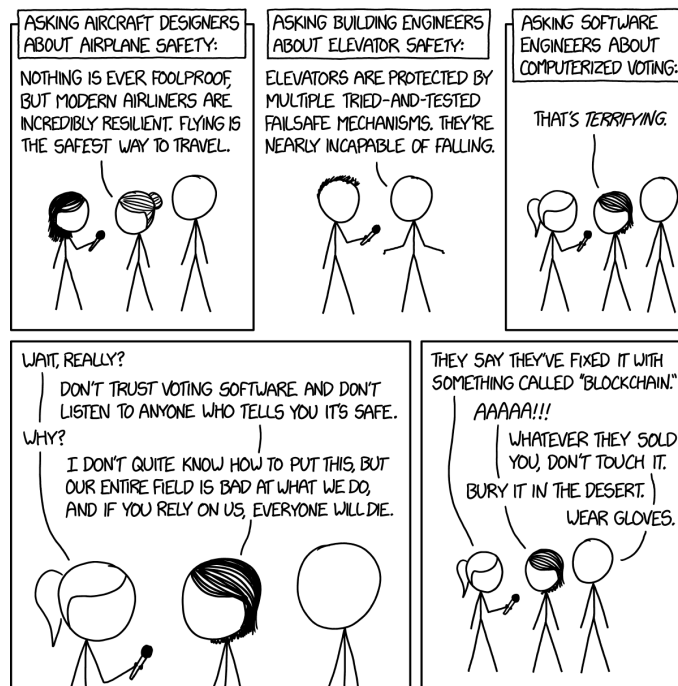


Vorkurs Informatik der Technischen Fakultät Universität Bielefeld

Der Mathematikteil

Skript September 2026

Dr Dirk Frettlöh
Technische Fakultät
Universität Bielefeld



Quelle: xkcd.com

Inhaltsverzeichnis

1	Das Handwerk: Notation und Rechentricks	5
1.1	Mengen und Zahlbereiche	5
1.2	Stenographie mit Sigma und Pi	7
1.3	Definition, Lemma, Satz, Beweis	9
1.4	Indexverschiebung und Summanden einzeln schreiben	10
1.5	Potenzen	12
1.6	Rechentricks	14
1.7	Wurzeln	16
1.8	Logarithmen	18
2	Die Kunst: Beweisen — vollständige Induktion	21
2.1	Vollständige Induktion	22
3	Formale Logik	26
3.1	Aussagenlogik	26
3.2	Quantoren	30
4	Folgen und Reihen	32
4.1	Das kleine Epsilon	32
4.2	Unendliche Reihen	37
4.3	Bonusmaterial: Potenzreihen	39
5	Abbildungen, aka Funktionen	42
5.1	Wichtige Vokabeln	42
5.2	Injektiv, surjektiv, bijektiv	44
5.3	Umkehrfunktionen	45
5.4	Prominente Funktionen: Polynome, Sinus, Kosinus und Kollegen	48
5.4.1	Polynome	48
5.4.2	Sinus, Kosinus und Kollegen	49
6	Zahlbereiche: $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$	56
6.1	Von Gruppen und Körpern — $\mathbb{Q}$	56
6.2	Das Füllen der Lücken — $\mathbb{R}$	58
6.3	Die fehlenden Wurzeln — $\mathbb{C}$	60

7 Überblick Mathe I Lineare Algebra (und diskrete Mathematik)	63
Lineare Gleichungssysteme	64
I Komplexe Zahlen	69
II Vektorräume	69
III Lineare Abbildungen und Matrizen	72
IV Matrizen und lineare Gleichungssysteme	73
8 Überblick Mathe I Analysis	75
I Folgen	75
II Stetige Funktionen	75
III Differentialrechnung	78
9 Bonusmaterial: Binomialkoeffizienten	82
9.1 Die Formel von Signore Binomi	82
9.2 Der Binomialkoeffizient	83
9.3 Der binomische Lehrsatz	83
10 Bonusmaterial: Polarkoordinaten	86
11 Bonusmaterial: Unendlichkeit	87
11.1 Endliche und unendliche Mengen	88
11.2 Überabzählbarkeit – noch mehr als $\mathbb{N}$	91

If people do not believe that mathematics is simple, it is only because they do not realize how complicated life is.

Mr Bean

Vorab:

Der Vorkurs Informatik ist ein Angebot der Technischen Fakultät (kurz TechFak) der Universität Bielefeld. Es richtet sich an Studierende eines Informatikstudiengangs der TechFak (also Informatik, Naturwissenschaftliche Informatik (NWI), Bioinformatik und Genomforschung (BIG) und Kognitive Informatik (KOI)) **vor dem ersten Semester**. Die Teilnahme ist freiwillig.

Gerne können auch Studienanfänger der Physik-Studiengänge, die “Mathe 1 für Informatik und Naturwissenschaften” hören, teilnehmen.

Da ein Informatikstudium gute mathematische Kenntnisse verlangt sowie Programmierkenntnisse, empfiehlt sich eine Teilnahme, um diese Kenntnisse aufzufrischen oder zu verbessern. Der Vorkurs hat daher einen Matheteil (erste und dritte Woche) und einen Informatikteil (zweite Woche).

An den Leser:

Mathematik hat mehr mit Können zu tun als mit Wissen. Ohne Aufgaben zu bearbeiten lernt man nichts. Es gibt in diesem Skript **Aufgaben** (leicht bis mittelschwer, werden live in den Tutorien bearbeitet) und Aufgaben mit *: **Aufgaben*** (mittelschwer bis schwer, werden in der Mittagspause/zu Hause/im Zug... selbständig bearbeitet und bei Bedarf am nächsten Tag in den Tutorien besprochen), die der Nachbereitung und Festigung des Behandelten dienen.

Die ersten 8 Kapitel dieses Skripts decken die Themen der zwei Wochen des Matheteils des Vorkurses ab. **Obacht: 2026** fielen drei Tage aus, es werden wohl ein bis zwei Kapitel weggelassen oder nur kurz angerissen. Kapitel 9, 10 und 11 enthalten Bonusmaterial. Die den Kapiteln vorangestellten Zitate sind alle echt, nur die angegebene Zuweisung ist gelogen.

Dank:

Dieses Skript beruht in Teilen auf dem von Lars Scheele vom Mathematikteil des Vorkurs Informatik 2007 der TechFak.

1 Das Handwerk: Notation und Rechentricks

Mathematik ist eine basisdemokratische Wissenschaft. Es gibt keine Geheimnisse, jeder kann eine logische Argumentation nachvollziehen.

Kim Jong Un

Mathematik ist vieles: Werkzeug, abstrakte Wissenschaft, Knobelei, Suche nach Mustern... nicht zuletzt ist Mathematik eine Sprache. In dieser Sprache lassen sich viele natürliche Phänomene formulieren. Grandioses Beispiel ist die Physik; die Entwicklung der Physik und der Mathematik gingen Hand in Hand. Heute wird Mathematik auch zunehmend wichtiger in vielen anderen Disziplinen, etwa Biologie oder Soziologie. Besonders wichtig ist die Mathematik in der Informatik. Es ist nicht übertrieben zu sagen, dass die Mathematik die Mutter der Informatik ist. Der Vater — oder die andere Mutter — ist die Elektrotechnik, ein Fach, in dem Physik eine große Rolle spielt und somit auch Mathematik. Viele Gesetze der Physik lassen sich am einfachsten in der Sprache der Mathematik ausdrücken. Wir fangen an mit einigen Elementen dieser Sprache.

1.1 Mengen und Zahlbereiche

Die natürlichen Zahlen hat der liebe Gott gemacht, alles andere ist Menschenwerk.

Papst Franziskus

Eine Menge ist ein grundlegendes Konzept in der Mathematik. Darunter versteht man einfach eine Zusammenfassung von Dingen (den Elementen) zu einem Ganzen. Allgemein werden in der Mathematik zur Beschreibung von Mengen geschweifte Klammern benutzt: $\{$ und $\}$. (Bei Programmiersprachen ist das dagegen sehr verschieden. In Haskell, das ihr im ersten Semester kennenlernen werdet, benutzt man für Mengen etwa $[$ und $]$.) Eine Menge ist also z.B.

$$M := \{1, 2, 34, 42\} \text{ oder } A := \{\text{Hund, Katze, Maus}\}$$

Dabei heißt das $:=$ soviel wie “Wird definiert als”. Die Elemente der ersten Menge M sind also 1, 2, 34 und 42, die Elemente der zweiten Menge sind Hund, Katze und Maus. Das kann man kurz schreiben als $1 \in M$, $2 \in M$ usw bzw $\text{Hund} \in A$ usw. Will man sagen “ist nicht Element” so schreibt man $\notin$. Also ist etwa $5 \notin M$. Wichtig ist die **leere Menge**, die kein Element enthält. Die kann man so schreiben: $\{\}$, aber öfter wird die so: $\emptyset$ geschrieben.

Die Reihenfolge spielt in einer Menge übrigens keine Rolle, es ist also $\{1, 2, 3\} = \{3, 2, 1\} = \{2, 3, 1\}$ usw. Will man die Reihenfolge beachten, so benutzt man meistens runde Klammern, also z.B. $(1, 2, 3)$ Dann ist $(1, 2, 3) \neq (3, 2, 1)$.

Im Folgenden soll die Menge der **natürlichen Zahlen** $\{1, 2, 3, 4, \dots\}$ mit $\mathbb{N}$ bezeichnet werden. Dabei wird stillschweigend davon ausgegangen, dass $0 \notin \mathbb{N}$. Die Menge der natürlichen Zahlen mit 0 soll mit $\mathbb{N}_0$ bezeichnet werden. Die Menge aller **ganzen Zahlen**, also die Menge $\{\dots, -2, -1, 0, 1, 2, 3, \dots\}$ soll mit $\mathbb{Z}$ bezeichnet werden.

Um noch komplexere Mengen zu beschreiben gibt es die folgende Notation:

$$\{ \text{Platzhalter für Objekte (oft mit "Datentyp")} \mid \text{Bedingungen an Objekte} \}$$

Die Menge aller geraden natürlichen Zahlen wird damit etwa so geschrieben: $\{n \in \mathbb{N} \mid n \text{ gerade}\}$; die Menge aller ganzen Zahlen, die größer als -10 sind, so: $\{n \in \mathbb{Z} \mid n > -10\}$. Hier heißt $>$ soviel wie “ist größer als”, $<$ heißt “ist kleiner als”. Analog heißt $\geq$ soviel wie “ist größer oder gleich”, $\leq$ heißt “ist kleiner oder gleich”. Also gilt etwa $3 > 2$, $5 < 1000$ oder auch $7 \leq 7$ und $7 \geq 7$ (aber nicht $7 > 7$). Genau so könnte man die geraden natürlichen Zahlen schreiben als $\{2 \cdot n \mid n \in \mathbb{N}\}$. Die Menge aller geraden ganzen Zahlen außer 6 kann man so schreiben: $\{2 \cdot n \mid n \in \mathbb{Z}, n \neq 3\}$. Dabei heißt $\neq$ “ist ungleich”. Die Menge aller Brüche aus ganzen Zahlen kann man so schreiben: $\{\frac{k}{n} \mid k, n \in \mathbb{Z}, n \neq 0\}$. Diese Zahlen heißen **rationale Zahlen**, und die Menge der rationalen Zahlen bezeichnet man auch kurz mit $\mathbb{Q}$.

Im Folgenden wird in den Beispielen die Menge der reellen Zahlen (geschrieben $\mathbb{R}$) eine wichtige Rolle spielen. Diese Menge wird später vielleicht noch genauer eingeführt; für viele Zwecke reicht es, sich auf Intuition oder Schulwissen zu verlassen. (Es ist OK sich die reellen Zahlen als Punkte auf einer unendlichen Geraden vorzustellen, jeder Punkt entspricht einer Zahl. Oder aber alle Zahlen, die ich als Dezimalzahlen mit beliebig vielen — auch unendlich vielen! — Nachkommastellen schreiben kann). Bestimmte wichtige Teilmengen von $\mathbb{R}$ werden wie folgt notiert:

$$\mathbb{R}_0^+ := \{x \in \mathbb{R} \mid x \geq 0\}; \quad \mathbb{R} \setminus \{0\} := \{x \in \mathbb{R} \mid x \neq 0\}; \quad \mathbb{R}^+ := \{x \in \mathbb{R} \mid x > 0\}.$$

Hier erklären wir auch **Intervalle**: Ein Abschnitt der reellen Zahlen, jeweils mit oder ohne Randwert. Formal:

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

Das erste, also $[a, b]$, heißt **abgeschlossenes** Intervall, das zweite, $]a, b[$ heißt **offenes** Intervall. Statt $]a, b[$ schreibt man auch oft (a, b) . Analog gibt es die **halboffenen** Intervalle

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

So ist etwa das Intervall $[-1; 2[$ (oder auch $[-1, 2)$) die Menge aller reellen Zahlen zwischen -1 und 2, ohne die 2, aber inklusive der -1.

Für Mengen möchte man noch weitere Zeichen erklären. Der **Schnitt** zweier Mengen A, B ist

$$A \cap B := \{x \mid x \in A \text{ und } x \in B\}.$$

Die **Vereinigung** zweier Mengen ist

$$A \cup B := \{x \mid x \in A \text{ oder } x \in B\}.$$

Die **Differenz** (“A ohne B”) zweier Mengen ist

$$A \setminus B := \{x \mid x \in A \text{ und } x \notin B\}.$$

Für $A := \{0, 2, 4, 6, 8, 10\}$ und $B := \{1, 2, 3, 4, 5\}$ ist also

$$A \cap B := \{2, 4\}, \quad A \cup B = \{0, 1, 2, 3, 4, 5, 6, 8, 10\}, \quad A \setminus B := \{0, 6, 8, 10\}.$$

Ein anderes Beispiel mit Intervallen:

$$[0, 2] \cap]1, 3[=]1, 2], \quad [0, 2] \cup]1, 3[= [0, 3[, \quad [0, 2] \setminus]1, 3[= [0, 1].$$

Außerdem gibt es noch eine Notationen dafür, dass eine Menge in einer anderen enthalten ist: Falls alle Elemente aus der Menge A auch in der Menge B liegen, so schreibt man $A \subset B$ (“ A ist Teilmenge von B ”). Das kann man auch umdrehen: $A \supset B$ heißt, dass B Teilmenge von A ist. Die entsprechenden Symbole durchgestrichen heißt, dass das jeweils nicht der Fall ist. $A \not\subset B$ heißt also, dass A nicht Teilmenge von B ist. Beispiele dazu: $\{1, 2, 3, 4, 5\} \not\subset \{1, 2, 3\}$, oder $[0, 2] \not\subset]1, 3[$, und auch umgekehrt $[0, 2] \not\supset]1, 3[$.

Aufgabe 1.0 Was ergeben sich hier für Intervalle?

$[-1, 1] \cup [0, 2]$, $[-1, 1] \cap [0, 2]$, $[-1, 1] \cap]0, 2[$, $([-2, 0] \setminus [0, 2]) \cap [-3, 0]$, $[-2, 0] \setminus ([0, 2] \cap [-3, 0])$.

Welche der folgenden Aussagen stimmen immer, bei welchen gibt es Ausnahmen? Finden Sie konkrete Beispiele für die Ausnahmen.

$$A \setminus B = B \setminus A, \quad (A \cap B) \cup C = A \cap (B \cup C), \quad A \setminus (B \cup C) = (A \setminus B) \cap (A \setminus C)$$

$$C \setminus (A \cap B) = (C \setminus A) \cup (C \setminus B), \quad (A \setminus B) \cap C = A \cap (B \setminus C), \quad A \setminus (B \setminus C) = (A \cap C) \cup (A \setminus B)$$

1.2 Stenographie mit Sigma und Pi

(Aus einem IQ-Test) Fügen Sie zu den Zahlenreihen diejenige Zahl zu, die logisch folgen müsste:

1, 1, 2, 3, 5, 8, ?

1, 3, 6, 10, ?

65536, 256, 16, ?

Bei dieser Art Tests geht es darum, ein System zu erkennen, das der Zahlenreihe zugrunde liegt. Das Erkennen eines solchen Systems wird gemeinhin als “intelligentes Verhalten” gedeutet – aber wie verhält es sich mit der Logik? Schauen wir uns diese Zahlenreihe an:

3, 5, 7, ?

Einerseits könnte man sagen, die nächste Zahl sei die 9 und dann käme 11, 13, 15 – der Abstand beträgt jeweils 2. Andererseits stehen dort die ersten drei ungeraden Primzahlen, so dass man diese Reihe auch mit 11, 13, 17 fortsetzen könnte.

Keine der beiden Begründungen hat logisch gesehen Vorrang vor der anderen. Kein System ist “besser” oder “naheliegender” als das andere. (Man kann sogar mathematisch begründen, dass die nächste Zahl immer 19 ist, egal wie die Zahlenreihe aussieht. Das können wir hier aber nicht ausführen, dazu brauchen wir erst etwas – nein, viel – mehr mathematische Kenntnisse; siehe Carl E. Lindholm: “Mathematics made difficult”, 1972.)

Hier noch ein vielleicht etwas abstruses Beispiel. Die folgende Reihe von Buchstaben soll “logisch” fortgesetzt werden:

$M, D, M, D, ?, ?, ?$

Eine Möglichkeit wäre natürlich die Folgende:

M, D, M, D, M, D, M

Eine ganz andere Art sieht so aus:

$$M, D, M, D, F, S, S$$

Hierbei interpretiert man die Buchstaben als die Anfangsbuchstaben der Wochentage Montag, Dienstag, Mittwoch und Donnerstag – worauf dann natürlich “logischerweise” Freitag, Samstag und Sonntag folgen.

Auch wenn dieses Beispiel nicht ganz ernst gemeint ist, verdeutlicht es doch ein grundlegendes Dilemma: wie kann man solche Zahlenfolgen, mit denen man es auch in der Mathematik zu tun hat, platzsparend aufschreiben ohne in die “Mehrdeutigkeitsfalle” zu tappen? Nehmen wir mal an, wir möchten die Summe der ersten 10 ungeraden Zahlen bilden:

$$1 + 3 + 5 + 7 + 9 + 11 + 13 + 15 + 17 + 19 = 100$$

(Wer’s nicht glaubt, rechnet es nach.)

Wenn es dann nicht mehr die ersten 10, sondern die ersten 30 (oder 500) ungeraden Zahlen sind, möchte man die Notation vielleicht etwas abkürzen. Aber eine Schreibweise wie

$$1 + 3 + 5 + \cdots + 19 = 100$$

hat es in sich, wie wir gesehen haben – woher kann man sicher sein, dass der Leser des Textes nicht ein anderes System findet und benutzt? Wie kann man deutlich machen, dass man in dieser Summe wirklich alle ungeraden Zahlen haben möchte und zum Beispiel nicht nur die Primzahlen?

Die Lösung liegt in der Benutzung einer besonderen Notation, die das Problem in den Griff bekommt und alle Mehrdeutigkeiten beseitigt. Man schreibt zum Beispiel für die obige Summe

$$\sum_{k=1}^{10} (2k-1) = 100$$

Das sieht auf den ersten Blick verwirrend aus (Wo kommt der Buchstabe k her?), aber ist ungemein praktisch. Das Zeichen ist ein griechischer Buchstabe, ein großes Sigma (Σ). Es soll an S wie Summe erinnern. Der Laufindex k wird eingeführt und soll bei 1 beginnend alle ganzen Zahlen bis einschließlich 10 durchlaufen. (In der Informatik kennt man das als *for*-Schleife mit Zählvariable k).

Für jeden ganzzahligen Wert von k wird der Ausdruck hinter dem Σ berechnet. Alle diese Ausdrücke werden dann aufaddiert.

Es wird sofort deutlich, dass die obige Formulierung tatsächlich das Gewünschte leistet. Für $k = 1$ ist der Klammerausdruck gleich 1 und jedes Mal, wenn k um eines größer wird, vergrößert sich der nächste Summand um 2, es werden also die ungeraden Zahlen durchlaufen. Für $k = 10$ kommt man beim letzten Summanden 19 an.

Hier zeigt sich ein weiterer Vorteil der Notation: nehmen wir mal an, dass wir uns nicht festlegen wollen, bis wohin unsere Summe geht, sondern einfach eine natürliche Zahl $n \in \mathbb{N}$ wählen und die Summe der ersten n ungeraden Zahlen bilden möchten. Dann schreiben wir einfach

$$\sum_{k=1}^n (2k-1)$$

und erhalten das Gewünschte, ganz gleich ob $n = 10$ oder $n = 500$ ist. Später werden wir zeigen, dass der Wert dieser Summe stets n^2 ist.

Analog zum Summenzeichen Σ verwendet man auch ein Zeichen, um Produkte zu bilden. Sollen die Ausdrücke (statt “mathematischer Ausdrücke” sagt man kurz “Terme”) miteinander multipliziert werden, verwendet man ein großes Pi (Π). Das Produkt der ersten n natürlichen Zahlen ($n!$ sprich: “ n **Fakultät**.”) kann also wie folgt geschrieben werden:

$$n! = \prod_{i=1}^n i \quad (1)$$

Wenn klar ist, über welchen Ausdruck zu summieren bzw. multiplizieren ist, werden die Klammern oft fortgelassen.

Es gibt noch eine wichtige Konvention: taucht ein Summenzeichen auf, das keine Summanden enthält, so wird dies die “leere Summe” genannt und bekommt per Definition den Wert 0. Dies macht Sinn, weil 0 das “neutrale Element” der Addition ist – die Addition einer 0 ändert den Wert einer Summe nicht.

Analog wird das “leere Produkt” als 1 definiert, da bei der Multiplikation die 1 das neutrale Element ist. Als Konsequenz erhalten wir:

$$0! = \prod_{i=1}^0 i = 1$$

Aufgabe 1.1. Berechne die folgenden Summen und Produkte:

$$\begin{array}{ccccc} \sum_{k=3}^7 (k^2 - k) & \sum_{j=5}^7 \left(\frac{1}{2^{j-7}} \right) & \sum_{i=-3}^4 i & \sum_{n=-1}^3 (2n+1) & \prod_{k=3}^6 k \\ \sum_{\alpha=0}^6 2^\alpha & \sum_{m=-1}^3 \frac{1}{2^{m+1}} & \prod_{i=8}^{10} i + 1 & \prod_{k=1}^0 k & \prod_{k=0}^0 k \\ \sum_{i=0}^9 10^i & \sum_{i=0}^9 10^i & \prod_{n=-1}^9 (n^3 + n) & \prod_{n=1}^9 m & \end{array}$$

1.3 Definition, Lemma, Satz, Beweis

Seit Jahrhunderten (eigentlich seit Euklid, siehe wikipedia) benutzen mathematische Texte einen strengen logischen Aufbau mit einer bestimmten Schreibweise aus wenigen Bausteinen: Definition, Satz, Beweis... So werden auch die Mathematikvorlesungen strukturiert sein, und – oft weniger strikt – auch Informatik- oder Physikvorlesungen. Am Anfang stehen oft eine **Definition** (oder mehrere). Als wir oben die Fakultät einer natürlichen Zahl definiert haben, hätten wir das auch mittels dieser Schreibweise so schreiben können.

Definition 1.1. Die **Fakultät** einer Zahl $n \in \mathbb{N}$ ist definiert als $n! = \prod_{i=1}^n i$

Ein zentrales Resultat hat den Namen **Satz** (oder **Theorem**). Ein Beispiel ist Satz 1.2 im nächsten Abschnitt. Hinter dem Satz steht oft der Beweis. Am Ende des Beweises schrieb man früher oft qed oder QED (lateinisch für *quod erat demonstrandum*). Heute schreibt man oft einfacher $\square$. Auch das sieht man im nächsten Abschnitt hinter Satz 1.2.

Ein **Lemma**, eine **Proposition** oder ein **Hilfssatz** sind kleinere Ergebnisse, die oft vorbereitend dem Beweis eines Satzes dienen. Ein **Korollar** ist eine Folgerung. Eine etwas andere Rolle hat eine **Vermutung**: das ist eine Aussage, die (noch) nicht bewiesen ist. In der Mathematik äußert man nicht leichtfertig Vermutungen. Manchmal lässt man sich aber hinreißen. Eine der wichtigsten unbewiesenen Vermutungen ist die vor ca 150 Jahren aufgestellte *Riemannsche Vermutung*. Was die genau besagt würde den Rahmen dieses Skripts sprengen. Aber wer sie beweist wird garantiert berühmt.

Außerdem braucht man, wie wir später sehen werden, viele verschiedene Buchstaben als Platzhalter für Funktionen, Variablen, Vektoren, Laufindizes... Da reichen die lateinischen Buchstaben (also die unseres normalen Alphabets) oft nicht aus. Ein Trick ist, die Buchstaben mit Strichen, Balken usw zu verzieren: x ist dann etwas anderes als x' , und etwas anderes als $\bar{x}$, und etwas anderes als $\tilde{x}$ usw. Daneben benutzt man oft tiefgestellte Indizes: x_0 ist was anderes als x_2 usw. Das reicht oft immer noch nicht. Daher benutzt man neben den lateinischen Buchstaben noch griechische:

α	alpha	ε	epsilon	κ	kappa	ξ	xi	τ	tau	ω	omega
β	beta	ζ	zeta	λ	lambda	π	pi	φ	phi		
γ	gamma	η	eta	μ	my	ϱ, ϱ	rho	χ	chi		
δ	delta	ϑ	theta	ι	iota	ν	ny	σ	sigma	ψ	psi

Γ	GAMMA	Θ	THETA	Ξ	XI	Σ	SIGMA	Ψ	PSI
Δ	DELTA	Λ	LAMBDA	Π	PI	Φ	PHI	Ω	OMEGA

Früher waren in der Mathematik auch deutsche Schreibschriftbuchstaben in Benutzung, das ist außer Mode geraten. Ich selbst fände es schön, kyrillische (russische) Buchstaben als Variablen zu benutzen: Sei $\mathbb{I}\mathbb{I} \subset \mathbb{N}$ und $\mathfrak{r} \in \mathbb{I}\mathbb{I}$. Das hat sich aber noch nicht durchgesetzt.

1.4 Indexverschiebung und Summanden einzeln schreiben

Indexverschiebung

Kommen wir nun zu einem etwas formaleren Aspekt der Notation mit Summen- bzw. Produktzeichen. Dabei soll im Folgenden nur das Summenzeichen betrachtet werden, die Überlegungen für das Produktzeichen laufen völlig analog.

Ausgangspunkt dieser Überlegung ist, dass die Wahl des Bereiches, die unsere Zählvariable durchläuft etwas willkürlich ist. Das folgende Beispiel soll das illustrieren:

$$\sum_{k=1}^4 (k+1)^2 = 2^2 + 3^2 + 4^2 + 5^2 = \sum_{k=2}^5 k^2$$

Es ist offensichtlich, dass beide Notationen mit dem Summenzeichen die gleiche Summe meinen. Der Unterschied besteht darin, dass der Laufindex k "verschoben" wurde. In der zweiten Summe läuft er nicht von 1 bis 4, sondern von 2 bis 5 – dafür muss zum Ausgleich in der Summe das k durch $k-1$ ersetzt werden (bzw. $k+1$ durch k).

Ein weiteres Beispiel, aus der obigen Übung:

$$\sum_{j=5}^7 \left(\frac{1}{2j-7} \right) = \sum_{j=1}^3 \left(\frac{1}{2(j+4)-7} \right) = \sum_{j=1}^3 \left(\frac{1}{2j+1} \right)$$

Die wichtige Regel bei der **Indexverschiebung** ist also: wird der Wertebereich des Laufindex nach unten verschoben, so muss der Index selbst zum Ausgleich vergrößert werden und umgekehrt.

Als Anwendung der Indexverschiebung soll für eine reelle Zahl $q \neq 1$ die folgende Formel bewiesen werden:

Satz 1.2. Für alle $n \in \mathbb{N}$ und für alle $q \in \mathbb{R}$ mit $q \neq 1$ gilt: $\sum_{k=0}^n q^k = \frac{1-q^{n+1}}{1-q}$

Die Summe links heißt **(endliche) geometrische Reihe**. Die und ihre unendliche Variante werden uns noch oft begegnen. Die Worte “Summe” und “Reihe” werden oft gleichbedeutend benutzt. Genauer heißt eine unendliche Summe (s. Kap. 4) meistens nicht mehr “Summe”, sondern “Reihe”.

Beweis. Nach Multiplikation mit dem Nenner genügt es zu zeigen:

$$(1 - q) \cdot \sum_{k=0}^n q^k = 1 - q^{n+1}$$

Multipliziere die linke Seite aus:

$$\begin{aligned} (1 - q) \cdot \sum_{k=0}^n q^k &= \sum_{k=0}^n (q^k) - q \cdot \sum_{k=0}^n q^k \\ &= \sum_{k=0}^n (q^k) - \sum_{k=0}^n q^{k+1} \\ &= \sum_{k=0}^n (q^k) - \sum_{k=1}^{n+1} q^k \\ &= 1 + \sum_{k=1}^n (q^k) - \left(\sum_{k=1}^n (q^k) + q^{n+1} \right) \\ &= 1 - q^{n+1} \end{aligned}$$

□

Summanden einzeln schreiben

In diesem Beweis wird ein weiteres wichtiges Prinzip beim Rechnen mit Summen deutlich: man kann **Summanden abspalten und einzeln hinschreiben**. So kann man zum Beispiel die folgende Summe umschreiben:

$$\sum_{k=1}^{n+1} k^2 = \left(\sum_{k=1}^n k^2 \right) + (n+1)^2$$

Oder die folgende elementare Formel für Fakultäten notieren:

$$(n+1) \cdot n! = (n+1) \cdot \prod_{i=1}^n i = \prod_{i=1}^{n+1} i = (n+1)!$$

Eine Schlussbemerkung: Beim Rechnen mit Summen oder Produkten kann es nützlich sein, die Summe auszuschreiben. Zum Beispiel ist

$$\sum_{k=0}^n ((k+1)^2 - k^2) = (n+1)^2.$$

Das sieht mancher vielleicht direkt. Falls nicht, sehen wir vielleicht mehr, wenn wir die Summe ausschreiben. Da das n keinen konkreten Wert hat, benutzen wir die Pünktchenschreibweise vom Anfang des Kapitels. Eine eventuelle Mehrdeutigkeit stellt jetzt kein Problem mehr dar, denn die Summe liegt ja oben in eindeutiger Schreibweise vor. Wenn n nicht allzu klein ist, dann können wir mal die ersten drei und die letzten zwei Summanden hinschreiben:

$$\sum_{k=0}^n (k+1)^2 - k^2 = (1^2 - 0^2) + (2^2 - 1^2) + (3^2 - 2^2) + \dots + (n^2 - (n-1)^2) + ((n+1)^2 - n^2)$$

Sortieren wir ein wenig um, dann ergibt sich:

$$= -0^2 + 1^2 - 1^2 + 2^2 - 2^2 + 3^2 - \dots + (n-1)^2 - (n-1)^2 + n^2 - n^2 + (n+1)^2.$$

Wir sehen, dass sich 1 und -1 genau aufheben, ebenso 2^2 und -2^2 , 3^3 und -3^3 usw. bis n^2 und $-n^2$. übrig bleiben nur der erste Summand (also -0) und der letzte (also $(n+1)^2$). Insgesamt ergibt sich $-0 + (n+1)^2$, also $(n+1)^2$.

Wenn eine Summe so in sich zusammenfällt, heißt sie auch **Teleskopsumme**. Wir haben oben schon eine Teleskopsumme gesehen, im Beweis zur endlichen geometrischen Reihe (Seite 11). Wer möchte kann sich diesen Beweis nochmal mit der Pünktchenschreibweise klar machen.

Aufgabe 1.2. Berechne folgende Ausdrücke:

$$\sum_{k=1}^{10} (k^7 - k^5 + k) + \sum_{k=21}^{30} ((k-20)^5 - (k-20)^7) \qquad \sum_{k=4}^{10} (k-2)^2 - \sum_{k=0}^7 (k+1)^2$$

Aufgabe* 1.3. Berechne folgende Ausdrücke (Die Lösung ist hier keine Zahl, sondern ein einfacher Ausdruck, in dem n vorkommt, aber kein Summenzeichen und kein k .)

$$\sum_{k=1}^n ((k+1)^3 - k^3), \quad \sum_{k=2}^{n+1} ((k-1)! - k!), \quad \prod_{k=1}^n \frac{k+1}{k}$$

Aufgabe* 1.4. Zeige, dass folgende Formel gilt für $n \in \mathbb{N}$.

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

(Man muss etwas tricksen, um diese Reihe auf die Form einer Teleskopsumme zu bekommen!)

1.5 Potenzen

Eine andere Kurzschreibweise sollte werdenden Informatikern in Fleisch und Blut übergehen: Potenzschreibweise. Ein Beispiel: Ein Bit (die kleinste theoretische Speichereinheit im Rechner) enthält die Information "0 oder 1". Also 2 Möglichkeiten. Ein Byte (eine der kleinsten

praktischen Speichereinheiten im Rechner) sind 8 Bit. Dafür gibt es $2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 = 256$ Möglichkeiten, etwas mit einem Byte zu speichern. (Könnten wir durchzählen: 0000 0000, 0000 0001, 0000 0010, 0000 0011 usw.) Ein Kilobyte (kurz KB) waren bis 1995 genau 1024 Byte¹. Das sind $2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2$ Byte bzw.

$$2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \text{ Bit}$$

Das sind 13 Zweien. Es bietet sich an, eine abkürzende Schreibweise zu verwenden. Statt 13 Zweien schreiben wir einfach 2^{13} (gesprochen "Zwei hoch dreizehn").

Wir sehen daran auch schon eine Rechenregel für Potenzen: 1024 sind genau 2^{10} . Also ist ein KB genau 2^{10} Byte, das sind $2^{10} \cdot 2^3$ Bit, und das sind 2^{13} Bit. Es ist

$$2^{10} \cdot 2^3 = 2^{10+3} = 2^{13}.$$

Das ist kein Zufall. Generell gilt: $2^a \cdot 2^b$, das sind a Zweien miteinander multipliziert mal b Zweien miteinander multipliziert, insgesamt also $a + b$ Zweien miteinander multipliziert, also insgesamt 2^{a+b} . Noch allgemeiner gelten die folgenden Regeln.

Satz 1.3. Für $a, b, c \in \mathbb{N}$ gelten die folgenden **Potenzregeln**:

1. $a^b \cdot a^c = a^{b+c}$
2. $a^0 = 1$
3. $(a^b)^c = a^{b \cdot c}$
4. $a^{-b} = \frac{1}{a^b}$
5. $(a \cdot b)^c = a^c \cdot b^c$

Weil a, b und c hier natürliche Zahlen sind, kann man sich diese Gesetze noch durch gesunden Menschenverstand erklären, in Analogie zu der Überlegung oben. Z.B. überlegt man sich etwa zu Regel 2 folgendes: Wegen Regel 1 ist

$$a^b \cdot a^0 = a^{b+0} = a^b.$$

Also muss $a^0 = 1$ sein, sonst stimmt die Gleichung nicht. Bei Regel 5 überlegt man sich z.B.

$$(a \cdot b)^c = a \cdot b \cdot a \cdot b \cdots a \cdot b = a \cdot a \cdots a \cdot b \cdot b \cdots b = a^c \cdot b^c$$

Aufgabe 1.5. Was ist eine Trillion durch 1000 Billionen? Was ist die Hälfte von 2^{30} ? Was ist die Hälfte von 10^{30} ?

Zusatzfrage (knifflig): Was ist $10^{(10^{10})}$ geteilt durch $10^{(10^9)}$?

Aufgabe 1.6. Was ist größer, $(7^7)^7$ oder $7^{(7^7)}$?

Sortiere die Ausdrücke der Größe nach: $((7^7)^7)^7$, $7^{(7^{(7^7)})}$, $(7^{(7^7)})^7$, $(7^7)^{(7^7)}$, $7^{((7^7)^7)}$.

¹Heute ist ein KB genau 1000 Byte, um Verwechslungen zu vermeiden: "Kilo" heißt "Tausend". 1024 Byte heißen heute "Kibibyte", kurz KiB.

Aufgabe 1.7. Vereinfache die folgenden Terme mit Hilfe der Potenzregeln (und anderen Tricks). Oft ist ein Problem, dass man nicht weiß, wann man fertig ist. Daher ist hier bei jeder Aufgabe in eckigen Klammern angegeben, wie viele Zeichen die Antwort benötigt. Ein Bruchstrich zählt als ein Zeichen. Ein “Mal” kann man weglassen.

- | | | | |
|---|-----|---|-----|
| (a) $\frac{2^{2n}}{2^n}$ | [2] | (b) $\frac{(a^n)^2}{a^n}$ | [2] |
| (c) $\frac{2^{1000} + 2^{1001}}{2^{999}}$ | [1] | (d) $(-x)^{25} + x^{25}$ | [1] |
| (e) $\frac{(x+x)^{50}}{2^{50}x^{49}}$ | [1] | (f) $\frac{a^a}{a^b}a^{b-a+1}$ | [1] |
| (g) $\frac{(a-b)^9}{(\frac{1}{a}-\frac{1}{b})^9}$ | [5] | (h) $\frac{(x^n+x^{n+1})(1-x)}{(x^n-x^{n-1})(x+1)}$ | [2] |
| (i) $\sum_{k=0}^{50} 2^k$ | [5] | (j) $\sum_{k=1}^{50} 2^k$ | [5] |
| (k) $\sum_{k=2}^{50} 2^k$ | [5] | (l) $\sum_{k=2}^{50} 3^k$ | [7] |

1.6 Rechentricks

Bei den letzten Aufgaben galt es, durch bestimmte Tricks, nämlich “legale” Umformungen mathematischer Ausdrücke (kurz: **Terme**) zum Ziel zu kommen. Um uns für diese Tricks, dieses Manipulieren mathematischer Ausdrücke, einen Namen auszudenken, nennen wir es vornehm “Terme umformen”.

