
1 Die reellen Zahlen

1. Ziele des Mathematikstudiums:

Die Studierenden sollen lernen,

- präzise und logisch zu denken,
- komplexe Strukturen schnell und gründlich zu erfassen,
- Dinge kritisch zu hinterfragen und niemals dem ersten Augenschein zu vertrauen,
- Probleme systematisch zu analysieren, erfolgsorientiert zu denken, Lösungen zielgerichtet selbst zu erarbeiten und dabei Teamfähigkeit zu entwickeln
- und sich mündlich und schriftlich klar und verständlich auszudrücken

Sie sollen am Ende ihres Studiums

- grundlegende mathematische Denkstrukturen beherrschen,
- einen Überblick über die wichtigsten mathematische Disziplinen und ihre Zusammenhänge gewonnen haben
- und erfolgreich mathematische Probleme bearbeiten können.

Dies alles ist nur zu erreichen, wenn man immer auf Neues gespannt ist, den Dingen auf den Grund gehen will und bereit ist, sich mit schwierigen Fragen leidenschaftlich auseinanderzusetzen.

2. Vorlesungen, Übungen und all das ...

Die **Vorlesung** findet mittwochs und freitags von 10 bis 12 in Hörsaal 12 statt. Alles, was Sie zur Analysis wissen müssen, wird in der Vorlesung erzählt, erklärt und – zum Mitschreiben – an die Tafel geschrieben. Die Inhalte sollten regelmäßig nachgearbeitet werden, denn es geht sehr schnell voran. Andernfalls hat man keine Chance. Der Besuch der Vorlesung wird zwar nicht kontrolliert, aber vorausgesetzt.

Jede Woche gibt es ein Übungsblatt mit Aufgaben, die schriftlich bearbeitet und abgegeben werden müssen. Die Aufgaben werden von studentischen Hilfsassistenten korrigiert und im Rahmen einer 2-stündigen **Übung** besprochen. Der Besuch der (in kleinen Gruppen organisierten) Übungen ist verpflichtend. Die Teilnahme an der Klausur am Ende des Semesters (am 3.3.2009) ist nur für diejenigen möglich, die genügend viele Aufgaben erfolgreich bearbeitet haben. Studierende im Kombi-BA müssen eine mündliche Prüfung absolvieren.

Das 2-stündige **Tutorium** wird in zwei größeren Gruppen abgehalten. Es ist ein Angebot, das beim Verständnis der Vorlesung und beim Lösen der Aufgaben helfen soll. Die Teilnahme wird allen Studierenden empfohlen, in einigen Studiengängen ist sie verpflichtend.

3. Modulschein:

Die Kriterien zum Erwerb des Modulscheines variieren von Studiengang zu Studiengang. Am Ende steht das Bestehen einer Klausur oder einer mündlichen Prüfung. Die Berechtigung zur Teilnahme an der Klausur wird durch das erfolgreiche Bearbeiten von Übungsaufgaben erworben.

Die Klausur bzw. mündliche Prüfung muss auf jeden Fall bestanden werden. Die Note kann durch besondere Leistungen in den Übungen verbessert werden.

1.1 Sprachregelungen

Es geht um eine Einführung in die Sprache der Mathematik, ausgehend von Schulkenntnissen.

Zunächst werden die verschiedenen Zahlenbereiche vorgestellt. Es folgt eine kurze Einführung in die Mengenlehre und dann eine Klärung der Begriffe bei Zahlen, Rechenoperationen und algebraischen Termen.

Mit den Grundbegriffen der formalen Logik werden Gleichungen und Ungleichungen behandelt.

Als **natürliche Zahlen** bezeichnet man die Zahlen $1, 2, 3, \dots$

Erweitert man die natürlichen Zahlen um die **Null** und die Zahlen $-1, -2, -3, \dots$, so spricht man von **ganzen Zahlen**.

Brüche von ganzen Zahlen wie etwa $\frac{1}{2}$ oder $-\frac{7}{12}$ bezeichnet man als **rationale Zahlen**.

Was gibt es noch? Rationale Zahlen kann man auch als endliche oder periodische unendliche Dezimalzahlen schreiben. Lässt man beliebige (unendliche) Dezimalbrüche zu, so spricht man von **reellen Zahlen**. Das liefert auch irrationale Zahlen wie $\sqrt{2}$ oder die Kreiszahl π .

Will man nicht nur über einzelne Zahlen sprechen, sondern auch über Gesamtheiten von Zahlen, so benötigt man den Mengenbegriff.

Nach Cantor versteht man unter einer **Menge** M die Zusammenfassung von wohlunterschiedenen (mathematischen) Objekten zu einem neuen Ganzen. Die dabei

zusammengefassten Objekte nennt man die **Elemente** von M , die Menge ist wieder ein mathematisches Objekt. Ist x ein Element der Menge M , so schreibt man: $x \in M$. Ist dies nicht der Fall, so schreibt man: $x \notin M$.

Mengen mit nur wenigen Elementen kann man beschreiben, indem man alle ihre Elemente angibt, etwa in der Form

$$M = \{2, 3, 5, 7, 11, 13, 17, 19\}.$$

Bei größeren oder gar unendlichen Mengen geht das nicht. Aber gerade dafür hat Cantor den Mengenbegriff eingeführt.

Zunächst folgen hier Symbole für die verschiedenen Zahlenbereiche:

- $\mathbb{N}$ ist die Menge der natürlichen Zahlen,
- $\mathbb{Z}$ die Menge der ganzen Zahlen,
- $\mathbb{Q}$ die Menge der rationalen Zahlen und
- $\mathbb{R}$ die Menge der reellen Zahlen.

Die oben angegebene Menge $M = \{2, 3, 5, 7, 11, 13, 17, 19\}$ ist die Menge aller Primzahlen, die kleiner als 20 sind, und deshalb kann man sie auch durch genau diese Eigenschaft beschreiben:

$$M = \{n \in \mathbb{N} : n \text{ ist Primzahl und kleiner als } 20\}.$$

Zwei Mengen heißen **gleich**, wenn sie die gleichen Elemente besitzen. So ist z.B.

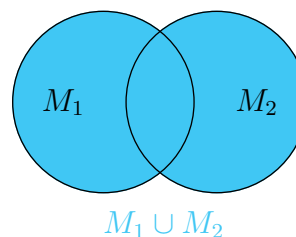
$$\{1, 2, 3\} = \{1, 1, 2, 2, 2, 3, 3, 3, 3\} \quad \text{und} \quad \{x \in \mathbb{R} : 2x + 1 = -5\} = \{-3\}.$$

Eine Menge T heißt **Teilmenge** einer Menge M , falls jedes Element von T auch Element von M ist. Man schreibt dann: $T \subset M$.

Zum Beispiel ist $\{1, 2, 3\} \subset \{1, 2, 3, 4, 5\}$ und $\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R}$.

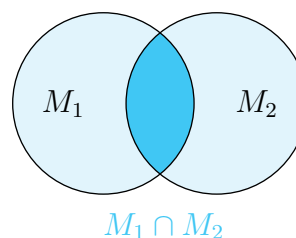
Fasst man die Elemente einer Menge M_1 und einer Menge M_2 zu einer neuen Gesamt-Menge zusammen, so bildet man die **Vereinigungsmenge**

$$M_1 \cup M_2 = \{x : x \in M_1 \text{ oder } x \in M_2\}.$$



Betrachtet man die Menge genau derjenigen Elemente, die in zwei Mengen M_1 und M_2 zugleich enthalten sind, so bildet man ihre **Schnittmenge**

$$M_1 \cap M_2 = \{x : x \in M_1 \text{ und } x \in M_2\}.$$



1.1. Beispiel

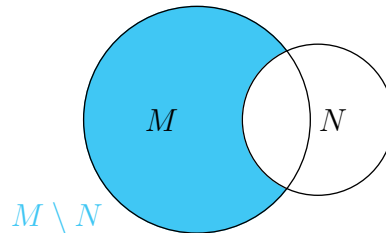
Mit $\mathbb{N}_0$ bezeichnet man die Menge $\mathbb{N} \cup \{0\}$.

Es ist $\mathbb{Z} = \mathbb{N} \cup \{0\} \cup \{n : -n \in \mathbb{N}\}$ und $\{1, 3, 5, 7, 9\} \cap \{3, 6, 9, 12\} = \{3, 9\}$.

Sind M und N zwei Mengen, so nennt man

$$M \setminus N = \{x \in M : x \notin N\}$$

die **Differenzmenge** von M und N .



Wir wollen uns noch mit etwas elementarer Algebra befassen.

Die einfachsten mathematischen Objekte, mit denen wir es zu tun haben, sind **Zahlen**. Wir unterscheiden zwischen positiven und negativen (reellen) Zahlen und schreiben jeweils $x > 0$ oder $x < 0$. Zum Beispiel ist $37.2 > 0$ und $-2/3 < 0$.

Häufig werden Zahlen durch **Variablen** beschrieben, dafür verwenden wir lateinische oder griechische Buchstaben. Steht a für eine reelle Zahl, so kann man a nicht ansehen, ob die Zahl positiv oder negativ ist. Man kann aber ein „Vorzeichen“ davor setzen: $+a$ oder $-a$. Dieses Vorzeichen sagt nichts darüber aus, ob wir es mit einer positiven oder negativen Zahl zu tun haben.

Wir kennen zwei „Rechenoperationen“ in $\mathbb{R}$, die **Addition** und die **Multiplikation**. Mit der Addition ist die Zahl **Null** verbunden: Addiert man die Null, so ändert sich nichts. Mit der Multiplikation ist die **Eins** verbunden: Multipliziert man mit der Eins, so ändert sich nichts. Ausgehend von diesen beiden neutralen Zahlen gewinnt man zu jeder anderen Zahl eine „inverse Zahl“:

1. Zu jeder Zahl a gibt es das „Negative“ $-a$, mit $a + (-a) = 0$.
2. Zu jeder Zahl a gibt es das „(multiplikative) Inverse“ a^{-1} , mit $a \cdot a^{-1} = 1$. Dabei darf a nicht die Null sein, denn dann müsste $a \cdot a^{-1} = 0$ sein.

Subtraktion und **Division** sind in diesem Sinne keine neuen Rechenoperationen, vielmehr ist

$$a - b := a + (-b) \quad \text{und} \quad a/b := a \cdot b^{-1}.$$

Potenzen gewinnt man durch fortgesetztes Multiplizieren: $a^n := \underbrace{a \cdot a \cdots a}_{n\text{-mal}}$.

Algebraische Terme sind zusammengesetzte mathematische Objekte, wie z.B.

$$\alpha, \quad x^2, \quad a + b, \quad \frac{x - y}{x + y}, \quad -3 \quad \text{oder} \quad \frac{3x(y^2 - z^2)}{5x + 7y - 2z}.$$

Von völlig anderer Natur sind logische Aussagen.

Aussagen sind sinnvolle mathematische Sätze, die (im Gegensatz zu Aussagen des Alltags) eindeutig in „wahre“ und „falsche“ Aussagen eingeteilt werden können. Vor Jahrtausenden wurde entschieden, dass es in der Logik keinen anderen Wahrheitswert gibt. Beispiele sind etwa die Aussagen

- „ $x^2 - x + 41$ ist für jedes $x \in \mathbb{N}$ eine Primzahl“ (falsch),
- „Die Menge $\mathbb{N}$ besteht aus unendlich vielen Elementen“ (wahr),
- „Wenn x gerade ist, dann ist $3x$ durch 6 teilbar“ (wahr).

Typische Aussagen in der Mathematik sind Gleichungen und Ungleichungen, z.B.

$$\begin{aligned}(a+b)^2 &= a^2 + 2ab + b^2, \\ (a+b)(a-b) &= a^2 - b^2, \\ x^{n+m} &= x^n \cdot x^m\end{aligned}$$

oder etwa

$$\begin{aligned}3 &< 5, \\ (x+y)^2 &\geq x^2 + 2xy, \\ u &\neq v.\end{aligned}$$

Zu jeder Aussage $\mathcal{A}$ gibt es die **logische Verneinung** „nicht $\mathcal{A}$ “, (in Zeichen: $\neg \mathcal{A}$) die genau dann wahr ist, wenn $\mathcal{A}$ falsch ist. Oft gibt es ein besonderes Symbol dafür. Die Verneinung der Aussage „ $x \in M$ “ ist die Aussage „ $x \notin M$ “.

Sind $\mathcal{A}, \mathcal{B}$ zwei Aussagen, so kann man sie zur **Disjunktion** „ $\mathcal{A}$ oder $\mathcal{B}$ “ (in Zeichen: $\mathcal{A} \vee \mathcal{B}$) bzw. zur **Konjunktion** „ $\mathcal{A}$ und $\mathcal{B}$ “ (in Zeichen: $\mathcal{A} \wedge \mathcal{B}$) verknüpfen. Der Wahrheitswert der zusammengesetzten Aussagen ergibt sich nach festen Regeln aus den Wahrheitswerten der einzelnen Aussagen. Am einfachsten lässt sich das mit Hilfe von Wahrheitstafeln beschreiben:

$\mathcal{A}$	$\mathcal{B}$	$\mathcal{A} \vee \mathcal{B}$	$\mathcal{A}$	$\mathcal{B}$	$\mathcal{A} \wedge \mathcal{B}$
w	w	w	w	w	w
w	f	w	w	f	f
f	w	w	f	w	f
f	f	f	f	f	f

Typischerweise treten „oder“-Verknüpfungen auf, wenn man „und“-Verknüpfungen verneint.

Dass zwei Aussagen logisch das Gleiche bedeuten, heißt nur, dass sie den gleichen Wahrheitswert besitzen. Sind Variable im Spiel, so muss das für alle eingesetzten Werte gelten. Diese logische Gleichheit von Aussagen wird durch das Äquivalenzzeichen „ $\iff$ “ ausgedrückt. Man kann es benutzen, um „logische Gesetze“ zu formulieren, wie etwa die Verneinungsregeln von de Morgan (einem Logiker des 19. Jahrhunderts):

$$\neg(\mathcal{A} \wedge \mathcal{B}) \iff (\neg \mathcal{A}) \vee (\neg \mathcal{B}),$$

$$\neg(\mathcal{A} \vee \mathcal{B}) \iff (\neg \mathcal{A}) \wedge (\neg \mathcal{B}),$$

$$\neg \neg \mathcal{A} \iff \mathcal{A}.$$

Wir wollen uns jetzt mit dem logischen Schließen beschäftigen. Üblicherweise geht man in der Mathematik von einem Axiomensystem aus, das aus einfachen, oft als bekannt und selbstverständlich erachteten Aussagen besteht. Daraus werden nach den Regeln der formalen Logik Schlüsse gezogen und nach und nach immer tiefere Aussagen hergeleitet. Wichtigstes Hilfsmittel ist dabei die **Implikation** oder **logische Folgerung**:

„wenn $\mathcal{A}$, dann $\mathcal{B}$ “ (in Zeichen: $\mathcal{A} \implies \mathcal{B}$).

Auf eine naive Weise ist jedem klar, was damit gemeint ist. Aber das reicht nicht. Es muss eine Vorschrift geben, wie der Wahrheitswert der Implikation $\mathcal{A} \implies \mathcal{B}$ aus den Wahrheitswerten für $\mathcal{A}$ und $\mathcal{B}$ hergeleitet werden kann. Die dabei auftretende Problematik soll an einem Beispiel demonstriert werden. Wir betrachten die Aussage

$$((x \in \mathbb{Z}) \wedge (x > 10)) \implies x^2 > 25.$$

Eigentlich ist dies keine Aussage, sondern eine sogenannte Aussageform. Erst wenn man für die Variable x eine Zahl eingesetzt hat, kann man den Wahrheitswert ermitteln und tatsächlich von einer Aussage sprechen. Das führt zu einer Fallunterscheidung:

1. Ist die **Prämisse** $((x \in \mathbb{Z}) \wedge (x > 10))$ wahr (z.B. $x = 11$), so ist mit Sicherheit auch $x^2 > 25$.
2. Ist die Prämisse falsch, so gibt es wiederum mehrere Möglichkeiten.
 - (a) Ist z.B. $x = 8$, so ist $x^2 = 64$, die Aussage $x^2 > 25$ also wahr.
 - (b) Ist $x = 4$, so ist $x^2 = 16$ kleiner als 25 und $x^2 > 25$ falsch.

Es scheint möglich zu sein, aus falschen Prämissen beliebige Schlüsse zu ziehen. Fest steht nur: Wenn die gefolgerte Aussage falsch ist und beim Beweis alles richtig gemacht wurde, dann muss schon die Prämisse falsch gewesen sein. Deshalb versteht man unter der Implikation $\mathcal{A} \implies \mathcal{B}$ einfach die Aussage

$$\mathcal{B} \vee \neg \mathcal{A}.$$

Das ergibt folgende Wahrheitstafel:

$\mathcal{A}$	$\mathcal{B}$	$\mathcal{A} \implies \mathcal{B}$
w	w	w
w	f	f
f	w	w
f	f	w

Überraschenderweise zeigt sich, dass man aus einer falschen Aussage alles folgern kann. Für die Mathematik hat das eine besondere Bedeutung, die an Hand der Teilmengenbeziehung demonstriert werden kann.

Dass $N \subset M$ ist, kann man durch die Aussage

$$\forall x : (x \in N \implies x \in M)$$

ausdrücken. Dabei steht das Symbol $\forall$ für „für alle“ oder „für jedes“.

Wir betrachten die Menge $\{x \in \mathbb{Z} : (x - 1)^2 = 5\}$. Es gibt keine ganze Zahl x , so dass das Quadrat von $x - 1$ die Zahl 5 ergibt. Also haben wir eine Menge angegeben, die kein Element besitzt. Man spricht von der **leeren Menge** und bezeichnet sie mit dem Symbol $\emptyset$. Da Mengen durch ihre Elemente festgelegt werden, kann es nur eine leere Menge geben! Nun können wir zeigen, dass die leere Menge Teilmenge jeder beliebigen Menge M ist:

Da die Aussage $x \in \emptyset$ für jedes x falsch ist, ist die Aussage $x \in \emptyset \implies x \in M$ immer wahr, und das bedeutet, dass die Aussage $\emptyset \subset M$ für jede Menge M wahr ist.

Wir kommen noch einmal auf die Begriffe der Schulmathematik zurück. Was bedeutet es, die quadratische Gleichung

$$x^2 - 6x - 16 = 0$$

zu lösen?

Verschiedene Fragen stellen sich:

- Existiert überhaupt eine Lösung?
- Wieviele Lösungen gibt es? Ist die Lösung womöglich eindeutig bestimmt?
- Wie findet man die Lösung oder die Lösungsmenge?

Man kann die Fragen in folgender Aufgabenstellung zusammenfassen:

Aufgabe: Bestimme die Menge $\{x \in \mathbb{R} : x^2 - 6x - 16 = 0\}$.

Lösung:

Wenn damit zu rechnen ist, dass die Lösung eindeutig bestimmt ist, kann man es mit einem Eindeutigkeitsbeweis versuchen, der unter günstigen Umständen die Lösung liefert. Aber wie soll man anfangen?

