

Mathematik für Wirtschaftswissenschaftler

Skript zur Vorlesung im WS 2013/2014

Andreas Kollross

Version: 26. Februar 2014

Inhaltsverzeichnis

Vorwort	4
Kapitel 1. Grundlagen	5
1. Mengen und Zahlen	5
1.1. Die Mengenschreibweise	5
1.2. Der Aufbau des Zahlensystems	8
1.3. Summen- und Produktschreibweise, Binomialkoeffizient	12
1.4. Vollständige Induktion	15
1.5. Betrag, Ungleichungen, Intervalle	16
1.6. Komplexe Zahlen	18
2. Vektoren	21
2.1. Der $\mathbb{R}^n$.	21
2.2. Skalarprodukt und Abstandsfunktion	21
3. Abbildungen	23
3.1. Injektivität und Surjektivität	24
3.2. Lineare Abbildungen	25
3.3. Matrizen	27
3.4. Polynome	29
4. Lineare Gleichungssysteme	32
4.1. Das Gaußverfahren	32
4.2. Lineare Gleichungssysteme und Matrizen	37
4.3. Die Inverse einer Matrix	38
4.4. Untervektorräume	41
4.5. Lineare Unabhängigkeit und Basen	43
5. Grenzwerte und Stetigkeit	45
5.1. Folgen und Konvergenz	45
5.2. Reihen	49
5.3. Funktionsgrenzwerte und Stetigkeit	55
6. Differentiation	62
6.1. Ableitung einer Funktion in einer Variablen	62

6.2.	Die Regel von L'Hôpital	65
6.3.	Partielle Ableitungen	66
6.4.	Totale Differenzierbarkeit und Jacobimatrix	68
Kapitel 2.	Optimierung	72
1.	Lokale Extrema (ohne Nebenbedingungen)	72
1.1.	Der Gradient	73
1.2.	Die Hessematrix	73
1.3.	Definitheit und Determinante	74
1.4.	Kriterium für lokale Extrema	78
2.	Lokale Extrema unter Nebenbedingungen	79
2.1.	Flachpunkte unter Nebenbedingungen	81
2.2.	Kriterium für Extrema unter Nebenbedingungen	84
3.	Kompakta	87
Kapitel 3.	Approximation	89
1.	Taylorpolynome	89
1.1.	Taylorpolynom in einer Variablen	89
1.2.	Taylorpolynom in mehreren Variablen	93
2.	Das Newton-Verfahren	93
2.1.	Newton-Verfahren, eindimensional	93
2.2.	Newton-Verfahren, mehrdimensional	96
Kapitel 4.	Differentialgleichungen	97
1.	Begriffsbestimmung und Beispiele	98
2.	Trennung der Variablen	101
Index		105

Vorwort

Dieses Skript entsteht zur Zeit parallel zu meiner Vorlesung Mathematik für Wirtschaftswissenschaftler an der Universität Stuttgart. Ich werde es abschnittsweise online stellen. Falls Sie Fehler und Unstimmigkeiten bemerken oder andere Kommentare haben, schicken Sie mir bitte eine E-Mail. Beachten Sie bitte, dass auch Kapitel, die schon eine Zeit lang online stehen, später noch verbessert oder ergänzt werden können.

Andreas Kollross

KAPITEL 1

Grundlagen

1. Mengen und Zahlen

1.1. Die Mengenschreibweise. Ein grundlegendes und unverzichtbares Konzept für die heutige Mathematik ist der Begriff der *Menge*. In Mengen werden einzelne Gegenstände oder Objekte zu einem größeren Ganzen zusammengefasst, in der Mathematik sind das oft Zahlbereiche, Lösungsmenge von Gleichungen oder etwa geometrische Figuren, also Mengen von Punkten in der Ebene oder im Raum, um nur einige Beispiele zu nennen.

Definition 1.1.1. Eine *Menge* ist eine wohldefinierte Gesamtheit von Objekten, den *Elementen* der Menge. Ist das Objekt x ein Element der Menge, so drückt man dies durch $x \in M$ aus, entsprechend bedeutet $x \notin M$, dass x kein Element von M ist.

Beispielsweise kann man die Menge $\mathbb{N}$ der natürlichen Zahlen betrachten. Dann gilt z.B.

$$1 \in \mathbb{N}, \quad 2 \in \mathbb{N}, \quad 9999 \in \mathbb{N}, \quad \frac{1}{2} \notin \mathbb{N}, \quad -1 \notin \mathbb{N}.$$

Mengen kann man auf verschiedene Arten festlegen:

Durch Beschreiben: Man erklärt in Worten, was in der Menge enthalten sein soll. Um deutlich zu machen, dass eine Menge definiert werden soll, setzt man die Beschreibung in geschweifte Klammern:

$$\mathbb{N} := \{\text{natürliche Zahlen}\}$$

Der Doppelpunkt auf der linken Seite des Gleichheitszeichen macht darauf aufmerksam, dass das Objekt auf der linken Seite (nämlich “ M ”) durch das, was auf der rechten Seite steht, definiert werden soll.

Durch Aufzählen: Man schreibt alle Elemente der Menge zwischen geschweifte Klammern, zum Beispiel:

$$M := \{1, 2, 3, 4, 5\}.$$

Man beachte, dass es auf die Reihenfolge dabei nicht ankommt, es gilt also

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

Außerdem gilt, dass ein Element entweder in der Menge enthalten ist oder nicht, insbesondere ist ein Element auch nur einmal enthalten, wenn es mehrfach aufgezählt wird, also gilt z.B.: $\{1, 1, 1, 2, 3\} = \{1, 2, 3\}$. Eine Variante ist die aufzählende Schreibweise mit Pünktchen als Auslassungszeichen:

$$M := \{1, 2, \dots, 5\}, \quad \mathbb{N} := \{1, 2, 3, 4, \dots\}.$$

Durch Aussondern aus einer anderen Menge: Es werden aus einer vorhandenen Menge alle Elemente x mit einer bestimmten Eigenschaft ausgewählt. Zum Beispiel kann man die Elemente der obigen Menge M als die natürlichen Zahlen kennzeichnen, die kleiner als 6 sind:

$$M = \{x \in \mathbb{N} \mid x < 6\}.$$

Durch Mengenoperationen: Die *Vereinigung* $A \cup B$ zweier Mengen A und B ist die Menge aller Objekte x , so dass x ein Element von A oder ein Element von B ist, zum Beispiel:

$$\{1, 2, 3\} \cup \{1, 8, 9\} = \{1, 2, 3, 8, 9\}.$$

Die *Schnittmenge* $A \cap B$ zweier Mengen A und B ist die Menge aller Objekte, die sowohl Elemente von A als auch von B sind, zum Beispiel:

$$\{1, 2, 3, 4\} \cap \{2, 4, 6, 8, 10\} = \{2, 4\}.$$

Die *Differenz* $A \setminus B$ zweier Mengen, besteht aus allen Elementen von A , die nicht in B enthalten sind, zum Beispiel:

$$\{1, 2, 3, 4, 5, 6\} \setminus \{3, 5, 7\} = \{1, 2, 4, 6\}.$$

Die *leere Menge* ist die Menge, die kein einziges Element enthält, man schreibt dafür das Symbol $\emptyset$ oder benutzt leere Mengenklammern, d.h.

$$\emptyset = \{\}.$$

Zwei Mengen sind genau dann gleich, wenn Sie, unabhängig von der Art, wie sie definiert sind, die gleichen Elemente umfassen; diese Eigenschaft von Mengen wird auch als *Extensionalität des Mengenbegriff* bezeichnet. Es gilt also etwa:

$$\mathbb{N} = \{1, 2, 3, \dots\} = \{\text{natürliche Zahlen}\} = \{\text{positive ganze Zahlen}\}.$$

Die Anzahl der Elemente einer Menge wird auch als die *Mächtigkeit* (auch Kardinalität) der Menge bezeichnet, man schreibt $|M|$ für die Mächtigkeit von M , z.B.:

$$A := \{7, 5, 3, 2\} \Rightarrow |A| = 4, \quad |\mathbb{N}| = \infty.$$

Eine Menge der Mächtigkeit n wird auch als *n-elementige Menge* bezeichnet.

Das *kartesische Produkt* $A \times B$ zweier Mengen A und B ist die Menge der geordneten Paare (a, b) , wobei $a \in A$ und $b \in B$ ist. Zum Beispiel gilt:

$$\{1, 2\} \times \{1, 2, 5\} = \{(1, 1), (1, 2), (1, 5), (2, 1), (2, 2), (2, 5)\}.$$

Man beachte, dass es bei der Schreibweise (a, b) , im Gegensatz zur Mengenschreibweise, sehr wohl auf die Reihenfolge ankommt, (a, b) ist im Allgemeinen *nicht* dasselbe wie (b, a) (es sei denn $a = b$) und es gilt auch nicht, dass (a, a) etwa gleich (a) wäre. Man kann ein kartesisches Produkt auch mit drei oder mehr Mengen bilden, etwa ist $A \times B \times C$ die Menge der geordneten *Tripel* (a, b, c) , wobei $a \in A$, $b \in B$ und $c \in C$ gilt. Allgemeiner bezeichnet man die Elemente eines kartesischen Produkts von n Mengen als (*geordnete*) *n-Tupel*.

Eine Menge A heißt *Teilmenge* einer Menge B , wenn jedes Element von A auch in B als Element enthalten ist. In diesem Fall schreibt man $A \subseteq B$ oder auch $A \subset B$. Will man darauf hinweisen, dass es sich um eine *echte Teilmenge* handelt, d.h. jedes Element von A ist auch Element von B , aber nicht umgekehrt, dann schreibt man auch manchmal $A \subsetneq B$ oder $A \subsetneqq B$.

Man kann auch Mengen von Mengen bilden. Etwa enthält $\{1, \{1\}\}$ die Zahl 1 und die Menge, die nur die Zahl 1 als Element enthält (das ist nicht dasselbe!). Weitaus nützlicher ist die folgende Konstruktion: Für irgendeine Menge A bezeichnet man mit $\text{Pot}(A)$ die *Potenzmenge* von A , das ist die Menge aller Teilmengen von A . Zum Beispiel gilt

$$\text{Pot}(\{1, 2, 3\}) = \{\{\}, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\}\}.$$

Bei einer endlichen Menge mit n Elementen gilt, dass ihre Potenzmenge 2^n Elemente hat.

1.2. Der Aufbau des Zahlensystems. Die natürlichen Zahlen

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

bilden die Ausgangsbasis unserer Konstruktion des Zahlensystems. Man erhält sie, wenn man eine endliche Menge von Gegenständen abzählt. Es gibt unendlich viele, hat man irgendeine natürliche Zahl, kann man sie immer noch um eins erhöhen, es kann also keine größte geben. (“Unendlich”, oder “ ∞ ”, ist *keine* natürliche Zahl.) Manchmal nimmt man noch die Null hinzu, und bezeichnet die so definierte Menge mit

$$\mathbb{N}_0 = \{0, 1, 2, 3, \dots\}.$$

Je zwei Zahlen aus $\mathbb{N}_0$ kann man addieren oder multiplizieren und erhält dabei wieder eine Zahl aus $\mathbb{N}_0$. Allerdings ist die Subtraktion nicht uneingeschränkt möglich, deshalb führen wir die *ganzen Zahlen* ein, wir erhalten sie aus $\mathbb{N}_0$, indem wir Vorzeichen zulassen,

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

auf diese Weise können wir z.B. Kontostände mit Soll und Haben ausdrücken. In den ganzen Zahlen ist Addition, Subtraktion und Multiplikation uneingeschränkt möglich, aber nicht die Division (ohne Rest). Diesen Mangel behebt man, indem man die *rationalen Zahlen* $\mathbb{Q}$ einführt, d.h. die Brüche, wobei auch negative Brüche und die Null eingeschlossen sind, etwa so:

$$\mathbb{Q} = \left\{ \frac{m}{n} \mid m \in \mathbb{Z}, n \in \mathbb{N} \right\}$$

Man erinnere sich daran, dass die Darstellung von rationalen Zahlen durch Brüche *nicht eindeutig* ist: $\frac{1}{2} = \frac{2}{4} = \frac{3}{6} = \frac{1000}{2000}$ bezeichnen alle dieselbe Zahl. Erst durch *Kürzen*, also Darstellen mit teilerfremdem Zähler und Nenner, kann man eine eindeutige Schreibweise erreichen.

Die obigen Darstellung einer rationalen Zahl mit Nenner und Zähler nennt man einen *echten Bruch*, daneben wird oft auch die Schreibweise als *gemischter Bruch* verwendet, wobei zunächst der ganzzahlige Anteil (die Zahl zur Null hin gerundet) geschrieben wird und rechts daneben der restliche Anteil als echter Bruch, z.B. $\frac{3}{2} = 1\frac{1}{2}$, oder $-\frac{12}{5} = -2\frac{2}{5}$. Diese Schreibweise ist nicht konform mit der üblichen Praxis, ein fehlendes Verknüpfungszeichen als Multiplikation zu interpretieren, so ist $a\frac{b}{c}$ als $a \cdot \frac{b}{c}$ zu interpretieren, während mit $1\frac{1}{2}$ die Zahl $1 + \frac{1}{2}$ gemeint ist.

Die rationalen Zahlen bilden einen Zahlbereich, in dem Addition, Subtraktion, Multiplikation und Division, außer durch Null, uneingeschränkt möglich sind. Dennoch haben die rationalen Zahlen “Lücken”! Man betrachte etwa ein Quadrat der

Seitenlänge eins. Nach dem Satz des Pythagoras ist die Länge der Diagonalen gleich $\sqrt{2}$. Diese Zahl, die Quadratwurzel von 2, ist *keine rationale Zahl*. Denn angenommen, es gäbe es eine positive rationale Zahl $\frac{m}{n}$, deren Quadrat gleich 2 ist, dann wäre $(\frac{m}{n})^2$ eine Zahl mit einer geraden Anzahl von Primfaktoren im Zähler als auch im Nenner, was $\frac{m^2}{n^2} = 2$ widerspricht, denn die (eindeutige) Primfaktorzerlegung von 2 besteht nur aus einem einzigen Faktor, also einer ungeraden Anzahl.

Um diese Lücken zu füllen, führt man die reellen Zahlen ein und ergänzt $\mathbb{Q}$ um die sogenannten *irrationalen Zahlen* zum Zahlbereich $\mathbb{R}$ der *reellen Zahlen*. Diese kann man darstellen als *unendliche Dezimalbrüche*, also Kommazahlen mit unendlich vielen Stellen nach dem Komma. Die rationalen Zahlen sind in dieser Darstellungsweise genau die *periodischen Dezimalbrüche*.

Das Umwandeln von gewöhnlichen Brüchen in (unendliche) Dezimalbrüche geschieht durch das übliche schriftliche Divisionsverfahren, erhält man dabei eine *Periode*, d.h. eine Folge sich wiederholender Ziffern hinter dem Komma, werden diese durch Überstreichen gekennzeichnet:

$$\frac{1}{7} = 0,\overline{142857}.$$

Will man umgekehrt einen periodischen Dezimalbruch in einen gewöhnlichen Bruch (auch *gemeinen Bruch* genannt) umwandeln, sollte man sich an folgenden Trick erinnern: Hat man einen periodischen Dezimalbruch, der vor dem Komma eine Null stehen hat und dessen Periode sofort hinter dem Komma beginnt, dann gilt z.B.:

$$999 \cdot 0,\overline{281} = (1000 - 1) \cdot 0,\overline{281} = 281,\overline{281} - 0,\overline{281} = 281.$$

Daraus folgt: $0,\overline{281} = \frac{281}{999}$. Insbesondere gilt $0,\overline{9} = 1$.

Eine alternative Konstruktion für die reellen Zahlen ist durch sogenannte *Intervallschachtelungen* gegeben. Dabei stellt man sich ineinander geschachtelte Intervalle mit rationalen Intervallgrenzen vor

$$[a_1, b_1] \supseteq [a_2, b_2] \supseteq [a_3, b_3] \supseteq \dots$$

vor, deren Längen beliebig klein werden. Die reellen Zahl sind *vollständig* in dem Sinn, dass durch jede solche Intervallschachtelung genau eine reelle Zahl gegeben ist, also eine, die in allen Intervallen enthalten ist. In diesem Sinne kann man sich die reellen Zahlen als eine Gerade (die reelle Zahlengerade) vorstellen, auf der es keine "Lücken" gibt.

Sowohl die rationalen als auch die reellen Zahlen bilden einen *Körper*, d.h. einen Zahlbereich, in dem Addition, Subtraktion, Multiplikation und Division (außer durch Null) uneingeschränkt möglich sind und in dem die üblichen Rechenregeln gelten. Eine genaue Definition folgt nun. Man beachte, dass man dabei, um die Formulierung nicht unnötig lang zu machen, nur die Addition und Multiplikation als Verknüpfungen fordert und die Subtraktion bzw. Division dadurch ersetzt, dass man das entsprechende inversen Element (die negative Zahl bzw. den Kehrwert) addiert bzw. mit ihm multipliziert.

Definition 1.2.1. Ein *Körper* ist ein Zahlbereich K (eine Menge von Zahlen), in dem es zwei Verknüpfungen, die Addition und die Multiplikation, gibt, für die folgende Gesetze gelten:

- (i) Es gibt ein *neutrales Element der Addition*, die Null 0, so dass $a + 0 = a$ für alle Elemente von K gilt.
- (ii) Zu jedem Element $a \in K$ gibt es ein *inverses Element bezüglich der Addition*, für das wir $-a$ schreiben, so dass gilt $a + (-a) = 0$.
- (iii) Es gibt *neutrales Element der Multiplikation*, die Eins 1, so dass $1 \cdot a = a$ für alle Elemente von K gilt.
- (iv) Zu jedem Element a von K außer der Null gibt es ein *inverses Element bezüglich der Multiplikation*, für das wir a^{-1} schreiben, so dass gilt $a \cdot a^{-1} = 1$.
- (v) Außerdem gelten die *Kommutativgesetze* für die Addition und Multiplikation:

$$a + b = b + a, \quad a \cdot b = b \cdot a,$$

die *Assoziativgesetze* jeweils für Addition und Multiplikation:

$$a + (b + c) = (a + b) + c, \quad a \cdot (b \cdot c) = (a \cdot b) \cdot c$$

und das *Distributivgesetz*

$$a \cdot (b + c) = a \cdot b + a \cdot c$$

für alle Elemente a, b, c von K .

Die obigen Gesetze werden auch *Körperaxiome* genannt. Aus ihnen kann man eine Reihe weiterer Rechenregeln ableiten. Der Vorteil dieser etwas abstrakten Sichtweise ist eine immense Arbeitersparnis, die darin besteht, dass man aus den Axiomen abgeleitete Rechenregeln für *alle* Körper gelten. Als Beispiel dafür wollen wir eine Aussage, die in jedem Körper gilt, beweisen.

Satz 1.2.2. *Sei K ein Körper. Seien a und b zwei Elemente von K . Dann hat die Gleichung $a + x = b$ genau eine Lösung und zwar $x = b - a$.*

BEWEIS. Natürlich ist diese Aussage für jede(n), der mit rationalen oder reellen Zahlen rechnen gelernt hat, “klar”. Man bringt a mit einem Minuszeichen auf die andere Seite, fertig! ... da wir aber diese Aussage ganz allgemein für *beliebige* Körper beweisen sollen, können wir hier nicht unser Schulwissen anwenden, sondern dürfen ausschließlich mit den Körperaxiomen aus Definition 1.2.1 arbeiten. Wir müssen also Schritt für Schritt aus der Aussage $a + x = b$ die Aussage $x = b - a$ herleiten. Dabei wollen wir in jedem Schritt nur ein einziges Axiom verwenden. Sei also

$$a + x = b$$

gegeben. Nach Axiom (iii) in Definition 1.2.1 gibt es zu a ein inverses Element bezüglich der Addition, das mit $(-a)$ bezeichnet wird. Wenn wir es auf beiden Seiten zur obigen Gleichung addieren, erhalten wir

$$(a + x) + (-a) = b + (-a).$$

Nun wenden wir auf der linken Seite in der linken Klammer das Kommutativgesetz für der Addition an und erhalten daraus

$$(x + a) + (-a) = b + (-a).$$

Auf die linke Seite können wir jetzt das Assoziativgesetz für die Addition anwenden, was

$$x + (a + (-a)) = b + (-a).$$

ergibt; wieder nach dem Axiom (iii) in Definition 1.2.1 gilt nun

$$x + 0 = b + (-a).$$

Dies führt mit Axiom über das neutrale Element der Addition zu

$$x = b + (-a),$$

was die Behauptung war, wenn man berücksichtigt, dass in der Sprache von Definition 1.2.1 die Subtraktion von a durch das Addieren von $-a$ ausgedrückt werden muss. $\square$

Abschließend halten wir für den Aufbau des Zahlensystems die folgende Kette von *Inklusionen* (=Teilmengenbeziehungen) fest:

$$\mathbb{N} \subset \mathbb{N}_0 \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R}.$$

1.3. Summen- und Produktschreibweise, Binomialkoeffizient. Sind Zahlen $a_m, \dots, a_n$ vorgegeben, so kann man ihre Summe durch das *Summenzeichen* notieren:

$$\sum_{i=m}^n a_i := a_m + \dots + a_n.$$

Hier nennt man $a_m, \dots, a_n$ die *Summanden*, i die *Summationsvariable* oder den *Laufindex*; m und n die *untere* bzw. *obere Summationsgrenze*. Gilt $m \geq n$, so ist die Summe *leer* und ihr Wert gleich Null. Zum Beispiel gilt

$$\sum_{i=1}^5 i = 1 + 2 + 3 + 4 + 5 = 15.$$

Satz 1.3.1. Für $n \in \mathbb{N}_0$ gilt

$$\sum_{i=1}^n i = \frac{1}{2}n(n+1).$$

BEWEIS. Diese Formel kann man sich recht leicht durch geometrische Anschauung überlegen. Man stellt sich ein Quadrat der Kantenlänge n , bestehend aus n mal n kleineren Kästchen vor. In der ersten Zeile ist ein Kästchen eingefärbt, in der zweiten zwei Kästchen usw. bis zur letzten Zeile, in der n Kästchen eingefärbt sind. Wie groß ist nun die gesamte eingefärbte Fläche? Die Fläche unterhalb der Diagonalen hat den Flächeninhalt $\frac{n^2}{2}$, da aber die Kästchen auf der Diagonalen ganz eingefärbt sind, muss man noch n halbe Kästchen dazu nehmen. Insgesamt ist also die Fläche $\frac{n^2}{2} + \frac{n}{2} = \frac{1}{2}n(n+1)$ eingefärbt. (Wir werden später noch eine andere Möglichkeit zum Beweis kennenlernen.) $\square$

Analog definiert man für Zahlen $a_m, \dots, a_n$ das Produkt durch das *Produktzeichen*:

$$\prod_{i=m}^n a_i := a_m \cdot \dots \cdot a_n.$$

Zu beachten ist, dass das *leere Produkt*, das man im Fall $m \geq n$ erhält, den Wert *eins* hat (und nicht etwa Null). Mit Hilfe des Produktzeichens können wir *Potenzen* von reellen Zahlen mit natürlichen Zahlen als Exponenten wie folgt schreiben:

$$x^n := \prod_{i=1}^n x.$$

Man beachte, dass mit unserer Definition $x^0 = 1$ für alle reellen Zahlen gilt, insbesondere gilt auch $0^0 = 1$. Für von Null verschiedene x definiert man negative

Exponenten durch

$$x^{-n} := (x^{-1})^n.$$

Mit diesen Definitionen gelten für $x \neq 0$ die üblichen Potenzrechengesetze

$$x^n x^m = x^{n+m}, \quad (x^n)^m = x^{nm}.$$

Für das Produkt der ersten n natürlichen Zahlen, $n \in \mathbb{N}_0$, gibt es eine besondere Bezeichnung, die *Fakultät* von n :

$$n! := \prod_{i=1}^n i = 1 \cdot 2 \cdot 3 \cdot \dots \cdot n,$$

gesprochen “ n Fakultät”. Die Fakultät von Null ist gleich dem leeren Produkt, also gilt $0! = 1$. Die Fakultät gibt die Anzahl *Permutationen* von n Objekten an, d.h. der Möglichkeiten an, n Objekte in verschiedene Reihenfolgen zu bringen. Z.B. gibt es $3! = 6$ Möglichkeiten, die Ziffern 1,2,3 zu permutieren:

$$123, 132, 213, 231, 312, 321.$$

Mit Hilfe der Fakultät lässt sich auch der *Binomialkoeffizient* definieren:

$$\binom{n}{k} := \frac{n!}{k!(n-k)!}$$

gesprochen “ n über k “. Er gibt die Anzahl der k -elementigen Teilmenge einer Menge mit n Elementen an. Diese Anzahl ist gleich der Anzahl der $(n-k)$ -elementigen Teilmengen und es gilt

$$\binom{n}{k} = \binom{n}{n-k}.$$

Seinen Namen hat der Binomialkoeffizient daher, dass er als Koeffizient im folgenden Satz auftritt.

Satz 1.3.2 (Binomischer Lehrsatz). *Für reelle Zahlen x, y und $n \in \mathbb{N}_0$ gilt*

$$(x+y)^n = \sum_{k=0}^n \binom{n}{k} x^{n-k} y^k.$$

Zum Beispiel erhalten wir aus dem Satz für $n = 2$ die wohlbekannte *binomische Formel* $(x+y)^2 = x^2 + 2xy + y^2$.

BEWEIS. Um sich die Gültigkeit der Formel zu überlegen, reicht es, die n Faktoren in $(x + y)^n$ auszuschreiben und (gedanklich) auszumultiplizieren:

$$\underbrace{(x + y)(x + y) \dots (x + y)}_{n \text{ Faktoren}}.$$

Da es $\binom{n}{k}$ Möglichkeiten gibt, in jeder Klammer den linken Summanden x genau k -mal und den rechten Summanden y genau $(n - k)$ -mal auszuwählen, ergibt sich die Aussage des Satzes. $\square$

Ein weitere Möglichkeit, Binomialkoeffizienten zu berechnen, ist das *Pascalsche Dreieck*. Dabei schreibt man in Dreiecksform, beginnend mit einer Eins, wie folgt Zahlen untereinander, wobei sich jede Zahl aus der Summe der beiden jeweils links und rechts über ihr stehenden ergeben, die fehlenden Zahlen am Rand werden als Nullen interpretiert:

$$\begin{array}{ccccccccccc}
 & & & & & & & & & & 1 \\
 & & & & & & & & & & & 1 \\
 & & & & & & & & & & 1 & 2 & 1 \\
 & & & & & & & & & 1 & 3 & 3 & 1 \\
 & & & & & & & & 1 & 4 & 6 & 4 & 1 \\
 & & & & 1 & 5 & 10 & 10 & 5 & 1 \\
 & & 1 & 6 & 15 & 20 & 15 & 6 & 1 \\
 & 1 & 7 & 21 & 35 & 35 & 21 & 7 & 1
 \end{array}$$

Das Pascalsche Dreieck, unendlich fortgesetzt, besteht genau aus allen Binomialkoeffizienten, wobei $\binom{n}{k}$ in der $(n + 1)$ -Zeile an der $(k + 1)$ -Stelle steht Diese Aussage, nämlich, dass sich jeder der Binomialkoeffizienten aus der Summe der beiden über ihm stehenden ergibt, wollen nun beweisen.

Satz 1.3.3. *Es gilt*

$$\binom{n + 1}{k + 1} = \binom{n}{k} + \binom{n}{k + 1}$$

BEWEIS. Wir rechnen nach:

$$\begin{aligned} \binom{n}{k} + \binom{n}{k+1} &= \frac{n!}{k!(n-k)!} + \frac{n!}{(k+1)!(n-k-1)!} = \\ &= \frac{n!}{k!(n-k)!} + \frac{n!(n-k)}{k!(k+1)(n-k)!} = \frac{n!(k+1) + n!(n-k)}{k!(k+1)(n-k)!} = \\ &= \frac{n!(n+1)}{(k+1)!(n-k)!} = \frac{(n+1)!}{(k+1)!(n-k)!} = \binom{n+1}{k+1}. \end{aligned}$$

□

1.4. Vollständige Induktion. Die *vollständige Induktion* ist ein Beweisprinzip, mit dem man Aussagen beweisen kann, die für alle natürlichen Zahlen richtig sind. Im Unterschied zu anderen Beweisverfahren wird dabei eine Aussage nicht sofort für beliebige natürliche Zahlen bewiesen, sondern nur für einen festen Startwert sowie für den Übergang von einer Zahl auf die nächsthöhere. Sei $A(n)$ für jede natürliche Zahl $n \in \mathbb{N}_0$ eine Aussage und sei $n_0 \in \mathbb{N}_0$ die kleinste Zahl, für die wir die Aussage zeigen wollen, dann funktioniert ein Beweis durch vollständige Induktion wie folgt:

- (i) Beweise, dass $A(n_0)$ gilt.
- (ii) Beweise: Falls $A(n)$ gilt, folgt daraus $A(n+1)$.

Damit sind dann unendlich viele Aussagen

$$A(n_0), A(n_0+1), A(n_0+2), \dots$$

bewiesen. Der Schritt (i) heißt *Induktionsanfang*, der Schritt (ii) wird *Induktionsschritt* genannt. Beim Beweis von (ii) geht man so vor: Man nimmt an, dass $A(n)$ gilt, diese Annahme ist die *Induktionsvoraussetzung* und leitet daraus die *Induktionsbehauptung* $A(n+1)$ her.

Ein Paradebeispiel für eine solche Aussage $A(n)$ ist der Satz 1.3.1. Wir wollen ihn noch einmal beweisen, diesmal mit vollständiger Induktion.

BEWEIS VON SATZ 1.3.1. mit vollständiger Induktion:

- (i) Induktionsanfang: Wir zeigen die Aussage $A(0)$, d.h. zu beweisen ist $\sum_i 1^0 = 0$. Diese Aussage ist offensichtlich richtig, es handelt sich ja um eine leere Summe auf der linken Seite.

- (ii) Induktionsschritt: Wir nehmen an, $A(n)$ sei richtig und leiten daraus $A(n+1)$ her. Es gilt also

$$\sum_{i=1}^n i = \frac{1}{2}n(n+1).$$

Indem wir zu diesem Wert $(n+1)$ dazu addieren, erhalten wir den Wert von

$$\begin{aligned} \sum_{i=1}^{n+1} i &= \frac{1}{2}n(n+1) + (n+1) = \\ &= \frac{n^2 + n + 2n + 2}{2} = \frac{(n+1)(n+2)}{2}. \end{aligned}$$

Dies ist genau die Aussage $A(n+1)$ und damit ist der Induktionsschritt gezeigt.

□

Als weitere Anwendung der vollständigen Induktion wollen wir den folgenden Satz beweisen.

Satz 1.4.1. *Die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge beträgt $\binom{n}{k}$.*

- BEWEIS. (i) Induktionsanfang: Die Anzahl der k -elementigen Teilmengen einer 0-elementigen, d.h. leeren Menge beträgt 1, falls $k = 0$, und 0 sonst.
- (ii) Induktionsschritt: Sei eine $(n+1)$ -elementige Menge M gegeben. In ihr sei ein Element x gewählt. Nach der Induktionsvoraussetzung ist die Anzahl aller k -elementigen Teilmengen von M , die x nicht enthalten, gleich $\binom{n}{k}$, die Anzahl aller k -elementigen Teilmengen, die x enthalten, gleich $\binom{n}{k-1}$. Insgesamt ist die Anzahl aller k -elementigen Teilmengen von M somit $\binom{n}{k} + \binom{n}{k-1} = \binom{n+1}{k}$.

□

1.5. Betrag, Ungleichungen, Intervalle. Sind a und b reelle Zahlen, so gilt stets genau einer der folgenden Fälle:

- $a < b$, d.h. a ist kleiner als b ,
- $a > b$, d.h. a ist größer als b
- $a = b$, d.h. a und b sind gleich.

Man kann sich die reellen Zahlen auf einer Zahlengerade angeordnet vorstellen. Um reelle Zahlen zu vergleichen, benutzen wir die Symbole

$<$ “kleiner als”,
 $>$ “größer als”,
 $\leq$ “kleiner oder gleich”,
 $\geq$ “größer oder gleich”,
 $=$ “gleich”,
 $\neq$ “ungleich”.

Beim *Rechnen mit Ungleichungen* ist daran zu denken, dass sich das Ungleichheitszeichen umdreht, wenn mit einer negativen Zahl multipliziert wird, d.h. es gilt

$$a < b \iff \begin{cases} ac < bc, & \text{falls } c > 0; \\ ac > bc, & \text{falls } c < 0. \end{cases}$$

(Hier haben wir das Zeichen $\iff$ verwendet, um die *Äquivalenz* zweier Ungleichungen auszudrücken. Äquivalenz von Gleichungen oder Ungleichungen bedeutet, dass sie dieselbe Lösungsmenge haben.) Gegebenenfalls muss beim Rechnen eine Fallunterscheidung vorgenommen werden.

