

Diskrete Strukturen

Wintersemester 2025/26

Prof. Dr. David Sabel

Stand: 13. Januar 2026

Theoretische Informatik
Fachbereich DCSM
Hochschule RheinMain
Unter den Eichen 5
65195 Wiesbaden

Email: david.sabel@hs-rm.de

Inhaltsverzeichnis

1	Einleitung	1
2	Logik	2
2.1	Aussagen und die Aussagenlogik	2
2.1.1	Was sind Aussagen?	2
2.1.2	Syntax der Aussagenlogik	3
2.1.3	Semantik der Aussagenlogik	5
2.2	Gesetze der Logik	9
2.2.1	Tautologien und Kontradiktionen	10
2.2.2	Logische Äquivalenz	11
2.2.3	Rechengesetze	13
2.2.4	Dualität und weitere Rechengesetze	14
2.3	Boolesche Funktionen und vollständige Mengen von Junktoren	16
2.3.1	Boolesche Funktionen	16
2.3.2	Vollständige Mengen an Junktoren	17
2.4	Anwendungen der Aussagenlogik	17
2.4.1	Schaltungsentwurf	17
2.4.2	Modelle suchen	19
2.5	Prädikatenlogische Formeln: Formeln mit Quantoren	19
2.5.1	Aussageformen	20
2.5.2	All- und Existenzaussagen	20
2.5.2.1	Allaussagen	20
2.5.2.2	Existenzaussagen	21
2.5.3	Rechenregeln für Quantoren	23
2.5.3.1	Negation von All- und Existenzaussagen	23
2.5.4	Quantorregeln	23
2.5.5	Vertauschungsregeln	24
2.6	Schlussbemerkungen	24
3	Mengen	25
3.1	Was sind Mengen?	25
3.1.1	Operationen auf Mengen	28
3.1.2	Mengen und Logik	30
3.2	Kartesisches Produkt	33
3.2.1	Allgemeines kartesisches Produkt	34

3.3	Mächtigkeiten und Zählformeln	35
3.4	Schlussbemerkungen	39
4	Relationen, Funktionen und Abzählbarkeit von Mengen	40
4.1	Relationen	40
4.1.1	Mehrstellige Relationen	43
4.1.2	Anwendung: Datenbanksysteme	43
4.1.3	Äquivalenzrelationen und weitere Eigenschaften von Relationen	44
4.2	Funktionen	49
4.3	Abzählbarkeit	53
4.4	Schlussbemerkungen	57
5	Beweise und Beweisen	58
5.1	Einführendes	58
5.1.1	Aufbau mathematischer Texte	58
5.1.2	Warum will man beweisen?	58
5.1.3	Was kennzeichnet Sätze und Beweise?	59
5.2	Beweisarten	59
5.2.1	Direkter Beweis	59
5.2.2	Kontraposition	60
5.2.3	Beweis durch Widerspruch	61
5.2.4	Äquivalenzen	61
5.3	Fallunterscheidung	62
5.4	Schubfachprinzip	63
5.5	Vollständige Induktion	65
5.5.1	Beispiele	66
5.5.2	Vollständige Induktion mit anderem Startwert	67
5.5.3	Starke vollständige Induktion, Fibonacci-Zahlen und der Goldene Schnitt	68
5.5.4	Fibonacci-Zahlen	69
5.6	Induktive Definitionen	71
5.7	Schlussbemerkungen	74
6	Grundlagen der Graphentheorie	75
6.1	Einführendes	75
6.2	Graphen und grundlegende Begriffe	75
6.2.1	Gleichheit und Isomorphie	79
6.3	Eulergraphen	79
6.4	Bäume	83
6.5	Planare Graphen	84
6.6	Bipartite Graphen	85
6.7	Färbungen	87

6.8	Repräsentation von Graphen	87
6.8.1	Adjazenzmatrix	87
6.8.2	Adjazenzlisten	88
6.8.3	Implizite Repräsentation	88
6.9	Schlussbemerkungen	88
7	Algebraische Grundstrukturen	89
7.1	Einführendes	89
7.2	Gruppen	89
7.2.1	Der erweiterte euklidische Algorithmus	93
7.2.2	Eine Anwendung aus der Kryptographie: Die RSA-Verschlüsselung	96
7.3	Körper	97
7.4	Schlussbemerkungen	98
	Literatur	99

1 Einleitung

Die Diskrete Mathematik behandelt solche Strukturen, die endlich oder abzählbar unendlich sind. Damit grenzt sie sich von mathematischen Gebieten ab, die kontinuierliche Mengen und Funktionen darauf betrachten, wie beispielsweise die reelle Analysis.

Diskrete Strukturen passen zur diskreten digitalen Welt der Informatik. In diesem Skript sollen die wesentlichen mathematischen Grundlagen der diskreten Mathematik für die Studierenden der Informatik vermittelt und Anwendungen der Informatik diskutiert werden. Wir befassen uns mit grundlegenden mathematischen Themen wie der Logik und dem Beweisen, mit Strukturen wie Mengen und Relationen. Neben einer Einführung in die Graphentheorie werden wir einführend die Grundbegriffe einer Gruppe und eines Körpers der abstrakten Algebra kennenlernen.

Es gibt viele Bücher zur Diskreten Mathematik und auch zur Einführung in das Themengebiet aus Sicht der Informatik. Daher ist es nicht sinnvoll an dieser Stelle eine Auflistung zu machen. Der präsentierte Inhalt orientiert sich am Material von Marc-Alexander Zschiegner, welches dankenswerterweise zur Verfügung gestellt wurde und vorwiegend an den Büchern (TT13; BZ14; Cum21). In den einzelnen Kapiteln werden weitere Referenzen genannt.

2 Logik

Die Logik als Teilgebiet der Mathematik ist eine wichtige Grundlage für die Informatik und tritt an vielen Stellen auf. Z.B. wird sie in Programmen verwendet, um den Programmfluss zu steuern, digitale Schaltungen basieren auf der Schaltungslogik, logische Programmiersprachen basieren auf Logiken, moderne Inferenzsysteme verwenden logische Schlüsse, in der Verifikation von Programmen wird die Spezifikation oft in einer Logik formuliert, u.s.w.

In diesem Kapitel betrachten wir im wesentlichen Aussagen und die Aussagenlogik. Schließlich führen wir Quantoren und Universen ein, die über die Aussagenlogik hinaus gehen.

Geht man formal vor, so legt man für eine Logik zunächst die *Syntax* und anschließend die *Semantik* fest. Die Syntax sagt, wie Formeln der Logik aufgebaut sind. Die Semantik definiert die Bedeutung der Formeln.

Es gibt viele Bücher, die einen Überblick über Aussagen, Aussagenlogik und Quantoren geben, und viele weiterführende Bücher zur Logik.

Z.B. ist in (TT13) eine knappe Einführung zu finden, ein Kapitel in (BZ14) widmet sich der Booleschen Algebra. Ein empfehlenswertes Buch zur Logik für Informatiker:innen (welches inhaltlich weit über den Inhalt dieses Kapitels hinaus geht) ist (KK06). Zur Einführung in die mathematische Logik sei (EFT18) empfohlen.

2.1 Aussagen und die Aussagenlogik

2.1.1 Was sind Aussagen?

Eine Aussage ist ein sprachlicher Satz, der entweder falsch oder wahr ist (aber nie beides gleichzeitig).

Eine wahre Aussage hat den Wahrheitswert „w“ (oder „1“ oder „true“), eine falsche Aussage hat den Wahrheitswert „f“ (oder „0“ oder „false“). In einem Schaltkreis entsprechen diese Wahrheitswerte den Zuständen „Strom fließt“ und „Strom fließt nicht“.

Aussagen müssen wahr oder falsch sein, aber wir müssen nicht notwendigerweise wissen, ob sie wahr oder falsch sind, sie sind trotzdem Aussagen.

Beispiel 2.1.1. *Beispiele für Aussagen und ihre Wahrheitswerte sind:*

- *Wiesbaden ist die Landeshauptstadt von Hessen.* (w)
- *Es gibt beliebig große Primzahlen.* (w)
- $10+20 = 42$ (f)

- *Alle Informatikstudierende lieben die Mathematik (unklar, vermutlich f)*
- *Jede gerade Zahl größer als 2 ist die Summe zweier Primzahlen (unbekannt). Dies ist die Goldbach-Vermutung, die bis heute unbewiesen ist.*
- *Jedes Java-Programm terminiert. (f)*

Beispiel 2.1.2. *Keine Aussagen sind zum Beispiel:*

- *Guten Morgen!*
- *$5+3$*
- *Wann ist die Vorlesung endlich vorbei?*
- *Keine Aussage ist folgendes Paradoxon: „Ich lüge gerade.“ Sowohl die Annahme, es sei wahr, als auch die Annahme, es sei falsch, führen zu einem Widerspruch!*

Übungsaufgabe 2.1.3. *Welche der folgenden Sätze sind Aussagen?*

- *Die Erde ist eine Scheibe.*
- *Lasst uns feiern!*
- *2 ist eine Primzahl.*
- *Wann machen wir Pause?*
- *Dieser Satz ist falsch.*
- *Die Kreiszahl π enthält jede beliebige Ziffernfolge.*

Können Sie bei allen Aussagen ihren Wahrheitswert angeben?