Einfacher ist es, mit „Äquivalenz-Umformungen“ zu arbeiten. Man ersetzt die zu lösende Gleichung schrittweise durch logisch äquivalente Aussagen, bis man eine einfache Form gefunden hat, die leichter zu bearbeiten ist. Im vorliegenden Fall sieht das folgendermaßen aus:

$$\begin{aligned}
x^2 - 6x - 16 = 0 &\iff x^2 - 2 \cdot 3 \cdot x + 3^2 = 16 + 3^2 \\
&\iff (x - 3)^2 = 25 \\
&\iff x - 3 = \pm 5 \\
&\quad \text{(die Schreibweise steht für eine „oder“-Verknüpfung)} \\
&\iff x = 8 \vee x = -2
\end{aligned}$$

Damit lautet die Lösung: $\{x \in \mathbb{R} : x^2 - 6x - 16 = 0\} = \{-2, 8\}$. Was hier benutzt wurde, nennt man bekanntlich die „Methode der quadratischen Ergänzung“, die auf der binomischen Formel $a^2 + 2ab + b^2 = (a + b)^2$ beruht.

Achtung! Bei Äquivalenzumformungen muss man durchaus vorsichtig sein, denn nicht jeder Schritt ist umkehrbar. Das ist eine beliebte Fehlerquelle.

Ein etwas schwierigeres Problem ist die Lösung einer quadratischen Ungleichung. Aber auch dabei kann man die Methode der quadratischen Ergänzung einsetzen. Wir betrachten z.B. die quadratische Ungleichung

$$x^2 - 4x \geq 140.$$

Per quadratischer Ergänzung erhält man die äquivalente Aussage

$$x^2 - 2 \cdot 2 \cdot x + 2^2 \geq 144,$$

also $(x - 2)^2 \geq 144$. Aber wie geht es jetzt weiter? Einfach auf beiden Seiten die Wurzel zu ziehen, ist nicht unbedingt richtig, dabei bereiten die negativen Zahlen Probleme. Tatsächlich gilt für jede reelle Zahl y :

$$\sqrt{y^2} = |y| = \begin{cases} y & \text{falls } y \geq 0, \\ -y & \text{falls } y < 0. \end{cases}$$

Damit ist unsere ursprüngliche quadratische Ungleichung äquivalent zu der Ungleichung

$$|x - 2| \geq 12,$$

denn aus Gründen, die wir hier noch nicht erläutern können, darf man bei einer Ungleichung zwischen positiven reellen Zahlen auf beiden Seiten die positive Wurzel ziehen.

Wir sind aber immer noch nicht am Ende. Was bedeutet der Ausdruck $|x - 2|$? Ist $x \geq 2$, also $x - 2 \geq 0$, so steht er für die positive Zahl $x - 2$. Ist $x < 2$, also $x - 2 < 0$, so ist $|x - 2| = -(x - 2) = 2 - x$. In beiden Fällen sieht man, dass $|x - 2|$ der Abstand der Zahl x von der Zahl 2 ist. Das bedeutet:

$$|x - 2| \geq 12 \iff x \leq -10 \text{ oder } x \geq 14.$$

Die Lösung lautet nun:

$$\{x \in \mathbb{R} : x^2 - 4x \geq 140\} = \{x \in \mathbb{R} : x \leq -10\} \cup \{x \in \mathbb{R} : x \geq 14\}.$$

Definition

Für $n \in \mathbb{N}_0$ wird die Zahl n -**Fakultät** definiert durch

$$0! := 1 \quad \text{und} \quad n! := 1 \cdot 2 \cdot \dots \cdot n \quad \text{für } n \geq 1.$$

Für $0 \leq k \leq n$ wird der **Binomialkoeffizient** $\binom{n}{k}$ definiert durch

$$\binom{n}{k} := \frac{n(n-1)(n-2) \cdots (n-k+1)}{1 \cdot 2 \cdots k} = \frac{n!}{k!(n-k)!}.$$

1.2. Satz

$n!$ ist die Anzahl der Möglichkeiten, die Elemente der Menge $\{1, \dots, n\}$ anzuordnen.

BEWEIS: Wir beginnen mit einfachen Fällen. Eine Zahl kann man nur auf **eine** Art und Weise anordnen. Die Zahlen 1 und 2 kann man auf zwei Möglichkeiten (in der Form 1, 2 und 2, 1) anordnen, die Zahlen 1, 2 und 3 in der Form

$$1, 2, 3, \quad 1, 3, 2, \quad 2, 1, 3, \quad 2, 3, 1, \quad 3, 1, 2 \quad \text{und} \quad 3, 2, 1.$$

Das sind $6 = 1 \cdot 2 \cdot 3$ Möglichkeiten. Und dieser Fall zeigt schon, wie es allgemein geht. Will man n Zahlen anordnen, so gibt es n Möglichkeiten für die erste Position, nur noch $n-1$ für die zweite, $n-2$ für die dritte, usw. Bei der letzten übriggebliebenen Zahl hat man überhaupt keine Wahl mehr, es gibt nur noch eine Möglichkeit. Insgesamt sind das $n \cdot (n-1) \cdot (n-2) \cdot \dots \cdot 2 \cdot 1 = n!$ Möglichkeiten. ■

Ist M eine beliebige Menge, so kann man die Menge $P(M)$ aller Teilmengen von M bilden. Sie wird als die **Potenzmenge** von M bezeichnet.

1.3. Beispiele

A. Sei $M := \{1, 2, 3\}$. Dann ist

$$P(M) = \{ \emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\} \}.$$

B. Da auch $\emptyset \subset \emptyset$ ist, folgt: $P(\emptyset) = \{\emptyset\}$ ist eine Menge mit einem Element. Weiter ist $P(P(\emptyset)) = \{ \emptyset, \{\emptyset\} \}$ eine Menge mit 2 Elementen.

1.4. Satz

$\binom{n}{k}$ ist die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge.

BEWEIS: Wir nehmen $M := \{1, 2, 3, \dots, n\}$ als n -elementige Menge und suchen eine k -elementige Teilmenge N . Für das erste Element haben wir n Zahlen zur Auswahl, für das zweite noch $n-1$, usw. So bekommen wir $n(n-1)(n-2) \cdots (n-k+1)$ verschiedene Möglichkeiten und erhalten alle k -elementigen Teilmengen von M . Allerdings haben wir angeordnete Teilmengen gebildet, und je zwei davon sind gleich, wenn sie die gleichen Elemente besitzen. Das bedeutet, dass wir jede Teilmenge $k!$ -mal gezählt haben, und die richtige Anzahl ist

$$N = \frac{n(n-1)(n-2) \cdots (n-k+1)}{k!} = \binom{n}{k}.$$

■

1.5. Satz

Es gelten folgende Formeln:

1. $\binom{n}{0} = 1$, $\binom{n}{1} = n$ und $\binom{n}{k} = \binom{n}{n-k}$.
2. $\binom{n}{k} = \binom{n-1}{k-1} + \binom{n-1}{k}$.

BEWEIS: Die Aussagen in (1) sind trivial (warum?).

Die Aussage (2) muss man nachrechnen:

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

■

Daraus ergibt sich das *Pascalsche Dreieck*:

$$\begin{array}{cccccccc} n=0 & & & & & & & & 1 \\ & 1 & & & & & & & 1 & 1 \\ & 2 & & & 1 & & 2 & & 1 \\ & 3 & & 1 & & 3 & & 3 & & 1 \\ & 4 & & 1 & & 4 & & 6 & & 4 & & 1 \\ & 5 & & 1 & & 5 & & 10 & & 10 & & 5 & & 1 \\ & & & & & & & \dots & & & & & & \end{array}$$

Wir haben schon den sogenannten „Allquantor“ $\forall$ kennengelernt. Das Gegenstück dazu ist der „Existenzquantor“ $\exists$, der „es gibt ein ...“ bedeutet und z.B. folgendermaßen benutzt wird:

$$\exists x : x^2 = 2.$$

Diese Aussage beschreibt die Existenz der Zahl $\sqrt{2}$.

Interessant ist das Verhalten der Quantoren bei der logischen Verneinung:

$$\neg(\exists x : \mathcal{A}(x)) \iff \forall x : \neg\mathcal{A}(x).$$

$$\neg(\forall x : \mathcal{B}(x)) \iff \exists x : \neg\mathcal{B}(x).$$

Beispiele werden wir in den nächsten Abschnitten in großer Zahl kennenlernen.

1.2 Die Axiome der reellen Zahlen

Wir wollen jetzt den Dingen auf den Grund gehen.

- Warum gelten die binomischen Formeln?
- Warum ist $(-1) \cdot (-1) = 1$?
- Warum darf man nicht durch Null dividieren?
- Was ist 0^0 ?
- Warum gibt es keine Wurzel aus -7 ?

Um vernünftig Schlüsse ziehen zu können, braucht man eine solide Ausgangsbasis. Deshalb führen wir jetzt die reellen Zahlen axiomatisch ein.

Die reellen Zahlen bilden eine Menge $\mathbb{R}$. Je zwei Elementen $x, y \in \mathbb{R}$ ist auf eindeutige Weise ein Element $x + y \in \mathbb{R}$ und ein Element $x \cdot y \in \mathbb{R}$ zugeordnet.

Axiome für die Addition und Multiplikation:

1. **Kommutativgesetz:** $a + b = b + a$ und $a \cdot b = b \cdot a$, für alle $a, b \in \mathbb{R}$.
2. **Assoziativgesetz:** $a + (b + c) = (a + b) + c$ und $a \cdot (b \cdot c) = (a \cdot b) \cdot c$, für alle $a, b, c \in \mathbb{R}$.
3. **Distributivgesetz:** $a \cdot (b + c) = a \cdot b + a \cdot c$, für alle $a, b, c \in \mathbb{R}$.
4. **Existenz der Null und des Negativen:**

a) Es gibt genau ein Element $0 \in \mathbb{R}$, so dass für alle $a \in \mathbb{R}$ gilt:

$$a + 0 = a.$$

b) Zu jedem Element $a \in \mathbb{R}$ gibt es ein Element $-a \in \mathbb{R}$ mit

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

5. **Existenz der Eins und des Inversen:**

a) Es gibt genau ein Element $1 \neq 0$ in $\mathbb{R}$, so dass für alle $a \in \mathbb{R}$ gilt:

$$a \cdot 1 = a.$$

b) Zu jedem Element $b \neq 0$ in $\mathbb{R}$ gibt es ein Element $b^{-1} \in \mathbb{R}$ mit

$$b \cdot b^{-1} = 1.$$

Exemplarisch wollen wir ein paar einfache Aussagen beweisen:

2.1. Satz

1. *Negatives und Inverses sind jeweils eindeutig bestimmt.*
2. *Ist $a \in \mathbb{R}$ beliebig, so ist $a \cdot 0 = 0$.*
3. *Sind $a, b \in \mathbb{R}$ mit $a \cdot b = 0$, so ist $a = 0$ oder $b = 0$.*
4. *Es ist $(-1) \cdot (-1) = 1$.*

BEWEIS: 1) Sei $a + (-a) = 0$. Ist außerdem auch $a + c = 0$, so ist

$$-a = -a + 0 = -a + (a + c) = ((-a) + a) + c = 0 + c = c.$$

Beim Inversen argumentiert man analog.

2) Wegen $a \cdot 0 = a \cdot (0 + 0) = a \cdot 0 + a \cdot 0$ und der Eindeutigkeit der Null muss $a \cdot 0 = 0$ sein. Hier wurde ein Standardtrick benutzt, 0 wurde durch $0 + 0$ ersetzt. Das Distributivgesetz führt dann zum Erfolg.

3) Sei $a \cdot b = 0$. Ist $a \neq 0$, so ist $0 = a^{-1}(a \cdot b) = (a^{-1}a)b = 1 \cdot b = b$. Der erste Trick besteht darin, zu sehen, dass man nur den Fall $a \neq 0$ betrachten muss. Der zweite Trick ist die Erkenntnis, dass man nun das Element a^{-1} zur Verfügung hat.

4) Es ist

$$(-1) + (-1)(-1) = (-1)(1 + (-1)) = (-1) \cdot 0 = 0.$$

Weil auch $(-1) + 1 = 0$ und nach (1) das Negative von -1 eindeutig bestimmt ist, muss $(-1)(-1) = 1$ sein. Wieder wurde das Distributivgesetz trickreich ausgenutzt.

■

Axiome der Anordnung:

In $\mathbb{R}$ ist eine Teilmenge $\mathbb{R}_+$ ausgezeichnet, die Menge der *positiven reellen Zahlen*. Für $x \in \mathbb{R}_+$ schreibt man: $x > 0$ („ x ist größer als 0“).

1. Für eine Zahl $a \in \mathbb{R}$ gilt immer genau eine der drei Beziehungen $a > 0$, $a = 0$ oder $-a > 0$.
2. Ist $a > 0$ und $b > 0$, so ist auch $a + b > 0$.
3. Ist $a > 0$ und $b > 0$, so ist auch $a \cdot b > 0$.

Ist $a - b > 0$, so schreibt man $a > b$ oder $b < a$. Ist $a < b$ oder $a = b$, so schreibt man $a \leq b$ („ a ist kleiner oder gleich b “).

Hier sind ein paar Folgerungen:

2.2. Satz

1. Ist $a > b$ und $b > c$, so ist auch $a > c$ (Transitivität).
2. Ist $a > b$ und c beliebig, so ist auch $a + c > b + c$.
3. Ist $a > b$ und $c < 0$, so ist $ac < bc$.
4. Ist $a \leq b + \varepsilon$ für alle $\varepsilon > 0$, so ist $a \leq b$.

BEWEIS: 1) Ist $a - b > 0$ und $b - c > 0$, so ist auch $a - c = (a - b) + (b - c) > 0$.

2) Ist $a - b > 0$ und c beliebig, so ist $(a + c) - (b + c) = a - b > 0$.

3) Sei $a - b > 0$. Wenn $c < 0$ ist, ist $-c > 0$, also $bc - ac = (a - b)(-c) > 0$ und damit $ac < bc$.

4) Wir nehmen an, es sei $a > b$. Dann ist $\varepsilon := (a - b)/2 > 0$ und

$$b + \varepsilon = \frac{2b + (a - b)}{2} = \frac{a + b}{2} < \frac{a + a}{2} = a.$$

Das ist ein Widerspruch.

Dies ist übrigens ein typisches Beispiel für eine Aussage, die **nur** mit Hilfe des Widerspruchsprinzips auf einfache Weise aus den Axiomen abgeleitet werden kann! Ein konstruktiver Beweis müsste z.B. die Dezimalbruch-Darstellung sehr intensiv benutzen. ■

Definition

Eine Teilmenge $M \subset \mathbb{R}$ heißt **induktiv**, falls gilt:

1. Die 1 gehört zu M .
2. Liegt x in M , so liegt auch $x + 1$ in M .

Auf den ersten Blick scheint diese Definition ziemlich nutzlos zu sein. Ganz $\mathbb{R}$ ist induktiv, die Menge der rationalen Zahlen ist induktiv, vielleicht ist ja jede Teilmenge von $\mathbb{R}$ induktiv? Schauen wir genauer hin, so sehen wir:

1. Ist $M \subset \mathbb{R}$ induktiv, so gehört die Zahl 1 zu M .
2. Mit 1 muss auch $2 = 1 + 1$ zu M gehören. Und mit 2 gehört $3 = 2 + 1$ zu M , usw.

Das motiviert die folgende Festlegung:

Definition

Eine Zahl $n \in \mathbb{R}$ heißt **natürliche Zahl**, falls sie in jeder induktiven Menge liegt.

Die Menge $\mathbb{N}$ aller natürlichen Zahlen ist die Schnittmenge aller induktiven Mengen (und damit die „kleinste“ induktive Menge).

Wir haben bisher nur Schnittmengen von endlich vielen Mengen betrachtet. Jetzt sollten wir die Notationen erweitern:

Ist I eine beliebige Menge und zu jedem Element $i \in I$ eine Menge M_i gegeben, so sprechen wir von einer **Familie von Mengen** und schreiben dafür $(M_i)_{i \in I}$. Gemeint ist damit einfach die Menge aller M_i , also die Menge $\{M_i : i \in I\}$. Die Aussage „ x ist Element jeder Menge M_i “ kann abgekürzt geschrieben werden:

$$\forall i \in I : x \in M_i.$$

Die gewünschte Schnittmenge ist dann die Menge

$$\bigcap_{i \in I} M_i := \{x : \forall i \in I \text{ ist } x \in M_i\}.$$

Man kann auch die Vereinigung einer ganzen Familie von Mengen bilden. Ein Element x liegt genau dann in $M_i \cup M_j \cup M_k$, wenn es in wenigstens einer der drei Mengen liegt. Dementsprechend liegt x in der Vereinigung aller M_i , wenn es wenigstens ein $i \in I$ gibt, so dass $x \in M_i$ gilt. Diese Aussage wird abgekürzt durch

$$\exists i \in I : x \in M_i.$$

Die Vereinigungsmenge ist die Menge

$$\bigcup_{i \in I} M_i := \{x : \exists i \in I \text{ mit } x \in M_i\}.$$

Wir kommen zurück zu den natürlichen Zahlen. Wenn wir die Menge $\mathbb{R}$ mit allen Rechenregeln als bekannt voraussetzen, können wir $\mathbb{N} \subset \mathbb{R}$ als kleinste induktive Menge definieren. So erhalten wir ein Modell für die uns intuitiv bekannte Menge der Zahlen 1, 2, 3, 4, ..., und es folgt das

2.3. Induktionsprinzip

Sei $M \subset \mathbb{N}$ eine Teilmenge. Ist $1 \in M$ und mit $n \in M$ stets auch $n + 1 \in M$, so ist $M = \mathbb{N}$.

BEWEIS: Definitionsgemäß ist $\mathbb{N}$ die kleinste induktive Teilmenge von $\mathbb{R}$. Die hier betrachtete Menge M ist ebenfalls induktiv und außerdem Teilmenge von $\mathbb{N}$. Das geht nur, wenn $M = \mathbb{N}$ ist. ■

Damit steht der **Beweis durch vollständige Induktion** als neues Beweisprinzip zur Verfügung. Soll eine Aussage $A(n)$ für alle natürlichen Zahlen bewiesen werden, so genügt es zu zeigen:

1. Es gilt $A(1)$ (Induktionsanfang),
2. Aus $A(n)$ folgt stets $A(n+1)$ (Induktionsschluss).

Wendet man das Induktionsprinzip auf die Menge $M = \{n \in \mathbb{N} : A(n)\}$ an, so folgt, dass $M = \mathbb{N}$ ist, also $A(n)$ wahr für alle $n \in \mathbb{N}$.

