Matrix ausgeführt werden, ohne dass sich am linken Teil der Matrix etwas ändert. Nach Induktionsannahme kann nun $\tilde{A}$ durch Zeilenumformungen in Zeilenstufenform gebracht werden. Damit hat die gesamte Matrix ebenfalls Zeilenstufenform.

Wir führen jetzt noch für $k = 1, 2, \dots, r$ und $i = 1, 2, \dots, k-1$ die Umformungen $\mathbf{II}_{i,k}(-a_{ij_k})$ aus (mit dem jeweils aktuellen Wert des Eintrags a_{ij_k}) und räumen auf diese Weise auch noch den Teil der Spalten oberhalb der führenden Einsen aus. Wir erhalten so die reduzierte Zeilenstufenform. $\square$

Dieser Beweis liefert uns sogar einen Algorithmus. Wir werden die Umformungen an einer Beispielmatrix durchführen.

12.8. Beispiel. Wir bestimmen die reduzierte Zeilenstufenform der folgenden Matrix über $\mathbb{R}$:

$$A = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix}$$

BSP
Umformung
in ZSF

Die erste Spalte ist keine Null-Spalte, also ist $j_1 = 1$. Der oberste Eintrag in der ersten Spalte ist bereits 1, also sind keine Umformungen vom Typ III oder I nötig. Wir räumen den Rest der Spalte aus, indem wir das Fünffache der ersten Zeile von der zweiten und das Neunfache der ersten Zeile von der dritten Zeile abziehen. Dann bekommen wir die neue Matrix

$$\begin{pmatrix} 1 & 2 & 3 & 4 \\ 0 & -4 & -8 & -12 \\ 0 & -8 & -16 & -24 \end{pmatrix}$$

Wir machen mit der rechten unteren 2×3 -Matrix weiter. Ihre erste Spalte $(-4, -8)$ ist keine Null-Spalte, also ist $j_2 = 2$. Wir multiplizieren die zweite Zeile der gesamten Matrix mit $-1/4$ und bekommen

$$\begin{pmatrix} 1 & 2 & 3 & 4 \\ 0 & 1 & 2 & 3 \\ 0 & -8 & -16 & -24 \end{pmatrix}.$$

Dann addieren wir das Achtfache der zweiten Zeile zur dritten:

$$\begin{pmatrix} 1 & 2 & 3 & 4 \\ 0 & 1 & 2 & 3 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Damit haben wir die Zeilenstufenform erreicht (mit $r = 2$). Für die reduzierte Zeilenstufenform müssen wir noch das Doppelte der zweiten Zeile von der ersten abziehen; das liefert schließlich

$$A' = \begin{pmatrix} 1 & 0 & -1 & -2 \\ 0 & 1 & 2 & 3 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Wir können jetzt eine Basis des Kerns von A (der gleich dem Kern von A' ist) ablesen, nämlich

$$b_3 = (1, -2, \mathbf{1}, \mathbf{0}) \quad \text{und} \quad b_4 = (2, -3, \mathbf{0}, \mathbf{1}). \quad \clubsuit$$

Elementare Zeilenumformungen lassen sich durch Multiplikation mit gewissen invertierbaren Matrizen von links beschreiben. Diese Matrizen sind die sogenannten *Elementarmatrizen* $E_i(\lambda)$ mit $\lambda \in K^\times$ und $i \in \{1, 2, \dots, m\}$ und $E_{ij}(\lambda)$ mit $\lambda \in K$ und $i, j \in \{1, 2, \dots, m\}$, $i \neq j$. Um sie zu definieren, führen wir $M_{kl} = (\delta_{ik}\delta_{jl})_{1 \leq i, j \leq m}$ ein; in dieser Matrix sind alle Einträge null bis auf den Eintrag in der k -ten Zeile und l -ten Spalte, der den Wert 1 hat. (Die Matrizen M_{kl} entsprechen der Basis $(\phi_{kl})_{1 \leq k, l \leq m}$ von $\text{Hom}(K^m, K^m)$ wie in Folgerung 10.23.) Dann ist

$$E_i(\lambda) = I_m + (\lambda - 1)M_{ii} \quad \text{und} \quad E_{ij}(\lambda) = I_m + \lambda M_{ij}.$$

$E_i(\lambda)$ unterscheidet sich von der Einheitsmatrix I_m dadurch, dass an der i -ten Position auf der Diagonalen statt 1 der Eintrag λ steht. In $E_{ij}(\lambda)$ steht außerhalb der Diagonalen an der Position (i, j) der Eintrag λ . Wegen

$$E_i(\lambda)E_i(\lambda^{-1}) = I_m \quad \text{und} \quad E_{ij}(\lambda)E_{ij}(-\lambda) = I_m$$

sind diese Elementarmatrizen invertierbar. Was bewirkt die Multiplikation von links mit so einer Elementarmatrix? Dazu überlegen wir, dass gilt

$$M_{kl}A = \left(\sum_{h=1}^m \delta_{ik}\delta_{hl}a_{hj} \right)_{1 \leq i \leq m, 1 \leq j \leq n} = (\delta_{ik}a_{lj})_{1 \leq i \leq m, 1 \leq j \leq n};$$

in dieser Matrix sind alle Zeilen null bis auf die k -te Zeile, in welcher sich die l -te Zeile von A befindet. Multiplikation von links mit M_{kl} setzt also die l -te Zeile von A in die k -te Zeile und löscht alle anderen Zeilen.

Damit ergibt sich, dass die Zeilen von $E_i(\lambda)A$ mit den entsprechenden Zeilen von A übereinstimmen bis auf die i -te Zeile, die mit λ multipliziert wird. Der Effekt ist also die elementare Zeilenumformung $\mathbf{I}_i(\lambda)$ vom Typ I. Ebenso stimmen die Zeilen von $E_{ij}(\lambda)A$ mit denen von A überein mit Ausnahme der i -ten Zeile, zu der das λ -fache der j -ten Zeile addiert wird. Der Effekt ist also die elementare Zeilenumformung $\mathbf{II}_{i,j}(\lambda)$ vom Typ II.

Wir veranschaulichen das für die 2×3 -Matrix $A = \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix}$:

$$\begin{aligned} E_1(\lambda)A &= \begin{pmatrix} \lambda & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix} = \begin{pmatrix} \lambda a & \lambda b & \lambda c \\ d & e & f \end{pmatrix} \\ E_2(\lambda)A &= \begin{pmatrix} 1 & 0 \\ 0 & \lambda \end{pmatrix} \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix} = \begin{pmatrix} a & b & c \\ \lambda d & \lambda e & \lambda f \end{pmatrix} \\ E_{12}(\lambda)A &= \begin{pmatrix} 1 & \lambda \\ 0 & 1 \end{pmatrix} \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix} = \begin{pmatrix} a + \lambda d & b + \lambda e & c + \lambda f \\ d & e & f \end{pmatrix} \\ E_{21}(\lambda)A &= \begin{pmatrix} 1 & 0 \\ \lambda & 1 \end{pmatrix} \begin{pmatrix} a & b & c \\ d & e & f \end{pmatrix} = \begin{pmatrix} a & b & c \\ d + \lambda a & e + \lambda b & f + \lambda c \end{pmatrix} \end{aligned}$$

Der Inhalt von Satz 12.7 ist also, dass es zu jeder Matrix $A \in \text{Mat}(m \times n, K)$ eine invertierbare Matrix $P \in \text{Mat}(m, K)$ gibt, sodass PA reduzierte Zeilenstufenform hat, wobei P ein Produkt von Elementarmatrizen ist.

Sei jetzt $A \in \text{Mat}(m, K)$ invertierbar. Wendet man diese Aussage an auf A^{-1} und beachtet, dass die reduzierte Zeilenstufenform einer invertierbaren $m \times m$ -Matrix gerade die Einheitsmatrix I_m ist (siehe Lemma 12.18 unten), dann erhält man ein Produkt P von Elementarmatrizen mit $PA^{-1} = I_m$. Es folgt $A = P$. Wir haben bewiesen:

Satz. Jede invertierbare Matrix ist ein Produkt von Elementarmatrizen.

Daraus folgt:

Zwei Matrizen $A, B \in \text{Mat}(m \times n, K)$ lassen sich genau dann durch Zeilenumformungen ineinander überführen, wenn es eine invertierbare Matrix $P \in \text{Mat}(m, K)$ gibt, sodass $B = PA$ ist.

Statt Zeilenumformungen kann man ganz analog *Spaltenumformungen* betrachten. Sie werden durch Multiplikation mit Elementarmatrizen von *rechts* bewirkt. Man hat dann die folgende analoge Aussage:

Zwei Matrizen $A, B \in \text{Mat}(m \times n, K)$ lassen sich genau dann durch Spaltenumformungen ineinander überführen, wenn es eine invertierbare Matrix $Q \in \text{Mat}(n, K)$ gibt, sodass $B = AQ$ ist.

Wir sprechen gelegentlich von „der“ reduzierten Zeilenstufenform einer Matrix. Tatsächlich ist das Ergebnis eines Verfahrens, das eine Matrix in reduzierte Zeilenstufenform überführt, eindeutig bestimmt, wie der folgende Satz zeigt. Es ist also letztlich ganz egal, welche Zeilenumformungen man in welcher Reihenfolge macht, um zur reduzierten Zeilenstufenform zu gelangen.

Satz. Sind $A, B \in \text{Mat}(m \times n, K)$ zwei Matrizen in reduzierter Zeilenstufenform mit demselben Zeilenraum, dann gilt $A = B$.

SATZ
Elementar-
matrizen
erzeugen
invertierbare
Matrizen

SATZ
Eindeutigkeit
der Zeilen-
stufenform

Beweis. Sei $U \subset K^n$ der Zeilenraum von A und B . Für $0 \leq k \leq n$ sei

$$V_k = \{(x_1, \dots, x_n) \in K^n \mid x_1 = x_2 = \dots = x_k = 0\} \subset K^n$$

und $d_k = \dim(U \cap V_k)$. Dann ist $d_0 = r = \dim U$, $d_n = 0$ und $d_k - 1 \leq d_{k+1} \leq d_k$. Es gibt also genau r „Sprungstellen“ j_i mit $d_{j_i-1} = r + 1 - i$ und $d_{j_i} = r - i$ für $i \in \{1, 2, \dots, r\}$. Aus der Definition der Zeilenstufenform ergibt sich, dass j_i genau die Position der führenden Eins in der i -ten Zeile von A und von B ist. Die lineare Abbildung

$$\phi: U \longrightarrow K^r, \quad (x_1, x_2, \dots, x_n) \longmapsto (x_{j_1}, x_{j_2}, \dots, x_{j_r})$$

ist dann ein Isomorphismus, und die ersten r Zeilen von A und B müssen die Urbilder $\phi^{-1}(\mathbf{e}_1), \phi^{-1}(\mathbf{e}_2), \dots, \phi^{-1}(\mathbf{e}_r)$ der Standard-Basisvektoren von K^r sein. Insbesondere sind A und B gleich. $\square$

Wir kommen zu linearen Gleichungen und Gleichungssystemen.

12.9. Definition. Seien V und W zwei K -Vektorräume und $f: V \rightarrow W$ eine lineare Abbildung. Ist $b \in W$ ein gegebener Vektor, dann heißt die Gleichung $f(x) = b$, deren Lösungen $x \in V$ gesucht sind, eine *lineare Gleichung*. Die Gleichung heißt *homogen*, wenn $b = \mathbf{0}$ ist, sonst *inhomogen*.

DEF
Lineare
Gleichung

Ist $V = K^n$ und $W = K^m$, dann kann die Gleichung unter Benutzung der zu f gehörenden Matrix $A \in \text{Mat}(m \times n, K)$ auch geschrieben werden als $A\mathbf{x} = \mathbf{b}$ mit Spaltenvektoren $\mathbf{x} \in K^n$ und $\mathbf{b} \in K^m$. In diesem Fall spricht man auch von einem *linearen Gleichungssystem* (mit m Gleichungen in n Unbekannten). $\diamond$

Wir können schon recht genau sagen, welche Struktur die Lösungsmenge einer linearen Gleichung hat.

* 12.10. **Satz.** Seien V und W zwei K -Vektorräume und $f: V \rightarrow W$ eine lineare Abbildung.

SATZ
Lösungs-
menge
einer
linearen
Gleichung

- (1) Die Lösungsmenge der homogenen linearen Gleichung $f(x) = \mathbf{0}$ ist ein Untervektorraum von V , nämlich der Kern von f .
- (2) Sei $\mathbf{0} \neq b \in W$. Ist $b \notin \text{im}(f)$, dann hat die inhomogene lineare Gleichung $f(x) = b$ keine Lösung. Anderenfalls sei $x_0 \in V$ mit $f(x_0) = b$. Dann ist die Lösungsmenge gegeben durch $x_0 + \ker(f) = \{x_0 + v \mid v \in \ker(f)\}$.

Beweis. Die erste Aussage folgt direkt aus der Definition des Kerns und der Tatsache, dass $\ker(f)$ ein Untervektorraum von V ist. In der zweiten Aussage ist klar, dass es genau dann Lösungen gibt, wenn $b \in \text{im}(f)$ ist (das ist die Definition von $\text{im}(f)$). Es bleibt die letzte Behauptung zu zeigen. Sei dazu $x \in V$. Dann gilt

$$\begin{aligned} f(x) = b &\iff f(x) = f(x_0) \iff f(x - x_0) = \mathbf{0} \\ &\iff x - x_0 \in \ker(f) \iff x \in x_0 + \ker(f). \quad \square \end{aligned}$$

Das allgemeine Rezept für die Lösung einer linearen Gleichung $f(x) = b$ lautet also:

- (1) Prüfe, ob $b \in \text{im}(f)$. Falls nein, dann gibt es keine Lösung.
- (2) Bestimme eine „spezielle Lösung“ $x_0 \in V$.
- (3) Bestimme $\ker(f)$.
- (4) Die Lösungsmenge ist $x_0 + \ker(f)$. Ist $\ker(f)$ endlich-dimensional mit Basis $(x_1, x_2, \dots, x_n)$, dann ist die „allgemeine Lösung“

$$x = x_0 + \lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_n x_n$$

mit $\lambda_1, \lambda_2, \dots, \lambda_n \in K$.

Die ersten beiden Schritte wird man im Regelfall zusammen ausführen, denn wenn man feststellt, dass $b \in \text{im}(f)$ ist, dann wird man meistens auch ein Urbild gefunden haben.

Im homogenen Fall (also $b = \mathbf{0}$) gilt stets $b \in \text{im}(f)$ und wir können $x_0 = \mathbf{0}$ nehmen; die Lösungsmenge ist dann $\ker(f)$.

12.11. **Beispiel.** Wir betrachten die folgende inhomogene lineare Differentialgleichung erster Ordnung:

$$y'(x) + y(x) = x.$$

BSP
inhomogene
lineare
Diff.gleichung

Dabei sei $y \in \mathcal{C}^1(\mathbb{R})$. Hier ist $K = \mathbb{R}$, $V = \mathcal{C}^1(\mathbb{R})$, $W = \mathcal{C}(\mathbb{R})$ und $f: y \mapsto y' + y$. Die Gleichung ist $f(y) = \text{id}_{\mathbb{R}}$. Wir suchen nach einer speziellen Lösung. Mit etwas Probieren finden wir $y_0(x) = x - 1$. (In der Vorlesung über *Gewöhnliche Differentialgleichungen* werden Sie lernen, wie man solche Lösungen systematisch findet.) Jetzt müssen wir den Kern von f bestimmen, also die Menge aller Funktionen y mit $y' + y = \mathbf{0}$. Ich behaupte, dass $\ker(f) = \langle x \mapsto e^{-x} \rangle$ ist; die Funktionen y mit $y' + y = \mathbf{0}$ haben also die Form $y(x) = C e^{-x}$ mit $C \in \mathbb{R}$. Zum Beweis betrachten wir $z(x) = e^x y(x)$; dann gilt

$$z'(x) = e^x y(x) + e^x y'(x) = e^x (y(x) + y'(x)) = \mathbf{0},$$

also ist $z(x) = C$ konstant und damit $y(x) = C e^{-x}$. Umgekehrt sind diese Funktionen auch Lösungen von $y' + y = \mathbf{0}$. Die allgemeine Lösung ist also

$$y(x) = x - 1 + C e^{-x}, \quad C \in \mathbb{R}.$$



Wie sieht das obige Rezept konkret aus, wenn wir ein lineares Gleichungssystem lösen wollen? Sei $A\mathbf{x} = \mathbf{b}$ ein lineares Gleichungssystem mit $A \in \text{Mat}(m \times n, K)$. Im homogenen Fall $\mathbf{b} = \mathbf{0}$ müssen wir eine Basis von $\ker(A)$ bestimmen. Dazu bringen wir A in reduzierte Zeilenstufenform und lesen eine Basis des Kerns ab wie in Lemma 12.3. Im inhomogenen Fall sei $A' = (A \mid \mathbf{b})$ die *erweiterte Matrix* des Systems; wir erhalten sie, indem wir an die Matrix A den Spaltenvektor $\mathbf{b}$ als $(n+1)$ -te Spalte anfügen.

DEF
erweiterte
Matrix

12.12. Satz. Sei K ein Körper, seien $m, n \in \mathbb{N}$, sei $A \in \text{Mat}(m \times n, K)$ und sei $\mathbf{b} \in K^m$ ein Spaltenvektor. Sei weiter $A' = (A \mid \mathbf{b})$ die erweiterte Matrix des linearen Gleichungssystems $A\mathbf{x} = \mathbf{b}$. Dann gilt

SATZ
inhom.
LGS

$$\mathbf{b} \in \text{im}(A) \iff \text{rk}(A') = \text{rk}(A).$$

Dies kann geprüft werden, indem A' in reduzierte Zeilenstufenform $\tilde{A}'$ gebracht wird. $\text{rk}(A') = \text{rk}(A)$ ist dann dazu äquivalent, dass die letzte Spalte von $\tilde{A}'$ keine führende Eins einer Zeile enthält (das bedeutet $j_r \leq n$ in der Notation von Definition 12.1). In diesem Fall kann eine spezielle Lösung von $A\mathbf{x} = \mathbf{b}$ aus $\tilde{A}'$ wie folgt abgelesen werden: Die letzte Spalte von $\tilde{A}'$ sei $(\tilde{b}_1, \dots, \tilde{b}_r, 0, \dots, 0)$. Dann ist

$$\mathbf{x}_0 = \sum_{i=1}^r \tilde{b}_i \mathbf{e}_{j_i}$$

eine Lösung des Gleichungssystems.

Beweis. Seien $A_1, \dots, A_n$ die Spalten von A . Es gilt

$$\begin{aligned} \mathbf{b} \in \text{im}(A) = \langle A_1, \dots, A_n \rangle &\iff \langle A_1, \dots, A_n, \mathbf{b} \rangle = \langle A_1, \dots, A_n \rangle \\ &\iff \text{im}(A') = \text{im}(A). \end{aligned}$$

Die letzte Aussage impliziert $\text{rk}(A') = \text{rk}(A)$. Es gilt immer $\text{im}(A) \subset \text{im}(A')$, also folgt aus $\text{rk}(A') = \text{rk}(A)$ auch die Gleichheit von $\text{im}(A')$ und $\text{im}(A)$. Damit ist die erste Behauptung gezeigt.

Sei nun $\tilde{A}'$ die reduzierte Zeilenstufenform von A' . Dann bilden die ersten n Spalten von $\tilde{A}'$ die reduzierte Zeilenstufenform $\tilde{A}$ von A . Der Rang von A' ist genau dann größer als der Rang von A , wenn $\tilde{A}'$ mehr Nichtnull-Zeilen hat als $\tilde{A}$. Das bedeutet aber gerade, dass die letzte Spalte von $\tilde{A}'$ eine führende Eins enthalten muss. Das zeigt die zweite Aussage. Für die letzte Aussage beachten wir, dass die Zeilenumformungen, die im Zuge der Herstellung der reduzierten Zeilenstufenform durchgeführt werden, die ursprünglichen Gleichungen durch äquivalente Gleichungen ersetzen. Mit $\tilde{A}' = (\tilde{A} \mid \tilde{\mathbf{b}})$ hat also das lineare Gleichungssystem $\tilde{A}\mathbf{x} = \tilde{\mathbf{b}}$ dieselben Lösungen wie das ursprüngliche Gleichungssystem. Da die j_i -te Spalte von $\tilde{A}$ gerade der Standard-Basisvektor $\mathbf{e}'_i$ ist, ergibt sich

$$\tilde{A}\mathbf{x}_0 = \sum_{i=1}^r \tilde{b}_i \tilde{A}\mathbf{e}_{j_i} = \sum_{i=1}^r \tilde{b}_i \mathbf{e}'_i = \tilde{\mathbf{b}}.$$

□

12.13. Beispiel. Wir lösen das folgende lineare Gleichungssystem (mit $K = \mathbb{Q}$ oder $\mathbb{R}$):

**BSP
LGS**

$$\begin{array}{cccccccl} x_1 & & & - & x_3 & - & 2x_4 & = & 3 \\ -x_1 & + & x_2 & + & x_3 & & & = & -2 \\ & & x_2 & & & - & x_4 & = & 0 \end{array}$$

oder, in Matrixschreibweise,

$$\begin{pmatrix} 1 & 0 & -1 & -2 \\ -1 & 1 & 1 & 0 \\ 0 & 1 & 0 & -1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = \begin{pmatrix} 3 \\ -2 \\ 0 \end{pmatrix}.$$

Die erweiterte Matrix ist

$$A' = \begin{pmatrix} 1 & 0 & -1 & -2 & 3 \\ -1 & 1 & 1 & 0 & -2 \\ 0 & 1 & 0 & -1 & 0 \end{pmatrix}.$$

Ihre reduzierte Zeilenstufenform ergibt sich als

$$\tilde{A}' = \begin{pmatrix} \mathbf{1} & 0 & -1 & 0 & 1 \\ 0 & \mathbf{1} & 0 & 0 & -1 \\ 0 & 0 & 0 & \mathbf{1} & -1 \end{pmatrix}.$$

Das zugehörige lineare Gleichungssystem ist:

$$\begin{array}{cccccccl} x_1 & & & - & x_3 & & & = & 1 \\ & & x_2 & & & & & = & -1 \\ & & & & & & x_4 & = & -1 \end{array}$$

Als spezielle Lösung erhalten wir daraus $\mathbf{x}_0 = (\mathbf{1}, -\mathbf{1}, 0, -\mathbf{1})$. Außerdem lesen wir ab: $\text{rk}(A) = 3$, $\dim \ker(A) = 4 - 3 = 1$, und eine Basis von $\ker(A)$ ist gegeben durch $\mathbf{x}_1 = (1, 0, \mathbf{1}, 0)$. Die allgemeine Lösung des Gleichungssystems ist also

$$\mathbf{x} = \mathbf{x}_0 + \lambda_1 \mathbf{x}_1 = (1 + \lambda_1, -1, \lambda_1, -1)$$

mit $\lambda_1 \in K$.