2.1.2 Syntax der Aussagenlogik

In diesem Abschnitt erläutern wir, wie wir zusammengesetzte Aussagen (Formeln) bilden können. Z.B. können wir aus den Aussagen „Es regnet.“ und „Die Straße ist nass.“ die Formel „Wenn es regnet, dann ist die Straße nass“ bilden. Damit es jedoch einfach und für jeden nachvollziehbar (und nicht zuletzt eindeutig) ist, verzichtet man auf sprachliche Konstruktionen, (wie „Wenn . . . , dann . . .“), sondern definiert eine Formelsprache¹.

Um Aussagen miteinander zu verknüpfen und eine handliche Syntax zu haben, abstrahieren wir im folgenden von den Sätzen als Aussagen und verwenden stattdessen Großbuchstaben für die Aussagen, wie A, B, C . Sie sind sozusagen Platzhalter für die konkreten Sätze. Wir nennen diese Platzhalter auch aussagenlogische Variablen (die je nach Wahrheitsgehalt mit wahr oder falsch belegt werden können).

Die Wahrheitswerte w und f und einzelne aussagenlogische Variablen bezeichnet man auch als *atomare Aussagen* (da man sie nicht weiter zerlegen kann).

¹Genau wie in der Arithmetik: Hier würde heutzutage niemand auf Idee kommen statt „ $3+5*2$ “ zu schreiben „addiere 3 zum Produkt aus 5 und 2“.

Aus den atomaren Aussagen kann man größere Aussagen zusammensetzen, indem man logische Operatoren auf sie anwendet². In der Aussagenlogik spricht man von aussagenlogischen Formeln für solche zusammengesetzten Aussagen.

Nun können wir die wichtigsten Operatoren (Verknüpfungen) einführen;

Definition 2.1.4 (Aussagenlogische Formeln). *Atomare Aussagen sind auch aussagenlogische Formeln. Wenn F und G aussagenlogische Formeln sind, dann sind auch*

- $(\neg F)$ (bezeichnet als Negation, gesprochen „nicht F “)
- $(F \wedge G)$ (bezeichnet als Konjunktion, auch logisches Und, gesprochen „ F und G “)
- $(F \vee G)$ (bezeichnet als Disjunktion, auch logisches Oder, gesprochen „ F oder G “)
- $(F \oplus G)$ (bezeichnet als Kontravalenz, auch exklusives Oder³, gesprochen „ F entweder oder G “)
- $(F \rightarrow G)$ (bezeichnet als Implikation, gesprochen „wenn F , dann G “)
- $(F \leftrightarrow G)$ (bezeichnet als Äquivalenz, auch Biimplikation, gesprochen „ F genau dann, wenn G “)

aussagenlogische Formeln.

Beispiel 2.1.5. *Man mache sich klar, dass $(A \vee (B \wedge \neg(C \rightarrow w)))$ eine aussagenlogische Formel ist, denn*

- A ist eine aussagenlogische Variable, daher eine atomare Aussage und daher auch eine aussagenlogische Formel.
- B ist eine aussagenlogische Variable, daher eine atomare Aussage und daher auch eine aussagenlogische Formel.
- C ist eine aussagenlogische Variable, daher eine atomare Aussage und daher auch eine aussagenlogische Formel.
- w ist eine atomare Aussage und daher auch eine aussagenlogische Formel
- Da C und w aussagenlogische Formeln sind, gilt dies auch für $(C \rightarrow w)$.
- Da $(C \rightarrow w)$ eine aussagenlogische Formel ist, gilt dies auch für $\neg(C \rightarrow w)$.
- Da B und $\neg(C \rightarrow w)$ aussagenlogische Formeln sind, gilt dies auch für $(B \wedge \neg(C \rightarrow w))$.
- Da A und $(B \wedge \neg(C \rightarrow w))$ aussagenlogische Formeln sind, gilt dies auch für $A \vee (B \wedge \neg(C \rightarrow w))$.

Die aussagenlogischen Operatoren $\neg, \vee, \wedge, \dots$ nennt man auch *Junktoren*. Dabei ist $\neg$ 1-stellig, während $\vee, \wedge, \oplus, \rightarrow, \leftrightarrow$ zweistellig sind. Wir zählen auch w und f zu den Junktoren, da sie 0-stellige Junktoren sind (sie erwarten gar keine Argumente).

²Analog zur Arithmetik, in der man aus 3 und 5 und dem Pluszeichen den arithmetischen Ausdruck $3 + 5$ zusammensetzt.

³Das Adjektiv *exklusiv* ist hier im Sinne von ausschließend aufzufassen.

2.1.3 Semantik der Aussagenlogik

Die Semantik der Junktoren ist wie folgt festgelegt: Sie sind Funktionen auf Wahrheitswerten. Der Einfachheit halber legen wir diese Funktionen mit *Wahrheitstabellen*⁴ fest:

Diese haben als Spalten zunächst die Eingaben und schließlich die Ausgabespalte für den entsprechenden Verknüpfung der Eingaben mit dem Junktor. Neben der Kopfzeile, die Eingaben und Ausgabe beschriften, gibt es je eine Zeile für jede mögliche Kombination der Eingaben

Definition 2.1.6. *Die Negation hat die folgende Semantik:*

A	$\neg A$
w	f
f	w

D.h. A ist genau dann eine wahre Aussage, wenn $(\neg A)$ falsch ist.

Z.B. ist $\neg(\text{Die Erde ist eine Scheibe})$ eine wahre Aussage.

Definition 2.1.7. *Die Konjunktion hat die folgende Semantik:*

A	B	$(A \wedge B)$
f	f	f
f	w	f
w	f	f
w	w	w

D.h. $(A \wedge B)$ ist genau dann eine wahre Aussage, wenn beide Aussagen A und B wahr sind.

Definition 2.1.8. *Die Disjunktion hat die folgende Semantik:*

A	B	$(A \vee B)$
f	f	f
f	w	w
w	f	w
w	w	w

D.h. $(A \vee B)$ ist genau dann eine wahre Aussage, wenn mindestens eine der beiden Aussagen A oder B wahr ist.

Im Gegensatz dazu verlangt die Kontravalenz, dass genau eine der Aussagen wahr ist:

⁴Wahrheitstabellen werden manchmal auch Wahrheitstafel, Wahrheitswerttabelle, Wahrheitswerttafel genannt.

Definition 2.1.9. Die Kontravalenz hat die folgende Semantik:

A	B	$(A \oplus B)$
f	f	f
f	w	w
w	f	w
w	w	f

Definition 2.1.10. Die Implikation hat die folgende Semantik:

A	B	$(A \rightarrow B)$
f	f	w
f	w	w
w	f	f
w	w	w

D.h. $(A \rightarrow B)$ ist genau dann wahr, wenn A falsch ist oder A und B beide wahr sind.

Definition 2.1.11. Die Äquivalenz hat die folgende Semantik:

A	B	$(A \leftrightarrow B)$
f	f	w
f	w	f
w	f	f
w	w	w

D.h. $(A \leftrightarrow B)$ ist genau dann wahr, wenn A und B denselben Wahrheitswert haben.

Der Wahrheitswert einer aussagenlogischen Formel hängt von den Wahrheitswerten der aussagenlogischen Variablen ab. Daher benötigt man zum Berechnen eines solchen Wahrheitswerts eine *Belegung* der Variablen. Eine Belegung $\mathcal{B}$ weist (endlich vielen) aussagenlogischen Variablen einen Wahrheitswert zu, d.h. für Variable A liefert $\mathcal{B}(A)$ den Wert f oder den Wert w .

Für eine Formel F ist eine *Belegung* $\mathcal{B}$ für F eine Belegung, die zumindest allen Variablen aus F einen Wahrheitswert zuweist. Hat man eine Formel F und eine Belegung $\mathcal{B}$ für F gegeben, so kann man den Wahrheitswert von F unter der Belegung $\mathcal{B}$ berechnen, indem man alle Variablen A in F durch ihre Belegung $\mathcal{B}(A)$ ersetzt und anschließend von innen nach außen den Wahrheitswert mittels der Semantik der Junktoren berechnet.

Wir notieren den Wahrheitswert einer Formel F unter einer Belegung $\mathcal{B}$ mit $\mathcal{B}(F)$. Wir gehen dabei stets davon aus, dass $\mathcal{B}$ eine Belegung für F ist, ohne dies jedes mal explizit auszuschreiben.