Terme umformen ist das wichtigste Handwerkszeug in der Mathematik.

Beim Vereinfachen von Ausdrücken ist das nötig, beim Bestimmen von Grenzwerten, beim Ableiten oder Integrieren von Funktionen oder das Führen von Beweisen, immer wieder werden wir Terme umformen müssen. Da uns das also sowieso begleitet — hier im Vorkurs oder später in den Mathevorlesungen — brauchen wir dem eigentlich kein eigenes (Unter-)Kapitel zu widmen. Das tun wir dennoch, weil wir erstens damit auf die Wichtigkeit dieses Werkzeugs hinweisen können, und zweitens ein paar Tricks und ein paar Grundprinzipien schildern.

Ein paar elementare Regeln, die beim Vereinfachen/Umformen von Termen immer wieder nützlich sind.

- Binomische Formeln: $(a+b)^2 = a^2+2ab+b^2$, $(a-b)^2 = a^2-2ab+b^2$, $a^2-b^2 = (a-b)(a+b)$.
- $p-q$ -Formel: Die Nullstellen von $x^2 + px + q$ sind $-\frac{p}{2} \pm \sqrt{\left(\frac{p}{2}\right)^2 - q}$ (hat keine (reelle) Lösung, falls unter der Wurzel eine negative Zahl steht, eine Lösung, falls unter der Wurzel eine 0 steht, zwei Lösungen sonst).
- Unfallfreies Bruchrechnen. (Das können wir in diesem Kurs nicht nachholen. Wenn ihr merkt, dass ihr da Nachholbedarf habt: es gibt viele Angebote im Netz. Das Problem ist, die guten vom Quatsch zu trennen. Ganz OK sind <http://www.bruchrechnen.de>, <http://mathematik.net> und <https://de.wikipedia.org/wiki/Bruchrechnung>. Es gibt ein gutes Buch zum Nachholen mathematischer Grundlagen, dass über die Unibibliothek online verfügbar ist: “Vorkurs Mathematik” von E. Cramer und J. Nešlehová. Kapitel 3.1 widmet sich dem Bruchrechnen. Übungsaufgaben findet ihr im Folgenden in diesem Text, denn Bruchrechnen kommt immer wieder vor.)

Zwei Vorgehensweisen beim Beweis von Gleichheit Wenn die Aufgabe die Form hat “Zeigen Sie, dass $\frac{\frac{1}{a+b}(a^2-b^2)}{a-b} + a\frac{a-b}{ab-a^2} = 0$ ist (für $a \neq b$)”, so gibt es im Wesentlichen zwei Vorgehensweisen:

Methode 1. Entweder ich fange mit der einen Seite an (hier: der komplizierte Term links vom “=”) und forme ihn solange legal um, bis die andere Seite da steht (hier: die rechte Seite, also 0). Oder

Methode 2. Ich schreibe die Gleichung hin und forme sie um, indem ich jeweils rechts und links dieselben Umformungen durchführe. Hier also auch: beide Seiten mit dem selben Term addieren/subtrahieren/multiplizieren usw.

Beispiel: Methode 1 würde in diesem Beispiel etwa so aussehen:

$$\begin{aligned} \frac{\frac{1}{a+b}(a^2-b^2)}{a-b} + a\frac{a-b}{ab-a^2} &= \frac{a^2-b^2}{(a+b)(a-b)} + \frac{a(a-b)}{a(b-a)} = \frac{a^2-b^2}{a^2-b^2} + \frac{(a-b)}{(b-a)} \\ &= 1 + \left(-\frac{b-a}{b-a}\right) = 1 - 1 = 0. \end{aligned}$$

Wir haben hier eine ununterbrochene Kette von Gleichheiten, also ist der Term ganz am Anfang gleich dem ganz am Ende.

Methode 2 würde etwa so aussehen. (Dabei heißt das Zeichen $\Leftrightarrow$ soviel wie “ist gleichbedeutend mit”.)

$$\begin{aligned} \frac{\frac{1}{a+b}(a^2-b^2)}{a-b} + a\frac{a-b}{ab-a^2} &= 0 & \Big| \text{ auf beiden Seiten } -a\frac{a-b}{ab-a^2} \\ \Leftrightarrow \frac{a^2-b^2}{(a+b)(a-b)} &= -\frac{a(a-b)}{a(b-a)} \\ \Leftrightarrow \frac{a^2-b^2}{a^2-b^2} &= -\frac{(a-b)}{(b-a)} \\ \Leftrightarrow 1 &= \frac{b-a}{b-a} \\ \Leftrightarrow 1 &= 1 & \text{ (wahr) } \end{aligned}$$

Hier haben wir die Gleichung, die wir nachweisen wollen (die oberste) solange mit legalen Mitteln umgeformt, bis wir eine Gleichung erhielten, die offensichtlich wahr ist, nämlich $1 = 1$. Da die letzte Gleichung wahr ist, und die anderen Gleichungen alle gleichbedeutend mit dieser sind, ist auch die erste Gleichung wahr.

Bei Methode 2 hat man mehr Möglichkeiten (man kann auf beiden Seiten Terme abziehen/addieren/usw), aber man muss wissen, was herauskommt. Wenn wir nur den linken Term gegeben hätten und die Aufgabe hätte gelautet: “Vereinfache so weit wie möglich”, dann hätten wir zu Methode 1 greifen müssen.

Aufgabe* 1.10. Zeige, dass (für $a \neq b, a \neq 0, b \neq 0$) gilt:

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

Vereinfache so weit wie möglich ($k \neq \pm 1$; das Ergebnis erfordert nur ein Zeichen)

$$\frac{k^2 - k}{k^2 - 1} + \frac{2k - 1}{(k + 1)^2} - \frac{k - 2}{k^2 + 2k + 1}$$

1.7 Wurzeln

Später werden in der Vorlesung die Potenzen anders erklärt als oben, nämlich über die Exponentialfunktion. Das ist viel komplizierter. Warum sollte man das also kompliziert erklären, wenn es auch einfach geht? Die Antwort ist, dass wir a^b bisher nur für natürliche Zahlen a, b erklären können. Wir können das problemlos auch für "krumme" Zahlen a erklären, solange b eine natürliche Zahl ist. Z.B. wäre π^3 ja einfach $\pi \cdot \pi \cdot \pi$. Dabei ist π die "Kreiszahl"². Umgekehrt gilt aber: Wie sollen wir 3^π erklären? Was soll es bedeuten, die Drei π -mal mit sich selbst malzunehmen? Wie man dieses Problem löst, wird (normalerweise) in der Vorlesung "Mathe I für Naturwissenschaften" im ersten Semester behandelt.

Dennoch bringt uns der einfache Ansatz noch etwas weiter. Wir können fragen, was etwa $4^{\frac{1}{2}}$ ist, so dass das mit den Regeln 1.-4. aus Satz 1.3 vernünftig zusammenpasst. Wir sehen etwa wegen Regel 1:

$$4^{\frac{1}{2}} \cdot 4^{\frac{1}{2}} = 4^{\frac{1}{2} + \frac{1}{2}} = 4^1 = 4.$$

$4^{\frac{1}{2}}$ ist also die Zahl, die mit sich selbst malgenommen 4 ergibt. Also 2. Oder auch -2 , aber wir machen es uns hier einfach (wir bestimmen die Regeln!)

In diesem Abschnitt sollen alle betrachteten Zahlen positiv sein!

Also ist $4^{\frac{1}{2}} = 2 = \sqrt{4}$. Genau so ist $8^{\frac{1}{3}} = 2 = \sqrt[3]{8}$ (also die dritte Wurzel aus 8). Allgemein gilt für $a, b \in \mathbb{N}$:

$$a^{\frac{1}{b}} = \sqrt[b]{a} \quad (\text{Also die } b\text{-te Wurzel aus } a) \quad (2)$$

Rechentechisch sind Wurzeln kompliziert und ärgerlich: Terme lassen sich nicht vereinfachen, oder beim Rechnen machen sie aus ganzen Zahlen (im Rechner: **integer**) oft krumme Zahlen (im Rechner: **float**, **double**...). Ein Weg zur Vereinfachung sind Potenzen. Entweder können wir z.B. mit dem Quadrat einer Zahl rechnen statt mit der Zahl selbst. Oder wir benutzen den Umstand, dass wir Wurzeln jetzt als Potenzen schreiben können: dann können wir die Potenzregeln aus Satz 1.3 benutzen.

Als erstes sehen wir, dass (Quadrat-)Wurzelziehen das Umkehren von Quadrieren ist. Eine Zahl $a \geq 0$ (oder allgemeiner ein Ausdruck) zum Quadrat genommen und dann die Wurzel gezogen liefert wieder a . Oder umgekehrt. Dabei machen wir hier regen Gebrauch von den Potenzregeln:

$$\sqrt{a^2} = (a^2)^{\frac{1}{2}} = a^{2 \cdot \frac{1}{2}} = a^1 = a = a^{\frac{1}{2} \cdot 2} = (\sqrt{a})^2$$

Das ist eigentlich bereits alles, was man sich zu Wurzeln aus positiven Zahlen merken sollte. Kennt man die Potenzregeln, so muss man sich sogar nur Gleichung (2) merken, der Rest folgt aus den Potenzregeln. Z.B. ist ja

$$\sqrt[n]{a^n} = (a^{\frac{1}{n}})^n = a^{\frac{1}{n} \cdot n} = a^1 = a.$$

wegen (2), Regel 3 und Regel 1.

²Also das Verhältnis zwischen Durchmesser und Umfang eines (perfekten) Kreises. Es ist $\pi = 3,1415926 \dots$, aber die Nachkommastellen gehorchen keinem bekannten Gesetz. Man kennt heute ein paar Billionen Nachkommastellen. Man weiß aber z.B. nicht, ob irgendwo die Sequenz 1111111111111111 vorkommt.

Beispiel 1.4. Ein weiteres Beispiel zum Benutzen der Potenzregeln beim Rechnen mit Wurzeln: Wir wollen prüfen, ob $\sqrt{\sqrt{2} + \sqrt{2}} = \sqrt{\sqrt{2^3}}$ ist.

Eine Möglichkeit, dass zu zeigen, ist beide Seiten der Gleichung wiederholt umzuformen und zu quadrieren. Wir wollen hier aber den Zusammenhang Wurzel $\leftrightarrow$ Potenzen nutzen. Also:

$$\begin{aligned} \sqrt{\sqrt{2} + \sqrt{2}} &= \sqrt{\sqrt{2^3}} \\ \Leftrightarrow (\sqrt{2} + \sqrt{2})^{\frac{1}{2}} &= (\sqrt{2^3})^{\frac{1}{2}} \\ \Leftrightarrow (2\sqrt{2})^{\frac{1}{2}} &= ((2^3)^{\frac{1}{2}})^{\frac{1}{2}} \\ \Leftrightarrow 2^{\frac{1}{2}} \sqrt{2}^{\frac{1}{2}} &= (2^3)^{\frac{1}{4}} \\ \Leftrightarrow 2^{\frac{1}{2}} 2^{\frac{1}{4}} &= 2^{3 \cdot \frac{1}{4}} \\ \Leftrightarrow 2^{\frac{1}{2} + \frac{1}{4}} &= 2^{\frac{3}{4}} \end{aligned}$$

Die letzte Gleichung stimmt (wegen Regel 2). Also stimmt auch die erste.

Was ist nun mit negativen Zahlen? Da muss man aufpassen (Wurzel des Quadrats von a muss nicht a sein):

$$\sqrt{(-2)^2} = \sqrt{4} = 2 \neq -2$$

Um unfallfrei mit negativen Ausdrücken hantieren zu dürfen, ist die folgende Notation nützlich: Der **Betrag** (auch “Absolutbetrag”) einer Zahl ist die Zahl selbst ohne ihr Vorzeichen. Geschrieben wird der Betrag von x als $|x|$. also ist z.B. $|2| = 2$, $|-5| = 5$, $|\frac{-1}{-2}| = \frac{1}{2}$, $|0| = 0$. Damit können wir uns alternativ zur obigen Regel auch merken

$$\sqrt{a^2} = \sqrt{|a|^2} = |a|.$$

Ein wichtiger Trick im Umgang mit Wurzeln ist die Entfernung von Summen von Wurzeln aus dem Nenner eines Bruchs.

Merkregel: Sieht man im Nenner eines Bruchs etwas von der Form $\sqrt{a} + \sqrt{b}$, so sollte man mit $\sqrt{a} - \sqrt{b}$ erweitern, damit keine Wurzeln mehr im Nenner auftauchen.

Wegen der dritten binomischen Formel (s. Seite 14) ist ja $(\sqrt{a} - \sqrt{b})(\sqrt{a} + \sqrt{b}) = \sqrt{a}^2 - \sqrt{b}^2 = a - b$. Die Wurzeln verschwinden also.

Beispiel 1.5. Vereinfachen von $\frac{\sqrt{12} - \sqrt{20}}{\sqrt{3} + \sqrt{5}}$:

$$\begin{aligned} \frac{\sqrt{12} - \sqrt{20}}{\sqrt{3} + \sqrt{5}} &= \frac{(2\sqrt{3} - 2\sqrt{5})(\sqrt{3} - \sqrt{5})}{(\sqrt{3} + \sqrt{5})(\sqrt{3} - \sqrt{5})} = \frac{2(\sqrt{3} - \sqrt{5})^2}{3 - 5} = -(3 - 2\sqrt{3}\sqrt{5} + 5) = -3 + 2\sqrt{15} - 5 \\ &= -8 + 2\sqrt{15} \end{aligned}$$

Aufgabe* 1.11. Leite die folgenden **Wurzelregeln** aus den Potenzregeln ab.

1. $\sqrt[n]{a^m} = \sqrt[n]{a^m}$
2. $\sqrt[n]{a} \sqrt[n]{b} = \sqrt[n]{ab}$
3. $\frac{\sqrt[n]{a}}{\sqrt[n]{b}} = \sqrt[n]{\frac{a}{b}}$

4. $\sqrt[n]{\sqrt[k]{a}} = \sqrt[nk]{a}$

Aufgabe 1.12. Zeige, dass $\sqrt{3\sqrt{3}} = \sqrt[4]{27}$ ist, $\sqrt{64x^4} = 8x^2$, $\sqrt[4]{81a^2} = 3\sqrt{a}$ sowie $\sqrt{x\sqrt{x}} = \sqrt[4]{x^3}$.

Aufgabe 1.13. Wieder kann die Antwort jeweils mit so wenigen Zeichen geschrieben werden wie in der Klammer angegeben.

Was ist $\sqrt{3^2 + 4^2}$? [1]

Was ist $\sqrt{8100}$? [2]

Was ist $\frac{12\sqrt{9}}{\sqrt[6]{3}}$? [1]

Was ist $\sqrt[6]{64a^2 729a^4}$? [2]

Was ist $\frac{a^{3/2}}{\sqrt{a}}$? [1]

Was ist $\frac{a-b}{\sqrt{a}-\sqrt{b}}$? [5]

Was ist $\frac{1-x^3}{1-\sqrt{x}}$? [5]

Aufgabe* 1.14. (Man muss etwas tricksen) Es gilt $x + 16 = 8\sqrt{x}$. Welchen Wert hat x ? Es gilt $x^4 = 5x^2 - 4$. Welche Werte kann x haben?

Aufgabe* 1.15. (Ziemlich länglich; für die, die die Herausforderung suchen) Zeige, dass für $-1 < x < 1$ gilt:

$$\frac{-\frac{1}{\sqrt{1-x^2}} + \frac{(1-x)x}{(1-x^2)^{\frac{3}{2}}}}{1 + \left(\frac{1-x}{\sqrt{1-x^2}}\right)^2} = -\frac{1}{2\sqrt{1-x^2}}$$

Vereinfache $\frac{\left(\frac{\sqrt{1-x^2} - \frac{-2x^2}{2\sqrt{1-x^2}}}{1-x^2}\right)}{1 + \left(\frac{x}{\sqrt{1-x^2}}\right)^2}$ so weit wie möglich. (Das Ergebnis kommt mit 7 Zeichen aus.

Wurzel und Bruchstrich gelten als je ein Zeichen.)

1.8 Logarithmen

Logarithmen sind in gewissem Sinne die Umkehrung von Potenzen. Bei Potenzen lautet die Frage z.B. “zwei hoch vier ist was?”. Bei Logarithmen lautet die Frage “zwei hoch was ist sechzehn?”.

Die kurze Antwort auf die zweite Frage ist vier. In vornehmen Worten — bzw auf mathematisch — würde die zweite Frage lauten: “Was ist der Logarithmus (zur Basis 2) von 16?” Lautet die Frage “zehn hoch was ist tausend?”, dann hieße das auf mathematisch: “Was ist der Logarithmus (zur Basis 10) von 1000?” Allgemeiner kann man sich merken:

$$a^x = b, \quad \text{gegeben } a \text{ und } b, \text{ gesucht } x, \quad \text{Antwort: } \log_a(b) = x.$$

Wenn es nicht zu Unklarheiten führt, lässt man die Klammern auch gerne weg und schreibt

$$\log_a b = x.$$

Offenbar gibt es nicht *den* Logarithmus, sondern die Zahl a spielt eine Rolle. Dieses a heißt **Basis** des jeweils betrachteten Logarithmus. In der Informatik kommen hauptsächlich zwei (bis drei) Basen im Zusammenhang mit Logarithmen vor: 2 und 10 (und $e = 2,71828\dots$, siehe unten).

Beispiel 1.6. Es ist $\log_2 8 = 3$, $\log_2 16 = 4$, $\log_2 32 = 5$, usw. Es ist $\log_3 3 = 1$, $\log_3 9 = 2$, $\log_3 27 = 3$, usw. Es ist $\log_5 25 = 2$, $\log_5 125 = 3$, $\log_5 625 = 4$, usw.

Offenbar ist $\log_{10}$ etwa die Zahl der Dezimalstellen einer Zahl. Genauer:

Zehner-Logarithmus von a abgerundet $+1$ = Zahl der Dezimalstellen von a .
Zweier-Logarithmus von a abgerundet $+1$ = Zahl der Stellen von a in Binärdarstellung.

Zum Rechnen mit Logarithmen muss man sich eigentlich wieder nur die Potenzgesetze merken, zusammen mit

$$a^{\log_a b} = b, \quad (\log_a a^b) = b$$

Aufgabe 1.16. $\log_2 4$, $\log_2 2$, $\log_2 1$, $\log_2 64$, $\log_2 \frac{1}{2}$, $\log_2 2^n$? Und $\log_2 0$?
 $\log_{10} 10.000$, $\log_{10} 1.000.000$, $\log_{10} 1.000.000.000.000$, $\log_{10} 10^{15}$, $\log_{10} 10^{2n}$?

Aufgabe 1.17. $\log_2 4^n$, $\log_2 64^n$, $\log_2 ((2^n)^n)$, $(\log_2(2^n))^n$, $\log_2(\frac{1}{4}2^{n+1} + 2^{n-1})$?
 $\log_{10}(\sqrt{10})$, $\log_{10}(10^9 \cdot 10^3)$?

Aufgabe* 1.18. Zeige die Gültigkeit der folgenden **Logarithmusregeln** unter Verwendung der Potenzregeln ($a, b, c \in \mathbb{N}$, $a \neq 1$).

1. $\log_a(b \cdot c) = \log_a b + \log_a c$
2. $\log_a 1 = 0$
3. $\log_a(b^c) = c \log_a b$
4. $\log_a(\frac{b}{c}) = \log_a b - \log_a c$
5. $\log_b x = \frac{\log_a x}{\log_a b}$

Tipp: Logarithmen ziehen ist eine legale Umformung. In anderen Worten: Die Gleichung $x = y$ ($x > 0, y > 0$) stimmt genau dann, wenn $\log_a x = \log_a y$. Ebenso stimmt $a^x = a^y$ genau dann, wenn $x = y$.

Aufgabe 1.19. Noch eine Aufgabe zum Handwerk, also zum Terme umformen: (Wieder in Klammern hinter jeder Teilaufgabe die Zahl der Zeichen, die die Antwort benötigt.)

- Was ist $\log_7 98 - \log_7 14$ [1]
 Was ist $\log_2 \frac{1}{8} + \log_2 \sqrt{2} + \log_2 \sqrt{8}$ [2]
 Was ist $\log_a(6) - \log_a(9) + \log_a(12) - \log_a(8)$ [1]
 Was ist $n^{\frac{5}{\log_2(n^5)}}$ [1]
 Was ist $2^x - 3^{x \log_3 2}$ [1]

Aufgabe* 1.20. Zeige, dass für $x \geq 1$, $a \in \mathbb{N}$ gilt:

$$\log_a(x + \sqrt{x^2 - 1}) + \log_a(x - \sqrt{x^2 - 1}) = 0$$

Tabelle mathematischer Symbole

$\{\}$	Mengenklammer
$:=$	wird definiert als
$\mathbb{N}$	Menge der natürlichen Zahlen $\{1, 2, 3, \dots\}$
$\mathbb{Z}$	Menge der ganzen Zahlen $\{\dots, -2, 1, 0, 1, 2, 3, \dots\}$
$\mathbb{Q}$	Menge der rationalen Zahlen $\{\frac{k}{m} \mid k, m \in \mathbb{Z}\}$
$\mathbb{R}$	Menge der reellen Zahlen
$\mathbb{R}^+$	Menge der positiven reellen Zahlen
$\mathbb{R}_0^+$	Menge der nicht-negativen reellen Zahlen
$<$	kleiner als
$>$	größer als
$\leq$	kleiner oder gleich
$\geq$	größer oder gleich
$[a, b]$	abgeschlossenes Intervall (inklusive a und b)
$]a, b[$	offenes Intervall (ohne a und b)
$]a, b]$	halboffenes Intervall (ohne a , mit b)
$[a, b[$	halboffenes Intervall (mit a , ohne b)
$\in$	ist Element von (einer Menge)
$\notin$	ist nicht Element von
$\subset$	ist Teilmenge von
$\not\subset$	ist nicht Teilmenge von
$\cup$	Vereinigung (zweier Mengen)
$\cap$	Schnitt (zweier Mengen)
$\setminus$	Differenz zweier Mengen ("A ohne B")
$\sum$	Summenzeichen
	(z.B. $\sum_{i=0}^6 2^i$ "Summe für i von 0 bis 6 über 2 hoch i ")
$\prod$	Produktzeichen (analog zum Summenzeichen)
$!$	Fakultät (z.B. $3! = 3 \cdot 2 \cdot 1$)
$\forall$	Für alle (Kap. 3)
$\exists$	Es gibt (Kap. 3)
$\Rightarrow$	daraus folgt (Kap. 3)
$\Leftrightarrow$	genau dann, wenn (Kap. 3)
$\wedge$	und (Kap. 3)
$\vee$	oder (Kap. 3)

2 Die Kunst: Beweisen — vollständige Induktion

Wir wollen wissen - wir werden wissen.

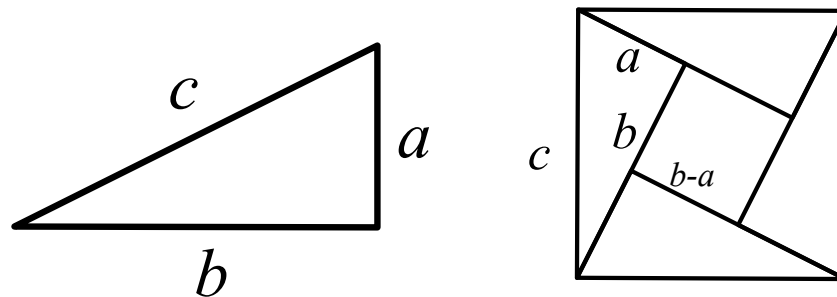
NSA

Nun von der Notation und einigen grundlegenden Rechengesetzen direkt zur Königsdisziplin der Mathematik: dem Beweis. Einen Beweis kann man sich in etwa so vorstellen wie in einem Polizeikrimi: Die Frage ist, wer der Mörder ist. Es gibt verschiedene Theorien, aber am Ende wird der echte Mörder überführt. Mittels eines Beweises, der keine Zweifel übrig lässt.

Das echte Leben ist leider etwas komplizierter. Weder bei einem Mord noch in der Physik noch in der Medizin noch den meisten anderen Wissenschaften (von Philosophie oder Theologie erst gar nicht zu reden) sind Beweise jemals hundertprozentig. Sichere Wahrheiten gibt es dort nicht, nicht im wissenschaftlichen Sinne. In der Mathematik ist das anders. Wenn etwas bewiesen wurde, dann ist das wahr. Der Satz des Pythagoras ist wahr, denn wir können ihn beweisen. Etwa so:

Satz 2.1 (Pythagoras). *In einem rechtwinkligen Dreieck, in der die längste Seite die Länge c hat und die anderen beiden Seiten die Längen a und b , gilt $a^2 + b^2 = c^2$.*

Beweis. Wir nehmen vier Kopien unseres Dreiecks (im Bild links) und legen sie zu einem großen Quadrat mit Seitenlänge c zusammen (im Bild rechts). In der Mitte entsteht ein kleines Quadrat mit Seitenlänge $b - a$. Jetzt berechnen wir die Fläche des großen Quadrats auf zwei Weisen:



Zum einen ist die Fläche des großen Quadrats einfach c^2 . Zum anderen ist die Fläche des großen Quadrats gleich der Fläche der vier Dreiecke plus der Fläche des kleinen Quadrats. Also $4 \cdot \frac{1}{2}ab + (b - a)^2$. Jetzt benutzen wir die binomische Formel und vereinfachen diesen Ausdruck etwas:

$$4 \cdot \frac{1}{2}ab + (b - a)^2 = 2ab + b^2 - 2ab + a^2 = b^2 + a^2.$$

Die Fläche des großen Quadrats ist also sowohl c^2 als auch $a^2 + b^2$. Also ist $c^2 = a^2 + b^2$. $\square$

Dieser Beweis funktioniert für alle möglichen Werte von a, b, c . (Na gut, b muss hier größer als a sein, sonst wird $b - a$ negativ, das würde die Argumentation kaputtmachen. Aber wir können immer die kürzeste Dreiecksseite a nennen und die zweitkürzeste b .) Somit haben wir den Satz für alle rechtwinkligen Dreiecke bewiesen. Wenn wir konkrete Zahlen als Seitenlängen nehmen, würden wir den Satz immer bestätigt finden. Aber egal, wie viele konkrete Zahlenbeispiele wir prüfen, theoretisch könnte es immer sein, dass beim nächsten Beispiel der Satz versagt. Ein wirklicher Beweis verlangt, dass wir alle (unendlich vielen) Fälle beweisen. Das haben wir hier

erreicht, indem wir die Zahlen durch Buchstaben ersetzt haben, die für jede beliebigen Zahl stehen können. Das ist einer der grundlegenden Tricks.

Wir wissen heute ohne jeden Zweifel, dass sich etwa π oder $\sqrt{2}$ nicht als Bruch zweier ganzer Zahlen schreiben lassen. Denn das wurde bewiesen. (Im Falle von $\sqrt{2}$ bereits von den antiken Griechen vor über 2000 Jahren, im Falle von π von dem Schweizer Mathematiker Lambert vor etwa 250 Jahren.) Der Umstand, dass Mathematik wahre Antworten liefern kann, verleiht der Mathematik eine besondere Rolle. übrigens auch der theoretischen Informatik, auch dort lassen sich Aussagen beweisen. Der Grund ist ihre enge Verwandtschaft mit der Mathematik, oder anders formuliert: Auch Programmieren findet in einer idealen Welt statt, der digitalen, und nicht in der realen Welt.

Mathematik und Informatik können also Fragen zweifelsfrei beantworten, anders als jede andere Wissenschaft. Das Problem ist natürlich, ob die Fragen jemanden interessieren. Das kann man aber auch umdrehen: Sobald sich eine Fragestellung aus der realen Welt in die Sprache der Mathematik oder der Informatik übersetzen lässt (innerhalb eines Modells, oder unter bestimmten Annahmen) liefern Mathematik oder Informatik exakte Antworten. Diese können sehr wertvoll oder relevant sein, wenn das Modell gut ist oder die Annahmen realistisch.

In diesem Kapitel nähern wir uns diesem Thema mit der Beweismethode der “vollständigen Induktion”. Die benutzt man häufig, wenn eine Aussage der Form “für alle natürlichen Zahlen gilt...” bewiesen werden soll.

2.1 Vollständige Induktion

Wir wollen ein Beispiel betrachten. Schauen wir uns die Summen der ersten ungeraden Zahlen an, die wir schon im ersten Kapitel gesehen haben:

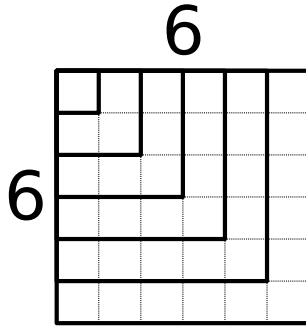
$$\begin{aligned}1 &= 1 \\4 &= 1 + 3 \\9 &= 1 + 3 + 5 \\16 &= 1 + 3 + 5 + 7 \\25 &= 1 + 3 + 5 + 7 + 9\end{aligned}$$

Es fällt auf, dass auf der linken Seite stets eine Quadratzahl steht. Diese Vermutung kann man jetzt mit Hilfe der Summennotation allgemein für beliebiges $n \in \mathbb{N}$ formulieren. Wir formulieren es sofort als Satz, denn — Spoilerwarnung — später wird es uns in der Tat gelingen, diese Formel für alle $n \in \mathbb{N}$ zu beweisen.

Satz 2.2. *Für alle $n \in \mathbb{N}$ gilt:*

$$\sum_{k=1}^n (2k-1) = n^2$$

Für $n \in \{1, 2, 3, 4, 5\}$ haben wir uns schon von der Richtigkeit der Vermutung überzeugt. Aber wie können wir wirklich sicher sein, dass das auch für $n = 1000$ noch klappt? Natürlich kann man nicht alle Fälle nachrechnen, weil das unendlich viele sind. Was ist der Ausweg aus diesem Dilemma? Ein möglicher geometrischer Beweis ginge so:



Beginnend mit einem Kästchen in der oberen linken Ecke, wird immer eine ungerade Anzahl an Kästchen hinzugefügt und es entsteht jeweils wieder ein Quadrat. (Das Bild zeigt die Situation für $n = 6$.) Solch ein Beweis ist durchaus zulässig, der Satz 2.2 ist damit bereits bewiesen. Wir wollen aber an diesem Problem ein sehr mächtiges Beweisverfahren schildern: die “vollständige Induktion”. Gegeben ist eine Aussage in Abhängigkeit von n , sagen wir $A(n)$. In unserem Fall ist also $A(n)$ die Aussage

$$\sum_{k=1}^n (2k - 1) = n^2$$

Für gegebenes $n \in \mathbb{N}$ kann $A(n)$ wahr oder falsch sein. Wir haben uns oben davon überzeugt, dass $A(1)$ bis $A(5)$ wahr sind, möchten aber gern beweisen, dass $A(n)$ für alle $n \in \mathbb{N}$ eine wahre Aussage darstellt. Die vollständige Induktion geht nun in zwei Schritten vor:

- (IA) *Der Induktionsanfang.* Man beweist, dass $A(1)$ gilt. (Oder $A(0)$ oder $A(4)$, generell für einen geeigneten Anfangswert.)
- (IS) *Der Induktionsschritt (oder Induktionsschluss).* Man beweist: falls für ein $n \in \mathbb{N}$ die Aussage $A(n)$ wahr ist, dann auch die Aussage $A(n + 1)$.

Die Annahme im Induktionsschritt, dass $A(n)$ wahr ist, hat auch einen Namen: *Induktionsvoraussetzung*, kurz (IV). Manchmal braucht man auch, dass $A(k)$ wahr ist für alle $k \leq n + 1$; egal, das macht keinen Unterschied.

Der Trick besteht also darin, dass man im Induktionsschritt von der Aussage für ein **beliebiges** $n \in \mathbb{N}$ auf den Nachfolger $n + 1$ schließt. Hat man beide Schritte gezeigt, ist man fertig und hat bewiesen, dass die Aussage für alle $n \in \mathbb{N}$ gilt.

Im Detail stelle man sich das so vor: der Induktionsanfang beinhaltet die Wahrheit der Aussage $A(1)$. Jetzt wenden wir den Induktionsschritt für $n = 1$ an und erhalten, dass $A(2)$ eine wahre Aussage ist. Jetzt aber können wir wieder den Induktionsschritt für $n = 2$ anwenden und erhalten $A(3)$ und so weiter.

Ein hilfreiches Bild ist vielleicht das Folgende: man stelle sich die Aussagen wie Dominosteine vor, die aneinander gereiht sind. Wenn ein Dominostein umfällt, dann soll das bedeuten, dass die Aussage wahr ist. Der Induktionsschritt formuliert nun das Gesetz, dass ein fallender Dominostein seinen Nachbarn umstößt. (Wenn $A(n)$ gilt, dann auch $A(n + 1)$.) Und der Induktionsanfang garantiert uns das Fallen des ersten Steins – und damit fallen alle um.

Soweit die graue Theorie. Widmen wir uns dem Beweis von Satz 2.2.

Beweis. (IA): der Induktionsanfang ist bereits geleistet, wir haben uns davon überzeugt, dass $A(1)$ wahr ist.

(IS): Induktionsschritt: Sei $n \in \mathbb{N}$ beliebig, so dass $A(n)$ wahr ist, d.h. es gelte

$$\sum_{k=1}^n (2k-1) = n^2$$

Zu zeigen ist, dass aus dieser Voraussetzung $A(n+1)$ folgt. Rechnen wir das nach:

$$\sum_{k=1}^{n+1} (2k-1) = \sum_{k=1}^n (2k-1) + (2(n+1)-1) \stackrel{IV}{=} n^2 + 2n + 1 = (n+1)^2$$

□

Das kleine "IV" (Induktionsvoraussetzung) deutet die Stelle an, an der benutzt wird, dass $A(n)$ laut Annahme gilt.

Das Prinzip der vollständigen Induktion ist oft nur ein Baustein in längeren Beweisen. Man kann viele Aussagen der Form „Für alle $n \in \mathbb{N}$ gilt [eine Formel]“ beweisen, Gleichungen, aber auch Ungleichungen wie

$$\text{Für alle } n \in \mathbb{N}, \text{ gilt: } n^2 < 2^n.$$

Klingt plausibel, oder? Wenn man genau hinsieht, merkt man aber:

$1^2 < 2^1$, wahr. $2^2 < 2^2$, falsch! $3^2 < 2^3$, falsch! $4^2 < 2^4$, auch falsch! $5^2 < 2^5$, wahr. Und danach scheint alles gut zu gehen: $36 < 64$, $49 < 128$, Das ist ein Beispiel einer vollständigen Induktion, wo der Induktionsanfang bei $n = 5$ liegen muss (denn für $n \in \{2, 3, 4\}$ ist die Aussage ja falsch.)

Es wird eine Anekdote über den deutschen Mathematiker Carl-Friedrich Gauß erzählt, der im Alter von neun Jahren in der Schule die Strafarbeit aufbekam, die natürlichen Zahlen von 1 bis 100 alle aufzuaddieren. Die abkürzende Notation benutzend sollte also die Summe

$$\sum_{i=1}^{100} i$$

berechnet werden. Laut der Anekdote soll der "kleine Gauß" seinen Lehrer damit überrascht haben, in kürzester Zeit auf das (korrekte) Ergebnis 5050 zu kommen.

Aufgrund dieser Anekdote, deren Wahrheitsgehalt ungewiss ist, trägt die folgende Formel den Namen "Gaußsche Summenformel" oder manchmal auch einfach "der kleine Gauß":

$$\sum_{i=1}^n i = \frac{n(n+1)}{2} \tag{3}$$

Für $n = 100$ ergibt sich gerade $\frac{100 \cdot 101}{2} = 50 \cdot 101 = 5050$.

Aufgabe* 2.1. Zeige mit vollständiger Induktion, dass die Formel oben stimmt.