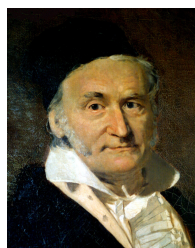


Hier ist das Rezept noch einmal ganz konkret:

- (1) Die erweiterte Matrix A' aufstellen.
- (2) A' in reduzierte Zeilenstufenform $\tilde{A}'$ bringen. Sei $r = \text{rk}(\tilde{A}')$ und seien $1 \leq j_1 < j_2 < \dots < j_r \leq n + 1$ die Positionen der führenden Einsen der ersten r Zeilen von $\tilde{A}'$.
- (3) $j_r = n + 1 \Rightarrow$ keine Lösung. Anderenfalls:
- (4) Sei $J = \{1, 2, \dots, n\} \setminus \{j_1, j_2, \dots, j_r\}$ die Menge der „freien“ Positionen. Setze $x_j = \lambda_j \in K$ beliebig für $j \in J$ und löse das der Matrix $\tilde{A}'$ entsprechende Gleichungssystem nach x_{j_i} , $i \in \{1, 2, \dots, r\}$ auf. Das ergibt die allgemeine Lösung.

Im Beispiel oben ist $r = 3$, $j_1 = 1$, $j_2 = 2$, $j_3 = 4 < 5 = n + 1$, $J = \{3\}$. Wir setzen also $x_3 = \lambda$ und lösen das System nach x_1, x_2, x_4 auf.

Diese Methode für die Lösung linearer Gleichungssysteme (und ihre Varianten) heißt *gaußsches Eliminationsverfahren* oder kürzer *Gauß-Elimination*. Eine Variante besteht darin, statt der reduzierten Zeilenstufenform nur die Zeilenstufenform herzustellen und dann das System schrittweise „von unten her“ durch Einsetzen



C.F. Gauß
(1777–1855)

zu lösen. Diese Version ist etwas effizienter im Hinblick auf die Zahl der nötigen Rechenoperationen, dafür aber auch etwas umständlicher durchzuführen.

Wir haben den Rang einer Matrix als den Rang der zugehörigen linearen Abbildung definiert, also als die Dimension ihres Spaltenraums. Man sollte also eigentlich genauer vom „Spaltenrang“ sprechen, denn man könnte genauso gut die Dimension des Zeilenraums, also den „Zeilenrang“ betrachten. Zum Glück macht das keinen Unterschied, wie wir jetzt zeigen werden.

12.14. Satz. *Sei K ein Körper, seien $m, n \in \mathbb{N}$ und $A \in \text{Mat}(m \times n, K)$. Dann ist die Dimension des Zeilenraums von A gleich der Dimension des Spaltenraums von A .*

SATZ
Zeilenrang =
Spaltenrang

Beweis. Nach Lemma 12.6 und Satz 12.7 können wir annehmen, dass A reduzierte Zeilenstufenform hat. Sei $r = \text{rk}(A)$ die Dimension des Spaltenraums von A . Dann hat A genau r Zeilen, die keine Null-Zeilen sind, und diese Zeilen sind linear unabhängig, denn das gilt bereits, wenn man nur die Spalten $j_1, j_2, \dots, j_r$ (Notation wie in Definition 12.1) betrachtet — die Matrix $(a_{i,j_k})_{1 \leq i, k \leq r}$ ist die Einheitsmatrix I_r . Da diese Zeilen den Zeilenraum erzeugen, bilden sie also eine Basis; damit ist die Dimension des Zeilenraums $r = \text{rk}(A)$. $\square$

Der Normalformalgorithmus aus Satz 12.7 berechnet demnach auch die Dimension und eine Basis des Zeilenraums der gegebenen Matrix. Wenn man also die Dimension und eine Basis des von Vektoren $v_1, \dots, v_m \in K^n$ erzeugten Untervektorraums bestimmen möchte, dann schreibt man diese Vektoren als Zeilen in eine Matrix und bestimmt ihre (reduzierte) Zeilenstufenform. Die von null verschiedenen Zeilen der resultierenden Matrix bilden dann eine Basis.

Man kann den obigen Satz kurz und elegant formulieren, wenn man folgende Definition verwendet.

*** 12.15. Definition.** Sei K ein Körper, seien $m, n \in \mathbb{N}$ und sei $A = (a_{ij}) \in \text{Mat}(m \times n, K)$. Die *Transponierte* von A oder die *zu A transponierte Matrix* ist $A^\top = (a_{ji})_{1 \leq i \leq n, 1 \leq j \leq m} \in \text{Mat}(n \times m, K)$.

DEF
Transponierte
Matrix $\diamond$

Da es leicht zu Verwirrung führt: Die Schreibweise

$$A^\top = (a_{ji})_{1 \leq i \leq n, 1 \leq j \leq m}$$

bedeutet Folgendes: Der erste Index unten hinter der Klammer (hier i) ist der Zeilenindex und der zweite (hier j) ist der Spaltenindex. Die Matrix $A^\top$ hat also n Zeilen und m Spalten. Der Eintrag in Zeile i und Spalte j ist a_{ji} und damit derselbe Eintrag wie in *Spalte i und Zeile j* der Matrix A . Gleichbedeutend könnte man auch

$$A^\top = (a_{ij})_{1 \leq j \leq n, 1 \leq i \leq m}$$

schreiben. In diesem Fall wäre j der Zeilen- und i der Spaltenindex.

Die Matrix wird also „an der Hauptdiagonale gespiegelt“.

12.16. **Beispiel.**

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}^\top = \begin{pmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{pmatrix}.$$

BSPTransponierte
Matrix

Der Zeilenraum von A ist der Spaltenraum von $A^\top$ und umgekehrt. Satz 12.14 sagt also

$$\text{rk}(A^\top) = \text{rk}(A).$$

Hier sind die wichtigsten Rechenregeln für transponierte Matrizen:

12.17. **Lemma.** Sei K ein Körper und seien $l, m, n \in \mathbb{N}$.

LEMMARechenregeln
für $A^\top$

- (1) Die Abbildung $\text{Mat}(m \times n, K) \rightarrow \text{Mat}(n \times m, K)$, $A \mapsto A^\top$ ist ein Isomorphismus (es gilt also $(A + B)^\top = A^\top + B^\top$ und $(\lambda A)^\top = \lambda A^\top$ für $A, B \in \text{Mat}(m \times n, K)$, $\lambda \in K$; die Bijektivität folgt aus Teil (3)).
- (2) Für $A \in \text{Mat}(l \times m, K)$ und $B \in \text{Mat}(m \times n, K)$ gilt $(AB)^\top = B^\top A^\top$.
- (3) Für $A \in \text{Mat}(m \times n, K)$ gilt $(A^\top)^\top = A$.

Beweis. Übung. □

Abschließend wollen wir noch Antworten auf die folgenden zwei Fragen überlegen:

Frage 1: Wie kann man feststellen, ob eine Matrix $A \in \text{Mat}(n, K)$ invertierbar ist?

Frage 2: Wie kann man zu einer invertierbaren Matrix $A \in \text{Mat}(n, K)$ die Inverse A^{-1} berechnen?

Die erste Frage wird durch das folgende Lemma beantwortet:

12.18. **Lemma.** Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$ eine quadratische Matrix. A ist genau dann invertierbar, wenn ihre reduzierte Zeilenstufenform die Einheitsmatrix I_n ist.

LEMMArZSF einer
invertierbaren
Matrix

Beweis. A ist genau dann invertierbar, wenn die zugehörige lineare Abbildung $f: K^n \rightarrow K^n$ ein Isomorphismus ist. Da Definitions- und Wertebereich dieselbe Dimension haben, ist das dazu äquivalent, dass f surjektiv ist, also Rang n hat (siehe Folgerung 10.14). Das bedeutet, dass es in der reduzierten Zeilenstufenform von A keine Null-Zeile gibt, also ist $r = n$ und $j_1 = 1, j_2 = 2, \dots, j_n = n$. In der j -ten Spalte steht also der j -te Standard-Basisvektor, und die Matrix ist die Einheitsmatrix. □

Jetzt zur zweiten Frage. Dazu beachten wir, dass ein lineares Gleichungssystem $A\mathbf{x} = \mathbf{b}$ für jedes $\mathbf{b}$ eine eindeutige Lösung hat, wenn A invertierbar ist; diese Lösung ist $\mathbf{x} = A^{-1}\mathbf{b}$. (Die Umkehrung gilt ebenfalls — gibt es für jedes $\mathbf{b}$ eine eindeutige Lösung, dann ist A invertierbar — Übung.) Wenn wir für $\mathbf{b}$ den Standard-Basisvektor $\mathbf{e}_j$ einsetzen, dann bekommen wir als Lösung gerade die j -te Spalte von A^{-1} . Wir können also A^{-1} finden, indem wir die linearen Gleichungssysteme $A\mathbf{x} = \mathbf{e}_j$ für $j \in \{1, 2, \dots, n\}$ alle lösen. Dies geht im Wesentlichen in einem Rutsch, wie im nächsten Satz beschrieben wird.

* **12.19. Satz.** Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$ eine quadratische Matrix. Sei weiter $A' = (A \mid I_n) \in \text{Mat}(n \times 2n, K)$ und $\tilde{A}'$ ihre reduzierte Zeilenstufenform. A ist genau dann invertierbar, wenn $\tilde{A}'$ die Form $(I_n \mid B)$ hat; in diesem Fall ist $B = A^{-1}$.

SATZ
Berechnung
von A^{-1}

Beweis. Sei $\tilde{A}' = (\tilde{A} \mid B)$, dann ist $\tilde{A}$ die reduzierte Zeilenstufenform von A . Nach Lemma 12.18 ist A genau dann invertierbar, wenn $\tilde{A} = I_n$ ist. Die Matrix A' repräsentiert das Gleichungssystem $A(\mathbf{x}_1 \mid \mathbf{x}_2 \mid \cdots \mid \mathbf{x}_n) = (\mathbf{e}_1 \mid \mathbf{e}_2 \mid \cdots \mid \mathbf{e}_n)$ oder kurz $AX = I_n$ mit $X = (\mathbf{x}_1 \mid \mathbf{x}_2 \mid \cdots \mid \mathbf{x}_n) \in \text{Mat}(n, K)$. Die Zeilenumformungen, die zur reduzierten Zeilenstufenform führen, ergeben das dazu äquivalente Gleichungssystem $I_n X = B$, also ist $X = B$ die Lösung von $AX = I_n$; damit ist $B = A^{-1}$. $\square$

12.20. Beispiel. Sei $K = \mathbb{Q}$ und

$$A = \begin{pmatrix} -1 & -2 & 2 & 2 \\ 2 & 5 & -4 & -4 \\ -1 & -1 & 1 & 1 \\ 1 & 1 & 0 & 1 \end{pmatrix}.$$

BSP
Berechnung
von A^{-1}

Wir überführen

$$A' = \begin{pmatrix} -1 & -2 & 2 & 2 & 1 & 0 & 0 & 0 \\ 2 & 5 & -4 & -4 & 0 & 1 & 0 & 0 \\ -1 & -1 & 1 & 1 & 0 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 \end{pmatrix}$$

in reduzierte Zeilenstufenform:

$$\begin{aligned} A' &\xrightarrow{\mathbf{I}_1(-1); \mathbf{II}_{2,1}(-2), \mathbf{III}_{3,1}(1), \mathbf{IV}_{4,1}(-1)} \begin{pmatrix} 1 & 2 & -2 & -2 & -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 2 & 1 & 0 & 0 \\ 0 & 1 & -1 & -1 & -1 & 0 & 1 & 0 \\ 0 & -1 & 2 & 3 & 1 & 0 & 0 & 1 \end{pmatrix} \\ &\xrightarrow{\mathbf{II}_{1,2}(-2), \mathbf{II}_{3,2}(-1), \mathbf{II}_{4,2}(1)} \begin{pmatrix} 1 & 0 & -2 & -2 & -5 & -2 & 0 & 0 \\ 0 & 1 & 0 & 0 & 2 & 1 & 0 & 0 \\ 0 & 0 & -1 & -1 & -3 & -1 & 1 & 0 \\ 0 & 0 & 2 & 3 & 3 & 1 & 0 & 1 \end{pmatrix} \\ &\xrightarrow{\mathbf{I}_3(-1); \mathbf{II}_{1,3}(2), \mathbf{II}_{4,3}(-2)} \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 0 & -2 & 0 \\ 0 & 1 & 0 & 0 & 2 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 & 3 & 1 & -1 & 0 \\ 0 & 0 & 0 & 1 & -3 & -1 & 2 & 1 \end{pmatrix} \\ &\xrightarrow{\mathbf{II}_{3,4}(-1)} \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 0 & -2 & 0 \\ 0 & 1 & 0 & 0 & 2 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 6 & 2 & -3 & -1 \\ 0 & 0 & 0 & 1 & -3 & -1 & 2 & 1 \end{pmatrix} \end{aligned}$$

Es folgt

$$A^{-1} = \begin{pmatrix} 1 & 0 & -2 & 0 \\ 2 & 1 & 0 & 0 \\ 6 & 2 & -3 & -1 \\ -3 & -1 & 2 & 1 \end{pmatrix}.$$



13. MATRIZEN UND LINEARE ABBILDUNGEN

Wir haben bisher Matrizen als zu linearen Abbildungen $K^n \rightarrow K^m$ gehörend betrachtet. Dabei war es aber eigentlich nur wichtig, dass wir in Definitions- und Wertebereich jeweils eine bestimmte Basis betrachten, in diesem Fall die Standard-Basis. Ganz analog können wir einer K -linearen Abbildung $f: V \rightarrow V'$ eine Matrix zuordnen, wenn wir Basen $B = (b_1, b_2, \dots, b_n)$ von V und $B' = (b'_1, b'_2, \dots, b'_m)$ von V' fixieren. Es gibt dann nämlich eindeutig bestimmte Skalare $a_{ij} \in K$, sodass

$$f(b_j) = a_{1j}b'_1 + a_{2j}b'_2 + \dots + a_{mj}b'_m$$

für alle $j \in \{1, 2, \dots, n\}$ gilt.

13.1. Definition. In der oben beschriebenen Situation heißt

$$\text{Mat}_{B,B'}(f) = (a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n} \in \text{Mat}(m \times n, K)$$

die *Matrix von f bezüglich der Basen B und B'* .

DEF $\text{Mat}_{B,B'}(f)$

◇

Wie vorher auch enthält die j -te Spalte der Matrix die Koeffizienten des Bildes $f(b_j)$ des j -ten Basisvektors in B , wenn es als Linearkombination der Basisvektoren in B' geschrieben wird.

Man kann $\text{Mat}_{B,B'}(f)$ auch beschreiben als die Matrix, die zu der linearen Abbildung $\phi_{B'}^{-1} \circ f \circ \phi_B: K^n \rightarrow K^m$ gehört.

$$\begin{array}{ccc} V & \xleftarrow{\phi_B} & K^n \\ f \downarrow & & \downarrow \text{Mat}_{B,B'}(f) \\ V' & \xleftarrow{\phi_{B'}} & K^m \end{array}$$

Dabei ist $\phi_B: K^n \rightarrow V$ (und analog $\phi_{B'}: K^m \rightarrow V'$) die zu B gehörende Linearkombinationenabbildung:

$$\phi_B(\lambda_1, \lambda_2, \dots, \lambda_n) = \lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n.$$

Da B und B' Basen sind, sind die Linearkombinationenabbildungen ϕ_B und $\phi_{B'}$ Isomorphismen; insbesondere gibt es die Umkehrabbildung $\phi_{B'}^{-1}$.

13.2. Beispiel. Wir betrachten $V = P_{<3}$, den $\mathbb{R}$ -Vektorraum der Polynomfunktionen vom Grad < 3 und die lineare Abbildung $D: V \rightarrow V$, $f \mapsto f'$. Seien weiter $B = (x \mapsto 1, x \mapsto x, x \mapsto x^2)$ und $B' = (x \mapsto 1, x \mapsto x - 1, x \mapsto (x - 1)(x - 2))$ zwei Basen von V . Dann ist

BSP

Matrix einer lin. Abb.

$$\begin{aligned} \text{Mat}_{B,B}(D) &= \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{pmatrix}, & \text{Mat}_{B',B}(D) &= \begin{pmatrix} 0 & 1 & -3 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{pmatrix}, \\ \text{Mat}_{B,B'}(D) &= \begin{pmatrix} 0 & 1 & 2 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{pmatrix} & \text{und} & \text{Mat}_{B',B'}(D) &= \begin{pmatrix} 0 & 1 & -1 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{pmatrix}. \end{aligned} \quad \clubsuit$$

Wie wir sehen, kann ein und dieselbe lineare Abbildung durch viele verschiedene Matrizen beschrieben werden. Wie hängen diese miteinander zusammen? Dazu erst eine einfache Aussage über Verknüpfungen von linearen Abbildungen.

13.3. Lemma. Seien $g: V \rightarrow V'$ und $f: V' \rightarrow V''$ zwei K -lineare Abbildungen zwischen endlich-dimensionalen Vektorräumen. Seien weiter B eine Basis von V , B' eine Basis von V' und B'' eine Basis von V'' . Dann gilt

LEMMA
Matrix
von $f \circ g$

$$\text{Mat}_{B,B''}(f \circ g) = \text{Mat}_{B',B''}(f) \text{Mat}_{B,B'}(g).$$

Beweis. Das folgt aus der Definition der Matrixmultiplikation. $\square$

Daraus ergibt sich sofort:

13.4. Folgerung. Sei $f: V \rightarrow V'$ eine K -lineare Abbildung zwischen endlich-dimensionalen Vektorräumen. Seien B und $\tilde{B}$ zwei Basen von V und B' und $\tilde{B}'$ zwei Basen von V' . Dann ist

FOLG
Basiswechsel

$$\text{Mat}_{\tilde{B},\tilde{B}'}(f) = \text{Mat}_{B',\tilde{B}'}(\text{id}_{V'}) \text{Mat}_{B,B'}(f) \text{Mat}_{\tilde{B},B}(\text{id}_V).$$

Beweis. Das folgt aus Lemma 13.3 und $f = \text{id}_{V'} \circ f \circ \text{id}_V$. Skizze:

$$\begin{array}{ccc} (V, \tilde{B}) & \xrightarrow{f} & (V', \tilde{B}') \\ \text{id}_V \downarrow & & \uparrow \text{id}_{V'} \\ (V, B) & \xrightarrow{f} & (V', B') \end{array}$$

$\square$

Da id_V und $\text{id}_{V'}$ Isomorphismen sind, sind die *Basiswechselmatrizen* $\text{Mat}_{\tilde{B},B}(\text{id}_V)$ und $\text{Mat}_{B',\tilde{B}'}(\text{id}_{V'})$ invertierbar.

DEF
Basiswechsel-
matrix

Umgekehrt kann jede invertierbare Matrix als eine Basiswechselmatrix auftreten, wobei eine der beiden Basen beliebig vorgegeben werden kann.

13.5. Lemma. Sei K ein Körper, sei $n \in \mathbb{N}$, sei V ein K -Vektorraum mit Basis $B = (b_1, b_2, \dots, b_n)$. Sei weiter $A \in \text{Mat}(n, K)$ invertierbar. Dann gibt es Basen B' und B'' von V , sodass

LEMMA
Basiswechsel-
matrizen

$$A = \text{Mat}_{B,B'}(\text{id}_V) = \text{Mat}_{B'',B}(\text{id}_V)$$

ist.

Beweis. Seien $A = (a_{ij})$ und $B'' = (b''_1, b''_2, \dots, b''_n)$. Die Aussage $A = \text{Mat}_{B'',B}(\text{id}_V)$ bedeutet $b''_j = a_{1j}b_1 + \dots + a_{nj}b_n$. Wir definieren b''_j durch diese Gleichung für $j \in \{1, 2, \dots, n\}$; dann gilt die gewünschte Aussage (B'' ist eine Basis, weil A invertierbar ist: die b_i lassen sich als Linearkombinationen der b''_j ausdrücken, deren Koeffizienten die Einträge von A^{-1} sind).

Es gibt dann auch eine Basis B' , sodass $A^{-1} = \text{Mat}_{B',B}(\text{id}_V)$ ist; damit folgt $A = \text{Mat}_{B,B'}(\text{id}_V)$, denn

$$\text{Mat}_{B,B'}(\text{id}_V) \text{Mat}_{B',B}(\text{id}_V) = \text{Mat}_{B',B'}(\text{id}_V) = I_n$$

nach Lemma 13.3. $\square$

13.6. Satz. Seien K ein Körper und $n \in \mathbb{N}$. Die Menge der invertierbaren Matrizen in $\text{Mat}(n, K)$ bildet eine Gruppe unter der Matrixmultiplikation.

SATZ
Gruppe der
invertierbaren
Matrizen

Beweis. Die Matrixmultiplikation ist assoziativ, die (invertierbare) Einheitsmatrix I_n ist neutrales Element. Jede invertierbare Matrix hat definitionsgemäß eine (selbst invertierbare) Inverse. Es bleibt zu zeigen, dass die Verknüpfung wohldefiniert ist, d.h., dass das Produkt zweier invertierbarer Matrizen wieder invertierbar ist. Das folgt aus

$$(AB)(B^{-1}A^{-1}) = A(BB^{-1})A^{-1} = AA^{-1} = I_n = (B^{-1}A^{-1})(AB);$$

die Inverse von AB ist also $B^{-1}A^{-1}$. $\square$

In jedem Monoid M gilt, dass die Menge der invertierbaren Elemente ein Untermonoid bildet, das eine Gruppe ist. Der Beweis ist identisch.