Intervalle sind durch Ungleichungen definierte zusammenhängende Teilmengen von $\mathbb{R}$.

$$\begin{aligned}
 [a, b] &:= \{x \in \mathbb{R} \mid a \leq x \leq b\}, \\
 (a, b] &:= \{x \in \mathbb{R} \mid a < x \leq b\}, \\
 [a, b) &:= \{x \in \mathbb{R} \mid a \leq x < b\}, \\
 (a, b) &:= \{x \in \mathbb{R} \mid a < x < b\}.
 \end{aligned}$$

Auch ∞ und $-\infty$ können als *Intervallgrenzen* vorkommen:

$$[a, \infty) := \{x \in \mathbb{R} \mid a \leq x\} \text{ etc.}$$

Der *Betrag* einer reellen Zahl ist definiert durch

$$|x| := \begin{cases} x, & \text{falls } x \geq 0; \\ -x, & \text{falls } x < 0. \end{cases}$$

Es gelten die Rechenregeln

$$|a| \cdot |b| = |ab|, \quad \left| \frac{c}{d} \right| = \frac{|c|}{|d|}, \text{ für } d \neq 0,$$

und die *Dreiecksungleichung*

$$|a + b| \leq |a| + |b|.$$

Beim Rechnen mit Beträgen sind oft *Fallunterscheidungen* nötig:

Beispiel 1.5.1. Gesucht sind alle Lösungen der Ungleichung

$$|3x| \leq |x + 1| - x.$$

“Kritische” Stellen, sind $x = -1$ und $x = 0$:

- Fall a) Angenommen, es gilt $x \geq 0$. Dann ist die Ungleichung äquivalent zu $3x \leq x + 1 - x \iff x \leq \frac{1}{3}$. Es ergibt sich die Menge von Lösungen $\mathbb{L}_a = [0, \frac{1}{3}]$.
- Fall b) Angenommen, $-1 \leq x < 0$. In diesem Fall ist die Ungleichung äquivalent zu $-3x \leq x + 1 - x \iff x \geq -\frac{1}{3}$. Es ergibt sich die Menge von Lösungen $\mathbb{L}_b = [-\frac{1}{3}, 0)$.
- Fall c) Sei nun $x < -1$. Dann ist die obige Ungleichung zu $-3x \leq -x - 1 - x \iff x \geq 1$ äquivalent. Da wir aber $x < -1$ angenommen haben, erhalten wir also keine Lösungen in diesem Fall, $\mathbb{L}_c = \emptyset$.

Insgesamt erhalten wir die Lösungsmenge der Ungleichung als

$$\mathbb{L} = \mathbb{L}_a \cup \mathbb{L}_b \cup \mathbb{L}_c = [0, \frac{1}{3}] \cup [-\frac{1}{3}, 0) \cup \emptyset = [-\frac{1}{3}, \frac{1}{3}].$$

1.6. Komplexe Zahlen. Die *komplexen Zahlen* sind ein Körper, den man erhält, wenn man auf $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$, also auf der Menge von Paaren reeller Zahlen eine Addition und Multiplikation wie folgt einführt: Die Addition ist komponentenweise definiert,

$$(a, b) + (c, d) := (a + c, b + d).$$

Die Multiplikation ist so definiert:

$$(a, b) \cdot (c, d) := (ac - bd, ad + bc).$$

Man kann nachprüfen, dass es sich um einen Körper handelt. In der Tat, $(0, 0)$ und $(1, 0)$ sind neutrale Elemente bezüglich der Addition bzw. Multiplikation, durch

$$-(a, b) := (-a, -b)$$

ist das Inverse eines Elements (a, b) bezüglich der Addition gegeben, durch

$$(a, b)^{-1} := \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right)$$

das Inverse eines Elements $(a, b) \neq (0, 0)$ bezüglich der Multiplikation. Distributiv- und Assoziativgesetze lassen sich durch eine Rechnung nachprüfen und dass beide Kommutativgesetze erfüllt sind, sieht man den Definitionen der Verknüpfungen direkt an. Wir schreiben $\mathbb{C}$ für den Körper der komplexen Zahlen.

Beschränkt man sich auf komplexe Zahlen der Form $(a, 0)$, erhält man die reellen Zahlen zurück, denn es gilt $(a, 0) + (c, 0) = (a + c, 0)$ und $(a, 0)(c, 0) = (ac, 0)$. Daher wird die Zahl a auch der *Realteil* der komplexen Zahl (a, b) genannt, in Zeichen $\operatorname{Re}(a, b) = a$. Die Zahl b nennt man die *Imaginärteil* von (a, b) , in Zeichen $\operatorname{Im}(a, b) = b$. Gilt für (a, b) , dass $b = 0$ ist, nennt man die Zahl (a, b) *reell*, gilt $a = 0$, nennt man sie *rein imaginär*.

Übersichtlicher und weiter verbreitet als die bis jetzt verwendete Paarschreibweise für komplexe Zahlen ist es, die *imaginäre Einheit* $i := (0, 1)$ einzuführen und komplexe Zahlen der Form $(x, 0)$ als reelle Zahlen aufzufassen. Man schreibt dann

$$a + ib = (a, b)$$

mit reellen Zahlen a, b . Unter Beachtung von Kommutativ-, Assoziativ- und Distributivgesetzen rechnet man auf diese Weise mit komplexen Zahlen genauso wie mit reellen Zahlen, wobei man beim Multiplizieren zweier rein imaginärer Zahlen lediglich $i^2 = i \cdot i = -1$ beachten muss. (Man schreibt deshalb auch manchmal $\sqrt{-1}$ anstatt i für die imaginäre Einheit). Zum Beispiel:

$$(3 - 5i)(-1 + 2i) = -3 + 6i + 5i - 10i^2 = -3 + 11i + 10 = 7 + 11i.$$

Sei $z = (a, b)$ eine komplexe Zahl. Dann ist durch $\bar{z} = (a, -b)$ die *komplex-konjugierte Zahl* zu z gegeben. Die komplexe Konjugation bewirkt also die Spiegelung an der x -Achse. Komplexe Konjugation ist mit der Addition und Multiplikation verträglich, d.h. es gilt $\overline{z + w} = \bar{z} + \bar{w}$ und $\overline{zw} = \bar{z}\bar{w}$. Den Real- und Imaginärteil einer komplexen Zahl kann man mit Hilfe der komplexen Konjugation auch durch

$$\operatorname{Re}(z) = \frac{1}{2}(z + \bar{z}), \quad \operatorname{Im}(z) = \frac{i}{2}(\bar{z} - z)$$

ausdrücken. Insbesondere ist eine komplexe Zahl z genau dann reell, wenn $z = \bar{z}$ gilt. Mit der komplexen Konjugation lässt sich das multiplikative Inverse von $z \neq 0$

schreiben als

$$z^{-1} = \frac{\bar{z}}{z\bar{z}}.$$

Durch

$$|z| := \sqrt{z\bar{z}} = \sqrt{(a+bi)(a-bi)} = \sqrt{a^2+b^2}$$

ist der *Betrag* der komplexen Zahl $z = (a, b)$ definiert.

Ist der $\mathbb{R}^2$ mit den Verknüpfungen der komplexen Zahlen ausgestattet, nennt man ihn auch die *Gaußsche Zahlenebene*. Man kann die Verknüpfungen der komplexen Zahlen dann als Operationen in der Ebene interpretieren. Die Addition der komplexen Zahlen ist die übliche Vektoraddition.

Die Multiplikation kann man besser verstehen, wenn *Polarkoordinaten* eingeführt werden. Dazu schreibt man jede komplexe Zahl in der Form

$$r \cdot (\cos(\varphi) + i \sin(\varphi)),$$

wobei $r \geq 0$ und φ reelle Zahlen sind. Für von Null verschiedene komplexe Zahlen kann man diese Darstellung eindeutig machen, indem man fordert, dass $\varphi \in [0, 2\pi)$ gilt. Die Zahl r ist der Betrag der komplexen Zahl und φ ist der Winkel, den der Vektor in der Gaußschen Zahlenebene mit der x -Achse einschließt. Bei der Multiplikation zweier komplexer Zahlen multiplizieren sich nun die Beträge, die Winkel addieren sich, es gilt also

$$r(\cos(\varphi) + i \sin(\varphi)) \cdot s(\cos(\vartheta) + i \sin(\vartheta)) = rs(\cos(\varphi + \vartheta) + i \sin(\varphi + \vartheta)).$$

Insbesondere bewirkt die Multiplikation mit einer Zahl vom Betrag 1 eine Drehung um den Nullpunkt. Der obige Identität folgt aus den Additionstheoremen für den Sinus und den Kosinus. Umgekehrt kann man sich die Additionstheoreme daraus herleiten, indem man nämlich die linke Seite nach den Rechenregeln für komplexe Zahlen multipliziert und $r = s = 1$ setzt:

$$\begin{aligned}\cos(\varphi + \vartheta) &= \cos(\varphi) \cos(\vartheta) - \sin(\varphi) \sin(\vartheta), \\ \sin(\varphi + \vartheta) &= \cos(\varphi) \sin(\vartheta) + \sin(\varphi) \cos(\vartheta).\end{aligned}$$

Das Additionstheorem für den Kosinus ergibt sich dann aus dem Realteil, das für den Sinus aus dem Imaginärteil des Produkts.

2. Vektoren

2.1. Der $\mathbb{R}^n$. Die mit kartesischen Koordinaten versehene Ebene bezeichnen wir mit $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$. Als Menge besteht der $\mathbb{R}^2$ aus allen geordneten Paaren von reellen Zahlen,

$$\mathbb{R}^2 = \{(x, y) \mid x, y \in \mathbb{R}\}.$$

Entsprechend bezeichnen wir mit $\mathbb{R}^3$ die Menge aller geordneten reellen Zahlentripel. Allgemeiner verwenden wir folgende Bezeichnung:

$$\mathbb{R}^n = \{(x_1, \dots, x_n) \mid x_1, \dots, x_n \in \mathbb{R}\},$$

d.h. wir bezeichnen mit $\mathbb{R}^n$ die Menge aller geordneten n -Tupel von reellen Zahlen. Das Wort *geordnet* betont hierbei, dass es auf die Reihenfolge ankommt, beispielsweise sind $(1, 0, 0, 0)$ und $(0, 1, 0, 0)$ zwei verschiedene 4-Tupel von reellen Zahlen. Die Elemente von $\mathbb{R}^n$ bezeichnen wir auch als *Vektoren*.

Vektoren im $\mathbb{R}^n$ kann man addieren und subtrahieren, dies geschieht *komponentenweise*, d.h. man definiert

$$(x_1, \dots, x_n) + (y_1, \dots, y_n) := (x_1 + y_1, \dots, x_n + y_n),$$

und

$$(x_1, \dots, x_n) - (y_1, \dots, y_n) := (x_1 - y_1, \dots, x_n - y_n).$$

Außerdem kann man Vektoren mit einem *Skalar*, d.h. mit einer reellen Zahl, multiplizieren, indem man jede Komponente des Vektors mit der Zahl multipliziert:

$$\lambda(x_1, \dots, x_n) := (\lambda x_1, \dots, \lambda x_n),$$

wobei λ eine reelle Zahl bezeichnet.

2.2. Skalarprodukt und Abstandsfunktion. Seien $x = (x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ zwei Vektoren im $\mathbb{R}^n$.

Die *Norm* oder *Länge* eines Vektors im $\mathbb{R}^n$ ist gegeben durch

$$\|x\| := \sqrt{x_1^2 + \dots + x_n^2}.$$

Im $\mathbb{R}^2$ ist die Norm des Vektors (a, b) nach Pythagoras gerade die Länge der Hypotenuse eines rechtwinkligen Dreiecks mit den Kathetenlängen a und b . Im $\mathbb{R}^3$ ist die Norm von (x, y, z) genau die Länge der Raumdiagonalen eines Quaders mit

den Kantenlängen x , y und z . Man sagt, x ist ein *Einheitsvektor*, wenn $\|x\| = 1$ gilt. Für die Norm gilt die *Dreiecksungleichung*, nämlich

$$\|x + y\| \leq \|x\| + \|y\| \quad \text{für alle } x, y \in \mathbb{R}^n,$$

die wir ohne Beweis verwenden werden.

Das *Skalarprodukt* von x und y ist wie folgt definiert

$$x \cdot y := x_1 y_1 + \cdots + x_n y_n.$$

Das Ergebnis des *Skalarprodukts* ist also eine reelle Zahl (nicht zu verwechseln mit der skalaren Multiplikation im letzten Abschnitt, wo das Ergebnis ein Vektor ist.) Zwei Vektoren x und y stehen aufeinander *senkrecht*, man sagt auch: sind *orthogonal* zueinander, wenn $x \cdot y = 0$ gilt. Sind x und y zwei Einheitsvektoren, dann gibt $x \cdot y$ den Kosinus des Winkels zwischen den beiden Vektoren an. Der Zusammenhang zwischen Norm und Skalarprodukt ist durch $\|x\| = \sqrt{x \cdot x}$ gegeben.

Der *Abstand* zwischen zwei Punkten x und y im $\mathbb{R}^n$ ist durch

$$d(x, y) := \|y - x\|,$$

also durch die Länge des Differenzvektors, definiert. Aus der Dreiecksungleichung für die Norm folgt die *Dreiecksungleichung* für den Abstand:

$$d(x, y) + d(y, z) \geq d(x, z).$$

Anschaulich gesprochen: Geht man in gerader Linie vom Punkt x zum Punkt z , so ist der Weg höchstens so lang wie der auf einem “Umweg” über einen dritten Punkt y .

Wenn man die genaue Lage eines Punktes im $\mathbb{R}^n$ nicht kennt, sondern nur weiß, dass er keinen größeren Abstand als vorgegebene Zahl ε von einem festen Referenzpunkt x hat, dann ist folgende Definition nützlich. Der (*offene*) ε -Ball um den Punkt x ist definiert durch

$$B_\varepsilon(x) := \{y \in \mathbb{R}^n \mid \|y - x\| < \varepsilon\}.$$

Er wird auch manchmal als ε -Umgebung von x bezeichnet. Mit dieser Schreibweise bedeutet $y \in B_\varepsilon(x)$, dass y nicht weiter als ε von x entfernt liegt. Im $\mathbb{R}^2$ handelt es sich um die Kreisscheibe ohne Rand vom Radius ε mit Mittelpunkt x , im $\mathbb{R}^3$ um die Kugel Radius ε mit Mittelpunkt x , ohne die Punkte auf der Kugeloberfläche. Im Eindimensionalen, also für eine reelle Zahl x , ist der ε -Ball um x das Intervall $(x - \varepsilon, x + \varepsilon)$.

3. Abbildungen

Eine *Abbildung* zwischen zwei Mengen M und N ist eine Zuordnungsvorschrift, die jedem Element von M ein Element von N zuweist. Man kann zum Ausdruck bringen, dass f so eine Abbildung ist, indem man

$$f: M \rightarrow N$$

schreibt. Dann nennt man M den *Definitionsbereich* und N die *Zielmenge* von f . Durch die Schreibweisen

$$f(m) = n \quad \text{oder} \quad m \mapsto n$$

wird erklärt, dass das Element m auf das Element n *abgebildet* wird. Man sagt in diesem Fall auch, n ist das *Bild* von m , oder n ist der *Wert* von f bei m .

Es ist auch üblich, eine Abbildungen auch als *Funktion* zu bezeichnen, dies tut man vor allem dann, wenn die Zielmenge aus Zahlen besteht, also etwa eine Teilmenge von $\mathbb{R}$ ist.

Ein Beispiel für eine solche Funktion ist

$$f: \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) = x^2,$$

die jede reelle Zahl auf ihr Quadrat abbildet. Man beachte, dass bei einer Abbildung verschiedene Element des Definitionsbereichs dasselbe Bild haben können, in unserem Beispiel gilt etwa $f(2) = 4$ und auch $f(-2) = 4$. Außerdem müssen die Elemente der Zielmenge nicht notwendigerweise auch alle als Bild irgendeines Elements im Definitionsbereich auftreten, etwa in unserem Beispiel wird kein Element von $\mathbb{R}$ durch f auf eine negative Zahl abgebildet.

Die Menge aller Elemente der Zielmenge, die als Bilder auftreten, heißt das *Bild* oder die *Wertmenge* $f(M)$ von f . Die Menge aller Elemente der Zielmenge, die als Bilder einer Teilmenge $A \subseteq M$ auftreten,

$$f(A) := \{n \in N \mid \text{es gibt ein } m \in A, \text{ so dass } f(m) = n\}$$

nennt man allgemeiner das *Bild* von A . Für eine Teilmenge $B \subseteq N$ der Zielmenge nennt man

$$f^{-1}(B) := \{m \in M \mid f(m) \in B\}$$

das *Urbild* von B . Alle Elemente die auf $n \in N$ abgebildet werden, heißen *Urbilder* von n .

Hat man zwei Abbildungen $f: M \rightarrow N$ und $g: N \rightarrow P$, so bezeichnet man mit $g \circ f$ die *Verkettung*, d.h. die Hintereinanderausführung von g nach f , in Zeichen:

$$g \circ f: M \rightarrow P, \quad g \circ f(m) := g(f(m)).$$

Man beachte, dass bei der Hintereinanderausführung von Abbildungen es auf die Reihenfolge ankommt, seien etwa die beiden Funktionen $f(x) = x^2$ und $g(x) = x+1$ von $\mathbb{R}$ nach $\mathbb{R}$ gegeben, dann gilt $f \circ g(x) = x^2 + 2x + 1$, aber $g \circ f(x) = x^2 + 1$.

Auf jeder Menge M ist die sogenannte *identische Abbildung*

$$\text{id}_M: M \rightarrow M, x \mapsto x$$

gegeben, die jedes Element von M auf sich selbst abbildet. Man nennt eine Abbildung *konstant*, wenn Sie alle Elemente des Definitionsbereichs auf dasselbe Element in der Zielmenge abbildet.

3.1. Injektivität und Surjektivität. Für eine gegebene Abbildung $f: M \rightarrow N$ kann jedes Element n des Zielbereichs N eines, mehrere oder gar kein Urbild haben. Bei der oben betrachteten Funktion $f: \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ etwa haben positive Zahlen zwei Urbilder, die Null genau ein Urbild und negative Zahlen gar keines.

Eine Abbildung $f: M \rightarrow N$ heißt *injektiv*, wenn jedes Element von N höchstens ein Urbild hat, oder mit anderen Worten, falls es keine zwei verschiedenen Elemente in M gibt, die auf dasselbe Element von N abgebildet werden.

Eine Abbildung $f: M \rightarrow N$ heißt *surjektiv*, wenn jedes Element von N mindestens ein Urbild hat, oder mit anderen Worten, falls $f(M) = N$ gilt.

Hat eine Abbildung beide Eigenschaften, wird sie *bijektiv* oder *umkehrbar* genannt. In diesem Fall hat jedes Element von N genau ein Urbild, dies ist gleichbedeutend damit, dass es eine *Umkehrabbildung* gibt, d.h. eine Abbildung $g: N \rightarrow M$, so dass $g(f(m)) = m$ und $f(g(n)) = n$ gilt.

Eine reelle Funktion heißt *monoton steigend*, falls aus $a \leq b$ stets $f(a) \leq f(b)$ folgt, *streng monoton steigend*, falls aus $a < b$ stets $f(a) < f(b)$ folgt. Eine streng monoton steigende Funktion ist insbesondere injektiv.

3.2. Lineare Abbildungen. Eine Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ heißt *linear*, falls

$$f(x + y) = f(x) + f(y) \quad \text{und} \quad f(\lambda x) = \lambda f(x)$$

für alle Vektoren x und y aus dem $\mathbb{R}^n$ und alle reellen Zahlen λ gilt.

Zum Beispiel ist die Abbildung

$$f: \mathbb{R}^2 \rightarrow \mathbb{R}, \quad f(x_1, x_2) = 2x_1 + 3x_2$$

linear. Denn es gilt $f(x_1 + y_1, x_2 + y_2) = 2(x_1 + y_1) + 3(x_2 + y_2) = 2x_1 + 3x_2 + 2y_1 + 3y_2 = f(x_1, x_2) + f(y_1, y_2)$ sowie $f(\lambda x_1, \lambda x_2) = 2\lambda x_1 + 3\lambda x_2 = \lambda f(x_1, x_2)$.

Andererseits ist die Abbildung

$$f: \mathbb{R}^2 \rightarrow \mathbb{R}, \quad f(x_1, x_2) = x_1^2$$

nicht linear. Denn es gilt $f((1, 0) + (2, 0)) = f(3, 0) = 9$, was nicht gleich $f(1, 0) + f(2, 0) = 1 + 4 = 5$ ist.

Lineare Abbildungen sind besonders einfache Abbildungen: Seien $e_1, \dots, e_n$, die sogenannten *Standardbasisvektoren* des $\mathbb{R}^n$, d.h. der Vektor e_k hat an der k -ten Stelle eine Eins, an den anderen Stellen Nullen:

$$e_1 = (1, 0, \dots, 0), \quad e_2 = (0, 1, 0, \dots, 0), \quad \dots, \quad e_n = (0, \dots, 0, 1).$$

Damit können wir jeden Vektor x als $(x_1, \dots, x_n) = x_1 e_1 + \dots + x_n e_n$ schreiben. Nun ist der Wert $f(x)$ der linearen Abbildung f bereits durch die Vektoren $f(e_1), \dots, f(e_n)$ und die Zahlen $x_1, \dots, x_n$ gegeben, denn man erhält, wenn man die Linearitätseigenschaften mehrfach anwendet: $f(x) = f(x_1 e_1 + \dots + x_n e_n) = x_1 f(e_1) + \dots + x_n f(e_n)$. Dies zeigt: *Eine lineare Abbildung ist durch ihre Werte auf den Standardbasisvektoren bereits eindeutig festgelegt.* Daher kann man die gesamte Information über eine lineare Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^m$ in einer $m \times n$ -Matrix, d.h. einem rechteckigen Zahlenschema mit m Zeilen und n Spalten, mit reellen Einträgen zusammenfassen:

Definition 3.2.1. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine lineare Abbildung. Dann heißt die Matrix

$$A = (f(e_1), \dots, f(e_n)),$$

deren Spalten durch die Werte der Abbildung f auf den Basisvektoren $e_1, \dots, e_n$ gegeben sind, die *darstellende Matrix* der linearen Abbildung f .

Ist

$$A = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix}$$

die darstellende Matrix einer linearen Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$, so erhält man den Wert $f(x)$ bei einem beliebigen Vektor $x \in \mathbb{R}^n$ als

$$Ax := (a_{11}x_1 + \dots + a_{1n}x_n, \dots, a_{m1}x_1 + \dots + a_{mn}x_n),$$

wodurch wir auch die *Multiplikation* oder *Anwendung* einer Matrix mit n Spalten auf einen Vektor mit n Einträgen definiert haben.

Umgekehrt erhält man aus jeder $m \times n$ -Matrix A mit reellen Einträgen eine lineare Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$, indem man $f(x) := Ax$ setzt. (Die Linearität dieser Abbildung lässt sich leicht nachrechnen). Insgesamt zeigt dies, dass lineare Abbildungen $\mathbb{R}^n \rightarrow \mathbb{R}^m$ und reelle $m \times n$ -Matrizen im wesentlichen dasselbe sind.

Beispiele 3.2.2. Zwei Beispiel für eine darstellende Matrix:

- (i) Sei die lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ durch $f(x_1, x_2, x_3) = (x_1 + 2x_3, 4x_2 - x_3)$ definiert. Dann ist

$$A = \begin{pmatrix} 1 & 0 & 2 \\ 0 & 4 & -1 \end{pmatrix}$$

die darstellende Matrix von f .

- (ii) Die identische Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$ ist eine lineare Abbildung. Ihre darstellende Matrix ist die $n \times n$ -Einheitsmatrix.

$$E := \begin{pmatrix} 1 & & \\ & \ddots & \\ & & 1 \end{pmatrix}$$

Dies ist die Matrix, deren Spalten von den Standardbasisvektoren des $\mathbb{R}^n$ gebildet wird. (Zur Erläuterung der Schreibweise dieser Matrix sei darauf hingewiesen, dass es eine allgemein üblich Konvention ist, fehlende Einträge in einer Matrix als Null anzunehmen, so dass man sich hier darauf beschränken kann, die Einsen auf der Diagonale hinzuschreiben, die Nullen in allen anderen Einträgen dagegen weglassen kann.)

3.3. Matrizen. Eine $m \times n$ -Matrix ist ein rechteckiges Zahlenschema mit m Zeilen und n Spalten wie im vorherigen Abschnitt eingeführt. Im folgenden werden wir reelle Matrizen, d.h. Matrizen mit reellen Einträgen betrachten.

Wir wollen einige Rechenregeln für Matrizen kennen lernen. Beim Rechnen mit Matrizen gibt es einige allgemein akzeptierte Übereinkünfte; dazu gehört, dass man Vektoren im $\mathbb{R}^n$ als *Spaltenvektoren* auffasst, d.h. man schreibt die Einträge von x untereinander:

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix},$$

oder, mit anderen Worten, man fasst Vektoren im $\mathbb{R}^n$ als $n \times 1$ -Matrizen auf. Diese Unterscheidung ist nur dann notwendig, wenn mit Matrizen gerechnet wird und wir werden sie sonst ignorieren, d.h. wo es keine Rolle spielt, werden wir Vektoren aus dem $\mathbb{R}^n$ weiterhin platzsparend als *Zeilenvektoren* $x = (x_1, \dots, x_n)$ schreiben.

Zwei $m \times n$ -Matrizen kann man addieren oder subtrahieren, indem man *komponentenweise* addiert oder subtrahiert. Gilt

$$A := \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \quad \text{und} \quad B := \begin{pmatrix} b_{11} & \dots & b_{1n} \\ \vdots & & \vdots \\ b_{m1} & \dots & b_{mn} \end{pmatrix},$$

dann ist die Summe der beiden Matrizen gegeben durch

$$A + B := \begin{pmatrix} a_{11} + b_{11} & \dots & a_{1n} + b_{1n} \\ \vdots & & \vdots \\ a_{m1} + b_{m1} & \dots & a_{mn} + b_{mn} \end{pmatrix}$$

und $A - B$ erhält man entsprechend, indem man in der obigen Definition die Pluszeichen durch Minuszeichen ersetzt.

Die *Matrixmultiplikation* geschieht nicht komponentenweise, sondern ist wie folgt definiert. Sei A eine $m \times n$ -Matrix und B eine $n \times p$ -Matrix,

$$A := \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \quad \text{und} \quad B := \begin{pmatrix} b_{11} & \dots & b_{1p} \\ \vdots & & \vdots \\ b_{n1} & \dots & b_{np} \end{pmatrix},$$

dann ist das Matrizenprodukt AB gegeben durch

$$AB := \begin{pmatrix} a_{11}b_{11} + \cdots + a_{1n}b_{n1} & \cdots & a_{11}b_{1p} + \cdots + a_{1n}b_{np} \\ \vdots & & \vdots \\ a_{m1}b_{11} + \cdots + a_{mn}b_{n1} & \cdots & a_{m1}b_{1p} + \cdots + a_{mn}b_{np} \end{pmatrix}.$$

Dies kann man auch so ausdrücken: Der Eintrag von AB in der i -ten Zeile und j -ten Spalte ist das Skalarprodukt der i -ten Zeile von A mit der j -ten Spalte von B . Kurze Merkregel für die Matrizenmultiplikation: “Zeile mal Spalte”.

Dazu ein Beispiel: Seien

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \quad \text{und} \quad B = \begin{pmatrix} 1 & 2 \\ 0 & -1 \\ -1 & 1 \end{pmatrix}.$$

Das Matrizenprodukt AB berechnet sich dann als

$$AB = \begin{pmatrix} 1 \cdot 1 + 2 \cdot 0 - 3 \cdot 1 & 1 \cdot 2 - 2 \cdot 1 + 3 \cdot 1 \\ 4 \cdot 1 + 5 \cdot 0 - 6 \cdot 1 & 4 \cdot 2 - 5 \cdot 1 + 6 \cdot 1 \end{pmatrix} = \begin{pmatrix} -2 & 3 \\ -2 & 9 \end{pmatrix}.$$

Man beachte, dass das Matrizenprodukt AB nur dann definiert ist, wenn die Spaltenanzahl von A und die Zeilenanzahl von B übereinstimmen. Das Matrizenprodukt ist außerdem nicht kommutativ, d.h. auch wenn beide Produkte AB und BA definiert sind, sind sie im allgemeinen nicht gleich. Allerdings ist die Multiplikation von Matrizen assoziativ (ohne Beweis): Ist das Produkt $(AB)C$ von drei Matrizen definiert, dann gilt $(AB)C = A(BC)$ und wir schreiben dafür ABC .

Der Sinn des Matrizenprodukts ist, dass es der Verkettung von linearen Abbildungen entspricht.

Satz 3.3.1. Seien $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ und $g: \mathbb{R}^p \rightarrow \mathbb{R}^n$ zwei lineare Abbildungen so dass A die darstellende Matrix von f ist und B die darstellende Matrix von g ist. Dann ist $f \circ g: \mathbb{R}^p \rightarrow \mathbb{R}^m$ eine lineare Abbildung, deren darstellende Matrix durch das Matrixprodukt AB gegeben ist.

BEWEIS. Die Linearität von $f \circ g$ folgt sofort aus $f(g(v+w)) = f(g(v)) + f(g(w)) = f(g(v)) + f(g(w))$ und $f(g(\lambda v)) = f(\lambda g(v)) = \lambda f(g(v))$. Die zweite Aussage folgt sofort, wenn man die Assoziativität der Matrixmultiplikation (die wir hier nicht bewiesen haben) verwendet: Es gilt $f(g(e_i)) = f(B(e_i)) = A(Be_i) = (AB)e_i$, d.h. die i -te Spalte von AB ist das Bild von e_i unter $f \circ g$. $\square$

Wir sehen nun, dass die vorher definierte Anwendung einer $m \times n$ -Matrix auf einen Vektor im $\mathbb{R}^n$ nur ein Spezialfall der Matrixmultiplikation ist, wenn man den Spaltenvektor im $\mathbb{R}^n$ als $n \times 1$ -Matrix auffasst:

$$Ax = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + \dots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n \end{pmatrix}.$$

Zum Abschluss dieses Unterabschnitts wollen wir noch einige Rechenregeln für das Rechnen mit Matrizen (ohne Beweise) angeben. Bereits oben erwähnt wurde die Assoziativität der Matrixmultiplikation, es gelten auch folgende Distributivgesetze:

$$A(B + C) = AB + AC, \quad (A + B)C = AC + BC.$$

Bisweilen nützlich ist die folgende Notation. Sei A eine $m \times n$ -Matrix. Dann ist die zu A *transponierte Matrix* A^t definiert als die $n \times m$ -Matrix, deren Eintrag in der i -ten Zeile und j -ten Spalte gerade der Eintrag a_{ji} aus der j -ten Zeile und i -ten Spalte von A ist, zum Beispiel:

$$\begin{pmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{pmatrix}^t = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}$$

Für transponierte Matrizen gelten die Regeln

$$(A^t)^t = A, \quad (AB)^t = B^t A^t.$$

3.4. Polynome. Seien $a_0, \dots, a_k$ reelle oder komplexe Zahlen. Ein Ausdruck von der Form

$$a_k x^k + a_{k-1} x^{k-1} + \dots + a_2 x^2 + a_1 x + a_0$$

heißt *Polynom* in der Variablen x . Die Zahlen $a_0, \dots, a_k$ heißen die *Koeffizienten* des Polynoms. Oft wird auch die durch ein solches Polynom gegebene *Polynomfunktion*

$$P(x) = \sum_{j=0}^k a_j x^j$$

einfach als Polynom bezeichnet. Wenn wir $a_k \neq 0$ annehmen, dann wird a_k als der *Leitkoeffizient* des Polynoms bezeichnet, allgemein ist es derjenige unter den von Null verschiedenen Koeffizienten, der vor der höchsten Potenz von x steht. Ist a_k der Leitkoeffizient, dann sagt man, das Polynom hat den *Grad* k . Ist kein einziger Koeffizient von Null verschieden, dann handelt es sich um das *Nullpolynom*,

dessen Grad man per Definition gleich $-\infty$ setzt. Die Zahl a_0 wird auch als das *Absolutglied* des Polynoms bezeichnet.

Lässt man für die Koeffizienten und die Variable ausschließlich reelle Zahlen zu, so spricht man von einem *reellen Polynom*, lässt man komplexe Zahlen zu, dann handelt es sich um ein *komplexes Polynom*. Für Polynome niedrigen Grades gibt es die folgenden Bezeichnungen

- Polynome vom Grad 0 nennt man *konstante Polynome*.
- Polynome vom Grad 1, also Ausdrücke der Form $ax + b$ werden auch *lineare Polynome* genannt.
- Polynome vom Grad 2, also Ausdrücke der Form $ax^2 + bx + c$ nennt man *quadratische Polynome*.
- Polynome vom Grad 3 heißen *kubische Polynome*.

Entsprechend nennt man bei einem Polynom die Summanden a_0 auch *konstantes Glied*, a_1x *lineares Glied*, a_2x^2 *quadratisches Glied* und a_3x^3 *kubisches Glied*.