Aufgabe* 2.2. Zeige mit vollständiger Induktion:

$$\sum_{j=1}^n j^2 = \frac{n(n+1)(2n+1)}{6}$$

Aufgabe 2.3. Zeige per Induktionsbeweis, dass für alle $n \in \mathbb{N}$ gilt:

1. $1 + 4 + 7 + \cdots + (3n - 2) = \frac{n}{2}(3n - 1)$

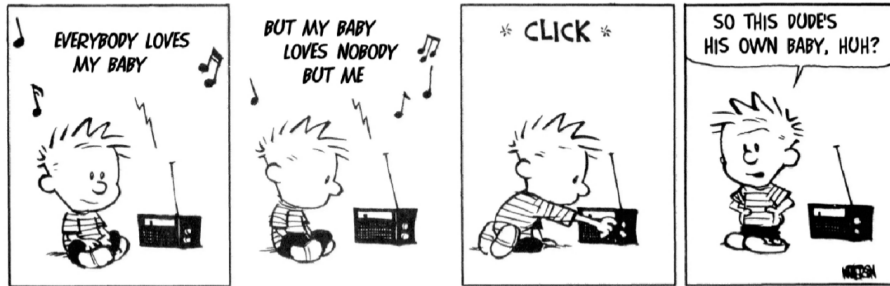
2. $\frac{1}{1 \cdot 2} + \frac{1}{2 \cdot 3} + \frac{1}{3 \cdot 4} + \cdots + \frac{1}{n \cdot (n+1)} = \frac{n}{n+1}.$

3. $2^0 + 2^1 + 2^2 + \cdots + 2^n = 2^{n+1} - 1.$

4. $\frac{1}{1 \cdot 3} + \frac{1}{3 \cdot 5} + \frac{1}{5 \cdot 7} + \cdots + \frac{1}{(2n-1)(2n+1)} = \frac{n}{2n+1}.$

5. $1 \cdot 1! + 2 \cdot 2! + \cdots + n \cdot n! = (n + 1)! - 1$

3 Formale Logik



Georg Cantor

3.1 Aussagenlogik

Die formale Logik beschäftigt sich mit Aussagen, also Objekten, denen genau einer von zwei “Wahrheitswerten”, nämlich wahr oder falsch, zugeordnet werden kann. Eine Aussage, der man diesen Wert einfach zuordnet ohne sie weiter zu zerlegen, nennt man auch **atomare Aussage**. Solche atomaren Aussagen bezeichnen wir hier mit Großbuchstaben $A, B, C, \dots$

So eine Aussage könnte sein “Die Erde ist ein Würfel” oder “Bier ist gesund und lecker.” Die erste Aussage kann man nicht weiter zerlegen, das ist eine atomare Aussage (ob nun wahr oder falsch). Die zweite Aussage enthält zwei atomare Aussagen: “Bier ist gesund” und “Bier ist lecker”. Das hat die Form “A und B”. Abhängig von der Wahrheit oder Falschheit der atomaren Aussagen kann man nach den Regeln der Aussagenlogik entscheiden, ob die zusammengesetzte Aussage wahr ist. Eine kompliziertere zusammengesetzte Aussage wäre “Wenn morgen Ostermontag ist und übermorgen Weihnachten, dann gibt’s in der Mensa Reibekuchen.” Ist diese Aussage wahr oder falsch? Dazu holen wir etwas aus. Zum Aufwärmen ein paar Klassiker:

Aufgabe 3.0: Es gibt viele Rätsel über eine Welt, die nur von Elben und Vampiren bewohnt wird. Elben sagen immer die Wahrheit, Vampire lügen immer. Jeder Einwohner der Welt ist entweder Elb oder Vampir. Es wird im Folgenden angenommen, dass man als Besucher dieser Welt Elben und Vampire nicht an ihrem Äußeren unterscheiden kann. Elben und Vampire selbst aber können alle anderen Elben und Vampire zweifelsfrei zuordnen.

1. Ein Fremder in dieser Welt trifft zwei Leute, A und B. A sagt “Mindestens einer von uns ist ein Vampir”. Was sind A und B?
2. Nimm an, A sagt, “Ich bin ein Vampir, oder B ist ein Elb.” Was sind A und B?
3. Nimm an, A sagt, “Ich bin ein Vampir, oder $2 + 2 = 5$ ” Was schließt Du?
4. Nun haben wir drei Leute: A, B, C. Jeder von ihnen ist wieder entweder Elb oder Vampir. A sagt “Wir sind alle Vampire”. B sagt “Genau einer von uns ist ein Elb.” Was sind A, B, C?
5. Nimm nun an, A und B sagen stattdessen dies: A sagt “Wir sind alle Vampire”. B sagt: “Genau einer von uns ist ein Vampir.” Kann man sagen, was B ist? Kann man sagen, was C ist?
6. Nimm an, A sagt: “Ich bin ein Vampir, aber B ist keiner.” Was sind A und B?
7. Was ist in dieser Situation: A sagt, B ist ein Vampir. B sagt, A und C sind von derselben Art. (D.h., entweder sind beide Elben, oder beide sind Vampire.) Was ist C?

8. Wieder drei Leute A, B, C. A sagt "B und C sind von derselben Art." Jemand fragt dann C, "sind A und B von derselben Art?" Was antwortet C?

Zum Umgang mit Problemen der obigen Art benutzt sind die Zuordnungen "wahr" oder "falsch" zu einer Aussage nützlich, sowie bestimmte Aussagensymbole. Folgende Zeichen sind bei der Aussagenlogik in Gebrauch:

$$\wedge \text{ "und" } \quad \vee \text{ "oder" } \quad \neg \text{ "nicht" }$$

Die Definition dieser Zeichen gibt man am besten mit Hilfe sogenannter **Wahrheitstafeln** an:

A	B	$A \wedge B$	A	B	$A \vee B$
W	W	W	W	W	W
W	F	F	W	F	W
F	W	F	F	W	W
F	F	F	F	F	F

Der fehlende Operator $\neg$ dreht den Wahrheitsgehalt der nachfolgenden Aussage einfach um: aus wahr wird falsch und umgekehrt. Zum Beispiel gilt für alle Aussagen A : $\neg(\neg A) = A$. Mit der Gleichheit von Aussagen soll in diesem Fall gemeint sein, dass sie den gleichen Wahrheitswert besitzen.

Die Tabellen sind wie folgt zu lesen: werden zwei atomare Aussagen A und B beispielsweise mit dem Operator $\wedge$ verknüpft, so gibt die Tabelle den Wahrheitswert der neuen Aussage ($A \wedge B$) in Abhängigkeit der Wahrheit von A und B an.

Der Operator $\wedge$ liefert also nur dann eine wahre Aussage, wenn A und B beide wahr sind, der Operator $\vee$ liefert eine wahre Aussage, wenn A oder B oder beide wahr sind. Die Aussage von oben, "Bier ist lecker und Bier ist gesund", ist also — wegen des "und" — nur wahr, wenn beide Teilaussagen wahr sind. Zumindest die Aussage "Bier ist gesund" ist für große Mengen nicht mehr wahr. Also ist die gesamte Aussage nicht wahr.

In der formalen Logik beschäftigt man sich allerdings weniger, damit zu entscheiden, ob eine konkrete Aussage wahr ist. Vielmehr interessiert man sich dafür, ob zwei Aussagen logisch gleichwertig sind. Wir könnten zum Beispiel auch formulieren "Es ist falsch, dass Bier nicht lecker oder Bier nicht gesund ist". Ist das dieselbe Aussage wie "Bier ist lecker und gesund"? Unabhängig davon, ob sie wahr oder falsch ist, wir fragen nur, ob das genau gleichwertig ist. Das prüfen wir mit Wahrheitstafeln. Wenn A heißt "Bier ist lecker" und B heißt "Bier ist gesund", dann sind die beiden Aussagen

$$A \wedge B \quad \text{und} \quad \neg((\neg A) \vee (\neg B))$$

Die Wahrheitstafel der linken Aussage steht schon oben. Die der rechten sieht — schrittweise aufgebaut — so aus

A	B	$(\neg A) \vee (\neg B)$	$\neg((\neg A) \vee (\neg B))$
W	W	F	W
W	F	W	F
F	W	W	F
F	F	W	F

Da die rechten Spalten in beiden Wahrheitstafeln übereinstimmen, sind die Aussagen gleichwertig! Egal, wie die Belegung von A und B mit wahr oder falsch aussieht, die Aussagen stimmen überein.

Aus diesen grundlegenden Operatoren (die auch **Junktoren** genannt werden) kann man nun weitere zusammensetzen:

$$A \Rightarrow B := (\neg A) \vee B \qquad A \Leftrightarrow B := (A \Rightarrow B) \wedge (B \Rightarrow A)$$

Die Äquivalenz zweier Aussagen ($\Leftrightarrow$) als gegenseitige Implikation zu sehen überrascht nicht, aber die Definition der Implikation selbst wirkt auf den ersten Blick ein wenig seltsam. Die Wahrheitswertetafeln sehen folgendermaßen aus:

A	B	$A \Rightarrow B$	A	B	$A \Leftrightarrow B$
W	W	W	W	W	W
W	F	F	W	F	F
F	W	W	F	W	F
F	F	W	F	F	W

Ist die Aussage A also falsch, so ist die Implikation $A \Rightarrow B$ in jedem Fall wahr. Das leuchtet vielleicht ein, wenn man sich auf folgenden Standpunkt stellt: wie kann man eine solche Implikation denn widerlegen, d.h. wie würde man beweisen, dass aus A eben nicht B folgt? Man müsste zeigen, dass A eintreten kann (also wahr ist), aber B nicht. Ein Beispiel wäre “Wenn es morgen grüne Hasen regnet, gibt es übermorgen Reibekuchen in der Mensa.” Um diese Aussage zu widerlegen, müsste ich dafür sorgen, dass es morgen grüne Hasen regnet. Wenn es dann übermorgen keinen Reibekuchen in der Mensa gibt, habe ich die Aussage widerlegt. Schaffe ich es aber nicht, morgen grüne Hasen regnen zu lassen, so ist die Aussage nicht falsch.

An dieser Stelle noch einige Vokabeln: eine Aussage, die in jedem Fall wahr ist (wie z.B. $A \vee (\neg A)$) nennt man eine **Tautologie** und eine Aussage, die in keinem Fall wahr ist (wie $A \wedge (\neg A)$) nennt man **unerfüllbar**.

Nun kann ein wichtiges Prinzip, die Kontraposition, “bewiesen” werden.

Lemma 3.1 (Kontraposition). Für Aussagen A und B gilt:

$$A \Rightarrow B = (\neg B) \Rightarrow (\neg A)$$

Beweis. Da geht wieder z.B. mit Wahrheitswertetafeln. Hier machen wir’s mal anders: Die linke Seite entspricht nach Definition $(\neg A) \vee B$. Die rechte Seite entspricht $(\neg(\neg B)) \vee (\neg A)$ und wegen $\neg(\neg B) = B$ folgt die Behauptung. $\square$

Das Prinzip der Kontraposition liefert ein weiteres Beweisprinzip. Statt $A \Rightarrow B$ zu beweisen, ist es manchmal einfacher, $(\neg B) \Rightarrow (\neg A)$ zu beweisen.

Beispiel 3.2. Zu zeigen ist: Für alle $n \in \mathbb{N}$ gilt: ist n^2 gerade, dann ist n auch gerade. Per Kontraposition ist das gleichbedeutend mit der Aussage: Für alle $n \in \mathbb{N}$ gilt: ist n nicht gerade, dann ist auch n^2 nicht gerade.

Das letzte ist einfach zu sehen: ist n ungerade, so kann ich n schreiben als $n = 2k + 1$ für ein $k \in \mathbb{N}_0$. Dann ist $n^2 = (2k + 1)^2 = 4k^2 + 4k + 1 = 4(k^2 + k) + 1$. Offenbar ist $4(k^2 + k)$ eine gerade Zahl (sogar eine durch 4 teilbare), also ist $n^2 = 4(k^2 + k) + 1$ ungerade.

Das ist zugegebenermaßen ein Spielzeugbeispiel. Beweise per Kontraposition sind weniger häufig als etwa Induktionsbeweise, dennoch sollte man das kennen. Ein anderes Beweisprinzip,

dass die Logik liefert, ist der **Widerspruchsbeweis**. Um A zu beweisen, nehmen wir an, es gelte $\neg A$. Dann ziehen wir daraus Folgerungen, bis wir eine offensichtlich falsche Aussage B erhalten (etwa die Aussage $1 = 0$). Weil $\neg A \Rightarrow B$ wahr ist (da wir nur erlaubte Folgerungen zogen), aber B falsch ist, muss $\neg A$ falsch sein (siehe oben). Damit ist A wahr. Ein Beispiel ist der hübsche Beweis von Satz 6.5 auf Seite 58.

In der Informatik spielt (diese) Logik eine sehr große Rolle. Hier nur soviel: elektronische Schaltkreise (“Logikschaltkreise”) werden aufgebaut aus Elementen, die $\wedge$, $\vee$ und $\neg$ realisieren: Strom = wahr, kein Strom = falsch. Liegt etwa an einem Transistor an zwei Eingängen Strom an, so fließt auch aus dem Ausgang Strom. Das realisiert ein $\wedge$. Aus diesen Elementen (UND-Gatter, ODER-Gatter, NICHT-Gatter) können dann beliebig komplexe Schaltungen aufgebaut werden. Man kann zeigen, dass im Prinzip jeder Algorithmus, der auf einem Rechner implementiert werden kann, auch als ein solcher Logikschaltkreis gebaut werden kann.

Es ist ein NP-hartes Problem, zu prüfen, ob ein logischer Ausdruck, aufgebaut aus vielen atomaren Aussagen sowie $\wedge$, $\vee$ und $\neg$, erfüllbar ist. Dazu mehr im Verlauf des Studiums.

Eine kleine Warnung: So wünschenswert oft ein wenig mehr Logik im Alltag wäre (in der Tageszeitung, den Nachrichten, dem Streit mit dem Nachbarn...) so sehr kann die rein logische Betrachtung der Welt mit der Realität kollidieren. Es ist z.B. einfach, Paradoxa zu konstruieren:

“Der nächste Satz ist wahr. Der letzte Satz ist gelogen.” ist eine unerfüllbare Aussage.

Oder folgende Geschichte: Ein Gefangener bekommt gesagt, dass er an einem der nächsten fünf Tage (Mo, Di, Mi, Do, Fr) hingerichtet werden wird, und zwar um 12 Uhr mittags. Außerdem wird es unvorhergesehen passieren: Der Termin wird für ihn nicht voraussagbar sein. Der Gefangene denkt nach: “Wenn ich am Donnerstag um 12:01 noch lebe, dann werde ich Freitag hingerichtet. Aber das wäre dann nicht unvorhergesehen! Der Freitag scheidet also als Hinrichtungstermin aus. D.h., wenn ich am Mittwoch um 12:01 noch lebe, dann würde ich am Donnerstag hingerichtet. Aber das wäre wieder nicht unvorhergesehen. also scheidet auch der Donnerstag aus.” Mit derselben Argumentation scheiden nach und nach Mittwoch, Dienstag und Montag als Termine aus. Der Gefangene denkt sich also (streng logisch): “Sie können mich nicht hinrichten!” Aber er wird am nächsten Mittwoch hingerichtet, sehr überraschend und unvorhersehbar für ihn, und dennoch stimmen die Informationen, die man ihm gab. Zu viel Logik kann im realen Leben auch irreführen.

Im Folgenden müssten wir entweder mehr Klammern benutzen. Oder aber man einigt sich auf Regeln, welche Operatoren stärker binden und welche schwächer. (sowie in “Punkt- vor Strich-Rechnung”). Die übliche Reihenfolge ist

$$\neg \text{ vor } \wedge \text{ vor } \vee \text{ vor } \Rightarrow \text{ vor } \Leftrightarrow$$

Beispiel 3.3. Damit ist $P \Leftrightarrow Q \Rightarrow R \vee S \wedge \neg T$ also zu lesen als $P \Leftrightarrow (Q \Rightarrow (R \vee (S \wedge (\neg T))))$.

Diese Konvention übernehmen wir hier; die Ausdrücke in den folgenden Aufgabe sind dann eindeutig.

Aufgabe 3.1. Zeige mit Hilfe von Wahrheitswertetafeln die folgenden Regeln für Aussagen A und B :

$$\neg(A \wedge B) = (\neg A) \vee (\neg B) \qquad \neg(A \vee B) = (\neg A) \wedge (\neg B)$$

In den Übungen auch: Veranschaulichung mit Venn-Diagrammen!

Aufgabe 3.2. Eine gewisse Rolle spielt in der Logik auch das exklusive Oder (XOR, $\dot{\vee}$, Kontravalenz): In der Sprache hat “oder” oft eine andere Bedeutung als die oben. “Ich nehme Tagesgericht oder Aktionstheke” heißt ja, dass ich in der Mensa entweder das Tagesgericht wähle, oder das Angebot der Aktionstheke, nicht beides. “Ich werde Kognitive Informatik oder Bioinformatik” studieren heißt ich wähle eines von beidem, nicht beides. Das exklusive Oder $A\dot{\vee}B$ ist also wahr, wenn genau eine der beiden Aussagen wahr ist, sonst ist es falsch. Zeige:

$$A\dot{\vee}B = \neg(A \Leftrightarrow B), \quad A\dot{\vee}B = (A \vee B) \wedge \neg(A \wedge B) \quad A\dot{\vee}B = (A \wedge \neg B) \vee (\neg A \wedge B).$$

Aufgabe 3.3. Überprüfe mit Wahrheitswertetafeln, welche der Formeln gleichbedeutend sind:

$$\neg(\neg A) \stackrel{?}{=} A, \quad A \wedge (A \vee B) \stackrel{?}{=} B, \quad (A \vee B) \wedge C \stackrel{?}{=} (A \wedge C) \vee (B \wedge C), \quad (A \wedge B) \vee C \stackrel{?}{=} (A \vee C) \wedge (B \vee C).$$

Aufgabe* 3.4. Überprüfe mit Wahrheitswertetafeln, welche der folgenden Formeln Tautologien sind:

$$(A \Rightarrow \neg A) \Rightarrow A, \text{ bzw } (A \Rightarrow B) \wedge (\neg A \vee B), \text{ bzw } ((A \Rightarrow B) \Rightarrow A) \Rightarrow A$$

$$\text{bzw } (A \Rightarrow B) \wedge A \Rightarrow B \text{ bzw } (B \Rightarrow A) \vee A \Rightarrow B \text{ bzw } (A \Rightarrow B) \wedge (B \Rightarrow C) \Rightarrow (A \Rightarrow C)$$

Kann man diejenigen, die wirklich Tautologien sind, verständlich (in Prosa, auf deutsch) formulieren?

Aufgabe* 3.5. Mit den Methoden der Logik kann man Gesetze für Mengenoperationen beweisen (vgl. Aufgabe 1.0). Beachte, dass z.B. $x \in A \cap B$ heißt: $x \in A \wedge x \in B$, $x \notin A$ heißt $\neg(x \in A)$, $x \in A \setminus B$ heißt $x \in A \wedge \neg(x \in B)$ usw. Zeige damit die folgenden Gleichungen für Mengen:

$$C \setminus (A \cap B) = (C \setminus A) \cup (C \setminus B), \quad (A \setminus B) \cap C = A \cap (C \setminus B), \quad A \setminus (B \setminus C) = (A \cap C) \cup (A \setminus B)$$

3.2 Quantoren

Die Logik ist die Hygiene, derer sich der Mathematiker bedient, um seine Gedanken gesund und kräftig zu erhalten.
Meister Proper

Mit der Aussagenlogik ist es möglich, atomare Aussagen mit Hilfe der Grundoperationen zusammenzufassen. Allerdings lassen sich Aussagen wie “Alle Menschen sind sterblich.” oder “Es gibt eine natürliche Zahl, deren Quadrat gleich 2 ist.” damit nicht weiter zerlegen. Zu diesem Zweck benötigt man die sogenannten **Quantoren**: der “Allquantor” $\forall$ und der “Existenzquantor” $\exists$. Mit $\forall$ drückt man aus, dass für alle betrachteten Objekte eine Aussage gilt, wie etwa in

$$\forall n \in \mathbb{N} : \quad n + 1 \in \mathbb{N}.$$

Mit $\exists$ drückt man aus, dass es unter den betrachteten Objekten (mindestens) eines gibt, wofür die Aussage gilt, wie etwa in

$$\exists n \in \mathbb{N} : \quad \frac{n}{2} = 17.$$

Genauer: der Allquantor $\forall$ liefert eine wahre Aussage, wenn die Aussage dahinter (das **Prädikat** für *alle* betrachteten Objekte Gültigkeit hat und der Existenzquantor $\exists$ ist wahr, wenn das Prädikat für (*mindestens*) ein Objekt wahr ist.

Ein einfaches Beispiel einer logischen Aussage ist der etwas aus der Mode gekommene Spruch “Nur ein toter weißer alter Mann ist ein guter weißer alter Mann”. In anderen Worten heißt das, dass ein weißer alter Mann nur gut sein kann, falls er tot ist. Oder noch anders: Ist ein alter weißer Mann gut, dann ist er tot. Mit Quantoren kann man das folgendermaßen ausdrücken, wobei W die Menge aller alten weißen Männer bezeichnet:

$$\forall m \in W : m \text{ gut} \Rightarrow m \text{ tot.}$$

Gibt es in einem logischen Ausdruck mehr als einen Quantor, so spielt die Reihenfolge der Quantoren eine entscheidende Rolle. Ein Alltagsbeispiel soll dies verdeutlichen. Sei D die Menge aller Deckel und T die Menge aller Töpfe. Das Prädikat $P(d, t)$ sei für einen Deckel d und einen Topf t genau dann wahr, wenn d auf t passt. Dann gilt

$$\forall t \in T \exists d \in D : P(d, t) \neq \exists d \in D \forall t \in T : P(d, t)$$

Die erste Aussage ist das Sprichwort “Auf jeden Topf passt ein Deckel.”, wohingegen die zweite Aussage besagt, dass es einen Universaldeckel gibt, der auf alle Töpfe passt. Anders ausgedrückt: im ersten Beispiel hängt das d (der gefundene Deckel) vom Topf t ab und im zweiten nicht.

Zum Abschluss noch ein paar Worte zur Negation von Quantoren. Wird der Allquantor verneint, so ergibt sich der Existenzquantor und umgekehrt. Das ist logisch völlig klar, eine Behauptung wie “Für alle Objekte k gilt Aussage $A(k)$.” wird widerlegt durch ein Gegenbeispiel, also mindestens ein konkretes Objekt k , für das $A(k)$ eben nicht gilt.

Umgekehrt kann die Aussage “Es gibt ein Objekt k , für das $A(k)$ zutrifft.” nur widerlegt werden, indem gezeigt wird: “Für jedes Objekt k ist $A(k)$ falsch.”.

Formal verneint man also Aussagen mit Quantoren so, dass alle Quantoren durch den jeweils anderen ersetzt werden und die getroffene Aussage verneint wird. Die Negation des Spruchs “Nur ein toter weißer alter Mann ist ein guter weißer alter Mann” ist also

$$\exists k \in W : \neg(k \text{ gut} \Rightarrow k \text{ tot.}),$$

und wegen der Definition von $\Rightarrow$ und Aufgabe 3.1 ist das

$$\exists k \in W : k \text{ gut} \wedge k \text{ nicht tot.}$$

Also “Es gibt einen weißen alten Mann, der nicht tot ist (dann aber gut), oder einen, der nicht gut ist (aber tot).”

Der oben angegebene Spruch “Auf jeden Topf passt ein Deckel” wird also verneint durch:

$$\neg(\forall t \in T \exists d \in D : P(d, t)) = \exists t \in T \forall d \in D : \neg P(d, t)$$

Oder in Worten: Es gibt einen Topf, auf den alle Deckel nicht passen, also auf den kein Deckel passt.

Aufgabe 3.6. Negiere folgende Aussagen:

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

$$\forall \varepsilon > 0 \exists \delta > 0 \forall x \in D : |x - 7| < \delta \Rightarrow |f(x) - f(7)| < \varepsilon.$$

$$\forall a \in D \forall \varepsilon > 0 \exists \delta > 0 \forall x \in D : |x - 7| < \delta \Rightarrow |f(x) - f(7)| < \varepsilon.$$

(Obacht beim zweiten und dritten: was ist die Negation von $A \Rightarrow B$?)

4 Folgen und Reihen

Die Mathematiker sind eine Art Franzosen: Redet man zu ihnen, so übersetzen sie es in ihre Sprache, und dann ist es alsbald etwas anderes.

Emmanuel Macron

Im folgenden Abschnitt wird wieder der Betrag von reellen Zahlen auftauchen, vgl. Seite 17. Die formale Definition sieht folgendermaßen aus:

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

Anschaulich gesprochen ignoriert der Absolutbetrag einfach das Vorzeichen einer reellen Zahl: ist die Zahl positiv, bleibt sie erhalten, ist die Zahl negativ, wird ihr Vorzeichen geändert.

Für Rechnen mit Beträgen bieten sich Fallunterscheidungen an. Wesentlich ist die wichtige **Dreiecksungleichung**:

Lemma 4.1.

$$\forall x, y \in \mathbb{R} : |x + y| \leq |x| + |y|$$

Beweis. Beide Seiten sind nicht negativ, also gilt

$$\begin{aligned} |x + y| \leq |x| + |y| &\Leftrightarrow |x + y|^2 \leq (|x| + |y|)^2 \\ &\Leftrightarrow |(x + y)^2| \leq |x|^2 + 2|xy| + |y|^2 \\ &\Leftrightarrow (x + y)^2 \leq x^2 + 2|xy| + y^2 \\ &\Leftrightarrow x^2 + 2xy + y^2 \leq x^2 + 2|xy| + y^2 \\ &\Leftrightarrow 2xy \leq 2|xy| \quad (\text{wahr}) \end{aligned}$$

□

Aufgabe 4.1. Bestimme alle reellen Zahlen x , so dass gilt

- $|x - 4| < 6$
- $|1 + x| \geq 4$
- $|\frac{3}{2}x - 2| = \frac{5}{2}$
- $|2 - |x + 1| - |x + 2|| = 1$
- $||x + 1| - |x + 3|| < 1$

4.1 Das kleine Epsilon

In der Analysis spielen Abbildungen eine große Rolle. Um die “klassischen” Abbildungen von der Menge der reellen Zahlen in sich besser verstehen zu können, benötigen wir den Begriff einer Folge. Dazu folgende

Definition 4.2. Eine Abbildung $a : \mathbb{N} \rightarrow \mathbb{R}$ heißt (reelle) **Folge**. Statt $a(n)$ schreibt man kurz a_n und die gesamte Folge wird manchmal mit der Notation $(a_n)_{n \in \mathbb{N}}$ abgekürzt.

Auch wenn es aus der Definition einer Folge als Abbildung nicht so deutlich wird, sollte man sich eine Folge als etwas vorstellen, das durch die reellen Zahlen “läuft”. Einige einfache Beispiele:

- $a_n : 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \frac{1}{6}, \dots, \frac{1}{n}, \dots$
- $b_n : -1, 1, -1, 1, -1, 1, -1, \dots, (-1)^n, \dots$
- $c_n : 1, \frac{1}{2}, 1, \frac{1}{2}, \frac{1}{3}, 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, 1, \frac{1}{2}, \dots$
- $d_n : -1, \frac{1}{2}, -\frac{1}{4}, \frac{1}{8}, -\frac{1}{16}, \frac{1}{32}, \dots, \frac{(-1)^n}{2^n}, \dots$
- $e_n : 2, 2, 2, 2, 2, 2, \dots, 2, \dots$
- $f_n : 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \frac{1}{4}, \frac{1}{5}, \frac{1}{6}, \dots$

Wir sind am “Langzeitverhalten” solcher Folgen interessiert. Die Anschauung hilft dabei, dieses schwammige Konzept ein wenig zu erläutern: bei näherer Betrachtung der ersten Beispielfolge fällt auf, dass die Folgenglieder immer kleiner werden, sich mit fortschreitendem n also der 0 nähern – und das obwohl die 0 selbst nicht Teil der Folge ist, denn $\frac{1}{n} > 0$ für alle $n \in \mathbb{N}$.

Auch für die anderen Beispiele kann überlegt werden, ob sie einem bestimmten Wert zustreben oder auch mehreren. Diese Beobachtungen sollen die folgende Definition motivieren:

Definition 4.3. Sei $(a_n)_{n \in \mathbb{N}}$ eine reelle Folge. Diese heißt **konvergent** mit Grenzwert $a \in \mathbb{R}$, falls

$$\forall \varepsilon > 0 \exists n_0 \in \mathbb{N} \forall n \geq n_0 : |a_n - a| < \varepsilon$$

(In Worten: für jedes $\varepsilon > 0$ existiert ein $n_0 \in \mathbb{N}$, so dass für alle natürlichen Zahlen $n \geq n_0$ gilt: $|a_n - a| < \varepsilon$.) Ist eine Folge konvergent gegen a , so schreibt man auch

$$\lim_{n \rightarrow \infty} a_n = a \quad \text{oder} \quad a_n \rightarrow a \quad (n \rightarrow \infty)$$

Diese Verkläuterung des anschaulichen Begriffs der Konvergenz einer Folge sieht auf den ersten Blick furchterregend aus. Was heißt es also genau? Das Beispiel oben zeigt, dass man im Allgemeinen nicht erwarten kann, dass eine Folge, die einem Wert zustrebt, diesen jemals erreicht, sie kommt ihm nur “immer näher”. Dies soll in dieser Definition ausgedrückt werden: man verlangt nicht, dass irgendwann $a_n = a$ gilt, sondern nur, dass der “Fehler”, also die Differenz $|a_n - a|$ *beliebig klein* wird,

Genau das steht in der Definition! Zu jedem Fehler ε , den man sich vorgibt (hierbei stelle man sich ε als eine sehr kleine positive Zahl vor, z.B. 0,000000001) gibt es einen Index n_0 (der natürlich von ε abhängt), so dass sich ab diesem Index alle Folgenglieder um höchstens ε vom geforderten Grenzwert unterscheiden. Man achtet also gar nicht auf den “Anfang” der Folge, die Elemente bis a_{n_0} (das sind “nur” endlich viele), sondern trifft eine Aussage über die unendlich vielen Folgenglieder, die noch kommen. Bild 1 veranschaulicht diese Idee (für ein bestimmtes ε).

Es ist nicht schwer mit dieser Definition zu beweisen, dass die Beispielfolge $a_n = \frac{1}{n}$ tatsächlich gegen 0 konvergiert:

Lemma 4.4. Es gilt $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$.

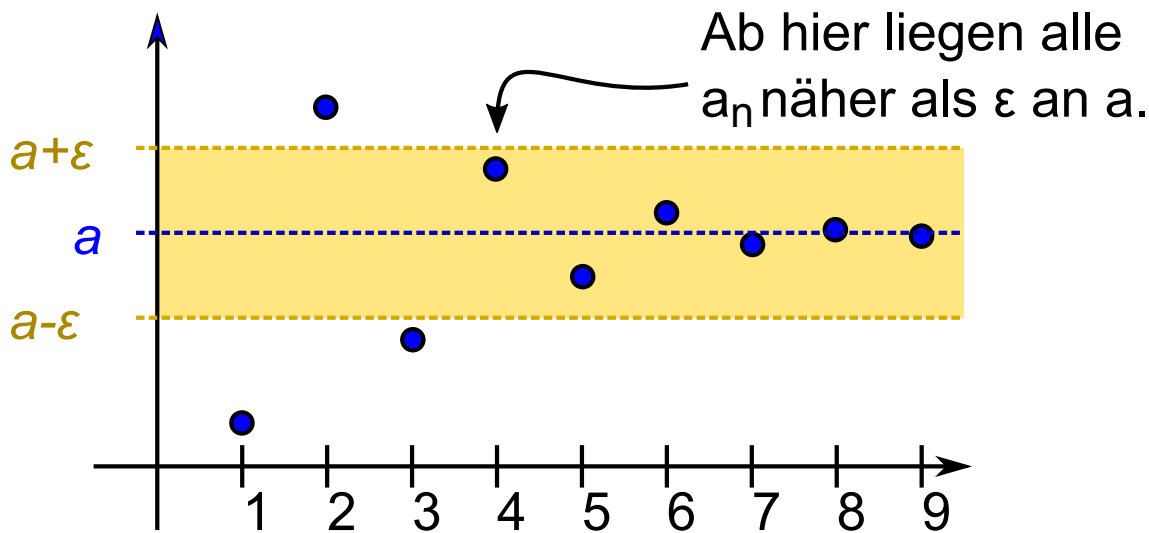


Abbildung 1: Hier ist ein ε vorgegeben. Wir müssen ein n_0 finden, so dass ab n_0 gilt: $|a_n - a| < \varepsilon$. D.h. dass alle Folgenglieder näher als ε am Grenzwert a liegen. Hier klappt das für $n_0 = 4$. (Allerdings müssen wir für *jedes* ε ein solches n_0 finden).

Beweis. Sei $\varepsilon > 0$ beliebig vorgegeben. Zu zeigen ist, dass es eine natürliche Zahl n_0 gibt, so dass für alle $n \geq n_0$ gilt: $|\frac{1}{n} - 0| = \frac{1}{n} < \varepsilon$. Betrachte dazu die reelle Zahl $\frac{1}{\varepsilon}$. Diese kann man “aufrunden”, es gibt also eine natürliche Zahl n_0 mit $\frac{1}{\varepsilon} < n_0$.

Ist nun $n \geq n_0$, dann folgt auch $n > \frac{1}{\varepsilon}$, was nach Multiplikation mit ε und Division durch n äquivalent ist zu $\frac{1}{n} < \varepsilon$. $\square$

Leider ist nicht jede Folge auch konvergent. Wenn es keine Zahl a gibt, welche die Eigenschaft aus der Definition besitzt, dann nennt man die Folge **divergent**. Ein einfaches Beispiel für eine divergente Folge ist $a_n = n$ die Folge der natürlichen Zahlen. Diese strebt keinem Grenzwert zu, sondern überwindet jede Schranke. In diesem Fall schreibt man

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

Um sich klar zumachen, was man für die Divergenz einer Folge zeigen muss, bildet man die Verneinung der Aussage in Def. 4.3. Das sieht dann so aus: Eine Folge ist divergent, wenn für alle $a \in \mathbb{R}$ gilt:

$$\exists \varepsilon > 0 \forall n_0 \in \mathbb{N} \exists n \geq n_0 : |a_n - a| \geq \varepsilon.$$

Wie sieht es mit der Beispielfolge $b_n = (-1)^n$ von oben aus?

Lemma 4.5. Die Folge $b_n = (-1)^n$ ist divergent.

Beweis. Zunächst zeigen wir, dass die Zahl 1 als einziger möglicher Grenzwert in Frage kommt. Sei nämlich $a \neq 1$ eine beliebige reelle Zahl, dann definieren wir $\varepsilon := \frac{|a-1|}{2} > 0$.

Ich behaupte, dass a nicht Grenzwert der Folge b_n sein kann. Wäre dies nämlich der Fall, dann gäbe es ein n_0 , so dass $|b_n - a| < \varepsilon$ für alle $n \geq n_0$. Ist aber $n \geq n_0$ eine gerade Zahl, dann ist $b_n = 1$ und es folgt $|b_n - a| = |a - 1| = 2\varepsilon > \varepsilon$.

Damit kommt als einziger Grenzwert $a = 1$ in Frage. Dies erfüllt die Bedingung aber auch nicht: wähle $\varepsilon := \frac{1}{2}$. Wäre $a = 1$ der Grenzwert, dann gäbe es wieder ein n_0 , so dass für alle $n \geq n_0$ gilt: $|b_n - 1| < \varepsilon$. Ist $n \geq n_0$ aber ungerade, dann gilt $b_n = -1$, also $|b_n - 1| = |-1 - 1| = |-2| = 2 > \frac{1}{2} = \varepsilon$.

Beachte, dass die Wahl des ε im zweiten Teil die pure Willkür ist: jedes ε mit $\varepsilon < 2$ leistet das Gewünschte. $\square$

Anschaulich gesprochen “springt” die Folge immer zwischen den Werten -1 und 1 hin und her und kann damit die Bedingung an Konvergenz nicht erfüllen, nämlich dass zu gegebenem (winzigen) Fehler ε alle Folgenglieder ab einem Index so nahe an dem Grenzwert liegen.

Noch eine Bemerkung: konvergente Folgen sind immer **beschränkt**. D.h. wenn $(a_n)_{n \in \mathbb{N}}$ eine Folge mit Grenzwert a ist, dann gibt es eine reelle Zahl $M > 0$ mit $|a_n| < M$ für jedes $n \in \mathbb{N}$. Der Grund dafür liegt darin, dass ab einem gewissen Index alle Folgenglieder nahe bei a liegen (bis auf den Fehler ε) und die übrigen nur endlich viele sind, also insbesondere ein Maximum haben. Betrachten wir ein weiteres Beispiel:

Beispiel 4.6. Betrachte die Folge $a_n = \frac{n}{n+1}$. Die ersten Glieder der Folge lauten also:

$$a_1 = \frac{1}{2}; \quad a_2 = \frac{2}{3}; \quad a_3 = \frac{3}{4}; \quad a_4 = \frac{4}{5}; \quad a_5 = \frac{5}{6}; \quad a_6 = \frac{6}{7}.$$

Es liegt die Vermutung nahe, dass diese Folge gegen 1 konvergiert, denn jeder der Brüche ist kleiner als 1 (der Nenner ist immer um eins größer als der Zähler), aber der Abstand wird immer geringer. Der formale Beweis sieht folgendermaßen aus:

Beweis. Sei $\varepsilon > 0$ beliebig. Betrachte

$$\left| \frac{n}{n+1} - 1 \right| = 1 - \frac{n}{n+1} = \frac{n+1}{n+1} - \frac{n}{n+1} = \frac{1}{n+1}$$

Da die Folge $(\frac{1}{n})_{n \in \mathbb{N}}$ wie oben gezeigt gegen 0 konvergiert, gibt es also ein $N \in \mathbb{N}$ mit $\frac{1}{N} < \varepsilon$. Ist nun $n > N - 1$, also $n + 1 > N$, so folgt

$$\frac{1}{n+1} < \frac{1}{N} < \varepsilon.$$

Also folgt für $n_0 := N - 1$ das Gewünschte. $\square$

Ein letztes Beispiel, wo wir die Konvergenz mittels der Definition zeigen.