13.7. Definition. Die Gruppe der invertierbaren Matrizen in $\text{Mat}(n, K)$ heißt *allgemeine lineare Gruppe* und wird mit $\text{GL}(n, K)$ bezeichnet.

DEF
 $\text{GL}(n, K)$

Die Abkürzung „GL“ kommt von englisch *general linear group*. Es ist auch die Notation $\text{GL}_n(K)$ gebräuchlich.

*** 13.8. Satz.** Sei K ein Körper, seien $m, n \in \mathbb{N}$, seien V ein n -dimensionaler und V' ein m -dimensionaler K -Vektorraum und sei $f: V \rightarrow V'$ linear. Seien weiter B eine Basis von V , B' eine Basis von V' und $A = \text{Mat}_{B, B'}(f)$. Dann gilt: Die Menge der Matrizen von f bezüglich beliebiger Basen von V und V' ist genau

SATZ
Matrizen
derselben
lin. Abb.

$$\{PAQ \mid P \in \text{GL}(m, K), Q \in \text{GL}(n, K)\}.$$

Beweis. Nach Folgerung 13.4 und der nachfolgenden Diskussion hat jede Matrix von f die Form PAQ mit invertierbaren Matrizen P und Q . Nach Lemma 13.5 gibt es zu beliebig vorgegebenen invertierbaren Matrizen $P \in \text{GL}(m, K)$ und $Q \in \text{GL}(n, K)$ Basen $\tilde{B}$ von V und $\tilde{B}'$ von V' , sodass $P = \text{Mat}_{B', \tilde{B}'}(\text{id}_{V'})$ und $Q = \text{Mat}_{\tilde{B}, B}(\text{id}_V)$. Dann ist

$$PAQ = \text{Mat}_{B', \tilde{B}'}(\text{id}_{V'}) \text{Mat}_{B, B'}(f) \text{Mat}_{\tilde{B}, B}(\text{id}_V) = \text{Mat}_{\tilde{B}', \tilde{B}}(f)$$

auch eine Matrix von f . $\square$

*** 13.9. Definition.** Seien K ein Körper, $m, n \in \mathbb{N}$ und $A, B \in \text{Mat}(m \times n, K)$. Die Matrizen A und B heißen *äquivalent*, wenn es Matrizen $P \in \text{GL}(m, K)$ und $Q \in \text{GL}(n, K)$ gibt mit $PAQ = B$.

DEF
Äquivalenz
von Matrizen

Zwei Matrizen in $\text{Mat}(m \times n, K)$ sind also genau dann äquivalent, wenn sie dieselbe lineare Abbildung (aber evtl. bezüglich verschiedener Basen) repräsentieren.

*

13.10. Satz. Seien K ein Körper, $m, n \in \mathbb{N}$ und $A, B \in \text{Mat}(m \times n, K)$. Die Matrizen A und B sind genau dann äquivalent, wenn $\text{rk}(A) = \text{rk}(B)$ ist. In diesem Fall sei $r = \text{rk}(A)$; dann sind beide Matrizen äquivalent zur Matrix

$$M_r = \left(\begin{array}{c|c} I_r & \mathbf{0}_{r, n-r} \\ \hline \mathbf{0}_{m-r, r} & \mathbf{0}_{m-r, n-r} \end{array} \right).$$

Dabei steht $\mathbf{0}_{k,l}$ für eine Nullmatrix mit k Zeilen und l Spalten.

SATZ
Klassifikation
von Matrizen
bis auf
Äquivalenz

Beweis. Sei $r = \text{rk}(A)$. Wir zeigen, dass A zu M_r äquivalent ist. Sei $f: K^n \rightarrow K^m$ die zugehörige lineare Abbildung; sie hat Rang r , also ist $\dim \ker(f) = n - r$. Wir wählen eine Basis $B = (b_1, \dots, b_n)$ von K^n , sodass $(b_{r+1}, \dots, b_n)$ eine Basis von $\ker(f)$ ist. Mit $b'_i = f(b_i)$ für $i \in \{1, 2, \dots, r\}$ ist dann $(b'_1, \dots, b'_r)$ eine Basis von $\text{im}(f)$: Wir haben $\dim \text{im}(f) = r$ und

$$\text{im}(f) = \langle f(b_1), \dots, f(b_r), f(b_{r+1}), \dots, f(b_n) \rangle = \langle b'_1, \dots, b'_r, \mathbf{0}, \dots, \mathbf{0} \rangle = \langle b'_1, \dots, b'_r \rangle;$$

$(b'_1, \dots, b'_r)$ ist also ein Erzeugendensystem von $\text{im}(f)$ der richtigen Länge r . Wir ergänzen diese Basis von $\text{im}(f)$ zu einer Basis $B' = (b'_1, \dots, b'_m)$ von K^m . Dann ist $\text{Mat}_{B, B'}(f)$ gerade M_r ; M_r ist damit äquivalent zu A .

Gilt auch $\text{rk}(B) = r$, dann ist B ebenfalls äquivalent zu M_r . Es folgt, dass A und B äquivalent sind: Es gibt $P, P' \in \text{GL}(m, K)$ und $Q, Q' \in \text{GL}(n, K)$ mit $M_r = PAQ = P'BQ'$. Dann ist $B = (P'^{-1}P)A(QQ'^{-1})$.

Umgekehrt gilt $\text{rk}(B) = r = \text{rk}(A)$ für jede zu A äquivalente Matrix B , denn der Rang einer Matrix ist gleich dem Rang jeder von ihr repräsentierten linearen Abbildung. $\square$

Mit den Resultaten aus dem Kleingedruckten von Seite 98 ergibt sich aus Satz 13.10:

Folgerung. Jede Matrix $A \in \text{Mat}(m \times n, K)$ lässt sich durch Zeilen- und Spaltenumformungen in die Matrix M_r mit $r = \text{rk}(A)$ überführen.

FOLG
Zeilen- und
Spaltenumf.

Die Äquivalenz von Matrizen ist ein Beispiel einer *Äquivalenzrelation*. Eine Relation R zwischen Mengen X und Y ist formal eine Teilmenge $R \subset X \times Y$. Ist für $x \in X$ und $y \in Y$ das Paar (x, y) ein Element von R , dann sagt man, x und y stehen in der Relation R zueinander und schreibt auch $x R y$ oder ähnlich. Im Fall $X = Y$ spricht man auch von einer Relation auf X . Eine solche Relation heißt

- reflexiv, wenn $\forall x \in X: x R x$,
- symmetrisch, wenn $\forall x, y \in X: x R y \Rightarrow y R x$, und
- transitiv, wenn $\forall x, y, z \in X: (x R y \wedge y R z) \Rightarrow x R z$.

Eine Relation auf X , die reflexiv, symmetrisch und transitiv ist, ist eine *Äquivalenzrelation* auf X . Beispiele sind die Gleichheitsrelation $x = y$ (das ist die „feinste“ Äquivalenzrelation auf X) oder auch die „Allrelation“ $R = X \times X$ (die „größte“ Äquivalenzrelation auf X).

Lemma. Für Matrizen $A, B \in \text{Mat}(m \times n, K)$ schreiben wir $A \sim B$, wenn A und B äquivalent sind, wenn es also $P \in \text{GL}(m, K)$ und $Q \in \text{GL}(n, K)$ gibt mit $B = PAQ$.

Die Relation $\sim$ ist eine Äquivalenzrelation auf $\text{Mat}(m \times n, K)$.

LEMMA
Äquivalenz
von Matrizen
ist Äqu.rel.

Beweis. Wir müssen die drei Eigenschaften nachprüfen.

- Reflexivität: $A \sim A$, denn man kann $P = I_m$, $Q = I_n$ wählen.

- Symmetrie: Es gelte $A \sim B$; dann gibt es $P \in \text{GL}(m, K)$ und $Q \in \text{GL}(n, K)$ mit $B = PAQ$. Dann sind auch $P^{-1} \in \text{GL}(m, K)$ und $Q^{-1} \in \text{GL}(n, K)$ und es gilt $B = P^{-1}AQ^{-1}$, also $B \sim A$.
- Transitivität: Es gelte $A \sim B$ und $B \sim C$. Dann gibt es $P_1, P_2 \in \text{GL}(m, K)$ und $Q_1, Q_2 \in \text{GL}(n, K)$ mit $B = P_1AQ_1$ und $C = P_2BQ_2$. Es sind $P_2P_1 \in \text{GL}(m, K)$ und $Q_1Q_2 \in \text{GL}(n, K)$ und es gilt $C = (P_2P_1)A(Q_1Q_2)$, also ist $A \sim C$. $\square$

Die wichtigste Eigenschaft einer Äquivalenzrelation auf einer Menge X ist, dass sie zu einer Einteilung von X in sogenannte Äquivalenzklassen führt. Ist $\sim$ eine Äquivalenzrelation auf X und $x \in X$, dann schreiben wir $[x]$ für die Menge $\{y \in X \mid x \sim y\}$ der zu x äquivalenten Elemente von X und nennen $[x]$ die *Äquivalenzklasse von x* . Jedes Element von $[x]$ heißt ein *Repräsentant* der Äquivalenzklasse.

Lemma. Sei $\sim$ eine Äquivalenzrelation auf einer Menge X und sei $x \in X$. Dann sind für $y \in X$ die folgenden Aussagen äquivalent:

LEMMA
Eigensch.
Äqu.rel.

- (1) $x \sim y$.
- (2) $y \in [x]$.
- (3) $[y] \cap [x] \neq \emptyset$.
- (4) $[y] = [x]$.

Insbesondere sind zwei Äquivalenzklassen $[x]$ und $[y]$ entweder gleich oder disjunkt.

Beweis. Die Äquivalenz von (1) und (2) folgt aus der Definition von $[x]$.

„(2) $\Rightarrow$ (3)“: Wegen der Reflexivität von $\sim$ ist $y \in [y]$, also folgt aus $y \in [x]$, dass $y \in [y] \cap [x]$.

„(3) $\Rightarrow$ (4)“: Sei $z \in [y] \cap [x]$ und $w \in [y]$. Dann gilt $y \sim w$, $y \sim z$ und $x \sim z$; mit Symmetrie und Transitivität von $\sim$ folgt daraus $x \sim w$, also $w \in [x]$. Da w beliebig war, gilt $[y] \subset [x]$. Genauso erhalten wir $[x] \subset [y]$.

„(4) $\Rightarrow$ (2)“: Aus $y \in [y]$ und $[y] = [x]$ folgt $y \in [x]$. $\square$

Wir können die Menge der Äquivalenzklassen $X/\sim = \{[x] \mid x \in X\}$ bilden. Dann gibt es eine natürliche (oder „kanonische“) surjektive Abbildung $f: X \rightarrow X/\sim$, $x \mapsto [x]$. Die Urbildmenge $f^{-1}(\{[x]\})$ ist nach dem gerade bewiesenen Lemma genau $[x]$. Umgekehrt führt jede surjektive Abbildung $f: X \rightarrow M$ zu einer Äquivalenzrelation auf X (man sagt auch, f induziert eine Äquivalenzrelation) durch $x \sim y \iff f(x) = f(y)$.

Die Aussage von Satz 13.10 bedeutet dann, dass die Äquivalenz von $m \times n$ -Matrizen mit der durch $\text{Mat}(m \times n, K) \rightarrow \{1, 2, \dots, \min\{m, n\}\}$, $A \mapsto \text{rk}(A)$, induzierten Äquivalenzrelation übereinstimmt und dass M_r ein Repräsentant der durch $\text{rk}(A) = r$ gegebenen Äquivalenzklasse ist.

Wenn man nur Zeilenumformungen betrachtet, dann wird man zwei Matrizen A und B als äquivalent betrachten, wenn $B = PA$ ist mit einer invertierbaren Matrix P , denn Zeilenumformungen entsprechen der Multiplikation von links mit einer invertierbaren Matrix; vergleiche die Diskussion im Kleingedruckten auf Seite 98. Die dortigen Ergebnisse lassen sich so interpretieren, dass jede Äquivalenzklasse dieser Äquivalenzrelation einen eindeutig bestimmten Repräsentanten in reduzierter Zeilenstufenform hat.

14. DIE DETERMINANTE

In diesem Abschnitt führen wir die *Determinante* einer quadratischen Matrix ein. Das ist ein Skalar, der darüber Auskunft gibt, ob die Matrix invertierbar ist oder nicht. Wir definieren die Determinante rekursiv.

*

14.1. Definition. Seien K ein Körper und $A = (a_{ij}) \in \text{Mat}(n, K)$ mit $n \in \mathbb{N}$. Wir definieren die *Determinante* von A , $\det(A)$, rekursiv wie folgt:

DEF
Determinante
einer Matrix

- (1) Im Fall $n = 0$ ist $\det(A) = 1$.
- (2) Im Fall $n > 0$ sei $A_{ij} \in \text{Mat}(n-1, K)$ (für $i, j \in \{1, 2, \dots, n\}$) die Matrix, die aus A entsteht, wenn man die i -te Zeile und die j -te Spalte entfernt. Wir definieren

$$\begin{aligned}\det(A) &= \sum_{j=1}^n (-1)^{j-1} a_{1j} \det(A_{1j}) \\ &= a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + \dots + (-1)^{n-1} a_{1n} \det(A_{1n}).\end{aligned}$$

Für die Determinante ist auch folgende Schreibweise üblich:

$$\det((a_{ij})) = \begin{vmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{vmatrix} \quad \diamond$$

14.2. Beispiele. Für kleine positive Werte von n erhalten wir folgende Formeln:

BSP
Determinante

$$\begin{aligned}\det((a)) &= a \\ \begin{vmatrix} a & b \\ c & d \end{vmatrix} &= ad - bc \\ \begin{vmatrix} a & b & c \\ d & e & f \\ g & h & i \end{vmatrix} &= aei - afh + bfg - bdi + cdh - ceg\end{aligned}$$

Die Formel für die 3×3 -Determinante lässt sich mit Hilfe der „Sarrus-Regel“ merken: Man schreibt die ersten beiden Spalten noch einmal hinter die Matrix und bildet die Summe der Produkte über die nach rechts fallenden Diagonalen minus die Summe der Produkte über die nach rechts steigenden Diagonalen.

$$\begin{vmatrix} a & b & c & a & b \\ d & e & f & d & e \\ g & h & i & g & h \end{vmatrix}$$

Für größere Determinanten gibt es allerdings keine solche Merkregel!



Welche Eigenschaften hat die Determinante?

14.3. Satz. Seien K ein Körper und $n \in \mathbb{N}$. Für $d = \det: \text{Mat}(n, K) \rightarrow K$ gilt:

SATZ
Eigensch.
der Deter-
minante

- (1) $d(A)$ ist linear als Funktion jeder Zeile von A (dabei werden die Einträge der übrigen Zeilen als fest angesehen).
- (2) Hat A zwei gleiche Zeilen, dann ist $d(A) = 0$.
- (3) $d(I_n) = 1$.

Ist $d: \text{Mat}(n, K) \rightarrow K$ eine Abbildung, die (1) und (2) erfüllt, dann gilt:

- (4) Geht A' aus A durch eine elementare Zeilenumformung $\mathbf{I}_i(\lambda)$ hervor, dann gilt $d(A') = \lambda d(A)$.
- (5) Geht A' aus A durch eine elementare Zeilenumformung $\mathbf{II}_{i,j}(\lambda)$ hervor, dann gilt $d(A') = d(A)$.
- (6) Geht A' aus A durch Vertauschen zweier Zeilen hervor, dann gilt $d(A') = -d(A)$.
- (7) Es gilt $d(A) = \det(A)d(I_n)$ für alle $A \in \text{Mat}(n, K)$. Insbesondere ist $\det: \text{Mat}(n, K) \rightarrow K$ die einzige Abbildung, die (1), (2) und (3) erfüllt.

Außerdem gilt

- (8) $\det(A) \neq 0 \iff \text{rk}(A) = n \iff A$ invertierbar.

Beweis. Wir beweisen zuerst (4)–(6).

- (4) Das ist ein Spezialfall von Eigenschaft (1).
- (5) Wir schreiben $d(v_1, \dots, v_n)$ für $d(A)$, wenn A die Zeilen $v_1, \dots, v_n \in K^n$ hat; für $1 \leq i < j \leq n$ sei $d_{ij}(v_i, v_j) = d(v_1, \dots, v_n)$, wobei die v_k mit $k \notin \{i, j\}$ fest gewählt sind. Dann gilt

$$d(A') = d_{ij}(v_i + \lambda v_j, v_j) \stackrel{(1)}{=} d_{ij}(v_i, v_j) + \lambda d_{ij}(v_j, v_j) \stackrel{(2)}{=} d_{ij}(v_i, v_j) = d(A).$$

- (6) In der Notation des Beweises von Teil (5) haben wir

$$\begin{aligned} 0 &\stackrel{(2)}{=} d_{ij}(v_i + v_j, v_i + v_j) \\ &\stackrel{(1)}{=} d_{ij}(v_i, v_i) + d_{ij}(v_i, v_j) + d_{ij}(v_j, v_i) + d_{ij}(v_j, v_j) \\ &\stackrel{(2)}{=} d_{ij}(v_i, v_j) + d_{ij}(v_j, v_i), \end{aligned}$$

$$\text{also ist } d(A') = d_{ij}(v_j, v_i) = -d_{ij}(v_i, v_j) = -d(A).$$

Der Beweis der ersten drei Aussagen für $d = \det$ erfolgt durch Induktion über n . Im Fall $n = 0$ sind die Aussagen trivialerweise richtig. Sei also $n > 0$ und die Aussagen seien für kleinere Werte von n gezeigt.

- (1) $\det(A)$ ist linear in der ersten Zeile von A , denn nach Definition ist $\det(A)$ eine Linearkombination der Einträge der ersten Zeile, deren Koeffizienten nicht von der ersten Zeile abhängen. Sei $k \in \{1, 2, \dots, n-1\}$. Nach Induktionsannahme sind alle $\det(A_{1j})$ linear in der k -ten Zeile von A_{1j} und damit linear in der $(k+1)$ -ten Zeile von A . $\det(A)$ ist somit eine Linearkombination von Abbildungen, die linear als Funktion der $(k+1)$ -ten Zeile von A sind (mit Koeffizienten, die nicht von der $(k+1)$ -ten Zeile abhängen) und somit ebenfalls linear in der $(k+1)$ -ten Zeile von A .

- (2) Sei A eine Matrix, in der die k -te und die l -te Zeile übereinstimmen, wobei $1 \leq k < l \leq n$ ist. Ist $k > 1$, dann stimmt in jeder Matrix A_{1j} die $(k-1)$ -te mit der $(l-1)$ -ten Zeile überein; nach Induktionsannahme gilt $\det(A_{1j}) = 0$ für alle j , also auch $\det(A) = 0$. Es bleibt der Fall $k = 1$ zu betrachten. Falls $l > 2$ ist, dann vertauschen wir die l -te mit der zweiten Zeile. Nach Induktionsannahme und der damit für $d = \det$ auf Matrizen der Größe $n-1$ geltenden Aussage (6) bewirkt das einen Vorzeichenwechsel in allen $\det(A_{1j})$, ändert also nichts daran, ob $\det(A) = 0$ ist oder nicht. Wir können daher annehmen, dass die beiden ersten Zeilen von A gleich sind. Wir schreiben $d_{jm} = d_{mj}$ für die Determinante der Matrix, die aus A durch Streichen der ersten beiden Zeilen und der Spalten j und m entsteht. Dann gilt (unter Beachtung von $a_{2j} = a_{1j}$)

$$\begin{aligned}
 \det(A) &= \sum_{j=1}^n (-1)^{j-1} a_{1j} \det(A_{1j}) \\
 &= \sum_{j=1}^n (-1)^{j-1} a_{1j} \left(\sum_{m=1}^{j-1} (-1)^{m-1} a_{2m} d_{jm} + \sum_{m=j+1}^n (-1)^m a_{2m} d_{jm} \right) \\
 &= \sum_{1 \leq m < j \leq n} (-1)^{j-m} a_{1j} a_{1m} d_{jm} + \sum_{1 \leq j < m \leq n} (-1)^{m-j-1} a_{1j} a_{1m} d_{jm} \\
 &\stackrel{(*)}{=} \sum_{1 \leq j < m \leq n} (-1)^{m-j} a_{1j} a_{1m} d_{jm} + \sum_{1 \leq j < m \leq n} (-1)^{m-j-1} a_{1j} a_{1m} d_{jm} \\
 &= \sum_{1 \leq j < m \leq n} ((-1)^{m-j} + (-1)^{m-j-1}) a_{1j} a_{1m} d_{jm} \\
 &= 0.
 \end{aligned}$$

(Bei $(*)$ haben wir in der ersten Summe j und m vertauscht und dabei ausgenutzt, dass $a_{1m} a_{1j} d_{mj} = a_{1j} a_{1m} d_{jm}$ ist.)

- (3) Nach der rekursiven Definition ist $\det(I_n) = 1 \cdot \det(I_{n-1}) = 1$.

Es verbleiben:

- (7) Sei A' die reduzierte Zeilenstufenform von A . Aus (4), (5) und (6) folgt, dass $d(A) = d_0(A)d(A')$ ist, wobei $d_0(A) \neq 0$ nur von A und nicht von d abhängt: $d_0(A)$ ist das Produkt der Skalare λ^{-1} und -1 , die aus den elementaren Zeilenumformungen $\mathbf{I}_i(\lambda)$ und $\mathbf{III}_{i,j}$ herrühren, die man ausführt, um von A zu A' zu gelangen. Wenn A invertierbar ist, dann ist $A' = I_n$ (Lemma 12.18), also $d(A) = d_0(A)d(I_n)$. Für $d = \det$ ergibt sich $d_0(A) = \det(A)$. Ist A nicht invertierbar, dann hat A' eine Null-Zeile, womit (aus Eigenschaft (1)) $d(A') = 0$ folgt. Für $d = \det$ hat man $\det(A') = 0$. In beiden Fällen erhalten wir $d(A) = \det(A)d(I_n)$ wie behauptet.
- (8) Aus den Teilen (4), (5) und (6) folgt, dass $\det(A)$ genau dann null ist, wenn $\det(A')$ null ist, wobei A' die reduzierte Zeilenstufenform von A ist. Gilt $\text{rk}(A) = n$, dann ist A invertierbar, und nach Lemma 12.18 ist $A' = I_n$ und damit $\det(A') = \det(I_n) = 1 \neq 0$ nach Teil (3). Gilt $\text{rk}(A) < n$, dann hat A' eine Null-Zeile und damit ist $\det(A') = 0$ nach Teil (1). Die zweite Äquivalenz folgt daraus, dass eine lineare Abbildung $K^n \rightarrow K^n$ genau dann surjektiv ist, wenn sie ein Isomorphismus ist; vgl. Folgerung 10.14. $\square$

Wenn wir die Determinante einer $n \times n$ -Matrix A als Funktion der n Zeilen von A betrachten, die selbst Vektoren in K^n sind, dann erhalten wir eine sogenannte *alternierende Multilinearform*. Das ist ein Spezialfall einer multilinearen Abbildung.