Eine Zahl x heißt *Nullstelle* eines Polynoms P , wenn $P(x) = 0$ gilt. Die Nullstellen eines reellen quadratischen Polynoms $ax^2 + bx + c$ erhält man, falls die *Diskriminante* $b^2 - 4ac$ nicht negativ ist, aus der wohlbekannten Formel

$$x_{1,2} = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a},$$

andernfalls hat das Polynome keine reellen Nullstellen.

Die Formel funktioniert auch für die Nullstellen eines komplexen quadratischen Polynoms (das nach dem unten folgenden Satz 3.4.2 immer mindestens eine Nullstelle hat). Zum Berechnen interpretiert man die Quadratwurzel einer komplexen Zahl $r(\cos(\varphi) + i\sin(\varphi))$ mit $r \geq 0$ und $\varphi \in [0, 2\pi)$ als die komplexe Zahl $\sqrt{r}(\cos(\frac{\varphi}{2}) + i\sin(\frac{\varphi}{2}))$. Ein komplexes quadratisches Polynom hat entweder eine oder zwei Nullstellen, je nachdem, ob die Diskriminante gleich Null oder von Null verschieden ist.

Beispiel 3.4.1. Das reelle quadratische Polynom

$$P(x) = x^2 + x + 1$$

hat keine reellen Nullstellen, denn die Diskriminante $b^2 - 4ac = -3$ ist negativ. Man kann es aber auch als komplexes Polynom auffassen, dann hat es die

Nullstellen

$$x_{1,2} = \frac{-1 \pm i\sqrt{3}}{2} = -\frac{1}{2} \pm i\frac{\sqrt{3}}{2}.$$

Um dies einzusehen, muss daran denken, dass die negative reelle Zahl -3 in der Gaußschen Zahlenebene den Winkel π mit der reellen Achse einschließt, die Quadratwurzel ist somit die Zahl $i\sqrt{3}$, also die komplexe Zahl vom Betrag $\sqrt{3}$, die mit der x -Achse den Winkel $\frac{\pi}{2}$ einschließt.

Satz 3.4.2 (Fundamentalsatz der Algebra). *Jedes nicht konstante komplexe Polynom besitzt mindestens eine Nullstelle.*

Bemerkung 3.4.3. Den obigen Satz werden wir ohne Beweis verwenden. Wir wollen einige wichtige Folgerungen daraus erwähnen. Durch *Polynomdivision* erhält man zu zwei Polynomen $P(x)$ und $Q(x)$ zwei weitere Polynome $S(x)$ und $R(x)$, so dass

$$P(x) = S(x) \cdot Q(x) + R(x)$$

gilt, wobei der Grad von $R(x)$ dabei echt kleiner als der von $Q(x)$ ist. Das Polynom $R(x)$ stellt dabei den bei der Division eventuell auftretenden Rest dar.

Ist nun λ eine Nullstelle des Polynoms $P(x)$, dann muss die Polynomdivision von $P(x)$ durch $Q(x) = x - \lambda$ ohne Rest aufgehen, denn der Grad von $Q(x)$ ist hier eins, also ist der Grad von $R(x)$ gleich Null oder gleich $-\infty$. Aber aus $0 = P(\lambda) = S(\lambda) \cdot (\lambda - \lambda) + R(\lambda) = R(\lambda)$ folgt, dass $R(x)$ das Nullpolynom ist.

Aus dem Fundamentalsatz der Algebra 3.4.2 folgt nun, dass man bei einem komplexen Polynom vom Grad $k \geq 1$ sukzessive k solche *Linearfaktoren* von der Form $x - \lambda_j$ herausdividieren kann (bis man ein konstantes Polynom erhält). Dies beweist, dass man jedes komplexe Polynom vom Grad k in der Form

$$P(x) = c \cdot (x - \lambda_1) \cdot (x - \lambda_2) \cdot \dots \cdot (x - \lambda_k)$$

schreiben kann, wobei $\lambda_1, \dots, \lambda_k$ (nicht notwendig verschiedene) Nullstellen von $P(x)$ sind und c eine komplexe Konstante ist. Wir erhalten daraus als Folgerung:

Satz 3.4.4. *Jedes komplexe oder reelle Polynom vom Grad $k \geq 0$ hat höchstens k verschiedene Nullstellen.*

Bis jetzt haben wir Polynome in einer Variablen betrachtet, man kann aber auch Polynome in mehreren Variablen, die manchmal auch als *multivariate Polynome* bezeichnet werden, betrachten. Für $n \in \mathbb{N}$ betrachten wir ein Polynom in den

n Variablen $x_1, \dots, x_n$. Ein solches Polynom ist eine Summe von sogenannten *Monomen*, das sind Ausdrücke von der Form

$$c \cdot x_1^{p_1} \cdot x_2^{p_2} \cdot \dots \cdot x_n^{p_n},$$

wobei c eine reelle (oder komplexe) Zahl ist und die für Exponenten $p_1, \dots, p_n \in \mathbb{N}_0$ gilt. Die Zahl $p_1 + \dots + p_n$ nennt man den *Totalgrad* des Monoms. Zum Beispiel ist

$$P(x, y, z) = 5x^7 + 4x^2y^3z - \frac{3}{7}yz^2 + 2z - 7$$

ein Polynom in drei Variablen. Der Totalgrad des als zweiten Summanden auftretenden Monoms $4x^2y^3z$ ist $2 + 3 + 1 = 6$. Ein solches Polynom in n Variablen lässt sich in der Form

$$\sum_{(p_1, \dots, p_n) \in \mathbb{N}_0^n} a_{p_1, p_2, \dots, p_n} \cdot x_1^{p_1} \cdot x_2^{p_2} \cdot \dots \cdot x_n^{p_n}$$

schreiben, wobei nur endlich viele der Koeffizienten $a_{p_1, p_2, \dots, p_n}$ von Null verschieden sein dürfen. Der *Grad* des Polynoms ergibt sich als der höchste Totalgrad eines Monoms mit von Null verschiedenem Koeffizient (er wird gleich $-\infty$ für das Nullpolynom gesetzt). Im Beispiel oben ist der Grad des Polynoms $P(x, y, z)$ gleich 7.

4. Lineare Gleichungssysteme

4.1. Das Gaußverfahren. Ein System von Gleichungen

$$\begin{aligned} a_{11}x_1 + \dots + a_{1n}x_n &= b_1, \\ \vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n &= b_m, \end{aligned} \tag{4.1}$$

wobei die Koeffizienten a_{ij} und die Konstanten b_i reelle (oder komplexe) Zahlen sind, heißt *lineares Gleichungssystem* in n reellen (oder komplexen) Unbekannten (oder Variablen). Eine *Lösung* ist ein n -Tupel von reellen (bzw. komplexen) Zahlen $(x_1, \dots, x_n)$, für die alle m Gleichungen simultan erfüllt sind. Die *Lösungsmenge* des Gleichungssystems ist die Menge aller Lösungen.

Beispiel 4.1.1. Wir betrachten das folgende Gleichungssystem im $\mathbb{R}^2$.

$$\begin{aligned} x_1 - 3x_2 &= -1, \\ 3x_1 - 4x_2 &= 0. \end{aligned}$$

Um die Menge der Lösungen zu bestimmen, stellen wir folgende Überlegungen an: Es ändert sich nichts an der Menge der Lösungen, wenn man eine Gleichung

mit einer Zahl ungleich Null multipliziert und es ändert sich ebenfalls nichts an der Menge der Lösungen, wenn man zu einer Gleichung eine andere dazu addiert. Im obigen Beispiel nutzen wir dies, indem wir die erste Gleichung mit -3 multiplizieren und dann zur zweiten Gleichung addieren. Wir erhalten:

$$\begin{aligned}x_1 - 3x_2 &= -1, \\5x_2 &= 3,\end{aligned}$$

wobei wir die erste Gleichung wieder in der ursprünglichen Form übernommen haben (und nicht mit -3 multipliziert). Jetzt kann man bereits die Lösungsmenge bestimmen: Aus der unteren Gleichung folgt, dass $x_2 = \frac{3}{5}$ gilt. Dies können wir in die obere Gleichung einsetzen und erhalten

$$x_1 - \frac{9}{5} = -1,$$

was zu $x_1 = \frac{4}{5}$ äquivalent ist. Insgesamt zeigt dies, dass das Gleichungssystem genau die eine Lösung $(\frac{3}{5}, \frac{4}{5})$ hat.

Beispiel 4.1.2. Wie das folgende Beispiel zeigt, gibt es auch lineare Gleichungssystem mit leerer Lösungsmenge:

$$\begin{aligned}x_1 - 2x_2 &= 1, \\2x_1 - 4x_2 &= 1.\end{aligned}$$

Wir addieren das (-2) -fache der oberen Zeile zur unteren und erhalten

$$\begin{aligned}x_1 - 2x_2 &= 1, \\0 &= -1.\end{aligned}$$

Die zweite Gleichung ist nicht erfüllbar, das Gleichungssystem hat somit keine einzige Lösung.

Beispiel 4.1.3. Ein drittes Beispiel:

$$\begin{aligned}2x_1 - 7x_2 &= 5, \\4x_1 - 14x_2 &= 10.\end{aligned}$$

Wir subtrahieren das 2-fache der oberen Gleichung von der unteren und erhalten:

$$\begin{aligned}2x_1 - 7x_2 &= 5, \\0 &= 0.\end{aligned}$$

Die zweite Gleichung ist somit überflüssig geworden, das gesamte System ist zu $2x_1 - 7x_2 = 5$ äquivalent. Damit sehen wir, dass das Gleichungssystem unendlich

viele Lösungen hat, wir können z.B. x_2 beliebig vorgeben und x_1 daraus bestimmen (oder umgekehrt), etwa durch $x_1 = \frac{5}{2} + \frac{7}{2}x_2$. Die Lösungsmenge ist

$$\mathbb{L} = \left\{ \left(\frac{5}{2} + \frac{7}{2}t, t \right) \in \mathbb{R}^2 \mid t \in \mathbb{R} \right\}.$$

Die drei obigen Beispiele zeigen, dass ein lineares Gleichungssystem eine eindeutig bestimmte Lösung, gar keine Lösungen oder unendlich viele Lösungen haben kann. Wir notieren uns eine Beobachtung.

Satz 4.1.4. *Die folgenden Umformungen ändern nichts an der Lösungsmenge eines linearen Gleichungssystems:*

- (i) *Das Multiplizieren einer Gleichung mit einer von Null verschiedenen Zahl.*
- (ii) *Das Addieren einer Gleichung zu einer anderen.*
- (iii) *Das Ändern der Reihenfolge der Gleichungen.*

BEWEIS. Gilt eine Gleichung, d.h. steht auf beiden Seiten der Gleichung dieselbe Zahl, so erhält man eine weitere gültige Gleichung, wenn man beide Seiten der Gleichung mit einer Zahl λ multipliziert. Alle Lösungen der ursprünglichen Gleichungen sind dann auch Lösungen der mit λ multiplizierten Gleichung. Ist $\lambda \neq 0$, so kann man die Gleichung ein weiteres Mal mit λ^{-1} multiplizieren und erhält so die ursprüngliche Gleichung zurück. Dies zeigt, dass das Multiplizieren mit einer von Null verschiedenen Zahl eine Äquivalenzumformung ist.

Ähnlich kann man beim Addieren einer Gleichung zu einer anderen argumentieren, dies kann man ebenfalls wieder rückgängig machen, was zeigt, dass es sich auch hier um eine Äquivalenzumformung handelt. $\square$

Aus dem Satz folgt insbesondere, dass sich die Lösungsmenge nicht ändert, wenn man ein Vielfaches von einer Gleichung zu einer anderen hinzuaddiert. Dass das Ändern der Reihenfolge der Gleichungen nichts an der Lösungsmenge ändert, ist eine Selbstverständlichkeit, die wir im Beweis gar nicht erwähnt haben. (Wir haben diese Umformung aber der Vollständigkeit halber mit in den Satz aufgenommen).

Man kann sich leicht überlegen, dass man jedes lineare Gleichungssystem durch die in Satz 4.1.4 beschriebenen Umformungen (wie in den Beispielen 4.1.1 bis 4.1.3

vorgeführt) auf die folgende Form, *Stufenform* genannt, bringen kann. Dieses Verfahren (und ähnliche, bei denen die genannten Umformungen verwendet werden) nennt man *Gaußverfahren*.

Definition 4.1.5. Wir sagen, ein lineares Gleichungssystem (4.1), bestehend aus m Gleichungen in n Unbekannten habe *Stufenform*, wenn es $k \in \{0, \dots, m\}$ gibt, so dass folgendes gilt:

- (i) Die Gleichungen in den ersten k Zeilen sind nach dem Auftreten des ersten von Null verschiedenen Koeffizienten a_{ij} geordnet und in jeder Zeile steht der erste von Null verschiedene Koeffizient weiter rechts als in der Zeile darüber. (Genauer formuliert: Für $i, i+1 \leq k$ gilt: Ist in der i -ten Gleichung a_{ij} der von Null verschiedene Koeffizient mit kleinstem zweiten Index j , dann gilt für die $(i+1)$ -te Zeile: $a_{i+1,1} = \dots = a_{i+1,j} = 0$.)
- (ii) In den letzten $m-k$ -Zeilen stehen auf den linken Seiten nur Nullen.

In diesem Fall heißen die in den ersten k -Zeilen am weitesten links stehenden von Null verschiedenen Koeffizienten $a_{1j_1}, a_{2j_2}, \dots, a_{kj_k}$ *Pivotelemente*, die zugehörigen Variablen $x_{j_1}, x_{j_2}, \dots, x_{j_k}$ heißen *Pivotvariablen*.

Die obige Definition soll an einem Beispiel erläutert werden.

Beispiel 4.1.6. Das folgende lineare Gleichungssystem liegt in Stufenform vor.

$$\begin{aligned} 2x_2 + x_3 + x_4 &+ x_6 + 8x_7 = \frac{1}{2}, \\ 7x_4 + 6x_5 + x_6 + 6x_7 &= 2, \\ -3x_5 - x_6 - 2x_7 &= 1, \end{aligned}$$

Die Pivotelemente sind 2, 7 und -3 . Die Pivotvariablen sind x_2, x_4 und x_5 . Nun setzt man für die Nichtpivotvariablen x_1, x_3, x_6, x_7 reelle Parameter t_1, t_2, t_3, t_4 ein, dann kann man, mit der untersten Gleichung beginnend, zuerst nach x_5 auflösen, dann nach x_4 , und schließlich x_2 bestimmen:

$$x_5 = -\frac{1}{3}(t_3 + 2t_4 + 1),$$

damit erhält man

$$x_4 = \frac{1}{7}(-6x_5 - x_6 - 6x_7 + 2) = \frac{1}{7}t_3 - \frac{2}{7}t_4 + \frac{4}{7}$$

und schließlich

$$x_2 = \frac{1}{2}(-x_3 - x_4 - x_6 - 8x_7 + \frac{1}{2}) = -\frac{1}{2}t_2 - \frac{4}{7}t_3 - \frac{27}{7}t_4 - \frac{1}{28}.$$

Man erhält als Lösungsmenge

$$\mathbb{L} = \left\{ \left(\begin{array}{c} t_1 \\ -\frac{1}{2}t_2 - \frac{4}{7}t_3 - \frac{27}{7}t_4 - \frac{1}{28} \\ t_2 \\ \frac{1}{7}t_3 - \frac{2}{7}t_4 + \frac{4}{7} \\ -\frac{1}{3}t_3 - \frac{2}{3}t_4 - \frac{1}{3} \\ t_3 \\ t_4 \end{array} \right) \middle| t_1, t_2, t_3, t_4 \in \mathbb{R} \right\}.$$

Das im Beispiel verwendete Verfahren nennt man auch *Rückwärtseinsetzen*, da man mit dem Bestimmen der letzten Pivotvariablen beginnt. Wie in Beispiel 4.1.2 kann es vorkommen, dass das Gleichungssystem eine leere Lösungsmenge hat (wenn in dem auf Stufenform gebrachten Gleichungssystem eine Gleichung der Form $0 = b_i$ mit einem von Null verschiedenen b_i vorkommt).

Wir fassen zusammen, was wir uns über das Lösen von linearen Gleichungssystemen überlegt haben:

Satz 4.1.7. *Jedes lineare Gleichungssystem lässt sich durch die in Satz 4.1.4 beschriebenen Umformungen auf Stufenform bringen. Dabei können folgende Fälle auftreten:*

- (i) *Es tritt in der Stufenform mindestens eine Gleichung der Form $0 = b_i$ mit einem von Null verschiedenen b_i auf. Dann hat das Gleichungssystem eine leere Lösungsmenge.*
- (ii) *Es tritt in der Stufenform keine Gleichung der Form $0 = b_i$ mit einem von Null verschiedenen b_i auf und die Anzahl der Pivotvariablen ist gleich der Anzahl der Unbekannten. Dann hat das Gleichungssystem genau eine Lösung (und diese kann durch Rückwärtseinsetzen bestimmt werden).*
- (iii) *Es tritt in der Stufenform keine Gleichung der Form $0 = b_i$ mit einem von Null verschiedenen b_i auf und die Anzahl k der Pivotvariablen ist kleiner als die Anzahl n der Unbekannten. Dann hat das Gleichungssystem unendlich viele Lösungen. Die Menge der Lösungen ist durch $n - k$ reelle (oder komplexe) Parameter $t_1, \dots, t_{n-k}$ parametrisiert, d.h. legt man die Werte der Nichtpivotvariablen durch $t_1, \dots, t_{n-k}$ fest, gibt es*

jeweils genau eine Lösung, bei der die Pivotvariablen die vorgegebenen Werte haben (diese kann man dann durch Rückwärtseinsetzen allgemein in Abhängigkeit von $t_1, \dots, t_{n-k}$ bestimmen).

Die drei Fälle in diesem Satz entsprechen Beispiel 4.1.2, Beispiel 4.1.1 und Beispiel 4.1.3 (in dieser Reihenfolge).

4.2. Lineare Gleichungssysteme und Matrizen. Gegeben sei das folgende Gleichungssystem im $\mathbb{R}^n$.

$$\begin{aligned} a_{11}x_1 + \dots + a_{1n}x_n &= b_1, \\ &\vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n &= b_m, \end{aligned}$$

Fasst man die Koeffizienten a_{ij} und die rechten Seiten b_i zu einer $m \times n$ -Matrix A und einem Vektor $b \in \mathbb{R}^m$ zusammen, dann lässt sich das obige Gleichungssystem in der Kurzform

$$Ax = b$$

für ein $x \in \mathbb{R}^n$ schreiben. Es hat sich als praktisch herausgestellt, die rechte Seite b des Gleichungssystems als eine weitere Spalte rechts an die Matrix A anzufügen, die so definierte Matrix $(A|b)$ nennt man auch die zu dem Gleichungssystem gehörige *erweiterte Matrix*. Anstatt das ursprüngliche Gleichungssystem zu lösen, kann man auch mit der erweiterten Matrix arbeiten. Man führt an ihr die in Satz 4.1.4 beschriebenen Umformungen entsprechend durch, wie im folgenden Beispiel:

Beispiel 4.2.1. Wir betrachten das folgende lineare Gleichungssystem in 3 reellen Variablen.

$$\begin{aligned} x_1 + x_2 + 2x_3 &= 1 \\ 2x_1 + 2x_2 - x_3 &= 1 \\ 4x_1 + x_3 &= 2 \end{aligned}$$

In Form einer erweiterten Matrix kann man es wie folgt auf Stufenform bringen.

$$\left(\begin{array}{ccc|c} 1 & 1 & 2 & 1 \\ 2 & 2 & -1 & 1 \\ 4 & 0 & 1 & 2 \end{array} \right) \rightsquigarrow \left(\begin{array}{ccc|c} 1 & 1 & 2 & 1 \\ 0 & 1 & -5 & -1 \\ 0 & -4 & -7 & -2 \end{array} \right) \rightsquigarrow \left(\begin{array}{ccc|c} 1 & 1 & 2 & 1 \\ 0 & 1 & -5 & -1 \\ 0 & 0 & -27 & -6 \end{array} \right)$$

Hier wurden im ersten Schritt das 2-fache der ersten Zeile von der zweiten Zeile und das 4-fache der ersten Zeile von der dritten abgezogen; im zweiten Schritt

wurde das 4-fache der zweiten Zeile zur dritten addiert. Damit ist das Gleichungssystem auf Stufenform gebracht und man kann nun, wie im letzten Abschnitt beschrieben, in dieser Reihenfolge, x_3 , x_2 und dann x_1 berechnen; in diesem Fall hat das Gleichungssystem eine eindeutig bestimmte Lösung, nämlich $(x_1, x_2, x_3) = (\frac{4}{9}, \frac{1}{9}, \frac{2}{9})$.

Ein manchmal nützlicher Trick ist es, ein Gleichungssystem für zwei oder mehrere verschiedene rechte Seiten simultan zu lösen, indem man die Matrix um mehrere rechte Seiten erweitert und dann nur einmal daran die Zeilenumformungen ausführt, d.h. man löst zum Beispiel $Ax = b$ und $Ax = c$ gleichzeitig, indem man mit der erweiterten Matrix $(A|b|c)$ arbeitet.

Die hier verwendeten Umformungen einer Matrix sind so nützlich, dass es eine eigene Bezeichnung dafür gibt.

Definition 4.2.2. Die folgenden Umformungen einer Matrix bezeichnen wir als *elementare Zeilenumformungen* :

- (i) Das Multiplizieren einer Zeile mit einer von Null verschiedenen Zahl.
- (ii) Das Addieren einer Zeile zu einer anderen.
- (iii) Das Ändern der Reihenfolge der Zeilen.

Man sagt, eine Matrix A bzw. eine erweiterte Matrix $(A|b)$ habe *Zeilenstufenform*, falls das zugehörige Gleichungssystem $Ax = 0$ bzw. $Ax = b$ Stufenform hat.

Die vielleicht am häufigsten verwendete Umformung, nämlich das Addieren eines Vielfachen einer Zeile zu einer anderen, besteht natürlich aus einer Hintereinanderausführung von Umformungen vom Typ (i) und (ii).

4.3. Die Inverse einer Matrix. Wir wollen nun lineare Gleichungssysteme $Ax = b$ in einem wichtigen Spezialfall betrachten, nämlich in dem Fall, in dem die Anzahl der Gleichungen mit der Anzahl der Variablen übereinstimmt. Dann ist A eine *quadratische Matrix*, d.h. die Zeilenanzahl stimmt mit der Spaltenanzahl überein. Es können dann zwei Fälle eintreten:

- (i) Das Gleichungssystem $Ax = b$ ist eindeutig lösbar, in diesem Fall ist bei einer Stufenform jede Variable eine Pivotvariable; in diesem Fall erhält

man zu jeder beliebigen rechten Seite b genau eine Lösung x , so dass $Ax = b$ ist.

- (ii) In der Zeilenstufenform von A kommen Nullzeilen vor, es hängt von der rechten Seite b ab, ob es Lösungen des Gleichungssystems $Ax = b$ gibt.

Nur im ersten Fall hat die lineare Abbildung $x \mapsto Ax$ eine Umkehrabbildung. Man kann zeigen, dass die Umkehrabbildung einer linearen Abbildung ebenfalls linear ist. Die darstellende Matrix der Umkehrabbildung nennen wir die zu A *inverse Matrix*, und bezeichnen sie mit A^{-1} . Es gilt

$$AA^{-1} = A^{-1}A = E$$

d.h. beide Matrixprodukte AA^{-1} und $A^{-1}A$ ergeben die $n \times n$ -Einheitsmatrix E , (die ja die darstellende Matrix der identischen Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$, $x \mapsto x$ ist).

Zum praktischen Berechnen der Inversen einer Matrix benutzt man das im folgenden Satz beschriebene Verfahren, wobei $(A|E)$ die Matrix mit n Zeilen und $2n$ Spalten bezeichnet, die man erhält, wenn man die $n \times n$ -Einheitsmatrix rechts an A hinzufügt.

Satz 4.3.1. *Eine reelle $n \times n$ -Matrix A ist genau dann invertierbar, wenn keine Nullzeilen entstehen, wenn Sie durch elementare Zeilenumformungen auf Zeilenstufenform gebracht wird. In diesem Fall kann man die Inverse A^{-1} wie folgt berechnen: Man führt an $(A|E)$ solange elementare Zeilenumformungen durch, bis links die $n \times n$ -Einheitsmatrix entstanden ist, also man eine Matrix der Form $(E|B)$ erhalten hat; dann ist $B = A^{-1}$, die inverse Matrix zu A .*

BEWEIS. Dass das Verfahren funktioniert, kann man sich klar machen, indem man sich überlegt, dass das Umformen mit der um die Einheitsmatrix erweiterten Matrix $(A|E)$ dem simultanen Lösen der Gleichungen

$$Ax = e_1, \dots, Ax = e_n$$

entspricht. Die Spalten auf der rechten Seite in der am Ende erzeugten Matrix sind deswegen $A^{-1}e_1, \dots, A^{-1}e_n$, wie behauptet. $\square$

Das Verfahren zum Invertieren von Matrizen kann man sich kurz auch so merken:

$$(A|E) \rightsquigarrow \dots \rightsquigarrow (E|A^{-1}),$$

wobei die Pfeile $\rightsquigarrow$ wie oben für Zeilenumformungen stehen. Wir verdeutlichen es an zwei Beispielen.

Beispiele 4.3.2. (i) Wir zeigen zunächst: Die Matrix

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix}$$

ist nicht invertierbar. Dazu bringen wir sie auf Zeilenstufenform und stellen fest, dass dabei Nullzeilen entstehen. (Wenn man, wie in diesem Beispiel, nur untersuchen will, ob eine quadratische Matrix invertierbar ist, ist es nicht erforderlich, mit $(A|E)$ zu arbeiten, es genügt, die Umformungen an A auszuführen.)

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} \rightsquigarrow \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & -6 \\ 0 & -6 & -12 \end{pmatrix} \rightsquigarrow \begin{pmatrix} 1 & 2 & 3 \\ 0 & -3 & -6 \\ 0 & 0 & 0 \end{pmatrix}.$$

(ii) Nun zeigen wir, dass die folgende Matrix invertierbar ist und bestimmen ihre Inverse.

$$A := \begin{pmatrix} 2 & 0 & 1 \\ 4 & 0 & 1 \\ 6 & 1 & 0 \end{pmatrix}$$

Dazu führen wir an der Matrix $(A|E)$ die folgenden Zeilenumformungen durch.

$$\begin{aligned} \left(\begin{array}{ccc|ccc} 2 & 0 & 1 & 1 & 0 & 0 \\ 4 & 0 & 1 & 0 & 1 & 0 \\ 6 & 1 & 0 & 0 & 0 & 1 \end{array} \right) &\rightsquigarrow \left(\begin{array}{ccc|ccc} 2 & 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & -1 & -2 & 1 & 0 \\ 0 & 1 & -3 & -3 & 0 & 1 \end{array} \right) \rightsquigarrow \\ \left(\begin{array}{ccc|ccc} 2 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & -3 & -3 & 0 & 1 \\ 0 & 0 & -1 & -2 & 1 & 0 \end{array} \right) &\rightsquigarrow \left(\begin{array}{ccc|ccc} 2 & 0 & 0 & -1 & 1 & 0 \\ 0 & 1 & 0 & 3 & -3 & 1 \\ 0 & 0 & -1 & -2 & 1 & 0 \end{array} \right) \rightsquigarrow \\ \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & -\frac{1}{2} & \frac{1}{2} & 0 \\ 0 & 1 & 0 & 3 & -3 & 1 \\ 0 & 0 & 1 & 2 & -1 & 0 \end{array} \right) \end{aligned}$$

Die Matrix A ist also invertierbar und

$$A^{-1} = \begin{pmatrix} -\frac{1}{2} & \frac{1}{2} & 0 \\ 3 & -3 & 1 \\ 2 & -1 & 0 \end{pmatrix}$$

ist ihre Inverse.

4.4. Untervektorräume.

Definition 4.4.1. Eine nichtleere Teilmenge U des $\mathbb{R}^n$ heißt *Untervektorraum* oder *linearer Unterraum* des $\mathbb{R}^n$, wenn sie abgeschlossen unter Addition und skalarer Multiplikation von Vektoren ist. Genauer: Falls $u + v \in U$ und $\lambda u \in U$ für alle $u, v \in U$ und alle $\lambda \in \mathbb{R}$ gilt.

- Beispiele 4.4.2.**
- (i) Anschauliche Beispiele für Untervektorräume des $\mathbb{R}^2$ sind Geraden durch den Ursprung, die Menge, die nur aus dem Nullvektor besteht und der $\mathbb{R}^2$ selbst.
 - (ii) Anschauliche Beispiele für Untervektorräume des $\mathbb{R}^3$ sind Geraden, die durch den Ursprung verlaufen, Ebenen, die den Ursprung enthalten, die Menge, die nur aus dem Nullvektor besteht und der $\mathbb{R}^3$ selbst.
 - (iii) Die Teilmenge

$$U := \{(x_1, x_2, x_3) \in \mathbb{R}^3 \mid x_1 + x_2 + x_3 = 0\}$$

ist ein Untervektorraum des $\mathbb{R}^3$, denn sind (x_1, x_2, x_3) und (y_1, y_2, y_3) aus U , dann gilt auch für ihre Summe $(x_1 + y_1, x_2 + y_2, x_3 + y_3)$, dass die Summe der Komponenten gleich Null ist, außerdem gilt $\lambda x_1 + \lambda x_2 + \lambda x_3 = \lambda(x_1 + x_2 + x_3) = 0$.

- (iv) *Kein* Untervektorraum ist zum Beispiel die obere Halbebene im $\mathbb{R}^2$,

$$H := \{(x_1, x_2) \in \mathbb{R}^2 \mid x_2 > 0\},$$

denn diese enthält etwa den Vektor $v = (0, 1)$, aber nicht λv für $\lambda = -1$.

Da ein Untervektorraum nach Definition keine leere Menge ist, enthält er mindestens einen Vektor v und somit, für $\lambda = 0$, auch $\lambda v = 0$. Dies zeigt, dass jeder Untervektorraum den Nullvektor enthält.

Definition 4.4.3. Seien $v_1, \dots, v_k$ Vektoren im $\mathbb{R}^n$. Eine *Linearkombination* von $v_1, \dots, v_k$ ist ein Vektor der Form

$$\lambda_1 v_1 + \dots + \lambda_k v_k$$

wobei die Koeffizienten $\lambda_1, \dots, \lambda_k$ reelle Zahlen sind. Für eine Teilmenge $U \subseteq \mathbb{R}^n$ bezeichnet man die Menge der Vektoren, die sich als Linearkombination von (endlich vielen) Elementen aus U schreiben lassen, als den *Spann* von U , man schreibt auch $\text{span}(U)$ dafür. Gilt $V = \text{span}(U)$, dann sagt man auch, dass V von der Menge U *aufgespannt* wird.

Man überlegt sich leicht, dass die Summe zweier Linearkombinationen von Elementen aus U und das λ -fache einer Linearkombination von Elementen aus U wieder Linearkombinationen von Elementen aus U sind. Dies zeigt:

Satz 4.4.4. *Der Spann einer Teilmenge von $\mathbb{R}^n$ ist ein Untervektorraum.*

Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine lineare Abbildung und sei A die darstellende Matrix. Wir nennen ein lineares Gleichungssystem $Ax = b$ *homogen*, wenn die rechte Seite b der Nullvektor ist, andernfalls *inhomogen*. Ist $Ax = b$ ein lineares Gleichungssystem, dann nennen wir $Ax = 0$ das *zugehörige homogene Gleichungssystem*.

Satz 4.4.5. *Die Lösungsmenge eines homogenen linearen Gleichungssystems ist ein Untervektorraum des $\mathbb{R}^n$.*

BEWEIS. Seien v und w zwei Lösungen von $f(x) = 0$. Wegen der Linearität von f gilt $f(v + w) = f(v) + f(w) = 0 + 0 = 0$. Außerdem gilt für $\lambda \in \mathbb{R}$, dass $f(\lambda v) = \lambda f(v) = \lambda 0 = 0$ ist. Somit ist die Lösungsmenge eines homogenen linearen Gleichungssystems abgeschlossen unter Addition und skalarer Multiplikation von Vektoren. Die Lösungsmenge eines homogenen linearen Gleichungssystems ist nicht leer, da sie immer mindestens den Nullvektor enthält. $\square$

Die Lösungsmenge des homogenen linearen Gleichungssystems $f(x) = 0$ bezeichnet man auch als den *Kern* der linearen Abbildung f , in Zeichen: $\ker(f) := f^{-1}(\{0\})$. Da wir Matrizen auch als lineare Abbildungen interpretieren, schreiben wir auch $\ker(A)$ für den Kern der durch A gegebenen linearen Abbildung.

Nach dem obigen Satz 4.4.5 ist der Kern einer linearen Abbildung stets ein Untervektorraum. Allgemeiner gilt:

Satz 4.4.6. *Sei $Ax = b$ ein lineares Gleichungssystem. Dann ist die Lösungsmenge entweder leer oder von der Form*

$$x^* + \ker(A) = \{x^* + v \mid v \in \ker(A)\},$$

wobei x^ eine beliebige Lösung des inhomogenen Gleichungssystems $Ax = b$ ist.*

BEWEIS. Wir zeigen zuerst, dass jedes Element von $x^* + \ker(A)$ das inhomogene Gleichungssystem löst: Sei dazu v ein beliebiges Element aus dem Kern von A ; dann gilt $A(x^* + v) = A(x^*) + A(v) = b + 0 = b$.