Bevor wir uns Beispiele ansehen, sei noch erwähnt, dass ein Induktionsbeweis auch bei 0 oder einer Zahl $n_0 > 1$ beginnen kann. Im ersten Fall ist die Aussage dann für alle Zahlen $n \in \mathbb{N}_0$ bewiesen, im zweiten Fall für alle natürlichen Zahlen $n \geq n_0$.

2.4. Beispiele

- A.** Wir zeigen, dass $2n+1 < 2^n$ für $n \geq 3$ ist (für $n=1$ und $n=2$ ist es falsch).

Im Falle $n=3$ ergibt die linke Seite $2 \cdot 3 + 1 = 7$ und die rechte Seite $2^3 = 8$.

Ist die Aussage für $n \geq 3$ bewiesen, so ist

$$2(n+1) + 1 = (2n+1) + 2 < 2^n + 2 < 2^n + 2^n = 2^{n+1}.$$

- B.** Es soll die Aussage $n^2 < 2^n$ bewiesen werden. Zur Vorsicht kann man ein paar einfache Fälle testen. Dann stellt man fest, dass die Aussage zwar für $n=1$ stimmt, für $n=2, 3, 4$ aber falsch ist. Wir behaupten also:

$$n^2 < 2^n \text{ für alle natürlichen Zahlen } n \geq 5.$$

Der BEWEIS durch Induktion nach n beginnt bei $n=5$. Tatsächlich ist $5^2 = 25 < 32 = 2^5$. Damit ist der Induktionsanfang geschafft.

Nun nehmen wir an, dass $n \geq 5$ und die Aussage für n bewiesen ist. Für den Induktionsschluss haben wir zu zeigen, dass $(n+1)^2 < 2^{n+1}$ ist. Tatsächlich ist

$$\begin{aligned} (n+1)^2 &= n^2 + 2n + 1 \\ &< 2^n + 2n + 1 \text{ (nach Induktionsvoraussetzung)} \\ &< 2^n + 2^n = 2^{n+1} \text{ (weil } 2n+1 < 2^n \text{ für } n \geq 3 \text{ gilt).} \end{aligned}$$

- C.** Die sogenannte *Bernoulli'sche Ungleichung* besagt:

$$(1+x)^n > 1+nx \text{ für } x > -1, x \neq 0 \text{ und } n \geq 2.$$

BEWEIS durch Induktion nach n :

Induktionsanfang: Im Falle $n = 2$ ist die Ungleichung $(1 + x)^2 = 1 + 2x + x^2 > 1 + 2x$ offensichtlich erfüllt.

Induktionsschluss: Nach Voraussetzung ist $1 + x > 0$. Ist die Behauptung für n bewiesen, so folgt:

$$\begin{aligned} (1 + x)^{n+1} &= (1 + x)(1 + x)^n \\ &> (1 + x)(1 + nx) \quad (\text{nach Induktionsvoraussetzung}) \\ &= 1 + (n + 1)x + nx^2 > 1 + (n + 1)x. \end{aligned}$$

Definition

Es sei $n \in \mathbb{N}$, und für jede natürliche Zahl i mit $1 \leq i \leq n$ sei eine reelle Zahl a_i gegeben. Dann beschreiben wir die Summe der a_i durch das Symbol

$$\sum_{i=1}^n a_i := a_1 + a_2 + \cdots + a_n.$$

Dabei soll die „leere“ Summe (im Falle $n < 1$) den Wert 0 erhalten.

Für beliebige Indizes $k, l \in \mathbb{Z}$ setzt man:

$$\sum_{i=k}^l a_i := \begin{cases} 0 & \text{falls } k > l, \\ a_k + a_{k+1} + \cdots + a_l & \text{sonst.} \end{cases}$$

Aus den Axiomen für die Grundrechenarten ergeben sich z.B. die folgenden Regeln für den Umgang mit dem Summenzeichen:

1) **Multiplikation mit einer Konstanten:** Ist $c \in \mathbb{R}$, so ist

$$c \cdot \sum_{i=k}^l a_i = \sum_{i=k}^l (c \cdot a_i).$$

2) **Summe von Summen:** Ist zu jedem i noch eine reelle Zahl b_i gegeben, so gilt:

$$\sum_{i=k}^l a_i + \sum_{i=k}^l b_i = \sum_{i=k}^l (a_i + b_i).$$

Es gibt auch ein Produktzeichen:

$$\prod_{i=1}^n a_i := a_1 \cdot a_2 \cdots a_n.$$

Hier muss man aber das leere Produkt $= 1$ setzen. Ansonsten gelten analoge Regeln wie beim Summenzeichen. Wir notieren nur, dass $x^n = \prod_{i=1}^n x$ und deshalb $x^0 = 1$ für jede reelle Zahl x gilt (also auch $0^0 = 1$).

2.5. Geometrische Summenformel

Ist $x \in \mathbb{R}$, $x \neq 1$ und $n \in \mathbb{N}$, so gilt:

$$\sum_{i=0}^n x^i = \frac{x^{n+1} - 1}{x - 1}.$$

BEWEIS: Es ist

$$(x - 1) \cdot \sum_{i=0}^n x^i = \sum_{i=0}^n x^{i+1} - \sum_{i=0}^n x^i = \sum_{i=1}^{n+1} x^i - \sum_{i=0}^n x^i = x^{n+1} - 1.$$

■

2.6. Folgerung

Sind $a, b \in \mathbb{R}$, mit $a \neq b$, so ist

$$\sum_{i=0}^n a^i b^{n-i} = \frac{a^{n+1} - b^{n+1}}{a - b}.$$

BEWEIS: Die Aussage ist trivial für $b = 0$. Es sei also o.B.d.A. (d.h. „ohne Beschränkung der Allgemeinheit“) $b \neq 0$. Wir setzen $x = a/b$ in der geometrischen Summenformel. Es ist

$$\frac{a^{n+1} - b^{n+1}}{a - b} = b^n \cdot \frac{(a/b)^{n+1} - 1}{(a/b) - 1} = \frac{b^{n+1}}{b} \cdot \sum_{i=0}^n (a/b)^i = \sum_{i=0}^n a^i b^{n-i}.$$

■

2.7. Die binomische Formel

Seien $a, b \in \mathbb{R}$ und $n \in \mathbb{N}$. Dann gilt:

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k = a^n + n a^{n-1} b + \frac{n(n-1)}{2} a^{n-2} b^2 + \dots + n a b^{n-1} + b^n.$$

BEWEIS: Wir beweisen zunächst den Spezialfall $(1 + x)^n = \sum_{k=0}^n \binom{n}{k} x^k$.

Multipliziert man ein Produkt $(1 + x_1) \cdots (1 + x_n)$ distributiv aus, so erhält man Summanden der Gestalt $1, x_i, x_i x_j$ (mit $i < j$) usw. Jeder k -elementigen Teilmenge $\{x_{i_1}, \dots, x_{i_k}\}$ (mit $i_1 < \dots < i_k$) von $\{1, \dots, n\}$ entspricht genau ein Summand $x_{i_1} \cdots x_{i_k}$. Ist nun $x_1 = \dots = x_n = x$, so folgt der Spezialfall.

Im allgemeinen Fall können wir $a \neq 0$ voraussetzen und erhalten

$$(a+b)^n = a^n \left(1 + \frac{b}{a}\right)^n = a^n \cdot \sum_{k=0}^n \binom{n}{k} \left(\frac{b}{a}\right)^k = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k.$$

Damit ist alles gezeigt. Wir hätten natürlich auch einen Induktionsbeweis führen können, aber der wäre unübersichtlicher gewesen. ■

Im Falle $n = 2$ und $n = 3$ erhält man speziell

$$(a+b)^2 = a^2 + 2ab + b^2 \quad \text{und} \quad (a+b)^3 = a^3 + 3a^2b + 3ab^2 + b^3.$$

2.8. Folgerung

$$\sum_{k=0}^n \binom{n}{k} = 2^n \quad \text{und} \quad \sum_{k=0}^n (-1)^k \binom{n}{k} = 0.$$

BEWEIS: Setze $a = b = 1$, bzw. $a = 1$ und $b = -1$. ■

Nun folgt auch: Besitzt eine Menge M n Elemente, so besitzt $P(M)$ 2^n Elemente. Das erklärt die Bezeichnung „Potenzmenge“.

1.3 Vollständigkeit und Konvergenz

Definition

Ist $x \in \mathbb{R}$, so heißt $|x| := \begin{cases} x & \text{falls } x \geq 0, \\ -x & \text{falls } x < 0. \end{cases}$ der **(Absolut-)Betrag** von x .

Stellen wir uns die reellen Zahlen a, b als Punkte auf einer Geraden vor, so ist $|a - b| = |b - a|$ der Abstand von a und b auf der Geraden. Speziell ist $|a|$ der Abstand der Zahl a vom Nullpunkt.

3.1. Satz

Sind a, b, c reelle Zahlen, so gilt:

1. $|a \cdot b| = |a| \cdot |b|$.
2. Es ist stets $-|a| \leq a \leq +|a|$.
3. Ist $c > 0$, so gilt: $|x| < c \iff -c < x < +c$.
4. Es ist $|a + b| \leq |a| + |b|$ (**Dreiecksungleichung**).
5. Es ist $|a - b| \geq |a| - |b|$.

Zum BEWEIS: (1) und (2) erhält man durch Fallunterscheidung.

3) Ist $|x| < c$, so ist $-|x| > -c$ und daher

$$-c < -|x| \leq x \leq |x| < c.$$

Ist umgekehrt $-c < x < +c$, so unterscheiden wir zwei Fälle: Ist $x \geq 0$, so ist $|x| = x < c$. Ist $x < 0$, so ist $|x| = -x < -(-c) = c$ (wegen $x > -c$).

4) Wegen (2) ist $-(|a| + |b|) = -|a| - |b| \leq a + b \leq |a| + |b|$. Wegen (3) folgt daraus die Dreiecksungleichung.

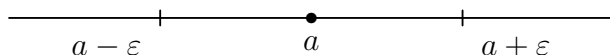
Zum Beweis von (5) benutzt man den beliebten Trick, eine Null einzufügen:

Es ist $|a| = |(a - b) + b| \leq |a - b| + |b|$. ■

Für beliebiges $a \in \mathbb{R}$ und $\varepsilon > 0$ nennt man die Menge

$$U_\varepsilon(a) := \{x \in \mathbb{R} \mid a - \varepsilon < x < a + \varepsilon\} = \{x \in \mathbb{R} : |x - a| < \varepsilon\}$$

die ε -**Umgebung** von a . Sie besteht aus allen Punkten x auf der Zahlengeraden, deren Abstand von a kleiner als ε ist.



Definition

Sind $a < b$ zwei reelle Zahlen, so heißt

$$(a, b) := \{x \in \mathbb{R} : a < x < b\}$$

das **offene Intervall** mit den Grenzen a und b . Die Menge

$$[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\}$$

nennt man das **abgeschlossene Intervall** mit den Grenzen a und b .

$[a, b) := \{x \in \mathbb{R} : a \leq x < b\}$ und $(a, b] := \{x \in \mathbb{R} : a < x \leq b\}$ heißen **halboffene Intervalle**.

Die ε -Umgebung von a ist also das offene Intervall $(a - \varepsilon, a + \varepsilon)$.

Man führt nun zwei (voneinander verschiedene) neue Objekte $-\infty$ und $+\infty$ ein, die nicht zu den reellen Zahlen gehören (so dass man mit ihnen auch nicht rechnen kann) und die folgende Eigenschaften besitzen:

- Es ist $-\infty < +\infty$.
- Für alle reellen Zahlen x ist $-\infty < x < +\infty$.

Die Menge $\overline{\mathbb{R}} := \mathbb{R} \cup \{-\infty, +\infty\}$ nennt man die **abgeschlossene Zahlengerade**. Die Mengen $(-\infty, a)$, $(-\infty, a]$, $[a, +\infty)$ und $(a, +\infty)$ nennt man **Halbgeraden** oder **entartete Intervalle**.

Definition

Eine Menge $M \subset \mathbb{R}$ heißt **nach oben beschränkt**, falls es eine reelle Zahl c gibt, so dass $x \leq c$ für alle $x \in M$ gilt. Die Zahl c nennt man dann eine **obere Schranke** für M .

Besitzt M eine **untere Schranke**, also eine reelle Zahl c , so dass $x \geq c$ für alle $x \in M$ ist, so heißt M nach **unten beschränkt**.

M heißt **beschränkt**, falls M nach unten und nach oben beschränkt ist.

Definition

Sei $M \subset \mathbb{R}$ eine nach oben beschränkte Menge. Wenn die Menge aller oberen Schranken von M ein kleinstes Element a besitzt, so nennt man diese kleinste obere Schranke das **Supremum** von M (in Zeichen: $a = \sup(M)$).

Ist M nach unten beschränkt, so nennt man die größte untere Schranke das **Infimum** von M (in Zeichen: $\inf(M)$).

3.2. Beispiel

Sei $M = (0, 1)$. Dann ist natürlich jede Zahl $c > 1$ eine obere Schranke von M . Und auch die 1 ist noch eine obere Schranke. Eine Zahl $c < 1$ kann dagegen keine obere Schranke sein, denn es gibt Zahlen d mit $c < d < 1$ (z.B. $d := (1 + c)/2$). Also ist $S := [1, \infty)$ die Menge der oberen Schranken von M . Tatsächlich hat S ein kleinstes Element, die 1. Damit ist $\sup(M) = 1$.

Für die Menge $N = (0, 1]$ erhalten wir die gleiche Menge von oberen Schranken. Deshalb ist auch $\sup(N) = 1$. Das Supremum einer Menge kann zu der Menge gehören, muss es aber nicht.

3.3. Vollständigkeits-Axiom

Jede nicht leere und nach oben beschränkte Menge besitzt ein Supremum.

Unbeschränkte Mengen besitzen definitionsgemäß kein Supremum oder kein Infimum. Diesen Mangel kann man aber künstlich beheben: Ist $M \subset \mathbb{R}$ nicht nach oben beschränkt, so setzt man $\sup(M) := +\infty$; ist M nicht nach unten beschränkt, so setzt man $\inf(M) := -\infty$.

Mit dieser Notation gilt jetzt:

$$M \subset \mathbb{R} \text{ beschränkt} \iff \sup(M) < +\infty \text{ und } \inf(M) > -\infty.$$

Eine besonders wichtige Anwendung betrifft die Verteilung der natürlichen Zahlen:

3.4. Satz von Archimedes

Zu jeder reellen Zahl x gibt es eine natürliche Zahl n , die größer als x ist.

BEWEIS: Mit Quantoren geschrieben, lautet die Behauptung:

$$\forall x \in \mathbb{R} \exists n \in \mathbb{N} \text{ mit } n > x.$$

Wir arbeiten nun mit dem Widerspruchsprinzip. Angenommen, es gibt ein $x_0 \in \mathbb{R}$, so dass $x_0 \geq n$ für alle $n \in \mathbb{N}$ gilt. Dann ist $\mathbb{N}$ nach oben beschränkt. Also existiert $a := \sup(\mathbb{N})$, die kleinste obere Schranke von $\mathbb{N}$. Dies ist eine reelle Zahl, und $a - 1$ ist keine obere Schranke mehr. Also gibt es ein $n_0 \in \mathbb{N}$ mit $a - 1 < n_0$. Dann ist $n_0 + 1 > a$. Da $n_0 + 1$ eine natürliche Zahl ist, widerspricht das der Supremums-Eigenschaft von a . ■

Jetzt zeigen wir mit Hilfe des Vollständigkeitsaxioms die Existenz von Wurzeln.

3.5. Satz von der Existenz der Quadratwurzel

Sei $a > 0$ eine reelle Zahl. Dann gibt es genau eine reelle Zahl $c > 0$ mit $c^2 = a$.

BEWEIS: Die Grundidee ist einfach: Wir betrachten die Menge aller positiven reellen Zahlen, deren Quadrat kleiner als a ist, und hoffen, dass das Supremum dieser Menge gerade die gesuchte Zahl x mit $x^2 = a$ ist. Zur Vereinfachung nehmen wir zunächst an, dass $a > 1$ ist. Sei $M := \{x \in \mathbb{R} \mid (x > 0) \wedge (x^2 < a)\}$.

- 1) Die Menge M ist nicht leer, denn 1 liegt in M .
- 2) Die Menge M ist nach oben beschränkt, denn a selbst ist eine obere Schranke: Ist nämlich $x > a$, so ist $x^2 > a^2 > a$, und x kann nicht in M liegen.
- 3) Sei $c := \sup(M)$. Offensichtlich muss $c \geq 1$ sein. Ist $c^2 = a$, so ist nichts mehr zu zeigen. Andernfalls gibt es 2 Möglichkeiten:

1. Fall: $c^2 < a$.

Was nun? Wir möchten zeigen, dass dieser Fall überhaupt nicht eintreten kann. Zwischen c^2 und a ist ja noch etwas Platz. Wenn wir also an c ein ganz kleines bisschen wackeln (d.h. c geringfügig vergrößern), dann können wir erwarten, dass sich auch c^2 nur wenig vergrößert und deshalb immer noch unterhalb von a bleibt. Wenn wir etwa eine sehr große natürliche Zahl n wählen, dann wird ihr Kehrwert sehr klein, und wir können hoffen, dass $(c + \frac{1}{n})^2 < a$ ist, im Widerspruch zu der Tatsache, dass $c = \sup(M)$ ist.

Aber wie groß müssen wir n wählen? Die folgende Überlegung hat den Rang einer Nebenrechnung. Weil wir vom Ziel ausgehen, gehen wir logisch in der falschen Richtung vor. Wenn wir also das geeignete n gefunden haben, müssen wir versuchen, alle Schlüsse umzukehren, damit die logische Richtung stimmt.

Hier kommt die Nebenrechnung: Wenn ein passendes n existiert, dann erhalten wir folgende Abschätzungen:

$$\begin{aligned} (c + \frac{1}{n})^2 < a &\implies c^2 + \frac{2c}{n} + \frac{1}{n^2} < a \\ &\implies \frac{1}{n}(2c + \frac{1}{n}) < a - c^2 \\ &\implies n > \frac{2c + \frac{1}{n}}{a - c^2}. \end{aligned}$$

Jetzt wollen wir daraus eine logisch saubere Deduktion machen: Wir wollen von der untersten Formel als Prämisse ausgehen; dass – bei gegebenem c – ein solches n existiert, müssen wir aber erst mal beweisen. Eine Anwendung des Satzes von Archimedes liegt nahe, aber dafür müssten wir n ganz auf einer Seite isolieren. Der Schönheitsfehler, dass n auf beiden Seiten der Ungleichung auftritt, ist jedoch schnell behoben, denn für $n \in \mathbb{N}$ gilt:

$$n > \frac{2c + 1}{a - c^2} \implies n > \frac{2c + \frac{1}{n}}{a - c^2}.$$

Finden wir nun ein n , so dass die linke Ungleichung erfüllt ist? Klar, denn weil $a - c^2 > 0$ ist, ist $\frac{2c + 1}{a - c^2}$ eine positive reelle Zahl, und nach dem Satz von Archimedes gibt es ein $n \in \mathbb{N}$ mit

$$n > \frac{2c+1}{a-c^2}.$$

Jetzt können wir in sauberer deduktiver Manier weiterschließen: Es ist $\frac{1}{n} \cdot (2c+1) < a - c^2$ und daraus folgt:

$$\begin{aligned} \left(c + \frac{1}{n}\right)^2 &= c^2 + 2c \cdot \frac{1}{n} + \frac{1}{n^2} \\ &\leq c^2 + \frac{1}{n} \cdot (2c+1) \\ &< c^2 + (a - c^2) = a. \end{aligned}$$

Das bedeutet, dass $c + \frac{1}{n}$ in M liegt und c keine obere Schranke sein kann. Dieser Fall kommt also nicht in Frage.