Beispiel 2.1.12. Betrachte die Formel $F = ((A \wedge \neg B) \rightarrow \neg(C \rightarrow A))$. Die Belegung $\mathcal{B}_1 = \{A \mapsto f, B \mapsto f, C \mapsto w\}$ ⁵ ist eine Belegung für F . Wir berechnen $\mathcal{B}_1(F)$:

⁵Wir schreiben hier die Belegung als Menge von einzelnen Abbildungen der Form „Variable $\mapsto$ Wahrheitswert“.

$$\begin{aligned}
\mathcal{B}_1(F) &= \mathcal{B}_1((A \wedge \neg B) \rightarrow \neg(C \rightarrow A)) \\
&= ((\mathcal{B}_1(A) \wedge \neg \mathcal{B}_1(B)) \rightarrow \neg(\mathcal{B}_1(C) \rightarrow \mathcal{B}_1(A))) \\
&= ((f \wedge \neg f) \rightarrow \neg(w \rightarrow f)) \\
&= ((f \wedge w) \rightarrow \neg(w \rightarrow f)) \\
&= (f \rightarrow \neg(w \rightarrow f)) \\
&= (f \rightarrow \neg f) \\
&= (f \rightarrow w) \\
&= w
\end{aligned}$$

Übungsaufgabe 2.1.13. Finde eine Belegung $\mathcal{B}_2$ für $F = ((A \wedge \neg B) \rightarrow \neg(C \rightarrow A))$, sodass $\mathcal{B}_2(F) = f$ gilt.

Möchte man alle Belegungen betrachten, so kann man wieder eine *Wahrheitstabelle* verwenden, welche den Wahrheitswertverlauf einer Formel tabellarisch auflistet: Die Tabelle hat dabei Spalten für jede Variable und für jede Teilformel und für die Formel selbst. Die Spalten für die Variablen werden Zeilenweise mit allen Belegungen für die Variablen gefüllt. Danach werden sukzessive die Wahrheitswerte der Teilformeln aus den bereits gefüllten Zellen berechnet bis schließlich in der letzten Spalte der jeweilige Wahrheitswert der Gesamtaussage berechnet wird.

Beispiel 2.1.14. Die Wahrheitstabelle für die Formel $(A \wedge \neg B) \rightarrow \neg(A \vee C)$ startet man mit einer Tabelle, welche Spalten für die Variablen A, B, C und für alle Teilformeln (ohne Doppelungen) und die Formel selbst enthält:

Schritt 1:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$

Anschließend kann man die Zeilen für die drei Variablen befüllen. Man sollte dies möglichst systematisch machen, damit man keine Belegung vergisst.

Schritt 2:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f					
f	f	w					
f	w	f					
f	w	w					
w	f	f					
w	f	w					
w	w	f					
w	w	w					

Anschließend kann man nun die Wahrheitswerte für die nächste der Teilformeln ausrechnen. Diese wendet immer einen Junktor auf Werte an, die man schon gefüllt hat. Für das Ausrechnen

kann man mit den entsprechenden Werten in der Wahrheitstabelle des Junktors nachschauen.
Wir schattieren zur Illustration die Spalten, deren Werte wir betrachten müssen:

Schritt 3:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f	w				
f	f	w	w				
f	w	f	f				
f	w	w	f				
w	f	f	w				
w	f	w	w				
w	w	f	f				
w	w	w	f				

Schritt 4:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f	w	f			
f	f	w	w	f			
f	w	f	f	f			
f	w	w	f	f			
w	f	f	w	w			
w	f	w	w	w			
w	w	f	f	f			
w	w	w	f	f			

Schritt 5:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f	w	f	f		
f	f	w	w	f	w		
f	w	f	f	f	f		
f	w	w	f	f	w		
w	f	f	w	w	w		
w	f	w	w	w	w		
w	w	f	f	f	w		
w	w	w	f	f	w		

Schritt 6:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f	w	f	f	w	
f	f	w	w	f	w	f	
f	w	f	f	f	f	w	
f	w	w	f	f	w	f	
w	f	f	w	w	w	f	
w	f	w	w	w	w	f	
w	w	f	f	f	w	f	
w	w	w	f	f	w	f	

Schritt 7:

A	B	C	$\neg B$	$(A \wedge \neg B)$	$A \vee C$	$\neg(A \vee C)$	$(A \wedge \neg B) \rightarrow \neg(A \vee C)$
f	f	f	w	f	f	w	w
f	f	w	w	f	w	f	w
f	w	f	f	f	f	w	w
f	w	w	f	f	w	f	w
w	f	f	w	w	w	f	f
w	f	w	w	w	w	f	f
w	w	f	f	f	w	f	w
w	w	w	f	f	w	f	w

Beachte, dass in jeder Wahrheitstabelle die Werte der Variablen nur einmal pro Zeile belegt werden, d.h. gleiche Vorkommen von Variablen in der Formel haben auch den gleichen Wert!

2.2 Gesetze der Logik

Um nicht immer Klammern setzen zu müssen, um eindeutige Formeln zu erhalten, legt man Prioritäten für die Junktoren fest.

Beispiel 2.2.1. Ohne solche Festlegungen ist z.B. nicht klar, welche Formel das Konstrukt $\neg A \vee B \wedge C$ meint. Es gibt verschiedene Möglichkeiten, was man sich klarmachen kann, indem man Klammern einfügt:

- $\neg((A \vee B) \wedge C)$
- $\neg(A \vee (B \wedge C))$
- $(\neg(A \vee B)) \wedge C$
- $((\neg A) \vee B) \wedge C$
- $((\neg A) \vee (B \wedge C))$

Wir legen fest: $\neg$ hat die höchste Priorität, $\wedge$ kommt vor $\oplus$, $\oplus$ kommt vor $\vee$, $\vee$ kommt vor $\rightarrow$, und $\rightarrow$ kommt vor $\leftrightarrow$. Hat ein Junktore höhere Priorität als ein anderer, so darf er seine Operanden zu erst an sich binden:

Z.B. in $\neg A \rightarrow B$ hat $\neg$ höhere Priorität als $\rightarrow$, daher wird erst A an $\neg$ gebunden, dann $(\neg A)$ und B an $\rightarrow$, d.h. das Konstrukt beschreibt die voll geklammerte Formel $((\neg A) \rightarrow B)$ und *nicht* die Formel $\neg(A \rightarrow B)$.

Die Formel $A \vee B \wedge C \rightarrow D$ meint die voll geklammerte Formel, $((A \vee (B \wedge C)) \rightarrow D$. Wenn man eine andere Klammerung möchte, muss man die Klammern explizit setzen. Wenn man sich nicht auskennt (z.B. die Prioritäten nicht kennt), kann man zusätzliche Klammern setzen⁶

Da wir keine weiteren Festlegungen getroffen haben, dürfen wir bisher z.B. $A \vee B \vee C$ nicht schreiben, da nicht klar ist, ob wir $((A \vee B) \vee C)$ oder $(A \vee (B \vee C))$ meinen, d.h. in solchen Fällen müssen immer Klammern gesetzt werden. Wir werden diese strikte Regel bald für die Operatoren $\wedge$, $\vee$, $\oplus$ aufheben aber für $\rightarrow$ und $\leftrightarrow$ stets beibehalten.

2.2.1 Tautologien und Kontradiktionen

Definition 2.2.2 (Tautologie). *Eine aussagenlogische Formel, die für jede Belegung wahr ist, nennt man eine Tautologie (oder auch allgemeingültige Formel)*

Tautologien sind daher die Sätze der Aussagenlogik, da sie stets gelten.

Man kann nachweisen, dass eine aussagenlogische Formel eine Tautologie ist, indem man die Wahrheitstabelle aufstellt und prüft, dass diese für jede Spalte den Wert w liefert.