Umgekehrt, sei w eine Lösung des inhomogenen Systems. Um zu zeigen, dass sich w als $x^* + v$ mit $v \in \ker(A)$ schreiben lässt, müssen wir zeigen, dass $w - x^*$ das zugehörige homogene System $Ax = 0$ löst: Es gilt $A(w - x^*) = A(w) - A(x^*) = b - b = 0$. $\square$

Beispiel 4.4.7. Die Lösungsmenge im Beispiel 4.1.6 ist von der Form

$$\begin{pmatrix} 0 \\ -\frac{1}{28} \\ 0 \\ \frac{4}{7} \\ -\frac{1}{3} \\ 0 \\ 0 \end{pmatrix} + \text{span} \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ -\frac{1}{2} \\ 1 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ -\frac{4}{7} \\ 0 \\ \frac{1}{7} \\ -\frac{1}{3} \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ -\frac{27}{7} \\ 0 \\ -\frac{2}{7} \\ -\frac{2}{3} \\ 0 \\ 1 \end{pmatrix} \right\}$$

Bemerkung 4.4.8. Eine Teilmenge von $\mathbb{R}^n$ heißt *affiner Unterraum* von $\mathbb{R}^n$, wenn sie entweder leer oder von der Form $x^* + V$ ist, wobei $x^* \in \mathbb{R}^n$ ein Vektor ist und V ein Untervektorraum von $\mathbb{R}^n$. Beispiele für affine Unterräume des $\mathbb{R}^3$ sind beliebige Punkte, Gerade und Ebenen (nicht nur solche, die den Nullpunkt enthalten). Die beiden obigen Sätze sagen also, dass die Lösungen von homogenen Gleichungssystemen stets Untervektorräume, die von inhomogenen Gleichungssystem affine Unterräume sind.

Satz 4.4.9. Das Bild einer linearen Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ ist ein Untervektorraum des $\mathbb{R}^m$.

BEWEIS. Sei A die darstellende Matrix von f . Dann sind die Elemente des Bildes von f genau die Vektoren, die sich als Ax für ein $x \in \mathbb{R}^n$ darstellen lassen. Dies zeigt, dass das Bild einer linearen Abbildung gleich dem Spann der Spaltenvektoren der darstellenden Matrix ist. $\square$

4.5. Lineare Unabhängigkeit und Basen. Man nennt k Vektoren $v_1, \dots, v_k$ im $\mathbb{R}^n$ *linear unabhängig*, wenn sich keiner der Vektoren als Linearkombination der anderen ausdrücken lässt, ansonsten nennt man sie *linear abhängig*. (Bei einem Vektor bedeutet lineare Unabhängigkeit, dass der Vektor von Null verschieden ist, sich also nicht als eine Linearkombination von null Vektoren – was eine leere Summe wäre – schreiben lässt.) Die Standardbasisvektoren des $\mathbb{R}^n$ sind beispielsweise offensichtlich linear unabhängig.

Beispiele 4.5.1. (i) Die folgenden Vektoren im $\mathbb{R}^3$ sind linear abhängig:

$$v_1 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}, v_2 = \begin{pmatrix} 1 \\ -1 \\ -2 \end{pmatrix}, v_3 = \begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix},$$

denn es gilt $v_3 = 2v_1 + v_2$.

(ii) Die folgenden Vektoren im $\mathbb{R}^3$ sind linear unabhängig:

$$v_1 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, v_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, v_3 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix},$$

wie man sich leicht überlegt.

Definition 4.5.2. Man sagt, dass die Vektoren $v_1, \dots, v_k$ eine *Basis* eines Untervektorraums V von $\mathbb{R}^n$ bilden, wenn sie linear unabhängig sind und V von ihnen aufgespannt wird.

Wie bestimmt man von einem Untervektorraum eine Basis, wenn dieser etwa als Spann einer Menge von Vektoren gegeben ist? Ähnlich wie in Beispiel 4.5.1 (ii) sieht man, dass bei einer Matrix in Zeilenstufenform die von Null verschiedenen Zeilen linear unabhängig sind. Außerdem gilt, wie man sich leicht klar macht: Führt man an einer Matrix elementare Zeilenumformungen durch, so ändert sich nichts am Spann der Zeilenvektoren. Diese beiden Aussagen zusammen liefern uns folgendes Verfahren, um eine Basis eines Untervektorraums V von $\mathbb{R}^n$ zu bestimmen:

- Seien die Vektoren $v_1, \dots, v_k$ mit $\text{span}\{v_1, \dots, v_k\} = V$ gegeben.
- Bilde die Matrix, deren Zeilen die Vektoren $v_1, \dots, v_k$ sind.
- Bringe diese mit elementaren Zeilenumformungen auf Zeilenstufenform.
- Die von Null verschiedenen Zeilen bilden dann eine Basis von V .

Beispiel 4.5.3. Wir verdeutlichen das Verfahren an einem Beispiel. Seien im $\mathbb{R}^4$ die Vektoren

$$\begin{pmatrix} -1 \\ -2 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ -1 \\ 0 \\ 2 \end{pmatrix}, \begin{pmatrix} 2 \\ 2 \\ 1 \\ 2 \end{pmatrix}, \begin{pmatrix} 0 \\ -1 \\ 1 \\ 2 \end{pmatrix}, \begin{pmatrix} -2 \\ -3 \\ -1 \\ 0 \end{pmatrix}$$

gegeben und sei $V \subseteq \mathbb{R}^4$ der Spann dieser Vektoren. Wir tragen diese Vektoren als Zeilen in eine Matrix ein und bringen diese auf Zeilenstufenform.

$$\begin{pmatrix} -1 & -2 & -1 & 1 \\ 0 & -1 & 0 & 2 \\ 2 & 2 & 1 & 2 \\ 0 & -1 & 1 & 2 \\ -2 & -3 & -1 & 0 \end{pmatrix} \rightsquigarrow \begin{pmatrix} -1 & -2 & -1 & 1 \\ 0 & -1 & 0 & 2 \\ 0 & -2 & -1 & 4 \\ 0 & -1 & 1 & 2 \\ 0 & 1 & 1 & -2 \end{pmatrix} \rightsquigarrow$$

$$\rightsquigarrow \begin{pmatrix} -1 & -2 & -1 & 1 \\ 0 & -1 & 0 & 2 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 \end{pmatrix} \rightsquigarrow \begin{pmatrix} -1 & -2 & -1 & 1 \\ 0 & -1 & 0 & 2 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Eine Basis von V ist also durch die Vektoren

$$\begin{pmatrix} -1 \\ -2 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ -1 \\ 0 \\ 2 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ -1 \\ 0 \end{pmatrix}$$

gegeben.

Teilweise ohne Beweise wollen wir noch einige nützliche Tatsachen notieren: Alle Basen eines Untervektorraums des $\mathbb{R}^n$ haben gleich viele Elemente, daher kann man die *Dimension* eines Untervektorraums V als diese Anzahl definieren, wir schreiben dafür kurz $\dim(V)$. Der $\mathbb{R}^n$ selbst hat somit die Dimension n , denn die Standardbasisvektoren bilden in der Tat eine Basis mit n Elementen. Die Dimension des Bildes einer linearen Abbildung f nennt man den *Rang* der Abbildung, man schreibt dafür $\text{rk}(f)$. Für lineare Abbildungen $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ gilt die *Dimensionsformel*

$$n = \text{rk}(f) + \dim(\ker(f)).$$

5. Grenzwerte und Stetigkeit

5.1. Folgen und Konvergenz.

Definition 5.1.1. Eine Zahlenfolge ist eine Abbildung $k \mapsto a_k, \mathbb{N} \rightarrow \mathbb{R}$, bzw. $\mathbb{N} \rightarrow \mathbb{C}$. Man spricht genauer auf von einer *reellen Zahlenfolge* bzw. *komplexen*

Zahlenfolge. Allgemeiner ist eine *Folge* eine Abbildung $\mathbb{N} \rightarrow \mathbb{R}^n$, $k \mapsto a_k$. Die Elemente a_k nennt man auch die *Folgenglieder*.

Man schreibt oft auch (a_k) , $(a_k)_{\mathbb{N}}$, oder $(a_k)_{k \in \mathbb{N}}$, um auszudrücken, dass es sich bei den Elementen a_k um die Glieder einer Folge handelt. Wir schreiben auch $a_n = (a_1, a_2, a_3, \dots)$, um die ersten Glieder einer Folge anzugeben.

- Beispiele 5.1.2.**
- (i) Die Folge der natürlichen Zahlen $a_k = k = (1, 2, 3, \dots)$.
 - (ii) Die Folge der Kehrwerte der natürlichen Zahlen $a_k = \frac{1}{k}$, die sogenannte *harmonische Folge*.
 - (iii) Eine *arithmetische Folge* ist eine Zahlenfolge, bei der die Differenz aufeinander folgender Glieder konstant bleibt, d.h. es gilt $a_{k+1} = a_k + d$. Dann gilt $a_{k+1} = a_1 + k \cdot d$.
 - (iv) Eine *geometrische Folge* ist eine Zahlenfolge, bei der der Quotient aufeinander folgender Glieder konstant bleibt, d.h. es gilt $a_{k+1} = q \cdot a_k \neq 0$. Dann gilt $a_{k+1} = q^k \cdot a_1$. Z.B. ist $a_n = (1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots)$ so eine Folge.
 - (v) *Rekursiv definierte Folgen* sind solche, bei denen ein oder mehrere Folgenglieder als Startwerte gegeben sind und sich die anderen durch ein rekursives Bildungsgesetz berechnen. Die bekannteste solche Folge dürfte die *Fibonacci-Folge* sein:

$$a_1 := 0, \quad a_2 := 1, \quad a_k := a_{k-1} + a_{k-2}$$

Für diese gilt: $a_k = (0, 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, \dots)$.

Nähert sich der Wert einer Folge a_k mit steigendem Index k immer mehr einem festen Wert an, so spricht man von einem Grenzwert. Zum Beispiel hat die Folge

$$a_k = \frac{1}{k}, \quad a_k = (1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots)$$

den Grenzwert 0, sie ist eine *Nullfolge*. Das “immer mehr” muss man präzisieren:

Definition 5.1.3. Ein Element $a \in \mathbb{R}^n$ heißt *Grenzwert* oder *Limes* der Folge a_k , wenn es für jede (auch noch so kleine) positive Zahl $\varepsilon \in \mathbb{R}$ eine Schranke $N \in \mathbb{N}$ gibt, so dass für a_N und alle nach a_N kommenden Folgenglieder a_k , $k \geq N$ gilt:

$$d(a_k, a) < \varepsilon,$$

d.h. ab dem Index N haben alle Folgenglieder von a einen kleineren Abstand als ε . Gibt es einen Grenzwert a , so sagt man die Folge ist *konvergent* oder sie *konvergiert*.

gegen a (auch: “ a_k geht gegen a ”) und schreibt dafür in Zeichen

$$\lim_{k \rightarrow \infty} a_k = a \quad \text{oder} \quad a_k \rightarrow a$$

Andernfalls sagt man, die Folge ist *divergent* oder sie *divergiert*.

Wir erinnern uns daran, dass die Abstandsfunktion im $\mathbb{R}^n$ durch

$$d(x, y) = \|y - x\|$$

gegeben ist, im Fall $n = 1$, im Fall einer reellen Zahlenfolge ist dies gleich $|y - x|$, im Fall einer komplexen Zahlenfolge kann man ebenfalls $|y - x|$ für den Abstand zweier Punkte x und y schreiben, da der komplexe Betrag mit der Norm im $\mathbb{R}^2$ übereinstimmt.

Beispiel 5.1.4. Als Beispiel geben wir einen Beweis dafür an, dass die Folge $a_k = \frac{1}{k}$ gegen Null konvergiert.

BEWEIS. Wir müssen beweisen, dass es zu jedem beliebigen $\varepsilon > 0$ ein $N \in \mathbb{N}$ gibt, so dass ab diesem Index die Folgenglieder a_k von 0 kleineren Abstand als ε haben, so dass also gilt

$$|a_k - 0| = \left| \frac{1}{k} - 0 \right| = \frac{1}{k} < \varepsilon.$$

Diese Ungleichung gilt dann, wenn k größer als $\frac{1}{\varepsilon}$ ist, eine passende Schranke wäre also $\lceil \frac{1}{\varepsilon} \rceil$, dies ist die Kurzschreibweise für die kleinste natürliche Zahl, die größer als $\frac{1}{\varepsilon}$ ist. $\square$

Bemerkungen 5.1.5. (i) Mit Hilfe von ε -Umgebungen (siehe Abschnitt 2.2)

kann man die Definition wie folgt umformulieren: Die Folge a_k konvergiert gegen a , wenn es zu jeder ε -Umgebung $B_\varepsilon(a)$ von a eine Schranke $N \in \mathbb{N}$ gibt, so dass ab dem Index N alle Folgenglieder a_k , $k \geq N$ in $B_\varepsilon(a)$ liegen.

(ii) Eine sehr elegante Formulierung der Definition von Konvergenz ist die folgende: *Eine Folge a_k konvergiert gegen a , wenn in jeder ε -Umgebung von a fast alle Folgenglieder a_k liegen.* Hierbei verwendet man “fast alle” als Kurzsprechweise für “alle bis auf endlich viele”.

(iii) Für reelle Zahlenfolgen benutzen wir auch folgende Bezeichnungen. Man schreibt $a_k \rightarrow +\infty$, wenn gilt: Für alle $R > 0$ gibt es ein $N \in \mathbb{N}$, so dass $a_k > R$ für alle $k \geq N$ gilt. Dies nennt man auch *bestimmte Divergenz*, man sagt dann auch die Folge “geht gegen unendlich”. (Analog schreibt man $a_k \rightarrow -\infty$, falls $-a_k \rightarrow \infty$.)

- (iv) Eine divergente Folge muss nicht notwendig gegen $+\infty$ oder $-\infty$ gehen, z.B. sind die Glieder der Folge $a_k = (-1)^k$ abwechselnd -1 und $+1$; es gilt $a_k = (-1, 1, -1, 1, \dots)$.
- (v) Die Konvergenz einer Folge im $\mathbb{R}^n$ ist dazu äquivalent, dass die n Komponenten der Folge alle einzeln konvergent sind.

BEWEIS. Gilt $a_k = (a_{1k}, \dots, a_{nk}) \rightarrow (a_1, \dots, a_n) = a$, so kann man zu jedem $\varepsilon > 0$ eine Schranke für die Indizes k finden, ab der $\|a_k - a\| = \sqrt{(a_{1k} - a_1)^2 + \dots + (a_{nk} - a_n)^2}$ kleiner als ε bleibt, daraus folgt, dass dies dann auch für $|a_{jk} - a_j|$ gilt.

Umgekehrt, konvergieren die Komponenten und ist ein $\varepsilon > 0$ vorgegeben, so finden wir für jede Komponente einzeln eine Indexschranke, ab der der Abstand $d(a_{jk}, a_j)$ kleiner als $\sqrt{\frac{\varepsilon^2}{n}}$ bleibt. Wählen wir nun das Maximum dieser n Schranken als neue Schranke, so folgt $\|a_k - a\| = \sqrt{(a_{1k} - a_1)^2 + \dots + (a_{nk} - a_n)^2} < \sqrt{n \cdot \frac{\varepsilon^2}{n}} = \varepsilon$ ab dieser Indexschranke. $\square$

Insbesondere ist die Konvergenz einer komplexen Zahlenfolge dazu äquivalent, dass Real- und Imaginärteil konvergent sind.

Beispiel 5.1.6. Die Folge $a_k \in (\frac{1}{k}, \frac{1}{k}, \dots, \frac{1}{k})$ im $\mathbb{R}^n$ konvergiert nach Bemerkung 5.1.5 (v) gegen den Nullvektor $(0, 0, \dots, 0)$, denn alle Komponenten sind Nullfolgen in $\mathbb{R}$.

Ohne Beweise geben wir die folgenden nützlichen Sätze über Konvergenz und Limites von Folgen an.

Satz 5.1.7. *Ist eine reelle Zahlenfolge monoton (d.h. gilt entweder $a_{k+1} \geq a_k$ für alle $k \in \mathbb{N}$ oder $a_{k+1} \leq a_k$ für alle $k \in \mathbb{N}$) und beschränkt (d.h. gibt es ein $C > 0$, so dass $|a_k| \leq C$ für alle Folgenglieder gilt) so ist sie konvergent.*

Satz 5.1.8 (Rechenregeln für Limiten). *Seien $a_k \rightarrow a$ und $b_k \rightarrow b$ konvergente (reelle oder komplexe) Zahlenfolgen und c eine feste (reelle oder komplexe) Zahl.*

Dann gelten folgende Rechenregeln:

$$\lim_{k \rightarrow \infty} (a_k \pm b_k) = a \pm b;$$

$$\lim_{k \rightarrow \infty} (a_k \cdot b_k) = a \cdot b;$$

$$\lim_{k \rightarrow \infty} (c \cdot a_k) = c \cdot a;$$

$$\lim_{k \rightarrow \infty} \frac{a_k}{b_k} = \frac{a}{b}, \quad \text{falls } b_k, b \neq 0.$$

Beispiele 5.1.9. (i) Ist a_k eine Zahlenfolge mit von Null verschiedenen Gliedern, so gilt

$$a_k \rightarrow 0 \iff \left| \frac{1}{a_k} \right| \rightarrow \infty.$$

(ii) Ist die Folge a_k eine *rationale Funktion* (=Quotient zweier Polynome) in k , dann kann man das Konvergenzverhalten untersuchen, indem man die höchste im Nenner vorkommende Potenz von k kürzt, z.B. wie folgt:

$$\frac{2k^3 + 5k^2 + k}{7k^3 - 2k} = \frac{k^3(2 + 5\frac{1}{k} + \frac{1}{k^2})}{k^3(7 - \frac{2}{k^2})} = \frac{2 + 5\frac{1}{k} + \frac{1}{k^2}}{7 - \frac{2}{k^2}} \longrightarrow \frac{2}{7}.$$

Ist der Grad des Nenners größer als der Grad des Zählers, so handelt es sich um eine Nullfolge. Ist der Grad des Nenners kleiner als der Grad des Zählers, so divergiert die Folge bestimmt, und zwar gegen $+\infty$, wenn die Leitkoeffizienten von Zähler und Nenner gleiches Vorzeichen haben und gegen $-\infty$, wenn sie verschiedenes Vorzeichen haben, z.B. gilt

$$a_k = \frac{k^7 + 9k}{k^4 + 2k^3 - 7} \longrightarrow +\infty,$$

$$b_k = \frac{k^7 + k^5}{-3k^4 + 9k^2} \longrightarrow -\infty.$$

$$c_k = \frac{5k^3 - 7k^2}{-4k^8 + 17k^2} \longrightarrow 0.$$

5.2. Reihen. Ist a_k eine Zahlenfolge, so erhält man daraus durch Aufsummieren der jeweils ersten k Folgenglieder eine neue Zahlenfolge

$$s_k := \sum_{n=1}^k a_n = a_1 + a_2 + \dots + a_k,$$

wir betrachten hier also die Folge

$$a_1, \quad a_1 + a_2, \quad a_1 + a_2 + a_3, \quad a_1 + a_2 + a_3 + a_4, \quad \dots$$

Eine auf diese Weise gebildete Folge nennt man *Reihe* über a_k ; ihre Folgenglieder s_k nennt man die *Partialsummen* der Reihe. Man schreibt auch

$$\sum_{n=1}^{\infty} a_n$$

für die Folge der Partialsummen. Diese Schreibweise wird auch als Abkürzung für den Grenzwert

$$\lim_{k \rightarrow \infty} \sum_{n=1}^k a_n$$

im Falle der Konvergenz verwendet. In diesem Falle hat das Symbol $\sum_{n=1}^{\infty} a_n$ also eine Doppelbedeutung. Die Zahlen a_n nennt man die *Glieder* der Reihe.

Ein wichtiges Beispiel ist die sogenannte *geometrische Reihe*

$$\sum_{n=0}^{\infty} q^n.$$

Man erhält z.B. für $q = \frac{1}{2}$ die Reihe

$$\sum_{n=0}^{\infty} q^n = 1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots$$

Die Partialsummenfolge dieser Reihe ist also $(1, 1\frac{1}{2}, 1\frac{3}{4}, 1\frac{7}{8}, \dots)$ und man sieht, dass sie gegen 2 konvergiert. Allgemein hat man die folgende Formel.

Satz 5.2.1 (Geometrische Summenformel). *Sei q eine reelle (oder auch komplexe) Zahl. Für die geometrische Summe*

$$\sum_{n=0}^k q^n = 1 + q + q^2 + q^3 + \dots + q^k$$

gilt, falls $q \neq 1$ ist

$$\sum_{n=0}^k q^n = \frac{1 - q^{k+1}}{1 - q}.$$

BEWEIS. Es gilt

$$\begin{aligned} (1 - q)(1 + q + \dots + q^k) &= \\ &= 1 + q + \dots + q^k - q - q^2 - \dots - q^{k+1} = \\ &= 1 - q^{k+1}. \end{aligned}$$

□

Beispiel 5.2.2 (Anwendung: Kapitalentwicklung/Rentenrechnung). Eine *Rente* ist eine regelmäßige Zahlung. Zum Beispiel werde in einen Sparvertrag jedes Jahr zum Jahresanfang 2000€ eingezahlt. Das Guthaben werde jedes Jahr am Jahresende mit 4% verzinst. Wie hoch ist das Guthaben nach 5 Jahren?

Sei $r = 2000$ die *Rate* und sei $q = 1 + \frac{4}{100} = 1,04$ der sogenannte *Zinsfaktor*. Die Verzinsung jeweils am Jahresende kann man rechnerisch dadurch erhalten, indem man das Kapital mit dem Faktor q multipliziert, die Einzahlung erhält man durch Addition von r . Für den Betrag nach 5 Jahren erhält man so:

$$\begin{aligned} K_5 &= (((r q + r) q + r) q + r) q + r = \\ &= r q^5 + r q^4 + r q^3 + r q^2 + r q = \\ &= r(q^5 + q^4 + q^3 + q^2 + q) = \\ &= r \cdot q \cdot \sum_{n=0}^4 q^n = r \cdot q \cdot \frac{1 - q^5}{1 - q}. \end{aligned}$$

In unserem Zahlenbeispiel erhält man als Kapital nach 5 Jahren

$$K_5 = 2000\text{€} \cdot 1,04 \cdot \frac{1 - 1,04^5}{-0,04} \approx 11.265\text{€}.$$

Allgemein halten wir die Formel

$$K_n = r \cdot q \cdot \frac{1 - q^n}{1 - q}$$

für die Kapitalentwicklung nach n Jahren bei *vorschüssiger Zahlweise* einer Rate r bei einem Zinsfaktor q fest. Bei *nachschüssiger Zahlweise*, also bei Zahlung am Ende jeder Zinsperiode (hier im Beispiel am Ende jedes Jahres) hat man die Formel

$$K_n = r \cdot \frac{1 - q^n}{1 - q}.$$

Eine weitere Anwendung der geometrischen Summenformel besteht darin, dass sie über das Konvergenzverhalten und gegebenenfalls den Grenzwert einer geometrischen Reihe Auskunft gibt.

Satz 5.2.3. Sei q eine reelle (oder komplexe) Zahl. Die geometrische Reihe

$$\sum_{n=0}^{\infty} q^n$$

konvergiert gegen $\frac{1}{1-q}$, falls $|q| < 1$ ist. Andernfalls divergiert die Reihe.

BEWEIS. Für die Partialsummen der geometrischen Reihe gilt nach Satz 5.2.1, falls $|q| < 1$ ist:

$$\sum_{n=0}^{\infty} q^n = \frac{1 - q^{k+1}}{1 - q} \rightarrow \frac{1}{1 - q},$$

denn es gilt $q^n \rightarrow 0$ für jede komplexe Zahl q , deren Betrag kleiner eins ist, was wir als bekannt voraussetzen. (Es genügt, dieses Resultat für reelle q zu kennen, da sich beim Potenzieren von komplexen Zahlen die Beträge potenzieren).

Die zweite Aussage folgt aus dem nachfolgenden Satz 5.2.4, da q^k für $|q| \geq 1$ keine Nullfolge ist. $\square$

Der folgende Satz gibt ein *notwendiges* Kriterium für die Konvergenz einer Reihe: Wenn die Glieder der Reihe keine Nullfolge bilden, dann divergiert die Reihe.

Satz 5.2.4. Ist $\sum_{n=0}^{\infty} a_n$ eine konvergente Reihe, so konvergiert die Folge a_n gegen Null.

BEWEIS. Es gilt für die Glieder einer Reihe $\sum_{n=0}^{\infty} a_n$:

$$a_k = \sum_{n=0}^k a_n - \sum_{n=0}^{k-1} a_n.$$

Konvergiert die Reihe, dann gehen beide Terme auf der rechten Seite gegen denselben Grenzwert, also konvergiert a_k nach den Rechenregeln für Limiten 5.1.8 gegen Null. $\square$

Das obige Kriterium ist allerdings *kein hinreichendes* Kriterium für die Konvergenz einer Reihe, wie das folgende Beispiel zeigt:

Beispiel 5.2.5. Die *harmonische Reihe*

$$\sum_{n=1}^{\infty} \frac{1}{n} = 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \dots$$

divergiert (bestimmt gegen $+\infty$). Dies macht man sich leicht klar, indem man die Reihenglieder wie folgt zusammenfasst:

$$1 + \frac{1}{2} + \left(\frac{1}{3} + \frac{1}{4}\right) + \left(\frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8}\right) + \left(\frac{1}{9} + \dots + \frac{1}{16}\right) + \dots$$

Jeder der zu Klammern zusammengefassten Terme ist größer als $\frac{1}{2}$, die Partialsummen übersteigen also jede reelle Zahl irgendwann. Die Reihe geht somit gegen $+\infty$, wenn auch sehr langsam.

Ohne Beweis verwenden wir:

Hilfssatz 5.2.6. *Konvergiert $\sum_{n=1}^{\infty} |a_n|$, so auch $\sum_{n=1}^{\infty} a_n$. (“Aus absoluter Konvergenz folgt Konvergenz.”)*

Man beachte, dass die Umkehrung nicht gilt, so konvergiert etwa die *alternierende harmonische Reihe* $\sum_{n=0}^{\infty} \frac{(-1)^n}{n}$.

Ein sehr nützliches Kriterium zum Beurteilen des Konvergenzverhaltens von Reihen ist der folgende Satz.

Satz 5.2.7 (Majorantenkriterium). *Ist eine $\sum_{n=0}^{\infty} b_n$ eine konvergente Reihe, deren Reihenglieder b_k nichtnegative reelle Zahlen sind und $\sum_{n=0}^{\infty} a_n$ eine weitere Reihe (mit reellen oder komplexen Reihengliedern), für die*

$$|a_n| \leq b_n$$

gilt (man sagt dann auch, $\sum_{n=0}^{\infty} b_n$ ist eine Majorante von $\sum_{n=0}^{\infty} a_n$), dann konvergiert $\sum_{n=0}^{\infty} a_n$.

BEWEIS. Wir betrachten zunächst die Partialsummenfolge $\sum_{n=0}^k |a_n|$. Wegen $|a_n| \leq b_n$ ist sie durch den Limes der Reihe $\sum_{n=0}^{\infty} b_n$ nach oben beschränkt, und da die Reihenglieder $|a_n|$ nichtnegativ sind, handelt es sich um eine monoton wachsende Folge. Nach Satz 5.1.7 konvergiert sie, nach Hilfssatz 5.2.6 konvergiert damit auch $\sum_{n=0}^k a_n$. $\square$

Damit können wir die Konvergenz einer sehr wichtigen Reihe beweisen.

Satz 5.2.8. *Für alle reellen (und komplexen) Zahlen konvergiert die Exponentialreihe*

$$\sum_{n=0}^{\infty} \frac{x^n}{n!} = 1 + x + \frac{x^2}{2} + \frac{x^3}{6} + \frac{x^4}{24} + \dots$$

BEWEIS. Zum Beweis berechnen wir zunächst den Betrag des Quotienten zweier aufeinanderfolgender Reihenglieder

$$\left| \frac{\frac{x^{n+1}}{(n+1)!}}{\frac{x^n}{n!}} \right| = \left| \frac{x^{n+1} \cdot n!}{x^n \cdot (n+1)!} \right| = \left| \frac{x}{n+1} \right|$$

Wir wählen ein $k \in \mathbb{N}$ größer als x , damit gilt dann

$$\left| \frac{x}{k+1} \right| < 1.$$

Nun setzen wir

$$q := \left| \frac{x}{k+1} \right|$$

Ab dem Index k sind dann die Glieder der Exponentialreihe beschränkt durch die der Reihe $\sum_{n=0}^{\infty} c \cdot q^n$ für ein geeignetes $c > 0$. Nach dem Majorantenkriterium in Satz 5.2.7 folgt die Konvergenz der Exponentialreihe. (Es spielt für die Konvergenzbetrachtung keine Rolle, dass $|a_n| \leq b_n$ erst ab einem gewissen Index gilt.) $\square$

Definition 5.2.9. Wir definieren die (komplexe) *Exponentialfunktion* $\exp: \mathbb{C} \rightarrow \mathbb{C}$ durch

$$e^x := \exp(x) := \sum_{n=0}^{\infty} \frac{x^n}{n!}.$$

Da für reelle x der Wert e^x auch reell ist, haben wir damit auch die reelle Exponentialfunktion $\exp: \mathbb{R} \rightarrow \mathbb{R}$ definiert.

Bemerkungen 5.2.10. (i) Es gilt $e^0 = 1 + 0 + 0 + \dots = 1$.

(ii) Der Wert $e := e^1$, den die Exponentialfunktion bei 1 annimmt heißt *Eulersche Zahl*, näherungsweise erhalten wir $e \approx \frac{1^0}{0!} + \frac{1^1}{1!} + \frac{1^2}{2!} + \frac{1^3}{3!} = 1 + \frac{1}{2} + \frac{1}{6} \approx 2,7$. Tatsächlich handelt es sich bei e um eine irrationale Zahl, eine genauere Annäherung ist

$$e \approx 2,718281828459045235360287471352662497757 \dots$$

(iii) Es gilt die *Funktionalgleichung der Exponentialfunktion*

$$e^{x+y} = e^x \cdot e^y$$

für alle komplexen Zahlen x und y .

(iv) Es gilt $e^x > 0$ für alle $x \in \mathbb{R}$. Dies ist aus der Definition der Exponentialreihe klar für $x \geq 0$ (für nichtnegative x sind alle Reihenglieder nichtnegativ) und folgt aus der Funktionalgleichung wegen $1 = e^{x-x} = e^x \cdot e^{-x}$. Außerdem ist aus der Definition der Exponentialreihe sofort ersichtlich, dass $e^x > 1$ gilt für $x > 0$. Damit folgt die strenge Monotonie der reellen Exponentialfunktion, denn es gilt $e^y = e^x \cdot e^{y-x} > e^x$ falls $y > x$.

(v) Für $z \in \mathbb{C}$ mit $\operatorname{Re}(z) = x$, $\operatorname{Im}(z) = y$ gilt die *Eulersche Formel*:

$$e^z = e^{x+iy} = e^x \cdot (\cos(y) + i \sin(y)).$$

Insbesondere gilt z.B. $e^{\pi i} = -1$.

5.3. Funktionsgrenzwerte und Stetigkeit. Stetigkeit ist ein Begriff der eng mit dem der Konvergenz verwandt ist. Stellt man eine reelle Funktion auf einem zusammenhängenden Definitionsintervall durch ihren Graphen dar, so bedeutet Stetigkeit anschaulich gesprochen, dass man den Graph ohne abzusetzen in einem Zug durchzeichnen kann.

- Beispiele 5.3.1.**
- (i) Die Funktion $x \mapsto x^2, \mathbb{R} \rightarrow \mathbb{R}$ ist stetig.
 - (ii) Die Funktion $x \mapsto |x|, \mathbb{R} \rightarrow \mathbb{R}$ ist stetig.
 - (iii) Die Funktion $x \mapsto \begin{cases} 0, & \text{falls } x \leq 0; \\ 1, & \text{falls } x > 0. \end{cases}, \mathbb{R} \rightarrow \mathbb{R}$ ist *unstetig*.

Für die formale Definition von Stetigkeit einer Abbildung definiert man zunächst die *Stetigkeit in einem Punkt*.

Definition 5.3.2 (ε - δ -Definition der Stetigkeit). Sei $D \subseteq \mathbb{R}^n$ der Definitionsbereich einer Abbildung $f: D \rightarrow \mathbb{R}^m$ und sei $x \in D$. Dann heißt f *stetig in x* , falls es zu jedem $\varepsilon > 0$ ein $\delta > 0$ gibt, so dass

$$\|x - x'\| < \delta \Rightarrow \|f(x) - f(x')\| < \varepsilon$$

für alle $x' \in D$ gilt. Die Abbildung f heißt *stetig*, falls sie in jedem Punkt $x \in D$ stetig ist.