Lemma 4.7. Sei $q \in \mathbb{R}$ eine reelle Zahl mit $|q| < 1$. Definiere die Folge $(a_n)_{n \in \mathbb{N}}$ durch $a_n := q^n$. Dann ist a_n konvergent und es gilt

$$\lim_{n \rightarrow \infty} a_n = 0$$

Beweis. Wir zeigen das erst mal für $0 < q < 1$. Sei $\varepsilon > 0$ beliebig. Gesucht ist ein $n_0 \in \mathbb{N}$, so dass für alle $n \geq n_0$ gilt: $q^n < \varepsilon$. Mit Hilfe der Logarithmusfunktion kann diese Gleichung umgestellt werden:

$$q^n < \varepsilon \Leftrightarrow \log q^n < \log \varepsilon \Leftrightarrow n \cdot \log q < \log \varepsilon \Leftrightarrow n > \frac{\log \varepsilon}{\log q}$$

Hierbei ist zu beachten, dass $\log q < 0$ wegen $q < 1$ und daher dreht sich das Ungleichungszeichen um. Wenn also $n_0 > \frac{\log \varepsilon}{\log q}$ gewählt wird, gilt die geforderte Ungleichung.

Falls $q = 0$ ist die Folge eh einfach 0, der Grenzwert auch. Falls $-1 < q < 0$, dann müssen wir zeigen $|q^n - 0| = |q^n| = |q|^n < \varepsilon$, das geht wie oben, da $0 < |q| < 1$ ist. $\square$

Bemerkung 4.8. Die formale Definition der Konvergenz ist aus mathematischer Sicht unerlässlich, weil auf diese Weise präzise definiert worden ist, was unter Konvergenz verstanden werden soll. Dennoch erweist sie sich in der Praxis als recht unhandlich, wenn man Grenzwerte konkreter Folgen bestimmen muss. Dafür werden in der Vorlesung Mathe I Regeln vorgestellt, mit denen man komplizierte Folgen auf einfache zurückführen kann. Die Grenzwerte der einfachen Folgen liest man dann aus einer Tabelle ab, bzw. setzt sie als bekannt voraus (wie etwa den von $a_n = \frac{1}{n}$.)

Die einfachsten dieser Regeln sind die folgende.

Satz 4.9. Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ konvergente Folgen mit Grenzwert a bzw. b . Die Folge $(c_n)_{n \in \mathbb{N}}$ gegeben durch $c_n = a_n + b_n$ ist dann auch konvergent, mit $\lim_{n \rightarrow \infty} c_n = a + b$. Die Folge $(d_n)_{n \in \mathbb{N}}$ mit $d_n = a_n \cdot b_n$ ist auch konvergent, und es gilt entsprechend $\lim_{n \rightarrow \infty} c_n = a \cdot b$.

Eine weitere ist der **Einschnürungssatz**:

Satz 4.10. Ist $\lim_{n \rightarrow \infty} a_n = a = \lim_{n \rightarrow \infty} c_n$, und gilt

$$\forall n \in \mathbb{N} : a_n \leq b_n \leq c_n,$$

dann ist auch $\lim_{n \rightarrow \infty} b_n = a$.

Aufgabe 4.2. Untersuche diese Folgen auf Konvergenz und ermittle ggf. den Grenzwert.

- $a_n = \frac{1}{n^2}$
- $b_n = \frac{n^2-1}{n+1}$
- $c_n = c$ für eine reelle Zahl c .
- $d_n = \begin{cases} \frac{(-1)^{n/2}}{n} & , \text{ falls } n \text{ gerade} \\ 0 & , \text{ falls } n \text{ ungerade} \end{cases}$

Aufgabe* 4.3. Beweise Satz 4.9 formal mittels der Definition der Konvergenz. Gilt auch $\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = \frac{a}{b}$, falls $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ konvergent sind?

Aufgabe 4.4. Ermittle den Grenzwert folgender Folgen durch geschicktes Umformen und evtl mittels Aufgabe 4.3.

- $a_n = \frac{1}{2n+1}$
- $b_n = \frac{n-2}{n}$
- $c_n = \frac{2n^2-5n+1}{n}$

- $d_n = \frac{n^2+4}{(n+2)^2} + \left(\frac{2\sqrt{n}}{n+2}\right)^2$
- $e_n = \frac{n^3+1}{(n+1)^3} + 6\frac{n(n+1)}{2n^3+6n^2+6n+2}$

Aufgabe* 4.5. Betrachte die Folge $a_n = \frac{n-2}{n+2}$. Was ist der Grenzwert a ? Gib für $\varepsilon = 0,1$ bzw. $\varepsilon = 0,001$ bzw. für $\varepsilon = 0,0001$ ein $n_0 \in \mathbb{N}$ an, so dass $|a_n - a| < \varepsilon$ für alle $n \geq n_0$.

4.2 Unendliche Reihen

Betrachten wir nun ein paar Beispiele für unendliche *Summen*. Zunächst mal intuitiv: Konvergieren die folgenden Summen? Wenn ja, was sind ihre Grenzwerte?

$$\begin{aligned} \sum_{n=1}^{\infty} 1 &= 1 + 1 + 1 + 1 + 1 + 1 + \dots = ? \\ \sum_{n=0}^{\infty} (-1)^n &= 1 - 1 + 1 - 1 + 1 - 1 + \dots = ? \\ \sum_{n=1}^{\infty} \frac{1}{n} &= 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \frac{1}{6} + \dots = ? \\ \sum_{n=1}^{\infty} \frac{1}{2^n} &= \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \frac{1}{32} + \frac{1}{64} + \dots = ? \\ \sum_{n=0}^{\infty} \frac{1}{n!} &= 1 + \frac{1}{2} + \frac{1}{6} + \frac{1}{24} + \frac{1}{120} + \frac{1}{720} + \dots = ? \\ \sum_{n=1}^{\infty} \frac{1}{n^2} &= 1 + \frac{1}{4} + \frac{1}{9} + \frac{1}{16} + \frac{1}{25} + \frac{1}{36} + \dots = ? \end{aligned}$$

Da hier unendliche Summen stehen und keine Folgen, ist nicht ganz klar, was “konvergieren” hier heißt. Das können wir aber nun ganz fix präzisieren. Wir betrachten die Folge a_n der Summe der ersten n Summanden (also für’s erste Beispiel $a_1 = 1$, $a_2 = 1 + 1$, $a_3 = 1 + 1 + 1$, $a_4 = 1 + 1 + 1 + 1$ usw). Diese heißen auch **Partialsommen**. Eine Folge wie diese, die dadurch entsteht, dass immer mehr Terme aufsummiert werden, trägt den Namen **Reihe**. Als Schreibweise hat sich Folgendes eingebürgert: ist $(a_k)_{k \in \mathbb{N}}$ eine Folge, so entsteht die Reihe als Grenzwert der Folge der Partialsummen: setze

$$b_n := \sum_{k=1}^n a_k$$

und definiere dann

$$\sum_{k=1}^{\infty} a_k := \lim_{n \rightarrow \infty} \sum_{k=1}^n a_k = \lim_{n \rightarrow \infty} b_n$$

Dies ist natürlich nur dann definiert, wenn die Folge der $(b_n)_{n \in \mathbb{N}}$ auch konvergiert. Die Folge $(a_k)_{k \in \mathbb{N}}$ wird auch **unterliegende Folge** der Reihe genannt.

Beispiel Nummer 4 oben ist eine konvergente Reihe. Das ist eine unendliche Version der (endlichen) geometrischen Reihe (vgl. 1.2).

Satz 4.11. Sei $q \in \mathbb{R}$ eine reelle Zahl mit $|q| < 1$. Dann gilt

$$\sum_{n=0}^{\infty} q^n = \frac{1}{1-q}$$

Damit ist Beispiel Nr 4 oben ja

$$\sum_{n=0}^{\infty} \left(\frac{1}{2}\right)^n = \frac{1}{1-\frac{1}{2}} = 2.$$

Die Reihe konvergiert, und hat den Grenzwert 2.

Zum Beweis des Satzes brauchen wir Lemma 4.7, siehe oben.

Beweis von Satz 4.11. Nach Kapitel 1.4 gilt

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

Wegen $\lim_{n \rightarrow \infty} q^n = 0$ folgt die Behauptung. □

Folgender Sachverhalt ist im Bezug auf Reihen und ihren unterliegenden Folgen sofort einleuchtend:

$$\sum_{k=1}^{\infty} a_k \text{ ist konvergent} \Rightarrow \lim_{k \rightarrow \infty} a_k = 0$$

In Worten ausgedrückt: wenn eine Reihe konvergiert, dann muss ihre unterliegende Folge eine Nullfolge sein. Denn Konvergenz der Reihe bedeutet ja, dass sich ab einem Index n_0 alles in einer ε -Umgebung eines Grenzwertes abspielt, insbesondere können sich die Folgenglieder der Folge der Partialsummen nicht mehr viel voneinander unterscheiden.

Die Frage, die man sich nun stellen kann (und sollte) ist die Folgende: gilt Äquivalenz, d.h. ist eine Reihe genau dann konvergent, wenn die unterliegende Folge eine Nullfolge ist? Kann man Konvergenz von Reihen an diesem griffigen Kriterium festmachen?

Die Antwort ist leider nein, wie Beispiel Nummer 3 von oben zeigt.

Satz 4.12. Die harmonische Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$ divergiert.

Beweis.

$$\begin{aligned} \sum_{n=1}^{\infty} \frac{1}{n} &= 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8} + \frac{1}{9} + \frac{1}{10} + \frac{1}{11} + \frac{1}{12} + \frac{1}{13} + \frac{1}{14} + \frac{1}{15} + \frac{1}{16} + \dots \\ &\geq 1 + \frac{1}{2} + \underbrace{\frac{1}{4} + \frac{1}{4}}_{=\frac{1}{2}} + \underbrace{\frac{1}{8} + \frac{1}{8} + \frac{1}{8} + \frac{1}{8}}_{=\frac{1}{2}} + \underbrace{\frac{1}{16} + \frac{1}{16} + \frac{1}{16} + \frac{1}{16} + \frac{1}{16} + \frac{1}{16} + \frac{1}{16} + \frac{1}{16}}_{=\frac{1}{2}} + \dots \\ &= 1 + \frac{1}{2} + \frac{1}{2} + \frac{1}{2} + \frac{1}{2} + \frac{1}{2} + \dots \end{aligned}$$

Damit ist die Folge der Partialsummen aber unbeschränkt und daraus folgt die Divergenz der harmonischen Reihe. $\square$

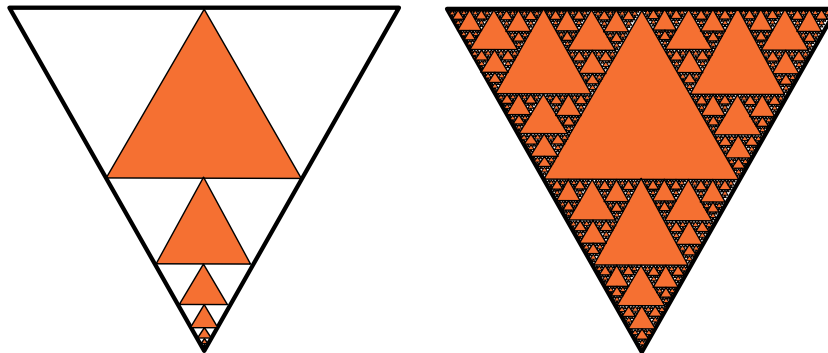
Die anderen Beispiele (Nr 5 und 6) von oben konvergieren. Die vorletzte Reihe wird in der Vorlesung Mathe I vorgestellt. Von der letzten Reihe wissen wir erst seit 1735, welchen Grenzwert sie hat, siehe en.wikipedia.org/wiki/Series_of_reciprocal_squares.

Die harmonische Reihe ist also ein Beispiel einer nicht konvergenten Reihe, die trotzdem eine Nullfolge als unterliegende Folge besitzt. Mit ihrer Hilfe kann man auch die folgende Aufgabe angehen:

Aufgabe* 4.6. Der Dämon und die Schnecke. Es war einmal eine Schnecke, die am Anfang eines Gummibandes saß, welches die Länge von $1m$ hatte. Eines Tages kroch sie los, um die andere Seite zu erreichen und legte dabei pro Tag $5cm$ zurück. Normalerweise wäre sie nach 20 Tagen am anderen Ende angekommen, hätte es nicht einen bösen Dämonen gegeben, der in jeder Nacht um Punkt Mitternacht das Gummiband mitsamt der Schnecke gleichförmig um $1m$ dehnte.

Kommt die Schnecke trotzdem ans Ziel, wenn man davon ausgeht, dass dieses Spezial-Gummi unbegrenzt dehnbar ist und sowohl Dämon als auch Schnecke lange genug leben? Und wenn ja, nach wie vielen Tagen wird es spätestens soweit sein?

Aufgabe* 4.7. Die beiden großen Dreiecke im Bild haben den Flächeninhalt 1. Wie kann man die jeweilige Gesamtfläche der (unendlich vielen, jeweils in der Zeichnung angedeuteten) orangefarbenen bzw. getönten Dreiecke errechnen?



Aufgabe 4.8. Was sind die Grenzwerte der folgenden Reihen?

- | | | |
|--|---|--|
| (a) $\sum_{n=0}^{\infty} \left(\frac{2}{5}\right)^n$ | (b) $\sum_{n=1}^{\infty} \left(\frac{-2}{5}\right)^n$ | (c) $\sum_{n=1}^{\infty} \frac{1}{3^n}$ |
| (d) $\sum_{n=0}^{\infty} \frac{2}{4^{2n}}$ | (e) $\sum_{n=2}^{\infty} \frac{1}{3^{n-1}}$ | (f) $\sum_{n=1}^{\infty} \frac{(-2)^{n-1}}{3^{n+1}}$ |
| (g) $\sum_{n=2}^{\infty} 2 \frac{5}{10^{n-1}}$ | (h) $\sum_{n=1}^{\infty} \frac{(-3)^{n+2}}{2^{2n+1}}$ | |

4.3 Bonusmaterial: Potenzreihen

Ein wichtiges Ziel in Mathe I ist: (Manche) **Funktionen kann man als Potenzreihen schreiben**. Darauf gehen wir hier nur kurz ein. Aber ein bis zwei Beispiele sollte man kennen.

Eines kennen wir schon: die (unendliche) geometrische Reihe:

$$f(x) = \frac{1}{1-x} = \sum_{n=0}^{\infty} x^n \quad (|x| < 1)$$

Da steht eine Funktion mit Definitionsmenge $\{x \in \mathbb{R} \mid |x| < 1\}$, also $f:]-1; 1[\rightarrow \mathbb{R}$, $f(x) = \frac{1}{1-x} = \sum_{n=0}^{\infty} x^n$. Der Funktionswert an der Stelle x ist der Grenzwert der Reihe, wenn x eingesetzt wird.

Eine Potenzreihe sieht allgemein so aus: $\sum_{n=0}^{\infty} a_n x^n$. Die Definitionsmenge ergibt sich aus den x -Werten, für die die Reihe konvergiert. Im Bsp oben konvergiert die Reihe ja nur für $|x| < 1$, also ist die Definitionsmenge $] - 1; 1[$. Ein anderes (wichtiges!) Beispiel ist die Potenzreihe $\sum_{n=0}^{\infty} \frac{1}{n!} x^n$. Die dadurch definierte Funktion ist (die) **Exponentialfunktion** (oder kurz **e-Funktion**). In der Vorlesung Mathe I wird gezeigt, dass diese Reihe für alle $x \in \mathbb{R}$ konvergiert. Also können wir als Definitionsmenge einfach ganz $\mathbb{R}$ nehmen:

$$\exp: \mathbb{R} \rightarrow \mathbb{R}, \quad \exp(x) = \sum_{n=0}^{\infty} \frac{1}{n!} x^n$$

So wird die Exponentialfunktion definiert. Statt $\exp(x)$ schreibt man auch e^x , wobei $e = 2,7182818\dots$ die **eulersche Zahl** ist. (Es gilt $e \notin \mathbb{Q}$, d.h. man kann e nicht als Bruch ganzer Zahlen schreiben.) Das passt, wenn man die naive Definition der Potenz auf e anwendet: für $x \in \mathbb{Q}$ ist dann wirklich die Potenz gleich dem Grenzwert der Reihe, also $e^x = \exp(x)$.

Rechnet man ganz viele Funktionswerte aus (also Grenzwerte der Reihe, etwa für $x = \frac{1}{2}, \frac{1}{4}, \frac{1}{3}, \frac{2}{3}, \frac{3}{5} \dots$) dann kann man den Graphen der Exponentialfunktion zeichnen: Am Graphen sieht man auch,

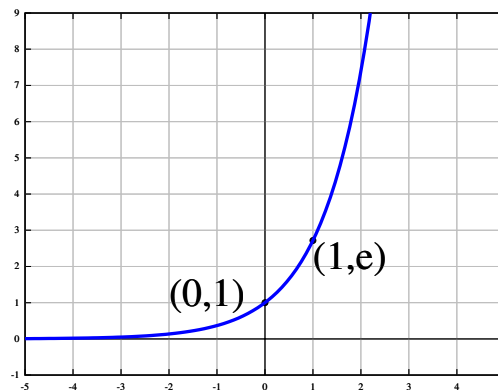


Abbildung 2: Der Graph der Exponentialfunktion.

was als Grenzwert der Reihe für $x = 0$ herauskommen soll: Für $x = 0$ ist die Reihe ja

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

0^n ist 0 für $n \neq 0$. Also sind alle außer dem ersten Summanden 0. Da wir an der Stelle $x = 0$ den Funktionswert 1 haben wollen. Da $0! = 1$ (s. Kapitel 1) ist der erste Summand 0^0 . In

Prinzip kann 0^0 alles sein, aber hier muss $0^0 = 1$, damit der Grenzwert für $x = 0$ wirklich 1 ist. (Das passt dann auch zu e^0 .)

Die Exponentialfunktion (und ihre Umkehrfunktion, der natürliche Logarithmus, dazu mehr im nächste Kapitel) tauchen sehr häufig in allen möglichen Gebieten auf. In der Wahrscheinlichkeitstheorie spielt sie eine extrem wichtige Rolle (bzw die Funktion $e^{-\frac{x^2}{2}}$), für Wachstumsmodelle, bei komplexen Zahlen und Differentialgleichungen und damit in der Physik.

Eine Anwendung ist, dass man mit Hilfe der e-Funktion Potenzen erklären kann, wo der Exponent irrational ist, z.B. $\sqrt{2}^{\sqrt{2}}$ oder für $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^x$. erinnert man sich an die Potenzgesetze, und benutzt man, dass $e^{\ln(x)} = x$, dann ist ja z.B.

$$\sqrt{2}^{\sqrt{2}} = (e^{\ln(\sqrt{2})})^{\sqrt{2}} = e^{\sqrt{2} \ln(\sqrt{2})} \quad \text{bzw} \quad x^x = (e^{\ln(x)})^x = e^{x \ln(x)},$$

und das kann man nun — zumindest im Prinzip — berechnen.

Etliche weitere Funktionen werden formal über ihre Potenzreihe definiert, z.B.

- $\sin(x) = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} = \frac{x}{1!} - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots$
- $\ln(1+x) = \sum_{k=1}^{\infty} (-1)^{k+1} \frac{x^k}{k} = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \dots$
- $\sinh x = \sum_{n=0}^{\infty} \frac{x^{2n+1}}{(2n+1)!} = x + \frac{x^3}{3!} + \frac{x^5}{5!} + \dots$ (“Sinus Hyperbolicus”)

Quizfrage: was ist die Potenzreihe von $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^2$? Was ist allgemein die eines Polynoms?

Falls man das umgekehrte Problem hat: Funktion f gegeben, gesucht seine Potenzreihe: Das Stichwort ist **Taylorreihe**. Dazu mehr in der Vorlesung Mathe I.

5 Abbildungen, aka Funktionen

In mathematics you don't understand things, you get used to them
Wall-E

5.1 Wichtige Vokabeln

Der Begriff der Abbildung spielt in allen Bereichen der Mathematik eine große Rolle. Anfangs mögen die vielen Begriffe, die neu gelernt werden wollen, etwas ungewohnt erscheinen, aber wie das obige Zitat verdeutlichen soll ist alles eine Frage der Gewohnheit.

Abbildungen sind sogenannte “eindeutige Zuordnungen”. Formal bedeutet dies das Folgende: sind zwei Mengen X und Y gegeben, so ist eine Abbildung $f : X \rightarrow Y$ eine Vorschrift, die jedem Element $x \in X$ **genau ein** Element in Y zuordnet, das normalerweise mit $f(x)$ bezeichnet wird. Die Menge X heißt dann **Definitionsmenge** oder **Definitionsbereich**, und die Menge Y wird **Zielmenge** oder **Zielbereich** genannt.

Dieses Konzept der eindeutigen Zuordnung ist uns aus dem Alltag vertraut. So kann man jedem Professor an der Uni sein Büro zuordnen, jedem Teilnehmer einer Klausur sein Ergebnis, etc. Beachtenswert ist hier, dass durchaus erlaubt ist, verschiedenen Elementen der Definitionsmenge das gleiche Element der Wertemenge zuzuordnen und dass es auch durchaus Elemente der Wertemenge geben kann, denen kein Element zugeordnet wird. Ein extremes Beispiel: wenn die Klausur so gut ausfällt, dass alle Studenten eine 1 schreiben, so wird jedem Studenten diese Note zugeordnet und andere Noten kommen nicht vor.

Ein konkreteres Beispiel wäre folgendes:

$$f : \{1, 2, 3\} \rightarrow \{1, 2, 3\}, \quad f(1) = 2, \quad f(2) = 2, \quad f(3) = 3.$$

In der Schule betrachtet man häufig Abbildungen von der Menge $\mathbb{R}$ der reellen Zahlen in sich (man spricht in dem Fall manchmal auch von *Funktionen*), zum Beispiel die Abbildung gegeben durch

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

Eine weitere Notation: will man direkt notieren, wohin ein gewisses Element der Definitionsmenge unter einer Abbildung abgebildet wird, ohne dass die Abbildung einen Namen bekommt (etwa f oder g), so schreibt man dies kurz so (hier am obigen Beispiel illustriert):

$$x \mapsto x^2$$

Das kann man natürlich nur machen, wenn klar ist, um welche Mengen es sich handelt!

Die Beispiele oben machen klar, dass nicht jedes Element der Zielmenge auch als Wert der Abbildung vorkommen muss: Bei f gibt es kein $x \in \{1, 2, 3\}$ mit $f(x) = 1$. Bei g gibt es kein $x \in \mathbb{R}$ mit $g(x) = -1$. Für die Werte, die auch “drankommen”, gibt es die Bezeichnung **Bildmenge** oder **Bild** von f . Kurz schreibt man die Bildmenge von $h : X \rightarrow Y$ als $h(X)$. Die Bildmenge von dem f oben ist also $f(\{1, 2, 3\}) = \{2, 3\}$. Die Bildmenge von dem g oben ist $g(\mathbb{R}) = \mathbb{R}_0^+$.

Noch genauer kann man folgendes tun: Ist nur eine Teilmenge $A \subseteq X$ gegeben, so kann man die Teilmenge von Y betrachten, auf die diese Menge A abgebildet wird. Dies liefert die Definition

der **Bildmenge**, oder kurz: des **Bildes** von A unter einer Abbildung:

$$f(A) := \{f(a) : a \in A\}$$

Es wirkt unglücklich, dass der Begriff Bildmenge eine doppelte Bedeutung hat. Aber das passt, man kann sich dran gewöhnen. Die Bildmenge von X unter f (im zweiten Sinne) ist eben die Bildmenge $f(X)$ im ersten Sinne.

ähnlich wie das Bild kann man für Teilmengen $B \subseteq Y$ das **Urbild** definieren:

$$f^{-1}(B) := \{x \in X : f(x) \in B\}$$

In Worten: das Urbild von B ist die Menge derjenigen Elemente aus X , die von f in die Menge B abgebildet werden. Es ist klar (oder?), dass $f^{-1}(Y) = X$ gilt.

Warnung! Auch wenn es von der Notation her so aussieht, darf das Urbild nicht mit der Umkehrabbildung verwechselt werden, die nicht in allen Fällen existiert! Denn das Urbild muss nicht immer ein Element enthalten: wieder mit obigen Beispielen gilt etwa

$$f^{-1}(\{2\}) = \{1, 2\}, \quad f^{-1}(\{2, 3\}) = \{1, 2, 3\}, \quad f^{-1}(\{1\}) = \emptyset$$

(hierbei bezeichnet $\emptyset$ die leere Menge) sowie im anderen Beispiel

$$g^{-1}(\{-1\}) = \emptyset, \quad g^{-1}(\{4\}) = \{-2, 2\}, \quad g^{-1}([0, 9]) = [-3, 3].$$

Zu beachten ist, dass Bild und Urbild keine Abbildungen der ursprünglichen Mengen sind, sondern vielmehr selbst **mengenwertig** sind!

Aufgabe 5.1. Gegeben sei die Abbildung

$$f : \{1, 2, 3\} \rightarrow \{0, 1, 2, 3, 4\} \quad \text{durch} \quad f(1) = 4, f(2) = 3, f(3) = 4.$$

Bestimme folgende Mengen:

$$f(\{1, 2\}), f(\{1, 2, 3\}), f(\{3\}), f^{-1}(\{4\}), f^{-1}(\{2\}), f^{-1}(\{1, 3\}), f^{-1}(\{0, 1, 2, 4\}).$$

Gegeben sei die Abbildung

$$g : \{-2, 0, 1, 2, 3\} \rightarrow \{-16, -1, 0, 1, 4, 16, 81, 256\} \quad \text{durch} \quad x \mapsto x^4$$

Bestimme folgende Mengen:

$$g(\{0, 3\}), g(\{-2, 2\}), g^{-1}(\{-16, 0, 16\}), g^{-1}(\{-1, 4, 256\}), g^{-1}(\{-1, 1, 81\}).$$

Aufgabe* 5.2. (Knifflig) Gebe Mengen X und Y und eine Abbildung $f : X \rightarrow Y$, sowie Teilmengen $A \subseteq X$ und $B \subseteq Y$ an mit

$$f^{-1}(f(A)) \neq A \quad \text{und} \quad f(f^{-1}(B)) \neq B$$

Dabei können X und Y klein sein, etwa nur drei Elemente haben.

Aufgabe* 5.3. (Zusatz, länglich, knifflig) Seien X und Y Mengen, $f : X \rightarrow Y$ eine Abbildung, $A, B \subseteq X$, $C, D \subseteq Y$. Zeige, oder widerlege durch ein Gegenbeispiel:

$$\begin{aligned} f(A \cup B) &= f(A) \cup f(B) & f(A \cap B) &= f(A) \cap f(B) \\ f^{-1}(C \cup D) &= f^{-1}(C) \cup f^{-1}(D) & f^{-1}(C \cap D) &= f^{-1}(C) \cap f^{-1}(D) \end{aligned}$$

5.2 Injektiv, surjektiv, bijektiv

Der Umgang mit Bild und Urbild erfordert etwas Übung, und wie Aufgabe 5.2 zeigt, ist die Intuition nicht immer der beste Ratgeber. Solche Phänomene müssen aber nicht auftreten, und wenn sie vermieden werden, hat man spezielle Namen für die Abbildungen.

Definition 5.1. Seien X und Y Mengen und $f : X \rightarrow Y$ eine Abbildung. Diese heißt **injektiv**, falls jedes $y \in Y$ höchstens ein Urbild hat. Mit anderen Worten: keinem y werden zwei oder mehr verschiedene Elemente zugeordnet.

Formal kann man die Injektivität einer Abbildung auch so formulieren: sind $x, z \in X$ gegeben mit $f(x) = f(z)$, dann folgt bei einer injektiven Abbildung automatisch $x = z$.

Ob eine Abbildung injektiv ist oder nicht, hängt allerdings nicht nur von der Abbildungsvorschrift, sondern auch von den beteiligten Mengen ab! Betrachte zum Beispiel die Abbildung $f : \mathbb{N} \rightarrow \mathbb{N}$ gegeben durch $f(n) = n^2$. Diese Abbildung ist injektiv, denn aus $n^2 = m^2$ folgt $n = m$ (Übung!). Betrachtet man aber die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$, so ist diese nicht injektiv, da zum Beispiel $f(3) = f(-3)$, aber $3 \neq -3$.

Definition 5.2. Seien wieder X und Y Mengen und $f : X \rightarrow Y$ eine Abbildung. Diese heißt **surjektiv**, wenn es zu jedem $y \in Y$ (mindestens) ein $x \in X$ gibt mit $f(x) = y$, d.h. jedes $y \in Y$ soll mindestens ein Urbild haben.

Die Surjektivität einer Abbildung garantiert nun, dass jedes Element im Wertebereich auch “getroffen” wird.

Als Beispiel betrachten wir die Funktion $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$ gegeben durch $f(x) = \sqrt{x}$. Diese ist surjektiv: sei $y \in \mathbb{R}_0^+$ beliebig gegeben, dann ist $x := y^2$ ein Urbild unter f , denn $f(x) = \sqrt{x} = \sqrt{y^2} = y$.

Injektivität und Surjektivität schließen einander nicht aus: wenn jedes $y \in Y$ **genau ein** Urbild unter f hat, dann ist die Abbildung sowohl injektiv als auch surjektiv und wird in diesem Fall **bijektiv** genannt. Dann (und nur dann!) kann man wirklich eine **Umkehrabbildung** definieren, also eine Abbildung $g : Y \rightarrow X$, die f in gewissem Sinn rückgängig macht, indem nämlich jedem $y \in Y$ sein eindeutig bestimmtes Urbild zugeordnet wird. Diese Umkehrabbildung wird oft auch mit f^{-1} bezeichnet, was allerdings nicht mit dem Urbild verwechselt werden darf!

Die Wurzelfunktion aus dem obigen Beispiel ist nicht nur surjektiv, sondern auch injektiv, denn falls $\sqrt{x} = \sqrt{x'}$, dann folgt durch Quadrieren sofort $x = x'$ (Beachte, dass aufgrund der Mengen $x \geq 0$ und $x' \geq 0$ gilt!). Die Abbildung ist also bijektiv und $y \mapsto y^2$ ist die Umkehrabbildung.

Zum Schluss des Abschnittes soll noch die Verkettung von Abbildungen definiert werden. Seien X, Y und Z Mengen und seien Abbildungen $f : X \rightarrow Y$ und $g : Y \rightarrow Z$ gegeben. Dann kann man die Verkettung (oder Verknüpfung) von f und g wie folgt definieren: die verkettete Abbildung wird mit $g \circ f : X \rightarrow Z$ bezeichnet und ist definiert durch

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

Dies kann man sich leicht anhand des folgenden Diagramms vorstellen:

$$\begin{array}{ccccc} X & \xrightarrow{f} & Y & \xrightarrow{g} & Z \\ X & \longrightarrow & \xrightarrow{g \circ f} & \longrightarrow & Z \end{array}$$

Zu guter Letzt noch eine spezielle Abbildung: ist X eine Menge, so gibt es stets die Abbildung $f : X \rightarrow X$ mit $f(x) = x$ für jedes $x \in X$. Diese Abbildung wird oft als **Identität** auf X bezeichnet und mit id_X notiert.

Aufgabe 5.4. Welche der folgenden Abbildungen sind injektiv, surjektiv bzw. bijektiv?

- a) $f : \mathbb{N} \rightarrow \mathbb{N}$ mit $f(n) = n + 1$.
- b) $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(n) = n + 1$.
- c) $f : \mathbb{N} \rightarrow \mathbb{N}$ mit $f(n) = n^2$.
- d) $f : \mathbb{R} \rightarrow \mathbb{R}_0^+$ mit $f(x) = x^2$.
- e) $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \sin(x)$.
- f) $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}$ mit $f(x) = 3x - 7$.

Aufgabe 5.5. Gegeben sei die Abbildung

$$f : \{1, 2, 3, 4, 5\} \rightarrow \{1, 2, 3, 4, 5\} \text{ durch } f(1) = 4, f(2) = 1, f(3) = 2, f(4) = 4, f(5) = 1.$$

Verkleinere Definitionsmenge und Wertemenge so, dass f bijektiv wird. (Dazu gibt es mehrere Möglichkeiten. Finde eine, wo Definitions- und Wertebereich möglichst groß bleiben.)

Gegeben sei die Abbildung

$$g : \{-5, -3, -2, -1, 0, 2, 5\} \rightarrow \{-7, 0, 2, 3, 4, 5, 7, 9, 12, 15, 16, 19, 25, 26, 28, 30\}$$

durch $g(x) := x^2 + 3$. Verkleinere Definitionsmenge und Wertemenge so, dass g bijektiv wird.

Aufgabe* 5.6. (Knifflig) Finde eine bijektive Abbildung von $\mathbb{N}$ nach $\mathbb{Z}$. Was ist die Umkehrabbildung? (Die Abbildungen dürfen auch mit **if**-Anweisungen beschrieben werden.)

Aufgabe* 5.7. (Knifflig) Gebe eine (möglichst einfache) Menge X und zwei Abbildungen $f : X \rightarrow X$ und $g : X \rightarrow X$ an, so dass gilt $f \circ g \neq g \circ f$. Ist es auch möglich, f und g bijektiv mit dieser Eigenschaft zu wählen?

5.3 Umkehrfunktionen

Mit Hilfe der Identität kann nun die Umkehrabbildung formal sauber definiert werden:

Definition 5.3. Seien X und Y Mengen und $f : X \rightarrow Y$ eine bijektive Abbildung. Die Umkehrabbildung $g : Y \rightarrow X$ ist durch die Eigenschaften

$$g \circ f = \text{id}_X \quad \text{und} \quad f \circ g = \text{id}_Y$$

eindeutig bestimmt.

Mit anderen Worten: verkettet man eine bijektive Abbildung mit ihrer Umkehrabbildung, erhält man immer die Identität.

Wir sahen in Kap. 5.2 bereits einfache Umkehrabbildungen. Abbildungen von $\mathbb{R}$ nach $\mathbb{R}$ heißen auch oft Funktionen. Deren Umkehrabbildungen heißen daher oft Umkehrfunktionen. Zu $f : \mathbb{R} \rightarrow \mathbb{R}$ lässt sich also die Umkehrfunktion bilden, wenn f bijektiv ist. Dazu müssen evtl Definitionsmenge und Bildmenge verkleinert werden. Dann können wir ein eindeutiges f^{-1} finden mit $f^{-1}(f(x)) = x$.

Manchmal sind die Umkehrabbildungen einfach zu finden. Z.B. ist für $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$, $f(x) = x^2$ die Umkehrfunktion $f^{-1} : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$, $f^{-1}(x) = \sqrt{x}$ (s. Kap 5). Was sind die Umkehrfunktionen von

- $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x + 1$?
- $g : \mathbb{R}^+ \rightarrow \mathbb{R}^+$, $g(x) = \sqrt[3]{x}$?
- $h : \mathbb{R}^+ \rightarrow [2, \infty[$, $h(x) = \sqrt[3]{x+1} + 1$?
- $\tilde{f} : \mathbb{R}^+ \rightarrow]0; 1]$, $\tilde{f}(x) = \frac{1}{x^2+1}$?

Im Allgemeinen berechnet man das so: Schreibe $y = f(x)$, stelle nach x um (d.h. forme um, bis da steht $x = \dots$, und rechts kommt kein x mehr vor.) Für's dritte Beispiel also:

$$y = \sqrt[3]{x+1} + 1 \Leftrightarrow (y-1)^3 = x+1 \Leftrightarrow (y-1)^3 - 1 = x$$

also $h^{-1} : \mathbb{R} \rightarrow \mathbb{R}$, $h^{-1}(x) = (x-1)^3 - 1$.

Aufgabe 5.8. Bilde die Umkehrfunktionen der folgenden Funktionen. Verkleinere bei Bedarf Definitions- oder Zielbereich.

1. $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = 2x - 1$
2. $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = \frac{x}{x+1}$
3. $f : \mathbb{R}^+ \rightarrow \mathbb{R}$, $f(x) = \frac{x^2-1}{2x}$
4. $f : \mathbb{R}^+ \rightarrow \mathbb{R}$, $f(x) = x^2 - 6x + 10$
5. $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = \frac{1}{x^2+2x+1}$
6. $f : \mathbb{R}^+ \rightarrow \mathbb{R}$, $f(x) = x + \sqrt{2x}$ (knifflig)
7. $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^4 + 4x^2 + 1$

Anschaulich kann man den Graphen der Umkehrfunktion f^{-1} vom Graphen der ursprünglichen Funktion f ableiten: Man muss den Graphen von f einfach an der Winkelhalbierenden zwischen (positiver) x - und y -Achse spiegeln. Vgl. Abb. 2.

Wie sieht das aus bei....

Exponentialfunktion: Beachte: die Exponentialfunktion $\exp(x) = e^x$ ist eigentlich durch eine Potenzreihe definiert (siehe Abschnitt 4.3). Es ist aber OK, sie sich als e^x vorzustellen, denn für alle x klappt das. (Dabei ist $e = \sum_{k=0}^{\infty} \frac{1}{k!} \approx 2,71828 \dots$).