Definition. Sei K ein Körper und seien $V_1, V_2, \dots, V_m$ und W K -Vektorräume. Eine Abbildung $f: V_1 \times V_2 \times \dots \times V_m \rightarrow W$ heißt *multilinear*, wenn f in jedem Argument K -linear ist, wenn also gilt

$$f(v_1, \dots, v_{i-1}, \lambda v_i, v_{i+1}, \dots, v_m) = \lambda f(v_1, \dots, v_{i-1}, v_i, v_{i+1}, \dots, v_m)$$

und

$$\begin{aligned} f(v_1, \dots, v_{i-1}, v_i + v'_i, v_{i+1}, \dots, v_m) \\ = f(v_1, \dots, v_{i-1}, v_i, v_{i+1}, \dots, v_m) + f(v_1, \dots, v_{i-1}, v'_i, v_{i+1}, \dots, v_m) \end{aligned}$$

für alle $i \in \{1, 2, \dots, m\}$, $v_j \in V_j$, $\lambda \in K$, $v'_i \in V_i$. Ist $W = K$, dann heißt f auch eine *Multilinearform*.

Eine Multilinearform $f: V^m = V \times V \times \dots \times V \rightarrow K$ heißt *alternierend*, wenn stets $f(v_1, \dots, v_m) = 0$ ist, sobald $v_i = v_j$ ist für gewisse $1 \leq i < j \leq m$. $\diamond$

Aussagen (1) und (2) in Satz 14.3 besagen gerade, dass $\det(A)$ eine alternierende Multilinearform der Zeilen von A ist. Da man einen K -Vektorraum V mit Basis $(b_1, \dots, b_n)$ mit K^n identifizieren kann, hat Aussage (7) in Satz 14.3 die folgende Interpretation:

Satz. Sei V ein K -Vektorraum mit Basis $(b_1, b_2, \dots, b_n)$. Dann gibt es genau eine alternierende Multilinearform $d: V^n \rightarrow K$ mit $d(b_1, b_2, \dots, b_n) = 1$.

Für praktische Zwecke wichtig sind die Aussagen in Satz 14.3, die zeigen, wie sich die Determinante unter elementaren Zeilenumformungen verhält. Das liefert ein praktisches Verfahren zur Berechnung auch größerer Determinanten.

14.4. Beispiel.

$$\begin{aligned} \begin{vmatrix} 1 & -1 & 1 & -1 \\ 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 1 \\ 1 & 2 & 4 & 8 \end{vmatrix} &= \begin{vmatrix} 1 & -1 & 1 & -1 \\ 0 & 1 & -1 & 1 \\ 0 & 2 & 0 & 2 \\ 0 & 3 & 3 & 9 \end{vmatrix} = \begin{vmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & -1 & 1 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 6 & 6 \end{vmatrix} \\ &= 2 \cdot \begin{vmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & -1 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 6 & 6 \end{vmatrix} = 2 \cdot \begin{vmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 6 \end{vmatrix} = 2 \cdot 6 \cdot \begin{vmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{vmatrix} = 12 \quad \clubsuit \end{aligned}$$

Die Eindeutigkeitsaussage (7) in Satz 14.3 ist wichtig, weil sie weitere Eigenschaften der Determinante zur Folge hat.

*

14.5. Satz. Sei K ein Körper, sei $n > 0$ und sei $A = (a_{ij}) \in \text{Mat}(n, K)$. Mit der in Definition 14.1 eingeführten Schreibweise A_{ij} gilt für jedes $i \in \{1, 2, \dots, n\}$:

$$\det(A) = \sum_{j=1}^n (-1)^{j-i} a_{ij} \det(A_{ij}).$$

DEF
multilineare
Abbildung
alternierende
Multilinear-
form

SATZ
Existenz u.
Eindeutigkeit
alternierender
Multilinear-
formen

BSP
Determi-
nanten-
berechnung

SATZ
Entwicklung
der Det.
nach der
 i -ten Zeile

Beweis. Wie im Beweis von Satz 14.3 zeigt man, dass die rechte Seite die Eigenschaften (1), (2) und (3) hat. Wegen der Eindeutigkeit folgt, dass die rechte Seite gleich $\det(A)$ sein muss.

(Alternativ könnte man auch die Zeilen 1 und i vertauschen und Satz 14.3 (6) verwenden.) $\square$

14.6. Beispiel. Die Berechnung der Determinante in Beispiel 14.4 lässt sich vereinfachen, indem man nach der zweiten Zeile entwickelt:

BSP

$$\begin{vmatrix} 1 & -1 & 1 & -1 \\ 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 1 \\ 1 & 2 & 4 & 8 \end{vmatrix} = - \begin{vmatrix} -1 & 1 & -1 \\ 1 & 1 & 1 \\ 2 & 4 & 8 \end{vmatrix} = \dots = 12$$



* **14.7. Satz.** Sei K ein Körper, sei $n \in \mathbb{N}$ und seien $A, B \in \text{Mat}(n, K)$. Dann gilt

$$\det(AB) = \det(A) \det(B).$$

Ist A invertierbar, dann ist $\det(A^{-1}) = \det(A)^{-1}$.

SATZ
Multiplika-
tivität
der Det.

Beweis. Wir fixieren B und betrachten A als variabel. Sei $d_B: \text{Mat}(n, K) \rightarrow K$, $A \mapsto \det(AB)$. Aus den Eigenschaften der Matrixmultiplikation folgt, dass die k -te Zeile von AB sich ergibt als die k -te Zeile von A (als Zeilenvektor) multipliziert mit B . Insbesondere ist die k -te Zeile von AB eine lineare Funktion der k -ten Zeile von A und hängt nicht von den Einträgen in den anderen Zeilen von A ab. Es folgt, dass d_B linear in den Zeilen von A ist. Ebenso gilt, dass aus der Gleichheit der k -ten und der l -ten Zeile von A die entsprechende Aussage für AB folgt. Damit erfüllt d_B auch die Eigenschaft (2) in Satz 14.3. Die Eindeutigkeitsaussage in Satz 14.3 liefert nun $\det(AB) = d_B(A) = \det(A)d_B(I_n) = \det(A)\det(B)$. Die letzte Aussage ergibt sich aus $\det(A)\det(A^{-1}) = \det(I_n) = 1$. $\square$

* **14.8. Satz.** Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Dann gilt

$$\det(A^\top) = \det(A).$$

SATZ
Symmetrie
der Det.

Beweis. Wir müssen zeigen, dass $\det(A^\top)$ die Eigenschaften (1), (2) und (3) aus Satz 14.3 hat. $\det(I_n^\top) = \det(I_n) = 1$ ist klar. Die beiden anderen Aussagen sind dazu äquivalent, dass $\det(A)$ linear in den *Spalten* von A ist und verschwindet, wenn A zwei gleiche *Spalten* hat. Die erste Aussage folgt leicht mit Induktion aus der rekursiven Definition der Determinante, denn für festes k ist jeder Term in der Summe linear in der k -ten Spalte von A (entweder durch a_{1k} oder durch $\det(A_{1j})$). Die zweite Aussage kann wie folgt gezeigt werden: Wenn A zwei gleiche Spalten hat, dann ist $\text{rk}(A) < n$, also $\det(A) = 0$ nach Satz 14.3, Teil (8). $\square$

Daraus folgt, dass man auch *Spaltenumformungen* bei der Berechnung der Determinante verwenden kann, auch mit *Zeilenumformungen* gemischt. Ebenso ergibt sich eine Formel zur Entwicklung der Determinante nach einer Spalte.

- * **14.9. Folgerung.** Sei K ein Körper, sei $n > 0$ und sei $A = (a_{ij}) \in \text{Mat}(n, K)$. **FOLG**
 Mit der in Definition 14.1 eingeführten Schreibweise A_{ij} gilt für jedes $j \in \{1, \dots, n\}$: **Entwicklung der Det. nach der j -ten Spalte**
- $$\det(A) = \sum_{i=1}^n (-1)^{j-i} a_{ij} \det(A_{ij}).$$

Beweis. Das folgt aus Satz 14.5, angewandt auf $A^\top$ und aus $\det(A^\top) = \det(A)$. $\square$

- 14.10. Beispiel.** Eine Matrix $A \in \text{Mat}(n, \mathbb{R})$ mit $AA^\top = I_n$ heißt *orthogonal*. **BSP**
 Was kann man über $\det(A)$ sagen?

Es gilt

$$1 = \det(I_n) = \det(AA^\top) = \det(A) \det(A^\top) = \det(A)^2,$$

also ist $\det(A) = \pm 1$. $\clubsuit$

- 14.11. Beispiel.** Wir berechnen die Determinante aus Beispiel 14.4 noch einmal. **BSP**
 Determinantenberechnung

$$\begin{vmatrix} 1 & -1 & 1 & -1 \\ 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 1 \\ 1 & 2 & 4 & 8 \end{vmatrix} = - \begin{vmatrix} -1 & 1 & -1 \\ 1 & 1 & 1 \\ 2 & 4 & 8 \end{vmatrix} = - \begin{vmatrix} -1 & 1 & 0 \\ 1 & 1 & 0 \\ 2 & 4 & 6 \end{vmatrix} = -6 \cdot \begin{vmatrix} -1 & 1 \\ 1 & 1 \end{vmatrix} \\ = -6((-1) \cdot 1 - 1 \cdot 1) = 12$$

(Entwicklung nach der zweiten Zeile, elementare Spaltenumformung $\mathbf{II}_{3,1}(-1)$, Entwicklung nach der dritten Spalte, Formel für 2×2 -Determinante). $\clubsuit$

Der Vollständigkeit halber halten wir noch fest, wie sich die Determinante bei Skalierung der ganzen Matrix ändert.

- 14.12. Folgerung.** Seien K ein Körper, $n \in \mathbb{N}$, $\lambda \in K$ und $A \in \text{Mat}(n, K)$. **FOLG**
 Dann gilt $\det(\lambda A)$

$$\det(\lambda A) = \lambda^n \det(A).$$

Beweis. Das folgt daraus, dass $\det(A)$ in jeder Spalte von A linear ist:

Mit $A = (\mathbf{a}_1 \mid \mathbf{a}_2 \mid \dots \mid \mathbf{a}_n)$ ergibt sich

$$\begin{aligned} \det(\lambda A) &= \det(\lambda \mathbf{a}_1 \mid \lambda \mathbf{a}_2 \mid \dots \mid \lambda \mathbf{a}_n) \\ &= \lambda \det(\mathbf{a}_1 \mid \lambda \mathbf{a}_2 \mid \dots \mid \lambda \mathbf{a}_n) = \dots = \lambda^n \det(\mathbf{a}_1 \mid \mathbf{a}_2 \mid \dots \mid \mathbf{a}_n). \quad \square \end{aligned}$$

Die Entwicklung der Determinante nach Zeilen und Spalten führt zu folgender „Formel“ für die Inverse einer Matrix.

- * **14.13. Definition.** Seien K ein Körper, $n > 0$ und $A \in \text{Mat}(n, K)$. Die Matrix $\tilde{A} \in \text{Mat}(n, K)$, deren Eintrag in der i -ten Zeile und j -ten Spalte durch $(-1)^{i-j} \det(A_{ji})$ (nicht A_{ij} !) gegeben ist, heißt die *adjunkte* (oder auch *adjungierte*) Matrix zu A . **DEF**
 Adjunkte Matrix $\diamond$

* 14.14. **Satz.** Seien K ein Körper, $n > 0$ und $A \in \text{Mat}(n, K)$. Dann gilt

$$A\tilde{A} = \tilde{A}A = \det(A)I_n.$$

Ist A invertierbar, dann ist $A^{-1} = \det(A)^{-1}\tilde{A}$.

SATZ
Adjunkte
Matrix

Beweis. Der Eintrag an der Stelle (i, k) im Produkt $A\tilde{A}$ ist

$$\sum_{j=1}^n a_{ij}(-1)^{j-k} \det(A_{kj}).$$

Im Fall $k = i$ ergibt das $\det(A)$ nach dem Satz 14.5 über die Entwicklung der Determinante nach der i -ten Zeile. Im Fall $k \neq i$ ergibt sich analog die Entwicklung nach der k -ten Zeile der Determinante der Matrix, die aus A entsteht, wenn man die k -te Zeile durch die i -te ersetzt. Da diese Matrix zwei gleiche Zeilen hat, ist ihre Determinante null. Das zeigt $A\tilde{A} = \det(A)I_n$. Die Aussage $\tilde{A}A = \det(A)I_n$ sieht man analog unter Verwendung von Folgerung 14.9. Die letzte Aussage folgt durch Multiplikation mit $\det(A)^{-1}A^{-1}$. $\square$

14.15. **Beispiel.** Für $n = 2$ bekommen wir die Formel

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

BSP

Inverse einer
2 × 2-Matrix



Hier ist ein Beispiel für eine allgemeine Formel für eine spezielle Determinante:

14.16. **Satz.** Sei K ein Körper, sei $n \in \mathbb{N}$, seien $a_1, a_2, \dots, a_n \in K$ und sei A folgende „Vandermonde-Matrix“ zu $a_1, a_2, \dots, a_n$:

$$A = (a_i^{j-1})_{1 \leq i, j \leq n} = \begin{pmatrix} 1 & a_1 & a_1^2 & \cdots & a_1^{n-1} \\ 1 & a_2 & a_2^2 & \cdots & a_2^{n-1} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & a_n & a_n^2 & \cdots & a_n^{n-1} \end{pmatrix} \in \text{Mat}(n, K).$$

SATZ

Vandermonde-
Determinante

Dann ist

$$\det(A) = \prod_{1 \leq i < j \leq n} (a_j - a_i).$$

Beweis. Durch Induktion nach n . Der Fall $n = 0$ ist klar ($\det(A) = 1$ und das Produkt ist leer). Sei also $n > 0$ und die Behauptung für $n - 1$ bewiesen. Wir subtrahieren die erste Zeile von den übrigen und entwickeln nach der ersten Spalte; das liefert

$$\det(A) = \begin{vmatrix} a_2 - a_1 & a_2^2 - a_1^2 & \cdots & a_2^{n-1} - a_1^{n-1} \\ \vdots & \vdots & & \vdots \\ a_n - a_1 & a_n^2 - a_1^2 & \cdots & a_n^{n-1} - a_1^{n-1} \end{vmatrix}.$$

Nun ist $x^m - y^m = (x - y)(x^{m-1} + x^{m-2}y + \dots + xy^{m-2} + y^{m-1})$. Wir können aus der ersten, $\dots$, $(n - 1)$ -ten Zeile also jeweils einen Faktor $(a_2 - a_1), \dots, (a_n - a_1)$ herausziehen:

$$\det(A) = \prod_{j=2}^n (a_j - a_1) \cdot \begin{vmatrix} 1 & a_2 + a_1 & \cdots & a_2^{n-2} + a_1 a_2^{n-3} + \dots + a_1^{n-2} \\ \vdots & \vdots & & \vdots \\ 1 & a_n + a_1 & \cdots & a_n^{n-2} + a_1 a_n^{n-3} + \dots + a_1^{n-2} \end{vmatrix}.$$

Durch Subtraktion von a_1 -mal der vorletzten Spalte von der letzten, dann Subtraktion von a_1 -mal der drittletzten Spalte von der vorletzten, $\dots$, a_1 -mal der ersten Spalte von der zweiten erhält man aus der verbliebenen Matrix die Vandermonde-Matrix zu $a_2, a_3, \dots, a_n$. Die Behauptung folgt dann aus der Induktionsvoraussetzung. $\square$

Die Vandermonde-Matrix ist genau die Matrix der linearen Abbildung

$$\phi: P_{<n} \longrightarrow \mathbb{R}^n, \quad f \longmapsto (f(a_1), f(a_2), \dots, f(a_n))$$

(wo $P_{<n}$ der $\mathbb{R}$ -Vektorraum der Polynomfunktionen vom Grad $< n$ ist), die bei der Interpolation von gegebenen Werten durch Polynome eine Rolle spielt, bezüglich der Basis $(x \mapsto 1, x \mapsto x, x \mapsto x^2, \dots, x \mapsto x^{n-1})$ von $P_{<n}$ und der Standardbasis von $\mathbb{R}^n$, siehe Beispiele 9.18 und 10.15.

Nach Vandermonde ist auch die „Vandermonde-Identität“

$$\sum_{j=0}^k \binom{m}{j} \binom{n}{k-j} = \binom{m+n}{k}$$

benannt.

Wir werden uns jetzt mit einer Verallgemeinerung der Formeln für die Determinante wie in Beispiel 14.2 beschäftigen. Diese Formeln erhält man aus der rekursiven Definition der Determinante wie in Definition 14.1. Das Resultat ist eine Summe von Termen der Form $\pm a_{1,\sigma(1)} a_{2,\sigma(2)} \cdots a_{n,\sigma(n)}$, wobei die Spaltenindizes $\sigma(1), \sigma(2), \dots, \sigma(n)$ paarweise verschieden sind (denn jede „verbrauchte“ Spalte wird in der weiteren Entwicklung entfernt). Die Abbildung $\sigma: \{1, 2, \dots, n\} \rightarrow \{1, 2, \dots, n\}$ ist demnach bijektiv, also eine Permutation. Wie man sich leicht überlegt, kommt auch jede Permutation in der Entwicklung der Determinante vor. Wir erinnern uns daran, dass die Permutationen von $\{1, 2, \dots, n\}$ die Elemente der *symmetrischen Gruppe* S_n sind; die Verknüpfung in dieser Gruppe ist die Komposition von Abbildungen. Damit haben wir Folgendes gezeigt:



G.W. Leibniz
(1646–1716)

* 14.17. **Satz.** Sei $n \in \mathbb{N}$. Es gibt eine eindeutig bestimmte Abbildung

$$\varepsilon: S_n \rightarrow \{-1, 1\},$$

sodass für jeden Körper K und jede Matrix $A = (a_{ij}) \in \text{Mat}(n, K)$ gilt

$$\det(A) = \sum_{\sigma \in S_n} \varepsilon(\sigma) a_{1,\sigma(1)} a_{2,\sigma(2)} \cdots a_{n,\sigma(n)}.$$

SATZ
Leibniz-
Formel

Diese Formel hat $\#S_n = n! = 1 \cdot 2 \cdot 3 \cdots (n-1) \cdot n$ Terme und taugt damit außer für sehr kleine Werte von n nicht zur praktischen Berechnung der Determinante! Sie ist aber nützlich für theoretische Überlegungen. Zum Beispiel folgt sofort, dass die Determinante einer Matrix mit ganzzahligen Einträgen eine ganze Zahl ist.



14.18. **Definition.** $\varepsilon(\sigma) \in \{-1, 1\}$ heißt das *Signum* oder *Vorzeichen* der Permutation $\sigma \in S_n$. σ heißt *gerade*, wenn $\varepsilon(\sigma) = 1$ und *ungerade*, wenn $\varepsilon(\sigma) = -1$ ist. $\diamond$

DEF
Signum einer
Permutation

Um etwas über diese Vorzeichenfunktion herauszufinden, führen wir die zu σ gehörende Permutationsmatrix ein.

14.19. Definition. Seien K ein Körper, $n \in \mathbb{N}$ und $\sigma \in S_n$. Dann bezeichnen wir mit $P(\sigma)$ die Matrix $(\delta_{i,\sigma(j)})_{1 \leq i,j \leq n} \in \text{Mat}(n, K)$ und nennen $P(\sigma)$ die zu σ gehörende *Permutationsmatrix*. $\diamond$

DEF
Permutations-
matrix

Die Einträge von $P(\sigma)$ sind 1 an Positionen der Form $(\sigma(j), j)$ und sonst 0. Das bedeutet, dass in der j -ten Spalte von $P(\sigma)$ gerade der Standard-Basisvektor $\mathbf{e}_{\sigma(j)}$ steht; die zu P gehörende lineare Abbildung $K^n \rightarrow K^n$ ist also genau die, die die Standardbasis entsprechend der Permutation σ vertauscht: $P(\sigma)\mathbf{e}_j = \mathbf{e}_{\sigma(j)}$, wenn wir $\mathbf{e}_j$ als Spaltenvektor auffassen.

14.20. Lemma. Seien K ein Körper und $n \in \mathbb{N}$.

LEMMA
Eigensch.
Permutations-
matrix

- (1) Für $\sigma, \tau \in S_n$ gilt $P(\sigma \circ \tau) = P(\sigma)P(\tau)$.
- (2) Für $\sigma \in S_n$ gilt $\varepsilon(\sigma) = \det(P(\sigma))$.
- (3) Für $\sigma, \tau \in S_n$ gilt $\varepsilon(\sigma \circ \tau) = \varepsilon(\sigma)\varepsilon(\tau)$.
- (4) Ist σ eine Transposition (also eine Permutation, die zwei Elemente vertauscht und alle anderen nicht ändert), dann ist $\varepsilon(\sigma) = -1$.

Beweis.

- (1) Für alle $j \in \{1, 2, \dots, n\}$ gilt

$$P(\sigma)P(\tau)\mathbf{e}_j = P(\sigma)\mathbf{e}_{\tau(j)} = \mathbf{e}_{\sigma(\tau(j))} = \mathbf{e}_{(\sigma \circ \tau)(j)} = P(\sigma \circ \tau)\mathbf{e}_j,$$

also ist $P(\sigma \circ \tau) = P(\sigma)P(\tau)$.