Etwas anschaulicher ist folgende, äquivalente, Definition.

Definition 5.3.3 (Alternative Definition der Stetigkeit). Sei $D \subseteq \mathbb{R}^n$ der Definitionsbereich einer Abbildung $f: D \rightarrow \mathbb{R}^m$ und sei $x \in D$. Dann heißt f *stetig in x* , falls es zu jedem ε -Ball $B_\varepsilon(f(x))$ einen δ -Ball $B_\delta(x)$ um x gibt, so dass

$$f(B_\delta(x) \cap D) \subseteq B_\varepsilon(f(x))$$

gilt. Die Abbildung f heißt *stetig*, falls sie in jedem Punkt $x \in D$ stetig ist.

Als ein Beispiel für die obigen Definition wollen wir einmal explizit die Stetigkeit der reellen Funktion $x \mapsto x^2$ im Punkt $x = 0$ nachweisen.

Beispiel 5.3.4. Sei die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$ wir wollen die Stetigkeit von f im Punkt $x = 0$ nachweisen. Sei $\varepsilon > 0$ vorgegeben. Wir wollen ein $\delta > 0$ finden, so dass $|x - x'| < \delta$ impliziert, dass

$$\varepsilon > |f(x) - f(x')| = |x^2 - x'^2| = |0 - x'^2| = x'^2.$$

Dies zeigt: wählt man den Abstand zwischen $x = 0$ und x' kleiner als $\sqrt{\varepsilon}$, d.h. wählt man x' betragsmäßig kleiner als $\sqrt{\varepsilon}$, so ist der Abstand zwischen $f(x) = 0$ und $f(x')$ kleiner als ε . Damit haben wir $\delta := \sqrt{\varepsilon}$ gefunden.

Dieses Beispiel soll nur dazu dienen, die ε - δ -Definition der Stetigkeit zu veranschaulichen. In der Praxis ist es oft nicht nötig, Stetigkeit nachzuprüfen, für viele Funktionen folgt Stetigkeit direkt und ohne Rechnung aus den folgenden Regeln.

Bemerkungen 5.3.5.

- (i) Konstante Funktionen, die identische Abbildung $\text{id}_{\mathbb{R}^n}: \mathbb{R}^n \rightarrow \mathbb{R}^n, x \mapsto x$, die Norm- und die Betragsfunktion $x \mapsto \|x\|$, bzw. $x \mapsto |x|$ sind stetig.
- (ii) Eine Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m, x \mapsto (f_1(x), \dots, f_m(x))$ ist genau dann stetig, wenn die Komponentenfunktionen $f_1, \dots, f_m: \mathbb{R}^n \rightarrow \mathbb{R}$ alle stetig sind. (Dies folgt aus der Bemerkung 5.1.5 (v) und dem Satz 5.3.6 unten.)
- (iii) *Summen, Differenzen, Produkte und Quotienten stetiger Funktionen sind stetig*, d.h. sind $f, g: D \rightarrow \mathbb{R}$ stetig, so auch $x \mapsto f(x) + g(x)$, $x \mapsto f(x) - g(x)$, $x \mapsto f(x) \cdot g(x)$ und $x \mapsto \frac{f(x)}{g(x)}$, sofern $g(x) \neq 0$ für alle $x \in D$.
- (iv) Verkettungen stetiger Abbildungen sind stetig.
- (v) Die Exponentialfunktion $\exp$ ist stetig (auch die komplexe Exponentialfunktion, die man ja als Abbildung $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ auffassen kann).
- (vi) Aus (i) und (iii) folgt insbesondere, dass Polynome (auch Polynome in mehreren Variablen) und rationale Funktionen stetig sind.
- (vii) Trigonometrische Funktionen wie $\sin, \cos, \tan$ sind stetig.
- (viii) Aus (ii) und (vi) folgt, dass lineare Abbildungen $\mathbb{R}^n \rightarrow \mathbb{R}^m$ stetig sind.

Der Zusammenhang zwischen Stetigkeit und Konvergenz wird durch den folgenden Satz hergestellt.

Satz 5.3.6. *Eine Abbildung $f: D \rightarrow \mathbb{R}^m$ mit Definitionsbereich $D \subseteq \mathbb{R}^n$ ist stetig in $x \in D$ genau dann, wenn für jede konvergente Folge $a_k \rightarrow a$ mit $a_k, a \in D$ gilt, dass $f(a_k) \rightarrow f(a)$.*

Die in Satz 5.3.6 gegebene alternative Charakterisierung der Stetigkeit nennt man auch *Folgenstetigkeit*. Kurz gefasst kann man dieses Kriterium auch so ausdrücken: Eine Abbildung f ist stetig genau dann, wenn sie jede gegen ein Element a ihres

Definitionsbereichs konvergente Folge auf eine gegen $f(a)$ konvergente Folge abbildet.

BEWEIS VON SATZ 5.3.6. Angenommen f sei stetig im Sinne der Definition 5.3.2 und $a_k \rightarrow a$ sei eine konvergente Folge, deren Glieder und deren Limes in D liegen. Wir müssen zeigen, dass $f(a_k) \rightarrow f(a)$ gilt. Sei $\varepsilon > 0$ vorgegeben. Wegen der Stetigkeit von f gibt es ein $\delta > 0$, so dass der δ -Ball um a in den ε -Ball um $f(a)$ abgebildet wird. Wegen der Konvergenz von a_k gegen a liegen fast alle Folgenglieder im δ -Ball um a , also liegen auch fast alle $f(a_k)$ im ε -Ball um $f(a)$.

Umgekehrt, angenommen f sei *nicht* stetig im Sinne der Definition 5.3.2. Dann gibt es ein $\varepsilon > 0$, so dass es zu jedem $\delta > 0$ ein Element aus $B_\delta(a)$ gibt, dessen Bild unter f nicht innerhalb von $B_\varepsilon(f(a))$ liegt. Insbesondere gibt es ein solches Element, wir bezeichnen es mit a_k , für jedes $\delta = \frac{1}{k}$ für $k \in \mathbb{N}$. Somit haben wir eine Folge $a_k \rightarrow a$ gefunden, so dass $f(a_k)$ nicht gegen $f(a)$ konvergiert. $\square$

Den obigen Satz kann man auch so ausdrücken: Eine Abbildung ist stetig genau dann, wenn man Limesbildung und Funktionsauswertung vertauschen kann, wenn also

$$\lim_{k \rightarrow \infty} f(a_k) = f(\lim_{k \rightarrow \infty} a_k)$$

für alle konvergenten Folgen $a_k \rightarrow a$ mit $a_k, a \in D$ gilt.

Manchmal betrachtet man das Verhalten einer Funktion in der Nähe einer Definitionslücke, wie im folgenden Beispiel.

Beispiel 5.3.7. Wir betrachten die Funktion $f: \mathbb{R} \setminus \{-1, +1\} \rightarrow \mathbb{R}$,

$$f(x) = \frac{(x-1)^2}{x^2-1},$$

Sie ist bei $x = -1$ und $x = +1$ nicht definiert. Betrachten wir nun das Verhalten der Funktion an ihren beiden Definitionslücken. Man kann die Funktion etwas anders aufschreiben:

$$f(x) = \frac{(x-1)^2}{(x-1)(x+1)} = \frac{x-1}{x+1}.$$

Man sieht, dass man die Funktion durch den Wert 0 auf den Punkt $x = 1$ *stetig fortsetzen* kann.

Am Punkt $x = -1$ ist dies andererseits nicht möglich. Für $x > -1$, die nahe bei -1 liegen erhält man negative Werte, die betragsmäßig umso größer werden, je weiter sich das x der Zahl -1 annähert. Für $x < -1$ erhält man positive Werte, die umso größer werden, je weiter sich das x der Zahl -1 annähert.

Um das Verhalten einer Funktion oder Abbildung auch in der Nähe eines Punktes, auch einer Definitionslücke, beschreiben zu können, verwendet man den folgenden Begriff.

Definition 5.3.8 (Funktionsgrenzwerte). Wir schreiben

$$\lim_{x \rightarrow a} f(x) = b,$$

falls für *jede* Folge $a_k \rightarrow a$, deren Folgenglieder im Definitionsbereich von f liegen, $f(a_k) \rightarrow b$ gilt.

Mit dieser Definition kann man sagen, dass eine Abbildung bei a stetig ist, genau dann, wenn $\lim_{x \rightarrow a} f(x) = f(a)$ gilt.

Bemerkungen 5.3.9. (i) Die Definition gilt allgemein für eine Abbildung $f: D \rightarrow \mathbb{R}^m$ mit $D \subseteq \mathbb{R}^n$. Für $n = 1$, also Abbildungen in einer Variablen, sind auch $a = \infty$ und $a = -\infty$, für $m = 1$ auch $b = \infty$ und $b = -\infty$ zugelassen.

(ii) Für $n = 1$, also Abbildungen in einer Variablen gibt es folgende Varianten der Definition. Man schreibt

$$\lim_{x \xrightarrow{>} a} f(x) = b,$$

wenn für alle Folgen $a_k \rightarrow a$, deren Folgenglieder größer als a sind, $f(a_k) \rightarrow b$ gilt, analog schreibt man

$$\lim_{x \xrightarrow{<} a} f(x) = b,$$

wenn für alle Folgen $a_k \rightarrow a$, deren Folgenglieder kleiner als a sind, $f(a_k) \rightarrow b$ gilt. Dies wird *rechts-* bzw. *linksseitiger Limes* genannt.

(iii) Also gilt für die Funktion in Beispiel 5.3.7:

$$\begin{aligned} \lim_{x \rightarrow 1} f(x) &= 0, & \lim_{x \xrightarrow{>} -1} f(x) &= -\infty, & \lim_{x \xrightarrow{<} -1} f(x) &= +\infty, \\ \lim_{x \rightarrow \infty} f(x) &= 1, & \lim_{x \rightarrow -\infty} f(x) &= 1. \end{aligned}$$

- Beispiele 5.3.10.** (i) Es gilt $\lim_{x \rightarrow 0} x^2 = 0$, denn die reelle Funktion $x \mapsto x^2$ ist stetig bei $x = 0$.
- (ii) Es gilt für die reelle Exponentialfunktion $\lim_{x \rightarrow \infty} e^x = \infty$, da aus der Definition der Exponentialreihe sofort $e^x \geq 1 + x$ für positive x folgt. Wegen $e^{-x} = \frac{1}{e^x}$ folgt damit $\lim_{x \rightarrow -\infty} e^x = 0$.
- (iii) Die Werte

$$\lim_{x \rightarrow \infty} \sin(x) \quad \text{und} \quad \lim_{x \rightarrow 0} \sin\left(\frac{1}{x}\right)$$

sind *nicht* definiert. Die Funktion $x \mapsto \sin\left(\frac{1}{x}\right)$, $\mathbb{R} \setminus \{0\}$ ist eine stetige Funktion, die aber nicht stetig auf $x = 0$ fortgesetzt werden kann.

- (iv) Wir haben bereits in Bemerkung 5.3.5 (ii) gesehen, dass eine Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^m$ genau dann stetig ist, wenn alle ihre Komponentenfunktionen stetig sind. Somit genügt es für Stetigkeitsfragen, den Fall $m = 1$, also Funktionen in eventuell mehreren Variablen zu betrachten. Dies kann man aber nicht einfach auf den Fall $n = 1$, also Funktionen in einer Variablen zurückführen, wie das folgende instruktive Beispiel zeigt: Sei die Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ definiert durch

$$f(x, y) = \begin{cases} 0, & \text{für } (x, y) = (0, 0); \\ \frac{xy}{x^2 + y^2}, & \text{sonst.} \end{cases}$$

Diese Funktion ist im Nullpunkt unstetig, dies kann man aber nicht herausfinden, indem man etwa nur die Funktionen $x \mapsto f(x, y)$ und $y \mapsto f(x, y)$ betrachtet, also f auf achsenparallele Geraden einschränkt, denn auf den Koordinatenachsen ist die Funktion konstant gleich Null, und außerhalb des Nullpunkts ist sie stetig.

Man kann die Unstetigkeit bei $(x, y) = (0, 0)$ aber nachweisen, indem man eine Folge betrachtet, die nicht etwa entlang der Achse sich auf den Nullpunkt zubewegt, sondern etwa auf einer Winkelhalbierenden zwischen den Achsen. Sei etwa $a_k = \left(\frac{1}{k}, \frac{1}{k}\right)$, dann gilt

$$f\left(\frac{1}{k}, \frac{1}{k}\right) = \frac{\frac{1}{k^2}}{2 \frac{1}{k^2}} = \frac{1}{2},$$

die Folge $f(a_k)$ ist also konstant gleich $\frac{1}{2}$ und geht somit auch gegen $\frac{1}{2}$, während $f(0, 0) = 0$ ist.

Ein wichtiger Satz ist der folgende.

Satz 5.3.11 (Zwischenwertsatz). *Sei f eine auf einem Intervall definierte stetige Funktion. Hat f zwei verschiedene reelle Zahlen als Werte, so nimmt f auch alle Zahlen dazwischen als Werte an.*

BEWEISSKIZZE. Angenommen, es gilt $f(a) < c < f(b)$. Dann müssen wir zeigen, dass es eine Zahl x im Definitionsintervall von f gibt mit $f(x) = c$. Durch Abziehen der Konstante c von der Funktion f kann man den Beweis auf den Fall zurückführen, dass $f(a) < 0$ und $f(b) > 0$ ist. Dann ist zu zeigen, dass es eine Nullstelle x von f gibt. Dies beweist man, indem man das Intervall $[a, b]$ fortwährend in Teilintervalle halber Länge aufteilt. Ist $f(\frac{a+b}{2})$ negativ, betrachten wir das Intervall $[\frac{a+b}{2}, b]$, ist es positiv, das Intervall $[a, \frac{a+b}{2}]$ (falls es eine Nullstelle ist, sind wir fertig). Auf diese Weise erhält man eine Intervallschachtelung, die sich auf eine Nullstelle von f zusammenzieht. $\square$

Diesen Satz kann man dafür benutzen, um die Existenz gewisser Umkehrfunktionen zu zeigen.

Satz 5.3.12. *Sei $f: I \rightarrow \mathbb{R}$ eine auf einem Intervall definierte stetige und streng monoton wachsende Funktion und sei $J := f(I)$ das Bild von f . Dann gibt es eine stetige monoton wachsende Umkehrfunktion $f^{-1}: J \rightarrow \mathbb{R}$ von f .*

BEWEIS. Wegen der strengen Monotonie ist f injektiv. Schränken wir die Zielmenge auf J ein, betrachten wir also $f: I \rightarrow J$, dann ist f auch surjektiv und besitzt eine Umkehrfunktion $f^{-1}: J \rightarrow \mathbb{R}$. Es folgt sofort, dass auch f^{-1} streng monoton wachsend ist. Auf den Beweis der Stetigkeit der Umkehrfunktion verzichten wir hier. $\square$

Als Beispiel für die Anwendung des obigen Satzes zeigen wir die Stetigkeit und Monotonie der Umkehrfunktion von $\exp: (0, \infty)$. In Bemerkung 5.2.10 wurde gezeigt, dass $\exp$ streng monoton wachsend ist. Mit der Stetigkeit von $\exp$ folgt, dass die Umkehrfunktion ebenfalls streng monoton wachsend und stetig ist.

Definition 5.3.13. Die Funktion $\ln := \exp^{-1}: (0, \infty) \rightarrow \mathbb{R}$ wird *natürlicher Logarithmus* genannt.

Es gilt $\ln(1) = 0$ und $\ln(e) = 1$. Aus der Funktionalgleichung der Exponentialfunktion folgt, dass

$$\ln(a \cdot b) = \ln(a) + \ln(b)$$

für alle positiven reellen Zahlen a, b gilt. Aus den Eigenschaften von $\exp$ folgt

$$\lim_{x \rightarrow \infty} \ln(x) = \infty, \quad \text{und} \quad \lim_{x \rightarrow 0} \ln(x) = -\infty.$$

Mit Hilfe von $\ln$ kann man Potenzen von positiven reellen Zahlen für beliebige Exponenten definieren:

Definition 5.3.14. Sei a eine positive reelle Zahl und $x \in \mathbb{R}$. Wir definieren

$$a^x := \exp(x \ln a) = e^{x \ln a}.$$

Wegen $\ln(e) = 1$ ist dies mit der Schreibweise $\exp(x) = e^x$ vereinbar. Aus der Funktionalgleichung der Exponentialfunktion folgt, dass für $x \in \mathbb{Z}$ diese Definition mit der früheren übereinstimmt. Es gelten die üblichen Rechenregeln für Potenzen

$$a^{x+y} = a^x \cdot a^y, \quad a^{-x} = \frac{1}{a^x}, \quad \ln(a^x) = x \ln(a), \quad (a^x)^y = a^{xy}.$$

Außerdem definieren wir für $n \in \mathbb{N}$ die n -te Wurzel durch

$$\sqrt[n]{a} := a^{\frac{1}{n}}.$$

Wir bemerken, dass man die n -te Wurzel auch mit Hilfe von Satz 5.3.12 einführen könnte, indem man die Monotonie (und Stetigkeit) der Funktion $x \mapsto x^n$ zeigt.

Definition 5.3.15. Wir definieren für eine reelle Zahl $b > 0$ und $x \in \mathbb{R}$ den *Logarithmus zur Basis b* durch

$$\log_b(x) = \frac{\ln(x)}{\ln(b)}.$$

Der Logarithmus zur Basis b von x ist die eindeutig bestimmte reelle Zahl a , so dass $b^a = x$.

Beispiel 5.3.16. Auf ein Konto werden zum Anfang jedes Jahres 1.000€ eingezahlt und am Ende des Jahres mit 3% verzinst. Nach wie vielen Jahren übersteigt das Guthaben zum ersten Mal den Wert 25.000€?

Wir erinnern uns an die Formel

$$K_n = r \cdot q \cdot \frac{1 - q^n}{1 - q}$$

für die Kapitalentwicklung bei vorschüssiger Zahlweise. Wir stellen um und erhalten

$$q^n = 1 - K_n \cdot \frac{1 - q}{rq}.$$

Anwenden des Logarithmus zur Basis q auf beiden Seiten liefert die Formel

$$n = \log_q \left(1 - K_n \cdot \frac{1 - q}{rq} \right).$$

In unserem Beispiel erhalten wir

$$n = \ln \left(1 + 25.000 \cdot \frac{0,03}{1.030} \right) : \ln(1,03) \approx 18,5.$$

Aus der Monotonie des Logarithmus folgt, dass nach 19 Jahren erstmals der Betrag von 25.000€ überschritten wird.

Beispiel 5.3.17 (“Null hoch Null”). Was ist 0^0 ? Wir betrachten die Funktion

$$(x, y) \mapsto x^y, \quad (0, \infty) \times \mathbb{R} \rightarrow \mathbb{R}.$$

Für $y = 0$ gilt $x^y = \exp(y \ln(x)) = \exp(0 \ln(x)) = \exp(0) = 1$, nach unserer Definition gilt also $x^0 = 1$, egal, welche reelle Zahl x ist. Andererseits gilt für $y > 0$:

$$\lim_{x \searrow 0} x^y = \lim_{x \searrow 0} \exp(y \ln(x)) = 0,$$

denn $\ln(x) \rightarrow -\infty$ für $x \rightarrow 0$. Die obige Funktion kann also nicht stetig auf den Punkt $(0, 0)$ fortgesetzt werden. Wir bleiben weiterhin dabei, dass $0^0 = 1$ gilt; ansonsten würde z.B. der binomische Lehrsatz nicht richtig sein, denn es soll gelten $(x + 0)^2 = x^2 \cdot 0^0 + 2 \cdot x \cdot 0^1 + x^0 \cdot 0^2 = x^2$.

6. Differentiation

6.1. Ableitung einer Funktion in einer Variablen. Bei der Differentiation geht es darum, die Änderung einer Größe zu beschreiben. Für eine reelle Funktion f in einer reellen Variablen kann man anschaulich für die Änderung des Funktionswertes zwischen zwei reellen Zahlen x und $x + h$ die Steigung der Sekante an den Graphen durch die Punkte $(x, f(x))$ und $(x + h, f(x + h))$, also den Wert

$$\frac{f(x + h) - f(x)}{h}$$

betrachten. Will man die Steigung in der Nähe eines Punktes betrachten, so muss man den Abstand h der beiden Zahlen und x und $x + h$ klein wählen. Dies führt auf die folgende Definition.

Definition 6.1.1. Eine reelle Funktion $f: D \rightarrow \mathbb{R}$, definiert auf einem Intervall $D \subseteq \mathbb{R}$, heißt *differenzierbar* bei $x \in D$, falls

$$\lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

existiert. In diesem Fall wird der Wert des obigen Limes mit $f'(x)$ bezeichnet. Die Funktion heißt *differenzierbar*, falls sie bei jedem $x \in D$ differenzierbar ist. In diesem Fall nennt man die Funktion $x \mapsto f'(x)$ die *Ableitung* von f .

Anschaulich ist $f'(x)$ die Steigung der Tangente an den Graphen von f im Punkt $(x, f(x))$.

Beispiele 6.1.2. (i) Konstante Funktionen $f(x) = c$ haben Ableitung Null, denn

$$\frac{f(x+h) - f(x)}{h} = \frac{c - c}{h} = 0.$$

(ii) Wir zeigen, dass die Funktion $f: x \mapsto x^2, \mathbb{R} \rightarrow \mathbb{R}$, bei jedem $x \in \mathbb{R}$ differenzierbar ist und ermitteln die Ableitung. Dazu berechnen wir die Sekantensteigung

$$\frac{f(x+h) - f(x)}{h} = \frac{x^2 + 2xh + h^2 - x^2}{h} = 2x + h.$$

Daraus folgt

$$\lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = 2x.$$

Die Funktion ist also differenzierbar und für die Ableitung gilt $f'(x) = 2x$.

(iii) Allgemeiner erhalten wir für die Funktion $x \mapsto x^k$, mit $k \in \mathbb{N}$ unter Verwendung des binomischen Lehrsatzes als Sekantensteigung

$$\frac{x^k + \binom{k}{1}x^{k-1}h + \binom{k}{2}x^{k-2}h^2 + \dots + \binom{k}{k}x^0h^k - x^k}{h}.$$

Dies geht für $h \rightarrow 0$ gegen $\binom{k}{1}x^{k-1} = kx^{k-1}$, was die Differenzierbarkeit zeigt. Für die Ableitung gilt also $f'(x) = kx^{k-1}$.

(iv) Für die Betragsfunktion $x \mapsto |x|, \mathbb{R} \rightarrow \mathbb{R}$, gilt bei $x = 0$:

$$\frac{f(x+h) - f(x)}{h} = \frac{|h|}{h} = \begin{cases} +1, & \text{falls } h \geq 0; \\ -1, & \text{falls } h < 0. \end{cases}$$

Die Betragsfunktion ist somit bei $x = 0$ nicht differenzierbar. (An allen anderen Stellen ist sie differenzierbar.)

Anschaulich kann man Differenzierbarkeit als Glattheit des Graphen verstehen, bei einer differenzierbaren Funktion hat der Graph keine “Knickstellen”.

Satz 6.1.3 (Rechenregeln für Ableitungen). *Sind die beiden auf dem Intervall I definierten Funktionen $f, g: I \rightarrow \mathbb{R}$ bei $x \in I$ differenzierbar und ist $c \in \mathbb{R}$ eine Konstante, sind auch die Funktionen $c \cdot f$, $f \pm g$, $f \cdot g$ bei x differenzierbar und es gilt*

$$\begin{aligned}(c \cdot f)'(x) &= c \cdot f'(x), \\ (g \pm f)'(x) &= f'(x) \pm g'(x), \\ (f \cdot g)'(x) &= f'(x)g(x) + f(x)g'(x).\end{aligned}$$

Hat außerdem g keine Nullstelle in I , dann ist f/g bei x differenzierbar und es gilt

$$\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}.$$

Für die Ableitung einer Verkettung von differenzierbaren Funktionen gilt die *Kettenregel*.

Satz 6.1.4 (Kettenregel, eindimensional). *Seien I, J zwei Intervalle und seien $g: I \rightarrow J$ und $f: J \rightarrow \mathbb{R}$ differenzierbare Funktionen. Dann ist auch die Verkettung $f \circ g: I \rightarrow \mathbb{R}$ differenzierbar und es gilt*

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

Satz 6.1.5 (Ableitung der Umkehrfunktion). *Seien I, J zwei Intervalle, $f: I \rightarrow J$ eine differenzierbare Funktion, deren Ableitung keine Nullstellen hat (deren Graph somit keine waagerechten Tangenten hat) und $f^{-1}: J \rightarrow I$ ihre Umkehrfunktion. Dann ist f^{-1} differenzierbar und es gilt für ihre Ableitung*

$$(f^{-1})'(y) = \frac{1}{f'(f^{-1}(y))}.$$

Wir wollen den Satz nicht beweisen, bemerken aber, dass die Formel für die Ableitung der Umkehrfunktion aus der Kettenregel folgt, wenn man die Differenzierbarkeit von f^{-1} als gegeben annimmt; denn es gilt $f(f^{-1}(y)) = y$, somit gilt nach Kettenregel $f'(f^{-1}(y)) \cdot (f^{-1})'(y) = 1$.

Ohne Beweis erinnern wir außerdem daran, dass $\exp'(x) = \exp(x)$, $\sin'(x) = \cos(x)$, $\cos'(x) = -\sin(x)$ gelten.

6.2. Die Regel von L'Hôpital. Die Regel von L'Hôpital erlaubt es oft, Funktionsgrenzwerte der Form $\lim_{x \rightarrow c} \frac{f(x)}{g(x)}$ auszuwerten, auch wenn beide Funktionen $f(x)$ und $g(x)$ gegen Null oder beide gegen $\pm\infty$ gehen.

Satz 6.2.1 (L'Hôpital'sche Regel). Sei $c \in \mathbb{R} \cup \{+\infty, -\infty\}$. Seien die Funktionen f und g differenzierbar auf einem offenen Intervall mit Endpunkt c . Sei $\lim_{x \rightarrow c} \frac{f'(x)}{g'(x)} = L \in \mathbb{R} \cup \{+\infty, -\infty\}$ (d.h. es liegt Konvergenz oder bestimmte Divergenz vor). Falls entweder

$$\lim_{x \rightarrow c} f(x) = \lim_{x \rightarrow c} g(x) = 0 \quad \text{oder} \quad \lim_{x \rightarrow c} |f(x)| = \lim_{x \rightarrow c} |g(x)| = \infty$$

gilt, dann folgt

$$\lim_{x \rightarrow c} \frac{f(x)}{g(x)} = L.$$

Beispiele 6.2.2. (i) Betrachte

$$\lim_{x \rightarrow 1} \frac{\ln(x)}{x - 1}.$$

Nenner und Zähler gehen gegen Null und es gilt $\lim_{x \rightarrow 1} \ln'(x) = \lim_{x \rightarrow 1} \frac{1}{x} = 1$ sowie $\lim_{x \rightarrow 1} (x - 1)' = 1$. Es folgt

$$\lim_{x \rightarrow 1} \frac{\ln(x)}{x - 1} = \frac{1}{1} = 1.$$

(ii)

$$\lim_{x \rightarrow \infty} \frac{\sqrt{x}}{\ln(x)} = \lim_{x \rightarrow \infty} \frac{\frac{1}{2\sqrt{x}}}{\frac{1}{x}} = \lim_{x \rightarrow \infty} \frac{1}{2} \sqrt{x} = \infty.$$

(iii) Man kann die Regel auch mehrmals anwenden:

$$\lim_{x \rightarrow \infty} \frac{x^3}{e^x} = \lim_{x \rightarrow \infty} \frac{3x^2}{e^x} = \lim_{x \rightarrow \infty} \frac{6x}{e^x} = \lim_{x \rightarrow \infty} \frac{6}{e^x} = 0.$$

(iv) Man kann die Regel manchmal auch anwenden, wenn die zu untersuchende Funktion kein Quotient ist. Wir wollen die Existenz des folgenden Limes untersuchen und seinen Wert bestimmen.

$$\lim_{x \rightarrow 0} x^x = \lim_{x \rightarrow 0} \exp(x \ln(x)).$$

Wegen der Stetigkeit von $\exp$ genügt es, die Funktion $x \mapsto x \ln(x)$ zu untersuchen.

$$\lim_{x \rightarrow 0} x \ln(x) = \lim_{x \rightarrow 0} \frac{\ln(x)}{\frac{1}{x}} = \lim_{x \rightarrow 0} \frac{\frac{1}{x}}{-\frac{1}{x^2}} = \lim_{x \rightarrow 0} -x = 0.$$

Somit gilt

$$\lim_{x \rightarrow 0} x^x = \exp\left(\lim_{x \rightarrow 0} x \ln(x)\right) = \exp(0) = 1.$$

6.3. Partielle Ableitungen. Nun wollen wir auch Funktionen in mehreren Variablen differenzieren. Im folgenden betrachten wir stets Abbildungen $f: D \rightarrow \mathbb{R}^m$, wobei der Definitionsbereich $D \subseteq \mathbb{R}^n$ folgende Eigenschaft haben soll.

Wir nennen eine Teilmenge $D \subseteq \mathbb{R}^n$ *offen*, wenn es um jeden Punkt $x \in D$ einen ε -Ball $B_\varepsilon(x)$ mit positivem ε gibt, der ganz in D enthalten ist. Eine offene Menge ist also eine, die keine Randpunkte enthält. Ein *Randpunkt* einer Teilmenge $D \subseteq \mathbb{R}^n$ ist ein Punkt, der die Eigenschaft hat, dass in jeder ε -Umgebung von ihm sowohl Punkte von D als auch von $\mathbb{R}^n \setminus D$ liegen. Ein Randpunkt einer Menge D muss nicht notwendig zu D dazugehören, so sind z.B. die Zahlen 1 und 2 Randpunkte des halboffenen Intervalls $(1, 2]$. Eine Teilmenge des $\mathbb{R}^n$ ist genau dann offen, wenn sie keinen ihrer Randpunkte enthält.

Beispiele für offene Mengen sind offene Intervalle, oder auch der offene ε -Ball um einen Punkt. Der ganze $\mathbb{R}^n$ ist ebenfalls offen. Nicht offen sind z.B. das Intervall $[0, \infty)$, da der linke Randpunkt dazugehört oder der abgeschlossene Einheitsball um die Null im $\mathbb{R}^n$, $\{x \in \mathbb{R}^n \mid \|x\| \leq 1\}$. Als einfache Merkregel kann man sagen: Teilmengen des $\mathbb{R}^n$ sind offen, wenn sie durch echte Ungleichungen, also durch Zeichen wie “ $<$ ” definiert sind, nicht durch “ $\leq$ ” etc.

Im folgenden wollen wir annehmen, dass $f: D \rightarrow \mathbb{R}$ eine Funktion ist, wobei $D \subseteq \mathbb{R}^n$ eine offene Teilmenge ist. Dann heißt f bei $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ nach der Variablen x_i *partiell differenzierbar*, falls die Funktion

$$t \mapsto f(x_1, \dots, x_{i-1}, t, x_{i+1}, \dots, x_n)$$

(die wegen der Offenheit von D auf einem Intervall um die Null definiert ist) bei $t = x_i$ differenzierbar ist und der Wert der Ableitung heißt in diesem Fall die *partielle Ableitung* von f bei $x = (x_1, \dots, x_n)$ nach x_i . Für diese partielle Ableitung schreiben wir

$$\frac{\partial f}{\partial x_i}(x) \quad \text{oder} \quad \frac{\partial}{\partial x_i} f(x).$$

Ist f bei jedem x aus dem Definitionsbereich D partiell nach allen Variablen $x_1, \dots, x_n$ partiell differenzierbar, so nennen wir f *partiell differenzierbar*. Sind die partiellen Ableitungen $\frac{\partial f}{\partial x_1}(x), \dots, \frac{\partial f}{\partial x_n}(x)$ stetig von x abhängig, so nennen wir f *stetig partiell differenzierbar*.

Beispiel 6.3.1. Wir betrachten die Funktion $f: \mathbb{R}^3 \rightarrow \mathbb{R}$, $f(x, y, z) = e^x \cdot x \cdot y + z^3$. Die Funktion ist stetig partiell differenzierbar und es gilt:

$$\begin{aligned}\frac{\partial f}{\partial x}(x, y, z) &= e^x \cdot x \cdot y + e^x \cdot 1 \cdot y = (x+1)e^x y, \\ \frac{\partial f}{\partial y}(x, y, z) &= e^x \cdot x, \\ \frac{\partial f}{\partial z}(x, y, z) &= 3z^2.\end{aligned}$$

Sind die partiellen Ableitungen von f selbst differenzierbar, dann kann man *partielle Ableitungen von höherer Ordnung* bilden, indem weiter partiell differenziert, man bildet dann mehrfache partielle Ableitungen, etwa k -ter Ordnung, die man so notiert:

$$\frac{\partial}{\partial x_{i_1}} \frac{\partial}{\partial x_{i_2}} \dots \frac{\partial}{\partial x_{i_k}} f(x) = \frac{\partial^k}{\partial x_{i_1} \partial x_{i_2} \dots \partial x_{i_k}} f(x).$$

Existieren alle partielle Ableitungen k -ter Ordnung, so nennt man die Funktion *k -mal partiell differenzierbar*, sind sie stetig, so sagt man, die Funktion ist *k -mal stetig partiell differenzierbar*. Gilt eine dieser beiden Bedingungen für beliebige $k \in \mathbb{N}$, so sagt man, f sei *beliebig oft (stetig) partiell differenzierbar* oder auch *unendlich oft (stetig) partiell differenzierbar*.