Ist $\exp : \mathbb{R} \rightarrow \mathbb{R}$ injektiv? Ja (Vgl. Bild 3). Ist $\exp$ surjektiv? Nein, denn die Funktionswerte sind alle positiv. Um $\exp$ surjektiv zu machen, müssen wir also die Zielmenge verkleinern:

$$\exp : \mathbb{R} \rightarrow \mathbb{R}^+, \exp(x) = \dots$$

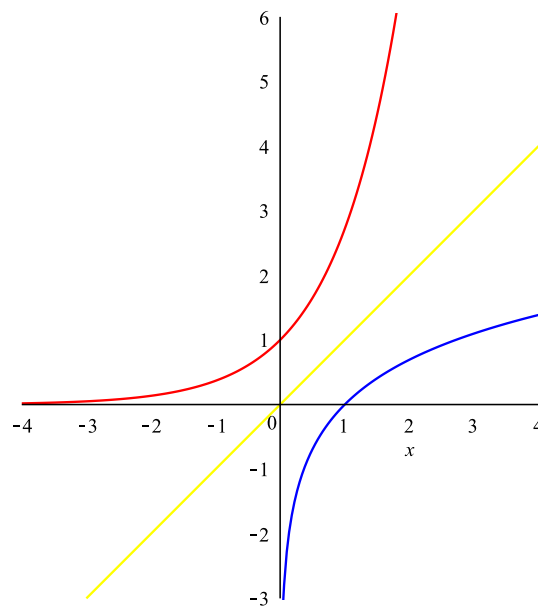


Abbildung 3: Rot bzw oben: die Exponentialfunktion $\exp$. Blau bzw unten: ihre Umkehrfunktion $\ln$. In gelb (bzw hell) die Winkelhalbierende der Koordinatenachsen.

Jetzt lässt sich f^{-1} bilden. Allerdings können wir f^{-1} zwar hinmalen, aber nicht mittels $\exp$, Polynomen, $\sin$, $\cos$... ausdrücken. Die Umkehrfunktion der Exponentialfunktion bekommt einen Namen: **natürlicher Logarithmus**, kurz: $\ln$ (für “logarithmus naturalis”) und ist *definiert* als Umkehrfunktion der Exponentialfunktion. Definitionsmenge und Bildmenge tauschen ihre Rollen, also ist

$$\ln : \mathbb{R}^+ \rightarrow \mathbb{R}, \ln(x) := \exp^{-1}(x)$$

(wobei $\cdot^{-1}$ hier wieder bedeutet: Umkehrfunktion, nicht “1 durch ...”).

Umkehrfunktion sein heißt ja, die Wirkung der Originalfunktion aufzuheben (im Sinne von $f \circ g = \text{id}$). Also ist dadurch, dass wir $\ln$ als $\exp^{-1}$ *definiert* haben, automatisch klar:

$$e^{\ln(x)} = \exp(\ln(x)) = x \text{ und } \ln(e^x) = \ln(\exp(x)) = x.$$

Ein paar Funktionswerte sind damit auch klar:

- $\ln(1) = 0$, denn $1 = e^0$, also $\ln(1) = \ln(e^0) = 0$.
- $\ln(e) = 1$, denn $e^1 = e$, also $\ln(e) = \ln(e^1) = 1$.
- $\ln(\frac{1}{e}) = -1$, denn $e^{-1} = \frac{1}{e}$, also $\ln(\frac{1}{e}) = \ln(e^{-1}) = -1$.

Allerdings lässt sich der Logarithmus als Potenzreihe schreiben: (etwas getrickst $\ln(x+1)$, sonst nicht konvergent)

$$\ln(1+x) = \sum_{k=1}^{\infty} (-1)^{k+1} \frac{x^k}{k} = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \dots \quad (\text{für } |x| < 1)$$

5.4 Prominente Funktionen: Polynome, Sinus, Kosinus und Kollegen

Do not worry about your difficulties in Mathematics. I can assure you mine are still greater.

Daniela Katzenberger

5.4.1 Polynome

In der Analysis spielen Funktionen, deren Definitions- und Zielbereich $\mathbb{R}$ ist (oder Teilmengen davon, etwa Intervalle), eine zentrale Rolle. Es gibt ein paar Funktionen, die immer wieder auftauchen. Eine große Klasse einfacher Beispiele sind Polynome: ein **Polynom** (oder auch **ganzrationale Funktion**) ist eine Funktion der Form $\sum_{k=0}^n a_k x^k$. Also $a_n x^n + a_{n-1} x^{n-1} + \dots + a_2 x^2 + a_1 x + a_0$, wobei $a_i \in \mathbb{R}$. Z.B. sind $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^2 + 1$, $g: \mathbb{R} \rightarrow \mathbb{R}$, $g(x) = -7x^3 + 15x^2 - \frac{7}{3}x + \pi$ und $h: \mathbb{R} \rightarrow \mathbb{R}$, $h(x) = 0,1x^5 - 10^{17}x^3 - 1$ Polynome. Die sollten in etwa aus der Schule bekannt sein.

Eine kompliziertere Klasse von Funktionen sind **gebrochen rationale Funktionen**, die haben die Form $\frac{\text{Polynom}}{\text{anderes Polynom}}$. Z.B. sind $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = \frac{x^2+1}{x+1}$, $g: \mathbb{R} \rightarrow \mathbb{R}$, $g(x) = \frac{x^7+100x^5-x^3+1}{x^4+x^3-x^2+x+1}$ und $h: \mathbb{R} \rightarrow \mathbb{R}$, $h(x) = \frac{\sqrt{2}x^5-x^2}{x^4+4x^3+x}$ gebrochen rationale Funktionen. Die sind oft sehr viel komplizierter im Umgang als Polynome! (Z.B. beim Ableiten, oder Stammfunktion finden.) Weiter unten wollen wir auf ein paar weitere Funktionen eingehen, die keine Polynome oder gebrochen rationale Funktionen sind. Aber ein wichtiges Werkzeug zum Umgang mit Polynomen und gebrochen rationalen Funktionen behandeln wir jetzt: die **Polynomdivision**. Manchmal nämlich lässt sich eine scheinbar gebrochen rationale Funktion wie $f: \mathbb{R}^+ \rightarrow \mathbb{R}$, $f(x) = \frac{x^2-1}{x+1}$ als Polynom schreiben. In diesem Beispiel ist $f(x) = x - 1$ (warum?).

Oder wir wollen die Nullstelle eines Polynoms finden, z.B. von $x^2 - 2x - 1$, oder von $x^3 - 4x^2 + 2x + 4$. Das erste ist sehr einfach (p-q-Formel liefert $1 \pm \sqrt{2}$), das zweite sehr schwierig! Es gibt keine einfache p-q-Formel für kubische Gleichungen (also $x^3 + px^2 + qx + r$). Es klingt verrückt, aber die Methode hier ist: raten! Genauer: es gibt eine Regel.

Satz 5.4 (Vieta). *Hat ein Polynom der Form $\pm x^n + a_{n-1}x^{n-1} + \dots + a_1x + a_0$ ($a_i \in \mathbb{Z}$) eine ganze Zahl $k \in \mathbb{Z}$ als Nullstelle, so gilt: $|k|$ ist Teiler von a_0 .*

Falls wir also Glück haben, hat unser Polynom $x^3 - 4x^2 + 2x + 4$ eine ganzzahlige Nullstelle $k \in \mathbb{Z}$. Dann muss $|k|$ ein Teiler von 4 sein. Das heißt, wir müssen testen: 1, -1, 2, -2, 4, -4. Nacheinander einsetzen zeigt: 1 und -1 sind keine Nullstellen, aber 2 schon! ($2^3 - 4 \cdot 2^2 + 2 \cdot 2 + 4 = 0$). Das hilft uns jetzt, alle Nullstellen zu finden. Dazu führen wir eine Polynomdivision durch. Das Prinzip zeigen wir hier an unserem Beispiel: Wir dividieren schriftlich unser Polynom durch das Polynom $x - 2$ (im Allgemeinen immer: x minus Nullstelle).

$$\begin{array}{r} (x^3 - 4x^2 + 2x + 4) : (x - 2) = x^2 - 2x - 2 \\ \underline{-x^3 + 2x^2} \\ -2x^2 + 2x \\ \underline{2x^2 - 4x} \\ -2x + 4 \\ \underline{2x - 4} \\ 0 \end{array}$$

Wir wissen nun also, dass $x^3 - 4x^2 + 2x + 4 = (x - 2)(x^2 - 2x - 2)$ ist. Die beiden anderen Nullstellen finden wir nun, indem wir die p-q-Formel auf $x^2 - 2x - 2$ anwenden. Das liefert $1 \pm \sqrt{3}$. Die drei(!) Nullstellen von $x^3 - 4x^2 + 2x + 4$ sind also $2, 1 + \sqrt{3}$ und $1 - \sqrt{3}$.

Im Beispiel eben ging die Division glatt auf (so wie $12 : 4$). Der Grund ist, dass 2 wirklich eine Nullstelle ist. Oft geht die Division nicht glatt auf (so wie $12 : 5$).

$$\begin{array}{r}
 (x^3 - 4x^2 + 2x + 4) : (x - 1) = x^2 - 3x - 1 + \frac{3}{x - 1} \\
 \underline{-x^3 \quad +x^2} \\
 -3x^2 + 2x \\
 \underline{3x^2 - 3x} \\
 -x + 4 \\
 \underline{x - 1} \\
 3
 \end{array}$$

In dem Fall schreibt man an der Stelle, wo es nicht weitergeht (hier bei der “3” unten), einfach den Rest zum Ergebnis dazu (hier $+\frac{3}{x-1}$).

In Kapitel 6 sehen wir, warum genau die erste Division aufgeht und die zweite nicht. Grob gesagt kann man jedes Polynom $x^n + a_{n-1}x^{n-1} + \dots + a_1x + a_0$ schreiben als Produkt $(x - \lambda_1)(x - \lambda_2) \dots (x - \lambda_n)$, wobei die λ_i die Nullstellen des Polynoms sind. Das ist aber erst dann wahr, wenn wir komplexe Zahlen zulassen. Die lernen wir in Kapitel 6 kennen.

Aufgabe 5.9. Was ist $x^3 + 2x^2 - 3x - 6$ geteilt durch $x - 1$? Was ist $x^3 + 2x^2 - 3x - 6$ geteilt durch $x + 2$? Was sind also die Nullstellen von $x^3 + 2x^2 - 3x - 6$?

Was ist $x^4 - 2x^3 - 2x^2 + 6x - 3$ geteilt durch $x^2 - 2x + 1$? Was sind die Nullstellen von $x^4 - 2x^3 - 2x^2 + 6x - 3$?

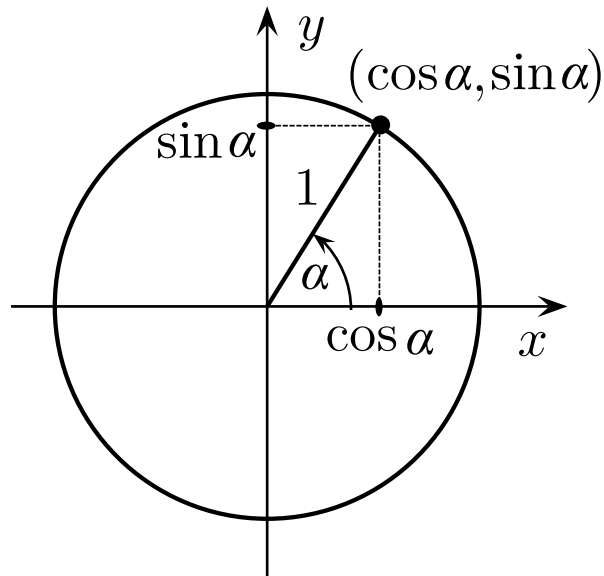
Zusatz: Was sind die Nullstellen von $x^4 + 3x^3 - x^2 - 5x + 2$?

5.4.2 Sinus, Kosinus und Kollegen

Eine weitere wichtige Funktion ist natürlich die **Exponentialfunktion** aus dem letzten Kapitel. Aus all diesen kann man durch Kombination (plus, mal, geteilt, ineinander einsetzen...) weitere Funktionen gewinnen, wie etwa $\frac{e^{x^2+x+1}+x^3}{x^2+e^{-2x}}$.

Die trigonometrischen oder Winkelfunktionen Sinus, Cosinus und Tangens (um die wichtigsten zu nennen) sind vielleicht schon aus der Schule bekannt. Allerdings ist die Notation und der Gebrauch an der Uni etwas anders als gewohnt, weshalb hier kurz darauf eingegangen werden soll.

Es gibt eine ganze Reihe verschiedener Möglichkeiten, wie die Winkelfunktionen definiert werden. Oben haben wir schon angedeutet, dass sin und cos als Potenzreihen definiert werden können. Der Transparenz halber sollen sie jetzt aber als Funktionen am Einheitskreis verstanden werden:



Der Vorteil hiervon ist, dass die Funktion $\sin$ und $\cos$ damit automatisch für alle reellen Werte für α erklärt werden können, also auch für Winkel, die größer als 90° sind oder sogar für negative Winkel. Hierbei gilt, wie man direkt ablesen kann:

$$\sin(-\alpha) = -\sin \alpha \quad \cos(-\alpha) = \cos \alpha$$

Bleibt die Frage, in welcher Einheit der Winkel α angegeben werden soll. Gewöhnt ist man hierbei das sogenannte Gradmaß, bei dem der Vollkreis in 360 Teile eingeteilt wird und die Winkel entsprechend in “Grad” angegeben werden, ein rechter Winkel entspricht also zum Beispiel 90° .

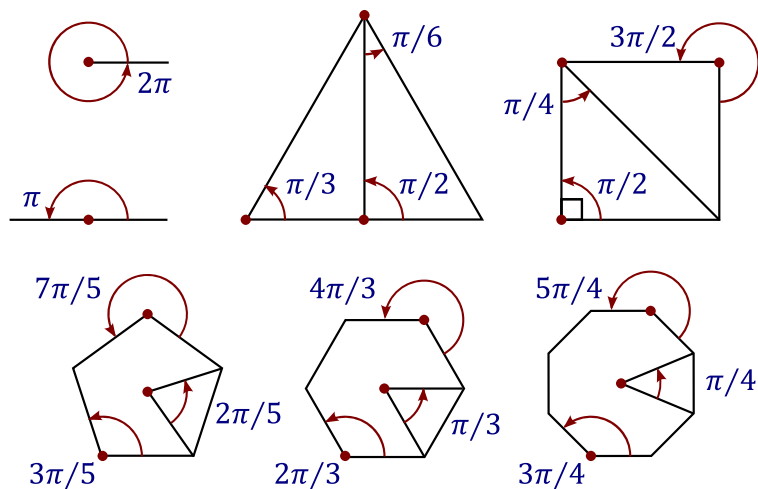
Diese Art Winkel zu messen ist an der Uni allerdings unüblich. Statt dessen begegnet einem hier in erster Linie das Bogenmaß. Bei dieser Einheit werden Winkel in der Länge des Kreisbogens notiert, den sie abstecken, wobei angenommen wird, dass der Kreis den Radius 1 hat. Dem Vollkreis entspricht also im Bogenmaß der Winkel 2π (denn das ist genau der Umfang eines Kreises mit Radius 1) und ein rechter Winkel wird als $\frac{\pi}{2}$ notiert.

Das erscheint anfangs gewöhnungsbedürftig, ist aber keine große Sache, selbst der Taschenrechner lässt sich per Knopfdruck umschalten. Und es ist ja auch nicht schwer, zwischen den Maßen umzurechnen, der Faktor beträgt nämlich einfach

$$\frac{2\pi}{360} = \frac{\pi}{180}$$

Mit anderen Worten: ist ein Winkel im Gradmaß gegeben, so muss er nur mit $\frac{\pi}{180}$ multipliziert werden und ist damit ins Bogenmaß umgerechnet. Umgekehrt geht es natürlich genauso: ist ein Winkel im Bogenmaß notiert, so rechnet man ihn ins Gradmaß um, indem man ihn mit $\frac{180}{\pi}$ multipliziert.

Einige Winkel kann man sich vielleicht so sogar leichter merken:



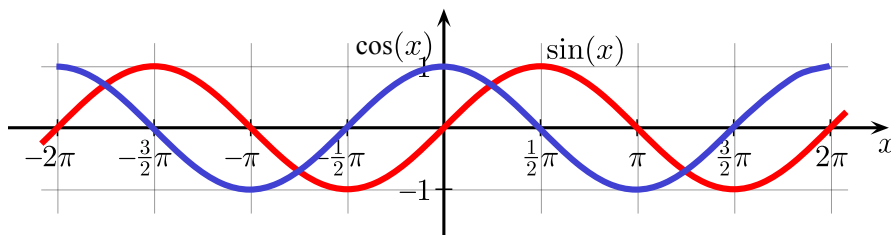
Allgemeiner ist der Innenwinkel in einem regulären n -Eck immer $\frac{n-2}{n}\pi$. ("regulär" heißt: Alle Seiten gleich lang, alle Innenwinkel gleich.)

Ein weiterer Vorteil der Definition von Sinus und Cosinus am Einheitskreis ist, dass der Satz des Pythagoras eine wichtige Identität liefert. Für alle reellen Zahlen x gilt nämlich:

$$\sin^2 x + \cos^2 x = 1$$

Hierbei ist $\sin^2 x$ eine Kurzschreibweise für $(\sin(x))^2$. Diese Vereinbarung hilft, Klammern zu sparen.

Weil Sinus und Kosinus ja nun für alle reellen Zahlen definiert ist, können wir den Graph dieser Funktionen zeichnen. Der sieht so aus:



Ein paar Werte lohnt es sich zu merken: Offenbar ist $\sin(0) = 0$, allgemeiner $\sin(\pi) = \sin(2\pi) = 0$, noch allgemeiner $\sin(k\pi) = 0$ für alle $k \in \mathbb{Z}$. Außerdem wiederholt sich der Graph periodisch: Zwischen 2π und 4π sieht er genau so aus wie zwischen 0 und 2π . Genauer gilt:

$$\sin(x + k2\pi) = \sin(x) \quad \text{für } k \in \mathbb{Z}.$$

Der Graph des Kosinus sieht genau so aus wie der des Sinus, ist aber um ein Stück nach links verschoben. Genauer: um $\frac{\pi}{2}$.

$$\cos\left(x - \frac{\pi}{2}\right) = \sin(x)$$

Außerdem gilt (was man am Graph sieht, aber auch am Einheitskreis)

$$\cos(-x) = \cos(x), \quad \sin(-x) = -\sin(x).$$

Merkt man sich dazu noch die Werte im folgenden Merkschema, dann kann man sich einen Sinus-Graphen malen, wann immer man einen braucht.

α	0	$\frac{\pi}{6}$	$\frac{\pi}{4}$	$\frac{\pi}{3}$	$\frac{\pi}{2}$
$\sin(\alpha)$	$\frac{\sqrt{0}}{2}$	$\frac{\sqrt{1}}{2}$	$\frac{\sqrt{2}}{2}$	$\frac{\sqrt{3}}{2}$	$\frac{\sqrt{4}}{2}$

Der **Tangens** ist nun einfach zu erklären: er ist das Verhältnis der Längen “Gegenkathete durch Ankathete”, also

$$\tan x = \frac{\sin x}{\cos x}$$

Wieder eine Vereinbarung, um Klammern zu sparen: $\sin x$ heißt dasselbe wie $\sin(x)$, wenn Missverständnisse ausgeschlossen sind.

Zum Schluss noch die Additionstheoreme von Sinus und Cosinus:

$$\begin{aligned}\sin(x+y) &= \sin x \cdot \cos y + \cos x \cdot \sin y \\ \cos(x+y) &= \cos x \cdot \cos y - \sin x \cdot \sin y\end{aligned}$$

Wer sich diese Formeln nicht ohne weiteres merken kann, der sei beruhigt: ich auch nicht. Diese Formeln finden sich aber in jedem Nachschlagewerk, sei es auf Papier, sei es online.

Es ist möglich (aber relativ aufwändig), diese mit geometrischen Mitteln zu beweisen. Einfacher und eleganter sind Beweise mit Hilfe der komplexen Exponentialfunktion oder Matrizenrechnung, auf die wir evtl. später noch eingehen.

Aufgabe 5.10. Zeige die Gültigkeit der folgenden Gleichungen:

- $1 + \tan^2 x = \frac{1}{\cos^2 x}.$
- $\tan x = \frac{\sin x}{\sqrt{1-\sin^2 x}} \quad (\text{für } x \in]-\frac{\pi}{2}, \frac{\pi}{2}[)$
- $\tan(x+y) = \frac{\tan x + \tan y}{1 - \tan x \tan y}, \tan(x-y) = \frac{\tan x - \tan y}{1 + \tan x \tan y}$
- $\sin 2x = 2 \sin x \cos x = \frac{2 \tan x}{1 + \tan^2 x}$
- $\sin^2 \theta = \frac{1 - \cos 2\theta}{2}$

Zusatz: Aus einem Abschlussexamen des Jahres 1837: (“Trigonometry and Popular Astronomy” der Universität Durham, England): Zeige

- $\cos 2\theta = 2 \cos^2 \theta - 1 \quad (\text{als Hilfe ist nur angegeben: } \sin(x+y) = \sin x \cdot \cos y + \cos x \cdot \sin y)$
- $\tan \theta = \sqrt{\frac{1 - \cos 2\theta}{1 + \cos 2\theta}}$

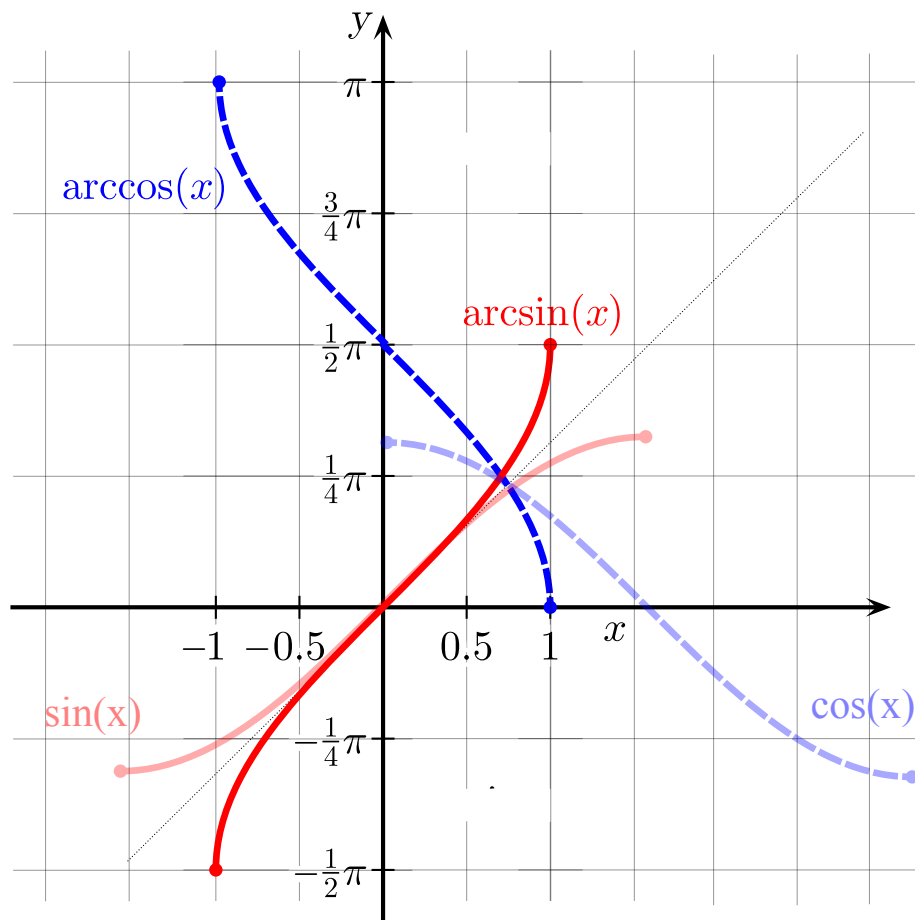
Außerdem gab es in dem Examen noch welche, die schwieriger sind (bei 3 habe ich keine Idee, wie es geht; Nr 4 geht mit vollständiger Induktion, aber die $\sqrt{-1}$ wird erst im nächsten Kapitel erklärt):

- Falls $x + \frac{1}{x} = 2 \cos \alpha$, so ist $x^{m/n} + \frac{1}{x^{m/n}} = 2 \cos(\frac{m}{n} \alpha)$ (mit $m, n \in \mathbb{N}$)
- Zeigen Sie de Moivres Formel: $(\cos \alpha \pm \sqrt{-1} \sin \alpha)^m = \cos(m\alpha) \pm \sqrt{-1} \sin(m\alpha)$; sowie ein paar weitere.

Bemerkung 5.5. Neben Sinus, Kosinus und Tangens gibt es noch etliche weitere trigonometrische Funktionen. Die anderen sind aber leicht auf diese drei zurückzuführen. Es ist z.B. der Kotangens $\cot(x) = \frac{1}{\tan x}$, oder der Sekans $\sec(x) = \frac{1}{\sin x}$. Die braucht man nicht wirklich zu lernen, wenn man die drei ursprünglichen kennt. Anders sieht es aus bei:

Umkehrfunktionen von Sinus und Kosinus: Sind $\sin : \mathbb{R} \rightarrow \mathbb{R}$ und $\cos : \mathbb{R} \rightarrow \mathbb{R}$ bijektiv? Kaum. Beide sind weder injektiv noch surjektiv.

Also schränkt man die Definitions- und Zielmengen geeignet ein: $\sin : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow [-1, 1]$, $\cos : [0, \pi] \rightarrow [-1, 1]$. Dann sind die Umkehrfunktionen dadurch gegeben. Die tauchen öfter mal auf, daher haben die auch spezielle Namen: **arccos** $:= \cos^{-1}$, **arcsin** $:= \sin^{-1}$. Die Graphen sehen so aus:



Definitions- und Zielmengen sehen also so aus: $\arccos : [-1, 1] \rightarrow [0, \pi]$, $\arcsin : [-1, 1] \rightarrow [-\frac{\pi}{2}, \frac{\pi}{2}]$. Die Funktionen arccos und arcsin kann man nicht einfach durch andere Funktionen ausdrücken. Man kann sie allerdings als Potenzreihen schreiben. Das muss man sich nicht

merken, aber um zu zeigen, dass es geht:

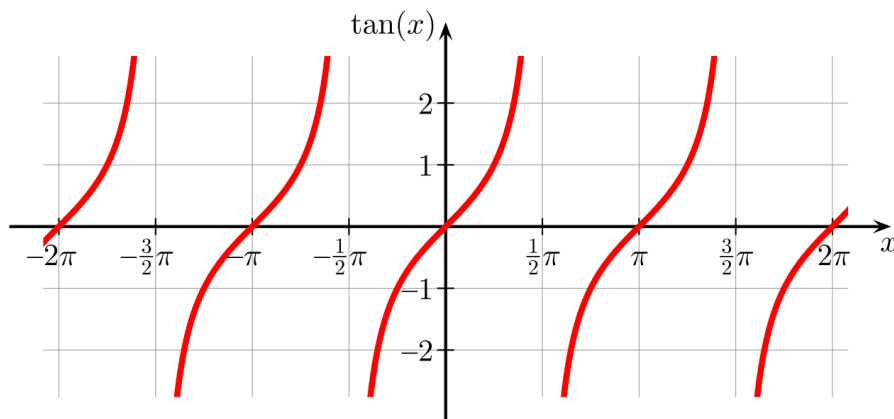
$$\begin{aligned}\arcsin(x) &= \sum_{k=0}^{\infty} \frac{(2k-1)!!}{(2k)!!} \frac{x^{2k+1}}{2k+1} = \sum_{k=0}^{\infty} \binom{2k}{k} \frac{x^{2k+1}}{4^k(2k+1)} \\ &= x + \frac{1}{2} \cdot \frac{x^3}{3} + \frac{1 \cdot 3}{2 \cdot 4} \cdot \frac{x^5}{5} + \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6} \cdot \frac{x^7}{7} + \dots\end{aligned}$$

$$\arccos(x) = \frac{\pi}{2} - \sum_{k=0}^{\infty} \frac{(2k-1)!!}{(2k)!!} \frac{x^{2k+1}}{2k+1} = \frac{\pi}{2} - \sum_{k=0}^{\infty} \binom{2k}{k} \frac{x^{2k+1}}{4^k(2k+1)}$$

Dabei ist $n!!$ die “Doppelfakultät”,

$$n!! = \begin{cases} n \cdot (n-2) \cdot (n-4) \cdot \dots \cdot 2 & \text{für } n \text{ gerade,} \\ n \cdot (n-2) \cdot (n-4) \cdot \dots \cdot 1 & \text{für } n \text{ ungerade,} \end{cases}$$

Tangens Auch der Tangens ist nicht bijektiv: er ist zwar surjektiv (also: alle reellen Zahlen kommen als Funktionswert des Tangens dran), aber nicht injektiv:

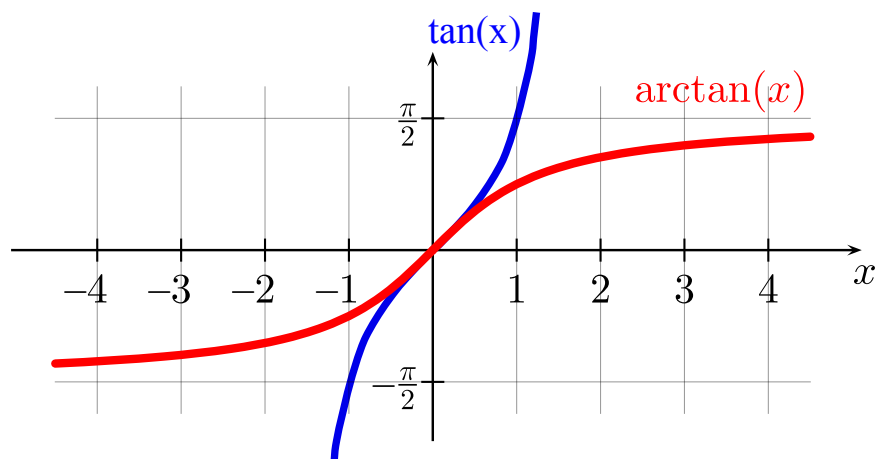


Offenbar gilt

$$\tan(x) = \tan(x + k\pi) \quad (k \in \mathbb{Z})$$

Also insbesondere z.B. $\tan(0) = 0 = \tan(\pi)$ (also nicht injektiv). Man kann den Tangens etwa folgendermaßen bijektiv machen: $\tan : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow \mathbb{R}$. Die Umkehrfunktion heißt **Arcustan-gens**, kurz **arctan**. Definitions- und Zielmenge sehen so aus: $\arctan : \mathbb{R} \rightarrow]-\frac{\pi}{2}, \frac{\pi}{2}[$. Der Graph

sieht so aus:



Auch den Arcustangens kann man als Potenzreihe schreiben, die ist sogar recht simpel:

$$\arctan x = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{2k+1} = x - \frac{1}{3}x^3 + \frac{1}{5}x^5 - \frac{1}{7}x^7 + \dots$$

Damit kann man übrigens den Wert von π berechnen: Wegen $\tan(\frac{\pi}{4}) = 1$ (Wieso?) ist ja $\arctan(1) = \frac{\pi}{4}$. Also ist

$$\pi = 4 \arctan(1) = 4 \left(x - \frac{1}{3}x^3 + \frac{1}{5}x^5 - \frac{1}{7}x^7 + \dots \right)$$

6 Zahlbereiche: $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$

Du wolltest Algebra - da hast Du den Salat

Tim Mälzer

6.1 Von Gruppen und Körpern — $\mathbb{Q}$

Im folgenden Kapitel wollen wir Verallgemeinerungen des Zahlbegriffes ansehen, die im Laufe der mathematischen Entwicklung stattfanden. Die natürlichen Zahlen, mit denen man ganze Größen zählt, stoßen schnell an ihre Grenzen.

Es gibt auf der Menge $\mathbb{N}_0$ der natürlichen Zahlen inklusive 0 eine sogenannte Verknüpfung, nämlich $+$. Das bedeutet, dass zu zwei Zahlen $n, m \in \mathbb{N}_0$ eine dritte Zahl $(n + m) \in \mathbb{N}_0$ existiert, die eindeutig bestimmt ist. Die 0 spielt für diese Verknüpfung eine Sonderrolle: es handelt sich um ein sogenanntes “neutrales Element”, d.h. die Addition von 0 ändert eine natürliche Zahl nicht. Für alle $n \in \mathbb{N}_0$ gilt

$$n + 0 = 0 + n = n. \quad \text{Deshalb gilt:}$$

$\mathbb{N}_0$ (mit $+$) ist ein Monoid

Die natürlichen Zahlen mit 0 bilden bezüglich $+$ eine Struktur, die man **Monoid** nennt. Die allgemeine Definition sieht folgendermaßen aus:

Definition 6.1. Sei M eine Menge und $\circ : M \times M \rightarrow M$ eine Verknüpfung auf M . Dann nennt man M einen Monoid, falls folgende Eigenschaften erfüllt sind:

- i) Assoziativität: Es gilt $(a \circ b) \circ c = a \circ (b \circ c)$ für alle $a, b, c \in M$.
- ii) Neutrales Element: Es gibt ein Element $e \in M$, so dass für alle $m \in M$ gilt: $e \circ m = m \circ e = m$.

Diese Definition verallgemeinert und abstrahiert einige der Eigenschaften der Addition von natürlichen Zahlen. Warum diese Abstraktion wichtig ist, zeigt sich daran, dass diese beiden formalen Eigenschaften auch von Strukturen erfüllt werden, von denen man es vielleicht nicht erwartet und die wir zum Teil auch schon kennen.

Ist zum Beispiel X eine Menge, dann kann man die Menge der Abbildungen von X in sich betrachten:

$$\text{Abb}(X, X) = \{f : X \rightarrow X\}$$

Auf dieser Menge gibt es eine Verknüpfung $\circ$, die wir in Abschnitt 5.2 schon kennengelernt haben, die Verkettung von Abbildungen. Diese ist assoziativ (Übung!) und das neutrale Element ist die Identität id_X . Die Menge $\text{Abb}(X, X)$ ist also ein Monoid.

Ein anderes Beispiel spielt in der theoretischen Informatik eine gewisse Rolle:

Beispiel 6.2. Sei Σ eine endliche Menge von Zeichen (das sogenannte Alphabet), dann betrachte folgende Menge sogenannter **Worte über Σ** :

$$W(\Sigma) = \{a_1 \dots a_k : k \in \mathbb{N}_0, a_i \in \Sigma\}$$

Diese Menge kann man sich als Strings (also Zeichenketten) über dem gegebenen Alphabet Σ vorstellen. Beachte, dass auch die leere Zeichenkette, also das Wort der Länge 0 zugelassen ist. Zur Verdeutlichung bezeichnet man dieses oft mit ε .

Um $W(\Sigma)$ als Monoiden zu identifizieren benötigen wir noch eine Verknüpfung: die **Konkatenation** (Hintereinanderreihung) von Worten. Sind also $a_1 \dots a_k$ und $b_1 \dots b_l$ Worte über Σ , so definiere

$$(a_1 \dots a_k) \circ (b_1 \dots b_l) = a_1 \dots a_k b_1 \dots b_l$$

Es ist nicht schwer zu sehen (oder?), dass diese Verknüpfung assoziativ ist und dass das leere Wort ε das neutrale Element ist.

Zurück zu den natürlichen Zahlen und der Verknüpfung $+$. Um die Operation des Addierens “rückgängig” machen zu können, braucht man negative Zahlen. Diese werden in gewisser Weise einfach “hinzuerfunden”: zu jeder natürlichen Zahl $n \in \mathbb{N}_0$ erklärt man die Existenz einer negativen Zahl $-n$, für die dann gelten soll:

$$n + (-n) = (-n) + n = 0$$

Die Menge aller positiven und negativen Zahlen nennt man dann schlicht die Menge der “ganzen Zahlen” und kürzt sie mit dem Buchstaben $\mathbb{Z}$ ab:

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

$\mathbb{Z}$ (mit $+$) ist eine Gruppe

Definition 6.3. Sei G eine Menge mit einer Verknüpfung $\circ : G \times G \rightarrow G$. Dann heißt G eine Gruppe, falls gilt:

- i) Assoziativität: Es gilt $(a \circ b) \circ c = a \circ (b \circ c)$ für alle $a, b, c \in G$.
- ii) Neutrales Element: Es gibt ein Element $e \in G$, so dass für alle $g \in G$ gilt: $e \circ g = g \circ e = g$.
- iii) Inverses Element: Zu jedem $g \in G$ gibt es ein Element $g^{-1} \in G$, so dass gilt: $g \circ g^{-1} = g^{-1} \circ g = e$.

G heißt außerdem **abelsch**, falls gilt: $\forall a, b \in G : a \circ b = b \circ a$.

In den Vorlesungen wird es später oft euch selbst überlassen, Wichtiges von Unwichtigem zu unterscheiden. Lernt dies! Hier helfen wir noch:
Monoide sind eher unwichtig. Gruppen sind ziemlich wichtig.