- (2) Der einzige von null verschiedene Term in der Formel 14.17 für $\det(P(\sigma^{-1}))$ ist $\varepsilon(\sigma)\delta_{1,1} \cdots \delta_{n,n} = \varepsilon(\sigma)$. Es folgt (wegen $P(\sigma)^{-1} = P(\sigma^{-1})$ nach Teil (1) und $\varepsilon(\sigma) = \pm 1$)

$$\det(P(\sigma)) = \det(P(\sigma)^{-1})^{-1} = \det(P(\sigma^{-1}))^{-1} = \varepsilon(\sigma)^{-1} = \varepsilon(\sigma).$$

- (3) Das folgt aus (1) und (2) und der Multiplikativität der Determinante.
- (4) In diesem Fall erhält man $P(\sigma)$ aus I_n durch Vertauschen zweier Spalten (oder Zeilen), also ist $\varepsilon(\sigma) = \det(P(\sigma)) = -\det(I_n) = -1$. $\square$

Da sich (wie man sich leicht überlegen kann) jede Permutation als Komposition von Transpositionen schreiben lässt, ist ε durch die Eigenschaften (3) und (4) in Lemma 14.20 eindeutig festgelegt: Ist σ Komposition einer geraden Anzahl von Transpositionen, dann ist σ gerade, sonst ungerade.

Es gibt eine Art Formel für $\varepsilon(\sigma)$. Dazu eine kleine Definition:

Definition. Sei $\sigma \in S_n$. Ein Paar (i, j) mit $1 \leq i < j \leq n$ heißt *Fehlstand* von σ , wenn $\sigma(i) > \sigma(j)$ ist. $\diamond$

DEF
Fehlstand

Dann gilt der folgende Satz.

Satz. Sei $\sigma \in S_n$ und sei m die Anzahl der Fehlstände von σ . Dann ist $\varepsilon(\sigma) = (-1)^m$.

SATZ
Signum und
Fehlstände

Beweis. Sei $\varepsilon'(\sigma)$ die durch $(-1)^{\text{Anzahl Fehlstände von } \sigma}$ definierte Funktion. Die Transposition τ , die k und l vertauscht (mit $k < l$), hat genau $m = 1 + 2(l - k - 1)$ Fehlstände (nämlich (k, l) sowie (k, j) und (j, l) für alle $k < j < l$). Da m ungerade ist, ist $\varepsilon'(\tau) = (-1)^m = -1 = \varepsilon(\tau)$.

Außerdem gilt

$$\varepsilon'(\sigma) = \prod_{1 \leq i < j \leq n} \frac{\sigma(j) - \sigma(i)}{j - i},$$

denn das rechts stehende Produkt hat Betrag 1 (jeder Faktor im Nenner tritt bis aufs Vorzeichen auch im Zähler auf) und $\sigma(j) - \sigma(i)$ ist genau dann negativ, wenn (i, j) ein Fehlstand von σ ist. Es folgt

$$\begin{aligned} \varepsilon'(\sigma \circ \tau) &= \prod_{1 \leq i < j \leq n} \frac{\sigma(\tau(j)) - \sigma(\tau(i))}{j - i} \\ &= \prod_{1 \leq i < j \leq n} \frac{\sigma(\tau(j)) - \sigma(\tau(i))}{\tau(j) - \tau(i)} \prod_{1 \leq i < j \leq n} \frac{\tau(j) - \tau(i)}{j - i} \\ &= \prod_{1 \leq k < l \leq n} \frac{\sigma(l) - \sigma(k)}{l - k} \prod_{1 \leq i < j \leq n} \frac{\tau(j) - \tau(i)}{j - i} = \varepsilon'(\sigma) \varepsilon'(\tau). \end{aligned}$$

Die Funktion ε' hat also die Eigenschaften (3) und (4) aus Lemma 14.20 und muss daher mit ε übereinstimmen. $\square$

Zum Abschluss dieses Kapitels wollen wir noch die geometrische Bedeutung der Determinante untersuchen. Wir betrachten den $\mathbb{R}^n$ und definieren erst einmal den Begriff der (positiven oder negativen) Orientierung einer Basis.

*

14.21. Definition. Sei $(\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ eine Basis des $\mathbb{R}^n$ und $A = (\mathbf{b}_1 \mid \dots \mid \mathbf{b}_n)$ die Matrix, deren Spalten die Basisvektoren sind. Wir sagen, die Basis sei *positiv orientiert*, wenn $\det(A) > 0$ ist, und *negativ orientiert*, wenn $\det(A) < 0$ ist. $\diamond$

DEF
Orientierung
einer Basis

Die Standardbasis ist positiv orientiert. Im Fall $n = 2$ ist eine Basis genau dann positiv orientiert, wenn der gegen den Uhrzeigersinn gemessene Winkel vom ersten zum zweiten Basisvektor kleiner ist als π ($= 180^\circ$).

Der Vergleich der Orientierung einer Basis und ihres Bildes führt zum Begriff der orientierungserhaltenden bzw. -umkehrenden linearen Abbildung.

14.22. Definition. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ ein Automorphismus (also ein invertierbarer Endomorphismus). Dann heißt f *orientierungserhaltend*, wenn f positiv orientierte Basen auf positiv orientierte Basen abbildet, und *orientierungsumkehrend*, wenn f positiv orientierte Basen auf negativ orientierte Basen abbildet. $\diamond$

DEF
orientierungs-
erhaltend,
-umkehrend

Man sieht leicht, dass f genau dann orientierungserhaltend (-umkehrend) ist, wenn $\det(A) > 0$ (< 0) ist, wobei A die zu f gehörende Matrix ist: Ist B die Matrix, deren Spalten die Vektoren einer positiv orientierten Basis bilden, dann sind die Spalten von AB die Bilder der Basisvektoren. Wegen $\det(B) > 0$ gilt dann

$$\det(AB) = \det(A) \det(B) > 0 \iff \det(A) > 0.$$

Wir wollen jetzt das Volumen von „linear verzerrten Würfeln“ betrachten. In der Ebene $\mathbb{R}^2$ sind das Parallelogramme. Allgemeiner definieren wir:

14.23. Definition. Ein *Parallelotop* im $\mathbb{R}^n$ ist die Menge

$$P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) = \{t_1 \mathbf{x}_1 + t_2 \mathbf{x}_2 + \dots + t_n \mathbf{x}_n \mid t_1, t_2, \dots, t_n \in [0, 1]\} \subset \mathbb{R}^n$$

für ein n -Tupel $(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ von Vektoren im $\mathbb{R}^n$. Das Parallelotop heißt *ausgeartet*, wenn die es aufspannenden Vektoren $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ linear abhängig sind (dann ist $P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ im echten Untervektorraum $\langle \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \rangle$ von $\mathbb{R}^n$ enthalten). $\diamond$

DEF
Parallelotop

Wir wollen jetzt untersuchen, wie man das „orientierte Volumen“ solcher Parallelotope definieren kann. Es sollte folgende Eigenschaften haben:

- (1) $\text{vol } P(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n) = 1$ (der n -dimensionale Einheitswürfel hat Volumen 1).
- (2) $\text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ ist positiv (bzw. negativ), wenn $(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ eine positiv (bzw. negativ) orientierte Basis von $\mathbb{R}^n$ ist.
- (3) $\text{vol } P(\mathbf{x}_1, \dots, \mathbf{x}_{j-1}, \lambda \mathbf{x}_j, \mathbf{x}_{j+1}, \dots, \mathbf{x}_n) = \lambda \text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ für $\lambda \in \mathbb{R}$.
- (4) $\text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) = 0$, wenn $P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ ausgeartet ist.
- (5) $\text{vol } P(\mathbf{x}_1, \dots, \mathbf{x}_{j-1}, \mathbf{x}_j + \mathbf{x}'_j, \mathbf{x}_{j+1}, \dots, \mathbf{x}_n)$
 $= \text{vol } P(\mathbf{x}_1, \dots, \mathbf{x}_{j-1}, \mathbf{x}_j, \mathbf{x}_{j+1}, \dots, \mathbf{x}_n)$
 $+ \text{vol } P(\mathbf{x}_1, \dots, \mathbf{x}_{j-1}, \mathbf{x}'_j, \mathbf{x}_{j+1}, \dots, \mathbf{x}_n).$

Diese letzte Eigenschaft kann man sich plausibel machen, wenn man an die Formel „Grundfläche mal Höhe“ denkt: Die Höhe von $\mathbf{x}_j + \mathbf{x}'_j$ über der „Grundfläche“, die durch das von den übrigen Vektoren aufgespannte Parallelotop gegeben ist, ist die Summe der (orientierten) Höhen von $\mathbf{x}_j$ und $\mathbf{x}'_j$.

Man beachte, dass das orientierte Volumen nicht nur von $P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ als Teilmenge des $\mathbb{R}^n$ abhängt, sondern auch von der Reihenfolge der aufspannenden Vektoren (genauer gesagt hängt das *Vorzeichen* davon ab). Daher betrachten wir das orientiert Volumen als Funktion des n -Tupels $(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$.

14.24. Satz. Die einzige Abbildung vol von der Menge der n -Tupel $(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ im $\mathbb{R}^n$ nach $\mathbb{R}$, die die obigen Eigenschaften hat, ist

$$\text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) = \det(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n),$$

wobei $\det(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)$ für die Determinante der Matrix steht, deren Spalten die Vektoren $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ sind.

SATZ
Determinante
ist orientiertes
Volumen

Beweis. Die Determinante hat jedenfalls die geforderten Eigenschaften (Satz 14.3 und Satz 14.8, sowie Definition 14.21). Aus der Eindeutigkeitsaussage in Satz 14.3, zusammen mit Satz 14.8, der besagt, dass die analoge Aussage auch für Spalten statt Zeilen gilt, folgt, dass die Determinante die einzige Abbildung ist, die die Eigenschaften (1), (3), (4) und (5) hat. $\square$

Man kann also mit Hilfe der Determinante Volumen messen.

14.25. Beispiel. Die Fläche des Dreiecks mit den Eckpunkten (x_1, y_1) , (x_2, y_2) und (x_3, y_3) ist

$$\frac{1}{2} \left| \det \begin{pmatrix} x_1 & x_2 & x_3 \\ y_1 & y_2 & y_3 \\ 1 & 1 & 1 \end{pmatrix} \right|.$$

BSP
Fläche eines
Dreiecks

Wir können das Dreieck so verschieben, dass die erste Ecke im Ursprung zu liegen kommt. Dann ist die gesuchte Fläche die Hälfte der Fläche des von $(x_2 - x_1, y_2 - y_1)$

und $(x_3 - x_1, y_3 - y_1)$ aufgespannten Parallelogramms. Die orientierte Fläche dieses Parallelogramms ist

$$\det \begin{pmatrix} x_2 - x_1 & x_3 - x_1 \\ y_2 - y_1 & y_3 - y_1 \end{pmatrix}.$$

Die obige Determinante lässt sich durch die Spaltenoperationen $\mathbf{II}_{2,1}(-1)$ und $\mathbf{II}_{3,1}(-1)$ und nachfolgender Entwicklung nach der dritten Zeile auf diese Form bringen. Durch den Absolutbetrag erhalten wir die Fläche statt der orientierten Fläche. ♣

Aus der Multiplikativität der Determinante folgt eine Interpretation der Determinante eines Endomorphismus, die für Anwendungen in der Analysis (z.B. die Transformationsformel für mehrdimensionale Integrale) relevant ist.

*

14.26. Satz. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ linear mit zugehöriger Matrix A und seien außerdem $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^n$. Dann gilt

$$\text{vol } P(f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_n)) = \det(A) \text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n).$$

SATZ
Determinante
ist Skalierung
des Volumens

Die Determinante eines Endomorphismus gibt also an, mit welchem Faktor das orientierte Volumen bei seiner Anwendung multipliziert wird.

Als Teilmengen des $\mathbb{R}^n$ gilt

$$f(P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n)) = P(f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_n))$$

(da f linear ist, kann die Linearkombination hineingezogen werden). Das (nicht orientierte) Volumen skaliert sich dann mit $|\det(A)|$.

Beweis. Sei X die Matrix, deren Spalten die Vektoren $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ sind. Es gilt

$$\begin{aligned} \text{vol } P(f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_n)) &= \det(f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_n)) \\ &= \det(A\mathbf{x}_1, A\mathbf{x}_2, \dots, A\mathbf{x}_n) \\ &= \det(AX) = \det(A) \det(X) \\ &= \det(A) \text{vol } P(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n). \quad \square \end{aligned}$$

15. EIGENWERTE UND EIGENVEKTOREN

Im vorletzten Abschnitt haben wir einen Klassifikationssatz bewiesen (Satz 13.10). Man kann ihn so interpretieren, dass die einzige Eigenschaft, die lineare Abbildungen zwischen zwei gegebenen endlich-dimensionalen Vektorräumen voneinander unterscheidet, der *Rang* ist: Ist $f: V \rightarrow W$ linear mit $\text{rk}(f) = r$, dann können wir Basen von V und W so wählen, dass f durch die Matrix M_r gegeben ist, und M_r hängt nur von r (und den Dimensionen von V und W) ab.

Wir werden jetzt statt linearen Abbildungen zwischen verschiedenen Vektorräumen V und W *Endomorphismen* $f: V \rightarrow V$ betrachten. Da wir nur einen Vektorraum haben, können wir auch nur *eine* Basis wählen. Wir haben also deutlich weniger Spielraum, was sich in einem erheblich schwierigeren Klassifikationsproblem niederschlägt.

Natürlich kann man lineare Abbildungen $f: V \rightarrow V$ auch als Spezialfall von linearen Abbildungen $V \rightarrow W$ betrachten, wo „zufällig“ $W = V$ ist. Dann wird man die Wahl von verschiedenen Basen von V auf der Quell- und der Zielseite zulassen. Zum Beispiel erhält man so die Basiswechselmatrizen $\text{Mat}_{B,B'}(\text{id}_V)$. Auf der anderen Seite geht so aber die Information verloren, dass es sich wirklich auf beiden Seiten um *denselben* Vektorraum handelt und nicht um zwei Vektorräume, die zufällig isomorph sind (d.h., dieselbe Dimension haben). Für das Klassifikationsproblem, das wir in diesem Abschnitt (und dann weiter in der *Linearen Algebra II*) studieren wollen, ist es aber wesentlich, dass f als Endomorphismus von V betrachtet wird. Anderenfalls wäre eine Aussage der Form $f(v) = \lambda v$ (siehe Definition 15.3 unten) nicht sinnvoll, bzw. sie würde sich nicht auf die f beschreibenden Matrizen übertragen.

Wir schreiben erst einmal auf, wie die Matrizen von f bezüglich verschiedener Basen von V miteinander zusammenhängen.

15.1. Satz. Sei V ein K -Vektorraum mit Basis $B = (b_1, b_2, \dots, b_n)$ und sei f ein Endomorphismus von V . Sei weiter $A = \text{Mat}_{B,B}(f) \in \text{Mat}(n, K)$ die Matrix von f bezüglich B . Dann ist

$$\{\text{Mat}_{B',B'}(f) \mid B' \text{ Basis von } V\} = \{PAP^{-1} \mid P \in \text{GL}(n, K)\}.$$

SATZ
Matrizen
eines Endo-
morphismus

Beweis. Es ist $\text{Mat}_{B',B'}(f) = \text{Mat}_{B,B'}(\text{id}_V) \text{Mat}_{B,B}(f) \text{Mat}_{B',B}(\text{id}_V) = PAP^{-1}$ mit $P = \text{Mat}_{B,B'}(\text{id}_V) \in \text{GL}(n, K)$. Umgekehrt lässt sich jede Matrix $P \in \text{GL}(n, K)$ in dieser Form schreiben (Folgerung 13.4). $\square$

* **15.2. Definition.** Seien K ein Körper, $n \in \mathbb{N}$. Zwei Matrizen $A, A' \in \text{Mat}(n, K)$ heißen *ähnlich*, wenn es eine Matrix $P \in \text{GL}(n, K)$ gibt mit $A' = PAP^{-1}$. $\diamond$

DEF
Ähnlichkeit
von Matrizen

Analog wie für die Äquivalenz von Matrizen zeigt man, dass die Ähnlichkeit von Matrizen eine Äquivalenzrelation ist.

Wenn A eine Matrix eines Endomorphismus f von V ist, dann sind die Matrizen von f bezüglich beliebiger Basen von V also gerade die zu A ähnlichen Matrizen.

Die Klassifikation von Matrizen bis auf Ähnlichkeit (und damit die Klassifikation der Endomorphismen endlich-dimensionaler Vektorräume) ist relativ kompliziert. Sie wird durch die *Jordan-Normalform* geleistet, die wir im nächsten Semester besprechen werden. Hier werden wir uns erst einmal auf die Diskussion einfacher „Invarianten“ (also Daten, die nur von f und nicht von der Basis abhängen) beschränken.

Die Idee ist, den Endomorphismus $f: V \rightarrow V$ mit anderen besonders einfachen Endomorphismen zu vergleichen. Die einfachsten Endomorphismen sind sicher die Multiplikationen mit einem Skalar $\lambda \in K$: $v \mapsto \lambda v$. Wir können uns fragen, ob es Elemente von V gibt, die sich unter f und dieser Abbildung gleich verhalten. Das führt auf folgende Definition.

- * **15.3. Definition.** Seien V ein K -Vektorraum und $f \in \text{End}(V)$. Ein Skalar $\lambda \in K$ heißt ein *Eigenwert* von f , wenn es einen Vektor $\mathbf{0} \neq v \in V$ gibt, sodass $f(v) = \lambda v$ ist. Jeder solche Vektor heißt ein *Eigenvektor* von f zum Eigenwert λ . $\diamond$

DEF
Eigenwert
Eigenvektor

Man beachte die Bedingung $v \neq \mathbf{0}$! Ohne sie wäre die Definition sinnlos, weil dann jedes λ ein Eigenwert wäre (denn $f(\mathbf{0}) = \mathbf{0} = \lambda \mathbf{0}$).



- * **15.4. Definition.** Seien V ein K -Vektorraum, $\lambda \in K$ und $f \in \text{End}(V)$. Der Untervektorraum

DEF
Eigenraum
geom. Vielfachheit

$$E_\lambda(f) = \{v \in V \mid f(v) = \lambda v\} = \ker(\lambda \text{id}_V - f)$$

von V heißt der λ -Eigenraum oder der Eigenraum zum Eigenwert λ von f .

Die Dimension $\dim E_\lambda(f)$ des λ -Eigenraums von f heißt die *geometrische Vielfachheit* des Eigenwerts λ von f . $\diamond$

$E_\lambda(f)$ besteht also aus dem Nullvektor und den Eigenvektoren zum Eigenwert λ . λ ist genau dann ein Eigenwert von f , wenn $E_\lambda(f) \neq \{\mathbf{0}\}$, also die geometrische Vielfachheit positiv ist.

15.5. Beispiele.

BSP
Eigenwerte
Eigenräume

- (1) $E_0(f)$ ist gerade der Kern von f . Null ist also genau dann ein Eigenwert von f , wenn f nicht injektiv ist. Ist V endlich-dimensional, dann ist das auch dazu äquivalent, dass f kein Isomorphismus ist.
- (2) Sei $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $(x, y) \mapsto (y, x)$. Dann hat f die Eigenwerte 1 und -1 , denn $v_1 = (1, 1) \in E_1(f)$ und $v_{-1} = (1, -1) \in E_{-1}(f)$. Man sieht leicht, dass beide Eigenräume eindimensional sind.
- (3) Sei $f: \mathcal{C}^\infty(\mathbb{R}) \rightarrow \mathcal{C}^\infty(\mathbb{R})$, $h \mapsto h'$. Dann hat f jedes $\lambda \in \mathbb{R}$ als Eigenwert, und es gilt $E_\lambda(f) = \langle x \mapsto e^{\lambda x} \rangle$. (Beweis wie in Beispiel 12.11.) $\clubsuit$

Wir werden sehen, dass die Situation von Beispiel (2) oben für Endomorphismen von endlich-dimensionalen Vektorräumen recht typisch ist. Wir zeigen erst einmal, dass es nicht zu viele Eigenwerte geben kann.

- * **15.6. Satz.** Seien V ein K -Vektorraum und $f \in \text{End}(V)$. Seien $\lambda_1, \dots, \lambda_m \in K$ paarweise verschieden und für $j \in \{1, 2, \dots, m\}$ sei $v_j \in V$ ein Eigenvektor von f zum Eigenwert λ_j . Dann ist $(v_1, v_2, \dots, v_m)$ linear unabhängig.

SATZ
Lin. Unabh.
von Eigenvektoren

Beweis. Wir verwenden Induktion über m . Der Fall $m = 0$ ist klar (null Vektoren sind stets linear unabhängig). Sei also $m > 0$, und die Behauptung für $m - 1$ sei

schon bewiesen. Seien $\alpha_1, \alpha_2, \dots, \alpha_m \in K$ mit $\alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_m v_m = \mathbf{0}$. Wir müssen zeigen, dass alle $\alpha_j = 0$ sind. Es ist

$$\begin{aligned} \mathbf{0} &= \lambda_m \mathbf{0} - f(\mathbf{0}) \\ &= \lambda_m(\alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_m v_m) - f(\alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_m v_m) \\ &= \lambda_m \alpha_1 v_1 + \lambda_m \alpha_2 v_2 + \dots + \lambda_m \alpha_m v_m - \alpha_1 \lambda_1 v_1 - \alpha_2 \lambda_2 v_2 - \dots - \alpha_m \lambda_m v_m \\ &= \alpha_1(\lambda_m - \lambda_1)v_1 + \alpha_2(\lambda_m - \lambda_2)v_2 + \dots + \alpha_{m-1}(\lambda_m - \lambda_{m-1})v_{m-1}. \end{aligned}$$

Aus der Induktionsannahme folgt

$$\alpha_1(\lambda_m - \lambda_1) = \alpha_2(\lambda_m - \lambda_2) = \dots = \alpha_{m-1}(\lambda_m - \lambda_{m-1}) = 0.$$

Weil $\lambda_m \neq \lambda_1, \lambda_2, \dots, \lambda_{m-1}$ ist, ergibt sich daraus $\alpha_1 = \alpha_2 = \dots = \alpha_{m-1} = 0$. Die ursprüngliche Gleichung reduziert sich also auf $\alpha_m v_m = \mathbf{0}$. Weil $v_m \neq \mathbf{0}$ ist, folgt auch $\alpha_m = 0$. $\square$

15.7. Folgerung. Seien V ein K -Vektorraum und $f \in \text{End}(V)$. Seien weiter $\lambda_1, \lambda_2, \dots, \lambda_m$ paarweise verschiedene Elemente von K und für $j \in \{1, 2, \dots, m\}$ seien $v_{j1}, v_{j2}, \dots, v_{jn_j}$ (mit $n_j \in \mathbb{N}$) linear unabhängige Elemente von $E_{\lambda_j}(f)$. Dann sind die v_{ji} (mit $j \in \{1, 2, \dots, m\}$ und $i \in \{1, 2, \dots, n_j\}$) linear unabhängig. Insbesondere gilt

FOLG
Dimension
von Eigen-
räumen

$$\dim E_{\lambda_1}(f) + \dim E_{\lambda_2}(f) + \dots + \dim E_{\lambda_m}(f) \leq \dim V.$$

Ist V endlich-dimensional, dann kann f also höchstens $\dim V$ Eigenwerte haben. Genauer gilt, dass die Summe der geometrischen Vielfachheiten der Eigenwerte höchstens $\dim V$ sein kann.