Mehrfache partielle Ableitungen nach derselben Variablen, z.B. k -mal nach x_i abgeleitet, kann man auch so schreiben:

$$\frac{\partial^k f}{\partial x_i^k}(x).$$

Beispiel 6.3.2. Wir betrachten die Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, $f(x, y) = x^3 y^2$. Die Funktion ist beliebig oft stetig partiell differenzierbar und es gilt für die partiellen Ableitungen erster Ordnung:

$$\begin{aligned}\frac{\partial f}{\partial x}(x, y) &= 3x^2 y^2, \\ \frac{\partial f}{\partial y}(x, y) &= 2x^3 y.\end{aligned}$$

Für die partiellen Ableitungen zweiter Ordnung gilt:

$$\begin{aligned}\frac{\partial^2 f}{\partial x^2}(x, y) &= 6x y^2, & \frac{\partial^2 f}{\partial x \partial y}(x, y) &= 6x^2 y, \\ \frac{\partial^2 f}{\partial y \partial x}(x, y) &= 6x^2 y, & \frac{\partial^2 f}{\partial y^2}(x, y) &= 2x^3,\end{aligned}$$

Wir machen hier die Beobachtung, dass es bei den beiden “gemischten” Ableitungen nicht darauf ankommt, ob zuerst nach x und dann nach y abgeleitet wird. In der Tat gilt der folgende Satz.

Satz 6.3.3. *Ist $D \subseteq \mathbb{R}^n$ eine offene Menge und $f: D \rightarrow \mathbb{R}$ zweimal stetig partiell differenzierbar, so gilt*

$$\frac{\partial^2 f}{\partial x_i \partial x_j} = \frac{\partial^2 f}{\partial x_j \partial x_i}.$$

Aus dem Satz folgt, dass für f beliebig oft stetig partiell differenzierbar auch bei höheren partiellen Ableitungen die Reihenfolge, in der die partiellen Differentiationen nach den einzelnen Variablen vorgenommen werden, keine Rolle spielt.

6.4. Totale Differenzierbarkeit und Jacobimatrix. Wir wollen nun definieren, was Differenzierbarkeit für Abbildungen $\mathbb{R}^n \rightarrow \mathbb{R}^m$ bedeutet – das ist nicht dasselbe wie partielle Differenzierbarkeit! Dazu stellen wir folgende Überlegung an:

Die Definition 6.1.1 der Differenzierbarkeit einer auf einem Intervall D definierten Funktion $f: D \rightarrow \mathbb{R}$ bei einem Punkt $x \in D$ ist äquivalent dazu, dass f sich um x durch eine affin-lineare Funktion $x \mapsto ax + t$ approximieren lässt, wobei der Fehler von höherer als erster Ordnung gegen Null geht, genauer: es gibt ein $a \in \mathbb{R}$, so dass

$$f(x + h) = f(x) + a \cdot h + r(h),$$

so dass für den Fehlerterm $r(h)$ gilt: $\frac{r(h)}{h} \rightarrow 0$ für $h \rightarrow 0$.

Die Ableitung bei x ist dann a und die approximierende affin-lineare Funktion ist $t \mapsto a(t - x) + f(x)$, ihr Graph ist die Tangente an den Graph von f bei x .

Dies kann man als alternative Definition der Differenzierbarkeit und der Ableitung einer Funktion benutzen. Sie erscheint zwar zunächst eventuell weniger naheliegend, hat aber den Vorteil, dass sie sich direkt auf den mehrdimensionalen Fall verallgemeinern lässt.

Definition 6.4.1. Sei $D \subseteq \mathbb{R}^n$ eine offene Teilmenge und $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine Abbildung. Wir sagen f ist bei $x \in D$ (total) differenzierbar, falls es eine reelle $m \times n$ -Matrix A gibt, so dass gilt

$$f(x + h) = f(x) + Ah + r(h),$$

wobei für den Fehlerterm gilt: $\frac{r(h)}{h} \rightarrow 0$ für $h \rightarrow 0$. In diesem Fall heißt die lineare Abbildung $x \mapsto Ax$, $\mathbb{R}^n \rightarrow \mathbb{R}^m$ das (*totale*) *Differential* von f bei x .

Unmittelbar aus dem Satz folgt, dass lineare Abbildungen $\mathbb{R}^n \rightarrow \mathbb{R}^m$ in jedem Punkt total differenzierbar sind und gleich ihrem Differential in jedem Punkt. Außerdem folgt sofort, dass jede bei x total differenzierbare Abbildung bei x auch stetig ist.

Um in der Praxis festzustellen, ob eine Abbildung total differenzierbar ist, benutzen wir meist den folgenden, sehr nützlichen, Satz. Bei einer Abbildung nach $\mathbb{R}^m$ bedeutet “partiell differenzierbar”, dass alle m Komponentenfunktionen partiell differenzierbar sind (die zugehörigen partiellen Ableitungen sind Vektoren im $\mathbb{R}^m$).

Satz 6.4.2. *Ist $f: D \rightarrow \mathbb{R}^m$, $D \subseteq \mathbb{R}^n$ offen, stetig partiell differenzierbar, dann ist f total differenzierbar (in jedem Punkt).*

Für eine auf einer offenen Menge D im $\mathbb{R}^n$ definierte Abbildung $f: D \rightarrow \mathbb{R}^m$, $f(x) = (f_1(x), \dots, f_m(x))$ bei der alle Komponentenfunktionen $f_1, \dots, f_m$ bei x partiell differenzierbar sind, definieren wir die *Jacobimatrix* $Df(x) = f'(x)$ wie folgt

$$Df(x) := f'(x) := \begin{pmatrix} \frac{\partial f_1}{\partial x_1}(x) & \dots & \frac{\partial f_1}{\partial x_n}(x) \\ \vdots & & \vdots \\ \frac{\partial f_m}{\partial x_1}(x) & \dots & \frac{\partial f_m}{\partial x_n}(x) \end{pmatrix}.$$

Mit anderen Worten, die Spalten von $Df(x)$ sind die partiellen Ableitungen von f . Es gilt:

Satz 6.4.3. *Ist $f: D \rightarrow \mathbb{R}^m$, $D \subseteq \mathbb{R}^n$ offen, bei $x \in D$ total differenzierbar, dann ist f bei x auch partiell differenzierbar und das Differential von f bei x ist durch die Jacobimatrix $Df(x)$ gegeben.*

Um einem konkreten Fall die totale Differenzierbarkeit zu testen, genügt es also, die partiellen Ableitungen (d.h. die Jacobimatrix) zu bestimmen und auf Stetigkeit zu untersuchen, wie im folgenden Beispiel. Das Differential ist dann durch die Jacobimatrix gegeben.

Beispiel 6.4.4. Sei $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$, $f(x, y, z) = (e^x yz, \cos(x) \sin(y))$ gegeben. Offensichtlich sind beide Komponenten partiell differenzierbar; wir berechnen die

Jacobimatrix

$$Df(x, y, z) = \begin{pmatrix} e^x y z & e^x z & e^x y \\ -\sin(x) \sin(y) & \cos(x) \cos(y) & 0 \end{pmatrix}.$$

Da alle Einträge stetig von (x, y, z) abhängen, ist die Funktion f also auch stetig partiell differenzierbar (sogar beliebig oft) und damit total differenzierbar. Das totale Differential bei (x, y, z) ist also durch die obige Jacobimatrix gegeben.

Satz 6.4.5 (Kettenregel, mehrdimensional). *Sind $f: D \rightarrow \mathbb{R}^m$ und $g: E \rightarrow \mathbb{R}^n$, wobei $D \subseteq \mathbb{R}^n$ und $E \subseteq \mathbb{R}^p$ offene Teilmengen sind total differenzierbar und gilt $g(E) \subseteq D$, dann ist die Verkettung $f \circ g: E \rightarrow \mathbb{R}^m$ total differenzierbar und für das Differential bei $x \in E$ gilt*

$$D(f \circ g)(x) = Df(g(x))Dg(x),$$

bzw., in anderer Schreibweise,

$$(f \circ g)'(x) = f'(g(x))g'(x),$$

wobei das Produkt auf der rechten Seite als Matrizenmultiplikation zu verstehen ist.

Man muss sich also keine neue Formel für die mehrdimensionale Kettenregel einprägen, die wohlbekannte Formel aus dem Eindimensionalen verallgemeinert sich direkt.

Beispiel 6.4.6. Wir betrachten die beiden Abbildungen

$$f: \mathbb{R}^2 \rightarrow \mathbb{R}^2, \quad f(x, y) = (x^2, x + 2y)$$

und

$$g: \mathbb{R}^3 \rightarrow \mathbb{R}^2, \quad g(x, y, z) = (x + y, z^2).$$

Ihre Verkettung $f \circ g$ ist die Abbildung

$$\begin{aligned} f \circ g: \mathbb{R}^3 &\rightarrow \mathbb{R}^2, \\ f \circ g(x, y, z) &= f(x + y, z^2) = \\ &= ((x + y)^2, x + y + 2z^2) = \\ &= (x^2 + 2xy + y^2, x + y + 2z^2). \end{aligned}$$

(Alle beteiligten Abbildungen sind offensichtlich stetig partiell differenzierbar und damit nach Satz 6.4.2 total differenzierbar.)

Wir berechnen zunächst die Jacobimatrix von $f \circ g$.

$$D(f \circ g)(x) = \begin{pmatrix} 2x + 2y & 2x + 2y & 0 \\ 1 & 1 & 4z \end{pmatrix}.$$

Nach der Kettenregel lässt sich diese auch als $Df(g(x, y, z))Dg(x, y, z)$ berechnen. Um das im Beispiel nachzurechnen, berechnen wir zunächst allgemein die Jacobimatrix von f :

$$Df(x, y) = \begin{pmatrix} 2x & 0 \\ 1 & 2 \end{pmatrix},$$

damit ergibt sich

$$Df(g(x, y, z)) = \begin{pmatrix} 2x + 2y & 0 \\ 1 & 2 \end{pmatrix}.$$

Für die Jacobimatrix von g erhalten wir

$$Dg(x, y, z) = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 2z \end{pmatrix}.$$

Es gilt in der Tat $D(f \circ g)(x) = Df(g(x))Dg(x)$, denn:

$$\begin{pmatrix} 2x + 2y & 0 \\ 1 & 2 \end{pmatrix} \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 2z \end{pmatrix} = \begin{pmatrix} 2x + 2y & 2x + 2y & 0 \\ 1 & 1 & 4z \end{pmatrix}.$$

KAPITEL 2

Optimierung

In diesem Kapitel werden wir uns damit beschäftigen, wie man Maximal- und Minimalstellen von Funktionen mehrerer Veränderlicher findet.

1. Lokale Extrema (ohne Nebenbedingungen)

Wir beginnen mit einem einfachen Beispiel im Eindimensionalen. Wir betrachten die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^3 - 3x$. Um lokale Extremstellen zu finden, betrachten wir die Ableitung $f'(x) = 3x^2 - 3$. Ihre Nullstelle liegen bei $x = -1$ und $x = +1$. Aus $f''(x) = 6x$ sehen wir, dass die zweite Ableitung von f bei $x = -1$ negativ ist, bei $x = +1$ positiv. (Mit anderen Worten, bei $x = -1$ wechselt das Vorzeichen der ersten Ableitung von positiv auf negativ, d.h. die Funktion ist links von $x = -1$ streng monoton steigend, rechts von $x = -1$ streng monoton fallend.) Wir schließen daraus, dass f bei $x = -1$ ein lokales Maximum hat. Analog kann man schließen, dass f bei $x = +1$ ein lokales Minimum hat.

Wir wollen ein ähnliches Kriterium für Funktionen in mehreren Variablen einführen.

Sei $f: D \rightarrow \mathbb{R}$ eine Funktion. Ein Punkt $x \in D$ heißt *lokale Maximalstelle* von f , wenn es eine ε -Umgebung von x gibt, in der keine Punkte von D liegen, auf denen f einen größeren Wert als $f(x)$ annimmt.

Sei $f: D \rightarrow \mathbb{R}$ eine Funktion. Ein Punkt $x \in D$ heißt *lokale Minimalstelle* von f , wenn es eine ε -Umgebung von x gibt, in der keine Punkte von D liegen, auf denen f einen kleineren Wert als $f(x)$ annimmt.

Man sagt, $x \in D$ ist *lokale Extremalstelle* von f , wenn x eine lokale Maximalstelle oder eine lokale Minimalstelle ist.

1.1. Der Gradient. Sei $f: D \rightarrow \mathbb{R}$ auf einer offenen Teilmenge $D \subseteq \mathbb{R}^n$ definiert und differenzierbar. Anstelle der Jacobimatrix (die hier eine $1 \times n$ -Matrix ist) ist es auch üblich, den sogenannten *Gradient* von f zu betrachten. Dabei handelt es sich lediglich um die zu $Df(x)$ transponierte Matrix, die man dann als Spaltenvektor im $\mathbb{R}^n$ auffasst. Man schreibt dafür $\nabla f(x)$ oder auch $\text{grad}f(x)$. Es gilt also:

$$\nabla f(x) = \text{grad}f(x) = \begin{pmatrix} \frac{\partial f}{\partial x_1}(x) \\ \vdots \\ \frac{\partial f}{\partial x_n}(x) \end{pmatrix}.$$

Der Gradient $\nabla f(x)$ ist ein Vektor im $\mathbb{R}^n$, der an jedem Punkt des Definitionsbereichs von f definiert ist, es handelt sich um ein sogenanntes *Vektorfeld*. Er zeigt in die Richtung des steilsten Anstiegs von f , bei lokalen Extrema verschwindet er. Punkte x , in denen $\nabla f(x) = 0$ gilt, nennt man *Flachpunkte* der Funktion.

Satz 1.1.1. Sei $f: D \rightarrow \mathbb{R}$ eine differenzierbare Funktion, die auf der offenen Teilmenge $D \subseteq \mathbb{R}^n$ definiert ist. Hat f in $x \in D$ eine lokale Extremalstelle, dann hat f in x einen Flachpunkt, d.h. es gilt $\nabla f(x) = 0$.

BEWEIS. Folgt aus dem bekannten Kriterium für Extremalstellen im Eindimensionalen: Hat nämlich f bei $x = (x_1, \dots, x_n)$ eine Extremalstelle, so haben auch die Funktionen $t \rightarrow f(x_1, \dots, x_{i-1}, t, x_{i+1}, \dots, x_n)$ bei $t = x_i$ eine Extremalstelle. $\square$

Beispiel 1.1.2. Sei $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ definiert durch $f(x, y) = x^2 + y^2$. Der Graph von f ist ein Rotationsparaboloid, ihr Minimum nimmt die Funktion bei $(x, y) = (0, 0)$ an. Der Gradient berechnet sich zu

$$\nabla f(x, y) = \begin{pmatrix} 2x \\ 2y \end{pmatrix}.$$

1.2. Die Hessematrix. Sei $f: D \rightarrow \mathbb{R}$, $D \subseteq \mathbb{R}^n$ offen, eine zweimal stetig partiell differenzierbare Funktion. Wir definieren die Hessematrix $\text{Hess}f(x)$ von f am Punkt x als die $n \times n$ -Matrix, bei der die partielle Ableitung zweiter Ordnung

$$\frac{\partial^2 f}{\partial x_i \partial x_j}$$

in der i -ten Zeile und j -ten Spalte steht:

$$\text{Hess}f(x) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1 \partial x_1}(x) & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n}(x) \\ \vdots & & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1}(x) & \cdots & \frac{\partial^2 f}{\partial x_n \partial x_n}(x) \end{pmatrix}.$$

Beispiel 1.2.1. Sei $f: \mathbb{R}^3 \rightarrow \mathbb{R}$ durch $f(x, y, z) = e^x \sin(y) z^3$ definiert. Die Hessematrix von f ist

$$\text{Hess}f(x) = \begin{pmatrix} e^x \sin(y) z^3 & e^x \cos(y) z^3 & 3e^x \sin(y) z^2 \\ e^x \cos(y) z^3 & -e^x \sin(y) z^3 & 3e^x \cos(y) z^2 \\ 3e^x \sin(y) z^2 & 3e^x \cos(y) z^2 & 6e^x \sin(y) z \end{pmatrix}.$$

1.3. Definitheit und Determinante. Sei eine reelle Matrix A gegeben, die *quadratisch* ist, d.h. gleich viele Zeilen wie Spalten hat. Man sagt, die Matrix A sei *symmetrisch*, wenn $A^t = A$ ist, dh., wenn für die Einträge a_{ij} der Matrix gilt: $a_{ij} = a_{ji}$.

Beispiel 1.3.1. Die Matrix A ist symmetrisch, die Matrix B nicht:

$$A = \begin{pmatrix} 1 & -2 & 3 \\ -2 & 4 & -5 \\ 3 & -5 & 6 \end{pmatrix}, \quad B = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}.$$

Nach Satz 6.3.3 ist die Hessematrix einer zweimal stetig partiell differenzierbaren Funktion stets symmetrisch.

Sei A eine reelle (symmetrische) $n \times n$ -Matrix und $x \in \mathbb{R}^n$ ein Vektor. Dann ist $x^t A x$ eine 1×1 -Matrix, also eine reelle Zahl. Z.B. sei

$$A = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 3 & 2 \\ 1 & 2 & 3 \end{pmatrix}, \quad x = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix},$$

dann ist

$$x^t A x = (1, 0, 1) \cdot \begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix} = 3 + 0 + 4 = 7.$$

In der folgenden Definition geht es um das Vorzeichen der Zahl $x^t A x$.

Definition 1.3.2. Ein symmetrische $n \times n$ -Matrix A heißt *positiv definit*, wenn für jeden von Null verschiedenen Vektor $x \in \mathbb{R}^n$ gilt, dass $x^t A x > 0$ ist. Sie heißt

negativ definit, wenn für jeden von Null verschiedenen Vektor $x \in \mathbb{R}^n$ gilt, dass $x^t A x < 0$ ist.

Bemerkung 1.3.3. Ein notwendiges (aber nicht hinreichendes) Kriterium für positive bzw. negative Definitheit ist, dass alle Diagonaleinträge a_{ii} der Matrix positiv bzw. negativ sind; denn es gilt $e_i^t A e_i = a_{ii}$ für die Standardbasisvektoren e_i .

Beispiele 1.3.4. (i) Die Einheitsmatrix ist positiv definit, denn es gilt

$$(x_1, \dots, x_n) \begin{pmatrix} 1 & & \\ & \ddots & \\ & & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = x_1^2 + \dots + x_n^2.$$

Der Ausdruck $x_1^2 + \dots + x_n^2$ wird nur Null, wenn x der Nullvektor ist, ansonsten ist er positiv.

(ii) Allgemeiner gilt: Eine *Diagonalmatrix* (eine Matrix, die nur auf der Diagonalen von Null verschiedene Einträge hat, d.h. es gilt $a_{ij} = 0$ für $i \neq j$) ist genau dann positiv definit, wenn alle Diagonaleinträge a_{ii} positiv sind und genau dann negativ definit, wenn alle Diagonaleinträge a_{ii} negativ sind, denn es gilt

$$(x_1, \dots, x_n) \begin{pmatrix} a_{11} & & \\ & \ddots & \\ & & a_{nn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = a_{11}x_1^2 + \dots + a_{nn}x_n^2.$$

(iii) Auch wenn alle Einträge positiv sind, ist eine symmetrische Matrix nicht unbedingt positiv definit, z.B. gilt:

$$(1, -1) \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = (1, -1) \begin{pmatrix} -1 \\ 1 \end{pmatrix} = -1 - 1 = -2.$$

Das letzte Beispiel zeigt, dass man es einer Matrix, falls sie keine Diagonalmatrix ist, nicht ohne weiteres ansehen kann, ob sie definit ist. Wir werden nun ein Kriterium kennen lernen, mit dem man entscheiden kann, ob eine symmetrische Matrix positiv (oder negativ) definit ist. Um das Kriterium formulieren zu können, benötigen wir den Begriff der Determinante.

Definition 1.3.5. Jeder reellen $n \times n$ -Matrix A werde eine reelle Zahl $\det(A)$, die *Determinante* von A zugeordnet. Dabei ist die Zahl $\det(A)$ durch folgende Eigenschaften festgelegt:

- (i) Ist A eine *obere Dreiecksmatrix*, d.h. gilt $a_{ij} = 0$ falls $i > j$, dann ist $\det(A)$ das Produkt der Diagonalelemente, in Formeln:

$$\det \begin{pmatrix} a_{11} & \dots & a_{1n} \\ & \ddots & \vdots \\ & & a_{nn} \end{pmatrix} = a_{11} \cdot \dots \cdot a_{nn}.$$

- (ii) Addiert man das Vielfache einer Zeile zu einer anderen, so ändert sich die Determinante nicht.
 (iii) Vertauscht man zwei beliebige Zeilen von A , so multipliziert sich die Determinante mit -1 .
 (iv) Multipliziert man eine Zeile mit λ , so multipliziert sich die Determinante ebenfalls mit λ .

Die Definition ist keine explizite Formel für die Determinante, aber sie sagt uns trotzdem, wie wir sie berechnen können: Man bringe die Matrix durch Zeilenumformungen auf obere Dreiecksgestalt, wobei man sich merken muss, wie oft man Zeilen vertauscht hat und mit welchen Faktoren $\lambda_1, \dots, \lambda_k$ man Zeilen malgenommen hat. Die Determinante der dabei entstandenen oberen Dreiecksmatrix ist das Produkt ihrer Diagonalelemente. Um die Determinante der ursprünglichen Matrix zu erhalten, muss man noch durch $\lambda_1 \cdot \dots \cdot \lambda_k$ teilen und mit -1 multiplizieren, falls man eine ungerade Anzahl von Zeilenvertauschungen vorgenommen hat.

Beispiel 1.3.6. Wir führen eine Berechnung einer Determinante nach dem gerade beschriebenen Rezept durch:

$$\begin{aligned} \det \begin{pmatrix} 0 & 3 & 1 \\ 1 & 3 & 1 \\ 1 & 1 & 0 \end{pmatrix} &= -\det \begin{pmatrix} 1 & 3 & 1 \\ 0 & 3 & 1 \\ 1 & 1 & 0 \end{pmatrix} = -\det \begin{pmatrix} 1 & 3 & 1 \\ 0 & 3 & 1 \\ 0 & -2 & -1 \end{pmatrix} = \\ &= -\det \begin{pmatrix} 1 & 2 & 1 \\ 0 & 3 & 1 \\ 0 & 0 & -\frac{1}{3} \end{pmatrix} = -1 \cdot 3 \cdot \left(-\frac{1}{3}\right) = 1. \end{aligned}$$

Bemerkungen 1.3.7. Für kleine n gibt es nützliche Formeln für die Determinate einer $n \times n$ -Matrix.

- (i) Bei einer 1×1 -Matrix ist die Determinante gleich ihrem einzigen Eintrag, d.h. es gilt $\det(A) = a_{11}$, falls $A = (a_{11})$.

- (ii) Für 2×2 -Matrizen berechnet sich die Determinante als das Produkt der Diagonalelemente minus das Produkt der beiden Nichtdiagonalelemente:

$$\det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = a_{11}a_{22} - a_{12}a_{21}.$$

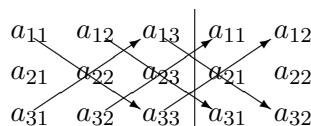
Beispielsweise gilt

$$\det \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} = 1 \cdot 4 - 2 \cdot 3 = 4 - 6 = -2.$$

- (iii) Für 3×3 -Matrizen gibt es die folgende Formel, die *Regel von Sarrus* genannt wird.

$$\det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} = a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{11}a_{23}a_{32} - a_{13}a_{22}a_{31} - a_{12}a_{21}a_{33}$$

Diese Formel wird auch manchmal als *Jägerzaunregel* bezeichnet, da man sie sich wie folgt leicht einprägen kann.



Man erweitert dabei die 3×3 -Matrix, indem man die beiden ersten Spalten noch einmal rechts anfügt. Man bildet jeweils die Produkte von den drei Einträgen, durch die ein Pfeil verläuft, addiert die Summe der derjenigen drei Produkte, die zu den nach rechts unten zeigenden Pfeilen gehören und subtrahiert davon die Summe der derjenigen drei Produkte, die zu den nach rechts oben zeigenden Pfeilen gehören.

Ein Beispiel hierzu: Die Determinante der Matrix

$$\begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 5 \\ 4 & 2 & 1 \end{pmatrix}$$

berechnet sich zu

$$\begin{aligned} & 1 \cdot 1 \cdot 1 + 2 \cdot 5 \cdot 4 + 3 \cdot 3 \cdot 2 - 4 \cdot 1 \cdot 3 - 2 \cdot 5 \cdot 1 - 1 \cdot 3 \cdot 2 = \\ & = 1 + 40 + 18 - 12 - 10 - 6 = 31. \end{aligned}$$

- (iv) *Warnung:* Für $n \times n$ -Matrizen mit $n \geq 4$ ist eine entsprechende Regel *nicht* gültig.

Mit Hilfe von Determinanten lässt sich ein Kriterium für Definitheit von Matrizen formulieren. Dabei muss man die sogenannten *Hauptminoren* der Matrix berechnen. Ist A eine reelle $n \times n$ -Matrix, dann sind ihre Hauptminoren die Determinanten der quadratischen Teilmatrizen der Rahmengrößen $1, \dots, n$ in der linken oberen Ecke. Die Hauptminoren einer $n \times n$ -Matrix mit den Einträgen a_{ij} sind also:

$$\det(a_{11}), \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}, \det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}, \dots, \det(A).$$

Satz 1.3.8 (Kriterium für Definitheit). *Eine $n \times n$ -Matrix A ist genau dann positiv definit, wenn alle ihre Hauptminoren positiv sind, sie ist genau dann negativ definit, wenn ihre Hauptminoren abwechselnd negativ und positiv sind und der erste Hauptminor a_{11} negativ ist.*

Beispiele 1.3.9. Wir betrachten die Matrizen

$$A = \begin{pmatrix} 3 & -1 \\ -1 & 1 \end{pmatrix}, \quad B = \begin{pmatrix} -1 & 2 \\ 2 & -5 \end{pmatrix}, \quad C = \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}.$$

- (i) Die Matrix A hat die Hauptminoren 3 und 2, also ist sie positiv definit.
- (ii) Die Matrix B hat die Hauptminoren -1 und 1, also ist sie negativ definit.
- (iii) Die Matrix C hat die Hauptminoren 1 und -3 , also ist sie weder positiv noch negativ definit.

1.4. Kriterium für lokale Extrema. Theorem 1.1.1 besagt, dass eine differenzierbare Funktion an einer Extremalstelle einen Flachpunkt hat. Umgekehrt gilt:

Satz 1.4.1. *Sei $D \subseteq \mathbb{R}^n$ eine offene Teilmenge und sei $f: D \rightarrow \mathbb{R}$ eine zweimal stetig partiell differenzierbare Funktion, die in $x \in D$ einen Flachpunkt hat, d.h. gilt es $\nabla f(x) = 0$. Dann gilt folgendes.*

- (i) *Ist $\text{Hess}f(x)$ positiv definit, so hat f bei x ein lokales Minimum.*
- (ii) *Ist $\text{Hess}f(x)$ negativ definit, so hat f bei x ein lokales Maximum.*
- (iii) *Gilt $\det(\text{Hess}f(x)) \neq 0$ und ist $\text{Hess}f(x)$ weder positiv noch negativ definit, so hat f bei x einen Sattelpunkt, d.h. f hat in x einen Flachpunkt und in jeder ε -Umgebung von x gibt es sowohl Punkte des Definitionsbereichs, in denen die Funktion größere Werte als $f(x)$ annimmt, als auch Punkte des Definitionsbereichs, in denen die Funktion kleinere Werte als $f(x)$ annimmt.*

Beispiel 1.4.2. Wir wollen alle lokalen Extrema der Funktion $f: \mathbb{R}^3 \rightarrow \mathbb{R}$, $f(x, y, z) = x^2 + y^2 - \cos(z)$ finden. Wir berechnen zunächst den Gradienten:

$$\text{grad} f(x, y, z) = \begin{pmatrix} 2x \\ 2y \\ \sin(z) \end{pmatrix}.$$

Die Flachpunkte sind also $(x, y, z) = (0, 0, k\pi)$. Wir berechnen die Hessematrix an diesen Punkten. Es gilt:

$$\text{Hess} f(x, y, z) = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & \cos(z) \end{pmatrix},$$

somit

$$\text{Hess} f(0, 0, 0) = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & \cos(k\pi) \end{pmatrix}.$$

Die Hessematrix ist also positiv definit bei $(0, 0, 2k\pi)$, es handelt sich bei diesen Punkten um lokale Minima. Sie ist weder positiv noch negativ definit bei $(0, 0, (2k + 1)\pi)$ und ihre Determinante ist dort von Null verschieden, es handelt sich bei diesen Punkten um Sattelpunkte.

Bemerkung 1.4.3. Für das Kriterium ist es wesentlich, dass der Definitionsbereich D der Funktion eine offene Teilmenge des $\mathbb{R}^n$ ist. Enthält D auch Randpunkte, dann kann eine Funktion dort auch ein Extremum annehmen, die man mit diesem Kriterium nicht findet. Etwa nimmt die Funktion $f: [-1, 1] \times [-1, 1] \rightarrow \mathbb{R}$, $f(x, y) = x$ ihr Maximum am Rand des Definitionsbereichs an. Ist die Menge D also nicht offen, so muss man auch die Randpunkte von D , auf denen die Funktion definiert ist, auf Extrema untersuchen.

2. Lokale Extrema unter Nebenbedingungen

Im vorangehenden Abschnitt haben wir Extrema von Funktionen auf offenen Teilmengen des $\mathbb{R}^n$ untersucht. Oft ist es jedoch so, dass man zum Maximieren einer Funktion nicht auf alle Punkte des Definitionsbereichs zurückgreifen darf, sondern nur auf Punkte, die gewisse Gleichungen (oder Ungleichungen) erfüllen. Den Fall von *Gleichungsnebenbedingungen* wollen wir in diesem Abschnitt untersuchen.

Beispiel 2.0.4. Wir wollen dies an einem in der Volkswirtschaftslehre vorkommenden Modell erläutern. Eine *Cobb-Douglas-Funktion* ist eine für positive $x_1, \dots, x_n$ erklärte Funktion der Form

$$f(x_1, \dots, x_n) = a \cdot x_1^{\alpha_1} \cdot \dots \cdot x_n^{\alpha_n},$$

wobei der Vorfaktor a und die Exponenten $\alpha_1, \dots, \alpha_n$ positive Konstanten sind. Solche Funktionen werden oft als *Produktions-* oder *Nutzenfunktionen* verwendet, d.h. sie gibt im ersten Fall die Menge eines produzierten Gutes an, wenn jeweils die Mengen $x_1, \dots, x_n$ der n verschiedenen Inputs eingesetzt werden; im Fall einer Nutzenfunktion ermöglicht es die Funktion, den Nutzen etwa für einen Verbraucher, der jeweils die Menge x_i des Gutes i konsumiert, anzugeben.

Auf dem ganzen Raum $(0, \infty)^n$ betrachtet, hat die Funktion keine lokalen Maxima oder Minima (denn erhöht man ein x_i , erhöht sich der Wert der Funktion, verringert man ein x_i , verringert sich auch der Funktionswert). Hat man aber eine zusätzliche Einschränkung (eine Nebenbedingung) an die zulässigen Punkte $x = (x_1, \dots, x_n)$, kann es lokale Maxima oder Minima unter dieser Nebenbedingung geben. So könnte es bei der Nutzenfunktion eines Verbrauchers etwa eine Budgetbeschränkung geben.

Für $n = 2$ wollen wir ein konkretes Zahlenbeispiel betrachten. Sei $f(x, y) = x^{\frac{1}{2}} \cdot y^{\frac{1}{2}}$ die Nutzenfunktion eines Verbrauchers und sei durch $x + 2y = 1$ eine Budgetbeschränkung, die durch die Funktion $g(x, y) = x + 2y - 1$ darstellen wollen. Wir wollen dann die lokalen Maxima von f unter der Nebenbedingung $g = 0$ bestimmen.

Zunächst wollen wir die Begriffe klären. Wir wollen als Nebenbedingung nicht nur eine Funktion zulassen, sondern mehrere Funktionen $g_1, \dots, g_m$, die wir zu einer Abbildung

$$g: D \rightarrow \mathbb{R}^m, g(x) = (g_1(x), \dots, g_m(x))$$

zusammenfassen.

Sei $f: D \rightarrow \mathbb{R}$ eine Funktion und sei $g: D \rightarrow \mathbb{R}^m$ eine Abbildung. Ein Punkt $x \in D$ heißt *lokale Maximalstelle von f unter der Nebenbedingung $g = 0$* , wenn es eine ε -Umgebung von x gibt, in der keine Punkte von $g^{-1}(\{0\})$ liegen, auf denen f einen größeren Wert als $f(x)$ annimmt.