Auf der Menge der ganzen Zahlen gibt es allerdings noch eine zweite Verknüpfung, nämlich “ $\cdot$ ”: für beliebige ganze Zahlen $u, v \in \mathbb{Z}$ ist das Produkt $u \cdot v$ definiert und wieder eine ganze Zahl. Für diese Verknüpfung übernimmt die 1 die Rolle des neutralen Elementes. Für alle $u \in \mathbb{Z}$ gilt

$$1 \cdot u = u \cdot 1 = u$$

Und wieder fehlt in $\mathbb{Z}$ für die meisten Zahlen das sogenannte *Inverse*, um die Operation rückgängig zu machen. Daher führt man eine symbolische Schreibweise ein und erklärt für ein $u \in \mathbb{Z}$ mit $u \neq 0$ einfach die Existenz einer “Zahl” $\frac{1}{u}$, für die gelten soll:

$$u \cdot \frac{1}{u} = \frac{1}{u} \cdot u = 1$$

Die 0 muss man hierbei fortlassen, da $a \cdot 0 = 0$ für alle Zahlen a gilt und die 0 folglich kein Inverses besitzen kann.

Wenn diese Überlegungen konsequent fortgesetzt werden, erhält man die Menge der rationalen Zahlen, die mit $\mathbb{Q}$ abgekürzt wird:

$$\mathbb{Q} = \left\{ \frac{u}{v} : u, v \in \mathbb{Z}, v \neq 0 \right\}$$

$\mathbb{Q}$ (mit $+$ und $\cdot$) ist ein Körper

Definition 6.4. Eine Menge K mit zwei Rechenoperationen $\oplus$ und $\odot$ heißt **Körper**, falls

1. $(K, \oplus)$ ist eine abelsche Gruppe
2. $(K \setminus \{0\}, \odot)$ ist eine abelsche Gruppe
3. $\forall a, b, c \in K : \quad a \odot (b \oplus c) = a \odot b \oplus a \odot c$

Dabei steht “0” für das neutrale Element in $(K, \oplus)$.

6.2 Das Füllen der Lücken – $\mathbb{R}$

Die rationalen Zahlen kann man sich als Verhältnis von Größen vorstellen. Den Griechen war die Menge der positiven rationalen Zahlen bekannt und sie glaubten auch, dass sich jede vorstellbare Größe als Verhältnis (also als Bruch) darstellen lässt.

Leider ist das nicht richtig: der Satz des Pythagoras sagt, dass die Länge der Diagonale im Einheitsquadrat eine Größe a ist mit der Eigenschaft $a^2 = 2$. Es gibt aber keine rationale Zahl mit dieser Eigenschaft. Der Beweis funktioniert am besten über einen Widerspruch.

Satz 6.5. $\sqrt{2} \notin \mathbb{Q}$.

Beweis. Angenommen es gibt eine positive Zahl $a \in \mathbb{Q}$ mit $a^2 = 2$. Dann lässt sich a als gekürzter Bruch schreiben:

$$a = \frac{p}{q} \quad \text{mit } p, q \in \mathbb{Z} \text{ teilerfremd}$$

Es gilt aber

$$a^2 = \left(\frac{p}{q} \right)^2 = \frac{p^2}{q^2} = 2 \iff p^2 = 2q^2$$

Damit ist p^2 eine gerade Zahl. Dann muss aber auch p gerade sein, denn das Quadrat einer ungeraden Zahl ist stets ungerade:

$$(2n+1)^2 = 4n^2 + 4n + 1$$

Da also p selbst eine gerade Zahl ist, gibt es eine Zahl $r \in \mathbb{Z}$ mit $p = 2r$. Dies eingesetzt ergibt:

$$(2r)^2 = 2q^2 \iff 4r^2 = 2q^2 \iff 2r^2 = q^2$$

Also ist q^2 gerade und demnach, mit demselben Argument wie oben, auch q .

Das aber ist ein Widerspruch! Denn wenn p und q beide gerade sind, sind beide durch 2 teilbar - wir hatten aber vorausgesetzt, dass die Darstellung von a als Bruch gekürzt ist, d.h. p und q sollten teilerfremd sein. Widerspruch. $\square$

Wieder haben wir also eine Größe, die in unserer bisherigen Zahlenmenge nicht vorkommt. Aber die Prozedur ist die gleiche wie vorher: wir erfinden ein Symbol für diese Größe (in diesem Fall $\sqrt{2}$) und nehmen solcherlei Elemente einfach mit auf. Die größere Menge, die wir erhalten, nennen wir dann $\mathbb{R}$: die Menge der reellen Zahlen.

In Wahrheit ist die Konstruktion der reellen Zahlen aus den rationalen um einiges komplizierter. Man kann es sich ungefähr so vorstellen: jede der neuen reellen Zahlen kann man beliebig gut durch rationale Zahlen approximieren, d.h. annähern. Konkret heißt dies, dass in jeder noch so kleinen Umgebung einer reellen Zahl auf der Zahlengeraden unendlich viele rationale Zahlen liegen. Die reellen Zahlen füllen quasi die Lücken zwischen den rationalen Zahlen aus.

Genauer gesprochen gilt: betrachtet man Folgen in $\mathbb{Q}$, so gibt es darunter solche, die zwar konvergieren sollten, weil die Abstände der Folgenglieder immer geringer werden, es aber dennoch nicht tun, weil ihr Grenzwert keine rationale Zahl ist. So kann man beispielsweise eine Folge rationaler Zahlen angeben, die gegen $\sqrt{2}$ konvergiert. Man erhält die reellen Zahlen nun aus den rationalen, indem man alle diese "fehlenden Grenzwerte" hinzunimmt, man spricht in diesem Zusammenhang auch von einer **Vervollständigung**.

Eine andere Art, die reellen Zahlen aufzufassen, ist die Entwicklung als Dezimalbrüche: man betrachtet die Menge aller "Zahlen" der Form

$$a_k a_{k-1} \dots a_1 a_0, a_{-1} a_{-2} \dots$$

Dabei sind die $a_i \in \{0, 1, \dots, 9\}$ die Ziffern. Gemeint ist mit obigem Ausdruck natürlich eigentlich folgende Reihe:

$$\sum_{i=-k}^{\infty} \frac{a_{-i}}{10^i}$$

Nicht alle Zahlen dieser Form sind rational: wann immer man von einem Bruch ausgeht, entsteht eine Dezimalzahl, die periodisch ist, also ab einem Punkt wiederholen sich die Ziffernfolgen. Die reellen Zahlen kann man also als Menge aller (periodischer und nicht periodischer) Dezimalzahlen auffassen.

Eine detailliertere Behandlung der reellen Zahlen würde den Rahmen des Vorkurses bei Weitem sprengen. In beinahe jedem Lehrbuch der Analysis wird dieses Thema mehr oder weniger ausführlich behandelt, Interessierte seien also darauf verwiesen. Außerdem wird die Konstruktion der reellen Zahlen mit an Sicherheit grenzender Wahrscheinlichkeit auch Gegenstand der Mathematikvorlesung im kommenden Semester sein.

Aufgabe* 6.1. (zum Knobeln) Gibt es irrationale Zahlen $a, b \in \mathbb{R} \setminus \mathbb{Q}$ mit $a^b \in \mathbb{Q}$?

6.3 Die fehlenden Wurzeln – $\mathbb{C}$

Schon in der Schule lernt man, dass das Produkt zweier negativer Zahlen positiv ist, ebenso wie das Produkt zweier positiver Zahlen. Es folgt, dass für jede reelle Zahl $a \in \mathbb{R}$ gilt:

$$a^2 \geq 0$$

Das bedeutet auch, dass die Gleichung $x^2 + 1 = 0$ keine Lösungen besitzt, oder anders ausgedrückt: es existiert keine reelle Zahl, deren Quadrat gleich (-1) ist.

Diese Lücke schließen wir nun auf folgende Art: wir führen eine neue Zahl i ein, für die gelten soll:

$$i^2 = -1$$

Diese Zahl heißt *imaginäre Einheit*. Den neu erschlossenen Zahlbereich nennen wir dann Menge der komplexen Zahlen und kürzen diese mit dem Buchstaben $\mathbb{C}$ ab:

$$\mathbb{C} = \{a + bi : a, b \in \mathbb{R}\}$$

Man kann also jede komplexe Zahl als Summe $a + b \cdot i$ darstellen, wobei a und b reelle Zahlen sind. Man nennt a den *Realteil* und b den *Imaginärteil* der komplexen Zahl $z = a + ib$.

Ist das Problem des Wurzelziehens denn nun vollständig gelöst? Die Wurzel aus (-1) können wir ziehen, wie sieht es mit anderen negativen Zahlen aus? Sei dazu $a \in \mathbb{R}_0^+$ eine beliebige positive Zahl. Dann können wir nun die Wurzel aus $-a$ wie folgt ziehen:

$$\sqrt{-a} = \sqrt{(-1) \cdot a} = \sqrt{-1} \cdot \sqrt{a} = i \cdot \sqrt{a} \in \mathbb{C}$$

(Obacht, manche Leute mögen die Schreibweise nicht: die *Wurzelfunktion* ist eine Funktion, und darf daher nur einen Wert haben. Daher ist die Wurzel als Funktion strenggenommen nur auf den positiven reellen Zahlen definiert. Aber die Gleichung $x^2 = 1$ hat ja zwei Lösungen: neben i auch $-i$. Manche Dozenten trennen das, und man darf nicht $\sqrt{a}$ schreiben, falls $a < 0$.)

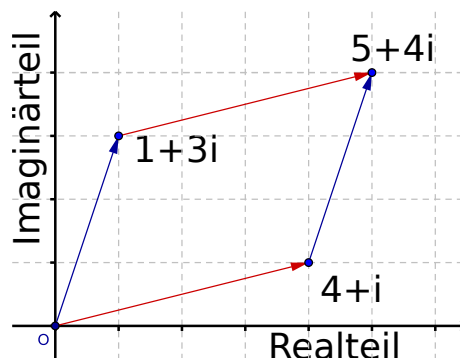
Und wie sieht es mit unseren Verknüpfungen aus? Wenn man die üblichen Rechenregeln zugrundelegt, kann man Summe und Produkt von komplexen Zahlen ganz einfach erklären:

$$(a + ib) + (c + id) = (a + c) + i \cdot (b + d)$$

Und ebenso

$$(a + ib) \cdot (c + id) = ac + iad + ibc + i^2bd = (ac - bd) + i \cdot (ad + bc)$$

Es gibt auch die Möglichkeit, sich die komplexen Zahlen graphisch vorzustellen. Wo die reellen Zahlen eine Zahlengerade bilden, sind die komplexen Zahlen Punkte in einer Ebene:



Die Addition der komplexen Zahlen ist dann die gewöhnliche Vektoraddition, wie sie im Kapitel über Lineare Algebra noch behandelt wird. Die Multiplikation hingegen ist ein wenig komplizierter, aber so, wie man sich das denken würde:

$$(a + bi)(c + di) = ac + bci + adi + bdi^2 = (ac - bd) + (bc + ad)i$$

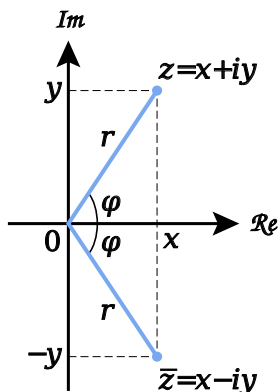
Also etwa $(2 - i)(3 + 2i) = 6 + 4i - 3i - 2i^2 = 6 + 2 + 1i = 8 + i$.

Um bequem mit komplexen Zahlen rechnen zu können, behandeln wir am Schluss noch ein paar Tricks. Eine wichtige Frage ist oft die Folgende: ist $z = a + ib$ gegeben, wie bestimme ich dann $\frac{1}{z}$? Oder konkreter ausgedrückt möchte man auch $\frac{1}{z}$ gern in der Form $c + id$ schreiben, aber wie ermittelt man c und d ?

Zu diesem Zweck ist es nützlich, die sogenannte *komplexe Konjugation* einer Zahl zu betrachten. Dahinter verbirgt sich nichts Mysteriöses, es handelt sich lediglich um eine Spiegelung an der reellen Achse: ist $z = a + ib$, so ist das komplex Konjugierte $\bar{z}$ definiert durch

$$\bar{z} = \overline{a + ib} = a - ib$$

Mit anderen Worten: der Realteil bleibt erhalten, der Imaginärteil wird durch sein Negatives ersetzt.



Auf den ersten Blick ist das nicht so spektakulär, aber die Wichtigkeit wird durch folgende Formel deutlich:

$$z \cdot \bar{z} = (a + ib) \cdot (a - ib) = a^2 + b^2 = |z|^2 \in \mathbb{R}$$

Das Produkt einer Zahl mit ihrem komplex Konjugierten ist also reell. Das ist geometrisch sofort einleuchtend, die Winkel heben sich so auf, dass das Ergebnis auf der reellen Achse liegt. Mit Hilfe dieses Tricks lässt sich die obige Aufgabe aber leicht lösen: wir erweitern den Bruch $\frac{1}{z}$ einfach mit $\bar{z}$:

$$\frac{1}{a + ib} = \frac{1}{z} = \frac{\bar{z}}{z \cdot \bar{z}} = \frac{\bar{z}}{|z|^2} = \frac{a}{a^2 + b^2} + i \frac{-b}{a^2 + b^2}$$

Das Beste an $\mathbb{C}$ aber ist, dass es den folgenden Satz wahr macht.

Theorem 6.6 (Fundamentalsatz der Algebra). *Ein Polynom $a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$ hat genau n Nullstellen, wenn man (a) komplexe Nullstellen mitzählt, und (b) Nullstellen mit ihrer Vielfachheit zählt.*

Punkt (b) ist so zu verstehen: $x^3 - 3x^2 + 3x - 1$ hat eine einzige Nullstelle, nämlich 1. Aber wegen $x^3 - 3x^2 + 3x - 1 = (x - 1)^3$ sollen wir diese Nullstelle dreifach zählen.

Aufgabe 6.2. Vereinfache die folgenden komplexen Ausdrücke, genauer: bringe sie auf die Form $a + bi$.

$$i(1 - i), \quad (2 + i)(2 - i), \quad (7 + 4i)(7 - 4i) - (8 + i)(8 - i), \quad (i - 1)^2, \quad (1 - i)^4, \quad (1 + \sqrt{3}i)^3$$

Aufgabe* 6.3. Löse die Gleichung $(a + ib)^2 = i$. Wie viele Lösungen gibt es? Wo liegen sie in der Ebene?

Aufgabe 6.4. Bestimme alle Lösungen der Gleichung $z^4 = 1$ in $\mathbb{C}$. Was fällt auf? Wo liegen die Ergebnisse geometrisch?

Aufgabe 6.5. Schreibe folgende komplexe Zahlen in der Form $a + ib$:

$$\frac{1+i}{1-i}, \quad \frac{1}{3-4i}, \quad \frac{2+7i}{-1-2i}, \quad \frac{1-7i}{1-2i}, \quad \frac{5-i}{2-3i}, \quad \frac{4i+3}{(1-2i)^2}, \quad \frac{(2i+1)^2+1}{(i+1)^2}, \quad \overline{\left(\frac{1}{i}\right)}, \quad \sqrt{\frac{1+i}{\sqrt{2}}}.$$

(Die letzte ist kniffliger als die anderen. Dort gibt es auch mehr als eine Lösung. Dann reicht eine beliebige.)

Aufgabe* 6.6. (Zusatz) Sei $\zeta \in \mathbb{C}$ (Zeta) eine komplexe Zahl mit $\zeta^n = 1$ für ein $n \in \mathbb{N}$ und $\zeta \neq 1$. (Eine solche Zahl heißt auch n -te Einheitswurzel.) Berechne den folgenden Ausdruck:

$$\sum_{k=0}^{n-1} \zeta^k$$

7 Überblick Mathe I Lineare Algebra (und diskrete Mathematik)

Eine mathematische Aufgabe kann manchmal genauso unterhaltsam sein wie ein Kreuzworträtsel und angespannte geistige Arbeit kann eine ebenso wünschenswerte Übung sein wie ein schnelles Tennisspiel.

Boris Becker

Die Mathematikvorlesung Mathe I im ersten Semester in der (genauer: “Mathematik für Naturwissenschaften I” laut ekVV, “Mathematik für Informatik I” laut Modulhandbuch) ist seit 2023 thematisch klar in drei Teile gegliedert: Diskrete Mathematik (Graphen, Zählen, Modulo-Rechnen), Lineare Algebra (Lineare Gleichungssysteme, Vektorräume, Matrizen) und Analysis (Folgen und Reihen, stetige Funktionen, Ableitung, Integral). Früher waren es nur zwei Teile, lineare Algebra und Analysis.

Die diskrete Mathematik bereitet erfahrungsgemäß am wenigsten Probleme — eigentlich nur dadurch, dass es die erste Universitäts-Mathematik ist, die Sie sehen werden. Wenn Sie später darauf zurückblicken, werden Sie sehen, dass es im Wesentlichen logisches Denken plus Knobeln (wie Sudoku) ist. Daher behandeln wir in diesem Vorkurs nur Lineare Algebra und Analysis.

Im Lineare-Algebra-Teil der Vorlesung passiert in etwa folgendes:

I Reelle und Komplexe Zahlen

II Vektorräume

- Gruppen
- Körper
- Vektorräume
- Dazu: Linearkombination, lineare Hülle (=Spann), linear unabhängig, Erzeugendensystem, Basis, Untervektorraum

III Lineare Abbildungen und Matrizen

- lineare Abbildung, Isomorphismus
- Matrizen, Typen von Matrizen, Rang, inverse Matrix,...
- Satz: jede lineare Abbildung wird von Matrix beschrieben, jede Matrix liefert lineare Abbildungen

IV Matrizen und lineare Gleichungssysteme

- Weitere Zusammenhänge Matrizen, lineare Abbildungen, Inverse, Rang...

Das wollen wir im Folgenden etwas beleuchten. So richtig interessant wird Lineare Algebra erst im zweiten Semester, da kommen Determinanten, Eigenwerte, Orthogonalität; das ist das, was man hinterher im Studium (und in der Anwendung) wirklich braucht (Z.B. für neuronale Netze, Computergrafik oder Biomathematik).

Das wichtigste Recheninstrument ist am Anfang das Lösen von linearen Gleichungssystemen, daher zunächst:

Lineare Gleichungssysteme

Quizfrage: Ein Hotel hat insgesamt 90 Zimmer. (Mit “Zimmer” ist ab jetzt immer Gästezimmer gemeint.) Es gibt Einbett-, Zweibett- und Dreibett-Zimmer. Insgesamt hat das Hotel 133 Gästebetten. Es gibt doppelt so viele Einbettzimmer wie Zwei- und Dreibettzimmer zusammen. Wie viele Ein-, Zwei- und Dreibettzimmer gibt’s?

Dieses Problem führt auf ein **Lineares Gleichungssystem** (kurz LGS). allgemein sieht ein LGS so aus:

$$\begin{aligned} a_{1,1}x_1 + a_{1,2}x_2 + \cdots + a_{1,n}x_n &= b_1 \\ \wedge \quad a_{2,1}x_1 + a_{2,2}x_2 + \cdots + a_{2,n}x_n &= b_2 \\ &\vdots \\ \wedge \quad a_{m,1}x_1 + a_{m,2}x_2 + \cdots + a_{m,n}x_n &= b_m \end{aligned}$$

Hierbei sind $a_{i,j} \in \mathbb{R}$ und $b_i \in \mathbb{R}$ gegeben. Gesucht sind alle Werte für $x_1, \dots, x_n$, die alle Gleichungen simultan erfüllen.

Beim Hotel-Beispiel oben ergeben sich etwa folgende drei Gleichungen:

$$x_1 + x_2 + x_3 = 90 \tag{4}$$

$$\wedge \quad x_1 + 2x_2 + 3x_3 = 133 \tag{5}$$

$$\wedge \quad x_1 = 2(x_2 + x_3), \quad \text{also} \quad x_1 - 2x_2 - 2x_3 = 0 \tag{6}$$

Um sich Schreibarbeit zu sparen, schreibt man auch oft nur die Koeffizienten in eine Tabelle und nennt dieses Konstrukt dann **Matrix**. Allgemein sieht das dann so aus:

$$\begin{pmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,n} & b_1 \\ a_{2,1} & a_{2,2} & \cdots & a_{2,n} & b_2 \\ \vdots & & \ddots & & \vdots \\ a_{m,1} & a_{m,2} & \cdots & a_{m,n} & b_m \end{pmatrix}$$

Für das Hotel-Beispiel also

$$\begin{pmatrix} 1 & 1 & 1 & 90 \\ 1 & 2 & 3 & 133 \\ 1 & -2 & -2 & 0 \end{pmatrix}$$

Diese Matrix heißt auch **erweiterte Koeffizientenmatrix**. Erweitert deshalb, weil die Spalte mit den b_i aufgenommen wurde, kommt diese nicht vor, spricht man schlicht von der Koeffizientenmatrix.

Es handelt sich hierbei lediglich um eine geschickte Art der Notation, die überflüssige Schreibarbeit ersparen soll. Um ein Gleichungssystem aber lösen zu können, muss man mit den Gleichungen, also den Matrizen rechnen können. Zu diesem Zweck betrachtet man sogenannte **elementare Zeilenumformungen**, die in vier Typen daher kommen.

Typ 1: Multiplikation einer Zeile der Matrix mit einer reellen Zahl $\lambda \neq 0$. Diese Operation entspricht der Multiplikation der entsprechenden Gleichung mit λ und führt zu einer äquivalenten Gleichung.

Typ 2: Addition einer Zeile zu einer anderen. Diese Operation entspricht der Addition der Gleichungen.

Typ 3: Addition des λ -fachen einer Zeile zu einer anderen. Dies ist keine wirklich neue Zeilenoperation, sondern schlicht eine Kombination der ersten beiden Typen.

Typ 4: Vertauschung zweier Zeilen. Dies entspricht lediglich einer Vertauschung zweier Gleichungen.

Wie man sofort sieht (besonders bei 1 und 4) ändern diese Operationen zwar die Matrix, nicht aber die Lösungsmenge des Gleichungssystems. Die Frage ist jetzt, ob es eine Möglichkeit gibt, mit Hilfe der Zeilenumformungen die Matrix in eine Gestalt zu bringen, die es direkt erlaubt, die Lösungen zu bestimmen. Es gibt verschiedene Varianten, dies zu tun. Eine davon ist der sogenannte **Gauß-Algorithmus**:

Gegeben sei ein Gleichungssystem mit m Gleichungen und n Variablen und die erweiterte Koeffizientenmatrix wie oben. Der Algorithmus funktioniert folgendermaßen.

Schritt 1: Suche die erste Spalte mit einem von 0 verschiedenen Eintrag. Sollte es keine geben, ist der Algorithmus am Ende, die erweiterte Koeffizientenmatrix ist dann die sogenannte Nullmatrix, in der alle Einträge gleich 0 sind. Angenommen die erste Spalte mit einem von 0 verschiedenen Eintrag ist die Spalte j .

Schritt 2: Falls $a_{1,j} = 0$ führe eine Vertauschung von Zeilen (Typ 4) durch, so dass $a_{1,j} \neq 0$.

Schritt 3: Führe wiederholt Zeilenumformungen vom Typ 3 durch, um alle Einträge $a_{i,j}$ mit $i \geq 2$ zu 0 zu machen, indem zur Zeile i das $-\frac{a_{i,j}}{a_{1,j}}$ -fache von Zeile 1 addiert wird.

Schritt 4: Wende den Algorithmus auf die Untermatrix an, die aus der erweiterten Koeffizientenmatrix entsteht, indem die erste Zeile und die ersten j Spalten fortgelassen werden.

Am Ende des Algorithmus hat die Matrix eine spezielle Gestalt, die **Zeilenstufenform** genannt wird. Schreibt man diese wieder in Gleichungen um, so sieht man, dass sich die Zahl der Unbekannten von Gleichung zu Gleichung jedes Mal um mindestens eine reduziert. Hier ein Beispiel einer Matrix in Zeilenstufenform:

$$\begin{pmatrix} 0 & 0 & \mathbf{4} & 1 & 0 & -5 & 9 \\ 0 & 0 & 0 & 0 & \mathbf{-2} & 0 & -10 \\ 0 & 0 & 0 & 0 & 0 & \mathbf{14} & 2 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Die Ermittlung der Lösung geschieht mittels "rückwärts einsetzen": Betrachte die unterste Zeile, die keine Nullzeile ist (sagen wir Zeile i). Dann ergibt sich die Gleichung

$$a_{i,j}x_j + a_{i,j+1}x_{j+1} + \cdots + a_{i,n}x_n = b_i \quad \text{mit } a_{i,j} \neq 0$$

Werden nun x_{j+1} bis x_n beliebig gewählt (als freie Variablen), so ist x_j dadurch eindeutig bestimmt:

$$x_j = \frac{b_i - a_{i,j+1}x_{j+1} - \cdots - a_{i,n}x_n}{a_{i,j}}$$

Auf diese Weise erhält man alle Lösungen des Gleichungssystems.

Beispiel 7.1. Sehen wir uns dieses Verfahren am Hotelbeispiel an:

$$\begin{array}{ccc|c} 1 & 1 & 1 & 90 \\ 1 & 2 & 3 & 133 \\ 1 & -2 & -2 & 0 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & -3 & -3 & -90 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & 0 & 3 & 39 \end{array}$$

"Erste Zeile mal -1, dann zur dritten Zeile addieren"

An der dritten Zeile (im untersten Block) lesen wir ab: $3x_3 = 39$, also $x_3 = 13$. An der zweiten Zeile lesen wir (mit $x_3 = 13$) dann ab: $x_2 + 2 \cdot 13 = 43$, also $x_2 = 43 - 26 = 17$. Damit sehen wir in der ersten Zeile $x_1 + 17 + 13 = 90$, also $x_1 = 60$. Es gibt genau eine Lösung (weder keine noch zwei noch drei...) nämlich 60 Einbettzimmer, 17 Zweibettzimmer, 13 Dreibettzimmer.

Was passiert, hätte die dritte Bedingung in der Aufgabe gelautet "es gibt gleich viele Einzel- und Dreibettzimmer"? Die dritte Bedingung übersetzt sich dann zu $x_1 = x_3$ bzw $x_1 - x_3 = 0$. Das Vorgehen wie oben liefert

$$\begin{array}{ccc|c} 1 & 1 & 1 & 90 \\ 1 & 2 & 3 & 133 \\ 1 & 0 & -1 & 0 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & -1 & -2 & -90 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & 0 & 0 & -47 \end{array}$$

"Zweite Zeile zur dritten Zeile addieren"

Die dritte Zeile heißt $0 \cdot x_2 + 0 \cdot x_3 = -47$. Es gibt keine Werte für x_2 und x_3 , so dass diese Gleichung stimmt. Es gibt keine Lösung. (Ein solches Hotel kann es nicht geben!) Was passiert wiederum, wenn die dritte Bedingung lautete "es gibt 47 mehr Einzelzimmer als Doppelzimmer"? Also $x_1 = x_3 + 47$, bzw $x_1 - x_3 = 47$:

$$\begin{array}{ccc|c} 1 & 1 & 1 & 90 \\ 1 & 2 & 3 & 133 \\ 1 & 0 & -1 & 47 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & -1 & -2 & -43 \\ \hline 1 & 1 & 1 & 90 \\ 0 & 1 & 2 & 43 \\ 0 & 0 & 0 & 0 \end{array}$$

Die dritte Zeile heißt jetzt $0 \cdot x_2 + 0 \cdot x_3 = 0$. Das sagt gar nix. Diese Gleichung gilt für alle Werte von x_2 und x_3 . Also liefert diese Zeile keine Information. Die zweite Zeile sagt: $x_2 = 43 - 2x_3$, und die erste damit $x_1 = -x_2 - x_3 + 90 = -43 + 2x_3 - x_3 + 90 = 47 + x_3$. Damit erhalten wir viele Lösungen:

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 47 \\ 43 \\ 0 \end{pmatrix} \text{ oder } \begin{pmatrix} 48 \\ 41 \\ 1 \end{pmatrix} \text{ oder } \begin{pmatrix} 49 \\ 39 \\ 2 \end{pmatrix} \text{ oder } \begin{pmatrix} 50 \\ 37 \\ 3 \end{pmatrix} \text{ oder } \dots$$

Es kommt drauf an, welchen Wert wir (z.B.) dem x_3 zuweisen: Für $x_3 = 0$ ergibt sich die erste Lösung oben, für $x_3 = 1$ die zweite usw. Allgemein könnten wir ja auch schreiben

$$\begin{pmatrix} 47 + x_3 \\ 43 - 2x_3 \\ x_3 \end{pmatrix}$$

Hier müssen wir wegen des Kontexts noch aufpassen, dass keine negativen Zimmerzahlen auftreten. Wenn wir uns von der Hotel-Interpretation lösen, dann wären auch negative oder nicht-ganzzahlige Lösungen OK.

Eine Methode, die unendlich vielen Lösungen aufzuschreiben ist so:

$$\{(47 + x_3, 43 - 2x_3, x_3) \mid x_3 \in \mathbb{R}\}$$

Da steht also “die Menge aller Tripel, wo der letzte Wert x_3 eine beliebige (frei wählbare,...) reelle Zahl ist, und die anderen...” naja, in der beschriebenen Weise von x_3 abhängen.

Beispiel 7.2. Sehen wir uns ein Beispiel in den komplexen Zahlen an. Gegeben sei das LGS $x_1 + ix_2 - ix_3 = 0$, $(1 - i)x_1 + (1 + i)x_2 - (1 + i)x_3 = 0$, $-ix_1 + x_2 - x_3 = 0$. Diesmal sind gesucht alle Lösungen mit $x_1, x_2, x_3 \in \mathbb{C}$. Wir erhalten

$$\begin{array}{ccc|c} 1 & i & -i & 0 \\ 1-i & 1+i & -1-i & 0 \\ -i & 1 & -1 & 0 \\ \hline 1 & i & -i & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array}$$

Die letzten beiden Zeilen liefern uns die nutzlose Information $0 = 0$. In der ersten Zeile hängen die Variablen voneinander ab: $x_1 + ix_2 - ix_3 = 0$, also $x_1 = -ix_2 + ix_3$. Hier dürfen wir zwei Werte frei wählen. Nehmen wir x_2 und x_3 beliebig, dann ist x_1 dadurch festgelegt. (Z.B. wäre mit $x_2 = 1, x_3 = 1$ dann $x_1 = -i + i = 0$. Oder mit $x_2 = i, x_3 = 0$ wäre $x_1 = 1$.) Die Menge aller Lösungen dieses LGS können wir analog zu oben nun so schreiben:

$$\{(-ix_2 + ix_3, x_2, x_3) \mid x_2, x_3 \in \mathbb{C}\}$$

Im LGS aus dem Hotelbeispiel standen in der letzten Spalte Elemente, die nicht alle Null waren. Im letzten Beispiel mit den komplexen Zahlen standen rechts nur Nullen. Geben wir dem Unterschied zunächst einen Namen.

Definition 7.3. Ein LGS, dessen erweiterte Koeffizientenmatrix in der rechten Spalte nur Nullen enthält, heißt **homogenes** LGS.

Ein LGS, dessen erweiterte Koeffizientenmatrix in der rechten Spalte Zahlen enthält, die nicht alle Null sind, heißt **inhomogenes** LGS.

Dieser Unterschied hat Konsequenzen für die möglichen Lösungen.

- Ein homogenes LGS hat immer die Lösung $(0, 0, \dots, 0)$. Ein homogenes LGS hat entweder genau eine Lösung (dann also $(0, 0, \dots, 0)$) oder unendlich viele Lösungen. Bei unendlich vielen Lösungen gibt es einen frei wählbaren Parameter, oder zwei frei wählbare Parameter, oder... aber maximal so viele, wie es Unbekannte x_i gibt.
- Ein inhomogenes LGS hat entweder gar keine Lösung, oder genau eine Lösung, oder unendlich viele Lösungen. Bei unendlich vielen Lösungen gibt es einen frei wählbaren Parameter, oder zwei frei wählbare Parameter, oder... aber maximal soviel, wie es Unbekannte x_i gibt. Die Lösungen eines inhomogenen LGS sind von der Form

eine Lösung des inhomogenen LGS + allgemeine Lösung des zugehörigen homogenen LGS

Dabei heißt “zugehöriges homogenes LGS” einfach, dass die Zahlen in der rechten Spalte der Koeffizientenmatrix alle auf Null gesetzt werden.

Die Bedeutung von “frei wählbare Parameter” kann man sich auch an den Beispielen oben klar machen. Im (dritten) Hotelbeispiel konnte das x_3 beliebig gewählt werden (ein frei wählbarer Parameter), die anderen Parameter x_1 und x_2 waren damit dann festgelegt. Im komplexen LGS konnten zwei Parameter, nämlich x_2 und x_3 , beliebig gewählt werden, x_1 ist dann dadurch festgelegt.

Im Hotelbeispiel ist eine Lösung des LGS das Tripel $(47, 43, 0)$. Die allgemeine Lösung hat die Form

$$(47, 43, 0) + (x_3, -2x_3, x_3).$$

Das heißt also, dass $(x_3, -2x_3, x_3)$, oder besser $\{(x_3, -2x_3, x_3) \mid x_3 \in \mathbb{R}\}$, die allgemeine Lösung des zugehörigen homogenen LGS sein muss. Das stimmt auch: nachrechnen liefert genau das.

Aufgabe 7.1. Eine Mutter und ihre Tochter sind zusammen 62 Jahre alt. Vor sechs Jahren war die Mutter viermal so alt wie damals die Tochter. Wie alt ist jeder?

Eine Mutter ist 21 Jahre älter als ihr Kind. In sechs Jahren wird die Mutter fünfmal so alt sein wie ihr Kind. Wo ist der Vater?

Aufgabe 7.2. Drei Schwerter kosten soviel wie 12 Langbögen und 40 Dolche zusammen. Ein Schwert und ein Langbogen kosten zusammen soviel wie 80 Dolche. Ein Schwert kostet doppelt soviel wie 1 Langbogen und 20 Dolche zusammen. Kann man daraus schließen, was jede einzelne Waffe kostet? Kann man auf das Verhältnis zwischen den Preisen schließen? Kann man diese Fragen beantworten, *bevor* man losrechnet?

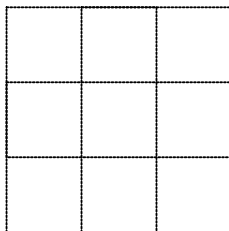
Aufgabe 7.3. Finde alle Lösungen mit $x_1, x_2, x_3 \in \mathbb{C}$ für folgendes Gleichungssystem:

$$\begin{aligned}(1 + i)x_1 + ix_2 + x_3 &= 1 \\ x_1 + x_2 + x_3 &= 1\end{aligned}$$

Aufgabe* 7.4. Finde alle Lösungen mit $x_1, x_2, x_3, x_4 \in \mathbb{R}$ für folgendes Gleichungssystem.

$$\begin{aligned} & x_2 + 3x_3 + 3x_4 = 1 \\ \wedge & x_1 + 3x_2 - 5x_3 + 9x_4 = -3 \\ \wedge & -2x_1 - 5x_2 + 12x_3 - 23x_4 = 8 \\ \wedge & 3x_1 + 11x_2 - 8x_3 + 41x_4 = -8 \\ \wedge & 2x_1 + 7x_2 - 7x_3 + 22x_4 = 0 \end{aligned}$$

Aufgabe* 7.5. Gegeben sei ein Quadrat mit 3×3 Kästchen:



Diese Kästchen sollen mit rationalen Zahlen gefüllt werden, so dass die Summe in horizontaler, vertikaler und diagonaler Richtung stets den gleichen Wert $a \in \mathbb{Q}$ ergibt. Geht das für jeden Wert a ? Wie viele Möglichkeiten gibt es für gegebenes a ? Für welche a können die Einträge ganzzahlig gewählt werden?

Noch zwei zum Knobeln:

1. Drei Damen unterschiedlichen Alters sitzen in der Stadtbahn und unterhalten sich. Die erste sagt zur zweiten: "Dein Alter ist ja genau das Quadrat von meinem Alter." Daraufhin die zweite zur dritten: "In drei Jahren bin ich gerade mal halb so alt, wie du jetzt bist!". "In drei Jahren", ruft die erste begeistert der dritten zu, "wird dein Alter genau das Quadrat von meinem Alter sein!" Wie alt sind die drei? (Dies ist ein Beispiel für ein *nicht-lineares* Gleichungssystem, weil nicht nur $a, b, c, \dots$ vorkommen, sondern auch a^2 .)
2. Daneben sitzen drei Herren, die sich ebenfalls unterhalten. Der erste: "Wenn ich so alt bin, wie ihr beiden jetzt zusammen seid, wirst du", damit zeigt er auf den zweiten, "doppelt so alt sein, wie ich jetzt bin." Der daraufhin: "Und wenn ich so alt bin, wie ihr beiden zusammen jetzt seid, wirst du", er zeigt auf den dritten, "doppelt so alt sein, wie ich jetzt bin." "Das ist ja lustig!", bemerkt da die Jüngste vom Nachbarsitz. Warum?

I Reelle Komplexe Zahlen

Siehe Kap 6.