Beweis. Sei

$$\sum_{j=1}^m \sum_{i=1}^{n_j} \alpha_{ji} v_{ji} = \mathbf{0}$$

mit $\alpha_{ji} \in K$. Sei $v_j = \sum_{i=1}^{n_j} \alpha_{ji} v_{ji} \in E_{\lambda_j}(f)$, dann gilt $v_1 + v_2 + \dots + v_m = \mathbf{0}$. Aus Satz 15.6 folgt dann $v_1 = v_2 = \dots = v_m = \mathbf{0}$, denn eventuell vorkommende Vektoren $\neq \mathbf{0}$ müssten linear unabhängig sein und könnten sich also nicht zum Nullvektor addieren. Da $(v_{j1}, v_{j2}, \dots, v_{jn_j})$ linear unabhängig ist, folgt dann aus $v_j = \mathbf{0}$ auch $\alpha_{ji} = 0$ für alle $i \in \{1, 2, \dots, n_j\}$. Da das für jedes $j \in \{1, 2, \dots, m\}$ gilt, sind alle $\alpha_{ji} = 0$. Das zeigt die lineare Unabhängigkeit.

Im Fall, dass $\dim E_{\lambda_j}(f) = \infty$ ist für ein j , folgt die letzte Aussage daraus, dass $E_{\lambda_j}(f)$ ein Untervektorraum von V ist, woraus sich $\dim V = \infty$ ergibt. Anderenfalls können wir für jedes j das Tupel $(v_{j1}, v_{j2}, \dots, v_{jn_j})$ als eine Basis von $E_{\lambda_j}(f)$ wählen. Wegen der linearen Unabhängigkeit der v_{ji} gilt dann

$$\begin{aligned} \dim V &\geq \#\{v_{ji} \mid j \in \{1, 2, \dots, m\}, i \in \{1, 2, \dots, n_j\}\} \\ &= n_1 + n_2 + \dots + n_m \\ &= \dim E_{\lambda_1}(f) + \dim E_{\lambda_2}(f) + \dots + \dim E_{\lambda_m}(f). \end{aligned} \quad \square$$

Wie können wir die Eigenwerte (und dann die Eigenräume) finden? Dazu wählen wir eine Basis und bestimmen die Matrix A von f bezüglich dieser Basis. Wir übertragen die Begriffe Eigenwert usw. auf Matrizen.

15.8. Definition. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Sei $\lambda \in K$. Dann heißt λ ein *Eigenwert* von A , wenn es einen Spaltenvektor $\mathbf{0} \neq \mathbf{x} \in K^n$ gibt mit $A\mathbf{x} = \lambda\mathbf{x}$. In diesem Fall heißt $\mathbf{x}$ ein *Eigenvektor* von A zum Eigenwert λ . Der Untervektorraum

DEF
Eigenwert
etc. für
Matrizen

$$E_\lambda(A) = \{\mathbf{x} \in K^n \mid A\mathbf{x} = \lambda\mathbf{x}\} = \ker(\lambda I_n - A)$$

heißt der λ -*Eigenraum* von A oder der *Eigenraum* von A zum Eigenwert λ ; seine Dimension heißt die *geometrische Vielfachheit* des Eigenwerts λ von A . $\diamond$

Die Eigenwerte von A und ihre geometrischen Vielfachheiten entsprechen dann denen von f .

Der Schlüssel zur Bestimmung der Eigenwerte ist folgende einfache Beobachtung.

15.9. Lemma. Seien K ein Körper, $\lambda \in K$, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. λ ist genau dann ein Eigenwert von A , wenn $\det(\lambda I_n - A) = 0$ ist. Die geometrische Vielfachheit des Eigenwerts λ ist $\dim \ker(\lambda I_n - A) = n - \text{rk}(\lambda I_n - A)$.

LEMMA
Charakterisierung von
Eigenwerten

Beweis. Wir haben folgende Kette von Äquivalenzen:

$$\begin{aligned} \lambda \text{ ist Eigenwert von } A &\iff E_\lambda(A) \neq \{\mathbf{0}\} \\ &\iff \ker(\lambda I_n - A) \neq \{\mathbf{0}\} \\ &\iff \det(\lambda I_n - A) = 0 \end{aligned}$$

Die letzte Aussage folgt aus der Definition der geometrischen Vielfachheit und dem Rangsatz 10.17. $\square$

15.10. Beispiel. Wir betrachten

$$A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \in \text{Mat}(2, \mathbb{R}).$$

BSP
Bestimmung
der
Eigenwerte

(Das ist die Matrix zu $f: (x, y) \mapsto (y, x)$ wie in Beispiel 15.5 (2).) Dann ist für $\lambda \in \mathbb{R}$

$$\det(\lambda I_2 - A) = \begin{vmatrix} \lambda & -1 \\ -1 & \lambda \end{vmatrix} = \lambda^2 - 1 = (\lambda - 1)(\lambda + 1).$$

Das verschwindet genau für $\lambda = 1$ und $\lambda = -1$, also sind das die Eigenwerte von A . Wir können Basen der Eigenräume $E_\lambda(A)$ mit dem Zeilenstufenform-Algorithmus, angewandt auf $\lambda I_2 - A$, berechnen. Für $\lambda = 1$ haben wir

$$\lambda I_2 - A = I_2 - A = \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \xrightarrow{\mathbf{II}_{2,1}(1)} \begin{pmatrix} 1 & -1 \\ 0 & 0 \end{pmatrix};$$

das liefert die (eielementige) Basis $((1, 1))$ von $E_1(A)$. Für $\lambda = -1$ sieht es wie folgt aus:

$$\lambda I_2 - A = -I_2 - A = \begin{pmatrix} -1 & -1 \\ -1 & -1 \end{pmatrix} \xrightarrow{\mathbf{I}_1(-1)} \begin{pmatrix} 1 & 1 \\ -1 & -1 \end{pmatrix} \xrightarrow{\mathbf{II}_{2,1}(1)} \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix};$$

die Basis ist $((-1, 1))$. $\clubsuit$

An diesem Beispiel sieht man, dass die Determinante, deren Verschwinden anzeigt, dass λ ein Eigenwert ist, ein *Polynom* in λ (mit Koeffizienten in K) ist. Wir müssen daher etwas ausholen und ein wenig über Polynome sprechen.

Exkurs: Polynome.

15.11. Definition. Sei K ein Körper. Ein *Polynom* in der *Variablen* (oder *Unbestimmten*) X über K ist ein Ausdruck der Form

$$p = a_n X^n + a_{n-1} X^{n-1} + \dots + a_1 X + a_0$$

mit $n \in \mathbb{N}$ und $a_0, a_1, \dots, a_n \in K$. a_j heißt der *j-te Koeffizient* von p oder der *Koeffizient von X^j in p* . Wir setzen $a_j = 0$ für $j > n$. Ist $a_n \neq 0$, dann hat das Polynom *Grad n* : $\deg(p) = n$ (englisch „degree“). In diesem Fall heißt a_n der *Leitkoeffizient* von p . Ist $a_n = 0$, dann heißt p *normiert*. Sind alle $a_j = 0$, dann ist p das *Nullpolynom*; sein Grad ist definiert als $\deg(\mathbf{0}) = -\infty$. Ist $n = 0$, dann heißt p *konstant* (d.h., $p = \mathbf{0}$ oder $\deg(p) = 0$). Wir schreiben $K[X]$ für die Menge der Polynome in X über K .

Sei $q = b_m X^m + \dots + b_1 X + b_0$ ein weiteres Polynom. Die Polynome p und q sind genau dann *gleich*, $p = q$, wenn ihre Koeffizienten übereinstimmen: $a_j = b_j$ für alle $j \in \mathbb{N}$ (mit der Konvention $a_j = 0$ für $j > n$ und $b_j = 0$ für $j > m$). Die *Summe* von p und q ist

$$p + q = \sum_{j=0}^{\max\{m,n\}} (a_j + b_j) X^j;$$

es gilt $\deg(p + q) \leq \max\{\deg(p), \deg(q)\}$. Das *Produkt* von p und q ist

$$p \cdot q = \sum_{k=0}^{m+n} \left(\sum_{i,j: i+j=k} a_i b_j \right) X^k;$$

es gilt $\deg(pq) = \deg(p) + \deg(q)$. Wir identifizieren K mit der Teilmenge der konstanten Polynome: $K \subset K[X]$. Die Menge $K[X]$ wird mit der eben definierten Addition und Multiplikation ein kommutativer Ring; er heißt der *Polynomring* in X über K . Die Einschränkung der Multiplikation auf $K \times K[X]$ macht $K[X]$ zu einem unendlich-dimensionalen K -Vektorraum mit Basis $(1, X, X^2, X^3, \dots)$. $\diamond$

Wenn es Sie stört, dass in der Definition von einem „Ausdruck der Form ...“ gesprochen wird, ohne dass gesagt wird, was das eigentlich „ist“, dann lesen Sie hier weiter.

Formal kann man die Definition auf stabile Füße stellen, indem man setzt

$$K[X] = \{(a_n)_{n \in \mathbb{N}} \in K^{\mathbb{N}} \mid \exists N \in \mathbb{N} \forall n > N: a_n = 0\}.$$

Das sind also die endlichen Folgen von Elementen von K , in dem Sinne, dass alle bis auf endlich viele Folgenglieder null sind. Man setzt weiter $X = (0, 1, 0, 0, 0, \dots) \in K[X]$ und definiert die Abbildung $i: K \rightarrow K[X]$, $a \mapsto (a, 0, 0, 0, \dots)$. Die Addition in $K[X]$ wird komponentenweise definiert, die Multiplikation mit X durch

$$X \cdot (a_0, a_1, a_2, \dots) = (0, a_0, a_1, a_2, \dots)$$

und die mit $i(a)$ durch

$$i(a) \cdot (a_0, a_1, a_2, \dots) = (aa_0, aa_1, aa_2, \dots).$$

Dann ist

$$(a_0, a_1, \dots, a_n, 0, 0, 0, \dots) = i(a_0) + i(a_1)X + i(a_2)X^2 + \dots + i(a_n)X^n,$$

und die Multiplikation damit wird so definiert, dass das Assoziativ- und das Distributivgesetz gelten. Mittels der Abbildung i wird K mit seinem Bild in $K[X]$ identifiziert; man schreibt also einfach a statt $i(a)$. Die Ringaxiome muss man dann noch nachprüfen. Die K -Vektorraum-Struktur von $K[X]$ ist einfach die als Untervektorraum von $K^{\mathbb{N}}$.

Das funktioniert auch dann noch, wenn man die Endlichkeitsbedingung in der Definition weglässt. Man erhält dann den Ring $K[[X]]$ der *formalen Potenzreihen* in X über K . Für

DEF
Polynom
Polynomring
Grad
Leitkoeffizient
normiert

$K = \mathbb{R}$ oder $\mathbb{C}$ spielen diese Potenzreihen eine wichtige Rolle in der Analysis (bzw. Funktionentheorie).

In Polynome kann man einsetzen:

15.12. Definition. Seien K ein Körper und $p = a_n X^n + \dots + a_1 X + a_0 \in K[X]$ ein Polynom. Für $\lambda \in K$ ist der Wert von p bei λ gegeben durch

$$p(\lambda) = a_n \lambda^n + \dots + a_1 \lambda + a_0.$$

λ heißt eine Nullstelle von p , wenn $p(\lambda) = 0$ ist. Für $p, q \in K[X]$ und $\lambda \in K$ gilt dann $(p + q)(\lambda) = p(\lambda) + q(\lambda)$ und $(p \cdot q)(\lambda) = p(\lambda) \cdot q(\lambda)$. $\diamond$

DEF
Werte und
Nullstellen
von
Polynomen

Ein Polynom $p \in K[X]$ führt also zu einer Polynomfunktion $K \rightarrow K$, $\lambda \mapsto p(\lambda)$. Die Abbildung $K[X] \rightarrow \text{Abb}(K, K)$, die einem Polynom die zugehörige Polynomfunktion zuordnet, ist injektiv, wenn der Körper K unendlich ist. Das ergibt sich aus dem folgenden Satz.

15.13. Satz. Seien K ein Körper, $n \in \mathbb{N}$ und $x_1, x_2, \dots, x_n \in K$ paarweise verschieden. Seien weiter $y_1, y_2, \dots, y_n \in K$. Dann gibt es ein eindeutig bestimmtes Polynom $p \in K[X]$ mit $\deg(p) < n$, sodass $p(x_j) = y_j$ ist für alle $j \in \{1, 2, \dots, n\}$.

SATZ
Eindeutigkeit
von
Polynomen

Beweis. Diese Aussage ist die Verallgemeinerung des Resultats von Beispiel 10.15 von $\mathbb{R}$ auf beliebige Körper. Der Beweis geht genauso wie dort. $\square$

15.14. Folgerung. Ein Polynom $p \in K[X]$ mit $\deg(p) = n \in \mathbb{N}$ kann nicht mehr als n Nullstellen in K haben.

FOLG
Nullstellen
von
Polynomen

Beweis. Angenommen, p hat $n + 1$ Nullstellen $x_1, x_2, \dots, x_{n+1} \in K$. Dann muss p das eindeutig bestimmte Polynom von Grad $< n + 1$ sein, das $p(a_j) = 0$ erfüllt für alle $j \in \{1, 2, \dots, n + 1\}$. Das Nullpolynom hat diese Eigenschaft, also muss $p = 0$ sein. Das ist aber ein Widerspruch zur Voraussetzung $\deg(p) = n$. $\square$

15.15. Folgerung. Hat der Körper K unendlich viele Elemente, dann ist die Abbildung $K[X] \rightarrow \text{Abb}(K, K)$, die einem Polynom die zugehörige Polynomfunktion zuordnet, injektiv.

FOLG
Polynome
und
Polynom-
funktionen

Ein Polynom $p \in K[X]$ ist dann also durch seine Werte $p(\lambda)$ für $\lambda \in K$ eindeutig bestimmt. Wir können also zum Beispiel den Vektorraum P der Polynomfunktionen mit $\mathbb{R}[X]$ identifizieren.

Beweis. Wir schreiben Φ für die Abbildung $K[X] \rightarrow \text{Abb}(K, K)$. Φ ist linear, also genügt es zu zeigen, dass $\ker(\Phi) = \{0\}$ ist. Sei also $p \in \ker(\Phi)$. Dann ist $p(\lambda) = 0$ für alle $\lambda \in K$, also hat (da K unendlich ist) das Polynom p unendlich viele Nullstellen in K . Nach Folgerung 15.14 muss p das Nullpolynom sein. $\square$

Für endliche Körper K ist die Aussage falsch: Ist $\#K = q < \infty$, dann ist $\dim_K \text{Abb}(K, K) = q$, denn eine Abbildung $f: K \rightarrow K$ ist durch die q Werte $f(a)$ für $a \in K$ eindeutig bestimmt. Auf der anderen Seite ist $\dim_K K[X] = \infty$, und damit kann es keine injektive lineare Abbildung $K[X] \rightarrow \text{Abb}(K, K)$ geben.

Der Kern von Φ besteht in diesem Fall aus allen Polynomen, die alle Elemente von K als Nullstellen haben. Man kann zeigen, dass

$$\prod_{a \in K} (X - a) = X^q - X$$

ist; der Kern besteht demnach genau aus den Vielfachen von $X^q - X$.

So wie man ganze Zahlen mit Rest durcheinander dividieren kann, gibt es auch für Polynome eine Division mit Rest („Polynomdivision“).

15.16. Satz. Sei K ein Körper und seien $f, g \in K[X]$ mit g normiert. Dann gibt es eindeutig bestimmte Polynome q („Quotient“) und r („Rest“) in $K[X]$ mit $f = qg + r$ und $\deg(r) < \deg(g)$.

SATZ
Polynom-
division

Beweis. Wir beweisen zunächst die Existenz. Sei $\deg(g) = m$, also

$$g = X^m + b_{m-1}X^{m-1} + \dots + b_1X + b_0.$$

Wir betrachten g als fest und führen den Beweis durch Induktion über $\deg(f)$. Ist $\deg(f) < m$, dann erfüllen $q = 0$ und $r = f$ die Bedingungen. Wir können also annehmen, dass $n = \deg(f) \geq m$ ist; die Existenzaussage sei für $\deg(f) < n$ bereits bewiesen. Es ist $f = a_nX^n + a_{n-1}X^{n-1} + \dots + a_0$, also ist

$$\tilde{f} = f - a_nX^{n-m}g = (a_{n-1} - a_nb_{m-1})X^{n-1} + \dots;$$

der Grad von $\tilde{f}$ ist damit kleiner als n . Nach Induktionsvoraussetzung gibt es $\tilde{q}, r \in K[X]$ mit $\tilde{f} = \tilde{q}g + r$ und $\deg(r) < m$. Wir setzen $q = a_nX^{n-m} + \tilde{q}$; dann ist

$$f = a_nX^{n-m}g + \tilde{f} = a_nX^{n-m}g + \tilde{q}g + r = qg + r$$

wie gewünscht.

Zur Eindeutigkeit: Seien $q_1, q_2, r_1, r_2 \in K[X]$ mit $f = q_1g + r_1 = q_2g + r_2$ und $\deg(r_1), \deg(r_2) < \deg(g)$. Es folgt

$$(q_1 - q_2)g = r_2 - r_1.$$

Die rechte Seite hat Grad $< \deg(g)$. Wäre $q_1 \neq q_2$, dann hätte die linke Seite Grad $\deg(q_1 - q_2) + \deg(g) \geq \deg(g)$, ein Widerspruch. Also ist $q_1 = q_2$ und damit auch $r_1 = r_2$. $\square$

Der Beweis übersetzt sich direkt in den üblichen Algorithmus zur Polynomdivision: Man subtrahiert geeignete Vielfache von g solange von f , bis man ein Polynom zurückbehält, dessen Grad kleiner als der von g ist.

15.17. Folgerung. Seien K ein Körper, $p \in K[X]$ und $\lambda \in K$. Wir schreiben

$$p = q \cdot (X - \lambda) + r$$

FOLG
Nullstellen

wie in Satz 15.16. Dann ist $r = p(\lambda)$ konstant. Insbesondere ist λ genau dann eine Nullstelle von p , wenn $r = 0$ ist.

Beweis. r ist konstant, da $\deg(r) < 1 = \deg(X - \lambda)$ ist. Außerdem gilt

$$p(\lambda) = q(\lambda)(\lambda - \lambda) + r = r.$$

$\square$

Ist λ eine Nullstelle von p , dann ist demnach $p = (X - \lambda) \cdot q$ mit einem Polynom $q \in K[X]$. Ist λ auch eine Nullstelle von q , dann ist $p = (X - \lambda)^2 \cdot \tilde{q}$ und so fort. Das führt zu folgender Definition.

15.18. Definition. Seien K ein Körper, $0 \neq p \in K[X]$ ein Polynom und $\lambda \in K$. Die *Vielfachheit* der Nullstelle λ von p ist die größte Zahl $n \in \mathbb{N}$, sodass man p in der Form $p = (X - \lambda)^n \cdot q$ schreiben kann mit einem Polynom $q \in K[X]$. In diesem Fall ist $q(\lambda) \neq 0$. $\diamond$

DEF
Vielfachheit
einer
Nullstelle

λ ist also genau dann eine Nullstelle, wenn die Vielfachheit von λ als Nullstelle positiv ist.

15.19. Beispiele. Wir betrachten $K = \mathbb{R}$. Für $p = X^3 - X^2 - X + 1 \in \mathbb{R}[X]$ gilt

BSP
Vielfachheiten

$$p = (X - 1)^2(X + 1),$$

also hat p die Nullstellen 1 (mit Vielfachheit 2: eine „doppelte“ Nullstelle) und -1 (mit Vielfachheit 1: eine „einfache“ Nullstelle).

Für $q = X^3 + X^2 + X + 1 \in \mathbb{R}[X]$ gilt dagegen

$$q = (X + 1)(X^2 + 1),$$

also hat q nur die (einfache) Nullstelle -1 in $\mathbb{R}$, denn der zweite Faktor $X^2 + 1$ nimmt nur positive Werte an und hat daher keine reelle Nullstelle. Wenn wir q aber als Polynom in $\mathbb{C}[X]$ betrachten, dann haben wir

$$q = (X + 1)(X + i)(X - i);$$

q hat also die drei komplexen (einfachen) Nullstellen -1 , i und $-i$. $\clubsuit$

15.20. Beispiel. Wie kann man die Nullstellen eines Polynoms p finden? Das ist einfach im Fall $\deg(p) = 1$, und im Fall $\deg(p) = 2$ kann man die Lösungsformel für quadratische Gleichungen verwenden. Für allgemeine Polynome ist die Bestimmung der Nullstellen allerdings ein schwieriges Problem. In einfach gelagerten Fällen (wie üblicherweise etwa bei Klausuraufgaben) kann man häufig eine Nullstelle α durch Probieren erraten. Dann kann man p durch $X - \alpha$ teilen, um ein Polynom kleineren Grades zu erhalten, mit dem man dann weitermacht.