2. Fall: $c^2 > a$.

Dieser Fall wird völlig analog erledigt.

4) Ist $a < 1$, so löst man zunächst $x^2 = 1/a$ und bildet dann den Kehrwert.

5) Schließlich zeigen wir noch die Eindeutigkeit: Ist $c_1^2 = c_2^2 = a$, für zwei positive reelle Zahlen c_1 und c_2 , so ist

$$0 = c_1^2 - c_2^2 = (c_1 - c_2)(c_1 + c_2), \quad \text{also } c_1 = c_2.$$

■

Definition

Sei $a > 0$ eine reelle Zahl. Die eindeutig bestimmte reelle Zahl $c > 0$ mit $c^2 = a$ nennt man die **(Quadrat-) Wurzel** von a und man schreibt:

$$c = \sqrt{a}$$

Zusätzlich definiert man noch $\sqrt{0} := 0$.

Ist x eine beliebige reelle Zahl $\neq 0$, so ist x oder $-x$ positiv. Weil $x^2 = x \cdot x = (-x) \cdot (-x)$ ist, folgt: Das Quadrat einer reellen Zahl ist niemals negativ. Ist also $a < 0$, so existiert die Quadratwurzel aus a nicht!

Wie oben zeigt man auch: Ist $a \geq 0$ und $n \in \mathbb{N}$, so gibt es genau ein $y \geq 0$ mit $y^n = a$. Man nennt $y = \sqrt[n]{a}$ die **n -te Wurzel** aus a . Ist n gerade, so darf a nicht negativ sein. Ist dagegen n ungerade, so ist $\sqrt[n]{a} = -\sqrt[n]{|a|}$ die eindeutig bestimmte Zahl, deren n -te Potenz a ergibt.

Insbesondere gilt dann:

$$\sqrt[n]{x^n} = \begin{cases} |x| & \text{falls } n \text{ gerade,} \\ x & \text{falls } n \text{ ungerade.} \end{cases}$$

Es gibt viele irrationale Zahlen, z.B. alle Zahlen $\sqrt[n]{p}$, $n \geq 2$ und p prim. Andererseits gilt:

3.6. Die rationalen Zahlen liegen „dicht“ in $\mathbb{R}$

Sei $a \in \mathbb{R}$ und $\varepsilon > 0$. Dann gibt es eine Zahl $q \in \mathbb{Q}$, so dass $|q - a| < \varepsilon$ ist.

BEWEIS: Ist a selbst rational, so ist die Aussage trivial. Außerdem können wir uns auf den Fall $a > 0$ beschränken.

Zu einem vorgegebenen $\varepsilon > 0$ kann man dann ein $n \in \mathbb{N}$ finden, so dass $1/n < \varepsilon$ ist. Weiter kann man nach Archimedes ein $m \in \mathbb{N}$ finden, so dass $m > n \cdot a$ ist, und da jede Teilmenge von $\mathbb{N}$ ein kleinstes Element besitzt, können wir m minimal wählen. Dann ist $m - 1 \leq n \cdot a < m$, und es folgt:

$$\frac{m}{n} - \varepsilon < \frac{m}{n} - \frac{1}{n} \leq a < \frac{m}{n} < \frac{m}{n} + \frac{1}{n} < \frac{m}{n} + \varepsilon.$$

Für $q := m/n$ ist also $|a - q| < \varepsilon$. ■

Ist für jedes $n \in \mathbb{N}$ eine reelle Zahl a_n gegeben, so sprechen wir von einer (unendlichen) **Zahlenfolge** und bezeichnen diese Folge mit (a_n) . Die Zahlen a_n selbst nennt man die *Glieder* der Folge. Man darf die Folge nicht mit der Menge ihrer Glieder verwechseln. So ist z.B. $a_n = (-1)^n$ eine unendliche Zahlenfolge, aber $\{a_n : n \in \mathbb{N}\} = \{1, -1\}$ besteht nur aus zwei Elementen.

Häufig verwendet man die folgende suggestive Sprechweise: Eine Eigenschaft der Folgenglieder a_n gilt für **fast alle** $n \in \mathbb{N}$, falls es ein $n_0 \in \mathbb{N}$ gibt, so dass alle a_n mit $n \geq n_0$ die fragliche Eigenschaft besitzen.

Wir haben schon gesehen, dass es viele irrationale Zahlen gibt, und wollen uns nun davon überzeugen, dass es sogar viel mehr irrationale als rationale Zahlen gibt.

Definition

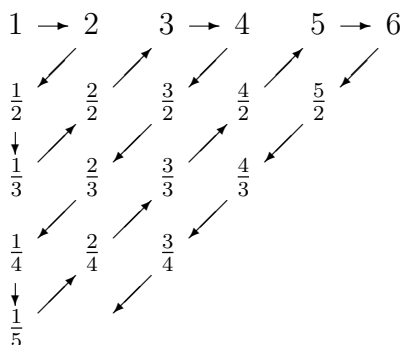
Eine Teilmenge $M \subset \mathbb{R}$ heißt **abzählbar**, wenn es eine Folge (a_n) gibt, deren zugehörige Menge genau M ist.

Damit sind auch endliche Mengen abzählbar!

3.7. Satz

Die Menge $\mathbb{Q}$ der rationalen Zahlen ist abzählbar.

Wir verwenden das **Cantorsche Diagonalverfahren**:



3.8. Satz

Die Menge $\mathbb{R}$ der reellen Zahlen ist nicht abzählbar.

BEWEIS: Wir beschränken uns auf reelle Zahlen zwischen 0 und 1 und führen den Beweis durch Widerspruch. Wäre die Menge der reellen Zahlen zwischen 0 und 1 abzählbar, so könnte man alle diese Zahlen in einer unendlichen Kolonne hintereinander aufschreiben:

$$\begin{aligned} x_1 &= 0.a_{11}a_{12}a_{13} \dots, \\ x_2 &= 0.a_{21}a_{22}a_{23} \dots, \\ x_3 &= 0.a_{31}a_{32}a_{33} \dots, \\ &\vdots \end{aligned}$$

Die Ziffern a_{ij} nehmen dabei wie üblich Werte zwischen 0 und 9 an.

Nun konstruieren wir eine reelle Zahl $y = 0.c_1c_2c_3 \dots$ wie folgt:

$$\text{Es sei } c_i := \begin{cases} 5 & \text{falls } a_{ii} \neq 5 \\ 4 & \text{falls } a_{ii} = 5 \end{cases}$$

Offensichtlich liegt y zwischen 0 und 1 und muss unter den Folgengliedern $x_1, x_2, x_3, \dots$ vorkommen. Es gibt also ein $n \in \mathbb{N}$, so dass $y = x_n$ ist. Dann ist $c_n = a_{nn}$, im Widerspruch zur Definition. ■

Definition

Eine Folge (a_n) **konvergiert** gegen eine reelle Zahl a , falls es zu jedem $\varepsilon > 0$ ein $n_0 \in \mathbb{N}$ gibt, so dass alle a_n mit $n \geq n_0$ in $U_\varepsilon(a)$ liegen. Man bezeichnet dann a als den **Grenzwert** oder **Limes** der Folge (a_n) und schreibt:

$$\lim_{n \rightarrow \infty} a_n = a.$$

Kurz gesagt konvergiert (a_n) genau dann gegen a , wenn in jeder ε -Umgebung von a fast alle a_n liegen. In der Formelsprache bedeutet das:

$$\forall \varepsilon > 0 \exists n_0 \in \mathbb{N}, \text{ so dass } \forall n \geq n_0 \text{ gilt: } |a_n - a| < \varepsilon.$$

Ist $a = 0$, so spricht man von einer „Nullfolge“

3.9. Die Eindeutigkeit des Grenzwertes

Der Grenzwert einer konvergenten Folge ist eindeutig bestimmt.

BEWEIS: Wir nehmen an, es gibt zwei Zahlen a und a' , die beide die Bedingungen der Definition erfüllen.

Zunächst nutzen wir die Voraussetzung aus. Ist ein $\varepsilon > 0$ beliebig vorgegeben, so gibt es Zahlen n_1 und n_2 , so dass $|a_n - a| < \varepsilon$ für $n \geq n_1$ und $|a_n - a'| < \varepsilon$ für $n \geq n_2$ ist. Wir setzen $n_0 := \max(n_1, n_2)$ und versuchen, den Abstand von a und a' nach oben abzuschätzen. Dazu benutzen wir den uralten Trick, eine dritte Zahl – hier ein a_n – zu addieren und gleich wieder zu subtrahieren, so dass man anschließend die Dreiecksungleichung anwenden kann: Für $n \geq n_0$ ist

$$|a - a'| = |(a_n - a') - (a_n - a)| \leq |a_n - a'| + |a_n - a| < 2\varepsilon.$$

Aber eine nicht-negative Zahl, die kleiner als jede positive Zahl der Gestalt 2ε ist, kann nur $= 0$ sein. Also ist $a = a'$. ■

Wir möchten gerne zeigen, dass die Folge $a_n = 1/n$ eine Nullfolge ist. Mit dem Satz des Archimedes ist das kein Problem:

3.10. Satz

$$\lim_{n \rightarrow \infty} \frac{1}{n} = 0.$$

BEWEIS: Sei $\varepsilon > 0$ eine (beliebig kleine) reelle Zahl. Dann gibt es zu der (eventuell recht großen) reellen Zahl $1/\varepsilon$ nach Archimedes immer noch ein $n_0 \in \mathbb{N}$ mit $n_0 > 1/\varepsilon$. Ist $n \geq n_0$, so ist $1/n \leq 1/n_0 < \varepsilon$, also $|1/n - 0| < \varepsilon$. Damit ist alles gezeigt. ■

Bemerkung: Der obige Beweis kann als Muster für viele Konvergenzbeweise dienen.

3.11. Beispiele

- A.** Es soll die Folge $a_n := n/(n+1)$ auf Konvergenz untersucht werden. Dabei sind wir im mathematischen Alltagsgeschäft angekommen: Niemand sagt uns, was der Grenzwert ist. Zum Glück lässt sich der in diesem Fall relativ leicht

erraten, indem man etwa für n einige Werte einsetzt. Wir vermuten, $a = 1$ ist der Grenzwert.

Es kommt nun darauf an, $|a_n - 1| = |1 - n/(n+1)| = 1/(n+1)$ möglichst klein zu bekommen. Das ist aber ein Kinderspiel, weil wir schon wissen, dass die Folge $1/n$ gegen Null konvergiert.

In der richtigen Reihenfolge aufgeschrieben sieht der Beweis nun folgendermaßen aus:

- Sei $\varepsilon > 0$ (beliebig) vorgegeben.
- Nach Archimedes gibt es ein $n_0 \in \mathbb{N}$ mit $n_0 > 1/\varepsilon$.
- Sei $n \geq n_0$. Dann ist $1/(n+1) \leq 1/(n_0+1) < 1/n_0 < \varepsilon$.
- Es folgt, dass $|a_n - 1| < \varepsilon$ ist.

Damit ist gezeigt, dass (a_n) gegen 1 konvergiert.

- B.** Sei $0 < q < 1$. Wir betrachten die Folge $a_n := q^n$. Hier ergeben einige Beispielrechnungen, dass es sich um eine Nullfolge handeln könnte.

Um zu gegebenem ε das richtige n_0 zu finden, versuchen wir es zunächst wieder mit der falschen Schlussrichtung. Die führt allerdings auf die Ungleichung $q^n < \varepsilon$ und damit auf $(1/q)^n > 1/\varepsilon$. Wollte man nun nach n auflösen, so müsste man logarithmieren. Dafür fehlen uns im Augenblick die Grundlagen. Eine bessere Lösung erhält man durch einen kleinen Trick: Weil $1/q > 1$ ist, gibt es ein $x > 0$, so dass $1/q = 1 + x$ ist. Dann liefert die Bernoulli'sche Ungleichung die Abschätzung

$$\left(\frac{1}{q}\right)^n = (1+x)^n \geq 1 + nx.$$

Wir müssen also nur ein n finden, so dass $1 + nx > 1/\varepsilon$ ist.

Jetzt zum eigentlichen Beweis in korrekter Schlussrichtung: Sei $\varepsilon > 0$ vorgegeben. Nach Archimedes gibt es ein $n_0 > (1/\varepsilon - 1)/x$. Für $n \geq n_0$ ist dann $nx \geq n_0x > 1/\varepsilon - 1$, also $(1/q)^n \geq 1 + nx > 1/\varepsilon$ und damit $q^n < \varepsilon$. Die Folge (q^n) konvergiert für $0 < q < 1$ tatsächlich gegen Null.

3.12. Die Monotonie des Grenzwertes

Es seien (a_n) , (b_n) und (c_n) drei Folgen.

1. Ist $a_n \leq b_n$, $\lim_{n \rightarrow \infty} a_n = a$ und $\lim_{n \rightarrow \infty} b_n = b$, so ist auch $a \leq b$.
2. Ist $a_n \leq c_n \leq b_n$ und $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n$, so konvergiert auch (c_n) gegen den gleichen Grenzwert.

BEWEIS: 1) Wir nehmen an, es sei $a > b$, und versuchen, einen Widerspruch herbeizuführen. Dazu benutzen wir die Tatsache, dass die Glieder einer Folge dem Grenzwert beliebig nahe kommen. Sei etwa $\varepsilon := (a - b)/2$. Weil (a_n) gegen a und (b_n) gegen b konvergiert, gibt es ein n_0 , so dass für $n \geq n_0$ gilt:

$$|a_n - a| < \varepsilon \text{ und } |b_n - b| < \varepsilon.$$

Daraus folgt $a_n > a - \varepsilon$ und $b_n < b + \varepsilon$, also

$$a_n - b_n > (a - \varepsilon) - (b + \varepsilon) = a - b - 2\varepsilon = 0.$$

Demnach wäre $a_n > b_n$ für genügend großes n , im Widerspruch zur Voraussetzung.

2) Es sei a der gemeinsame Grenzwert von (a_n) und (b_n) , und es sei ein $\varepsilon > 0$ vorgegeben. Dann gibt es ein n_0 , so dass für $n \geq n_0$ gilt:

$$|a_n - a| < \varepsilon \quad \text{und} \quad |b_n - a| < \varepsilon.$$

Also ist $a - \varepsilon < a_n \leq c_n \leq b_n < a + \varepsilon$ und damit auch $|c_n - a| < \varepsilon$. ■

Definition

Eine Folge (a_n) heißt **beschränkt** (bzw. **nach oben** oder **nach unten beschränkt**), falls die Menge der a_n diese Eigenschaft besitzt.

3.13. Satz

Ist (a_n) konvergent, so ist (a_n) beschränkt.

BEWEIS: Sei a der Grenzwert der Folge. Dann gibt es ein n_0 , so dass $|a_n - a| < 1$ für $n \geq n_0$ ist, also

$$a - 1 < a_n < a + 1 \quad \text{für } n \geq n_0.$$

Da auch die endlich vielen Zahlen $a_1, \dots, a_{n_0}$ eine beschränkte Menge bilden, ist (a_n) insgesamt beschränkt. ■

3.14. Regeln für die Berechnung von Grenzwerten

1. Die Folgen (a_n) bzw. (b_n) seien konvergent gegen a bzw. b . Dann gilt:

$$\lim_{n \rightarrow \infty} (a_n \pm b_n) = a \pm b \quad \text{und} \quad \lim_{n \rightarrow \infty} (a_n \cdot b_n) = a \cdot b.$$

2. Ist (a_n) konvergent gegen eine Zahl $a > 0$, so ist $a_n > a/2$ für fast alle n .

3. Ist (a_n) konvergent gegen a und sind a und fast alle $a_n \neq 0$, so ist

$$\lim_{n \rightarrow \infty} 1/a_n = 1/a.$$

BEWEIS: 1) ($\varepsilon/2$ -Methode):

Sei $\varepsilon > 0$ vorgegeben. Dann gilt für fast alle n : $|a_n - a| < \varepsilon/2$ und $|b_n - b| < \varepsilon/2$, also

$$|(a_n + b_n) - (a + b)| \leq |a_n - a| + |b_n - b| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Hätten wir mit $\dots < \varepsilon$ begonnen, so hätten wir am Ende $\dots < 2\varepsilon$ erhalten. Das wird auch beliebig klein und ist deshalb genauso gut.

Beim Beweis der Produktregel setzen wir gleich $|a_n - a| < \varepsilon$ und $|b_n - b| < \varepsilon$ voraus und verzichten auf besondere Eleganz. Weil (a_n) konvergiert, ist $|a_n|$ durch eine positive Konstante c beschränkt. Für fast alle n ist dann

$$|a_n b_n - ab| = |(a_n(b_n - b) + (a_n - a)b| \leq c \cdot \varepsilon + |b|\varepsilon,$$

und das wird beliebig klein.

2) Ist $a > 0$, so ist auch $\varepsilon := a/2 > 0$ und $|a_n - a| < \varepsilon$ für fast alle n , also

$$-a/2 < a_n - a < a/2,$$

und damit $a/2 < a_n < 3a/2$ für fast alle n .

3) Weil $a \neq 0$ ist, ist $|a_n| > |a|/2$ für fast alle n , und daher

$$|1/a_n - 1/a| = |a - a_n|/|aa_n| < (2/|a|^2) \cdot \varepsilon.$$

Auch hier wird die rechte Seite beliebig klein. ■

Wir werden diese Eigenschaften künftig als **Grenzwertsätze** zitieren.

3.15. Beispiele

A. Wir betrachten die Folge

$$a_n := \frac{3n + 1}{5n - 2} = \frac{3 + 1/n}{5 - 2/n}.$$

Weil $1/n$ und $2/n$ jeweils gegen Null konvergieren, konvergiert $3 + 1/n$ gegen 3 und $5 - 2/n$ gegen 5, also a_n gegen $3/5$.

B. Ein weiteres typisches Anwendungsbeispiel ist die Folge

$$a_n := \frac{n(n - 2)}{5n^2 + 3}.$$

Kürzen durch n^2 ergibt $a_n = (1 - 2/n)/(5 + 3/n^2)$.

Weil $1/n$ eine Nullfolge ist, konvergiert der Zähler gegen 1 und der Nenner gegen 5, also (a_n) gegen $1/5$.