II Vektorräume

Zwei Standardbeispiele für Vektorräume sind

$$\mathbb{R}^2 = \left\{ \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} \mid v_1, v_2 \in \mathbb{R} \right\} \quad \text{bzw.} \quad \mathbb{R}^3 = \left\{ \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \mid v_1, v_2, v_3 \in \mathbb{R} \right\}$$

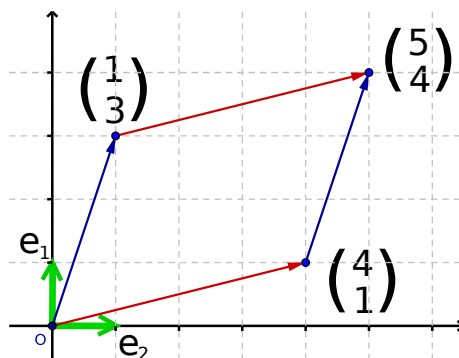
Es ist für spätere Zwecke günstig, Vektoren nicht als Zeilenvektoren zu schreiben: $(2, 0, -1)$, sondern als Spaltenvektoren: $\begin{pmatrix} 2 \\ 0 \\ -1 \end{pmatrix}$.

Elemente im $\mathbb{R}^2$ kann man sich als Punkte in einer (unendlichen) Ebene vorstellen. Die Koordinaten sind dann x - und y -Koordinate bzw Ost-West- und Nord-Süd-Koordinate. Elemente im $\mathbb{R}^3$ kann man sich als Punkte im Raum vorstellen, die dritte Koordinate ist dann etwa für oben-unten. Dabei muss es einen “Nullpunkt” geben, von wo aus man misst. Dieser Nullpunkt heißt auch oft “Ursprung”. Dann bedeutet der Vektor $\begin{pmatrix} 3 \\ 7 \end{pmatrix}$ einfach eine Verschiebung um 3 Einheiten in x -Richtung und um 7 Einheiten in y -Richtung.

Graphisch wird dies oft durch einen Pfeil symbolisiert. Mit diesem Bild im Kopf ist auch die graphische Veranschaulichung der Vektoraddition nicht schwer: sind v und w Vektoren, z.B.

$$v = \begin{pmatrix} 7 \\ 1 \end{pmatrix} \quad w = \begin{pmatrix} -3 \\ 2 \end{pmatrix} \quad \Rightarrow \quad v + w = \begin{pmatrix} 4 \\ 3 \end{pmatrix}$$

dann beschreibt der Summenvektor $v + w$ was geschieht, wenn die beiden Verschiebungen nacheinander (gleich in welcher Reihenfolge) ausgeführt werden.



Zeichnerisch stellt man das so dar, dass die beiden Vektoren aneinander gezeichnet werden, als “durchliefe man die Strecke”. Für $\mathbb{R}^3$ gilt dasselbe Bild im Raum.

Vektoren kann man also addieren. Man möchte auch Vektoren mit (reellen) Zahlen multiplizieren dürfen: 2-mal ein Vektor soll halt gerade das doppelte sein, (-1)-mal ein Vektor soll der Vektor in Gegenrichtung sein. Es ist also einfach

$$2 \begin{pmatrix} 3 \\ 7 \end{pmatrix} = \begin{pmatrix} 6 \\ 14 \end{pmatrix}, \quad \text{oder} \quad (-1) \cdot \begin{pmatrix} 3 \\ 7 \end{pmatrix} = \begin{pmatrix} -3 \\ -7 \end{pmatrix} \quad \text{oder} \quad \frac{1}{2} \cdot \begin{pmatrix} 3 \\ 6 \end{pmatrix} = \begin{pmatrix} \frac{3}{2} \\ 3 \end{pmatrix}$$

Allgemein ist

$$a \cdot \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} a \cdot v_1 \\ a \cdot v_2 \end{pmatrix}.$$

Formal definiert man einen Vektorraum dann so:

Definition 7.4. Gegeben ein Körper $\mathbb{K}$ (meistens $\mathbb{R}$ oder $\mathbb{C}$). Eine Menge V mit einer Verknüpfung $+$: $V \times V \rightarrow V$ und einer Verknüpfung $\cdot$: $\mathbb{K} \times V \rightarrow V$ heißt **Vektorraum**, falls

- $(V, +)$ ist abelsche Gruppe.
- $\forall a \in \mathbb{K}, v, w \in V : a(v + w) = av + aw.$
- $\forall a, b \in \mathbb{K}, v \in V : (a + b)v = av + bv.$

- $\forall a, b \in \mathbb{K}, v \in V : (ab)v = a(bv)$.
- $\forall v \in V : 1 \cdot v = v$ (wobei 1 das neutrale Element bzgl. Multiplikation in $\mathbb{K}$ ist)

Ein wichtiger Begriff ist nun der der **Basis**. Man möchte eine möglichst kleine Menge $\{b_1, b_2, \dots, b_n\}$ von Vektoren finden, so dass man alle Vektoren im Vektorraum damit darstellen kann. “Darstellen” heißt: als Linearkombination schreiben. Ein Vektor v heißt **Linearkombination** von $w_1, \dots, w_n$ falls es $a_i \in \mathbb{R}$ (oder allgemeiner in $\mathbb{K}$) gibt mit

$$v = \sum_{i=1}^n a_i w_i = a_1 w_1 + a_2 w_2 + \dots + a_n w_n$$

Z.B. ist $\begin{pmatrix} 3 \\ 7 \end{pmatrix}$ Linearkombination von $\begin{pmatrix} 3 \\ -1 \end{pmatrix}$ und $\begin{pmatrix} -2 \\ 2 \end{pmatrix}$, denn

$$\begin{pmatrix} 3 \\ 7 \end{pmatrix} = 5 \begin{pmatrix} 3 \\ -1 \end{pmatrix} + 6 \begin{pmatrix} -2 \\ 2 \end{pmatrix}$$

Eine Basis ist eine Menge von Vektoren, die (1.) groß genug ist, um alle Vektoren in V als Linearkombination davon zu schreiben, und (2.) so klein wie möglich.

Eine Basis für den Vektorraum $\mathbb{R}^2$ ist z.B.

$$\left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\}.$$

Denn jeden Vektor $\begin{pmatrix} a \\ b \end{pmatrix} \in \mathbb{R}^2$ kann ich schreiben als $a \begin{pmatrix} 1 \\ 0 \end{pmatrix} + b \begin{pmatrix} 0 \\ 1 \end{pmatrix}$. Und nicht anders! Ich kann jeden Vektor nur auf eine Weise als Linearkombination der zwei Basisvektoren schreiben.

Eine Basis für den Vektorraum $\mathbb{R}^3$ ist z.B.

$$\left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\}.$$

Denn jeden Vektor $\begin{pmatrix} a \\ b \\ c \end{pmatrix} \in \mathbb{R}^3$ kann ich schreiben als $a \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + b \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + c \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$. Und nicht anders! Ich kann jeden Vektor nur auf eine Weise als Linearkombination der drei Basisvektoren schreiben.

Das “nicht zu groß” präzisiert man genau so: Bzgl einer Basis soll jeder Vektor nur auf eine einzige Weise als Linearkombination geschrieben werden können.

Definition 7.5. Eine Menge $\{b_1, \dots, b_n\}$ von Vektoren ist eine **Basis** eines Vektorraums V , wenn alle $v \in V$ eindeutig als Linearkombination der Vektoren $b_1, \dots, b_n$ geschrieben werden können.

Entfällt das “eindeutig”, dann heißt $\{b_1, \dots, b_n\}$ nur **Erzeugendensystem**. In der Vorlesung wird eine Basis später etwas anders lautend (aber natürlich gleichbedeutend) definiert als *linear unabhängiges* Erzeugendensystem. Vektoren $\{b_1, \dots, b_n\}$ heißen **linear unabhängig**, wenn sich keines der b_i als Linearkombination der anderen schreiben lässt. Oder gleichbedeutend: $\{b_1, \dots, b_n\}$ sind linear unabhängig, falls

$$\alpha_1 b_1 + \alpha_2 b_2 + \dots + \alpha_n b_n = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \Rightarrow \alpha_1 = \alpha_2 = \dots = \alpha_n = 0.$$

Gerade dieser letzte Begriff ist gewöhnungsbedürftig. Daher ist es evtl hilfreich, die Erklärung oben parat zu haben. Letztendlich gewöhnt man sich an diese Begriffe durch Beispiele.

Beispiel 7.6. (Ins Publikum fragen) Welche der folgenden Mengen sind Basen von $\mathbb{R}^2$ bzw $\mathbb{R}^3$? Welche sind Erzeugendensysteme?

$$\left\{\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}\right\}, \quad \left\{\begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}\right\}, \quad \left\{\begin{pmatrix} -1 \\ -1 \end{pmatrix}, \begin{pmatrix} 2 \\ 2 \end{pmatrix}\right\}, \quad \left\{\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}\right\}?$$

Bei größeren Beispielen ist es kniffliger, das zu sehen. In der Vorlesung werden Methoden gezeigt, wie man prüft, ob eine Menge von Vektoren eine Basis ist. Weiterhin kommen dann viele Resultate über Basen, z.B.

Satz 7.7 (und Definition). *Jede Basis des $\mathbb{R}^n$ (oder allgemeiner $\mathbb{K}^n$) hat n Elemente. Die Zahl n heißt **Dimension** des Vektorraums.*

Leider gibt es auch unendlich-dimensionale Vektorräume, aber die kommen normalerweise erst in höheren Semestern dran (wenn überhaupt).

III Lineare Abbildungen und Matrizen

Im Analysisteil von Mathe I tauchen sehr schöne, aber auch sehr komplizierte Funktionen auf wie $\arcsin$, $\frac{e^x}{\sqrt{1-x^2}}$ oder $\ln(x\sqrt{x}+1)x^2 \exp(-x^2)$. In Linearer Algebra sind die Mengen etwas komplizierter ($\mathbb{R}^n$ als Definitionsbereich), aber die Abbildungen, um die es geht, sind sehr viel einfacher: es sind **lineare Abbildungen**.

Definition 7.8. Eine Abbildung $f : V \rightarrow W$ heißt **linear**, wenn für alle $u, v \in V$ und alle $a \in \mathbb{R}$ gilt:

- $f(u + v) = f(u) + f(v)$, und
- $f(av) = af(v)$.

(Publikum fragen:) Sind die folgenden Abbildungen linear?

- $f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = x + 1$.
- $f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = 2x$.
- $f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$.
- $f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = \sin(x)$.

In Linearer Algebra geht es nun um lineare Abbildungen von $\mathbb{R}^n$ nach $\mathbb{R}^m$. Es wird sich herausstellen:

Satz 7.9. *Jede lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$ wird von einer $m \times n$ -Matrix beschrieben. Jede $m \times n$ -Matrix liefert eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$.*

Also brauchen wir jetzt Matrizen.

IV Matrizen und lineare Gleichungssysteme

Eine Matrix haben wir oben schon gesehen im Zusammenhang mit LGS, die *erweiterte Koeffizientenmatrix*. Allgemein ist eine **Matrix** einfach eine Tabelle mit n Zeilen und m Spalten, deren Einträge reelle Zahlen sind. Zum Beispiel sind

$$\begin{pmatrix} 1 & 2 & 3 \\ -2 & 12 & 23 \end{pmatrix}, \quad \begin{pmatrix} 5 & \pi & -0,1 \\ 7 & 1000 & -\pi \\ 6 & 8 & -10 \end{pmatrix}, \quad \text{und} \quad \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 0 & 0 \end{pmatrix}$$

Matrizen. Auch Matrizen lassen sich addieren und subtrahieren, wenn sie das selbe Format haben. Zum Beispiel

$$\begin{pmatrix} 1 & -3 & 2 \\ 1 & 2 & 7 \end{pmatrix} + \begin{pmatrix} 0 & 3 & 5 \\ 2 & 1 & -1 \end{pmatrix} = \begin{pmatrix} 1+0 & -3+3 & 2+5 \\ 1+2 & 2+1 & 7+(-1) \end{pmatrix} = \begin{pmatrix} 1 & 0 & 7 \\ 3 & 3 & 6 \end{pmatrix}.$$

Die Menge aller Matrizen mit n Zeilen und m Spalten, deren Einträge reelle Zahlen sind, bezeichnet man als $\mathbb{R}^{n \times m}$. Matrizen werden wir hier oft mit Großbuchstaben $A, B, \dots$ benennen. Die Einträge der Matrix A in Zeile i und Spalte j bezeichnen wir mit $a_{i,j}$. So ist zum Beispiel in der Matrix $\begin{pmatrix} 1 & 2 & 3 \\ -2 & 12 & 23 \end{pmatrix}$ etwa $a_{1,1} = 1$, $a_{1,3} = 3$ und $a_{2,2} = 12$. Wenn die Zahl der Spalten und Zeilen (wie hier) klein ist, dann lassen wir das Komma zwischen i und j weg. Also a_{12} statt $a_{1,2}$. Manchmal ist es praktisch, statt der Bezeichnung A für eine Matrix diese so zu schreiben: $(a_{ij})_{i=1,\dots,n; j=1,\dots,m}$, oder kurz (a_{ij}) . Damit lässt sich etwa die Regel zur Addition von Matrizen sehr kompakt hinschreiben:

$$A + B := (a_{ij} + b_{ij})_{i=1,\dots,n; j=1,\dots,m}$$

Die Regel zur Skalarmultiplikation lautet in dieser kompakten Schreibweise einfach $\lambda \cdot A := (\lambda \cdot a_{ij})_{i=1,\dots,m; j=1,\dots,n}$. Zum Beispiel:

$$5 \cdot \begin{pmatrix} 1 & -3 & 2 \\ 1 & 2 & 7 \end{pmatrix} = \begin{pmatrix} 5 \cdot 1 & 5 \cdot (-3) & 5 \cdot 2 \\ 5 \cdot 1 & 5 \cdot 2 & 5 \cdot 7 \end{pmatrix} = \begin{pmatrix} 5 & -15 & 10 \\ 5 & 10 & 35 \end{pmatrix}.$$

Im Gegensatz zu Vektoren kann man Matrizen auch miteinander multiplizieren, also “Matrix mal Matrix”, nicht nur “Zahl mal Matrix”. Zumindest, wenn ihr Format passt. Genauer:

Zwei Matrizen können multipliziert werden, wenn die Spaltenanzahl der linken mit der Zeilenanzahl der rechten Matrix übereinstimmt.

Ist also A aus $\mathbb{R}^{\ell \times m}$ und B aus $\mathbb{R}^{m \times n}$, so kann man “A mal B” rechnen. Nennen wir das Ergebnis C , also $C = A \cdot B$, so ist

$$c_{ij} = a_{i1}b_{1j} + a_{i2}b_{2j} + \dots + a_{im}b_{mj}.$$

Ein Beispiel:

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} \cdot \begin{pmatrix} 6 & -1 \\ 3 & 2 \\ 0 & -3 \end{pmatrix} = \begin{pmatrix} 1 \cdot 6 + 2 \cdot 3 + 3 \cdot 0 & 1 \cdot (-1) + 2 \cdot 2 + 3 \cdot (-3) \\ 4 \cdot 6 + 5 \cdot 3 + 6 \cdot 0 & 4 \cdot (-1) + 5 \cdot 2 + 6 \cdot (-3) \end{pmatrix} = \begin{pmatrix} 12 & -6 \\ 39 & -12 \end{pmatrix}$$

- Achtung: Im Allg. ist $A \cdot B \neq B \cdot A$! (Matrizenmultiplikation ist nicht *kommutativ*.)
- Immerhin gilt $A \cdot (B \cdot C) = (A \cdot B) \cdot C$. (Matrizenmultiplikation ist *assoziativ*.)

Wie erwähnt ist ein zentrales Ziel der Vorlesung Mathe I Satz 7.9 auf Seite 72. Da wird noch viel mehr gemacht, aber hier wollen wir nur einen groben Überblick geben.

Aufgabe* 7.6. Ist $(\mathbb{R}^{n \times n}, +)$ eine Gruppe? Ist $(\mathbb{R}^{n \times n}, \cdot)$ eine Gruppe? Begründe Deine Antwort, im positiven Fall durch den Nachweis der Gruppenaxiome, im negativen Fall durch ein konkretes Beispiel, welches Gruppenaxiom verletzt ist. Im positiven Fall: Ist die Gruppe abelsch?

In jedem Fall: Finde zwei Matrizen A, B in $\mathbb{R}^{n \times n}$, so dass $A \cdot B \neq B \cdot A$ ist.

Aufgabe 7.7. Berechne für $A = \begin{pmatrix} 1 & -2 & 4 \\ -2 & 3 & -5 \end{pmatrix}$, $B = \begin{pmatrix} 2 & 4 \\ 3 & 6 \\ 1 & 2 \end{pmatrix}$, $C = \begin{pmatrix} -2 & 1 \\ 0 & 3 \end{pmatrix}$ die Ausdrücke AB , BA , AC , CA , BC , CB , $A^T C$, $C^T A$, (zu A^T s.u.: "Transponierte"), ABC und CBA , falls es möglich ist.

Die **Transponierte** einer $n \times m$ -Matrix $A = (a_{ij})$ ist die $n \times m$ -Matrix $A^T = (a_{ji})$. Das heißt:
Zu

$$A = \begin{pmatrix} a_{11} & \dots & a_{1m} \\ \vdots & \ddots & \vdots \\ a_{n1} & \dots & a_{nm} \end{pmatrix} \quad \text{ist} \quad A^T = \begin{pmatrix} a_{11} & \dots & a_{n1} \\ \vdots & \ddots & \vdots \\ a_{1m} & \dots & a_{nm} \end{pmatrix}$$

die Transponierte. Man schreibt also die erste Zeile als erste Spalte, die zweite Zeile als zweite Spalte usw.

8 Überblick Mathe I Analysis

Es gibt nichts Praktischeres als eine gute Theorie
Bob der Baumeister

Im Analysisteil der Vorlesung Mathe I passiert in etwa Folgendes:

I Konvergenz von Folgen und Reihen

- Konvergenz
- (Unendliche) Reihen, Exponentialfunktion
- Cauchyfolgen, Vollständigkeit

II Stetige Funktionen

- Abbildungen, in-, sur- und bijektiv
- ε - δ -Kriterium, folgenstetig
- gleichmäßig stetig, Funktionenfolgen
- Potenzreihen

III Differenzierbarkeit

- Definition Ableitung, Regeln (Kettenregel, Quotientenregel, Ableitung der Umkehrfunktion...)
- Extrema, Mittelwertsatz, ℓ' Hôpital
- Taylorreihen

I Folgen

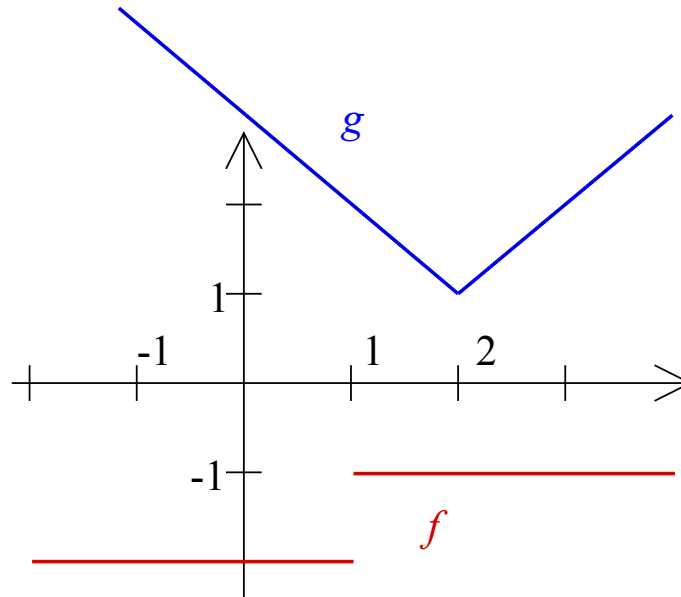
Siehe Kapitel 4. In der Vorlesung werden auch Cauchyfolgen behandelt. Das ist ein alternativer Konvergenzbegriff. In $\mathbb{R}$ ist der äquivalent zu unserem. Der Witz bei Cauchyfolgen wird erst viel später klar (wenn überhaupt, als Informatiker braucht man eher selten Cauchyfolgen).

II Stetige Funktionen

Betrachten wir zunächst folgende Abbildungen:

$$\begin{aligned} f : \mathbb{R} \rightarrow \mathbb{R} \quad \text{gegeben durch} \quad f(x) &= \begin{cases} -2 & , \text{ falls } x < 1 \\ -1 & , \text{ falls } x \geq 1 \end{cases} \\ g : \mathbb{R} \rightarrow \mathbb{R} \quad \text{gegeben durch} \quad g(x) &= \begin{cases} -x + 3 & , \text{ falls } x < 2 \\ x - 1 & , \text{ falls } x \geq 2 \end{cases} \end{aligned}$$

Beide Graphen sind in folgendem Koordinatensystem dargestellt.



Auffällig ist, dass zwar beide Funktionen “stückweise” definiert sind, aber der Graph von g eher “zusammenhängt”, d.h. in einem Stück gezeichnet werden kann, während man beim Graphen von f eine Art Sprungstelle hat. Die Funktion f ist also unstetig am Punkt 1, wohingegen g überall stetig ist. Formal definiert man den Begriff der Stetigkeit folgendermaßen:

Definition 8.1. Sei $D \subseteq \mathbb{R}$, $f : D \rightarrow \mathbb{R}$ eine Abbildung und $a \in D$. Die Abbildung f heißt **stetig** an der Stelle a , falls

$$\forall \varepsilon > 0 \exists \delta > 0 \forall x \in D : |x - a| < \delta \Rightarrow |f(x) - f(a)| < \varepsilon.$$

Ist eine Funktion f stetig an allen Punkten $a \in D$, so spricht man kurz von einer stetigen Funktion.

In Worten: Stetigkeit an der Stelle a bedeutet, dass die Bilder von Punkten “in der Nähe” von a auch nahe am Bildpunkt von a liegen. Das δ spielt in der Definition in gewisser Weise die Rolle von n_0 in der Definition von Folgenkonvergenz.

Es ist nicht schwer, formal nachzuweisen, dass die obige Beispielfunktion f nicht stetig im Punkt 1 ist:

Beweis. Sei $a = 1$ und wähle $\varepsilon := \frac{1}{2}$. Es genügt zu zeigen, dass es für dieses spezielle ε kein δ gibt, das die obige Bedingung erfüllt. Sei also $\delta > 0$ beliebig. Betrachte den Punkt $x := 1 - \frac{\delta}{2}$. Es gilt

$$x < 1 \text{ und } |x - a| = |x - 1| = \left| 1 - \frac{\delta}{2} - 1 \right| = \left| -\frac{\delta}{2} \right| = \frac{\delta}{2} < \delta.$$

Aber $f(x) = -2$ und $f(a) = -1$, also folgt

$$|f(x) - f(a)| = |-2 + 1| = |-1| = 1 > \frac{1}{2} = \varepsilon.$$

Damit ist die Bedingung der Stetigkeit verletzt. □

Ein weiteres Beispiel: die Wurzelfunktion $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$ gegeben durch $f(x) = \sqrt{x}$ ist stetig auf ganz $\mathbb{R}_0^+$.

Beweis. Zeige zunächst Stetigkeit im Punkt $a = 0$. Sei $\varepsilon > 0$ beliebig und wähle $\delta := \varepsilon^2$. Dann ist $\delta > 0$ und für $x \in \mathbb{R}_0^+$ mit $|x - a| = |x| = x < \delta$ gilt:

$$|\sqrt{x} - \sqrt{0}| = \sqrt{x} < \sqrt{\delta} = \sqrt{\varepsilon^2} = \varepsilon$$

Also ist f stetig im Punkt $a = 0$.

Sei nun $a > 0$ beliebig, ebenso wie $\varepsilon > 0$. Wähle nun $\delta := \varepsilon\sqrt{a}$. Dann ist $\delta > 0$. Sei nun $x \in \mathbb{R}_0^+$ mit $|x - a| < \delta$. Es folgt

$$|f(x) - f(a)| = |\sqrt{x} - \sqrt{a}| = \left| \frac{x - a}{\sqrt{x} + \sqrt{a}} \right| \leq \frac{|x - a|}{\sqrt{a}} < \frac{\delta}{\sqrt{a}} = \frac{\varepsilon\sqrt{a}}{\sqrt{a}} = \varepsilon.$$

Also ist f stetig im Punkt a und da $a > 0$ beliebig war, ist f damit stetig auf ganz $\mathbb{R}_0^+$. $\square$

Die meisten der in der Analysis betrachteten Funktionen sind stetig, z.B. alle Polynome. Der Beweis ist nicht ganz so einfach und wird in der Vorlesung Mathe I geführt (wahrscheinlich). Dabei werden Sätze der folgenden Art benutzt.

Satz 8.2. Seien $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ stetige Funktionen und $\alpha \in \mathbb{R}$ eine beliebige Zahl. Definiere dann

$$\begin{aligned} (f + g) : \mathbb{R} &\rightarrow \mathbb{R} && \text{gegeben durch} && (f + g)(x) = f(x) + g(x) \\ (f \cdot g) : \mathbb{R} &\rightarrow \mathbb{R} && \text{gegeben durch} && (f \cdot g)(x) = f(x) \cdot g(x) \\ (\alpha \cdot f) : \mathbb{R} &\rightarrow \mathbb{R} && \text{gegeben durch} && (\alpha \cdot f)(x) = \alpha \cdot f(x). \end{aligned}$$

Diese Funktionen sind dann alle stetig.

Nun haben wir fast alles beisammen, um die Stetigkeit aller Polynomfunktionen beweisen zu können. Es fehlen nur noch zwei Beobachtungen, die im folgenden Lemma zusammengefasst sind.

Lemma 8.3. Sei $c \in \mathbb{R}$ eine reelle Zahl. Betrachte die Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ gegeben durch $f(x) = c$ und $\text{id}_{\mathbb{R}} : \mathbb{R} \rightarrow \mathbb{R}$ gegeben durch $\text{id}_{\mathbb{R}}(x) = x$. Dann sind f und $\text{id}_{\mathbb{R}}$ stetig.

Insgesamt folgt nun für eine beliebige Polynomfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \sum_{k=0}^n a_k x^k$, dass f stetig ist, denn f lässt sich schreiben als Summe von Produkten von konstanten Funktionen und der Funktion $\text{id}_{\mathbb{R}}$ und nach Satz 8.2 folgt die Stetigkeit von f .

Folgenstetigkeit

Das ε - δ -Kriterium für die Stetigkeit einer Funktion ist zwar sehr nützlich, manchmal aber etwas unhandlich. In den meisten Fällen bietet sich die alternative Formulierung der Folgenstetigkeit an.

Dazu erst mal eine Schreibweise: ist $f : \mathbb{R} \rightarrow \mathbb{R}$ eine Funktion und $a \in \mathbb{R}$, so kann für eine beliebige Folge $(x_n)_{n \in \mathbb{N}}$ mit Grenzwert a der folgende Grenzwert betrachtet werden:

$$\lim_{n \rightarrow \infty} f(x_n).$$

Für den Fall, dass dieser Grenzwert existiert und unabhängig ist von der gewählten Folge $(x_n)_{n \in \mathbb{N}}$, d.h. wenn dieser Wert für jede Folge mit Grenzwert a gleich ist, dann sagt man, dass der Grenzwert “für x gegen a ” existiert und schreibt kurz

$$\lim_{x \rightarrow a} f(x).$$

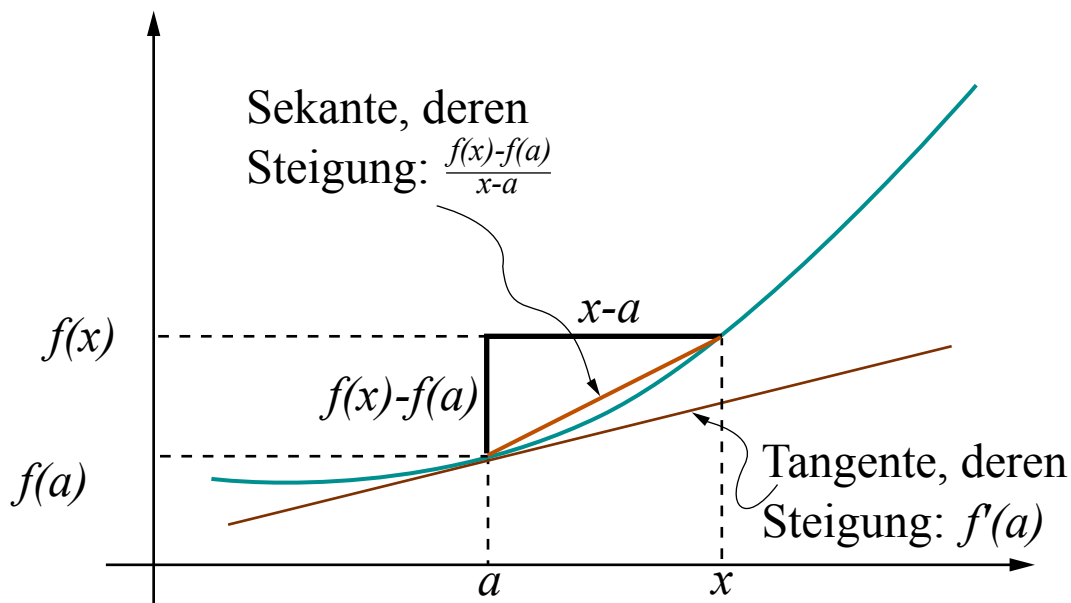
Satz 8.4. Sei $D \subseteq \mathbb{R}$ und $f : D \rightarrow \mathbb{R}$ eine Funktion. Sei weiter $a \in D$. Dann gilt: f ist genau dann stetig im Punkt a , wenn

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

Dieser Satz enthält eine Äquivalenz, also eine “Genau dann, wenn”-Aussage. Insofern zeigt sich, dass man Stetigkeit auch über das Folgenkriterium hätte definieren können – beide Begriffe sind absolut gleichwertig.

III Differentialrechnung

Für viele Zwecke (bestes Beispiel: Bewegungsgesetze in der Physik) ist es gewinnbringend, die Steigung einer Funktion (bzw. ihres Graphen) in jedem einzelnen Punkt zu kennen. Die Steigung zwischen zwei Punkten bekommt man aus dem *Steigungsdreieck*: Die (durchschnittliche) Steigung der Funktion f zwischen den Stellen x und a bekommt man als “Höhe durch Breite des Steigungsdreiecks”. Ist etwa $a = 3$, $x = 5$ sowie $f(a) = f(3) = 1$, $f(x) = f(5) = 2$, dann ist die (durchschnittliche) Steigung von f zwischen den Stellen 3 und 5 gerade $\frac{2-1}{5-3} = \frac{1}{2}$. (Nebenbei: Was heißt dann wohl “eine Straße hat 14% Steigung” in Grad? Was heißt “ein Geländewagen bewältigt 80% Steigung?”).



Allgemein ist die Steigung von f zwischen x und a also $\frac{f(x)-f(a)}{x-a}$. Wenn wir nun an der Steigung von f in einem Punkt interessiert sind, liegt es doch nahe, einfach x gegen a laufen zu lassen, genauer also: den Grenzwert von $\frac{f(x)-f(a)}{x-a}$ für x gegen a zu betrachten:

$$\lim_{x \rightarrow a} \frac{f(x) - f(a)}{x - a}.$$

Falls dieser existiert (das heißt hier: für alle Folgen $(x_n)_{n \in \mathbb{N}}$, die gegen a konvergieren, muss der Grenzwert existieren und alle diese Folgen müssen den gleichen Grenzwert liefern), so wird er als Steigung der Tangente oder auch **Ableitung** im Punkt a bezeichnet. Die Schreibweise dafür ist $f'(a)$ und die Funktion f heißt dann im Punkt a **differenzierbar**. Wie bei Stetigkeit auch nennt man eine Funktion schlicht differenzierbar, wenn sie in jedem ihrer Punkte differenzierbar ist.

Angenommen, die Funktion oben wäre $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = \frac{x^2}{4} - x + \frac{7}{4}$. Dann ergibt sich ganz allgemein für $f'(x)$:

$$\frac{f(x) - f(a)}{x - a} = \frac{\frac{x^2}{4} - x + \frac{7}{4} - (\frac{a^2}{4} - a + \frac{7}{4})}{x - a} = \frac{1}{4} \cdot \frac{x^2 - 4x - a^2 + 4a}{x - a} = \frac{1}{4} \frac{x^2 - a^2}{x - a} - \frac{x - a}{x - a} = \frac{1}{4}(x+a) - 1.$$

Also gilt

$$f'(a) = \lim_{x \rightarrow a} \frac{1}{4}(x+a) - 1 = \frac{1}{4}(a+a) - 1 = \frac{1}{2}a - 1.$$

Oder (gleichbedeutend, aber schöner) $f'(x) = \frac{1}{2}x - 1$.

Das passt zu den Regeln, die in der Schule gelernt werden. Aber hier sieht man auch, warum. Ein paar weitere Beispiele enthalten die Aufgaben unten. Ein ganz einfaches Beispiel ist die Ableitung einer konstanten Funktion: Sei $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = c$, wobei $c \in \mathbb{R}$. Dann ist

$$\frac{f(x) - f(a)}{x - a} = \frac{c - c}{x - a} = \frac{0}{x - a} = 0,$$

und $\lim_{x \rightarrow a} 0 = 0$. Also ist $f'(x) = 0$.

Differenzierbar ist übrigens eine stärkere Eigenschaft als stetig. Genauer:

Satz 8.5. Sei $D \subseteq \mathbb{R}$ und $f : D \rightarrow \mathbb{R}$ eine Funktion, die im Punkt $a \in D$ differenzierbar ist. Dann ist f stetig in a .

Bemerkung 8.6. Noch eine Bemerkung: manchmal nennt man den Term $x - a$ einfach h , man definiert also $h := x - a$. Dann sieht der Differenzenquotient etwas anders aus, beschreibt aber genau den gleichen Sachverhalt, es ist nämlich

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

Aufgabe 8.1. Gegeben sind $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x$; $g : \mathbb{R} \rightarrow \mathbb{R}$, $g(x) = -5x$ und $h : \mathbb{R} \rightarrow \mathbb{R}$, $h(x) = c \cdot x$ ($c \in \mathbb{R}$). Bestimme für einen beliebigen Punkt $a \in \mathbb{R}$ die Ableitung von f , g und h nur mit Hilfe des Differenzenquotienten.

Aufgabe 8.2. Gegeben sind $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^2$; $g : \mathbb{R} \rightarrow \mathbb{R}$, $g(x) = -2 \cdot x^2 + 5$ und $h : \mathbb{R} \rightarrow \mathbb{R}$, $h(x) = c \cdot x^2 + d \cdot x + e$ ($c, d, e \in \mathbb{R}$). Bestimme für einen beliebigen Punkt $a \in \mathbb{R}$ die Ableitung von f , g und h nur mit Hilfe des Differenzenquotienten.

Rechenregeln für Ableitungen

Schon Übung 8.1 zeigt, wie mühsam es manchmal ist, Ableitungen direkt mit dem Differenzenquotienten zu bestimmen. Zum Glück gibt es einige Rechenregeln, die einem das Leben erleichtern. Diese werden alle in der Vorlesung Mathe I bewiesen (wahrscheinlich).

Satz 8.7. Seien $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbare Funktionen und $\alpha \in \mathbb{R}$ beliebig. Dann gilt:

$$(f + g)' = f' + g' \text{ (Summenregel)} \quad (\alpha \cdot f)' = \alpha \cdot f' \text{ (Faktorregel)}$$

Nicht ganz so elementar ist die folgende Regel:

Satz 8.8 (Produktregel). Seien $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbare Funktionen. Dann gilt:

$$(f \cdot g)' = f' \cdot g + f \cdot g'$$

Aufgabe 8.3. Berechne die Ableitung der Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ gegeben durch $f(x) = x^n$ für alle $n \in \mathbb{N}$ mit Hilfe der Produktregel und vollständiger Induktion nach n .

Eine weitere wichtige Regel zur Bestimmung von Ableitungen ist die **Quotientenregel**.

Satz 8.9 (Quotientenregel). Seien $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar und betrachte wieder $D = \{x \in \mathbb{R} : g(x) \neq 0\}$. Dann ist die Funktion

$$\frac{f}{g} : D \rightarrow \mathbb{R} \quad \left(\frac{f}{g}\right)(x) = \frac{f(x)}{g(x)}$$

differenzierbar und es gilt

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

Jetzt haben wir fast alle wichtigen Regeln, beisammen, es fehlt allerdings noch die Regel zur Ableitung von verketteten Abbildungen. Recall: $f \circ g$ ist die Funktion $f(g(x))$.

Satz 8.10 (Kettenregel). Seien $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar, dann auch $(f \circ g)$ und es gilt:

$$(f \circ g)'(a) = g'(a) \cdot f'(g(a))$$

Zum Abschluss betrachten wir die Ableitung der Umkehrfunktion.

Satz 8.11. Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar und bijektiv, so dass die Umkehrabbildung f^{-1} auch differenzierbar ist. Dann gilt:

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

Beweis. Nach Definition der Umkehrabbildung gilt $(f \circ f^{-1}) = \text{id}_{\mathbb{R}}$, also folgt nach der Kettenregel, wenn auf beiden Seiten die Ableitung gebildet wird für alle $a \in \mathbb{R}$:

$$(f^{-1})'(a) \cdot f'(f^{-1}(a)) = 1 \quad \Leftrightarrow \quad (f^{-1})'(a) = \frac{1}{(f' \circ f^{-1})(a)}$$

□

Als Anwendung betrachte die Exponentialfunktion $\exp : \mathbb{R} \rightarrow \mathbb{R}^+$ gegeben durch $\exp(x) = e^x$. Diese ist dadurch gekennzeichnet, dass gilt $\exp' = \exp$, siehe oben. Ihre Umkehrabbildung ist der (natürliche) Logarithmus $\ln : \mathbb{R}^+ \rightarrow \mathbb{R}$. (“Logarithmus naturalis”).