BSP
Nullstellen

Als Beispiel betrachten wir

$$p = X^4 + X^3 - 6X^2 - 2X + 4 \in \mathbb{R}[X].$$

Wir probieren, ob 0, 1, -1 Nullstellen sind, und finden $p(-1) = 0$. Division von p durch $X + 1$ ergibt

$$p = (X + 1)(X^3 - 6X + 4),$$

und wir müssen noch die Nullstellen von $p_1 = X^3 - 6X + 4$ finden. Weiteres Probieren liefert die Nullstelle 2. Division von p_1 durch $X - 2$ ergibt

$$p_1 = (X - 2)(X^2 + 2X - 2).$$

Auf den verbleibenden Faktor können wir jetzt die Lösungsformel für quadratische Gleichungen anwenden; sie liefert die weiteren Nullstellen

$$\frac{-2 \pm \sqrt{2^2 - 4 \cdot (-2)}}{2} = -1 \pm \sqrt{3}. \quad \clubsuit$$

Zurück zu Eigenwerten und Eigenräumen. Wir haben gesehen, dass der Ausdruck

$$\det(\lambda I_n - A)$$

darüber entscheidet, ob λ ein Eigenwert von A ist oder nicht. Ist $A = (a_{ij})$, dann hat diese Determinante die folgende Form:

$$\begin{vmatrix} \lambda - a_{11} & -a_{12} & -a_{13} & \cdots & -a_{1n} \\ -a_{21} & \lambda - a_{22} & -a_{23} & \cdots & -a_{2n} \\ -a_{31} & -a_{32} & \lambda - a_{33} & \cdots & -a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -a_{n1} & -a_{n2} & -a_{n3} & \cdots & \lambda - a_{nn} \end{vmatrix}$$

Wenn wir das in die Leibniz-Formel einsetzen, dann bekommen wir

$$\begin{aligned} & (\lambda - a_{11})(\lambda - a_{22}) \cdots (\lambda - a_{nn}) + \text{Terme mit } \leq n-2 \text{ Faktoren } \lambda - a_{jj} \\ & = \lambda^n - (a_{11} + \dots + a_{nn})\lambda^{n-1} + \dots + (-1)^n \det(A). \end{aligned}$$

Das hat die Form $p(\lambda)$ mit einem normierten Polynom $p \in K[X]$ vom Grad n .

*

15.21. Definition. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Das Polynom $\chi_A = \det(XI_n - A) \in K[X]$ heißt das *charakteristische Polynom* von A . $\diamond$

DEF
Charakteristisches Polynom

Die Leibniz-Formel zeigt, dass man die Determinante allgemeiner für Matrizen mit Einträgen in einem kommutativen Ring R definieren kann; die Determinante ist dann ebenfalls ein Element von R . Die Multiplikativität und die Symmetrie der Determinante gelten auch in diesem erweiterten Kontext. In der obigen Definition wenden wir das für den Ring $K[X]$ und die Matrix $XI_n - A \in \text{Mat}(n, K[X])$ an.

In der Literatur findet man auch häufig die Definition $\det(A - XI_n)$ für das charakteristische Polynom von A . Diese Definition unterscheidet sich von der hier gegebenen nur durch einen Faktor $(-1)^n$. Das ändert nichts an den Nullstellen, sodass es für den Zweck der Eigenwert-Berechnung oder für die Definition der algebraischen Vielfachheit (siehe unten) keinen Unterschied macht. Der Nachteil der anderen Variante ist aus meiner Sicht, dass das Polynom für ungerades n dann nicht normiert ist (sondern Leitkoeffizient -1 hat).

Wir haben gesehen, dass die Eigenwerte von A genau die Nullstellen des charakteristischen Polynoms von A sind.

15.22. Beispiel. Was sind die Eigenwerte der Matrix

$$A = \begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & 1 \\ 1 & 2 & 3 \end{pmatrix} \in \text{Mat}(3, \mathbb{R}) ?$$

BSP
Eigenwerte

Wir bestimmen das charakteristische Polynom:

$$\begin{aligned} \det(XI_3 - A) &= \begin{vmatrix} X-1 & 0 & 0 \\ -1 & X-1 & -1 \\ -1 & -2 & X-3 \end{vmatrix} \\ &= (X-1) \begin{vmatrix} X-1 & -1 \\ -2 & X-3 \end{vmatrix} \\ &= (X-1)((X-1)(X-3) - 1 \cdot 2) \\ &= (X-1)(X^2 - 4X + 1) \end{aligned}$$

Ein Eigenwert ist $\lambda_1 = 1$, die anderen beiden finden wir mit Hilfe der Lösungsformel für quadratische Gleichungen:

$$\lambda_2 = 2 + \sqrt{3} \quad \text{und} \quad \lambda_3 = 2 - \sqrt{3}.$$



Wir können die Definitionen von Determinante und charakteristischem Polynom auch auf Endomorphismen übertragen.

15.23. Definition. Seien K ein Körper, V ein endlich-dimensionaler K -Vektorraum und $f \in \text{End}(V)$. Sei B eine beliebige Basis von V und sei $A = \text{Mat}_{B,B}(f)$. Dann ist die *Determinante* von f definiert als $\det(f) = \det(A)$ und das *charakteristische Polynom* χ_f von f ist das charakteristische Polynom von A . $\diamond$

DEF
Determinante,
char. Pol.
von Endo-
morphisimen

Die Definition ist sinnvoll, weil sie nicht von der Wahl der Basis B abhängt: Ist $A' = \text{Mat}_{B',B'}(f)$ mit einer anderen Basis B' von V , dann gibt es eine Matrix $P \in \text{GL}(n, K)$ (mit $n = \dim V$), sodass $A' = PAP^{-1}$ ist. Dann ist

$$\begin{aligned}\det(A') &= \det(PAP^{-1}) = \det(P) \det(A) \det(P^{-1}) \\ &= \det(A) \det(P P^{-1}) = \det(A) \det(I_n) = \det(A).\end{aligned}$$

Es ist auch $P(XI_n - A)P^{-1} = XPI_nP^{-1} - PAP^{-1} = XI_n - A'$, und die gleiche Rechnung wie eben zeigt, dass A und A' dasselbe charakteristische Polynom haben.

*** 15.24. Definition.** Seien K ein Körper, $n \in \mathbb{N}$, $A \in \text{Mat}(n, K)$ und $\lambda \in K$. Die *algebraische Vielfachheit* von λ als Eigenwert von A ist die Vielfachheit von λ als Nullstelle des charakteristischen Polynoms von A . Entsprechend definieren wir die *algebraische Vielfachheit* von λ als Eigenwert eines Endomorphismus f eines endlich-dimensionalen K -Vektorraums V . $\diamond$

DEF
algebraische
Vielfachheit

Wir haben jetzt also zwei Vielfachheiten von Eigenwerten definiert, die geometrische und die algebraische. In welcher Beziehung stehen sie zueinander? Wir wissen bisher Folgendes:

$$\text{geom. Vielfachheit} > 0 \iff \text{Eigenwert} \iff \text{alg. Vielfachheit} > 0$$

Müssen die beiden Vielfachheiten immer gleich sein?

15.25. Beispiel. Sei

$$A = \begin{pmatrix} \alpha & 1 \\ 0 & \alpha \end{pmatrix} \in \text{Mat}(2, K).$$

BSP
alg. $\neq$ geom.
Vielfachheit

Das charakteristische Polynom von A ist $(X - \alpha)^2$, also hat α die algebraische Vielfachheit 2. Auf der anderen Seite ist $E_\alpha(A) = \langle (1, 0) \rangle$ (denn $\text{rk}(\alpha I_2 - A) = 1$), also hat α die geometrische Vielfachheit 1. $\clubsuit$

Die Vielfachheiten können also verschieden sein. Eine Beziehung gilt jedoch.

*** 15.26. Satz.** Seien K ein Körper, V ein endlich-dimensionaler K -Vektorraum, $f \in \text{End}(V)$ und $\lambda \in K$. Dann ist die geometrische Vielfachheit von λ als Eigenwert von f nicht größer als seine algebraische Vielfachheit.

SATZ
geom. $\leq$ alg.
Vielfachheit

Die analoge Aussage gilt dann natürlich auch für Matrizen $A \in \text{Mat}(n, K)$.

Beweis. Sei $m = \dim E_\lambda(f)$ die geometrische Vielfachheit, sei $n = \dim V$ und sei $(b_1, b_2, \dots, b_m)$ eine Basis von $E_\lambda(f)$. Wir können diese Basis zu einer Basis $B = (b_1, b_2, \dots, b_n)$ von V erweitern. Dann ist

$$A = \text{Mat}_{B,B}(f) = \left(\begin{array}{c|c} \lambda I_m & D \\ \hline \mathbf{0}_{n-m,m} & C \end{array} \right)$$

mit Matrizen $D \in \text{Mat}(m \times (n - m), K)$ und $C \in \text{Mat}(n - m, K)$, denn für $j \in \{1, 2, \dots, m\}$ ist $f(b_j) = \lambda b_j$; in der j -ten Spalte von A kommt also das λ -fache des j -ten Standard-Basisvektors zu stehen. Das charakteristische Polynom von f ist dann

$$\det(XI_n - A) = \det \left(\begin{array}{c|c} (X - \lambda)I_m & -D \\ \hline \mathbf{0}_{n-m, m} & XI_{n-m} - C \end{array} \right) = (X - \lambda)^m \det(XI_{n-m} - C).$$

Dabei haben wir die Determinante m -mal jeweils nach der ersten Spalte entwickelt. Das zeigt, dass die Vielfachheit von λ als Nullstelle des charakteristischen Polynoms mindestens m ist. $\square$

Man kann die obige Formel auf allgemeinere Blockmatrizen ausdehnen (Beweis als Übung): Sind $A \in \text{Mat}(m, K)$, $B \in \text{Mat}(m \times n, K)$ und $C \in \text{Mat}(n, K)$, dann ist

$$\det \left(\begin{array}{c|c} A & B \\ \hline \mathbf{0}_{n, m} & C \end{array} \right) = \det(A) \det(C).$$

16. DIAGONALISIERBARKEIT

Wir haben in Folgerung 15.7 gesehen, dass die Summe der geometrischen Vielfachheiten der Eigenwerte eines Endomorphismus f eines n -dimensionalen Vektorraums (oder einer $n \times n$ -Matrix) höchstens n ist. Das macht den Fall interessant, in dem diese Schranke erreicht wird. Wir formulieren zunächst eine Definition und werden dann sehen, was sie mit dieser Frage zu tun hat.

*

16.1. Definition. Seien K ein Körper, $n \in \mathbb{N}$. Eine Matrix $A = (a_{ij}) \in \text{Mat}(n, K)$ ist eine *Diagonalmatrix* oder *diagonal*, wenn $a_{ij} = 0$ ist für alle $i, j \in \{1, 2, \dots, n\}$ mit $i \neq j$. Ist $a_{ii} = d_i$ für $i \in \{1, 2, \dots, n\}$, dann schreiben wir $\text{diag}(d_1, d_2, \dots, d_n)$ für A :

DEF
Diagonal-
matrix

$$\text{diag}(d_1, d_2, \dots, d_n) = \begin{pmatrix} d_1 & 0 & \cdots & 0 \\ 0 & d_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & d_n \end{pmatrix}. \quad \diamond$$

16.2. Lemma. Seien K ein Körper, $n \in \mathbb{N}$, V ein n -dimensionaler K -Vektorraum und $f \in \text{End}(V)$. Dann sind die folgenden Aussagen äquivalent:

LEMMA
Diagonal-
sierbarkeit

- (1) V hat eine Basis B , die aus Eigenvektoren von f besteht.
- (2) Die Summe der geometrischen Vielfachheiten der Eigenwerte von f ist n .
- (3) Sei $A = \text{Mat}_{B', B'}(f)$ die Matrix von f bezüglich einer beliebigen Basis B' von V . Dann ist A ähnlich zu einer Diagonalmatrix.

Beweis.

„(1) $\Rightarrow$ (2)“: Seien $\lambda_1, \lambda_2, \dots, \lambda_k \in K$ die (paarweise verschiedenen) Eigenwerte von f , sei $n_j = \dim E_{\lambda_j}(f)$ die geometrische Vielfachheit von λ_j und sei m_j die Anzahl der Basisvektoren in B , die Eigenvektoren zum Eigenwert λ_j sind. Dann haben wir jeweils m_j linear unabhängige Vektoren in $E_{\lambda_j}(f)$, also ist $m_j \leq n_j$. Auf der anderen Seite ist

$$n = \#B = m_1 + m_2 + \dots + m_k \leq n_1 + n_2 + \dots + n_k \leq n,$$

also haben wir Gleichheit; insbesondere ist $n_1 + n_2 + \dots + n_k = n$.

„(2) $\Rightarrow$ (1)“: Wir behalten die Bezeichnungen bei. Sei B_j eine Basis von $E_{\lambda_j}(f)$. Nach Folgerung 15.7 bilden die B_j zusammen eine linear unabhängige Familie B in V . Wegen

$$\#B = \#B_1 + \#B_2 + \dots + \#B_k = n_1 + n_2 + \dots + n_k = n = \dim V$$

ist B dann eine Basis von V , die nach Konstruktion aus Eigenvektoren von f besteht.

„(1) $\Rightarrow$ (3)“: Sei $B = (b_1, b_2, \dots, b_n)$ und sei α_j der Eigenwert von f , sodass $f(b_j) = \alpha_j b_j$. Dann ist

$$\text{Mat}_{B, B}(f) = \text{diag}(\alpha_1, \alpha_2, \dots, \alpha_n)$$

eine Diagonalmatrix, und nach Satz 15.1 ist A zu ihr ähnlich.

„(3) $\Rightarrow$ (1)“: Ist A ähnlich zu $D = \text{diag}(\alpha_1, \alpha_2, \dots, \alpha_n)$, dann gibt es nach Satz 15.1 eine Basis $B = (b_1, b_2, \dots, b_n)$ von V , sodass $D = \text{Mat}_{B, B}(f)$ ist. Daran liest man ab, dass $f(b_j) = \alpha_j b_j$ ist (und $b_j \neq \mathbf{0}$), also besteht B aus Eigenvektoren von f . $\square$

* **16.3. Definition.** Sei f ein Endomorphismus eines endlich-dimensionalen K -Vektorraums V . Dann heißt f *diagonalisierbar*, wenn V eine Basis hat, die aus Eigenvektoren von f besteht.

DEF
diagonalisierbar

Eine Matrix $A \in \text{Mat}(n, K)$ heißt *diagonalisierbar*, wenn sie zu einer Diagonalmatrix ähnlich ist, wenn es also eine Diagonalmatrix $D \in \text{Mat}(n, K)$ und eine invertierbare Matrix $P \in \text{GL}(n, K)$ gibt mit $PAP^{-1} = D$. $\diamond$

Aus dem Lemma ergibt sich, dass ein Endomorphismus genau dann diagonalisierbar ist, wenn die zugeordnete Matrix (bezüglich irgendeiner Basis) diagonalisierbar ist.

Aus dem Beweis ergibt sich auch, dass die Einträge auf der Diagonalen der Diagonalmatrix gerade die Eigenwerte sind; sie kommen so oft vor, wie es ihrer geometrischen Vielfachheit entspricht.

Da die geometrische Vielfachheit eines Eigenwerts höchstens so groß ist wie seine algebraische Vielfachheit, ist eine *notwendige* Bedingung für die Diagonalisierbarkeit, dass die Summe der *algebraischen* Vielfachheiten der Eigenwerte n ist. Das ist eine Eigenschaft des charakteristischen Polynoms.

16.4. Definition. Seien K ein Körper und $p \in K[X]$ ein normiertes Polynom vom Grad n . Wir sagen, p *zerfällt in Linearfaktoren* über K , wenn es $\alpha_1, \alpha_2, \dots, \alpha_n \in K$ gibt, sodass

DEF
Zerlegung in Linearfaktoren

$$p = \prod_{j=1}^n (X - \alpha_j) = (X - \alpha_1)(X - \alpha_2) \cdots (X - \alpha_n)$$

ist. $\diamond$

Das bedeutet also, dass die Summe der Vielfachheiten der Nullstellen von p in K gleich dem Grad n von p ist.

* **16.5. Folgerung.** Seien K ein Körper und f ein Endomorphismus eines endlich-dimensionalen K -Vektorraums V . Dann sind die folgenden Aussagen äquivalent:

FOLG
Charakterisierung von Diagonalisierbarkeit

- (1) f ist diagonalisierbar.
- (2) Das charakteristische Polynom von f zerfällt in Linearfaktoren über K und für jeden Eigenwert von f stimmt die geometrische mit der algebraischen Vielfachheit überein.

Insbesondere gilt: Hat f genau $n = \dim V$ verschiedene Eigenwerte, dann ist f diagonalisierbar.

Wichtig: Die Umkehrung des letzten Satzes gilt nicht. Zum Beispiel sind die Einheitsmatrix oder die Nullmatrix diagonalisierbar (weil schon diagonal), sie haben aber jeweils nur einen Eigenwert (nämlich 1 bzw. 0).



Beweis. Seien $\lambda_1, \lambda_2, \dots, \lambda_k \in K$ die verschiedenen Eigenwerte von f und seien n_j bzw. m_j ihre geometrischen bzw. algebraischen Vielfachheiten. Sei $n = \dim V$. Dann gilt wegen $n_j \leq m_j$ (Satz 15.26) und $m_1 + m_2 + \dots + m_k \leq \deg(\chi_f) = n$:

$$\begin{aligned} f \text{ diagonalisierbar} &\iff n_1 + n_2 + \dots + n_k = n \\ &\iff m_1 + m_2 + \dots + m_k = n \quad \text{und} \quad \forall j: m_j = n_j \\ &\iff \chi_f \text{ zerfällt in Linearfaktoren über } K \text{ und } \forall j: m_j = n_j. \end{aligned}$$

Hat f n verschiedene Eigenwerte, dann ist $k = n$ und $m_j = 1$. Dann muss auch die geometrische Vielfachheit $n_j = 1$ sein, also ist $n_1 + n_2 + \dots + n_k = n$. $\square$

16.6. Beispiel. Die Matrix aus Beispiel 15.22 ist diagonalisierbar, weil sie die drei verschiedenen Eigenwerte 1 , $2 + \sqrt{3}$ und $2 - \sqrt{3}$ hat. $\clubsuit$

16.7. Beispiel. Eine Matrix muss nicht diagonalisierbar sein, wenn ihr charakteristisches Polynom in Linearfaktoren zerfällt. Das hatten wir (siehe Beispiel 15.25) an Hand der Matrix

$$\begin{pmatrix} \alpha & 1 \\ 0 & \alpha \end{pmatrix}$$

gesehen, deren charakteristisches Polynom $(X - \alpha)^2$ in Linearfaktoren zerfällt, für die aber die geometrische Vielfachheit von α (nämlich 1) kleiner ist als die algebraische Vielfachheit (nämlich 2). $\clubsuit$

Die Gleichheit der geometrischen und algebraischen Vielfachheit ist also eine wesentliche Bedingung. Die Bedingung, dass das charakteristische Polynom in Linearfaktoren zerfällt, können wir hingegen erfüllen, wenn unser Körper „groß genug“ ist. Wir erinnern uns daran (Satz 4.4), dass der Körper $\mathbb{C}$ der komplexen Zahlen *algebraisch abgeschlossen* ist. Das bedeutet, dass jedes nicht-konstante Polynom $p \in \mathbb{C}[X]$ eine Nullstelle in $\mathbb{C}$ hat. Daraus folgt, dass jedes normierte Polynom über $\mathbb{C}$ in Linearfaktoren zerfällt:

16.8. Folgerung. Sei $p \in \mathbb{C}[X]$ normiert. Dann zerfällt p in Linearfaktoren über $\mathbb{C}$.

Beweis. Wir beweisen die Aussage durch Induktion über den Grad von p . Im Fall $\deg(p) = 0$ ist $p = 1$ und damit gleich dem leeren Produkt (anders ausgedrückt, p hat genau $\deg(p) = 0$ Nullstellen in $\mathbb{C}$), also zerfällt p trivialerweise in Linearfaktoren. Die Aussage gelte für Polynome vom Grad n , und p habe Grad $n + 1$. Nach dem Fundamentalsatz der Algebra 4.4 hat p eine Nullstelle $\alpha_1 \in \mathbb{C}$. Dann ist $p = (X - \alpha_1)q$ mit $q \in \mathbb{C}[X]$ normiert und $\deg(q) = n$. Nach der Induktionsvoraussetzung zerfällt q in Linearfaktoren:

$$q = (X - \alpha_2)(X - \alpha_3) \cdots (X - \alpha_{n+1}),$$

also gilt das auch für p :

$$p = (X - \alpha_1)q = (X - \alpha_1)(X - \alpha_2) \cdots (X - \alpha_{n+1}). \quad \square$$

Für den Beweis haben wir nur verwendet, dass $\mathbb{C}$ algebraisch abgeschlossen ist; die Aussage gilt also entsprechend für jeden algebraisch abgeschlossenen Körper. Außerdem kann man ganz allgemein zeigen, dass es zu jedem Körper K einen algebraisch abgeschlossenen Körper $\bar{K}$ gibt, der K als Teilkörper enthält (also sodass die Addition und Multiplikation in K die Einschränkungen derjenigen von $\bar{K}$ sind). Durch Übergang von K zu $\bar{K}$ kann man dann also immer erreichen, dass das charakteristische Polynom einer Matrix (oder eines Endomorphismus) in Linearfaktoren zerfällt. Die Bedingung an die Gleichheit der geometrischen und algebraischen Vielfachheiten der Eigenwerte ist also die eigentlich entscheidende für die Diagonalisierbarkeit.