Definition

Eine Folge (a_n) heißt **monoton wachsend** (bzw. **monoton fallend**), falls gilt:

$$a_n \leq a_{n+1} \quad (\text{bzw. } a_n \geq a_{n+1}) \quad \text{für (fast) alle } n.$$

3.16. Satz von der monotonen Konvergenz

Ist (a_n) monoton wachsend und nach oben beschränkt, so ist (a_n) konvergent.

BEWEIS: Die Menge $M := \{a_n : n \in \mathbb{N}\}$ ist nicht leer und nach oben beschränkt. Also ist $a := \sup(M)$ eine reelle Zahl. Sie ist unser Kandidat für den Grenzwert.

Sei $\varepsilon > 0$ vorgegeben. Dann ist $a - \varepsilon$ keine obere Schranke mehr, und es gibt ein n_0 mit $a - \varepsilon < a_{n_0}$. Nun benutzen wir die Monotonie. Für $n \geq n_0$ ist $a_{n_0} \leq a_n \leq a$, also

$$0 \leq a - a_n \leq a - a_{n_0} < \varepsilon.$$

Das bedeutet, dass (a_n) gegen a konvergiert. ■

Genauso zeigt man, dass eine monoton fallende nach unten beschränkte Folge konvergiert. Natürlich reicht es, wenn die Monotonie erst ab einem gewissen n_0 gilt.

Das Interessante am Satz von der monotonen Konvergenz ist, dass man die Konvergenz einer Folge zeigen kann, ohne den Grenzwert kennen zu müssen.

3.17. Beispiele

- A. Sei $a > 0$. Wir definieren rekursiv eine Folge (x_n) . Die Zahl $x_0 > 0$ kann beliebig gewählt werden, und dann sei

$$x_{n+1} := \frac{1}{2} \left(x_n + \frac{a}{x_n} \right).$$

Wir werden mit Hilfe des Satzes von der monotonen Konvergenz zeigen, dass (x_n) konvergiert, und anschließend den Grenzwert bestimmen.

Offensichtlich ist $x_n > 0$ für alle n , also (x_n) nach unten beschränkt. Außerdem ist

$$\begin{aligned} x_n - x_{n+1} &= x_n - \frac{1}{2} \left(x_n + \frac{a}{x_n} \right) \\ &= \frac{2x_n^2 - x_n^2 - a}{2x_n} = \frac{x_n^2 - a}{2x_n} \geq 0, \end{aligned}$$

denn es ist

$$\begin{aligned}
x_n^2 - a &= \frac{1}{4} \left(x_{n-1} + \frac{a}{x_{n-1}} \right)^2 - a \\
&= \frac{1}{4} \left(x_{n-1}^2 + 2a + \frac{a^2}{x_{n-1}^2} - 4a \right) \\
&= \frac{1}{4} \left(x_{n-1} - \frac{a}{x_{n-1}} \right)^2 \geq 0.
\end{aligned}$$

Damit ist (x_n) monoton fallend, also konvergent gegen eine reelle Zahl c .

Offensichtlich muss $c \geq 0$ sein. Wäre $c = 0$, so würde auch (x_n^2) gegen Null konvergieren. Das ist aber nicht möglich, da stets $x_n^2 \geq a > 0$ ist. Also ist $c > 0$. Da auch (x_{n+1}) gegen c konvergiert, folgt die Gleichung

$$c = \frac{1}{2} \left(c + \frac{a}{c} \right).$$

Es ergibt sich $2c^2 = c^2 + a$, also $c^2 = a$ und damit $c = \sqrt{a}$.

- B.** Die Folge $a_n := (1 + 1/n)^n$ konvergiert **nicht** gegen 1, wie man vermuten könnte (obwohl $1/n$ gegen Null, also $1 + 1/n$ gegen 1 konvergiert und 1^n stets = 1 ist). In Wirklichkeit kann man in diesem Fall nicht mit den Grenzwertsätzen argumentieren.

Mit Hilfe der Bernoulli'schen Ungleichung erhält man für alle $n \geq 2$

$$\begin{aligned}
\frac{a_n}{a_{n-1}} &= \left(\frac{n+1}{n} \right)^n \cdot \left(\frac{n-1}{n} \right)^{n-1} = \left(\frac{n^2-1}{n^2} \right)^n \cdot \frac{n}{n-1} \\
&= \left(1 - \frac{1}{n^2} \right)^n \cdot \frac{n}{n-1} > \left(1 - n \cdot \frac{1}{n^2} \right) \cdot \frac{n}{n-1} = 1.
\end{aligned}$$

Also ist (a_n) monoton wachsend.

Mit der Abschätzung

$$\frac{1}{k!} = \frac{1}{2} \cdot \frac{1}{3} \cdot \dots \cdot \frac{1}{k} < \left(\frac{1}{2} \right)^{k-1} \text{ für } k \geq 3$$

und der binomischen Formel folgt für $n \geq 3$:

$$\begin{aligned}
\left(1 + \frac{1}{n} \right)^n &= \sum_{k=0}^n \binom{n}{k} \frac{1}{n^k} = 1 + \sum_{k=1}^n \frac{n!}{k!(n-k)!} \cdot \frac{1}{n^k} \\
&= 1 + \sum_{k=1}^n \frac{1}{k!} \cdot \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{n \cdot n \cdot \dots \cdot n} \\
&< 1 + \sum_{k=1}^n \left(\frac{1}{2} \right)^{k-1} = 1 + \frac{(1/2)^n - 1}{(1/2) - 1} < 3.
\end{aligned}$$

Damit ist (a_n) nach oben beschränkt, also konvergent. Der Grenzwert

$$e = 2.718281 \dots$$

wird als **Euler'sche Zahl** bezeichnet. Die Existenz dieser Zahl, von der wir vorher nichts ahnen konnten, ergibt sich letztendlich aus dem Vollständigkeitsaxiom (wie die Existenz der Wurzeln). Ohne Vollständigkeitsaxiom gäbe es keine irrationalen Zahlen!

Betrachten wir jetzt die Folge

$$a_n := (-1)^n = \begin{cases} 1 & \text{falls } n \text{ gerade,} \\ -1 & \text{sonst.} \end{cases}$$

Diese Folge kann nicht konvergieren. Es ist aber etwas mühsam, das zu begründen. Wir suchen deshalb nach einfachen Kriterien, die belegen, dass eine Folge **nicht** konvergiert.

Wir wissen schon, dass jede konvergente Folge beschränkt ist. Im Umkehrschluss kann eine unbeschränkte Folge niemals konvergent sein. Leider hilft uns das bei unserem Beispiel nicht weiter, denn die Folge $a_n = (-1)^n$ ist beschränkt.

Definition

Eine Zahl $a \in \mathbb{R}$ heißt **Häufungspunkt** der Folge (a_n) , falls in jeder ε -Umgebung von a unendlich viele a_n liegen.

Jede konvergente Folge hat einen Häufungspunkt, nämlich ihren Grenzwert. Die Folge $a_n = (-1)^n$ besitzt **zwei** Häufungspunkte, nämlich $+1$ und -1 . Das scheint der Knackpunkt zu sein! Zunächst stellen wir folgende erstaunliche Tatsache fest:

3.18. Satz von Bolzano–Weierstraß

Jede beschränkte Folge besitzt mindestens einen Häufungspunkt.

BEWEIS: Sei (a_n) eine beschränkte Folge. Dann ist $c_n := \inf\{a_n, a_{n+1}, \dots\}$ eine monoton wachsende Folge, die nach oben beschränkt ist. Nach dem Satz von der monotonen Konvergenz strebt (c_n) gegen eine reelle Zahl c .

Sei $\varepsilon > 0$. Dann gibt es ein n_0 , so dass $c - \varepsilon < c_n \leq c$ für $n \geq n_0$ gilt. Nach Definition von c_n ist $c_n + \varepsilon$ keine untere Schranke der Menge $\{a_n, a_{n+1}, \dots\}$. Insbesondere gibt es zu jedem $n \geq n_0$ ein $m \geq n$, so dass $c_n \leq a_m < c_n + \varepsilon \leq c + \varepsilon$ ist, also $|a_m - c| < \varepsilon$. Das heißt, dass c ein Häufungspunkt von (a_n) ist. ■

3.19. Satz

Eine konvergente Folge hat nur einen Häufungspunkt.

BEWEIS: Sei (a_n) eine konvergente Folge mit Grenzwert a . Dann ist a auch ein Häufungspunkt. Annahme, es gibt einen weiteren Häufungspunkt $b \neq a$. Wählen wir ein ε mit $0 < \varepsilon < |b - a|/2$, so ist $U_\varepsilon(a) \cap U_\varepsilon(b) = \emptyset$ und fast alle a_n liegen in $U_\varepsilon(a)$. Da dann für $U_\varepsilon(b)$ nur noch höchstens endlich viele a_n übrigbleiben, ist das ein Widerspruch. ■

Zusammengefasst erhalten wir das folgende

3.20. Divergenzkriterium

Eine Folge (a_n) ist genau dann divergent (d.h. nicht konvergent), wenn eine der beiden folgenden Bedingungen erfüllt ist:

1. (a_n) ist unbeschränkt.
2. (a_n) ist beschränkt und hat mindestens zwei verschiedene Häufungspunkte.

BEWEIS: Wenn das Kriterium erfüllt ist, kann (a_n) nicht konvergieren.

Sei umgekehrt die Folge (a_n) divergent, aber beschränkt. Hätte (a_n) nur einen Häufungspunkt, der nach Voraussetzung nicht der Limes der Folge sein kann, so müsste es ein $\varepsilon > 0$ und unendlich viele a_n mit $|a_n - a| \geq \varepsilon$ geben. Diese a_n besitzen nach dem Satz von Bolzano-Weierstraß mindestens einen Häufungspunkt, der dann auch Häufungspunkt der Ausgangsfolge und zugleich $\neq a$ ist. Das widerspricht der Annahme. ■

Ist (a_ν) eine Folge reeller Zahlen und $(\nu(i))$ eine Folge von natürlichen Zahlen mit $\nu(i) < \nu(i+1)$ für alle i , so nennt man die Folge $(a_{\nu(i)})$ (oft auch in der Form (a_{ν_i}) geschrieben) eine **Teilfolge** der ursprünglichen Folge. Ist (a_ν) konvergent, so konvergiert auch jede Teilfolge von (a_ν) gegen den gleichen Grenzwert.

3.21. Satz

Ist a ein Häufungspunkt der Folge (a_n) , so gibt es eine Teilfolge von (a_n) , die gegen a konvergiert.

BEWEIS: Sei $U_n := U_{1/n}(a)$. Wir konstruieren eine Teilfolge $b_k := a_{n_k}$ induktiv wie folgt:

Es gibt einen kleinsten Index n_1 , so dass a_{n_1} in U_1 liegt. Dann sei $b_1 := a_{n_1}$.

Sind Indizes $n_1 < n_2 < \dots < n_k$ gefunden, so dass jeweils $b_\mu := a_{n_\mu}$ in U_μ liegt, für $\mu = 1, \dots, k$, dann liegen in U_{k+1} noch immer unendlich viele Folgenglieder a_n mit $n > n_k$. Ist $n_{k+1} > n_k$ der kleinste Index, so dass $a_{n_{k+1}}$ in U_{k+1} liegt, so setzen wir $b_{k+1} := a_{n_{k+1}}$.

Die Folge (b_k) konvergiert gegen a und ist nach Konstruktion eine Teilfolge von (a_n) . ■

Definition

Sei $M \subset \mathbb{R}$ eine beliebige Teilmenge. Ein Punkt $a \in M$ heißt **innerer Punkt** von M , falls es ein $\varepsilon > 0$ gibt, so dass die ε -Umgebung von a ganz in M liegt.

Die Menge M heißt **offen**, falls sie nur aus inneren Punkten besteht.

3.22. Satz

Eine Menge $M \subset \mathbb{R}$ ist genau dann offen, wenn M Vereinigung von offenen Intervallen ist.

BEWEIS: 1) Jedes offene Intervall (a, b) ist offen:

Ist $y \in (a, b)$, so ist $a < y < b$. Wählt man $0 < \varepsilon < \min(y - a, b - y)$, so liegt $U_\varepsilon(y)$ in (a, b) .

Damit ist (a, b) und auch jede Vereinigung von solchen Intervallen offen.

2) Sei $M \subset \mathbb{R}$ offen. Dann gibt es zu jedem $y \in M$ ein $\varepsilon = \varepsilon(y) > 0$, so dass $I_\varepsilon = (y - \varepsilon, y + \varepsilon)$ in M liegt. Die Vereinigung aller dieser Intervalle ist die Menge M . ■

Ist die offene Menge M Vereinigung einer Familie von Intervallen I_j , $j \in J$, so kann man in jedem I_j eine rationale Zahl q_j finden. Da die Menge der rationalen Zahlen abzählbar ist, kann es höchstens abzählbar viele „Komponenten“ I_j geben.

Definition

Eine Menge $K \subset \mathbb{R}$ heißt **kompakt**, falls jede Folge von Zahlen $x_\nu \in K$ eine Teilfolge $(x_{\nu(i)})$ besitzt, die gegen ein $x_0 \in K$ konvergiert.

3.23. Satz

Eine Menge $K \subset \mathbb{R}$ ist genau dann kompakt, wenn sie beschränkt und die Menge $\mathbb{R} \setminus K$ offen ist.¹

BEWEIS: 1) Sei K beschränkt, $\mathbb{R} \setminus K$ offen und (x_ν) eine Folge in K . Nach dem Satz von Bolzano-Weierstraß besitzt (x_ν) einen Häufungspunkt x_0 , und dann gibt es eine Teilfolge $(x_{\nu(i)})$, die gegen x_0 konvergiert.

Wäre x_0 ein Punkt der offenen Menge $\mathbb{R} \setminus K$, so lägen schon fast alle $x_{\nu(i)}$ in $\mathbb{R} \setminus K$. Das kann nicht sein, also muss $x_0 \in K$ liegen. Damit ist gezeigt, dass K kompakt ist.

¹Eine Menge, deren Komplement offen ist, nennt man auch **abgeschlossen**. Da die Geometrie abgeschlossener Mengen deutlich komplizierter als die der offenen Mengen ist, verschieben wir deren Behandlung auf den Anfang des 2. Semesters.

2) Sei umgekehrt K kompakt.

Wäre K unbeschränkt, so gäbe es zu jedem $\nu \in \mathbb{N}$ ein Element $x_\nu \in K$ mit $|x_\nu| > \nu$. Dann ist auch jede Teilfolge von (x_ν) unbeschränkt, kann also nicht konvergent sein. Das ist ein Widerspruch.

Sei $x_0 \in \mathbb{R} \setminus K$. Ist x_0 kein innerer Punkt von $\mathbb{R} \setminus K$, so enthält jede $(1/\nu)$ -Umgebung von x_0 ein Element $x_\nu \in K$. Diese Folge konvergiert offensichtlich gegen x_0 . Aber weil K kompakt ist, muss es auch eine Teilfolge $(x_{\nu(i)})$ geben, die gegen ein $y_0 \in K$ konvergiert. Das kann nicht sein, denn es wäre $x_0 = y_0$. Damit ist gezeigt, dass $\mathbb{R} \setminus K$ offen ist. ■

3.24. Beispiele

- A. Jedes abgeschlossene Intervall $[a, b]$ ist kompakt, denn es ist beschränkt, und die Menge $\mathbb{R} \setminus [a, b] = (-\infty, a) \cup (b, +\infty)$ ist offen.
- B. Jede endliche Menge $M = \{x_1, \dots, x_N\}$ (mit $x_1 < x_2 < \dots < x_N$) ist kompakt. Offensichtlich ist sie beschränkt, und

$$\mathbb{R} \setminus M = (-\infty, x_1) \cup (x_1, x_2) \cup \dots \cup (x_N, +\infty)$$

ist offen.

- C. Sei (x_ν) eine konvergente Folge mit Grenzwert x_0 . Dann ist

$$K := \{x_0\} \cup \{x_\nu : \nu \in \mathbb{N}\}$$

kompakt. Das sieht man folgendermaßen:

Sei (y_i) eine Folge in K . Enthält sie unendlich viele x_ν , so ist sie eine Teilfolge von (x_ν) und konvergiert gegen $x_0 \in K$. Enthält sie nur endlich viele x_ν , so muss eine dieser Zahlen unendlich oft vorkommen. Diese ist dann ein Häufungspunkt von (y_i) und zugleich ein Element von K , und es gibt eine Teilfolge, die dagegen konvergiert.

1.4 Funktionen

Unter dem **Paar** (x, y) versteht man die Menge $\{\{x, y\}, x\}$, bestehend aus einer 2-elementigen Menge $\{x, y\}$ und dem ersten Element $\{x\}$. Diese Definition liefert genau das, was man sich unter einem Paar vorstellt. Insbesondere ist $(x, y) = (x', y')$ genau dann wahr, wenn $x = x'$ und $y = y'$ ist. Die Elemente x und y nennen wir die beiden **Komponenten** (oder **Koordinaten**) des Paares (x, y) .

Definition

Die Menge $A \times B$ aller Paare (x, y) mit $x \in A$ und $y \in B$ wird als **kartesisches Produkt** oder **Produktmenge** von A und B bezeichnet.

Statt $A \times A$ kann man auch A^2 schreiben. Es bleibt trotzdem dabei, dass man bei einem Paar nicht die Komponenten vertauschen darf.

Im Falle $A = \mathbb{R}$ erhält man mit $\mathbb{R}^2$ ein Modell für die Anschauungsebene, die wir für graphische Darstellungen benutzen.

Definition

A und B seien zwei nicht leere Mengen. Unter einer **Funktion** oder **Abbildung** $f : A \rightarrow B$ versteht man eine Vorschrift, nach der jedem Element $x \in A$ genau ein Element $y \in B$ zugeordnet wird. Die Menge A nennt man den **Definitionsbereich** der Funktion f und bezeichnet sie auch mit D_f . Die Menge B nennt man den **Wertebereich** von f .

Wird dem einzelnen Element x durch f das Element y zugeordnet, so schreibt man

$$y = f(x), \quad f : x \mapsto y \quad \text{oder} \quad x \xrightarrow{f} y.$$

Das für uns wichtigste Beispiel sind die reellen Funktionen von einer Veränderlichen, also Funktionen, deren Definitionsbereich eine Teilmenge von $\mathbb{R}$ (z.B. ein Intervall) und deren Wertebereich ganz $\mathbb{R}$ ist.