Mit obiger Regel fällt es nun leicht, die Ableitung des Logarithmus zu bestimmen:

$$(\ln)'(a) = \frac{1}{\exp'(\ln(a))} = \frac{1}{\exp(\ln(a))} = \frac{1}{a}$$

Die Ableitung des Logarithmus ist also die Funktion, die jedem positiven Wert x seinen Kehrwert $\frac{1}{x}$ zuordnet.

Aufgabe 8.4. Sei $c \in \mathbb{R}^+$ beliebig und betrachte

$$f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = c^x \quad \text{und} \quad g : \mathbb{R} \rightarrow \mathbb{R}, g(x) = x^x.$$

Bestimme die Ableitungen von f und g . Tipp/Spoiler:

www.math.uni-bielefeld.de/~frettlow/teach/ueb/vorkurs2015/spoilera94.html

Damit sind die Grundlagen für das 2. Semester geschaffen, wo es im Analyseteil um höherdimensionale Funktionen geht (z.B. $f : \mathbb{R}^3 \rightarrow \mathbb{R}$, $f(x, y, z) = (x^2 + y^2)ze^{xyz}$), Extrema von diesen, mehrdimensionale Integrale und Differentialgleichungen. Damit ist dieser Vorkurs am Ende; es folgt noch Bonusmaterial zu Unendlichkeit.

9 Bonusmaterial: Binomialkoeffizienten

Die Algebra ist nichts anderes als eine symbolische Notierung geometrischer Sachverhalte, die Geometrie ihrerseits – das ist in Figuren verkörperte Algebra
Sophie Germain (?)

9.1 Die Formel von Signore Binomi

Wir haben weiter oben schon die binomischen Formeln gesehen. Betrachten wir die erste davon:

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

Was passiert, wenn wir nicht $(a+b)^2$ ausrechnen wollen, sondern $(a+b)^3$, $(a+b)^4$, $(a+b)^5$...?

Aufgabe 9.1. Berechne $(a+b)^0$, $(a+b)^1$, $(a+b)^2$, $(a+b)^3$, $(a+b)^4$, $(a+b)^5$ und fasse soweit wie möglich zusammen.

Wenn man das tut, sieht man

$$\begin{array}{rcl}
 (a+b)^0 & = & 1 \\
 (a+b)^1 & = & a + b \\
 (a+b)^2 & = & a^2 + 2ab + b^2 \\
 (a+b)^3 & = & a^3 + 3a^2b + 3ab^2 + b^3 \\
 (a+b)^4 & = & a^4 + 4a^3b + 6a^2b^2 + 4ab^3 + b^4 \\
 (a+b)^5 & = & a^5 + 5a^4b + 10a^3b^2 + 10a^2b^3 + 5ab^4 + b^5 \\
 (a+b)^6 & = & a^6 + 6a^5b + 15a^4b^2 + 20a^3b^3 + 15a^2b^4 + 6ab^5 + b^6 \\
 & & \vdots
 \end{array}$$

Betrachten wir nur die Zahlen im obigen Schema, so erhalten wir folgendes Muster:

$$\begin{array}{ccccccc}
& & & 1 \\
& & 1 & & 1 \\
& & & 2 & & 1 \\
& 1 & & 3 & & 3 & & 1 \\
& & 1 & & 4 & & 6 & & 4 & & 1 \\
& & & 5 & & 10 & & 10 & & 5 & & 1 \\
1 & & 1 & & 6 & & 15 & & 20 & & 15 & & 6 & & 1 \\
& & & & \vdots & & & & & & & & & &
\end{array}$$

Es ergibt sich ein dreieckiges Muster, am Rand stehen lauter Einsen, und an den anderen Stellen ist jede Zahl die Summe der beiden darüber stehenden. Das ist das **Pascalsche Dreieck**. Es taucht in der Mathematik häufig auf, und man kann viele interessante Muster darin entdecken.

Eines ist etwa: Stellen wir uns vor, dass wir in dem Dreieck in der 1 ganz oben starten. In jedem Schritt gehen wir schräg nach unten, d.h. zu der Zahl, die rechts oder links unter unserer aktuellen Zahl steht. Wie viele Möglichkeiten gibt es, zu einer gegebenen Stelle zu kommen? Z.B. zu der 3, oder zu der 15, oder zu der 20?

Die binomische Formel stammt übrigens nicht von einem Herrn Binomi. Die heißt so, weil die beiden Variablen a und b darin vorkommen: “binomisch” ist lateinisch für “zwei-namig”. Die Zuschreibung an Herrn Binomi ist ein alter Mathematikerwitz.

9.2 Der Binomialkoeffizient

Die Einträge im Pascalschen Dreieck heißen Binomialkoeffizienten. Warum das so heißt ist nach dem obigen nun klar. Weil dieser Ausdruck so häufig vorkommt, gibt es eine abkürzende Schreibweise dafür: Der k -te Eintrag in Reihe n heißt $\binom{n}{k}$. (Man spricht das Symbol $\binom{n}{k}$ als “ n über k ”.) Dabei soll man aber bei Null anfangen zu zählen. Die oberste Reihe ist also die 0-te Reihe. Die vorderste Zahl in einer Zeile ist die 0-te Zahl. Also ist z.B. $\binom{4}{0} = 1$, $\binom{4}{1} = 4$, oder $\binom{5}{2} = 10$.

$$\binom{n}{k} := \frac{n!}{(n-k)! \cdot k!} \quad (7)$$

Der Binomialkoeffizient gibt auch eine Antwort auf eine Frage, die in der Mathematik manchmal von Bedeutung ist: Sei M eine Menge mit n Elementen und sei $k \leq n$. Wie viele Teilmengen von M gibt es, die genau k Elemente haben? Die Antwort ist $\binom{n}{k}$. Die Reihenfolge, in der die Elemente auftauchen spielt dabei keine Rolle.

Als erste Anwendung kann gleich das Lottospiel dienen: es gibt beim Lotto exakt

$$\binom{49}{6} = \frac{49!}{43! \cdot 6!} = \frac{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44}{6!} = 13983816$$

Möglichkeiten.

Aufgabe 9.2. In der Vorlesung “Mathematik für Informatiker” sitzen 200 Personen. Das Studentenwerk verlost 3 Mensagutscheine mit einem Wert von je 10 Euro. Wie viele Möglichkeiten gibt es, 3 Preisträger auszuwählen?

Aufgabe 9.3. Zeige:

$$\binom{n}{0} = \binom{n}{n} = 1 \quad \text{und} \quad \binom{n}{k} = \binom{n}{n-k}$$

Aufgabe 9.4. Beweise die Formel

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

elementar mit Hilfe der Definition von $\binom{n}{k}$ als $\frac{n!}{(n-k)! \cdot k!}$.

9.3 Der binomische Lehrsatz

Wir haben gesehen: Multipliziert man $(a+b)^n$ aus und sortiert die Terme, ergeben sich genau die Zahlen aus dem Pascalschen Dreieck. Das ist natürlich kein Zufall. Die Überlegung dahinter ist die Folgende: ein Ausdruck der Form

$$(x+y)^n = (x+y) \cdot (x+y) \cdot (x+y) \cdot \dots \cdot (x+y)$$

soll ausmultipliziert werden. Das Endergebnis ist eine gigantische Summe, in der jeder einzelne Summand aus den einzelnen Summanden der Klammern gebildet wird, die miteinander multipliziert werden müssen. Wählt man zum Beispiel aus jeder Klammer den ersten Summanden x aus, erhält man im Ergebnis den Summanden x^n .

Bei diesem Verfahren kommen aber etliche Summanden mehrfach vor. Da die Multiplikation kommutativ ist, gilt ja zum Beispiel

$$x \cdot x \cdot y \cdot x = y \cdot x \cdot x \cdot x = x^3 y$$

Dieser Summand kann also auf mehrere Arten gewonnen werden. Und auf wie viele genau, das verrät uns die Kombinatorik.

In jedem Summanden kommen ja genau n Faktoren vor, einer aus jeder Klammer und jeder Faktor ist entweder x oder y . Also kann man jeden Summanden in der Form $x^k y^{n-k}$ für ein $k \leq n$ schreiben, wobei $k = 0$ auch zugelassen ist.

Und wie viele Möglichkeiten gibt es, sich den Summanden $x^k y^{n-k}$ aus den Klammern zusammenzusuchen? Genau $\binom{n}{k}!$ Man muss aus der n -elementigen Menge der Klammern eine k -elementige Teilmenge der Klammern auswählen, aus denen ein x kommen soll und dafür gibt es genau $\binom{n}{k}$ Möglichkeiten. Diese informellen Überlegungen begründen den folgenden

Satz 9.1 (Binomischer Lehrsatz).

Für $x, y \in \mathbb{R}$ und $n \in \mathbb{N}$ gilt:

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

Die obigen Überlegungen sind noch nicht als strenger Beweis formuliert. Da ich hier unter anderem eine Vorbildrolle spiele, soll dieser Satz nun mit Hilfe des Verfahrens der vollständigen Induktion bewiesen werden.

Beweis.

1) *Induktionsanfang:* Sei $n = 1$. Dann gilt

$$\sum_{k=0}^1 \binom{1}{k} x^k y^{1-k} = 1 \cdot y + 1 \cdot x = (x + y)^1$$

2) *Induktionsschritt:* Wir nehmen an, dass der Satz für ein $n \in \mathbb{N}$ bereits gilt. Nun soll er für $(n + 1)$ gezeigt werden. Es gilt aber

$$(x + y)^{n+1} = (x + y) \cdot (x + y)^n \stackrel{IV}{=} (x + y) \cdot \sum_{k=0}^n \binom{n}{k} x^k y^{n-k}$$

Multipliziert man die Klammer im rechten Ausdruck aus, so ergibt sich

$$(x + y) \cdot \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} = \sum_{k=0}^n \binom{n}{k} x^{k+1} y^{n-k} + \sum_{k=0}^n \binom{n}{k} x^k y^{n-k+1}$$

In der ersten Summe ist einfach ein x hinzugekommen und in der zweiten ein y . Um die beiden Ausdrücke addieren zu können, ist eine Indexverschiebung notwendig. Betrachten wir nur die erste Summe, so kann diese umgeschrieben werden:

$$\sum_{k=0}^n \binom{n}{k} x^{k+1} y^{n-k} = \sum_{k=1}^{n+1} \binom{n}{k-1} x^k y^{n-k+1}$$

Dies hat den immensen Vorteil, dass jetzt die Exponenten von x und y übereinstimmen. Also können wir die Summen wie folgt zusammenziehen:

$$(x+y)^{n+1} = \dots = x^{n+1} + y^{n+1} + \sum_{k=1}^n \left(\binom{n}{k-1} + \binom{n}{k} \right) x^k y^{n-k+1}$$

Hier wurde der letzte Ausdruck aus der ersten Summe (der für $k = n+1$) und der erste Ausdruck aus der letzten Summe (der für $k = 0$) gesondert aufgeschrieben und der Rest wurde addiert. Nach der Rechenregel für Binomialkoeffizienten ergibt sich aber gerade das Gewünschte:

$$(x+y)^{n+1} = \dots = \sum_{k=0}^{n+1} \binom{n+1}{k} x^k y^{n+1-k}$$

□

Aufgabe 9.5. Zeige mittels der Formel (7), Seite 83:

- $\binom{n}{k} = \frac{n}{k} \binom{n-1}{k-1}$
- $\binom{n}{h} \binom{n-h}{k} = \binom{n}{k} \binom{n-k}{h}$
- $\binom{n-1}{k} - \binom{n-1}{k-1} = \frac{n-2k}{n} \binom{n}{k}$

Aufgabe 9.6. Was ist die Summe der Einträge in der Zeile Nummer n des Pascalschen Dreiecks? Der binomische Lehrsatz erlaubt uns, das schnell zu beweisen. Zeige

$$\sum_{k=0}^n \binom{n}{k} = 2^n$$

(Hinweis: Wähle x und y im binomischen Lehrsatz geeignet.)

Aufgabe 9.7. Zeige per Induktionsbeweis:

1. $\sum_{k=1}^n k^3 = \binom{n+1}{2}^2$
2. $\sum_{j=k}^n \binom{j}{k} = \binom{n+1}{k+1}$ (Was bedeutet das anschaulich im Pascalschen Dreieck?)
3. $\sum_{j=0}^m \binom{m}{j}^2 = \binom{2m}{m}$

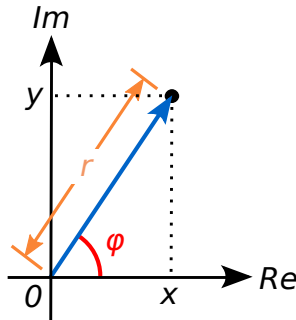
Aufgabe 9.8. Was passiert, wenn wir die Einträge im Pascalschen Dreieck folgendermaßen aufsummieren: Starte bei einer 1 links. Gehe immer einen Schritt nach rechts unten und einen nach rechts und addiere diese Zahl dazu. Mache so weiter, bis das Dreieck durchquert wird.

Es ergeben sich die **Fibonaccizahlen**. Die sind so definiert: $F(0) = F(1) = 1$, $F(n+1) = F(n) + F(n-1)$. Die Folge der Fibonaccizahlen fängt also so an: 1, 1, 2, 3, 5, 8, 13, 21, ... Zeige

$$\sum_{k=0}^{\lfloor \frac{n}{2} \rfloor} \binom{n-k}{k} = F(n+1).$$

10 Bonusmaterial: Polarkoordinaten

Um Multiplikation komplexer Zahlen besser zu verstehen, betrachten wir komplexe Zahlen mit Hilfe von sogenannten Polarkoordinaten. Dabei beschreiben wir eine komplexe Zahl (also anschaulich einen Punkt in der (komplexen) Ebene) nicht durch x - und y -Koordinaten (die heißen auch **kartesische Koordinaten**, heißt einfach “Ost-West” und “Nord-Süd”), sondern durch “Richtung” und “Entfernung”. Die Richtung geben wir durch den Winkel φ mit der x -Achse an, die Entfernung ist die Entfernung von 0 zum betrachteten Punkt.



Ist $z = a + ib$ eine komplexe Zahl, so nennt man die positive reelle Zahl $|z| = \sqrt{a^2 + b^2}$ den *Betrag* von z . Das ist gerade die Entfernung zwischen z und 0. Konkret kann diese komplexe Zahl also geschrieben werden als

$$|z| \cdot \cos \varphi + i \cdot |z| \cdot \sin \varphi = |z| \cdot (\cos \varphi + i \cdot \sin \varphi)$$

Ein wichtiger Zusammenhang zwischen Exponential- und (Ko-)Sinusfunktion ist die **Euler-sche Formel**:

$$\cos \varphi + i \cdot \sin \varphi = e^{i\varphi}$$

Damit kann man eine komplexe Zahl also kurz so schreiben:

$$|z|e^{i\varphi}$$

Ist eine komplexe Zahl in dieser Schreibweise gegeben, sagt man, die Zahl ist in **Polarkoordinaten** gegeben. Die suggestive Schreibweise als Potenz macht Sinn, da folgendes gilt:

$$e^{i\varphi} \cdot e^{i\beta} = e^{i(\varphi+\beta)}$$

Wer mag, kann dies mit Hilfe der Additionstheoreme von Sinus und Cosinus nachrechnen. Diese Formeln erlauben aber jetzt die Interpretation der Multiplikation komplexer Zahlen:

$$(r \cdot e^{i\varphi}) \cdot (s \cdot e^{i\beta}) = (r \cdot s) \cdot e^{i(\varphi+\beta)}$$

In Worten: werden zwei komplexe Zahlen miteinander multipliziert, so multiplizieren sich ihre Längen und ihre Winkel addieren sich.

Die **Umrechnungsformeln** von kartesischen in Polarkoordinaten und umgekehrt:

$$z = a + ib \text{ dann } z = \sqrt{a^2 + b^2} e^{i\varphi} \text{ mit } \varphi = \begin{cases} \arctan \frac{b}{a} & \text{für } a > 0 \\ \arctan \frac{b}{a} + \pi & \text{für } a < 0, b \geq 0 \\ \arctan \frac{b}{a} - \pi & \text{für } a < 0, b < 0 \\ +\pi/2 & \text{für } a = 0, b > 0 \\ -\pi/2 & \text{für } a = 0, b < 0 \end{cases}$$

$$z = r e^{i\varphi} \text{ dann } z = r \cos \varphi + (r \sin \varphi)i$$

Beim arctan muss man etwas aufpassen, welche Vorzeichen a und b haben. Nur daher sieht die Formel oben so kompliziert aus. Für $a = 0$ und $b = 0$ gibt's nix umzurechnen, 0 ist 0.

Polarkoordinaten erlauben uns, die Lösung des folgenden Problems zu erraten: was ist $\sqrt{i}$? Müssen wir hier schon wieder unsere Menge erweitern? Hört der Prozess erst auf, wenn uns die Buchstaben ausgehen?

Zum Glück ist $\sqrt{i}$ wieder eine komplexe Zahl. Und die Formel der Multiplikation oben verrät uns, dass sie die Länge 1 haben muss (denn i hat die Länge 1 und die Längen multiplizieren sich ja) und einen Winkel von 45 Grad (oder $\frac{\pi}{4}$ im Bogenmaß ausgedrückt.) Und eine solche Zahl gibt es wirklich:

$$\sqrt{i} = \frac{1+i}{\sqrt{2}}$$

Die Probe zeigt, dass wir richtig liegen:

$$\left(\frac{1+i}{\sqrt{2}}\right)^2 = \frac{(1+i)^2}{2} = \frac{1+2i-1}{2} = i$$

Aufgabe 10.1. 1. Rechne folgende komplexe Zahlen in Polarkoordinaten in kartesische Koordinaten um:

$$3 \cdot e^{\pi i}, 2 \cdot e^{-\pi i}, 5 \cdot e^{\frac{\pi}{2}i}, 2 \cdot e^{\frac{\pi}{6}i}, e^{1+\pi i}, i^i, -e^{3\pi i}.$$

2. Rechne folgende komplexe Zahlen in kartesischen Koordinaten in in Polarkoordinaten um:

$$-1+i, (1-i)^3, \frac{1+2i}{2-i}, \sqrt{3}+i, \sqrt{2}+\sqrt{2}i.$$

Aufgabe* 10.2. (Zusatz) Zeige: $e^{i(\varphi+\beta)} = e^{i\varphi} \cdot e^{i\beta}$. (Es ist hilfreich, Fakten aus den vorangegangenen Kapiteln zu nutzen.)

11 Bonusmaterial: Unendlichkeit

Es gibt zwei Dinge die unendlich sind: Das Universum und die menschliche Dummheit - beim Universum ist man sich noch nicht ganz sicher.

Lukas Podolski

11.1 Endliche und unendliche Mengen

Jeder Mensch kommt verhältnismäßig früh im Leben in Kontakt mit dem Begriff der Unendlichkeit: beim Zählen nämlich. Das schwer zu fassende Konzept, dass die natürlichen Zahlen “immer weiter” gehen, man also nie aufhören kann zu zählen ist anfangs gewöhnungsbedürftig, schließlich sind wir nur von endlichen Größen umgeben. Mit der Zeit gewöhnt man sich aber daran und akzeptiert es.

Eine wichtige Eigenschaft endlicher Mengen ist, dass es keine Bijektion auf eine echte Teilmenge geben kann.

Satz 11.1. *Seien n und m natürliche Zahlen mit $n > m$ und betrachte die Mengen*

$$N := \{k \in \mathbb{N} : k \leq n\} = \{1, \dots, n\} \quad M := \{k \in \mathbb{N} : k \leq m\} = \{1, \dots, m\}$$

Dann gilt $M \subsetneq N$ und es gibt keine bijektive Abbildung $f : N \rightarrow M$.

Beweis. Sei $f : N \rightarrow M$ surjektiv und wähle für $k \in M$ ein Urbild $a_k \in N$ aus. Falls diese nicht paarweise verschieden sind (d.h. $a_k \neq a_l$ für $k \neq l$), dann ist f keinesfalls injektiv. Nimm also an, dass diese paarweise verschieden sind und betrachte $N \setminus \{a_1, \dots, a_m\}$. Diese Menge ist nicht leer, da $n > m$ ist, sei also $b \in N \setminus \{a_1, \dots, a_m\}$.

Wegen $f(b) \in M$ gilt $f(b) = k$ für ein $k \leq m$, also folgt $f(b) = f(a_k)$, aber wegen $b \neq a_k$ ist f nicht injektiv. $\square$

Der Beweis wirkt tautologisch, aber genau diese Eigenschaft von endlichen Mengen ist es, die abhanden kommt, wenn man unendliche Mengen wie zum Beispiel die natürlichen Zahlen betrachtet: setze $M := \{n \in \mathbb{N} : n > 4\} = \{5, 6, 7, \dots\}$.

Dann ist $M \subsetneq \mathbb{N}$, aber $f : \mathbb{N} \rightarrow M$ gegeben durch $f(n) = n + 4$ ist eine Bijektion. Das kann man natürlich auch formal zeigen.

Beweis. Die Abbildung f ist injektiv, denn aus $f(n) = f(m)$ folgt $n + 4 = m + 4$ und daher $m = n$. Aber f ist auch surjektiv, denn falls $m \in M$, dann ist $m > 4$ und daher ist $m - 4 > 0$, also $m - 4 \in \mathbb{N}$ und $f(m - 4) = (m - 4) + 4 = m$. Daher ist $(m - 4)$ ein Urbild von m unter f . $\square$

Aus genau demselben Grund gibt es auch eine Bijektion zwischen den Mengen $\mathbb{N}$ und $\mathbb{N}_0$, obwohl erstere eine echte Teilmenge von letzterer ist.

Mengen zwischen denen eine Bijektion existiert heißen auch **gleichmächtig** (Im Zeichen: $N \cong M$, wenn N und M gleichmächtig sind). Die natürlichen Zahlen sind also gleichmächtig zu einer echten Teilmenge von sich. Diese Eigenschaft nahm der Mathematiker *Dedekind* als definierende Eigenschaft der Unendlichkeit.

Definition 11.2. Sei N eine Menge. Wenn es eine echte Teilmenge $M \subsetneq N$ und eine Bijektion $f : N \rightarrow M$ gibt, dann heißt N **unendlich**.

Eine alternative Definition ist die Folgende: eine Menge N heißt **endlich**, wenn sie entweder leer ist oder es eine natürliche Zahl n und eine Bijektion

$$f : N \rightarrow \{k \in \mathbb{N} : k \leq n\} = \{1, \dots, n\}$$

gibt. Alle Mengen, die nicht endlich sind heißen **unendlich**.

Beide Definitionen sind gleichwertig, der Beweis erfordert allerdings ein etwas umstrittenes Axiom der Mengenlehre, das Auswahlaxiom. Darauf wollen wir hier nicht näher eingehen.

In den behandelten Beispielen war die echte Teilmenge, die betrachtet wurde, nur um endlich viele Elemente kleiner als $\mathbb{N}$. Dass dies nicht immer so sein muss, zeigen folgende Beispiele: betrachte die Mengen

$$\begin{aligned} G &:= \{2n : n \in \mathbb{N}\} = \{2, 4, 6, 8, 10, \dots\} \subsetneq \mathbb{N} \\ \text{und} \quad Q &:= \{n^2 : n \in \mathbb{N}\} = \{1, 4, 9, 16, 25, \dots\} \subsetneq \mathbb{N} \end{aligned}$$

der geraden Zahlen bzw. der Quadratzahlen in $\mathbb{N}$. Beide Mengen sind gleichmächtig zu $\mathbb{N}$.

Beweis. Die Abbildungen $f : \mathbb{N} \rightarrow G$ und $g : \mathbb{N} \rightarrow Q$ gegeben durch

$$f(n) = 2n \quad \text{bzw.} \quad g(n) = n^2$$

sind Bijektionen. Ihre Umkehrabbildungen sind gegeben durch $f^{-1}(n) = \frac{n}{2}$ bzw. $g^{-1}(n) = \sqrt{n}$. $\square$

Mit diesem Verfahren kann man auch zeigen, dass die Menge der ganzen Zahlen $\mathbb{Z}$ zur Menge der natürlichen Zahlen $\mathbb{N}$ gleichmächtig ist.

Beweis. Es ist leicht zu sehen, dass “gleichmächtig sein” transitiv ist. D.h. wenn $X \cong Y$ und $Y \cong Z$, dann gilt auch $X \cong Z$. Wegen $\mathbb{N} \cong \mathbb{N}_0$ reicht es zu zeigen, dass gilt: $\mathbb{Z} \cong \mathbb{N}_0$. Betrachte dazu folgende Abbildung:

$$f : \mathbb{Z} \rightarrow \mathbb{N}_0 \quad f(x) = \begin{cases} 2x & , \text{ falls } x \geq 0 \\ -2x - 1 & , \text{ falls } x < 0 \end{cases}$$

Zunächst zur Surjektivität. Sei $n \in \mathbb{N}_0$ beliebig. Falls n gerade ist, dann ist $x := \frac{n}{2}$ eine ganze Zahl mit $x \geq 0$, also folgt $f(x) = 2x = n$. Falls n ungerade ist, dann ist $n+1$ gerade und wir setzen $x := -\frac{n+1}{2}$. Es folgt $x < 0$, also

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

Zur Injektivität: sind $x, y \in \mathbb{Z}$ mit $f(x) = f(y)$, so gibt es zwei Fälle: ist $f(x)$ gerade, dann folgt $x = \frac{f(x)}{2} = \frac{f(y)}{2} = y$. Ist $f(x)$ hingegen ungerade, dann ist nach obigem Argument $x = -\frac{f(x)+1}{2} = -\frac{f(y)+1}{2} = y$. $\square$

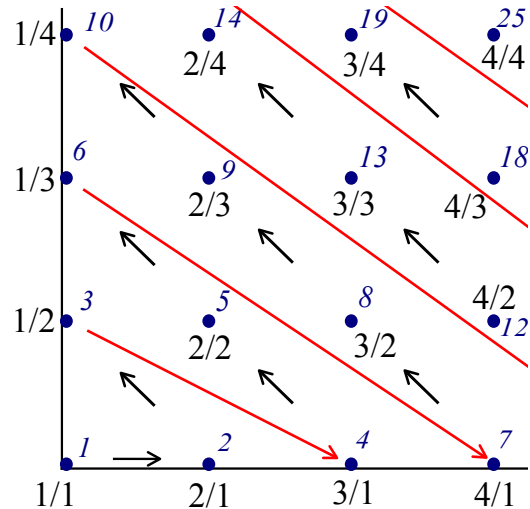
Die Abbildung f aus dem Beweis kann man sich folgendermaßen bildlich vorstellen:

$$\begin{array}{cccccccc} \dots & -3 & -2 & -1 & 0 & 1 & 2 & 3 & \dots \\ & & & & \downarrow & & & & \\ \dots & 5 & 3 & 1 & 0 & 2 & 4 & 6 & \dots \end{array}$$

Die negativen ganzen Zahlen werden auf die ungeraden Zahlen abgebildet, die positiven Zahlen gehen auf die geraden Zahlen und die 0 geht auf die 0.

Und was ist mit der Menge der rationalen Zahlen, $\mathbb{Q}$? Sind das denn wenigstens “mehr”? Die Antwort ist überraschenderweise nein: es gilt wieder $\mathbb{Q} \cong \mathbb{N}$.

Beweis. Zunächst halten wir fest, dass es reicht, $\mathbb{Q}_+^* = \{q \in \mathbb{Q} : q > 0\} \cong \mathbb{N}$ zu zeigen, denn $\mathbb{Q} \cong \mathbb{Q}_+^*$ folgt dann wie oben. Es reicht sogar eine surjektive Abbildung $f : \mathbb{N} \rightarrow \mathbb{Q}_+^*$ anzugeben, denn nach Kapitel 5 kann diese durch Verkleinern der Definitionsmenge injektiv, insgesamt also bijektiv gemacht werden und eine Bijektion zwischen $\mathbb{Q}_+^*$ und einer echten nicht endlichen Teilmenge der natürlichen Zahlen reicht aus. Wie findet man diese surjektive Abbildung f ? Betrachte folgendes Bild:



Das Bild ist so zu verstehen, dass die obere linke Ecke das Bild der Zahl 1 sein soll. Danach folgt man den Pfeilen und erhält so eine “Aufzählung” der positiven Brüche, also eine Abbildung $f : \mathbb{N} \rightarrow \mathbb{Q}_+^*$. Diese ist surjektiv, da jeder Bruch in dem Bild vorkommt, aber natürlich kommt jeder Bruch sogar mehrfach vor, da auch ungekürzte Schreibweisen auftauchen. $\square$

Definition 11.3. Sei M eine unendliche Menge. Falls $\mathbb{N} \cong M$, so heißt M **abzählbar**. Dieser Begriff leitet sich daraus her, dass eine bijektive Abbildung $f : \mathbb{N} \rightarrow M$ eine Abzählung der Elemente von M darstellt, man kann mit Hilfe dieser Abbildungen die Elemente von M also der Reihe nach hinschreiben.

Zum Beispiel kann man für $n \in \mathbb{N}$ folgende Schreibweise einführen: $a_n := f(n)$. Dann ist $a_n \in M$ für alle $n \in \mathbb{N}$ und es gilt aufgrund der Bijektivität von f :

$$M = \{a_1, a_2, a_3, a_4, a_5, a_6, \dots\}$$

Aufgabe 11.1. Hilberts Hotel. Die folgende Übung ist der interessante Vergleich von Hilbert zum Thema Unendlichkeit, der als “Hilberts Hotel” bekannt geworden ist. Dieses Hotel ist ein fiktiver Ort, der $\mathbb{N}$ Zimmer hat, also für jede natürliche Zahl ein Zimmer, mit anderen Worten: unendlich viele. Stellen wir uns vor, dass alle Zimmer belegt sind, das Hotel also restlos ausgebucht ist.

Trotzdem hat der Portier die Anweisung, alle neu ankommenden Gäste unterzubringen ohne bereits vorhandene Gäste ausquartieren zu müssen. Umquartieren innerhalb des Hotels ist ihm aber erlaubt. Wie kommt er mit folgenden Situationen klar?

- a) Ein neuer Gast kommt an.

- b) Es kommen n Gäste an, wobei $n \in \mathbb{N}$ eine natürliche Zahl ist.
- c) Ein Großraumbus mit $\mathbb{N}$ Gästen kommt an.
- d) Es kommen n Großraumbusse an.
- e) Eine Kolonne mit $\mathbb{N}$ Großraumbussen kommt an.

11.2 Überabzählbarkeit – noch mehr als $\mathbb{N}$

Das Unendliche ist weit, vor allem gegen Ende.

Batman

Im vorigen Abschnitt haben wir gesehen, dass $\mathbb{N} \cong \mathbb{Z} \cong \mathbb{Q}$ gilt. Da liegt doch die Vermutung nahe, dass jede unendliche Menge gleichmächtig zu den natürlichen Zahlen ist, also z.B. auch die Menge $\mathbb{R}$ der reellen Zahlen. Dies ist aber falsch: zunächst werden wir ein Argument erarbeiten, das zeigt, dass nicht automatisch je zwei unendliche Mengen bijektiv zueinander sind und danach zeigen, dass es keine surjektive Abbildung $f : \mathbb{N} \rightarrow \mathbb{R}$ gibt, die reellen Zahlen also tatsächlich mächtiger sind (man sagt: eine größere Kardinalität besitzen) als die natürlichen Zahlen.

Definition 11.4. Sei M eine beliebige Menge. Betrachte dazu die Menge aller Teilmengen von M :

$$\mathcal{P}(M) := \{A \subseteq M\}.$$

Diese wird als **Potenzmenge** von M bezeichnet. Beispielsweise gilt:

$$\mathcal{P}(\{a, b, c\}) = \{\emptyset, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}$$

Aufgabe 11.2. Sei M eine endliche Menge mit n Elementen. Zeige: die Potenzmenge von M hat 2^n Elemente.

Satz 11.5. Sei M eine beliebige Menge. Dann gibt es keine surjektive Abbildung

$$f : M \rightarrow \mathcal{P}(M).$$

Insbesondere gilt $M \not\cong \mathcal{P}(M)$.

Beweis. Wieder ein Beweis durch Widerspruch. Angenommen $f : M \rightarrow \mathcal{P}(M)$ wäre surjektiv. Betrachte dann folgende Menge

$$X := \{m \in M : m \notin f(m)\} \subseteq M$$

Da $X \subseteq M$ gilt $X \in \mathcal{P}(M)$. Wegen der Surjektivität von f gibt es also ein $x \in M$ mit $f(x) = X$. Nun unterscheiden wir zwei Fälle:

Fall 1: $x \in f(x)$. Da $f(x) = X$ folgt nach Definition von X : $x \notin f(x)$. Widerspruch.

Fall 2: $x \notin f(x)$. Dann folgt aber $x \in X = f(x)$. Widerspruch.

Beide Fälle führen also zu einem Widerspruch. Daher kann f nicht surjektiv sein, die Menge X kann kein Urbild haben. $\square$

Der Beweis basiert auf einem Paradoxon, das auch bekannt ist als der “Barbier von Sevilla”: das ist derjenige Mensch in dem Ort Sevilla, der alle Männer rasiert, die sich nicht selbst rasieren. Die Frage, ob er sich denn selbst rasieren führt auf einen ähnlichen unlösbaren Widerspruch.

Es folgt auch, dass es unendlich viele unendliche Mengen gibt, die alle untereinander nicht gleichmächtig sind: starte mit einer unendlichen Menge M , dann ist

$$M \not\cong \mathcal{P}(M) \not\cong \mathcal{P}(\mathcal{P}(M)) \not\cong \dots$$

Satz 11.6. *Es gibt keine surjektive Abbildung $f : \mathbb{N} \rightarrow \mathbb{R}$, d.h. $\mathbb{N} \not\cong \mathbb{R}$.*

Beweis. Es reicht zu zeigen, dass es keine surjektive Abbildung $f : \mathbb{N} \rightarrow [0, 1] = \{x \in \mathbb{R} : 0 \leq x \leq 1\}$ gibt, denn dann gibt es erst Recht keine Surjektion nach $\mathbb{R}$. Sei dazu eine beliebige Abbildung $f : \mathbb{N} \rightarrow [0, 1]$ gegeben und schreibe zu $n \in \mathbb{N}$ das Bild $f(n)$ als Dezimalzahl

$$f(n) = 0, a_{n,1}a_{n,2}a_{n,3}a_{n,4}a_{n,5}a_{n,6} \dots \quad a_{n,i} \in \{0, 1, \dots, 9\}$$

Die ersten Zahlen könnten beispielsweise so aussehen:

$$\begin{aligned} f(1) &= 0,4190318539\dots \\ f(2) &= 0,2901387503\dots \\ f(3) &= 0,8091244112\dots \\ f(4) &= 0,9933187609\dots \\ f(5) &= 0,0290047014\dots \end{aligned}$$

Definiere nun zu $n \in \mathbb{N}$ die Ziffer $b_n := a_{n,n} + 1$ falls $a_{n,n} \neq 9$, sonst setze $b_n := 0$. Dann ist b_n eine Ziffer zwischen 0 und 9, die man erhält, indem man aus der Zahl $f(n)$ die n -te Nachkommastelle um 1 erhöht bzw. zu 0 macht, wenn dort eine 9 steht.

Daraus kann man nun eine Zahl $x = 0, b_1b_2b_3b_4b_5b_6 \dots \in [0, 1]$ konstruieren. Diese liegt nicht im Bild von f ! Denn sie unterscheidet sich an der ersten Nachkommastelle von $f(1)$, an der zweiten von $f(2)$ usw. und an der n -ten Nachkommastelle von $f(n)$.

Da diese Zahl x nicht im Bild von f liegt, ist f daher nicht surjektiv. Dieses Verfahren heißt auch **Diagonalverfahren** nach Cantor. $\square$

Definition 11.7. Ist M eine unendliche Menge mit $M \not\cong \mathbb{N}$, so heißt M **überabzählbar**. Der eben bewiesene Satz zeigt die Überabzählbarkeit der reellen Zahlen.

Natürlich gibt es Mengen, die mächtiger sind als die reellen Zahlen, z.B. $\mathcal{P}(\mathbb{R})$. An dieser Stelle ist vor allem der Unterschied zwischen abzählbaren und überabzählbaren Mengen wichtig. Insbesondere folgt, dass es keine Möglichkeit gibt, die reellen Zahlen aufzuzählen, was für die rationalen Zahlen sehr wohl funktioniert.

Aufgabe 11.3. Betrachte die Menge

$$X = \{A \subseteq \mathbb{N} : A \text{ ist unendlich und abzählbar}\}$$

Zeige: X ist überabzählbar.

Was ist mit der Menge

$$Y = \{A \subseteq \mathbb{N} : A \text{ ist überabzählbar}\}?$$

Aufgabe 11.4. Sei I eine beliebige Menge (möglicherweise überabzählbar) und sei $f : I \rightarrow \mathbb{R}_0^+$ eine Abbildung. Betrachte

$$\sum_{i \in I} f(i)$$

und nimm an, dass diese Summe konvergent ist. Zeige: dann ist

$$X = \{i \in I : f(i) \neq 0\}$$

abzählbar oder endlich.