Sei $f: D \rightarrow \mathbb{R}$ eine Funktion und sei $g: D \rightarrow \mathbb{R}^m$ eine Abbildung. Ein Punkt $x \in D$ heißt *lokale Minimalstelle von f unter der Nebenbedingung $g = 0$* , wenn es eine ε -Umgebung von x gibt, in der keine Punkte von $g^{-1}(\{0\})$ liegen, auf denen f einen kleineren Wert als $f(x)$ annimmt.

In beiden Fällen sagt man auch, $x \in D$ ist eine *lokale Extremalstelle von f unter der Nebenbedingung $g = 0$* .

2.1. Flachpunkte unter Nebenbedingungen.

Definition 2.1.1. Seien $f, g_1, \dots, g_m: D \rightarrow \mathbb{R}$ stetig partiell differenzierbare Funktionen auf der offenen Teilmenge $D \subseteq \mathbb{R}^n$ mit der Eigenschaft, dass ihre Gradienten an allen Punkten linear unabhängig sind, d.h.

$$\nabla g_1(x), \dots, \nabla g_m(x)$$

seien linear unabhängig für alle $x \in D$. Sei die Abbildung $g: D \rightarrow \mathbb{R}^m$ definiert durch $g(x) = (g_1(x), \dots, g_m(x))$. Gilt an einem Punkt $x \in D$, dass sich der Gradient von f als Linearkombination der Gradienten der Funktionen $g_1, \dots, g_m$ darstellen lässt, d.h. gibt es $\lambda_1, \dots, \lambda_m \in \mathbb{R}$, so dass

$$\nabla f(x) = \lambda_1 \nabla g_1(x) + \dots + \lambda_m \nabla g_m(x), \quad (2.2)$$

dann sagt man, f habe einen *Flachpunkt bei x unter der Nebenbedingung $g = 0$* . In diesem Fall heißt das m -Tupel $(\lambda_1, \dots, \lambda_m)$ der *Lagrange-Multiplikator* im Punkt x .

Bemerkungen 2.1.2.

- (i) Für $m = 1$, also für den Fall, dass es nur eine Gleichung als Nebenbedingung gibt, vereinfacht sich die Formulierung der Definition: Wir schreiben dann $g: D \rightarrow \mathbb{R}$ statt g_1 und λ anstatt λ_1 , da der Lagrangemultiplikator in diesem Fall eine reelle Zahl ist. Die Bedingung, dass die Gradienten der Nebenbedingungsfunktionen linear unabhängig sind, bedeutet im Fall $m = 1$ lediglich, dass der Gradient von g nicht verschwindet und die Lagrange-Multiplikatorbedingung reduziert sich auf $\nabla f(x) = \lambda \nabla g(x)$, d.h. $\nabla f(x)$ ist ein Vielfaches von $\nabla g(x)$.
- (ii) Da die Gradienten von $g_1, \dots, g_m$ an jedem Punkt linear unabhängig sind, ist der Lagrange-Multiplikator in jedem Punkt, falls er existiert, eindeutig bestimmt.
- (iii) Die Bedingung, dass die Gradienten von $g_1, \dots, g_m$ überall linear unabhängig sind, garantiert, dass die Menge $g^{-1}(\{0\})$ eine "glatte" oder

“reguläre” $(n-m)$ -dimensionale “Fläche” ist (in mathematischer Terminologie: eine $(n-m)$ -dimensionale Untermannigfaltigkeit des $\mathbb{R}^n$). So lässt sich etwa die Einheitssphäre im $\mathbb{R}^3$ als Urbild der Funktion $g(x, y, z) = x^2 + y^2 + z^2 - 1$ schreiben, die x -Achse als Urbild von $g = (g_1, g_2)$, wobei $g_1(x, y, z) = y$, $g_2(x, y, z) = z$ ist. Im ersten Beispiel verschwindet $\nabla g(x, y, z)$ im Ursprung, d.h. $\nabla g(0, 0, 0) = 0$. Dies kann man dadurch umgehen, indem man die Null aus dem Definitionsbereich D entfernt ($\mathbb{R}^3$ ohne den Ursprung ist offen). Im zweiten Beispiel sind beide Gradienten überall linear abhängig.

Der folgende Satz ist analog zu Satz 1.1.1.

Satz 2.1.3. *Seien $f, g_1, \dots, g_m: D \rightarrow \mathbb{R}$ stetig partiell differenzierbare Funktionen auf der offenen Teilmenge $D \subseteq \mathbb{R}^n$, so dass $\nabla g_1(x), \dots, \nabla g_m(x)$ für alle $x \in D$ linear unabhängig sind. Sei die Abbildung $g: D \rightarrow \mathbb{R}^m$ definiert durch $g(x) = (g_1(x), \dots, g_m(x))$. Hat f in $x \in D$ eine lokale Extremalstelle unter der Nebenbedingung $g = 0$, dann hat f in x einen Flachpunkt unter der Nebenbedingung $g = 0$.*

Zu Beispiel 2.0.4 Wir wollen nun diesen Satz auf das obige Beispiel 2.0.4 anwenden. Sei $f(x, y) = x^{\frac{1}{2}} \cdot y^{\frac{1}{2}}$ und sei durch $g(x, y) = x + 2y - 1$ die Nebenbedingung gegeben. Wir berechnen die Gradienten von f und g :

$$\nabla f(x, y) = \begin{pmatrix} \frac{1}{2}x^{-\frac{1}{2}}y^{\frac{1}{2}} \\ \frac{1}{2}x^{\frac{1}{2}}y^{-\frac{1}{2}} \end{pmatrix}, \quad \nabla g(x, y) = \begin{pmatrix} 1 \\ 2 \end{pmatrix}.$$

Der Ansatz $\nabla f(x, y) = \lambda \nabla g(x, y)$ für den Lagrange-Multiplikator führt auf

$$\begin{pmatrix} \frac{1}{2}x^{-\frac{1}{2}}y^{\frac{1}{2}} \\ \frac{1}{2}x^{\frac{1}{2}}y^{-\frac{1}{2}} \end{pmatrix} = \lambda \begin{pmatrix} 1 \\ 2 \end{pmatrix},$$

d.h. auf die Gleichungen

$$\begin{aligned} \frac{1}{2}x^{-\frac{1}{2}}y^{\frac{1}{2}} &= \lambda, \\ \frac{1}{2}x^{\frac{1}{2}}y^{-\frac{1}{2}} &= 2\lambda. \end{aligned}$$

Die Existenz eines $\lambda \in \mathbb{R}$, für das beide Gleichungen erfüllbar sind, ist äquivalent zu

$$2x^{-\frac{1}{2}}y^{\frac{1}{2}} = x^{\frac{1}{2}}y^{-\frac{1}{2}} \iff 2y = x.$$

Zusammen mit der Nebenbedingung $2y + x = 1$ erhalten wir als einzigen lokalen Flachpunkt unter der Nebenbedingung $g = 0$ den Punkt $(x, y) = (\frac{1}{2}, \frac{1}{4})$.

Die notwendigen Werkzeuge, um diesen Punkt darauf zu untersuchen, ob er ein lokales Extremum unter der Nebenbedingung $g = 0$ ist, werden wir erst im nächsten Unterabschnitt kennenlernen (es handelt sich um eine lokales Maximum unter Nebenbedingung).

Im Fall dieses einfachen Beispiels kann man das Extremum auch mit der sogenannten *Substitutionsmethode* finden. Dazu parametrisiert man die Menge $g^{-1}(\{0\})$ durch die Abbildung $p: (0, 1) \rightarrow \mathbb{R}^2, x \mapsto (x, \frac{1}{2}(1 - x))$. Dabei wird das Intervall $(0, 1)$ bijektiv auf die Menge $g^{-1}(\{0\})$ abgebildet. Um lokale Extrema unter der Nebenbedingung $g = 0$ zu finden, kann man genauso gut die lokalen Extrema der Verkettung $f \circ p: (0, 1) \rightarrow \mathbb{R}, x \mapsto \frac{\sqrt{2}}{2}x^{\frac{1}{2}}(1 - x)^{\frac{1}{2}}$ bestimmen. Es gilt

$$(f \circ p)'(x) = \frac{\sqrt{2}}{4} \left(x^{-\frac{1}{2}}(1 - x)^{\frac{1}{2}} - x^{\frac{1}{2}}(1 - x)^{-\frac{1}{2}} \right).$$

Somit ist $(f \circ p)'(x) = 0$ äquivalent zu

$$x^{-\frac{1}{2}}(1 - x)^{\frac{1}{2}} = x^{\frac{1}{2}}(1 - x)^{-\frac{1}{2}} \iff x = 1 - x,$$

was uns wieder den Flachpunkt bei $x = \frac{1}{2}, y = \frac{1}{4}$ liefert. Um zu untersuchen, ob (und was für) ein Extremum vorliegt, berechnen wir

$$\begin{aligned} (f \circ p)''(x) &= \\ &= \frac{\sqrt{2}}{4} \left(-\frac{1}{2}x^{-\frac{3}{2}}(1 - x)^{\frac{1}{2}} - \frac{1}{2}x^{-\frac{1}{2}}(1 - x)^{-\frac{1}{2}} \right. \\ &\quad \left. - \frac{1}{2}x^{-\frac{1}{2}}(1 - x)^{-\frac{1}{2}} - \frac{1}{2}x^{\frac{1}{2}}(1 - x)^{-\frac{3}{2}} \right) = \\ &= \frac{\sqrt{2}}{8} \left(-x^{-\frac{3}{2}}(1 - x)^{\frac{1}{2}} - 2x^{-\frac{1}{2}}(1 - x)^{-\frac{1}{2}} - x^{\frac{1}{2}}(1 - x)^{-\frac{3}{2}} \right). \end{aligned}$$

Dieser Term ist für alle $x \in (0, 1)$ negativ und es gilt insbesondere $(f \circ p)''(\frac{1}{2}) < 0$, also handelt es sich um ein lokales Maximum unter Nebenbedingung.

Ist die Menge $g^{-1}(\{0\})$ von komplizierterer Gestalt und lässt sich nicht so einfach parametrisieren, ist es oft nicht ohne weiteres möglich, die Substitutionsmethode anzuwenden. Die Lagrangemultiplikatormethode wird in den Wirtschaftswissenschaften auch deswegen gerne verwendet, da sie zum einen oft einfacher anzuwenden ist als die Substitutionsmethode und da sich zum anderen die auftretenden Lagrangemultiplikatoren als ökonomische Größen interpretieren lassen.

2.2. Kriterium für Extrema unter Nebenbedingungen. Wir werden nun ein Kriterium analog zu Satz 1.4.1 kennen lernen, um Extrema unter Nebenbedingungen bestimmen zu können. Wir setzen in diesem Unterabschnitt voraus, dass $D \subseteq \mathbb{R}^n$ offen ist und dass die Funktionen $f, g_1, \dots, g_m: D \rightarrow \mathbb{R}$ alle zweimal stetig partiell differenzierbar sind.

Wir nehmen jetzt an, dass $x \in D$ ein Flachpunkt von f unter der Nebenbedingung $g = 0$ ist, dass also ein Lagrange-Multiplikator $(\lambda_1, \dots, \lambda_m)$ existiert, so dass (2.2) gilt. Wir definieren nun eine Hilfsfunktion $F: D \rightarrow \mathbb{R}$ wie folgt:

$$F(x) := f(x) - \lambda_1 g_1(x) - \dots - \lambda_m g_m(x).$$

Sei außerdem $N(x)$ die Matrix, deren Spalten von den Gradienten der Funktionen $g_1, \dots, g_m$ an der Stelle x :

$$N(x) := (\nabla g_1(x), \dots, \nabla g_m(x)).$$

Nun betrachten wir das folgende homogene lineare Gleichungssystem für Vektoren $u \in \mathbb{R}^n$, wobei man beachte, dass die Transponierte der Matrix $N(x)$ verwendet wird:

$$N(x)^t u = 0.$$

Dies kann man mit dem Gaußverfahren lösen und man erhält somit eine Basis des Lösungsraumes. Daraus, dass die Gradienten von $g_1, \dots, g_m$ bei x (und somit die Zeilen der Matrix $N(x)^t$) linear unabhängig sind, folgt, dass der Lösungsraum die Dimension $n - m$ hat, d.h. eine Basis des Lösungsraums hat $n - m$ Elemente $u_1, \dots, u_{n-m}$. Wir bilden nun die $n \times (n - m)$ -Matrix U , deren Spalten die Vektoren $u_1, \dots, u_{n-m}$ sind:

$$U := (u_1, \dots, u_{n-m}).$$

Im folgenden Satz betrachten wir die folgende $(n - m) \times (n - m)$ -Matrix, die mit Hilfe von U aus der Hessematrix $\text{Hess}F(x)$ der Hilfsfunktion F gebildet wird:

$$U^t \text{Hess}F(x) U.$$

Satz 2.2.1. *Mit den Bezeichnungen wie oben gilt:*

- (i) *Ist $U^t \text{Hess}F(x) U$ positiv definit, so hat f bei x ein lokales Minimum unter der Nebenbedingung $g = 0$.*
- (ii) *Ist $U^t \text{Hess}F(x) U$ negativ definit, so hat f bei x ein lokales Maximum unter der Nebenbedingung $g = 0$.*

- (iii) Gilt $\det(U^t \text{Hess} F(x) U) \neq 0$ und ist $U^t \text{Hess} F(x) U$ weder positiv noch negativ definit, so hat f bei x einen Sattelpunkt unter der Nebenbedingung $g = 0$.

Bemerkungen 2.2.2.

- (i) Die Aussage “ f hat bei x einen Sattelpunkt unter der Nebenbedingung $g = 0$ ” bedeutet: f hat in x einen Flachpunkt und in jeder ε -Umgebung von x gibt es sowohl Punkte von $D \cap g^{-1}(\{0\})$, in denen die Funktion größere Werte als $f(x)$ annimmt, als auch Punkte $D \cap g^{-1}(\{0\})$, in denen die Funktion kleinere Werte als $f(x)$ annimmt.
- (ii) Die Matrix U , deren Spalten Lösungen des Gleichungssystems $N(x)^t u = 0$ sind, ist durch diese Forderung nicht eindeutig bestimmt, im allgemeinen gibt es viele verschiedene Basen des Lösungsraums. Ersetzt man also die Spalten von U durch eine andere Basis des Lösungsraums, so erhält man im allgemeinen auch eine andere Matrix $U^t \text{Hess} F(x) U$, es ändert sich aber nichts an der Definitheit der Matrix oder daran, ob die Determinante von Null verschieden ist.
- (iii) Die geometrische Bedeutung der Matrix U besteht darin, dass ihre Spalten eine Basis des “Tangentialraum” an die Menge $g^{-1}(\{0\})$ im Punkt x bilden. Das Gleichungssystem $N(x)^t u = 0$ ist nämlich dazu äquivalent, alle Vektoren zu finden, die auf den Vektoren $\nabla g_1(x), \dots, \nabla g_m(x)$ senkrecht stehen.
- (iv) Der Fall von Extrema ohne Nebenbedingungen ist im obigen Satz als Spezialfall enthalten, nämlich indem man $m = 0$ setzt. Die Bedingung (2.2) bedeutet in diesem Fall genau $\nabla f(x) = 0$ und für die Matrix U kann man die $n \times n$ -Einheitsmatrix verwenden.

Beispiel 2.2.3. Wir betrachten die Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, $f(x, y) = xy$ mit der durch die Funktion $g: \mathbb{R}^2 \rightarrow \mathbb{R}$, $g(x, y) = x^2 + y^2 - 1$ gegebenen Nebenbedingung. Mit anderen Worten, wir suchen unter den Punkten auf dem Einheitskreis, diejenigen, an denen die Funktion lokale minimale oder maximale Werte annimmt. Wir berechnen die Gradienten beider Funktionen.

$$\nabla f(x, y) = \begin{pmatrix} y \\ x \end{pmatrix}, \quad \nabla g(x, y) = \begin{pmatrix} 2x \\ 2y \end{pmatrix}.$$

Die Bedingung $\nabla f(x, y) = \lambda \nabla g(x, y)$ führt auf die beiden Gleichungen

$$y = 2\lambda x \quad \text{und} \quad x = 2\lambda y.$$

Einsetzen einer Gleichung in die andere führt auf $y = 4\lambda^2 y$ oder $x = 4\lambda^2 x$. Da bei $g = 1$, also auf dem Einheitskreis, x und y nicht beide gleich Null sind, folgt

zum einen zunächst $\lambda = \pm \frac{1}{2}$. Außerdem folgt nun $y = \pm x$. Die Flachpunkte unter Nebenbedingungen sind also genau die vier Punkte, die sich als Schnittpunkte des Einheitskreises um den Ursprung mit den beiden Winkelhalbierenden des Koordinatensystems ergeben. Dies sind die Punkte

$$\frac{\sqrt{2}}{2} \begin{pmatrix} +1 \\ +1 \end{pmatrix}, \quad \frac{\sqrt{2}}{2} \begin{pmatrix} -1 \\ +1 \end{pmatrix}, \quad \frac{\sqrt{2}}{2} \begin{pmatrix} -1 \\ -1 \end{pmatrix}, \quad \frac{\sqrt{2}}{2} \begin{pmatrix} +1 \\ -1 \end{pmatrix}.$$

Die Funktion f nimmt an diesen Punkten die Werte $\frac{1}{2}, -\frac{1}{2}, \frac{1}{2}, -\frac{1}{2}$ an.

Betrachten wir zunächst die beiden Punkte

$$\pm \frac{\sqrt{2}}{2} \begin{pmatrix} +1 \\ +1 \end{pmatrix}.$$

Für diese Punkte haben wir $\lambda = \frac{1}{2}$. Wir erhalten als Hilfsfunktion $F(x, y) = xy - \frac{1}{2}(x^2 + y^2 - 1)$. Die Hessematrix ist die konstante Matrix

$$\text{Hess}F(x, y) = \begin{pmatrix} -1 & 1 \\ 1 & -1 \end{pmatrix}$$

Wir erhalten $N(x)^t = \pm(\sqrt{2}, \sqrt{2})$. Eine Lösung für U ist

$$U = \begin{pmatrix} -1 \\ 1 \end{pmatrix} \Rightarrow U^t \text{Hess}F(x, y) U = -4,$$

es handelt sich um zwei lokale Maxima unter Nebenbedingung.

Für die Punkte

$$\pm \frac{\sqrt{2}}{2} \begin{pmatrix} +1 \\ -1 \end{pmatrix}$$

haben wir $\lambda = \frac{1}{2}$ und erhalten als Hilfsfunktion $F(x, y) = xy + \frac{1}{2}(x^2 + y^2 - 1)$. Die Hessematrix ist die konstante Matrix

$$\text{Hess}F(x, y) = \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$$

Wir erhalten $N(x)^t = \pm(\sqrt{2}, -\sqrt{2})$. Eine Lösung für U ist

$$U = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \Rightarrow U^t \text{Hess}F(x, y) U = 4,$$

es handelt sich um zwei lokale Minima unter Nebenbedingung.

Fortsetzung von Beispiel 2.0.4 Wir wollen das lokale Extremum noch einmal, diesmal mit Hilfe von Satz 2.2.1, bestimmen. Aus der Rechnung oben folgt $\lambda = \frac{\sqrt{2}}{4}$ für den einzigen Flachpunkt $(x, y) = (\frac{1}{2}, \frac{1}{4})$ unter der Nebenbedingung $g = 0$. Wir haben $F(x, y) = x^{\frac{1}{2}}y^{\frac{1}{2}} - \frac{\sqrt{2}}{4}(x + 2y - 1)$. Wir erhalten

$$\text{Hess}F(x, y) = \frac{1}{4} \begin{pmatrix} x^{-\frac{3}{2}}y^{\frac{1}{2}} & x^{-\frac{1}{2}}y^{-\frac{1}{2}} \\ x^{-\frac{1}{2}}y^{-\frac{1}{2}} & x^{\frac{1}{2}}y^{-\frac{3}{2}} \end{pmatrix}$$

und daraus

$$\text{Hess}F(x, y) = \begin{pmatrix} -\frac{1}{4}\sqrt{2} & \frac{1}{2}\sqrt{2} \\ \frac{1}{2}\sqrt{2} & -\sqrt{2} \end{pmatrix}.$$

Mit $N(\frac{1}{2}, \frac{1}{4})^t = (1, 2)$ erhalten wir

$$U = \begin{pmatrix} -2 \\ 1 \end{pmatrix} \Rightarrow U^t \text{Hess}F(\frac{1}{2}, \frac{1}{4})U = -4\sqrt{2},$$

was noch einmal zeigt, dass ein lokales Maximum unter Nebenbedingung vorliegt.

3. Kompakta

Zum Abschluss dieses Kapitels wollen wir noch ein, zunächst sehr theoretisch klingendes, aber oft durchaus nützliches Kriterium für die Existenz von Extrema angeben. Dazu führen wir den Begriff einer kompakten Teilmenge des $\mathbb{R}^n$ ein.

Eine Teilmenge $A \subseteq \mathbb{R}^n$ heißt *abgeschlossen*, wenn ihr *Komplement* $\mathbb{R}^n \setminus A$ offen ist, oder, äquivalent, wenn sie alle ihre Randpunkte enthält. Beispiele für abgeschlossene Mengen sind abgeschlossene Intervalle $[a, b]$ in $\mathbb{R}$, oder etwa die x -Achse oder ein Kreis in $\mathbb{R}^2$. Als einfache Merkregel kann man sagen, dass Teilmengen, die durch Gleichungen definiert sind, z.B. $M = \{x \in \mathbb{R}^n \mid f(x) = 0\}$ oder *Subniveaus* $M = \{x \in \mathbb{R}^n \mid f(x) \leq 0\}$ abgeschlossen sind, wobei wir voraussetzen, dass f stetig ist.

Eine Teilmenge $A \subseteq \mathbb{R}^n$ heißt *beschränkt*, wenn es eine positive reelle Zahl C gibt, so dass $\|x\| \leq C$ für alle $x \in A$ gilt.

Eine Teilmenge $A \subseteq \mathbb{R}^n$ heißt *kompakt*, wenn sie sowohl abgeschlossen als auch beschränkt ist.

Satz 3.0.4. *Stetige Funktionen auf kompakten Teilmenge des $\mathbb{R}^n$ nehmen Maximum und Minimum an.*

Zum Beispiel folgt aus diesem Satz sofort, dass die Funktion f in Beispiel 2.2.3 mindestens ein Minimum und mindestens ein Maximum annimmt (denn f ist stetig und der Einheitskreis kompakt).

KAPITEL 3

Approximation

In diesem Kapitel beschäftigen wir uns mit der näherungsweisen Berechnung und Darstellung von Funktionen und Zahlenwerten.

1. Taylorpolynome

Die Taylorformel dient dazu, eine Funktion durch ein Polynom näherungsweise zu ersetzen. Die Idee dabei ist es, die Funktion in einer Umgebung eines Punktes anzunähern, indem man die Information über den Funktionswert und die Ableitungen an diesem Punkt verwendet. Wir wollen diese Vorgehensweise zunächst anhand von reellen Funktionen in einer Variablen kennenlernen.

1.1. Taylorpolynom in einer Variablen. Sei $f: I \rightarrow \mathbb{R}$ eine auf dem Intervall I definierte, beliebig oft differenzierbare Funktion. Mit $f^{(k)}(x)$ bezeichnen wir den Wert der k -ten Ableitung von f an der Stelle $x \in I$, d.h.

$$\begin{aligned} f^{(0)}(x) &= f(x) = && \text{Funktionswert bei } x, \\ f^{(1)}(x) &= f'(x) = && \text{Steigung des Graphen bei } x, \\ f^{(2)}(x) &= f''(x) = && \text{Wert der zweiten Ableitung bei } x, \\ &&& \text{etc.} \end{aligned}$$

Wir wollen nun ein Polynom finden, so dass der Wert von der nullten bis zur k -ten Ableitung des Polynoms an einem Punkt x mit den Werten der entsprechenden Ableitungen von f bei x übereinstimmt. Dazu überlegen wir uns, was die ersten k Ableitungen eines Polynoms sind. Zunächst betrachten wir den Spezialfall $x = 0$, d.h. wir wollen die Funktion zunächst in der Nähe der Null annähern (wir nehmen dazu $0 \in I$ an).

Beispiel 1.1.1. Wir berechnen die ersten Ableitungen für ein Beispiel, wir betrachten das Polynom $p(x) = x^4 + 8x^3 - 9x^2 + x + 7$.

k	k -te Ableitung	Wert bei $x = 0$
0	$x^4 + 8x^3 - 9x^2 + x + 7$	7
1	$4x^3 + 3 \cdot 8x^2 - 2 \cdot 9x + 1$	1
2	$3 \cdot 4x^2 + 2 \cdot 3 \cdot 8x - 2 \cdot 9$	-18
3	$2 \cdot 3 \cdot 4x + 2 \cdot 3 \cdot 8$	48
4	$2 \cdot 3 \cdot 4$	24
5	0	0

Natürlich gilt für ein Polynom n -ten Grades, dass die Ableitungen $f^{(n+1)}, f^{(n+2)}, \dots$ von höherer als n -ter Ordnung alle gleich Null sind. Außerdem gilt für $k \leq n$, dass der Wert der k -te Ableitung eines Polynoms $p(x) = a_n x^n + \dots + a_1 x + a_0$ vom Grad n gleich

$$p^{(k)}(0) = 1 \cdot 2 \cdot 3 \cdot \dots \cdot k \cdot a_k(0) = k! \cdot a_k(0)$$

ist. Umgekehrt bedeutet dies: Wählen wir ein Polynom n -ten Grades, dessen Koeffizienten durch

$$a_k = \frac{f^{(k)}(0)}{k!}$$

gegeben sind, dann gilt für dieses Polynom, dass sein Funktionswert bei $x = 0$ und seine erste bis k -te Ableitung bei $x = 0$ mit den entsprechenden Werten von f übereinstimmen. Dies führt auf die folgende Definition.

Definition 1.1.2. Das *Taylorpolynom* n -ter Ordnung von f um den Punkt $x = 0$ ist gegeben durch

$$\sum_{k=0}^n \frac{f^{(k)}(0)}{k!} x^k = f(0) + f'(0)x + \frac{f''(0)}{2}x^2 + \dots + \frac{f^{(n-1)}(0)}{(n-1)!}x^{n-1} + \frac{f^{(n)}(0)}{n!}x^n.$$

Beispiele 1.1.3. (i) Wir bestimmen das Taylorpolynom 5-ter Ordnung der Sinusfunktion um $x = 0$. Dazu berechnen wir den Funktionswert sowie die erste bis fünfte Ableitung von $f(x) = \sin(x)$ jeweils an der Stelle $x = 0$.

k	k -te Ableitung	Wert bei $x = 0$
0	$\sin^{(0)}(x) = \sin(x)$	$\sin(0) = 0$
1	$\sin^{(1)}(x) = \cos(x)$	$\cos(0) = 1$
2	$\sin^{(2)}(x) = -\sin(x)$	$-\sin(0) = 0$
3	$\sin^{(3)}(x) = -\cos(x)$	$-\cos(0) = -1$
4	$\sin^{(4)}(x) = \sin(x)$	$\sin(0) = 0$
5	$\sin^{(5)}(x) = \cos(x)$	$\cos(0) = 1$

Das Polynom ist also gegeben durch (die Koeffizienten mit geradem Index sind Null)

$$\frac{\sin^{(5)}}{5!}x^5 + \frac{\sin^{(3)}}{3!}x^3 + \frac{\sin^{(1)}}{1!}x = \frac{1}{120}x^5 - \frac{1}{6}x^3 + x.$$

Man sieht, dass das Taylorpolynom nur in der Nähe der Null den Funktionsverlauf gut wiedergibt, denn die Sinusfunktion ist periodisch, während das Polynom für $x \rightarrow \pm\infty$ gegen $\pm\infty$ geht.

Das Taylorpolynom kommt z.B. in Taschenrechnern zur Anwendung, um etwa numerische Werte der Sinusfunktion zu berechnen.

- (ii) Wegen $\exp'(x) = \exp(x)$ nehmen alle Ableitungen der Exponentialfunktion bei $x = 0$ den Wert $\exp(0) = 1$ an. Daran sieht man, dass die Partialsummen der Exponentialreihe

$$\exp(x) = \sum_{k=0}^n \frac{x^k}{k!}$$

nichts anderes sind als die Taylorpolynome der Exponentialfunktion bei $x = 0$.

Allgemeiner ist es manchmal wünschenswert, die Annäherung durch ein Taylorpolynom nicht bei $x = 0$ sondern um einen anderen Punkt x_0 durchzuführen. Die entsprechende Formel bekommt man, indem man x_0 in den Nullpunkt “verschiebt”.

Definition 1.1.4. Das *Taylorpolynom* n -ter Ordnung von f um den Punkt a ist gegeben durch

$$\begin{aligned} \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k &= \\ &= f(a) + f'(a)(x-a) + \frac{f''(a)}{2}(x-a)^2 + \cdots + \frac{f^{(n)}(a)}{n!}(x-a)^n. \end{aligned}$$

Beispiele 1.1.5. (i) Das Taylorpolynom erster Ordnung bei $x = a$ ist das lineare Polynom, dessen Graph die Tangente an den Graph von f bei $x = a$ ist. In diesem Sinne ist das Taylorpolynom eine Verallgemeinerung der Grundidee der Differentialrechnung, nämlich der Approximation durch eine lineare Abbildung. Beim Taylorpolynom wird eine Funktion durch Polynom beliebig hohen Grades angenähert.

- (ii) Das Taylorpolynom zweiten Grades liefert die Annäherung an den Graphen von f durch eine *Schmiegeparabel*. Diese ist der Graph eines quadratischen Polynoms, der bei $x = a$ denselben Funktionswert, dieselbe Tangentensteigung und denselben Wert der zweiten Ableitung hat.
- (iii) Wir berechnen das Taylorpolynom n -ter Ordnung von $f(x) = \ln(x)$ bei $x = 1$. Dazu berechnen wir den Wert der Ableitungen von f bei 1. Es gilt $\ln'(x) = x^{-1}$, daraus folgt $\ln^{(k)}(x) = (-1)^{k-1}(k-1)! \cdot x^{-k}$ und somit $\ln^{(k)}(1) = (-1)^{k-1}(k-1)!$ für $k \geq 1$. Wegen $\ln(1) = 0$ fällt das Absolutglied weg und wir erhalten für das n -te Taylorpolynom

$$\sum_{k=1}^n \frac{(-1)^{k-1}}{k} (x-1)^k.$$

Setzt man $x = 2$ ein, erhält man die Partialsummen der alternierenden harmonischen Reihe, siehe Kapitel 1, Unterabschnitt 5.2. Aus der Tatsache, dass die obige Summe für $n \rightarrow \infty$ gegen den Wert $\ln(x)$ konvergiert (was wir an dieser Stelle nicht beweisen) folgt, dass die alternierende harmonische Reihe gegen $\ln(2)$ konvergiert.

Bis jetzt haben wir uns noch nicht mit der Güte der Approximation durch ein Taylorpolynom befasst. Dazu kann man den folgenden Satz benutzen.

Satz 1.1.6 (Satz von Taylor). *Es gilt:*

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + \frac{1}{n!} \int_a^x f^{(n+1)}(t) (x-t)^n dt.$$

Der zweite Term auf der rechten Seite wird das (*Integral-*)*Restglied* genannt. Es gibt den Fehler an, der beim Ersetzen einer Funktion durch ihr Taylorpolynom n -ter Ordnung gemacht wird. Wenn man es abschätzen kann, d.h. garantieren kann, dass es betragsmäßig nicht größer als ein gewisser Wert werden kann, dann weiß man, welchen Fehler man schlimmstenfalls bei der Approximation macht. Kann man zeigen, dass für $n \rightarrow \infty$ für jedes $x \in I$ das Restglied gegen Null geht, dann hat man gezeigt, dass die *Taylorreihe*

$$\sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (x-a)^k$$

überall im Intervall I gegen den Wert von $f(x)$ konvergiert.

1.2. Taylorpolynom in mehreren Variablen. Zum Abschluss des Abschnitts wollen wir noch die Taylorformel für Funktionen mehrerer Variablen angeben. Dabei beschränken wir uns auf die zweite Ordnung. Sei $D \subseteq \mathbb{R}^n$ eine offene Teilmenge und sei $a \in D$. Dann ist das Taylorpolynom von f bei $x = a$ gegeben durch

$$f(a) + (x - a) \cdot \nabla f(a) + \frac{1}{2}(x - a)^t \text{Hess} f(a)(x - a),$$

wobei der Punkt zwischen x und $\nabla f(a)$ das Skalarprodukt bezeichnet.