Beispiel 2.2.3. $(A \wedge B) \rightarrow A$ ist eine Tautologie, denn

A	B	$A \wedge B$	$(A \wedge B) \rightarrow A$
f	f	f	w
f	w	f	w
w	f	f	w
w	w	w	w

Definition 2.2.4 (Widerspruch). *Eine aussagenlogische Formel, die für jede Belegung falsch ist, nennt man eine Kontradiktion (oder auch Widerspruch, widersprüchliche Formel, unerfüllbare Formel).*

Beispiel 2.2.5. $A \wedge \neg A$ ist eine Kontradiktion, denn

A	$\neg A$	$A \wedge \neg A$
f	w	f
w	f	f

⁶Kurz gesagt: „Zu viele Klammern setzen ist erlaubt, aber Klammern weglassen geht nur, wenn man die definierten Prioritäten verwendet.“

Es gibt Formeln, die weder Tautologien noch Widersprüche sind. Dies sind Formeln, für die es sowohl Belegungen gibt unter denen die Formel wahr ist als auch Belegungen unter denen die Formel falsch ist. Mit den Begriffen in der nächsten Definition, sind dies genau die Formeln, die sowohl erfüllbar als auch widerlegbar sind:

Definition 2.2.6 (Erfüllbarkeit und Widerlegbarkeit). *Eine aussagenlogische Formel, die für mindestens eine Belegung wahr ist, heißt erfüllbar. Eine aussagenlogische Formel, die für mindestens eine Belegung falsch ist, heißt widerlegbar.*

2.2.2 Logische Äquivalenz

Definition 2.2.7 (Logische Äquivalenz). *Zwei Formeln F und G heißen logisch äquivalent geschrieben $F \equiv G$, wenn sie für jede Belegung zum gleichen Wahrheitswert führen, d.h. für jede Belegung $\mathcal{B}$ gilt $\mathcal{B}(F) = \mathcal{B}(G)$.*

Man kann eine logische Äquivalenz $F \equiv G$ nachweisen, indem man eine Wahrheitstabelle mit Ausgabespalten für F und G aufstellt und dann vergleicht, dass die Ausgaben für jede Belegung identisch sind.

Beispiel 2.2.8. *Wir zeigen, dass $A \rightarrow B \equiv \neg A \vee B$ gilt:*

A	B	$\neg A$	$A \rightarrow B$	$\neg A \vee B$
f	f	w	w	w
f	w	w	w	w
w	f	f	f	f
w	w	f	w	w

Die Werte in den Spalten 4 und 5 sind identisch. Daher führen beide Formeln zum selben Wahrheitswert für jede Belegung. Damit ist gezeigt, dass sie logisch äquivalent sind.

Da die (syntaktische) Äquivalenz $\leftrightarrow$ verwendet werden kann, um Gleichheit von Wahrheitswerten zu prüfen, gilt offensichtlich:

Satz 2.2.9. *Zwei Formeln F und G sind genau dann logisch äquivalent, wenn $(F \leftrightarrow G)$ eine Tautologie ist.*

Beweis. Wir geben einen ausführlichen Beweis.

- Wir zeigen zunächst: „Wenn F und G logisch äquivalent sind, dann ist $(F \leftrightarrow G)$ eine Tautologie.“ Seien F und G logisch äquivalent. Sei $\mathcal{B}$ eine beliebige aber feste Belegung der aussagenlogischen Variablen. Da F und G logisch äquivalent, gilt $\mathcal{B}(F) = \mathcal{B}(G)$ und mit der Semantik von $\leftrightarrow$ (Definition 2.1.11) folgt $\mathcal{B}(F \leftrightarrow G) = w$. Da dies für jede Belegung gilt, folgt dass $F \leftrightarrow G$ eine Tautologie ist.

- Wir zeigen nun die Rückrichtung: „Wenn $(F \leftrightarrow G)$ eine Tautologie ist, dann sind F und G logisch äquivalent.“ Sei $(F \leftrightarrow G)$ eine Tautologie und sei $\mathcal{B}$ eine beliebige Belegung. Dann gilt $\mathcal{B}(F \leftrightarrow G) = w$. Dies kann aufgrund von Definition 2.1.11 nur sein, wenn einer der beiden folgenden Fälle gilt:

1. $\mathcal{B}(F) = w$ und $\mathcal{B}(G) = w$.
2. $\mathcal{B}(F) = f$ und $\mathcal{B}(G) = f$

Offensichtlich gilt in beiden Fällen $\mathcal{B}(F) = \mathcal{B}(G)$. Da $\mathcal{B}$ beliebige Belegung war, gilt dies für jede Belegung, was zeigt, dass F und G logisch äquivalent sind. $\square$

Logische äquivalente Formeln darf man austauschen, da sie den Wahrheitswert nicht verändern. Dies gilt auch für Teilformeln:

Satz 2.2.10. *Sei F eine aussagenlogische Formel, welche die Formel G als Teilformel enthält. Sei $G \equiv H$. Dann gilt $F \equiv F'$, wobei F' aus F entsteht, indem die Teilformel G durch H ersetzt wird.*

Beweisskizze. Sei $\mathcal{B}$ eine Belegung der Variablen aus F und F' . Die Wahrheitstabelle für F berechnet in einer Spalte den Wahrheitswert $\mathcal{B}(G)$ und verwendet diesen Wert, dann weiter, um $\mathcal{B}(F)$ zu berechnen. Die Wahrheitstabelle für F' kann daraus erzeugt werden, indem die Berechnung von $\mathcal{B}(G)$ durch die Berechnung von $\mathcal{B}(H)$ ersetzt wird und anschließend $\mathcal{B}(H)$ anstelle von $\mathcal{B}(G)$ verwendet wird. Da $G \equiv H$, gilt $\mathcal{B}(G) = \mathcal{B}(H)$, was zeigt, dass $\mathcal{B}(F) = \mathcal{B}(F')$ gelten wird. $\square$

Wir zeigen im folgenden, dass man in logisch äquivalenten Formeln Variablen durch Formeln ersetzen darf, ohne dass sich die logische Äquivalenz ändert. Wir definieren zunächst, was dieses Ersetzen genau meint.

Definition 2.2.11 (Substitution). *Seien F, G Formeln und A eine aussagenlogische Variable. Dann bezeichne $F[G/A]$ die Formel, die entsteht, indem alle Vorkommen von A in F durch die Formel G ersetzt werden.*

Wir verwenden dies auch für mehrere Variablen: Seien $F, G_1, \dots, G_n$ Formeln und $A_1, \dots, A_n$ aussagenlogische Variablen. Dann bezeichne $F[G_1/A_1, \dots, G_n/A_n]$ die Formel, die entsteht, indem jeweils für $1 \leq i \leq n$ alle Vorkommen von A_i in F durch die Formel G_i ersetzt werden.

Bemerkung 2.2.12. *Die Ersetzung bei mehreren Substitutionen sind parallel durchzuführen. Es gilt z.B. nicht immer $F[G/A][H/B] = F[G/A, H/B]$: Betrachte $F = A \wedge B$, $G = \neg B$, $H = C \vee D$. Dann ist $F[G/A, H/B] = A \wedge B[\neg B/A, C \vee D/B] = \neg B \wedge (C \vee D)$ aber $F[G/A][H/B] = A \wedge B[\neg B/A][C \vee D/B] = (\neg B \wedge B)[C \vee D/B] = \neg(C \vee D) \wedge (C \vee D)$.*

Satz 2.2.13. *Seien F, G, H aussagenlogische Formeln, A eine aussagenlogische Variable. Wenn $F \equiv G$ gilt, dann gilt auch $F[H/A] \equiv G[H/A]$*

Beweis. Sei $\mathcal{B}$ eine Belegung der Variablen aus $F[H/A]$ und $G[H/A]$. Betrachte die Belegung $\mathcal{B}'$ mit

$$\mathcal{B}'(X) = \begin{cases} \mathcal{B}(H) & \text{wenn } X = A \\ \mathcal{B}(X) & \text{wenn } X \neq A \end{cases}$$

Dann ist $\mathcal{B}'$ eine Belegung für F und G und offensichtlich gilt $\mathcal{B}(F[H/A]) = \mathcal{B}'(F)$ und $\mathcal{B}(G[H/A]) = \mathcal{B}'(G)$.

Da aus $F \equiv G$ auch $\mathcal{B}'(F) = \mathcal{B}'(G)$ folgt, ergibt sich insgesamt $\mathcal{B}(F[H/A]) = \mathcal{B}'(F) = \mathcal{B}'(G) = \mathcal{B}(G[H/A])$. $\square$

2.2.3 Rechengesetze

Durch Wahrheitstabellen kann man zahlreiche Rechengesetze für aussagenlogische Formeln als gültig nachweisen.