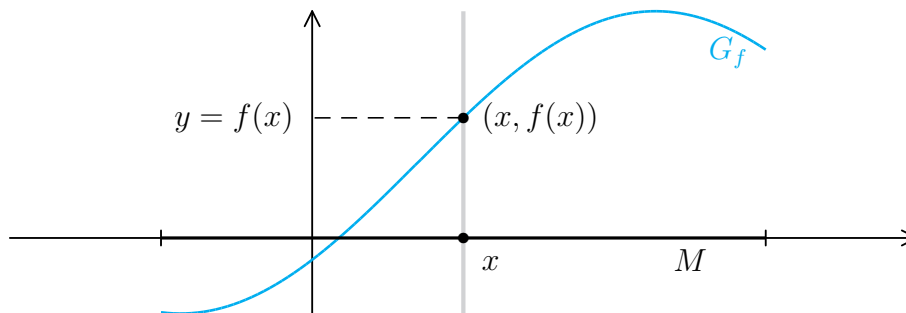
Definition

Sei $M \subset \mathbb{R}$ eine Teilmenge und $f : M \rightarrow \mathbb{R}$ eine Funktion. Dann nennt man die Menge

$$G_f := \{(x, y) \in M \times \mathbb{R} : y = f(x)\}$$

den **Graphen** von f .

Eine Teilmenge $G \subset M \times \mathbb{R}$ ist genau dann Graph einer Funktion $f : M \rightarrow \mathbb{R}$, wenn jede vertikale Gerade durch einen Punkt $(x, 0)$ (mit $x \in M$) die Menge G in genau einem Punkt (x, y) trifft. Denn das bedeutet ja gerade, dass jedem $x \in M$ genau ein $y \in \mathbb{R}$ zugeordnet wird, so dass (x, y) zu G gehört.



4.1. Beispiele

- A. Besonders einfach ist z.B. eine **affin-lineare Funktion** $f : \mathbb{R} \rightarrow \mathbb{R}$, definiert durch

$$f(x) := mx + c, \text{ mit } m \neq 0.$$

Hier ist $D_f = \mathbb{R}$, und die Zuordnung f wird durch den Rechenausdruck $mx + c$ beschrieben. Wenn $c = 0$ ist, spricht man von einer **linearen Funktion**.

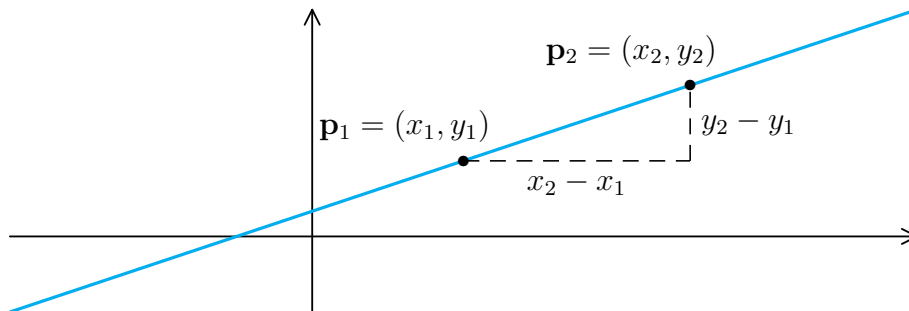
Der Graph einer affin-linearen Funktion f mit $D_f = \mathbb{R}$ und $f(x) := mx + c$ ist eine Gerade

$$L = \{(x, y) \in \mathbb{R}^2 : y = mx + c\}.$$

Der Faktor m wird **Steigung** der Geraden genannt. Ist $m = 0$, so liegt eine konstante Funktion (gegeben durch $f(x) \equiv c$) und bei L dann eine horizontale Gerade vor. Eine vertikale Gerade kann kein Graph sein. Die Funktion ist genau dann linear, wenn der Graph durch den Nullpunkt geht.

Sind $\mathbf{p}_1 = (x_1, y_1)$ und $\mathbf{p}_2 = (x_2, y_2)$ zwei Punkte auf der Geraden, so kann man die Steigung nach folgender Formel berechnen:

$$m = \frac{y_2 - y_1}{x_2 - x_1}.$$

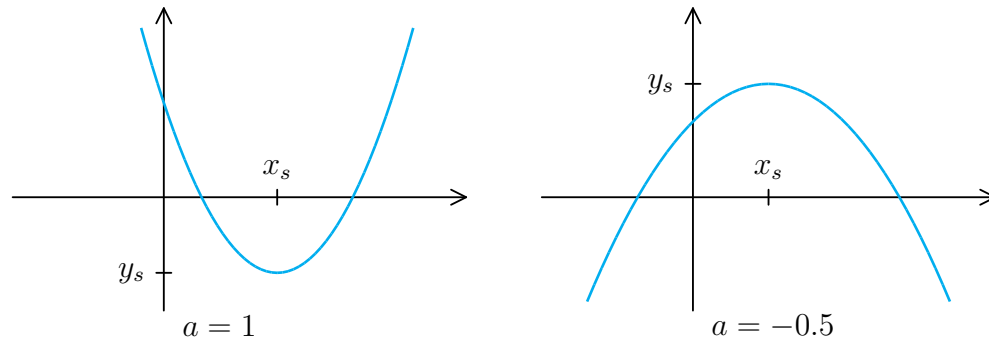


- B. Unter einer **quadratischen Funktion** versteht man eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

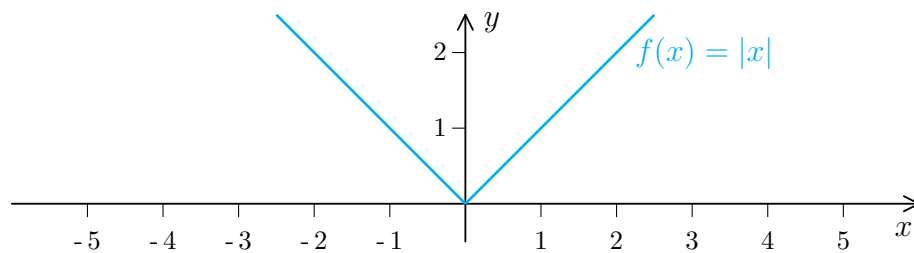
$$f(x) := ax^2 + bx + c, \text{ mit } a \neq 0.$$

Den Graphen einer quadratischen Funktion nennt man eine **Parabel**.

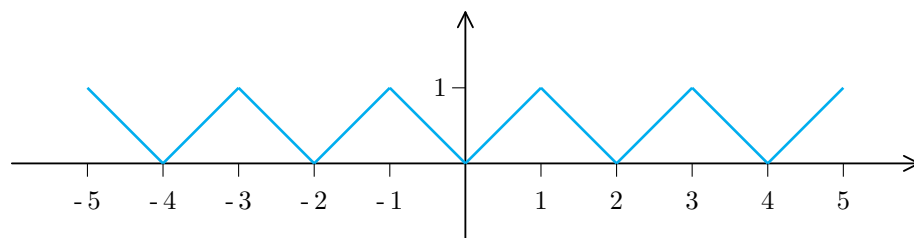
Ist $a > 0$, so ist die Parabel nach oben geöffnet. Ist $a < 0$, so ist sie nach unten geöffnet.



C. Die **Betragsfunktion** $f(x) := |x|$:



D. Die **Zickzackfunktion** Z :



Die Funktion Z ist ein typisches Beispiel einer „zusammengesetzten Funktion“. Der Graph von Z besteht aus Geradenstücken. Ist $2n \leq x < 2n + 1$, so ist $Z(x) = x - 2n$. Ist $2n + 1 \leq x < 2n + 2$, so ist $Z(x) = -x + (2n + 2)$. Zusammen ergibt das:

$$Z(x) = \begin{cases} x - 2n & \text{für } 2n \leq x < 2n + 1 \\ -x + (2n + 2) & \text{für } 2n + 1 \leq x < 2n + 2. \end{cases}$$

Definition

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **gerade** (bzw. **ungerade**), wenn $f(-x) = f(x)$ (bzw. $f(-x) = -f(x)$) für alle x ist.

Jede quadratische Funktion der Gestalt $f(x) = ax^2 + c$ ist ein typisches Beispiel für eine gerade Funktion. Die durch $f(x) := ax^3 + cx$ gegebene Funktion ist ungerade.

Weitere Beispiele für gerade Funktionen sind die Betragsfunktion und die Zickzackfunktion.

Definition

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **periodisch** (mit Periode T), falls $f(x+T) = f(x)$ für alle $x \in \mathbb{R}$ ist.

Die Zickzack-Funktion Z ist periodisch mit Periode 2.

Definition

Sei $M \subset \mathbb{R}$. Eine Funktion $f : M \rightarrow \mathbb{R}$ heißt **monoton wachsend** (bzw. **streng monoton wachsend**), wenn für $x_1, x_2 \in M$ mit $x_1 < x_2$ stets $f(x_1) \leq f(x_2)$ (bzw. $f(x_1) < f(x_2)$) ist.

Man nennt f **monoton fallend** (bzw. **streng monoton fallend**), wenn für $x_1, x_2 \in M$ mit $x_1 < x_2$ stets $f(x_1) \geq f(x_2)$ (bzw. $f(x_1) > f(x_2)$) ist.

Die affin-lineare Funktion $f(x) = mx + c$ ist genau dann streng monoton wachsend (bzw. fallend), wenn $m > 0$ (bzw. < 0) ist. Ist $m = 0$, also $f(x) = c$ eine **konstante Funktion**, so ist f gleichzeitig monoton wachsend und fallend. Die Zickzack-Funktion ist abwechselnd streng monoton wachsend und fallend.

Aus der Schule sind weitere Funktionen bekannt, wie etwa Sinus, Cosinus, Exponentialfunktion und Logarithmus. Eine exakte Definition dieser Funktionen können wir erst später geben.

Eine Zahlenfolge ist eigentlich nichts anderes als eine Abbildung $a : \mathbb{N} \rightarrow \mathbb{R}$. Statt $a(n)$ schreibt man a_n .

Eine **Permutation** ist eine Abbildung $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$ mit der Eigenschaft, dass $\{\sigma(1), \dots, \sigma(n)\} = \{1, \dots, n\}$ ist. Wir wissen schon, dass es genau $n!$ Permutationen von $\{1, \dots, n\}$ gibt.

Definition

Sind f und g reellwertige Funktionen mit Definitionsbereichen D_f bzw. D_g in $\mathbb{R}$, so sind die Funktionen $f + g$ und $f \cdot g$ auf $D_f \cap D_g$ definiert durch

$$(f + g)(x) := f(x) + g(x) \\ \text{und} \quad (f \cdot g)(x) := f(x) \cdot g(x).$$

Für alle $x \in D_f \cap D_g$, für die $g(x) \neq 0$ ist, definiert man außerdem die Funktion f/g durch

$$(f/g)(x) := f(x)/g(x).$$

Eine weitere Manipulation, die man mit Funktionen vornehmen kann, ist die Verkettung von Funktionen. Das untersuchen wir in einem allgemeineren Kontext.

Definition

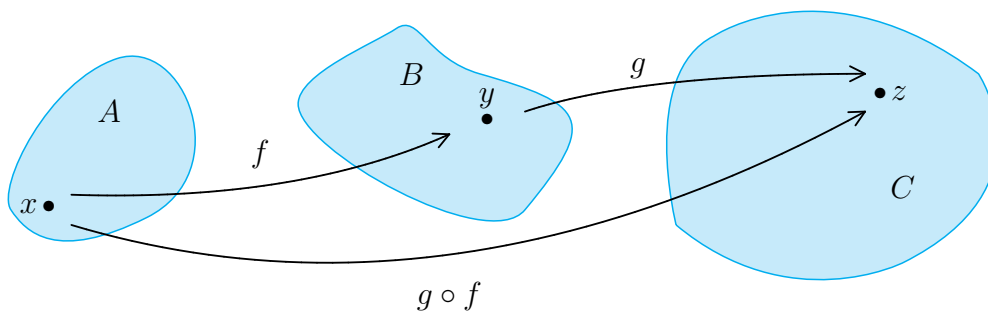
Es seien $f : A \rightarrow B$ und $g : B \rightarrow C$ zwei Abbildungen. Hintereinander ausgeführt ergeben sie eine neue Abbildung $g \circ f : A \rightarrow C$, die durch

$$(g \circ f)(x) := g(f(x)) \quad (\text{für } x \in A)$$

definiert wird.

Man nennt $g \circ f$ die **Verknüpfung** (oder **Verkettung**) von g mit f .

Eine Skizze kann die Situation vielleicht verdeutlichen:



Einem Element $x \in A$ wird also zunächst ein Element $y = f(x) \in B$ zugeordnet, und diesem y wird seinerseits das Element $z = g(y)$ zugeordnet. Insgesamt ist dann $z = g(f(x))$. Obwohl man **zuerst** die Zuordnung f und **dann** die Zuordnung g ausführt, schreibt man die Verknüpfung in der Form $g \circ f$. Hier muss man ausnahmsweise von rechts nach links lesen. Aus der Definition wird klar, dass das so sein muss, aber es ist auch immer ein wenig verwirrend. Dass die Reihenfolge eine wichtige Rolle spielt, kann man sofort sehen:

4.2. Beispiel

Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch $f(x) := x^2 + 1$ und $g : \{y \in \mathbb{R} : y \geq 2\} \rightarrow \mathbb{R}$ durch $g(y) := \sqrt{y - 2}$. Dann ist

$$\begin{aligned} (g \circ f)(x) &= \sqrt{x^2 - 1}, \quad \text{für } x \geq 1, \text{ und} \\ (f \circ g)(y) &= y - 1, \quad \text{für } y \geq 2. \end{aligned}$$

Ist $f : A \rightarrow B$ eine Abbildung und $M \subset A$, so heißt die Menge

$$f(M) := \{f(x) : x \in M\} = \{y \in B : \exists x \in M \text{ mit } y = f(x)\}$$

das *Bild* oder die **Bildmenge** von M unter f . Die Menge

$$f^{-1}(N) := \{x \in A : f(x) \in N\}$$

heißt das **Urbild** von N unter f .

4.3. Beispiel

Sei $f : [0, 2\pi] \rightarrow \mathbb{R}$ definiert durch $f(x) := 2 \sin x + 1$. Da $\sin x$ alle Werte zwischen -1 und $+1$ annimmt (wie wir später beweisen werden), ist $f(\mathbb{R}) = [-1, 3]$.

Außerdem ist

$$\begin{aligned} f^{-1}(\{y \in \mathbb{R} : y \geq 0\}) &= \{x \in [0, 2\pi] : f(x) \geq 0\} \\ &= \{x \in [0, 2\pi] : \sin x \geq -1/2\} \\ &= [0, 7\pi/6] \cup [11\pi/6, 2\pi]. \end{aligned}$$

und

$$\begin{aligned} f^{-1}([5, 7]) &= \{x \in [0, 2\pi] : 5 \leq f(x) \leq 7\} \\ &= \{x \in [0, 2\pi] : 2 \leq \sin x \leq 3\} = \emptyset. \end{aligned}$$

Das Urbild einer nicht leeren Menge kann also durchaus leer sein.

Definition

Eine Abbildung $f : A \rightarrow B$ heißt **surjektiv**, falls gilt:

$$f(A) = B, \text{ d.h. zu jedem } y \in B \text{ gibt es ein } x \in A \text{ mit } f(x) = y.$$

Eine Abbildung $f : A \rightarrow B$ heißt **injektiv**, falls für alle $x_1, x_2 \in A$ gilt:

$$\text{Ist } x_1 \neq x_2, \text{ so ist auch } f(x_1) \neq f(x_2).$$

f ist genau dann **surjektiv**, wenn die Gleichung $f(x) = y$ **immer lösbar** ist. Eindeutige Lösbarkeit ist dabei nicht erforderlich. Der Graph einer surjektiven Funktion von einer reellen Veränderlichen wird von jeder horizontalen Geraden mindestens einmal getroffen.

f ist genau dann **injektiv**, wenn die Gleichung $f(x) = y$ für jedes $y \in B$ **höchstens eine Lösung** besitzt, wenn also die Mengen $f^{-1}(\{y\})$ immer höchstens aus einem Element bestehen. Dass es überhaupt keine Lösung gibt, ist dabei durchaus erlaubt. Der Graph einer injektiven Funktion wird von jeder horizontalen Geraden höchstens einmal getroffen.

Den Nachweis der Injektivität einer Abbildung führt man meist durch Kontraposition, d.h. man zeigt: Ist $f(x_1) = f(x_2)$, so ist $x_1 = x_2$.

4.4. Beispiele

- A.** Ist $a \neq 0$, so ist die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) := ax + b$ surjektiv. Denn die Gleichung $y = ax + b$ wird durch $x = (y - b)/a$ gelöst. f ist auch injektiv: Sei etwa $f(x_1) = f(x_2)$. Dann ist $ax_1 + b = ax_2 + b$, also $a(x_1 - x_2) = 0$. Da $a \neq 0$ vorausgesetzt wurde, muss $x_1 = x_2$ sein.
- B.** Die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) := x^2$ ist nicht surjektiv, denn negative Zahlen können nicht als Bild vorkommen. Dagegen ist die gleiche Abbildung mit dem Wertebereich $\{y \in \mathbb{R} : y \geq 0\}$ surjektiv. Die Gleichung $y = x^2$ wird dann durch $x = \sqrt{y}$ und $x = -\sqrt{y}$ gelöst. f ist ebenfalls nicht injektiv! Für $x \neq 0$ ist nämlich $-x \neq x$, aber $f(-x) = f(x)$.

Ist allgemein $f : A \rightarrow B$ eine Abbildung und $M \subset A$, so definiert man die **Einschränkung** von f auf M (in Zeichen: $f|_M$) als diejenige Abbildung von M nach B , die durch $(f|_M)(x) := f(x)$ gegeben wird. Ist $f(x) = x^2$ und $M := \{x \in \mathbb{R} : x \geq 0\}$, so ist $f|_M$ injektiv, denn die Gleichung $y = x^2$ besitzt nur eine Lösung (nämlich $x = \sqrt{y}$) in M .

Allgemein gilt: Ist $f : A \rightarrow B$ injektiv und $M \subset A$, so ist auch $f|_M$ injektiv. Ist umgekehrt $f|_M$ surjektiv, so ist auch f surjektiv.

Abbildungen, die sowohl injektiv als auch surjektiv sind, bei denen also die Gleichung $f(x) = y$ für jedes $y \in B$ eindeutig lösbar ist, spielen eine ganz besondere Rolle:

Definition

Eine Abbildung $f : A \rightarrow B$ heißt **bijektiv**, falls sie injektiv und surjektiv ist.

Von den oben (in A bis C) betrachteten Beispielen sind nur $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) := ax + b$ (und $a \neq 0$) und $f : \mathbb{R}_{\geq 0} := \{x \in \mathbb{R} : x \geq 0\} \rightarrow \mathbb{R}_{\geq 0}$ mit $f(x) := x^2$ bijektiv.

Wenn es zwischen A und B eine bijektive Abbildung gibt, so ist nicht nur jedem $x \in A$ genau ein $y \in B$ zugeordnet, sondern umgekehrt auch jedem $y \in B$ genau ein $x \in A$, die eindeutig bestimmte Lösung x der Gleichung $f(x) = y$. So erhalten wir automatisch auch eine Abbildung von B nach A , die **Umkehrabbildung**

$$f^{-1} : B \rightarrow A, \text{ mit } f^{-1}(f(x)) = x.$$

Bemerkung: Ist $f : A \rightarrow B$ eine beliebige Abbildung, so kann man für jede Teilmenge $N \subset B$ die Urbildmenge $f^{-1}(N)$ bilden. Ist f injektiv, so ist $f^{-1}(\{y\})$ immer entweder leer oder eine Menge mit einem einzigen Element. Nur wenn f außerdem noch surjektiv ist, existiert die Umkehrabbildung $f^{-1} : B \rightarrow A$.