Wir erinnern uns daran, dass das charakteristische Polynom χ_A einer Matrix $A = (a_{ij}) \in \text{Mat}(n, K)$ die Form

$$\chi_A = X^n - (a_{11} + a_{22} + \dots + a_{nn})X^{n-1} + \dots + (-1)^n \det(A)$$

BSP
diagonalisierbare Matrix
BSP
Nicht diagonalisierbare Matrix

FOLG
Faktorisierung von Polynomen über $\mathbb{C}$

hat; der Koeffizient von X^0 ergibt sich dabei aus

$$\chi_A(0) = \det(0 \cdot I_n - A) = \det(-A) = (-1)^n \det(A).$$

Wenn χ_A in Linearfaktoren zerfällt:

$$\begin{aligned}\chi_A &= (X - \lambda_1)(X - \lambda_2) \cdots (X - \lambda_n) \\ &= X^n - (\lambda_1 + \lambda_2 + \dots + \lambda_n)X^{n-1} + \dots + (-1)^n \lambda_1 \lambda_2 \cdots \lambda_n,\end{aligned}$$

dann sehen wir durch Vergleich der beiden Darstellungen von χ_A , dass wir Summe und Produkt der Eigenwerte einfach von der Matrix ablesen können. Bevor wir das als Lemma formulieren, ist hier noch eine Definition:

16.9. Definition. Seien K ein Körper, $n \in \mathbb{N}$ und $A = (a_{ij}) \in \text{Mat}(n, K)$. Die *Spur* von A ist

DEF
Spur einer
Matrix

$$\text{Tr}(A) = \sum_{i=1}^n a_{ii} = a_{11} + a_{22} + \dots + a_{nn}.$$

(„Tr“ von englisch *trace*.)

◇

16.10. Lemma. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Das charakteristische Polynom von A zerfällt in Linearfaktoren. Dann ist $\text{Tr}(A)$ die Summe und $\det(A)$ das Produkt der Eigenwerte von A , jeweils entsprechend ihrer algebraischen Vielfachheit gezählt.

LEMMA
Spur und
Determinante
durch
Eigenwerte

Da die charakteristischen Polynome ähnlicher Matrizen gleich sind, folgt

$$\text{Tr}(PAP^{-1}) = \text{Tr}(A).$$

Die Spur erfüllt aber sogar noch eine etwas stärkere Aussage.

* **16.11. Satz.** Die Spur ist eine K -lineare Abbildung $\text{Mat}(n, K) \rightarrow K$. Für alle $A \in \text{Mat}(n, K)$ gilt $\text{Tr}(A^\top) = \text{Tr}(A)$.

SATZ
Eigensch.
der Spur

Sind $A \in \text{Mat}(m \times n, K)$ und $B \in \text{Mat}(n \times m, K)$, dann gilt

$$\text{Tr}(AB) = \text{Tr}(BA).$$

Man beachte, dass AB eine $m \times m$ -Matrix und BA eine $n \times n$ -Matrix ist.

Beweis. Die erste Aussage ist klar (die Spur ist eine Linearkombination der Matrix-Einträge), die zweite ebenfalls, da $A^\top$ und A dieselben Diagonaleinträge haben. Für die dritte Aussage seien $A = (a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ und $B = (b_{kl})_{1 \leq k \leq n, 1 \leq l \leq m}$. Der Diagonaleintrag in der i -ten Zeile und Spalte von $C = AB$ ist

$$c_{ii} = \sum_{j=1}^n a_{ij} b_{ji} \quad (i \in \{1, 2, \dots, m\})$$

und der Diagonaleintrag in der j -ten Zeile und Spalte von $C' = BA$ ist

$$c'_{jj} = \sum_{i=1}^m b_{ji} a_{ij} \quad (j \in \{1, 2, \dots, n\}).$$

Es folgt

$$\begin{aligned}\operatorname{Tr}(AB) &= \operatorname{Tr}(C) = \sum_{i=1}^m c_{ii} = \sum_{i=1}^m \sum_{j=1}^n a_{ij} b_{ji} \\ &= \sum_{j=1}^n \sum_{i=1}^m b_{ji} a_{ij} = \sum_{j=1}^n c'_{jj} = \operatorname{Tr}(C') = \operatorname{Tr}(BA). \quad \square\end{aligned}$$

Die Gleichung $\operatorname{Tr}(A) = \operatorname{Tr}(PAP^{-1})$ folgt daraus:

$$\operatorname{Tr}(PAP^{-1}) = \operatorname{Tr}(P(AP^{-1})) = \operatorname{Tr}((AP^{-1})P) = \operatorname{Tr}(A(P^{-1}P)) = \operatorname{Tr}(A).$$

Analog zu Definition 15.23 können wir daher auch die Spur eines Endomorphismus definieren.

16.12. Definition. Seien K ein Körper, V ein endlich-dimensionaler K -Vektorraum und $f \in \operatorname{End}(V)$. Seien B eine beliebige Basis von V und $A = \operatorname{Mat}_{B,B}(f)$. Dann ist die *Spur* von f definiert als $\operatorname{Tr}(f) = \operatorname{Tr}(A)$. $\diamond$

DEF
Spur
eines Endomorphismus

Die Aussage, dass Spur und Determinante die Summe und das Produkt der Eigenwerte sind (mit algebraischer Vielfachheit gezählt), gilt dann entsprechend auch für Endomorphismen.

Nun wollen wir uns noch überlegen, wie man, wenn die Matrix $A \in \operatorname{Mat}(n, K)$ diagonalisierbar ist, eine Matrix $P \in \operatorname{GL}(n, K)$ findet, die A diagonalisiert, also sodass $PAP^{-1} = D$ eine Diagonalmatrix ist.

16.13. Lemma. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \operatorname{Mat}(n, K)$ eine diagonalisierbare Matrix. Sei $(\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n)$ eine Basis von K^n , die aus Eigenvektoren von A besteht, mit $A\mathbf{b}_j = \lambda_j \mathbf{b}_j$ (wir betrachten $\mathbf{b}_j$ als Spaltenvektor). Sei $Q \in \operatorname{GL}(n, K)$ die Matrix, deren j -te Spalte $\mathbf{b}_j$ ist für $j \in \{1, 2, \dots, n\}$. Dann ist

$$Q^{-1}AQ = \operatorname{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$$

eine Diagonalmatrix.

LEMMA
Diagonalisierung

Man kann dann also $P = Q^{-1}$ nehmen.

Beweis. Sei $\mathbf{e}_j$ der j -te Standard-Basisvektor von K^n als Spaltenvektor. Dann gilt $Q\mathbf{e}_j = \mathbf{b}_j$ und damit auch $Q^{-1}\mathbf{b}_j = \mathbf{e}_j$. Es folgt

$$Q^{-1}AQ\mathbf{e}_j = Q^{-1}A\mathbf{b}_j = \lambda_j Q^{-1}\mathbf{b}_j = \lambda_j \mathbf{e}_j.$$

Das zeigt, dass die j -te Spalte von $Q^{-1}AQ$ gerade $\lambda_j \mathbf{e}_j$ ist, also ist

$$Q^{-1}AQ = \operatorname{diag}(\lambda_1, \lambda_2, \dots, \lambda_n). \quad \square$$

16.14. Beispiel. Sei $A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \in \operatorname{Mat}(2, \mathbb{R})$. Dann hat A die beiden Eigenwerte 1 und -1 und ist daher diagonalisierbar. Wir hatten in Beispiel 15.10 Basen der beiden Eigenräume gefunden:

$$E_1(A) = \left\langle \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\rangle \quad \text{und} \quad E_{-1}(A) = \left\langle \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right\rangle.$$

BSP
Diagonalisieren einer Matrix

Eine geeignete Matrix Q ist demnach $Q = \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}$; mit $P = Q^{-1} = \frac{1}{2} \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix}$ ist dann $PAP^{-1} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. $\clubsuit$

Wir haben gesehen, dass man Elemente des Grundkörpers in ein Polynom einsetzen kann (Definition 15.12). Man kann aber auch allgemeinere Objekte in Polynome einsetzen, zum Beispiel Matrizen oder Endomorphismen.

16.15. Definition. Sei K ein Körper, sei $n \in \mathbb{N}$ und seien $p \in K[X]$ ein Polynom und $A \in \text{Mat}(n, K)$. Wie üblich setzen wir $A^0 = I_n$ (das Einselement des Matrizenrings $\text{Mat}(n, K)$) und $A^{k+1} = A \cdot A^k$ für $k \in \mathbb{N}$. Ist

$$p = a_0 + a_1X + a_2X^2 + \dots + a_mX^m,$$

dann definieren wir

$$p(A) = \sum_{j=0}^m a_j A^j = a_0 I_n + a_1 A + a_2 A^2 + \dots + a_m A^m \in \text{Mat}(n, K).$$

Ist V ein K -Vektorraum und f ein Endomorphismus von V , dann setzen wir analog $f^{\circ 0} = \text{id}_V$ und $f^{\circ(k+1)} = f \circ f^{\circ k}$ und definieren

$$p(f) = \sum_{j=0}^m a_j f^{\circ j} = a_0 \text{id}_V + a_1 f + a_2 f^{\circ 2} + \dots + a_m f^{\circ m} \in \text{End}(V). \quad \diamond$$

(Die Schreibweise $f^{\circ k}$ anstatt von f^k verwenden wir, um Verwechslungen mit der punktweise definierten Potenz von (z.B.) reellen Funktionen zu vermeiden.)

Dann ist Folgendes klar: Ist A die Matrix von f bezüglich einer Basis B von V , dann ist $p(A)$ die Matrix von $p(f)$ bezüglich B .

Wir notieren noch, dass die Abbildungen $K[X] \rightarrow \text{Mat}(n, K)$, $p \mapsto p(A)$, und $K[X] \rightarrow \text{End}(V)$, $p \mapsto p(f)$, schöne Eigenschaften haben.

16.16. Lemma. In der Situation von Definition 16.15 gilt für $p, q \in K[X]$:

$$(p + q)(A) = p(A) + q(A) \quad \text{und} \quad (p \cdot q)(A) = p(A) \cdot q(A)$$

bzw.

$$(p + q)(f) = p(f) + q(f) \quad \text{und} \quad (p \cdot q)(f) = p(f) \circ q(f).$$

Außerdem ist $1(A) = I_n$ und $1(f) = \text{id}_V$, wobei 1 das konstante Polynom 1 bezeichnet.

Beweis. Das folgt aus den Rechenregeln für Matrizen bzw. Endomorphismen (also daraus, dass $\text{Mat}(n, K)$ und $\text{End}(V)$ Ringe sind), zusammen mit

$$\lambda A = (\lambda I_n)A = A(\lambda I_n) \quad \text{und} \quad \lambda f = (\lambda \text{id}_V) \circ f = f \circ (\lambda \text{id}_V). \quad \square$$

Sind R und R' zwei Ringe, dann heißt eine Abbildung $\phi: R \rightarrow R'$ ein *Ringhomomorphismus*, wenn $\phi(1_R) = 1_{R'}$ ist und für alle $r_1, r_2 \in R$ gilt $\phi(r_1 + r_2) = \phi(r_1) + \phi(r_2)$ und $\phi(r_1 r_2) = \phi(r_1) \phi(r_2)$ (d.h., ϕ bildet die Einselemente aufeinander ab und ist mit Addition und Multiplikation verträglich). Die Abbildungen $K[X] \rightarrow \text{Mat}(n, R)$, $p \mapsto p(A)$, und $K[X] \rightarrow \text{End}(V)$, $p \mapsto p(f)$, sind also Ringhomomorphismen. Weil sie dadurch gegeben sind, dass man etwas in die Polynome einsetzt, heißen sie auch *Einsetzungshomomorphismen*.

Mehr über Ringe im Allgemeinen und Polynomringe im Besonderen gibt es in der „Einführung in die Zahlentheorie und algebraische Strukturen“.

DEF
Einsetzen
von Matrizen
und Endo-
morphis-
men
in Polynome

LEMMA
Einsetzungs-
homomor-
phismus

16.17. Beispiel. Seien $A \in \text{Mat}(n, K)$ und $\lambda \in K$, seien weiter $v \in E_\lambda(A)$ und $p \in K[X]$. Dann gilt

$$p(A) \cdot v = p(\lambda)v :$$

Ist $p = a_0 + a_1X + a_2X^2 + \dots + a_mX^m$, dann rechnet man

$$\begin{aligned} p(A) \cdot v &= (a_0I_n + a_1A + a_2A^2 + \dots + a_mA^m) \cdot v \\ &= a_0v + a_1A \cdot v + a_2A^2 \cdot v + \dots + a_mA^m \cdot v \\ &= a_0v + a_1\lambda v + a_2\lambda^2v + \dots + a_m\lambda^mv \\ &= (a_0 + a_1\lambda + a_2\lambda^2 + \dots + a_m\lambda^m)v \\ &= p(\lambda)v. \end{aligned}$$



Sei eine Matrix $A \in \text{Mat}(n, K)$ gegeben. Dann kann man sich fragen, ob es stets ein Polynom $p \in K[X]$ gibt mit $p(A) = \mathbf{0}$ (und $p \neq 0$), bzw. was man über die Menge solcher Polynome mit „Nullstelle“ A aussagen kann.

Die erste Frage kann man leicht beantworten.

16.18. Lemma. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Dann gibt es ein normiertes Polynom $p \in K[X]$ mit $\deg(p) \leq n^2$, sodass $p(A) = \mathbf{0}$ ist.

LEMMA

Existenz
von $p \in K[X]$
mit $p(A) = \mathbf{0}$

Beweis. Der Beweis ist eine schöne Anwendung grundlegender Resultate der linearen Algebra. $\text{Mat}(n, K)$ ist ein K -Vektorraum der Dimension n^2 , also müssen die $n^2 + 1$ Elemente $A^0, A^1, A^2, \dots, A^{n^2} \in \text{Mat}(n, K)$ linear abhängig sein. Es gibt also $\lambda_0, \lambda_1, \dots, \lambda_{n^2} \in K$, nicht alle null, mit $\sum_{j=0}^{n^2} \lambda_j A^j = \mathbf{0}$. Wir setzen $m = \max\{j \mid \lambda_j \neq 0\}$. Nach eventueller Multiplikation mit λ_m^{-1} können wir $\lambda_m = 1$ annehmen. Die Behauptung folgt dann mit $p = \sum_{j=0}^m \lambda_j X^j$. $\square$

Man beachte, dass der Satz, dass ein Polynom vom Grad n höchstens n Nullstellen hat, nicht für Matrizen gilt. Zum Beispiel gilt für jede Matrix $A = \begin{pmatrix} a & b \\ c & -a \end{pmatrix}$ mit $a^2 + bc = -1$, dass $A^2 + I_2 = \mathbf{0}$ ist. Alle diese Matrizen (und davon gibt es unendlich viele, wenn K unendlich ist) sind also „Nullstellen“ von $X^2 + 1$.

Die zweite Frage kann man wie folgt beantworten:

16.19. Lemma. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Sei

$$P(A) = \{p \in K[X] \mid p(A) = \mathbf{0}\}.$$

LEMMA

Struktur
von $P(A)$

Dann gilt:

- (1) $\mathbf{0} \in P(A)$.
- (2) Aus $p, q \in P(A)$ folgt $p + q \in P(A)$.
- (3) Aus $p \in P(A)$, $q \in K[X]$ folgt $q \cdot p \in P(A)$.

Beweis. Die erste Aussage ist klar. Die beiden anderen sieht man so:

$$(p+q)(A) = p(A) + q(A) = \mathbf{0} + \mathbf{0} = \mathbf{0} \quad \text{und} \quad (q \cdot p)(A) = q(A) \cdot p(A) = q(A) \cdot \mathbf{0} = \mathbf{0}.$$



Sei R ein kommutativer Ring. Eine Teilmenge $I \subset R$ heißt ein *Ideal* von R , wenn I die obigen Eigenschaften hat:

$$0 \in I, \quad r, r' \in I \Rightarrow r + r' \in I, \quad r \in R, r' \in I \Rightarrow rr' \in I.$$

Ist $\phi : R \rightarrow R'$ ein Ringhomomorphismus, dann zeigt derselbe Beweis wie oben, dass sein *Kern* $\ker(\phi) = \{r \in R \mid \phi(r) = 0\}$ ein Ideal von R ist.

Die Tatsache, dass es im Polynomring eine „Division mit Rest“ gibt, führt zu einer einfachen Beschreibung von $P(A)$.

16.20. Satz. Seien K ein Körper, $n \in \mathbb{N}$ und $A \in \text{Mat}(n, K)$. Dann gibt es ein eindeutig bestimmtes normiertes Polynom m_A kleinsten Grades in $P(A)$, und

$$P(A) = \{q \cdot m_A \mid q \in K[X]\}$$

besteht genau aus den Vielfachen von m_A .

SATZ
Minimal-
polynom

Beweis. Nach Lemma 16.18 gibt es normierte Polynome in $P(A)$. Sei m_A ein solches mit minimalem Grad. Ist p ein weiteres normiertes Polynom in $P(A)$ mit $\deg(p) = \deg(m_A)$, dann ist $p - m_A \in P(A)$ entweder das Nullpolynom (und damit $p = m_A$) oder $0 \leq \deg(p - m_A) < \deg(m_A)$. Nach Division durch den Leitkoeffizienten würde man dann ein normiertes Polynom in $P(A)$ mit kleinerem Grad als $\deg(m_A)$ erhalten, im Widerspruch zur Wahl von m_A . Also ist m_A eindeutig bestimmt. Wir sehen auch, dass $\deg(m_A) = \min\{\deg(p) \mid 0 \neq p \in P(A)\}$ ist.

Nach Lemma 16.19 ist klar, dass $\{q \cdot m_A \mid q \in K[X]\} \subset P(A)$ ist. Sei umgekehrt $p \in P(A)$. Dann gibt es $q, r \in K[X]$ mit $p = q \cdot m_A + r$ und $\deg(r) < \deg(m_A)$ (Satz 15.16). Es folgt $r = p - q \cdot m_A \in P(A)$, und wegen $\deg(r) < \deg(m_A)$ muss (mit demselben Argument wie oben) $r = 0$ sein. Damit ist $p = q \cdot m_A$, also hat man auch die umgekehrte Inklusion. $\square$

Derselbe Beweis zeigt, dass jedes Ideal I von $K[X]$ die Form $I = \{p \cdot a \mid p \in K[X]\}$ hat mit einem geeigneten $a \in K[X]$ ($a = 0$ ist möglich, dann ist $I = \{0\}$). So ein Ideal, dessen Elemente genau die Vielfachen eines Elements a sind, heißt ein *Hauptideal*, und ein Ring, in dem jedes Ideal ein Hauptideal ist, ist ein *Hauptidealring*. Wir haben also gezeigt, dass $K[X]$ ein solcher Hauptidealring ist. Ein anderes Beispiel für einen Hauptidealring ist der Ring $\mathbb{Z}$ der ganzen Zahlen. Der Beweis ist im Wesentlichen derselbe, nur dass man Division mit Rest von ganzen Zahlen verwendet statt der Polynomdivision.

16.21. Definition. Das Polynom m_A in Satz 16.20 heißt das *Minimalpolynom* von A . Ist f ein Endomorphismus eines endlich-dimensionalen Vektorraums, der bezüglich einer geeigneten Basis durch A beschrieben wird, dann heißt $m_f = m_A$ das *Minimalpolynom* von f . $\diamond$

DEF
Minimal-
polynom

16.22. Beispiele. Seien $n \geq 1$ und $A \in \text{Mat}(n, K)$.

- (1) Sei $A = 0$ die Nullmatrix. Dann ist $m_A = X$.
- (2) Ist $A = \lambda I_n$, dann ist $m_A = X - \lambda$. Die Umkehrung gilt ebenfalls.
- (3) Ist A diagonalisierbar und hat die (paarweise verschiedenen) Eigenwerte $\lambda_1, \dots, \lambda_m$, dann ist

$$m_A = (X - \lambda_1)(X - \lambda_2) \cdots (X - \lambda_m).$$

Das Minimalpolynom zerfällt also in Linearfaktoren und hat keine mehrfachen Nullstellen. Das sieht man so: Jedes Element v von K^n lässt sich

BSP
Minimal-
polynome

schreiben als $v = v_1 + \dots + v_m$ mit $v_j \in E_{\lambda_j}(A)$. Wegen $Av_j = \lambda_j v_j$ ist dann mit $p = (X - \lambda_1)(X - \lambda_2) \cdots (X - \lambda_m)$ und Beispiel 16.17

$$\begin{aligned} p(A) \cdot v &= p(A) \cdot v_1 + p(A) \cdot v_2 + \dots + p(A) \cdot v_m \\ &= p(\lambda_1)v_1 + p(\lambda_2)v_2 + \dots + p(\lambda_m)v_m = \mathbf{0}, \end{aligned}$$

also ist p ein normiertes Polynom in $P(A)$. Auf der anderen Seite gibt es für jedes $j \in \{1, 2, \dots, m\}$ ein $\mathbf{0} \neq v_j \in E_{\lambda_j}(A)$, und es muss gelten

$$\mathbf{0} = m_A(A) \cdot v_j = m_A(\lambda_j)v_j, \quad \text{also ist } m_A(\lambda_j) = 0.$$

Das Minimalpolynom m_A muss also alle λ_j als Nullstellen haben. Es folgt $\deg(m_A) \geq m = \deg(p)$. Nach der Definition des Minimalpolynoms ist demnach $m_A = p$.

Da das charakteristische Polynom von A die Form

$$(X - \lambda_1)^{e_1}(X - \lambda_2)^{e_2} \cdots (X - \lambda_m)^{e_m}$$

hat, wobei die $e_j \geq 1$ (für $j \in \{1, 2, \dots, m\}$) die algebraischen Vielfachheiten der Eigenwerte sind, sehen wir, dass das charakteristische Polynom ein Vielfaches des Minimalpolynoms ist (also $\chi_A = q \cdot m_A$ mit einem Polynom $q \in K[X]$). Hat A n verschiedene Eigenwerte, dann sind die beiden Polynome gleich. 

Das dritte Beispiel oben wirft zwei Fragen auf:

- Für A diagonalisierbar gilt, dass χ_A ein Vielfaches von m_A ist. Ist das auch allgemein richtig?
- Gilt die Umkehrung der Beobachtung im Beispiel: Wenn m_A in Linearfaktoren zerfällt und keine mehrfachen Nullstellen hat, muss dann A diagonalisierbar sein?

Es stellt sich heraus, dass beide Fragen mit „Ja“ beantwortet werden können. Das werden wir zu Beginn des Sommersemesters sehen.