Satz 2.2.14 (Satz von de Morgan). *Seien F und G aussagenlogische Formeln. Dann gilt*

Erstes de Morgansches Gesetz: $\neg(F \wedge G) \equiv \neg F \vee \neg G$

Zweites de Morgansches Gesetz: $\neg(F \vee G) \equiv \neg F \wedge \neg G$

Beweis. Aus Satz 2.2.13 folgt, dass es es genügt, die Teilformeln F und G wie aussagenlogische Variablen zu behandeln und die Wahrheitstabelle für alle Belegungen dieser Variablen F und G zu berechnen:

F	G	$(F \wedge G)$	$\neg(F \wedge G)$	$\neg F$	$\neg G$	$\neg F \vee \neg G$
f	f	f	w	w	w	w
f	w	f	w	w	f	w
w	f	f	w	f	w	w
w	w	w	f	f	f	f

Da Spalten 4 und 7 identisch sind, folgt das erste de Morgansche Gesetz. Das zweite Gesetz lässt sich analog durch eine Wahrheitstabelle zeigen (Übungsaufgabe) $\square$

Die Junktoren $\neg$, $\wedge$, $\vee$ verhalten sich analog wie die Operation $-$, $\cdot$ und $+$ auf reellen Zahlen, d.h. bekannte Rechengesetze gelten für diese:

Satz 2.2.15 (Rechengesetze für $\wedge$, $\vee$, $\neg$). *Für alle aussagenlogischen Formeln F, G, H gelten die folgenden Gesetze:*

- *Kommutativgesetz: $F \wedge G \equiv G \wedge F$ und $F \vee G \equiv G \vee F$*
- *Assoziativgesetz: $(F \wedge G) \wedge H \equiv F \wedge (G \wedge H)$ und $(F \vee G) \vee H \equiv F \vee (G \vee H)$*
- *Distributivgesetz: $F \vee (G \wedge H) \equiv (F \vee G) \wedge (F \vee H)$ und $F \wedge (G \vee H) \equiv (F \wedge G) \vee (F \wedge H)$*
- *Existenz neutraler Elemente: $F \wedge w \equiv F$ und $F \vee f \equiv F$*
- *Existenz des Komplements: $F \wedge \neg F \equiv f$ und $F \vee \neg F \equiv w$*

Beweis. Sämtliche Gesetze lassen sich mit Wahrheitstabellen beweisen. Wir machen dies exemplarisch für das erste Distributivgesetz:

F	G	H	$G \wedge H$	$F \vee (G \wedge H)$	$F \vee G$	$F \vee H$	$(F \vee G) \wedge (F \vee H)$
f	f	f	f	f	f	f	f
f	f	w	f	f	f	w	f
f	w	f	f	f	w	f	f
f	w	w	w	w	w	w	w
w	f	f	f	w	w	w	w
w	f	w	f	w	w	w	w
w	w	f	f	w	w	w	w
w	w	w	w	w	w	w	w

Da die 5. und die letzte Spalte identische Wahrheitswerte haben, sind die Formeln logisch äquivalent und das Rechengesetz gilt. $\square$

Die Distributivgesetze sind von links nach rechts angewendet sind analog zum Ausmultiplizieren. Von rechts nach links angewendet, sind sie analog zum Ausklammern.

Übungsaufgabe 2.2.16. Zeige, dass $\oplus$ kommutativ und assoziativ ist.

Da $\vee, \wedge, \oplus$ assoziativ sind, lassen wir Klammern manchmal weg und schreiben $F \vee G \vee H$, $F \wedge G \wedge H$ bzw. $F \oplus G \oplus H$.

Eine Menge mit ausgezeichneten Elementen w, f und Operationen $\wedge, \vee, \neg$ auf dieser Menge, für welche die Rechengesetze aus Satz 2.2.15 gelten, nennt man *Boolesche Algebra* (benannt nach dem Mathematiker George Boole). D.h. insbesondere die Menge $\{f, w\}$ mit den bereits definierten Operationen $\wedge, \vee, \neg$ ist eine Boolesche Algebra. Es gibt auch weitere Boolesche Algebren, wie wir später noch sehen werden.

2.2.4 Dualität und weitere Rechengesetze

Betrachtet man die Gesetze der Booleschen Algebra, so stellt man fest, dass diese immer aus zwei Teilen bestehen, die symmetrisch sind: Man erhält den einen Teil aus dem Anderen, indem man $\wedge$ mit $\vee$ sowie f mit w vertauscht.

Diese Eigenschaft der Booleschen Algebra nennt man *Dualität*. Erhält man einen Satz durch einen anderen Satz, indem man $\wedge$ mit $\vee$ sowie f mit w vertauscht, so nennt man diesen Satz *dual* zu dem anderen Satz.

Wir verwenden die Dualität im Beweis der folgenden Gesetze:

Satz 2.2.17. Für alle aussagenlogischen Formeln F und G gelten die folgenden Gesetze:

- *Absorptionsgesetze:* $F \wedge (F \vee G) \equiv F$ und $F \vee (F \wedge G) \equiv F$
- *Idempotenzgesetze:* $F \vee F \equiv F$ und $F \wedge F \equiv F$

- *Involutionsgesetz (doppelte Negation):* $\neg(\neg F) \equiv F$
- *Extremalgesetze:* $F \vee w \equiv w$ und $F \wedge f \equiv f$

Beweis. Das Involutionsgesetz und das Extremalgesetz können mittels einer Wahrheitstabelle verifiziert werden.

Wir wenden die Rechengesetze aus Satz 2.2.15 an, um das erste Absorptionsgesetz nachzuweisen:

$$\begin{aligned}
 & F \wedge (F \vee G) \\
 \equiv & (F \vee f) \wedge (F \vee G) && \text{(Existenz neutraler Elemente)} \\
 \equiv & (F \vee (f \wedge G)) && \text{(Distributivgesetz)} \\
 \equiv & (F \vee (G \wedge f)) && \text{(Kommutativgesetz)} \\
 \equiv & (F \vee f) && \text{(Extremalgesetz)} \\
 \equiv & F && \text{(Existenz neutraler Elemente)}
 \end{aligned}$$

Aufgrund der Dualität können wir in allen Zwischenschritten $\wedge$ mit $\vee$ und f mit w vertauschen, was direkt das zweite Absorptionsgesetz zeigt.

Für den Beweis des 1. Idempotenzgesetz formen wir wie folgt um, wobei wir das 2. Absorptionsgesetz direkt verwenden:

$$\begin{aligned}
 & F \vee F \\
 \equiv & F \vee (F \wedge w) && \text{(Existenz neutraler Elemente)} \\
 \equiv & F && \text{(Absorptionsgesetz)}
 \end{aligned}$$

Da das 1. Absorptionsgesetz dual zum 2. Absorptionsgesetz ist, folgt das 2. Idempotenzgesetz mittels der Dualität.

□

Die Rechengesetze können verwendet werden, um aussagenlogische Formeln umzuformen und zu vereinfachen.

Beispiel 2.2.18. Wir vereinfachen die Formel $\neg(A \vee \neg B) \vee (A \wedge B)$

$$\begin{aligned}
 & \neg(A \vee \neg B) \vee (A \wedge B) \\
 \equiv & (\neg A \wedge \neg \neg B) \vee (A \wedge B) && \text{(De Morgansches Gesetz)} \\
 \equiv & (\neg A \wedge B) \vee (A \wedge B) && \text{(Involutionsgesetz)} \\
 \equiv & (B \wedge \neg A) \vee (B \wedge A) && \text{(Kommutativgesetz, 2 Mal)} \\
 \equiv & B \wedge (\neg A \vee A) && \text{(Distributivgesetz)} \\
 \equiv & B \wedge (A \vee \neg A) && \text{(Kommutativgesetz)} \\
 \equiv & B \wedge w && \text{(Komplement)} \\
 \equiv & B && \text{(Neutrale Elemente)}
 \end{aligned}$$

2.3 Boolesche Funktionen und vollständige Mengen von Junktoren

2.3.1 Boolesche Funktionen

Eine aussagenlogische Formel mit Variablen $A_1, \dots, A_n$ kann man als Boolesche Funktion h auffassen, welche die Booleschen Eingaben $A_1, \dots, A_n$ auf den Booleschen Wert $h(A_1, \dots, A_n)$ abbildet. Umgekehrt kann man sich Gedanken über solche Funktionen im allgemeinen machen. 0-stellige Boolesche Funktionen gibt es genau zwei: $h_1 = w$ und $h_2 = f$. Einstellige Funktionen gibt es vier: Je eine Funktion, die ihre Eingabe ignoriert und konstant auf f bzw. w abbildet, die Identitätsfunktion, und die Funktion, die ihre Eingabe negiert. Für die zweistelligen Funktionen gibt es bereits sechzehn Möglichkeiten, u.s.w..

Satz 2.3.1. *Für jede Boolesche Funktion h gibt es eine aussagenlogische Formel, welche h berechnet.*

Beweisskizze. Für den Sonderfall, dass die Boolesche Funktion h 0-stellig ist, also keine Eingaben erhält kann man direkt den Wahrheitswert f oder w als Formel nehmen. Anderenfalls sei h eine n -stellige Boolesche Funktion mit $n > 0$. Stelle eine Wahrheitstabelle für h auf, die für jede Belegung $\mathcal{B}$ der Eingaben $A_1, \dots, A_n$ den Funktionswert $h(\mathcal{B}(A_1), \dots, \mathcal{B}(A_n))$ enthält. Sei m die Anzahl an Zeilen in der Wahrheitstabelle. Für jede Zeile k der Wahrheitstabelle, sei $\mathcal{B}_k$ die Belegung der Eingaben $A_1, \dots, A_n$ in Zeile k . Wenn $h(\mathcal{B}_k(A_1), \dots, \mathcal{B}_k(A_n)) = w$ gilt, so erzeuge die Teilformel $F_k := (\mathcal{L}(A_1) \wedge \dots \wedge \mathcal{L}(A_n))$ wobei

$$\mathcal{L}(A_i) = \begin{cases} A_i & \text{wenn } \mathcal{B}(A_i) = w \\ \neg A_i & \text{wenn } \mathcal{B}(A_i) = f \end{cases}$$

Ansonsten erzeuge keine Formel F_k

Schließlich bilde die Formel F als Disjunktion aller Formeln F_i (gibt es keine Formel F_i , dann setze $F := f$).