4.5. Beispiele

- A. Sei $a \neq 0$. Die Umkehrabbildung der bijektiven Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) := ax + b$ ist gegeben durch $f^{-1}(y) := \frac{1}{a}(y - b)$.
- B. Die Umkehrabbildung der Abbildung $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ mit $f(x) := x^2$ ist gegeben durch $f^{-1}(y) := \sqrt{y}$.
- C. Sei A eine beliebige Menge. Dann wird die Abbildung $\text{id}_A : A \rightarrow A$ definiert durch $\text{id}_A(x) := x$. Man spricht von der **identischen Abbildung** oder der **Identität** auf A . Offensichtlich ist die Identität bijektiv und ihre Umkehrabbildung wieder die Identität.

Ist $f : A \rightarrow B$ bijektiv, so ist $f^{-1} \circ f = \text{id}_A$ und $f \circ f^{-1} = \text{id}_B$.

4.6. Satz

Es sei eine Abbildung $f : A \rightarrow B$ gegeben. Wenn es eine Abbildung $g : B \rightarrow A$ mit $g \circ f = \text{id}_A$ und $f \circ g = \text{id}_B$ gibt, so ist f bijektiv und $f^{-1} = g$.

BEWEIS: a) Sei $y \in B$ vorgegeben. Dann ist $x := g(y) \in A$ und $f(x) = f(g(y)) = y$. Also ist f surjektiv.

b) Seien $x_1, x_2 \in A$, mit $f(x_1) = f(x_2)$. Dann ist

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

Also ist f injektiv und damit sogar bijektiv.

Sei $y \in B$. Setzt man $x := g(y)$, so ist $f(x) = f \circ g(y) = y$, also $f^{-1}(y) = x = g(y)$.

■

4.7. Satz

Sind die Abbildungen $f : A \rightarrow B$ und $g : B \rightarrow C$ beide bijektiv, so ist auch $g \circ f : A \rightarrow C$ bijektiv, und

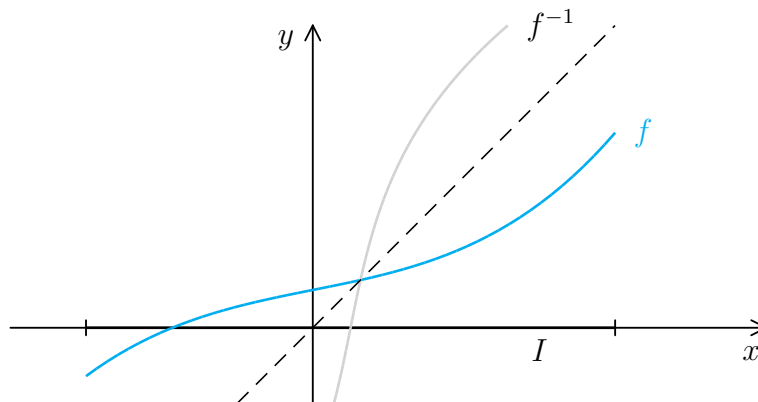
$$(g \circ f)^{-1} = f^{-1} \circ g^{-1}.$$

BEWEIS: Zu f und g existieren Umkehrabbildungen f^{-1} und g^{-1} . Dann sei $F := f^{-1} \circ g^{-1} : C \rightarrow A$. Offensichtlich ist $F \circ (g \circ f) = \text{id}_A$ und $(g \circ f) \circ F = \text{id}_C$, also $g \circ f$ bijektiv und $F = (g \circ f)^{-1}$. ■

4.8. Umkehrbarkeit und strenge Monotonie

Sei $I \subset \mathbb{R}$ ein Intervall und $f : I \rightarrow \mathbb{R}$ eine streng monoton wachsende Funktion. Dann ist f injektiv. Ist $J := f(I)$ die Bildmenge, so ist $f : I \rightarrow J$ sogar bijektiv und $f^{-1} : J \rightarrow I$ ebenfalls streng monoton wachsend.

Weiter ist $G_{f^{-1}} = \{(y, x) \in J \times I \mid (x, y) \in G_f\}$. Das ist der an der Winkelhalbierenden gespiegelte Graph von f .



BEWEIS: Seien $x_1, x_2 \in I$, $x_1 \neq x_2$. Dann ist eine der beiden Zahlen die kleinere, etwa $x_1 < x_2$. Aber dann ist $f(x_1) < f(x_2)$, wegen der strengen Monotonie, insbesondere also $f(x_1) \neq f(x_2)$. Das bedeutet, dass f injektiv ist.

Die Abbildung $f : I \rightarrow J := f(I)$ bleibt injektiv, aber sie ist zusätzlich surjektiv, und damit bijektiv. Sind etwa $y_1, y_2 \in J$, $y_1 < y_2$, so gibt es Elemente $x_1, x_2 \in I$ mit $f(x_i) = y_i$ für $i = 1, 2$. Wegen der Monotonie von f muss auch $x_1 < x_2$ gelten (wäre $x_1 \geq x_2$, also $x_2 \leq x_1$, so wäre auch $f(x_2) \leq f(x_1)$). Also ist f^{-1} streng monoton wachsend.

Schließlich ist $(y, x) \in G_{f^{-1}}$ genau dann, wenn $x = f^{-1}(y)$ ist, und das ist genau dann der Fall, wenn $f(x) = y$ ist, also $(x, y) \in G_f$. ■

Bemerkung: Ein analoger Satz gilt für streng monoton fallende Funktionen!

1.5 Komplexe Zahlen und Polynome

Definition

Die Menge $\mathbb{C}$ der **komplexen Zahlen** ist die Menge $\mathbb{R}^2$ der Paare (α, β) von reellen Zahlen, versehen mit der gewöhnlichen Vektoraddition

$$(\alpha_1, \beta_1) + (\alpha_2, \beta_2) := (\alpha_1 + \alpha_2, \beta_1 + \beta_2)$$

und der Multiplikation

$$(\alpha_1, \beta_1) \cdot (\alpha_2, \beta_2) := (\alpha_1\alpha_2 - \beta_1\beta_2, \alpha_1\beta_2 + \beta_1\alpha_2).$$

Statt $(1, 0)$ schreiben wir auch 1 , statt $(0, 1)$ schreiben wir i und sprechen auch von der **imaginären Einheit**. Statt $(0, 0)$ schreibt man einfach 0 .

Ist $z = (\alpha, \beta) = \alpha \cdot (1, 0) + \beta \cdot (0, 1) = \alpha + i\beta$ eine komplexe Zahl, so nennt man α den **Realteil** und β den **Imaginärteil** von z .

Offensichtlich ist $i^2 = (0, 1) \cdot (0, 1) = (-1, 0) = -1$.

5.1. Satz

Für die komplexen Zahlen gelten die folgenden Rechenregeln:

$$\begin{aligned} u + (v + w) &= (u + v) + w && \text{(Assoziativität der Addition),} \\ u + v &= v + u && \text{(Kommutativität der Addition),} \\ u + 0 &= u, \\ u + (-u) &= 0 && \text{(mit } -(\alpha, \beta) := (-\alpha, -\beta)), \\ u \cdot (v \cdot w) &= (u \cdot v) \cdot w && \text{(Assoziativität der Multiplikation),} \\ u \cdot v &= v \cdot u && \text{(Kommutativität der Multiplikation),} \\ u \cdot 1 &= u, \\ u \cdot (v + w) &= u \cdot v + u \cdot w && \text{(Distributivgesetz).} \end{aligned}$$

Die ersten 4 Regeln folgen direkt aus der Vektorrechnung, die anderen muss man nachrechnen. Ist z.B. $z_1 = \alpha_1 + i\beta_1$, $z_2 = \alpha_2 + i\beta_2$ und $z_3 = \alpha_3 + i\beta_3$, so ist

$$\begin{aligned} z_1 \cdot (z_2 + z_3) &= (\alpha_1 + i\beta_1) \cdot ((\alpha_2 + \alpha_3) + i(\beta_2 + \beta_3)) \\ &= (\alpha_1(\alpha_2 + \alpha_3) - \beta_1(\beta_2 + \beta_3)) + i(\beta_1(\alpha_2 + \alpha_3) + \alpha_1(\beta_2 + \beta_3)) \\ &= ((\alpha_1\alpha_2 - \beta_1\beta_2) + i(\alpha_1\beta_2 + \alpha_2\beta_1)) \\ &\quad + ((\alpha_1\alpha_3 - \beta_1\beta_3) + i(\alpha_1\beta_3 + \alpha_3\beta_1)) \\ &= (\alpha_1 + i\beta_1) \cdot (\alpha_2 + i\beta_2) + (\alpha_1 + i\beta_1)(\alpha_3 + i\beta_3) \\ &= z_1z_2 + z_1z_3. \end{aligned}$$

Jede reelle Zahl x kann in der Form $x = (x, 0)$ als komplexe Zahl aufgefasst werden. In diesem Sinne ist $\mathbb{R} \subset \mathbb{C}$.

Was noch fehlt, ist die Möglichkeit zu dividieren. Dazu müssen wir für festes z ein w mit $z \cdot w = 1$ finden. Das wird leichter, wenn wir konjugiert-komplexe Zahlen benutzen. Ist $z = \alpha + i\beta$ eine komplexe Zahl, so nennt man

$$\bar{z} := \alpha - i\beta$$

die dazu **konjugiert-komplexe** Zahl. Offensichtlich ist dann

$$z\bar{z} = \alpha^2 + \beta^2,$$

und das ist eine reelle Zahl ≥ 0 . Sie ist genau dann $= 0$, wenn $z = 0$ ist.

Ist $z \neq 0$, so ist $z \cdot \frac{\bar{z}}{z\bar{z}} = 1$, also $z^{-1} = \frac{1}{z\bar{z}} \bar{z}$. Das bedeutet:

$$(\alpha + i\beta)^{-1} = \frac{\alpha}{\alpha^2 + \beta^2} + i \frac{-\beta}{\alpha^2 + \beta^2}.$$

Man kann also mit komplexen Zahlen genauso wie mit reellen Zahlen rechnen. Hier sind ein paar Rechenbeispiele:

5.2. Beispiele

$$\text{A. } \frac{3+2i}{1-i} = \frac{(3+2i)(1+i)}{(1-i)(1+i)} = \frac{(3-2) + i(3+2)}{1^2 - i^2} = \frac{1+5i}{2} = \frac{1}{2} + i\frac{5}{2}.$$

Alle Formeln, die nur auf den elementaren Rechenregeln beruhen, gelten in $\mathbb{C}$ genauso wie in $\mathbb{R}$, das betrifft z.B. auch die binomischen Formeln.

$$\text{B. Es ist } i^{19} = i^{16} \cdot i^3 = (i^2)^8 \cdot i^2 \cdot i = -i.$$

C. Es ist

$$(2+3i)^3 = 2^3 + 3 \cdot 2^2 \cdot 3i + 3 \cdot 2 \cdot 3^2 i^2 + 3^3 i^3 = 8 + 36i - 54 - 27i = -46 + 9i.$$

Die komplexen Zahlen können nicht angeordnet werden, denn dann müsste $z^2 > 0$ für alle $z \in \mathbb{C}$ gelten. Wäre $i > 0$, so wäre auch $-1 = i^2 > 0$. Da auch $1 = (-1)(-1) > 0$ ist, ergibt das einen Widerspruch. Genauso folgt, dass $i < 0$ nicht gelten kann.

Immerhin können wir jeder komplexen Zahl einen Betrag zuordnen. Ist $z = \alpha + i\beta$, so heißt

$$|z| := \sqrt{z\bar{z}} = \sqrt{\alpha^2 + \beta^2}$$

der **Betrag** von z . Es gilt:

1. $|z| \geq 0$ und genau dann $= 0$, wenn $z = 0$ ist.

$$2. |z + w| \leq |z| + |w|.$$

$$3. |z \cdot w| = |z| \cdot |w|.$$

BEWEIS: 1) ist klar.

3) folgt aus der Beziehung

$$|zw|^2 = (zw)(\overline{zw}) = (z\overline{z})(w\overline{w}) = |z|^2 \cdot |w|^2.$$

2) Für jede komplexe Zahl $z = x + iy$ ist $\operatorname{Re}(z) = x \leq \sqrt{x^2 + y^2} = |z|$. Damit folgt:

$$|z + w|^2 = (z + w)(\overline{z + w}) = |z|^2 + |w|^2 + 2\operatorname{Re}(z\overline{w}) \leq |z|^2 + |w|^2 + 2|z| \cdot |w| = (|z| + |w|)^2.$$

■

5.3. Beispiel

Sei $z = 2 + i$ und $w = 3 - 2i$. Dann ist

$$|3z - 4w| = |(6 + 3i) - (12 - 8i)| = |-6 + 11i| = \sqrt{36 + 121} = \sqrt{157}.$$

Definition

Sei $z_0 \in \mathbb{C}$. Unter der ε -Umgebung von z_0 versteht man die Menge

$$U_\varepsilon(z_0) := \{z \in \mathbb{C} : |z - z_0| < \varepsilon\}.$$

Die ε -Umgebung $U_\varepsilon(z_0)$ ist eine (offene) Kreisscheibe mit Radius ε um z_0 , zu der der Kreisrand nicht gehört. Man bezeichnet eine solche Menge auch mit dem Symbol $D_\varepsilon(z_0)$ (für „Disk“).

Damit ist es möglich, den Konvergenzbegriff auch in $\mathbb{C}$ einzuführen.

Definition

Eine Folge (z_n) in $\mathbb{C}$ **konvergiert** gegen die komplexe Zahl z_0 , falls gilt:

$$\forall \varepsilon > 0 \exists n_0 \in \mathbb{N}, \text{ so dass } \forall n \geq n_0 \text{ gilt: } |z_n - z_0| < \varepsilon.$$

Offensichtlich gilt: $z_n = x_n + iy_n$ konvergiert genau dann gegen $z_0 = x_0 + iy_0$, wenn (x_n) gegen x_0 und (y_n) gegen y_0 konvergiert (denn es ist $|z_n - z_0| = \sqrt{(x_n - x_0)^2 + (y_n - y_0)^2}$).

Eine Folge (z_n) in $\mathbb{C}$ heißt **beschränkt**, falls es ein $R > 0$ gibt, so dass alle z_n in $D_R(0)$ liegen.

5.4. Satz von Bolzano-Weierstraß

Jede beschränkte Folge (z_ν) in $\mathbb{C}$ besitzt eine konvergente Teilfolge.

BEWEIS: Es gibt ein $R > 0$, so dass alle $z_\nu = x_\nu + iy_\nu$ in $D_R(0)$ liegen. Aber dann liegen sie erst recht in $I \times I$, mit $I := [-R, R]$. Die Folge (x_ν) besitzt eine konvergente Teilfolge (x_{ν_i}) mit einem Grenzwert $x_0 \in I$, die Folge (y_{ν_i}) besitzt eine konvergente Teilfolge $(y_{\nu_{i(k)}})$ mit einem Grenzwert $y_0 \in I$.

Die Teilfolge $(x_{\nu_{i(k)}} + iy_{\nu_{i(k)}})$ von (z_ν) konvergiert gegen $z_0 = x_0 + iy_0$. ■

Definition

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **Polynom(funktion)**, falls es reelle Zahlen $a_0, a_1, \dots, a_n$ gibt, so dass für alle $x \in \mathbb{R}$ gilt:

$$f(x) = a_0 + a_1x + \dots + a_nx^n.$$

Ist $f(x) = a_0 + a_1x + \dots + a_nx^n$ ein Polynom, so nennen wir $(a_0, \dots, a_n)$ ein **Koeffizientensystem** der Länge n für f . Im Augenblick wissen wir noch nicht, ob das Koeffizientensystem eindeutig bestimmt ist!

Da man $a_0 = a_1 = \dots = a_n = 0$ wählen kann, ist die Nullfunktion ein Polynom. Natürlich reicht dafür $a_0 = 0$ als Koeffizientensystem aus. Ist f nicht das Nullpolynom, so besitzt f ein Koeffizientensystem $(a_0, \dots, a_n)$ minimaler Länge mit $a_n \neq 0$. Dann bezeichnet man n als den **Grad** von f . Offensichtlich ist $\text{grad}(f) = 0$ genau dann, wenn f eine konstante Funktion $\neq 0$ ist. Das Nullpolynom erhält definitionsgemäß den Grad $-\infty$.

5.5. Abspaltung von Linearfaktoren

Sei $f(x) = a_0 + a_1x + \dots + a_nx^n$ ein Polynom vom Grad $n \geq 1$ und x_0 eine Nullstelle von f . Dann gibt es ein Polynom $g(x) = b_0 + b_1x + \dots + b_{n-1}x^{n-1}$ mit $b_{n-1} = a_n$, so dass gilt:

$$f(x) = (x - x_0) \cdot g(x) \quad \text{für alle } x \in \mathbb{R}.$$

BEWEIS: Weil $f(x_0) = 0$ ist, gilt für beliebiges $x \in \mathbb{R}$:

$$\begin{aligned}
f(x) &= f(x) - f(x_0) = \sum_{i=0}^n a_i x^i - \sum_{i=0}^n a_i x_0^i \\
&= \sum_{i=1}^n a_i (x^i - x_0^i) \quad (\text{die Terme mit } i=0 \text{ heben sich weg!}) \\
&= (x - x_0) \cdot \sum_{i=1}^n a_i \cdot \sum_{j=0}^{i-1} x^j x_0^{i-j-1} \\
&= (x - x_0) \cdot (a_1 + a_2(x_0 + x) + \cdots + a_n(x_0^{n-1} + x_0^{n-2}x + \cdots + x^{n-1})) \\
&= (x - x_0) \cdot g(x),
\end{aligned}$$

wobei $g(x) = b_0 + b_1x + \cdots + b_{n-1}x^{n-1}$ ein Polynom mit $b_{n-1} = a_n$ ist. ■

5.6. Folgerung (über die Anzahl der Nullstellen)

Ist $f(x) = a_0 + a_1x + \cdots + a_nx^n$ für alle $x \in \mathbb{R}$ und $a_n \neq 0$, so hat f höchstens n verschiedene Nullstellen.

BEWEIS: Wir führen Induktion nach n .

Im Falle $n = 0$ besitzt $f(x) = a_0$ (mit $a_0 \neq 0$) überhaupt keine Nullstelle. Sei nun $n \geq 1$, und die Behauptung für $n - 1$ schon bewiesen. Hat f keine Nullstelle, so ist nichts zu zeigen. Es sei also $f(x_0) = 0$. Dann gibt es ein Polynom

$$g(x) = b_0 + b_1x + \cdots + b_{n-1}x^{n-1}$$

mit $b_{n-1} = a_n$, so dass $f(x) = (x - x_0) \cdot g(x)$ für alle $x \in \mathbb{R}$ ist.