Beispiel 1.2.1. Sei $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ gegeben durch $f(x, y) = e^{x^2+y}$. Wir bestimmen das Taylorpolynom zweiter Ordnung beim Punkt $a = (0, 0)$. Dazu berechnen wir

$$f(0, 0) = 1,$$

sowie

$$\nabla f(x, y) = \begin{pmatrix} 2xe^{x^2+y} \\ e^{x^2+y} \end{pmatrix}, \text{ also } \nabla f(0, y) = \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

und

$$\text{Hess} f(x, y) = \begin{pmatrix} 2e^{x^2+y} + 4x^2e^{x^2+y} & 2xe^{x^2+y} \\ 2xe^{x^2+y} & e^{x^2+y} \end{pmatrix}, \text{ also } \text{Hess} f(0, 0) = \begin{pmatrix} 2 & 0 \\ 0 & 1 \end{pmatrix}.$$

Es ergibt sich

$$1 + y + x^2 + \frac{1}{2}y^2$$

für das Taylorpolynom.

2. Das Newton-Verfahren

Wir wollen noch ein weiteres numerisches Rechenverfahren kennenlernen. Mit dem Newtonverfahren kann man näherungsweise eine Gleichung lösen, bzw. eine Nullstelle einer Funktion bestimmen. Wir sehen uns das Verfahren zunächst im Eindimensionalen an.

2.1. Newton-Verfahren, eindimensional. Sei f eine auf dem Intervall $[a, b]$ definierte reelle stetige Funktion, die bei a einen negativen Wert und bei b einen positiven Wert annimmt. Nach dem Zwischenwertsatz 5.3.11 in Kapitel 1 hat f mindestens eine Nullstelle in $[a, b]$. Angenommen, f ist streng monoton steigend. Dann hat f genau eine Nullstelle in $[a, b]$.

Ein naheliegendes Anwendungsbeispiel ist es etwa, die Quadratwurzel einer positiven Zahl numerisch bis zur einer gewünschten Genauigkeit (ein gewisse Anzahl

Nachkommastellen) zu berechnen. Um z.B. eine Annäherung von $\sqrt{2}$ zu bekommen, können wir die Funktion $f(x) = x^2 - 2$ untersuchen. Da $f(1) = -1$ und $f(2) = 2$ gilt, können wir uns z.B. auf das Intervall $[a, b] = [1, 2]$ beschränken.

Ein einfaches Verfahren, um die Nullstelle zu suchen, ist die sogenannte *Bisektion*. Dabei wird das Intervall fortlaufend durch Teilintervalle der halben Länge ersetzt, je nachdem, ob die Nullstelle vor dem entsprechenden Unterteilungsschritt links oder rechts vom Mittelpunkt des Intervalls auftritt, d.h. je nachdem, ob $\frac{a+b}{2} \geq 0$ oder $\frac{a+b}{2} < 0$ gilt. Man erhält eine Intervallschachtelung, die sich auf die Nullstelle zusammenzieht.

Beispiel 2.1.1. Wir wollen dies anhand des gerade genannte Beispiels demonstrieren, also numerisch den Wert von $\sqrt{2}$ durch Bisektion approximieren. Wir betrachten also die Funktion $f(x) = x^2 - 2$ und starten mit dem Intervall $[1, 2]$. Im ersten Schritt berechnen wir $f(1,5) = 0,25$. Hier die ersten 10 Schritte des Verfahrens, mit einem Computerprogramm berechnet:

Schritt	a	b	$f(\frac{a+b}{2})$
1	1,000000000	2,000000000	0,250000000
2	1,000000000	1,500000000	-0,437500000
3	1,250000000	1,500000000	-0,109375000
4	1,375000000	1,500000000	0,066406250
5	1,375000000	1,437500000	-0,022460938
6	1,406250000	1,437500000	0,021728516
7	1,406250000	1,421875000	-0,000427246
8	1,414062500	1,421875000	0,010635376
9	1,414062500	1,417968750	0,005100250
10	1,414062500	1,416015625	0,002335547

Nach 10 Schritten sind zwei gültige Nachkommastellen von $\sqrt{2} = 1,41421356 \dots$ bestimmt.

Wir wollen nun ein anderes Verfahren betrachten, bei dem nicht nur der Funktionswert, sondern auch die Ableitung der Funktion an der untersuchten Stelle berücksichtigt wird. Dadurch kann man unter geeigneten Voraussetzungen eine höhere Konvergenzgeschwindigkeit erzielen.

Die Idee des Newton-Verfahrens ist die folgende. Man wähle einen Startpunkt x_1 in der Nähe einer zu ermittelnden Nullstelle und ersetze dann die Funktion durch ihr Taylorpolynom erster Ordnung, d.h. anschaulich, man ersetze den Graph der Funktion f durch seine Tangente bei x_1 . Anstatt der Nullstelle von f ermittelt man nun die Nullstelle x_2 des Taylorpolynoms (wir wollen annehmen, dass dieses nicht konstant ist und somit eine eindeutig bestimmte Nullstelle hat). Nun führt man das Verfahren mit x_2 an Stelle von x_1 fort.

In Formeln: Ersetze f durch die lineare Annäherung

$$x \mapsto f(x_1) + f'(x_1)(x - x_1)$$

Die Nullstelle x_2 dieser Funktion bestimmt sich (unter der Annahme, dass $f'(x_1) \neq 0$) aus

$$f(x_1) + f'(x_1)(x_2 - x_1) = 0 \iff x_2 = x_1 - \frac{f(x_1)}{f'(x_1)}.$$

Als Formel für das *Newton-Verfahren* erhalten wir allgemein die rekursive Formel

$$x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}. \quad (2.3)$$

Das Verfahren konvergiert unter den folgenden Voraussetzungen.

Satz 2.1.2. Sei $f: D \rightarrow \mathbb{R}$ eine auf dem offenen Intervall D definierte, beliebig oft differenzierbare Funktion. Seien $a, b \in D$ mit $a < b$ und seien $f'(x) > 0$, $f''(x) > 0$ für alle $x \in [a, b]$. Dann konvergiert das Newton-Verfahren, d.h. die durch (2.3) definierte Folge, mit $x_1 = b$, konvergiert gegen die (einzige) Nullstelle von f auf $[a, b]$.

Die Konvergenz dabei ist “lokal quadratisch”, d.h. grob gesprochen, in jedem Schritt verdoppelt sich die Anzahl der gültigen Nullstellen.

Beispiel 2.1.3. Wir wenden das Verfahren auf die numerische Berechnung von $\sqrt{2}$ an. Wir betrachten also die Funktion $f(x) = x^2 - 2$. Wegen $f'(x) = 2x$ lautet die Rekursionsformel hier

$$x_{n+1} = x_n - \frac{x_n^2 - 2}{2x_n} = \frac{x_n}{2} + \frac{1}{x_n}.$$

Unter Benutzung eines Computers erhalten wir die ersten 7 Schritte des Verfahrens. Um die Konvergenzgeschwindigkeit zu demonstrieren, geben wir jeweils 30 gültige Stellen an. Wir beginnen mit dem Startwert $x_1 = 2$.

n	x_n
1	2,000000000000000000000000000000
2	1,500000000000000000000000000000
3	1,416666666666666666666666666667
4	1,41421568627450980392156862745
5	1,41421356237468991062629557889
6	1,41421356237309504880168962350
7	1,41421356237309504880168872421
8	1,41421356237309504880168872421

Es fällt auf, dass das Verfahren viel schneller als das Bisektionsverfahren gegen $\sqrt{2} = 1,41421356237309504880168872421 \dots$ konvergiert.

Dieses Verfahren zum Bestimmen der Quadratwurzel einer positiven Zahl a durch Anwenden des Newton-Verfahrens auf die Funktion $f(x) = x^2 - a$, hier demonstriert für $a = 2$, nennt man auch *Babylonisches Wurzelziehen* oder *Heronverfahren*.

2.2. Newton-Verfahren, mehrdimensional. Hat man n reelle Funktionen in n reellen Variablen, so kann man wie folgt vorgehen. Man fasst die Funktionen zu einer Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ zusammen. Eine Nullstelle der Abbildung ist dann ein Punkt x im $\mathbb{R}^n$, für den $f(x)$ der Nullvektor im $\mathbb{R}^n$ ist. Wieder kann man die Abbildung f durch ihr erstes Taylorpolynom (d.h. durch die affin-lineare Abbildung, die durch das Differential von f gegeben ist) annähern, und dann ersatzweise nach einer Nullstelle dieser Abbildung

$$x \mapsto f(x_1) + Df(x_1)(x - x_1)$$

suchen. Damit die Nullstelle eindeutig bestimmt ist, muss das Differential $Df(x_1)$ eine invertierbare Matrix sein. Falls $Df(x_1)$ invertierbar ist, erhalten wir $x_2 = x_1 - Df(x_1)^{-1}f(x_1)$, bzw. die allgemeine Rekursionsformel

$$x_{k+1} = x_k - Df(x_k)^{-1}f(x_k).$$

KAPITEL 4

Differentialgleichungen

Differentialgleichungen sind Gleichungen, in denen eine Funktion und ihre Ableitungen vorkommen, wobei eine Funktion als Lösung gesucht ist. Ein Beispiel für eine Differentialgleichung (in typischer Schreibweise) ist

$$y' = y.$$

Diese Art der Kurzschreibweise, die sehr verbreitet ist, und die wir uns natürlich auch aneignen wollen, bedeutet: Es ist eine reelle Funktion in einer Variablen gesucht, die gleich ihrer eigenen Ableitung ist. Natürlich kennen Sie bereits eine Lösung dieser Differentialgleichung: Die Exponentialfunktion $x \mapsto \exp(x) = e^x$ hat genau diese Eigenschaft: Sie ist gleich ihrer eigenen Ableitung, denn es gilt stets $\exp'(x) = \exp(x)$.

Viele Vorgänge in den Natur- und Wirtschaftswissenschaften können durch Differentialgleichungen beschrieben werden. Ist die gesuchte Funktion eine Funktion in einer Variablen, so kann die Funktion etwa den zeitlichen Verlauf einer Größe beschreiben. Die obige Gleichung (bzw. eine Variante davon) etwa beschreibt Wachstums- und Zerfallsprozesse in der Natur. Nehmen wir an, wir haben etwa eine Bakterienkolonie in einer Nährlösung. Sei $y(x)$ die Anzahl der Bakterien zum Zeitpunkt x . Das Wachstum der Kolonie ist durch die Zellteilungsrate gegeben. Angenommen, pro Zeiteinheit und Bakterium finden durchschnittlich $\alpha > 0$ Zellteilungen statt. Dann wächst die Gesamtanzahl der Bakterien pro Zeiteinheit um $\alpha \cdot y(x)$. Mit anderen Worten, die Steigung der Funktion $y(x)$ ist durch $\alpha \cdot y(x)$ gegeben, oder, die Funktion $y(x)$ erfüllt die Differentialgleichung (in Kurzschreibweise)

$$y' = \alpha y.$$

Auch dafür kann man eine Lösung noch relativ einfach erraten: Die Funktion $x \mapsto e^{\alpha x}$ hat als Ableitung $\alpha e^{\alpha x}$, erfüllt also die obige Differentialgleichung. Es handelt sich hier um *exponentielles Wachstum*.

Ist dies die einzige Lösung der Differentialgleichung? Nein, man kann die Funktion $x \mapsto e^{\alpha x}$ mit einer beliebigen Konstante multiplizieren, also etwa $x \mapsto Ce^{\alpha x}$ mit einer reellen Zahl C betrachten. Damit kann man etwa die Lösung an einen Startwert anpassen. Ist C die Anzahl der Zellen zum Zeitpunkt $x = 0$, dann gilt $y(0) = Ce^{\alpha \cdot 0} = C$, d.h. die Funktion $x \mapsto Ce^{\alpha x}$ beschreibt das Wachstum, wenn zur Zeit $x = 0$ die Anzahl der Zellen C beträgt.

Analog kann man mit dieser Gleichung radioaktive Zerfallsprozesse beschreiben, wenn man ein negatives α wählt (Dann ist $-\alpha$ die Zerfallsrate). In den Wirtschaftswissenschaften ist die Gleichung unter dem Stichwort *stetige Verzinsung* von Bedeutung.

1. Begriffsbestimmung und Beispiele

Definition 1.0.1. Eine Gleichung, in der eine Funktion, ihre Ableitungen und eventuell die Variablen der Funktion, vorkommen, nennt man *Differentialgleichung*. Eine *Lösung* der Differentialgleichung ist eine Funktion, für die die Gleichung erfüllt ist.

Beispiel 1.0.2. Eine Differentialgleichung (in typischer Kurzschreibweise) ist etwa

$$2y'x = y.$$

Dies ist eine Differentialgleichung für eine Funktion $x \mapsto y(x)$ in einer reellen Variablen, nämlich x . Es kommen in der Gleichung die Funktion y , ihre Ableitung y' und ihre Variable x vor.

Noch einmal ausführlich, der Klarheit halber: Die oben angegebene Gleichung $2y'x = y$ ist ein *Kurzschreibweise* für die folgende Aufgabe: *Gesucht ist eine reelle Funktion $x \mapsto f(x)$, so dass*

$$2 \cdot f'(x) \cdot x = f(x)$$

für jedes x aus dem Definitionsbereich gilt.

Eine Lösung dieser Differentialgleichung ist die Funktion

$$f(x) = \sqrt{x}, \quad f: (0, \infty) \rightarrow \mathbb{R},$$

denn es gilt $f'(x) = \frac{1}{2}x^{-\frac{1}{2}}$ und somit in der Tat

$$2 \cdot f'(x) \cdot x = 2 \cdot \frac{1}{2}x^{-\frac{1}{2}} \cdot x = x^{\frac{1}{2}} = f(x).$$

Bemerkung 1.0.3. Die Lösung der obigen Differentialgleichung ist nicht eindeutig bestimmt. Z.B. ist für jedes $c \in \mathbb{R}$ auch die Funktion $x \mapsto c \cdot x^{\frac{1}{2}}$ eine Lösung; insbesondere ist auch die Nullfunktion $x \mapsto 0$ eine Lösung der Differentialgleichung.

Die Definition 1.0.1 am Anfang dieses Abschnitts ist bewusst vage gehalten. Das liegt daran, dass es viele verschiedene Typen von Differentialgleichungen gibt. Mit den folgenden Begriffsbestimmungen kann man eine vorliegende Differentialgleichung zumindest grob in gewisse Kategorien einordnen.

Definition 1.0.4. Man unterscheidet bei Differentialgleichungen grundsätzlich die beiden folgenden Typen:

- (i) Ist die gesuchte Funktion in einer Variablen, dann handelt es sich um eine *gewöhnliche Differentialgleichung*.
- (ii) Ist eine Funktion in mehreren Variablen gesucht (dann treten partielle Ableitungen in der Gleichung auf), so handelt es sich um eine *partielle Differentialgleichung*.

Außerdem unterscheidet man Differentialgleichungen nach ihrer Ordnung: Die höchste auftretende Ableitungsordnung in der Gleichung heißt *Ordnung* der Gleichung.

Das obige Beispiel $2y'x = y$ ist also eine *gewöhnliche Differentialgleichung erster Ordnung*. Wir werden im folgenden ausschließlich solche gewöhnlichen Differentialgleichungen erster Ordnung betrachten.

Beispiele 1.0.5. Einige weitere Beispiele und Bemerkungen.

- (i) Die als Beispiel am Anfang dieses Abschnitts verwendete Differentialgleichung kann man, für $x \neq 0$ auf in der Form

$$y' = \frac{y}{2x}$$

hinschreiben. Wir wollen nun $x > 0$ annehmen, d.h. die Gleichung auf der rechten Halbebene der xy -Ebene betrachten. Dann lässt sich die Gleichung als *Steigungsfeld* geometrisch interpretieren: Gesucht ist eine Funktion $f(x)$ für positive x , deren Graph an jedem seiner Punkte $(x, y) = (x, f(x))$ die vorgeschriebene Tangentensteigung $\frac{y}{2x}$ hat.

Dies ermöglicht eine zeichnerische Näherungslösung: Man zeichnet an genügend vielen Punkten kurze Geradenstücke mit der entsprechenden Steigung ein und versucht, zeichnerisch einen ungefähr passenden Funktionsgraph durch das Steigungsfeld zu legen.

- (ii) Betrachten wir noch einmal die Differentialgleichung $y' = y$. Sie hat die Lösung $x \mapsto e^x$. Aber auch $x \mapsto e^{x+a}$ ist eine Lösung der Differentialgleichung, denn auch für diese Funktion gilt, dass sie gleich ihrer eigenen Ableitung ist.

Eine gewöhnliche Differentialgleichung heißt *autonom*, wenn Sie von der Form $y'(x) = F(y(x))$ ist, d.h. die rechte Seite hängt nicht von x ab. Dann ist für jede Lösung $y(x)$ auch die entlang der x -Achse verschobene Funktion $y(x+a)$ eine Lösung.

Die Differentialgleichung hat noch eine weitere Eigenschaft, sie ist *linear*, d.h. die rechte Seite ist linear in $y(x)$. In diesem Fall gilt, dass für jede Lösung $y(x)$ auch $cy(x)$ für ein beliebiges $c \in \mathbb{R}$ wieder eine Lösung ist und dass für je zwei Lösungen auch deren Summe eine Lösung ist.

Wegen $e^{x+a} = e^x e^a$ bewirkt in diesem Fall das Verschieben längs der x -Achse nichts anderes als die Multiplikation mit einer positiven Konstante.

- (iii) Ein Beispiel für eine gewöhnliche Differentialgleichung zweiter Ordnung ist

$$y'' = -y.$$

Auch hier lassen sich Lösungen leicht erraten. Die Sinus- und die Kosinusfunktion sind etwa Lösungen, denn es gilt

$$\sin''(x) = \cos'(x) = -\sin(x)$$

bzw.

$$\cos''(x) = -\sin'(x) = -\cos(x).$$

Die Gleichung ist ebenfalls autonom und linear.

- (iv) Ist die Differentialgleichung von der Form $y' = F(x)$ für eine Funktion F , d.h. kommt $y(x)$ auf der rechten Seite der Differentialgleichung gar nicht vor, dann handelt es sich nicht um eine "echte" Differentialgleichung, sondern man erhält eine Lösung einfach durch Bilden einer Stammfunktion,

$$y(x) := \int_{x_0}^x F(x) dy + c.$$

Analog bezeichnet man auch das Lösen von allgemeinen Differentialgleichungen als *Integration* der Differentialgleichung. Wie beim Ermitteln von Stammfunktionen sind auch die Lösungen von Differentialgleichungen nicht immer in geschlossener Form darstellbar, Näherungsverfahren haben daher eine große Bedeutung.

Wir haben bereits Beispiele dafür gesehen, dass die Lösung einer Differentialgleichung durch die Gleichung selbst nicht eindeutig bestimmt ist. Bei der Gleichung $y' = y$ etwa haben wir die Schar von Lösungen $x \mapsto ce^x$, $c \in \mathbb{R}$. Es ist klar, dass man eine Zusatzbedingung stellen muss, um eine Lösung eindeutig festzulegen. Man kann zum Beispiel eine sogenannte *Anfangswertbedingung* stellen, das bedeutet, man fordert, dass die Lösung bei einem bestimmten Wert x_0 der Variablen x den *Anfangswert* y_0 annimmt. Eine gewöhnliche Differentialgleichung zusammen mit einer Anfangswertbedingung wird *Anfangswertproblem* genannt.

Ein Beispiel für ein Anfangswertproblem ist

$$y' = y, \quad y(0) = 2.$$

Eine Lösung für dieses Anfangswertproblem ist die Funktion $x \mapsto 2e^x$.

Existenz und Eindeutigkeit Inwieweit ist diese Lösung nun eindeutig? Und existiert immer eine Lösung? Ohne genau darauf einzugehen, halten wir als Merkregel fest: Ein Anfangswertproblem

$$y' = F(x, y), \quad y(x_0) = y_0,$$

besitzt immer eine eindeutig bestimmte Lösung, falls die rechte Seite, d.h. die Abbildung $F(x, y)$ stetig von x und y abhängig ist und differenzierbar in der zweiten Variablen y . Für diese Lösung gibt es ein *maximales Definitionsintervall*, das nicht unbedingt gleich ganz $\mathbb{R}$ sein muss. Lösungen können nämlich in endlicher Zeit gegen unendlich gehen (oder den Definitionsbereich der rechten Seite verlassen).

2. Trennung der Variablen

Wir werden nun ein Verfahren kennenlernen, mit dem man in vielen Fällen Lösungen für gewöhnliche Differentialgleichungen erster Ordnung bestimmen kann. Die Voraussetzung dafür ist, dass die rechte Seite der Differentialgleichung mit *getrennten Variablen* vorliegt, das bedeutet, wir betrachten ein Anfangswertproblem

der folgenden Form.

$$y' = f(x) \cdot g(y), \quad y(x_0) = y_0. \quad (*)$$

Hier sind f und g auf Intervallen definierte reelle Funktionen, x_0 ein Punkt im Definitionsbereich von f , y_0 ein Punkt im Definitionsbereich von g . Wir wollen zusätzlich annehmen, dass g auf seinem Definitionsbereich nicht den Wert Null annimmt.

Das besondere an der Form der Gleichung ist, dass auf der rechten Seite ein Produkt steht, wobei der eine Faktor nur von x abhängt, der andere nur von y . Eine solche Differentialgleichung wird auch *separierbar* genannt.

Lösungsverfahren: Trennung der Variablen

Wir setzen

$$F(x) := \int_{x_0}^x f(t) dt$$

und

$$H(y) := \int_{y_0}^y \frac{1}{g(s)} ds.$$

Es ist $y(x)$ genau dann eine Lösung von $(*)$ mit Anfangswertbedingung $y(x_0) = y_0$, wenn

$$H(y(x)) = F(x)$$

für alle x aus dem Definitionsbereich von y gilt. Ist dann H^{-1} eine differenzierbare Umkehrfunktion von H , so ist

$$y(x) := H^{-1}(F(x))$$

die gewünschte Lösung. Dieses Lösungsverfahren wird *Trennung der Variablen* genannt.

BEGRÜNDUNG. Angenommen, es gilt $H(y(x)) = F(x)$, dann folgt

$$H(y(x))' = F'(x) \iff \frac{y'(x)}{g(y(x))} = f(x),$$

was dazu äquivalent ist, dass die Differentialgleichung $y'(x) = f(x) \cdot g(y(x))$ erfüllt ist. Außerdem ist die Anfangswertbedingung erfüllt, denn falls insbesondere $H(y(x_0)) = F(x_0) = 0$ gilt, dann folgt, dass $y(x_0) = y_0$, denn H ist streng monoton. $\square$

Beispiele 2.0.6 (für Trennung der Variablen). (i) Löse das Anfangswertproblem

$$y' = y, \quad y(0) = 1.$$

Wir setzen $f(x) = 1$ und $g(y) = y$. Damit erhalten wir

$$F(x) = \int_{x_0}^x 1 \, dt = x - x_0 = x,$$

$$H(y) = \int_{y_0}^y \frac{1}{s} \, ds = \ln(y) - \ln(1) = \ln(y).$$

Eine differenzierbare Umkehrfunktion von H ist gegeben durch

$$H^{-1}(x) = \exp(x).$$

Die gesuchte Lösung ist somit

$$y(x) = H^{-1}(F(x)) = e^x.$$

(ii) Die Gleichung

$$y' = \frac{y}{2x}$$

vom Anfang dieses Abschnitts mit der Anfangsbedingung

$$y(1) = 1.$$

Wir setzen $f(x) = \frac{1}{2x}$ und $g(y) = y$ und erhalten

$$F(x) = \int_1^x \frac{1}{2} t^{-1} = \left[\frac{1}{2} \ln(t) \right]_1^x = \frac{1}{2} \ln(x) - \frac{1}{2} \ln(1) = \frac{1}{2} \ln(x),$$

$$H(y) = \int_1^y \frac{1}{s} \, ds = \ln(y) - \ln(1) = \ln(y).$$

Eine differenzierbare Umkehrfunktion von H ist gegeben durch

$$H^{-1}(x) = e^x.$$

Die gesuchte Lösung ist somit

$$y(x) = H^{-1}(F(x)) = e^{\frac{1}{2} \ln(x)} = x^{\frac{1}{2}} = \sqrt{x}.$$

(iii) Wir betrachten das Anfangswertproblem

$$y' = y^2 \cos(x), \quad y(0) = 1.$$

Setze $f(x) := \cos(x)$ und $g(y) := y^2$. Damit erhalten wir

$$F(x) := \int_0^x \cos(t) \, dt = [\sin(t)]_0^x = \sin(x)$$

und

$$H(y) := \int_1^y t^{-2} \, dt = [-t^{-1}]_1^y = -y^{-1} + 1.$$

die gesuchte Lösung erhält man, wenn man

$$H(y(x)) = 1 - \frac{1}{y(x)} = \sin(x) = F(x)$$

nach $y(x)$ auflöst:

$$y(x) = \frac{1}{1 - \sin(x)}.$$

Index

- ε -Ball, 22
- äquivalent, 17
- Abbildung, 23
 - identische, 24
 - konstante, 24
 - lineare, 25
- abgeschlossen, 87
- abhängig, linear, 43
- Ableitung, 63
 - partielle, 66
 - partiell von höherer Ordnung, 67
- Ableitungsregeln, 64
- Absolutglied, 30
- Abstand, 22
- Anfangswert, 101
- Anfangswertbedingung, 101
- Anfangswertproblem, 101
- Assoziativgesetz, 10
- Ball, 22
- Basis, 44
- beschränkt, 48, 87
- Betrag
 - einer komplexen Zahl, 20
 - einer reellen Zahl, 17
- bijektiv, 24
- Bild, 23
- Binomialkoeffizient, 13
- Bisektion, 94
- Bruch
 - echter, 8
 - gemeiner, 9
 - gemischter, 8
- Cobb-Douglas-Funktion, 80
- definit, 75
 - negativ, 75
 - positiv, 75
- Definitionsbereich, 23
- Determinante, 75
- Dezimalbrüche
 - periodische, 9
 - unendliche, 9
- Diagonalmatrix, 75
- Differential
 - totales, 69
- Differentialgleichung, 98
 - autonome, 100
 - gewöhnliche, 99
 - Lösung einer, 98
 - lineare, 100
 - partielle, 99
 - separierbare, 102
- Differenz
 - von Mengen, 6
- differenzierbar, 63
- Differenzierbarkeit, 63
 - partielle, 66
 - totale, 68
- Dimension, 45
- Dimensionsformel, 45
- Diskriminante, 30
- Distributivgesetz, 10
- divergent, 47
- Divergenz, 47
 - bestimmte, 47

- Dreieck
 - Pascalsches, 14
- Dreiecksmatrix, 76
- Dreiecksungleichung, 18, 22
- Einheit
 - imaginäre, 19
- Einheitsmatrix, 26
- Einheitsvektor, 22
- Element, 5
- Eulersche Zahl, 54
- Exponentialfunktion, 54
- Exponentialreihe, 53
- Extensionalität, 7
- Extremalstelle
 - lokale, 72, 81
- Fakultät, 13
- Fallunterscheidung, 18
- fast alle, 47
- Fibonacci-Folge, 46
- Flachpunkt, 73
 - unter Nebenbedingung, 81
- Folge, 46
 - arithmetische, 46
 - beschränkte, 48
 - geometrische, 46
 - harmonische, 46
 - monotone, 48
 - rekursiv definierte, 46
- Folgenglieder, 46
- Folgenstetigkeit, 56
- Formel
 - Eulersche, 54
- Fortsetzung
 - stetige, 57
- Funktion, 23
 - Cobb-Douglas, 80
 - konstante, 24
 - rationale, 49
- Funktionalgleichung
 - der Exponentialfunktion, 54
- Funktionsgrenzwert, 58
- Gaußverfahren, 35

- Gleichungssystem
 - homogenes, lineares, 42
 - inhomogenes, lineares, 42
 - lineares, 32
 - zugehöriges homogenes, 42
- Glied
 - einer Folge, 46
 - einer Reihe, 50
 - konstantes, 30
 - kubisches, 30
 - lineares, 30
 - quadratisches, 30
- Grad, 32
 - eines Polynoms, 29
- Gradient, 73
- Grenzwert, 46
- Hauptminoren, 78
- Heronverfahren, 96
- Imaginärteil, 19
- Induktion
 - vollständige, 15
- injektiv, 24
- Inklusion, 11
- Integration
 - einer Differentialgleichung, 101
- Intervall, 17
- Intervallgrenze, 17
- Intervallschachtelung, 9
- Jägerzaunregel, 77
- Jacobimatrix, 69
- Körper, 10
- Körperaxiome, 10
- Kürzen, 8
- Kapitalentwicklung, 51
- Kardinalität, 7
- Kern, 42
- Kettenregel
 - eindimensional, 64
- Koeffizient, 29, 32
- Kommutativgesetz, 10
- kompakt, 87

- Komplement, 87
- konstant, 24
- konvergent, 46
- Konvergenz, 46
- Länge
 - eines Vektors, 21
- Lösung, 32
- Lösungsmenge, 32
- Lagrange-Multiplikator, 81
- Laufindex, 12
- Leitkoeffizient, 29
- Limes, 46
 - linksseitiger, 58
 - rechtsseitiger, 58
- linear, 25
- linear abhängig, 43
- linear unabhängig, 43
- Linearfaktor, 31
- Linearkombination, 41
- Logarithmus
 - natürlicher, 60
 - zur Basis b , 61
- Mächtigkeit, 7
- Majorante, 53
- Majorantenkriterium, 53
- Matrix, 27
 - darstellende, 25
 - erweiterte, 37
 - inverse, 39
 - quadratische, 38, 74
 - symmetrische, 74
 - transponierte, 29
- Matrixmultiplikation, 27
- Maximalstelle
 - lokale, 72, 80
- Menge, 5
 - n -elementige, 7
 - leere, 6
- Minimalstelle
 - lokale, 72, 81
- Monome, 32
- monoton, 24, 48
- Newton-Verfahren
 - eindimensional, 95
 - mehrdimensional, 96
- Norm, 21
- Null hoch Null, 62
- Nullfolge, 46
- Nullpolynom, 29
- Nullstelle, 30
- Nutzenfunktionen, 80
- offen, 66
- Ordnung, 99
- Partialsumme, 50
- Periode, 9
- Permutation, 13
- Pivotelemente, 35
- Pivotvariable, 35
- Polarkoordinaten, 20
- Polynom, 29
 - komplexes, 30
 - konstantes, 30
 - kubisches, 30
 - lineares, 30
 - multivariates, 31
 - quadratisches, 30
 - reelles, 30
- Polynomdivision, 31
- Polynomfunktion, 29
- Potenz, 61
- Potenzmenge, 7
- Produkt
 - kartesisches, 7
 - leeres, 12
- Produktionsfunktionen, 80
- Produktzeichen, 12
- Quadratwurzel
 - einer komplexen Zahl, 30
- Rückwärtseinsetzen, 36
- Randpunkt, 66
- Rang, 45
- Rate, 51
- Realteil, 19

- Reihe, 50
 - alternierende harmonische, 53
 - geometrische, 50
 - harmonische, 52
- Reihenglied, 50
- Rente, 51
- Rentenrechnung, 51
- Restglied
 - Integral-, 92
- Sarrus
 - Regel von, 77
- Sattelpunkt, 78
 - unter Nebenbedingung, 85
- Schmiegeparabel, 92
- Schnittmenge, 6
- separierbar, 102
- Skalar, 21
- Skalarprodukt, 22
- Spaltenvektor, 27
- Spann, 41
- Standardbasisvektor, 25
- Steigung, 62
- Steigungsfeld, 99
- stetig, 55
- stetig fortsetzen, 57
- Stetigkeit, 55
 - in einem Punkt, 55
- Stufenform, 35
- Subniveau, 87
- Substitutionsmethode, 83
- Summationsvariable, 12
- Summe
 - geometrische, 50
 - leere, 12
- Summenzeichen, 12
- surjektiv, 24
- Taylor
 - Satz von, 92
- Taylorpolynom, 90, 91
- Taylorreihe, 92
- Teilmenge, 7
 - echte, 7
 - offene, 66
- Totalgrad, 32
- Tripel, 7
- Tupel, 7
- Umgebung, 22
- Umkehrabbildung, 24
- umkehrbar, 24
- unabhängig, linear, 43
- Unbekannte, 32
- Ungleichung, 17
- Untermannigfaltigkeit, 82
- Unterraum
 - affiner, 43
 - linearer, 41
- Untervektorraum, 41
- Urbild, 23
- Variable, 32
- Variablen
 - getrennte, 101
- Vektor, 21
- Vektorfeld, 73
- Vereinigung, 6
- Verzinsung
 - stetige, 98
- Vollständigkeit, 9
- Wachstum
 - exponentielles, 97
- Wert, 23
- Wertemenge, 23
- Winkel
 - zwischen zwei Vektoren, 22
- Wurzel, 61
- Wurzelziehen, Babylonisches, 96
- Zahl
 - Eulersche, 54
 - ganze, 8
 - irrationale, 9
 - komplex-konjugierte, 19
 - komplexe, 18
 - rationale, 8
 - reelle, 9
- Zahlenebene

- Gaußsche, 20
- Zahlweise
 - nachschüssige, 51
 - vorschüssige, 51
- Zeilenstufenform, 38
- Zeilenumformungen
 - elementare, 38
- Zeilenvektor, 27
- Zielmenge, 23
- Zinsfaktor, 51
- Zwischenwertsatz, 60