Dann gilt für alle Wahrheitswerte $b_1, \dots, b_n$: $h(b_1, \dots, b_n) \equiv F[b_1/A_1, \dots, b_n/A_n]$.

Wir sparen uns an dieser Stelle den genauen Beweis für und argumentieren nur: Wenn $h(b_1, \dots, b_n) = w$, dann gibt es für die Belegung $\mathcal{B} = \{A_1 \mapsto b_1, \dots, A_n \mapsto b_n\}$ eine Teilformel F_i , sodass $\mathcal{B}(\mathcal{L}(A_1)) = w, \dots, \mathcal{B}(\mathcal{L}(A_n)) \equiv w$ gilt. Daher ist $\mathcal{B}(F_i) \equiv w$ (mehrfach das Gesetz der Existenz neutraler Elemente anwenden) und auch $\mathcal{B}(F) = F[b_1/A_1, \dots, b_n/A_n] \equiv w$ (mehrfach das Extremalgesetz anwenden). Umgekehrt, wenn $h(b_1, \dots, b_n) = f$, dann ist jede der Teilformel F_i für die Belegung $\mathcal{B} = \{A_1 \mapsto b_1 \mapsto A_1, \dots, A_n \mapsto b_n\}$ logisch-äquivalent zu f , es mindestens ein j gibt mit $\mathcal{B}(\mathcal{L}(A_j)) = f$. Anwenden der Rechengesetze zeigt dann, dass auch $\mathcal{B}(F) \equiv f$ gelten muss. $\square$

Die im Beweis erzeugte Formel hat eine spezielle Form: Sie ist eine Disjunktion von Konjunktionen von Literalen (ein Literal ist Variable oder eine negierte Variable). Daher nennt man diese Form: *Disjunktive Normalform*.

2.3.2 Vollständige Mengen an Junktoren

Eine Menge M von Junktoren⁷ heißt vollständig, wenn für jede aussagenlogische Formel eine logisch äquivalente Formel existiert, die nur die Junktoren aus M benutzt.

Z.B. ist $\{w, f, \vee, \wedge\}$ keine vollständige Menge an Junktoren, da man z.B. für die Formel $\neg A$ keine logisch äquivalente Formel findet, die nur $\vee$ und $\wedge$ verwendet.

Die Menge $\{\vee, \wedge, \neg\}$ ist eine vollständige Menge an Junktoren:

Dies gilt, da $w \equiv A \vee \neg A$, $f \equiv A \wedge \neg A$, $F \rightarrow G \equiv \neg F \vee G$, $F \leftrightarrow G \equiv F \rightarrow G \wedge G \rightarrow F$ und $F \oplus G \equiv (F \vee G) \wedge \neg(F \wedge G)$.

Beispiel 2.3.2. Es gilt, dass $\{\neg, \vee\}$ und $\{\neg, \wedge\}$ jeweils vollständige Mengen von Junktoren sind: Es genügt zu zeigen, dass die Junktoren $\neg, \vee, \wedge$ ausgedrückt werden können. Um dies zu zeigen, ist es hinreichend zu zeigen, dass $\vee$ und $\wedge$ jeweils durch die anderen beiden Junktoren ausgedrückt werden können:

Es gilt $F \wedge G \equiv \neg(\neg F \vee \neg G)$ (dies kann leicht mit dem De Morganschen Gesetz und dem Involutionsgesetz verifiziert werden). Dual dazu gilt $F \vee G \equiv \neg(\neg F \wedge \neg G)$.

Übungsaufgabe 2.3.3. Zeige, dass $\{\oplus, \vee, w\}$ eine vollständige Menge von Junktoren ist.

2.4 Anwendungen der Aussagenlogik

2.4.1 Schaltungsentwurf

Eine wichtige Anwendung der Aussagenlogik in der Informatik ist der Schaltungsentwurf. Dabei stellen Schaltkreise aussagenlogische Formeln dar. Wir illustrieren dies mit einer 2-aus-3-Schaltung: Diese solle nur dann wahr liefern (z.B. einen Tresor öffnen), wenn mindestens 2 der 3 Eingänge (die z.B. Schlösser symbolisieren) wahr sind. Die Wahrheitstabelle der gesuchten Funktion h hat daher die Form:

A	B	C	$h(A, B, C)$
f	f	f	f
f	f	w	f
f	w	f	f
f	w	w	w
w	f	f	f
w	f	w	w
w	w	f	w
w	w	w	w

Mit dem in Abschnitt 2.3.1 vorgestellten Verfahren, gewinnen wir die Disjunktive Normalform:

⁷Beachte: Hier zählen die Wahrheitswerte w und f auch zu den Junktoren

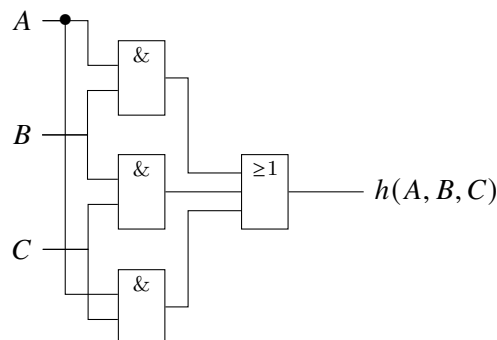
$$F := (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge C) \vee (A \wedge B \wedge \neg C) \vee (A \wedge B \wedge C)$$

Diese Formel kann man vereinfachen: Entweder versucht man sich selbst, mit den Rechengesetzen, oder man verwendet algorithmische Verfahren (z.B. sogenannte KV-Diagramme), oder man benutzt Software wie z.B. wolframalpha.com.

Der Vollständigkeit halber geben wir die Rechnung hier an (Kommutativgesetz und Assoziativgesetz wenden wir implizit an, ohne dies als einzelne Rechenschritte zu kennzeichnen):

$$\begin{aligned}
 & (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge C) \vee (A \wedge B \wedge \neg C) \vee (A \wedge B \wedge C) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge C) \vee ((A \wedge B) \vee (\neg C \wedge C)) \quad (\text{Ausklammern von } A \wedge B) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge C) \vee ((A \wedge B) \wedge w) \quad (\text{Existenz des Komplements}) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge \neg B \wedge C) \vee (A \wedge B) \quad (\text{Idempotenzgesetz}) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge ((\neg B \wedge C) \vee B)) \quad (\text{Ausklammern von } A) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge ((B \vee \neg B) \wedge (B \vee C))) \quad (\text{Ausmultiplizieren}) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge (w \wedge (B \vee C))) \quad (\text{Existenz des Komplements}) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge (B \vee C)) \quad (\text{Idempotenzgesetz}) \\
 \equiv & (\neg A \wedge B \wedge C) \vee (A \wedge B) \vee (A \wedge C) \quad (\text{Ausmultiplizieren}) \\
 \equiv & (C \wedge ((\neg A \wedge B) \vee A)) \vee (A \wedge B) \quad (\text{Ausklammern von } C) \\
 \equiv & (C \wedge ((A \vee \neg A) \wedge (A \vee B))) \vee (A \wedge B) \quad (\text{Ausmultiplizieren}) \\
 \equiv & (C \wedge (w \wedge (A \vee B))) \vee (A \wedge B) \quad (\text{Existenz des Komplements}) \\
 \equiv & (C \wedge ((A \vee B))) \vee (A \wedge B) \quad (\text{Idempotenzgesetz}) \\
 \equiv & (C \wedge A) \vee (C \wedge B) \vee (A \wedge B) \quad (\text{Ausmultiplizieren})
 \end{aligned}$$

Es gilt $F \equiv (A \wedge B) \vee (A \wedge C) \vee (B \wedge C)$. Hierfür können wir den Schaltkreis leicht aufzeichnen (ein Und-Gatter wird als $\boxed{\&}$ und ein Oder-Gatter wird als $\boxed{\geq 1}$ gezeichnet, links sind die Eingänge, rechts der Ausgang.)



Ein Ziel beim Schaltungsentwurf ist es, möglichst wenige Bauteile zu verwenden. Daher sucht man möglichst kleine Formeln.