Ist $x_1 \neq x_0$ eine Nullstelle von f , so ist

$$0 = f(x_1) = (x_1 - x_0) \cdot g(x_1),$$

also x_1 auch Nullstelle von g . Nach Induktionsvoraussetzung besitzt g aber höchstens $n - 1$ Nullstellen. Also hat f höchstens n Nullstellen. ■

5.7. Folgerung (Identitätssatz für Polynome)

Besitzt ein Polynom $f(x) = a_0 + a_1x + \cdots + a_nx^n$ mehr als n Nullstellen, so muss $a_0 = a_1 = \cdots = a_n = 0$ sein.

Der BEWEIS ist trivial.

Hieraus ergibt sich unmittelbar, dass die Koeffizienten eines Polynoms eindeutig bestimmt sind. Sind nämlich $(a_0, \dots, a_n)$ und $(b_0, \dots, b_n)$ zwei verschiedene Koeffizientensysteme für ein Polynom f (wobei man die Gleichheit der Längen ggf. durch Hinzunahme von Nullen erzwingen kann), so gibt es ein k mit $0 \leq k \leq n$, so dass $a_k \neq b_k$ und $a_i = b_i$ für $i > k$ ist. Dann gilt:

$$(a_0 - b_0) + (a_1 - b_1)x + \cdots + (a_k - b_k)x^k = 0$$

für alle $x \in \mathbb{R}$. Das kann aber nach der letzten Folgerung nicht sein.

Die Summe und das Produkt zweier Polynome ist wieder ein Polynom. Sei etwa $f(x) := \sum_{i=0}^n a_i x^i$ und $g(x) := \sum_{j=0}^m b_j x^j$. Ist $n \leq m$, so ist

$$(f + g)(x) = \sum_{i=0}^n (a_i + b_i) x^i + \sum_{i=n+1}^m b_i x^i.$$

Ist $n > m$, so verfährt man analog.

Das Produkt von f und g ist gegeben durch

$$\begin{aligned} (f \cdot g)(x) &:= \sum_{k=0}^{n+m} \left(\sum_{i+j=k} a_i b_j \right) x^k \\ &= a_0 b_0 + (a_1 b_0 + a_0 b_1) x + (a_2 b_0 + a_1 b_1 + a_0 b_2) x^2 + \cdots + a_n b_m x^{n+m}. \end{aligned}$$

Hieraus ergeben sich die Gradformeln:

1. $\text{grad}(f \cdot g) = \text{grad}(f) + \text{grad}(g).$
2. $\text{grad}(f + g) \leq \max(\text{grad}(f), \text{grad}(g))$

Dabei kann sich der Grad dadurch verringern, dass sich die höchsten Terme von f und g beim Addieren wegheben.

5.8. Division mit Rest für Polynome

Sei $g \neq 0$ ein Polynom. Dann gibt es zu jedem Polynom f eindeutig bestimmte Polynome q und r , so dass gilt:

1. $f = q \cdot g + r.$
2. $r = 0$ oder $\text{grad}(r) < \text{grad}(g).$

Auf den BEWEIS verzichten wir hier, der gehört in die Lineare Algebra.

5.9. Beispiel

Sei $f(x) = 2x^5 - 13x^4 + 17x^3 - x^2 + 10x + 8$ und $g(x) = 2x^2 - 3x$.

Will man $f(x)$ durch $g(x)$ mit Rest dividieren, so geht man folgendermaßen vor: Die Division des höchsten Termes von $f(x)$ (also $2x^5$) durch den höchsten Term von $g(x)$ (also $2x^2$) ergibt den höchsten Term $q_1(x)$ von $q(x)$ (also x^3). Dann subtrahiert man $g(x) \cdot q_1(x)$ von $f(x)$ und erhält ein Polynom $f_1(x)$. Mit dem beginnt man

die Prozedur erneut, so lange, bis nur noch ein Polynom vom Grad $< g(x)$ übrig bleibt. Das ist der gesuchte Rest $r(x)$.

$$\begin{array}{r}
 (2x^5 - 13x^4 + 17x^3 - x^2 + 10x + 8) : (2x^2 - 3x) = x^3 - 5x^2 + x + 1 \\
 \underline{2x^5 - 3x^4} \\
 -10x^4 + 17x^3 - x^2 + 10x + 8 \\
 \underline{-10x^4 + 15x^3} \\
 2x^3 - x^2 + 10x + 8 \\
 \underline{2x^3 - 3x^2} \\
 2x^2 + 10x + 8 \\
 \underline{2x^2 - 3x} \\
 13x + 8
 \end{array}$$

Hier ist $q(x) = x^3 - 5x^2 + x + 1$ und $r(x) = 13x + 8$.

Besitzt ein Polynom $f(x)$ eine Nullstelle x_0 , so benutzt man den Divisionsalgorithmus, um den Linearfaktor $x - x_0$ abzuspalten. Sei z.B. $p(x) := x^5 - 2x^4 - x + 2$. Durch Probieren stellt man schnell fest, dass p bei $x = 1$ eine Nullstelle besitzt. Durch Polynomdivision erhält man:

$$(x^5 - 2x^4 - x + 2) : (x - 1) = x^4 - x^3 - x^2 - x - 2.$$

Im allgemeinen ist es nicht möglich, Nullstellen von Polynomen höheren Grades ohne numerische Methoden zu bestimmen. In Spezialfällen gibt es aber ein paar nette kleine Hilfsmittel. Darüber soll im Ergänzungsteil gesprochen werden.

Definition

Sei f ein Polynom. Gibt es ein $x_0 \in \mathbb{R}$, ein $k \in \mathbb{N}$ und ein Polynom g mit $g(x_0) \neq 0$, so dass $f(x) = (x - x_0)^k \cdot g(x)$ ist, so nennt man x_0 eine Nullstelle der **Ordnung** (oder **Vielfachheit**) k .

5.10. Satz (über die Zerlegung in Linearfaktoren)

Sei $p(x)$ ein Polynom vom Grad $n \geq 0$.

1. p besitzt höchstens n Nullstellen (mit Vielfachheit gezählt).
2. Sind $c_1, \dots, c_m$ die verschiedenen Nullstellen von p , $m \leq n$, mit Vielfachheiten $k_1, \dots, k_m$, so ist

$$p(x) = (x - c_1)^{k_1} \cdot \dots \cdot (x - c_m)^{k_m} \cdot q(x),$$

mit $k_1 + \dots + k_m \leq n$ und einem Polynom q ohne Nullstellen.

Zum BEWEIS dividiert man sukzessive alle Nullstellen heraus, bis nur noch ein Polynom ohne Nullstellen übrigbleibt. Dessen Grad ist $= n - (k_1 + \dots + k_m) \geq 0$.

5.11. Beispiel

Sei $p(x) = x^7 - 2x^6 + 3x^5 - 4x^4 + 2x^3$.

Offensichtlich ist $p(x) = x^3 \cdot q_1(x)$, mit $q_1(x) := x^4 - 2x^3 + 3x^2 - 4x + 2$.

Da $q_1(0) \neq 0$ ist, hat p in $x = 0$ eine Nullstelle der Vielfachheit 3.

Probieren zeigt, dass $q_1(1) = 0$ ist. Polynomdivision ergibt:

$$q_1(x) = (x - 1) \cdot q_2(x), \text{ mit } q_2(x) := x^3 - x^2 + 2x - 2.$$

Da auch $q_2(1) = 0$ ist, kann man noch einmal durch $(x - 1)$ dividieren und erhält:

$$q_1(x) = (x - 1)^2 \cdot q_3(x), \text{ mit } q_3(x) := x^2 + 2.$$

$q_3(x)$ besitzt keine Nullstelle mehr. Also ist

$$p(x) = x^3 \cdot (x - 1)^2 \cdot (x^2 + 2)$$

die bestmögliche Zerlegung von $p(x)$.

Wir verlassen nun den Bereich der reellen Zahlen und lassen auch komplexe Nullstellen zu. Wir gehen sogar einen Schritt weiter und lassen gelegentlich auch Polynome mit komplexen Koeffizienten zu:

$$p(z) = c_0 + c_1 z + c_2 z^2 + \dots + c_n z^n,$$

mit $c_0, c_1, \dots, c_n \in \mathbb{C}$ und $c_n \neq 0$.

Wie im Reellen werden Grad und Nullstellen definiert, der Divisionsalgorithmus überträgt sich wörtlich auf komplexe Polynome, und wie im Reellen besteht auch im Komplexen der Zusammenhang zwischen Nullstellen und Linearfaktoren.

Besonders interessieren uns aber die komplexen Nullstellen reeller Polynome.

5.12. Komplexe Nullstellen eines reellen Polynoms

Sei $p(x)$ ein Polynom vom Grad n mit **reellen** Koeffizienten.

1. Ist $\alpha \in \mathbb{C}$ eine Nullstelle von p , so ist auch $\bar{\alpha}$ eine Nullstelle.
2. Die Anzahl der nicht-reellen Nullstellen von p ist gerade.
3. Ist $\text{grad}(p) = 2$, so besitzt p genau zwei Nullstellen (mit Vielfachheit gezählt).

BEWEIS: 1) Sei $p(x) = a_0 + a_1x + \cdots + a_nx^n$. Ist $\alpha \in \mathbb{C}$ eine Nullstelle von p , so ist

$$p(\alpha) = 0, \text{ und damit } 0 = \overline{p(\alpha)} = p(\bar{\alpha}).$$

2) folgt trivial aus (1).

3) Sei $p(x) = ax^2 + bx + c$, mit $a \neq 0$, und $\Delta_p = b^2 - 4ac$ die Diskriminante. Ist $\Delta_p \geq 0$, so ist nichts mehr zu zeigen. Ist $\Delta_p < 0$, so erhält man die beiden Lösungen

$$x_1 = \frac{-b + i\sqrt{-\Delta_p}}{2a} \quad \text{und} \quad x_2 = \frac{-b - i\sqrt{-\Delta_p}}{2a}.$$

■

Von entscheidender Bedeutung ist nun

5.13. Der Fundamentalsatz der Algebra

Jedes nicht konstante komplexe Polynom hat in $\mathbb{C}$ wenigstens eine Nullstelle.

Den BEWEIS können wir noch nicht führen. Mit Hilfe der Polynomdivision folgt aus dem Fundamentalsatz, dass ein Polynom vom Grad n in n Linearfaktoren zerlegt werden kann (wobei natürlich zugelassen ist, dass ein Faktor mehrfach auftritt).

5.14. Folgerung

*Jedes nicht konstante reelle Polynom **ungeraden Grades** besitzt wenigstens eine reelle Nullstelle.*

BEWEIS: Da das Polynom in Linearfaktoren zerfällt und die nicht reellen Nullstellen paarweise auftreten, bleibt mindestens eine reelle Nullstelle übrig. ■

Ist c eine komplexe Nullstelle eines reellen Polynoms $p(x)$ (mit $\text{Im}(c) \neq 0$), so ist auch $\bar{c}$ eine Nullstelle und $(x - c)(x - \bar{c}) = x^2 - (2\text{Re}(c))x + |c|^2$ ein quadratisches Polynom mit reellen Koeffizienten. Daraus folgt:

5.15. Folgerung

Jedes nicht konstante reelle Polynom zerfällt vollständig in (reelle) Linearfaktoren und quadratische Polynome ohne reelle Nullstelle.

Definition

Sind p, q zwei Polynome, von denen q nicht das Nullpolynom ist, so nennt man $f(x) := p(x)/q(x)$ eine **rationale Funktion**.

Der Definitionsbereich einer solchen rationalen Funktion ist die Menge $D_f := \{x \in \mathbb{R} : q(x) \neq 0\}$.

Da das Nennerpolynom q einen Grad ≥ 0 haben soll, besitzt es höchstens endlich viele Nullstellen. Eine rationale Funktion ist also fast überall definiert.

Was passiert in einem Punkt x_0 mit $q(x_0) = 0$?

Wir können den Faktor $(x - x_0)$ in der höchstmöglichen Potenz aus $p(x)$ und $q(x)$ herausziehen: Es gibt Zahlen $k \geq 0$ und $l > 0$, sowie Polynome $\tilde{p}$ und $\tilde{q}$, so dass gilt:

$$p(x) = (x - x_0)^k \cdot \tilde{p}(x) \text{ und } q(x) = (x - x_0)^l \cdot \tilde{q}(x),$$

mit $\tilde{p}(x_0) \neq 0$ und $\tilde{q}(x_0) \neq 0$. Dabei ist auch $k = 0$ möglich.

Nun sind zwei Fälle zu unterscheiden:

$$1. \text{ Ist } k \geq l, \text{ so ist } f(x) = (x - x_0)^{k-l} \cdot \frac{\tilde{p}(x)}{\tilde{q}(x)},$$

und die rechte Seite der Gleichung ist auch in $x = x_0$ definiert. Man nennt x_0 dann eine **hebbare Unbestimmtheitsstelle** von f . Der „unbestimmte Wert“ $f(x_0) = 0/0$ kann in diesem Fall durch eine bestimmte reelle Zahl ersetzt werden.

$$2. \text{ Ist } k < l, \text{ so ist } f(x) = \frac{1}{(x - x_0)^{l-k}} \cdot \frac{\tilde{p}(x)}{\tilde{q}(x)}.$$

Hier ist der erste Faktor bei x_0 nach wie vor nicht definiert, während der zweite Faktor einen bestimmten Wert annimmt. Man sagt in dieser Situation: f besitzt in x_0 eine **Polstelle** der Ordnung $l - k$.

Wir wollen nun eine beliebige rationale Funktion $f(x) = p(x)/q(x)$ in eine ganz bestimmte „Normalform“ bringen.

a) Ist $\text{grad}(p) \geq \text{grad}(q)$, so liefert der Satz von der Division mit Rest eine Zerlegung $p(x) = g(x) \cdot q(x) + r(x)$ mit $\text{grad}(r) < \text{grad}(q)$ und damit eine Darstellung

$$f(x) = g(x) + r(x)/q(x), \text{ mit einem Polynom } g.$$

b) Ist $\text{grad}(p) < \text{grad}(q)$ und zerfällt $q(x)$ vollständig in Linearfaktoren, so gibt es eine sogenannte *Partialbruchzerlegung*:

5.16. Partialbruchzerlegung

$p(x), q(x)$ seien zwei Polynome mit $\text{grad}(p) < \text{grad}(q)$. Außerdem sei

$$q(x) = (x - c_1)^{k_1} \cdots (x - c_r)^{k_r}$$

die Zerlegung von $q(x)$ in Linearfaktoren (über $\mathbb{C}$), mit paarweise verschiedenen komplexen Zahlen c_i . Dann gibt es eine eindeutig bestimmte Darstellung

$$\frac{p(x)}{q(x)} = \sum_{j=1}^r \sum_{k=1}^{k_j} \frac{a_{jk}}{(x - c_j)^k}, \text{ mit } a_{jk} \in \mathbb{C}.$$

Auch hier überlassen wir den Beweis der Linearen Algebra.

In einfachen Fällen kann man folgendermaßen vorgehen:

1. Zerlegung des Nenners in Linearfaktoren. Meistens scheitert man schon an dieser Stelle.
2. Ansatz mit „unbestimmten Koeffizienten“, so wie in der Formel vorgegeben.
3. Multiplikation beider Seiten mit dem Nenner der linken Seite. Das führt zu einer Polynomgleichung.
4. Vergleich der Koeffizienten bei den Potenzen von x . Wegen der Eindeutigkeit des Koeffizientensystems eines Polynoms liefert der Vergleich ein lineares Gleichungssystem, aus dem man die unbestimmten Koeffizienten gewinnen kann.

5.17. Beispiel

Sei $f(x) := \frac{x^2 + 5x + 2}{(x-1)(x+1)^2}$.

Ansatz: $\frac{x^2 + 5x + 2}{(x-1)(x+1)^2} = \frac{a_{11}}{x-1} + \frac{a_{21}}{x+1} + \frac{a_{22}}{(x+1)^2}$.

Multiplikation mit dem Nenner der linken Seite führt zu

$$\begin{aligned} x^2 + 5x + 2 &= a_{11}(x+1)^2 + a_{21}(x^2-1) + a_{22}(x-1) \\ &= (a_{11} + a_{21})x^2 + (2a_{11} + a_{22})x + (a_{11} - a_{21} - a_{22}). \end{aligned}$$

Koeffizientenvergleich liefert die Bestimmungsgleichungen

$$a_{11} + a_{21} = 1, \quad 2a_{11} + a_{22} = 5 \quad \text{und} \quad a_{11} - a_{21} - a_{22} = 2.$$

Als Lösung erhält man: $a_{11} = 2$, $a_{21} = -1$ und $a_{22} = 1$.

Man kann an dem Beispiel auch noch eine andere Methode demonstrieren. Multipliziert man am Anfang beide Seiten mit $(x+1)^2$ und setzt dann $x = -1$ ein, so erhält man sofort

$$a_{22} = \frac{(-1)^2 + 5(-1) + 2}{-2} = 1.$$

Dann ist

$$\frac{a_{11}}{x-1} + \frac{a_{21}}{x+1} = \frac{x^2 + 5x + 2}{(x-1)(x+1)^2} - \frac{1}{(x+1)^2} = \frac{x+3}{(x-1)(x+1)}$$

Hier kann man nun mit einer der beiden gezeigten Methoden weitermachen.

Sind $p(x)$ und $q(x)$ Polynome mit reellen Koeffizienten, so ist man oft an einer „reellen“ Partialbruchzerlegung interessiert. Sind alle Nullstellen von $q(x)$ reell, so erhält man reelle a_{jk} . Es bleibt das Problem der nicht-reellen Nullstellen.

Ist etwa c_j eine Nullstelle von $q(x)$, mit $\operatorname{Im}(c_j) \neq 0$, so ist auch $\overline{c_j}$ eine Nullstelle von $p(x)$. Es muss also ein $i \neq j$ mit $c_i = \overline{c_j}$ geben, und es ist dann $k_i = k_j$.

Weil $\overline{(p/q)(x)} = p(\overline{x})/q(\overline{x})$ ist, muss dann – wegen der Eindeutigkeit der Partialbruchzerlegung – gelten:

$$\frac{\overline{a_{jk}}}{(\overline{x} - c_i)^k} = \frac{a_{ik}}{(\overline{x} - c_i)^k}.$$

Daraus folgt: $a_{ik} = \overline{a_{jk}}$, für $k = 1, \dots, k_j$.

Nun gilt aber für beliebige komplexe Zahlen a und c :

$$\frac{a}{x - c} + \frac{\overline{a}}{x - \overline{c}} = \frac{\delta x + \varepsilon}{x^2 + \beta x + \gamma},$$

mit reellen Zahlen $\beta, \gamma, \delta, \varepsilon$. Behält man im Nenner quadratische Polynome, so muss man im Zähler affin-lineare Terme zulassen. Die Methode des Ansatzes und Koeffizientenvergleichs funktioniert dann genauso.