2.4.2 Modelle suchen

Für viele Anwendungen möchte man keine Tautologien finden, sondern man ist an erfüllbaren Formeln interessiert und sucht nach der erfüllenden Belegung (eine solche Belegung nennt man auch *Modell*). Es gibt spezielle Algorithmen und Software, um dies in der Praxis möglichst schnell zu tun. Man nennt diese Werkzeuge SAT-Solver, da sie versuchen nachzuweisen, dass eine Formel *satisfiable* (englisch für erfüllbar) ist. Man kann damit vielfältige Aufgaben lösen, z.B. auch Rätsel: Man modelliert die gegebenen Fakten als aussagenlogische Formel (oft als spezielle Normalform, die konjunktive Normalform wird hier öfter verwendet als die disjunktive) und steckt diese als Eingabe in den SAT-Solver. Dieser berechnet die erfüllende Belegung, welche eine Lösung des Rätsels liefert.

Wir demonstrieren dies an einem Rätsel von Raymond Smullyan, aber suchen die erfüllende Belegung durch Hinschauen:

Die Frage nach dem Pfefferdieb. Es gibt drei Verdächtige: Den Hutmacher, den Schnapphasen und die Haselmaus. Folgendes ist bekannt:

1. Genau einer von ihnen ist der Dieb.
2. Unschuldige sagen immer die Wahrheit
3. Schnappphase: Der Hutmacher ist unschuldig.
4. Hutmacher: Die Haselmaus ist unschuldig

Wir kodieren dies als aussagenlogische Formel und verwenden Variablen H , S und M , die wahr sein sollen, wenn der Hutmacher, der Schnappphase, oder die Maus schuldig sind. Kodieren der Fakten ergibt

- Für 1.: $(H \vee S \vee M) \wedge \neg(H \wedge S) \wedge \neg(H \wedge M) \wedge \neg(S \wedge M)$
- Für 2. und 3.: $\neg S \rightarrow \neg H$
- Für 2. und 4.: $\neg H \rightarrow \neg M$

Die Konjunktion der drei Formeln kann man nun vereinfachen, was zeigt, dass die Formel äquivalent zur folgenden Formel ist: Vereinfachen der Formeln ergibt

$$\neg H \wedge \neg M \wedge S$$

Hier sieht man nun leicht, dass die Formel nur eine erfüllende Belegung hat: $\{H \mapsto f, M \mapsto f, S \mapsto w\}$, d.h. der Schnappphase war der Dieb.

2.5 Prädikatenlogische Formeln: Formeln mit Quantoren

Während wir für die Aussagenlogik sehr formal Syntax und Semantik eingeführt haben, werden wir nun die prädikatenlogischen Formeln eher informell betrachten. Das wesentliche Lernziel ist es, den Umgang mit Quantoren zu erlernen.

2.5.1 Aussageformen

Bisher haben wir nur Formeln betrachtet, die Aussagen enthalten, die nicht variabel über Objekten sind. Z.B. hatten wir Aussagen wie „2 ist eine Primzahl“ aber keine Aussagen, die so etwas repräsentieren wie „ x ist eine Primzahl“ (wobei x eine Variable ist). Jetzt wollen wir dies tun, indem wir sogenannte Aussageformen erlauben. Dies sind Sätze mit Variablen für Objekte. Setzt man konkrete Werte für die Variablen ein, so erhält man wieder Aussagen. Das entspricht genau dem obigen Beispiel: „ x ist eine Primzahl“ ist noch keine Aussage (und der Wahrheitsgehalt ist auch nicht bestimmbar), aber setzen wir 2 für x ein, so erhalten wir wieder einen Satz „2 ist eine Primzahl“ (der in diesem Fall wahr ist). Setzen wir 4 für x ein, so erhalten wir „4 ist eine Primzahl“ – ein falscher Satz. Wir könnten aber auch „Grün“ einsetzen, doch ist der Satz „Grün ist eine Primzahl“ wahr oder falsch? Um dieses Problem zu umgehen, schränken wir die Werte ein, die für x eingesetzt werden. In unserem Beispiel würden nur natürliche Zahlen sinnvoll sein. Im Satz „ x ist eine Grundfarbe der additiven Farbmischung“ wären alle Farben eine sinnvolle Menge und im Satz „Man benötigt x Fußballspieler:innen pro Team, damit ein Spiel angepfiffen werden darf.“ wären die Zahlen $\{0, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11\}$ sinnvoll. Wir nennen eine solche Menge, aus denen man die Werte für die Variable x wählen darf, im folgenden ein *Universum*. Wir nehmen stets an, dass ein Universum mindestens ein Element enthält.

Definition 2.5.1 (Aussageform). *Seien $x_1, \dots, x_n$ Variablen, und $U_1, \dots, U_n$ Universen. Dann ist eine Aussageform ein Satz mit Variablen $x_1, \dots, x_n$, sodass nach Ersetzen aller Variablen x_i durch Objekte $u_i \in U_i$ ein Satz entsteht der wahr oder falsch ist.*

In den aussagenlogischen Formeln haben wir Aussagen durch aussagenlogische Variablen repräsentiert. Analog dazu repräsentieren wir nun Aussageformen durch Prädikatensymbole $P, Q, R, \dots$. Diese sehen wie Funktionsaufrufe aus, da sie als Argumente, die Variablen x_i erhalten, die sie benötigen. Semantisch sind sie wie Funktionen zu sehen, die nach Anwendung auf konkrete Werte aus den Universen zu einem Wahrheitswert werden. Z.B. können wir schreiben $P_1(x) = \text{„}x \text{ ist Primzahl“}$ und verwenden dann $P_1(x)$ in unseren Formeln und schreiben $P_1(2)$ für die Aussage „2 ist Primzahl“. Die Prädikate können von mehreren Variablen abhängen, z.B. $Q(x, y, z) := \text{„}x + y = z\text{“}$. Dann ist $Q(1, 2, 3)$ eine wahre Aussage, aber $Q(2, 2, 3)$ ist falsch. Auch Verknüpfungen mehreren Aussageformen zu größeren Formeln ist möglich und erlaubt, z.B. $P_1(x) \wedge Q(x, y, z)$ u.s.w..

2.5.2 All- und Existenzaussagen

Da wir nun Variablen für Objekte in den Formeln eingeführt haben, können wir auch Aussagen daraus machen, ohne konkrete Objekte einzusetzen, sondern stattdessen über die Objekte aus dem Universum zu quantifizieren.

2.5.2.1 Allaussagen

Allaussagen sind Aussagen über *alle Elemente eines Universums*. Beispiele für Allaussagen sind:

- Alle Primzahlen > 10 sind ungerade.
- In jedem Dreieck beträgt die Summe der Winkel 180 Grad.
- Alle Menschen sind sterblich.
- Jede der Zahlen 15, 27, 69 ist ungerade.

Solche Allaussagen können stets in die Form

Für alle x aus einem Universum U gilt

gebracht werden. Für die Beispiele:

- Für alle x aus $\mathbb{N}$ gilt: Wenn $x > 10$ und x Primzahl, dann ist x ungerade.
- Für alle x aus der Menge der Dreiecke gilt: Winkelsumme von x ist 180 Grad.
- Für alle x aus der Menge der Menschen gilt: x ist sterblich.
- Für alle x aus $\{15, 27, 69\}$ gilt: x ist ungerade.

Um dieses „Für alle x aus einem Universum U . . . “ symbolisch als Formel auszudrücken, wird der Allquantor $\forall$ verwendet. Die Syntax ist $\forall x \in U : P(x)$, wenn $P(x)$ eine Aussageform ist. Z.B. $\forall x \in \mathbb{R} : x^2 \geq 0$ oder $\forall x \in \{17, 39, 47\} : \text{ungerade}(x)$. Diese Semantik des Allquantors ist, dass die Formel $\forall x \in U : P(x)$ genau dann wahr ist, wenn $P(u)$ für jedes $u \in U$ eine wahre Aussage ist. Damit sind $\forall x \in \mathbb{R} : x^2 \geq 0$ und $\forall x \in \{17, 39, 47\} : \text{ungerade}(x)$ wahre Aussagen, aber $\forall x \in \mathbb{N} : \text{Wenn } x \text{ Primzahl, dann ist } x \text{ ungerade}$ als auch die Aussage $\forall x \in \{2, 3, 4, 5\} : \text{gerade}(x)$ sind beide falsch: „Wenn 2 Primzahl, dann ist 2 ungerade“ ist falsch und $\text{gerade}(2)$ und $\text{gerade}(4)$ sind beide falsch.

Wenn das Universum U eine *endliche* Menge ist, kann man den Allquantor als Abkürzung ansehen und man kann ihn eliminieren:

Satz 2.5.2. Sei $U = \{u_1, \dots, u_n\}$ ein endliches Universum. Dann gilt $\forall u \in U : P(u) \equiv P(u_1) \wedge \dots \wedge P(u_n)$.

D.h. in diesem Fall ist der Allquantor eine „lange“ Und-Aussage.

Beispiel 2.5.3. $\forall x \in \{2, 3, 4, 5\} : \text{gerade}(x) \equiv \text{gerade}(2) \wedge \text{gerade}(3) \wedge \text{gerade}(4) \wedge \text{gerade}(5)$

2.5.2.2 Existenzaussagen

Existenzaussagen sind Aussagen die fordern, dass *mindestens ein Element eines Universums* eine gewisse Eigenschaft hat. Beispiele für Existenzaussagen sind:

- Es gibt eine gerade Primzahl.
- Sei $U = \{1, 3, 5, 7, 8\}$. Dann gibt es ein $x \in U$, für das gilt: x ist gerade.

Solche Existenzaussagen können stets in die Form

Es gibt ein x aus einem Universum U , für das gilt

gebracht werden. Für die Beispiele:

- Es gibt ein $x \in \mathbb{N} : x$ ist gerade Primzahl
- Es gibt ein $x \in \{1, 3, 5, 7, 8\} : x$ ist gerade

Um dieses „Es gibt x aus einem Universum U , für das gilt:“ symbolisch als Formel auszudrücken, wird der Existenzquantor $\exists$ verwendet. Die Syntax ist $\exists x \in U : P(x)$, wenn $P(x)$ eine Aussageform ist. Z.B. $\exists x \in \mathbb{R} : x^2 = 2$ oder $\exists x \in \{2, 4, 5, 9\} : \text{gerade}(x)$. Die Semantik des Existenzquantors ist, dass die Formel $\exists x \in U : P(x)$ genau dann wahr ist, wenn $P(u)$ für *mindestens eine* $u \in U$ eine wahre Aussage ist. Damit sind $\exists x \in \mathbb{R} : x^2 = 2$ (denn $\sqrt{2}^2 = 2$ ist wahr) und $\exists x \in \{2, 4, 5, 9\} : \text{gerade}(x)$ (denn $\text{gerade}(2)$ ist wahr) wahre Aussagen aber $\exists x \in \mathbb{N} : x^2 = 2$ als auch die Aussage $\exists x \in \{3, 5, 7, 9\} : \text{gerade}(x)$ sind beide falsch.

Wenn das Universum U eine *endliche* Menge ist, kann man den Existenzquantor als Abkürzung ansehen und man kann ihn eliminieren:

Satz 2.5.4. Sei $U = \{u_1, \dots, u_n\}$ ein endliches Universum. Dann gilt $\exists u \in U : P(u) \equiv P(u_1) \vee \dots \vee P(u_n)$.

D.h. in diesem Fall ist der Existenzquantor eine „lange“ Oder-Aussage.

Beispiel 2.5.5. $\exists x \in \{2, 3, 4, 5\} : \text{gerade}(x) \equiv \text{gerade}(2) \vee \text{gerade}(3) \vee \text{gerade}(4) \vee \text{gerade}(5)$

Wir fassen die Syntax der prädikatenlogischen Formeln in der folgenden Definition zusammen:

Definition 2.5.6 (Prädikatenlogische Formeln). Atomare Aussagen sind auch prädikatenlogische Formeln. Wenn P ein n -stelliges Prädikat ist, dann ist $P(x_1, \dots, x_n)$ eine prädikatenlogische Formel. Wenn F und G prädikatenlogische Formeln sind und U ein Universum ist, dann sind auch

- $(\neg F)$
- $(F \wedge G)$
- $(F \vee G)$
- $(F \oplus G)$
- $(F \rightarrow G)$
- $(F \leftrightarrow G)$
- $\forall x \in U : F$
- $\exists x \in U : F$

prädikatenlogische Formeln.

Beachte, dass wir hier unter den Quantoren beliebige Formeln erlauben. Das schließt den Fall ein, dass x gar nicht vorkommt, aber auch den Fall, dass wir Verschachtelungen von Quantoren erlauben: Damit können wir z.B. ausdrücken

- Wenn zwei Zahlen nicht größer oder kleiner zueinander sind, dann müssen sie gleich sein:
 $\forall x \in \mathbb{N} : \forall y \in \mathbb{N} : (\neg(x < y) \wedge \neg(y < x)) \rightarrow x = y$

- Für jede natürliche Zahl, gibt es eine größere natürliche Zahl, die Primzahl ist $\forall x \in \mathbb{N} : \exists y \in \mathbb{N} : x < y \wedge y \text{ ist Primzahl}$

Wenn das Universum U klar ist, dann lassen wir es manchmal weg und schreiben $\forall x : F$ statt $\forall x \in U : F$ (analog für den Existenzquantor).

2.5.3 Rechenregeln für Quantoren

2.5.3.1 Negation von All- und Existenzaussagen

Negation vor einem Quantor, kann man über den Quantor schieben, wobei sich der Quantor umdreht: Aus einem Allquantor wird ein Existenzquantor und umgekehrt. D.h.

Satz 2.5.7. *Es gelten die logischen Äquivalenzen:*

- $\neg \forall x \in U : P(x) \equiv \exists x \in U : \neg P(x)$
- $\neg \exists x \in U : P(x) \equiv \forall x \in U : \neg P(x)$.

Beide Teile sind auch sprachlich begründbar:

- Wenn $P(x)$ nicht für alle $x \in U$ gilt, dann muss es ein $u \in U$ geben, für welches $P(u)$ falsch ist, d.h. $\neg P(u)$ gilt.
- Wenn es kein $u \in U$ gibt, sodass $P(u)$ gilt, dann ist $P(x)$ für alle $x \in U$ falsch.

Beispiel 2.5.8. *Wir zeigen einige Beispiele für das Schieben der Negation:*

- Die Negation von „Alle Menschen sind schlau“ ergibt „Es gibt (mindestens) einen nicht schlauen Menschen“
- Die Negation von „Es gibt eine schwere Klausur“ ergibt „Alle Klausuren sind nicht schwer“.
- $\neg(\exists x \in \mathbb{N} : x^2 = 2) \equiv \forall x \in \mathbb{N} : x^2 \neq 2$
- $\neg(\forall x \in \mathbb{R} : x^2 \neq 2) \equiv \exists x \in \mathbb{R} : x^2 = 2$

2.5.4 Quantorregeln

Man kann Allquantoren über die Konjunktion und Existenzquantoren über die Disjunktion ziehen, wenn sie über dasselbe Universum quantifizieren.

Satz 2.5.9. *Die folgenden logischen Äquivalenzen gelten:*

- $(\forall x \in U : P(x)) \wedge (\forall x \in U : Q(x)) \equiv (\forall x \in U : (P(x) \wedge Q(x)))$
- $(\exists x \in U : P(x)) \vee (\exists x \in U : Q(x)) \equiv (\exists x \in U : (P(x) \vee Q(x)))$

Beispiel 2.5.10. *Wir zeigen einige Beispiele für die Quantorregeln aus Satz 2.5.9: Betrachte die Sätze „Alle Autos haben ein Lenkrad und alle Autos haben mindestens drei Räder“. Diese Satz ist logisch äquivalent zu „Alle Autos haben ein Lenkrad und mindestens drei Räder.“. Die Formel $(\exists x \in \mathbb{N} : x > 10) \vee (\exists x \in \mathbb{N} : x < 5)$ ist logisch äquivalent zu $(\exists x \in \mathbb{N} : x > 10 \vee x < 5)$.*

Beachte, dass es im Allgemeinen nicht richtig ist, den Allquantor über eine Disjunktion oder den Existenzquantor über eine Konjunktion zu ziehen.

Beispiel 2.5.11. Betrachte den Satz „Es gibt ein Auto mit mehr als vier Rädern und es gibt ein Auto mit höchstens drei Rädern“. Der Satz ist sicher nicht äquivalent zu „Es gibt ein Auto mit mehr als vier und höchstens drei Rädern“.

„Für alle Primzahlen p gilt: p ist ungerade oder $p = 2$ “ ist eine wahre Aussage, aber „(Für alle Primzahlen p gilt: p ist ungerade) oder (Für alle Primzahlen p gilt: $p = 2$)“ ist falsch und nicht logisch äquivalent zur vorherigen.

2.5.5 Vertauschungsregeln

Man kann gleichartige Quantoren, die hintereinander stehen vertauschen:

Satz 2.5.12. Für Universen U_1, U_2 und Formeln F gelten die folgenden Äquivalenzen:

- $\forall x \in U_1 : \forall y \in U_2 : F \equiv \forall y \in U_2 : \forall x \in U_1 : F$
- $\exists x \in U_1 : \exists y \in U_2 : F \equiv \exists y \in U_2 : \exists x \in U_1 : F$

Beispiel 2.5.13. Zum Beispiel sind die Formeln $\forall x \in \mathbb{R} : \forall y \in \mathbb{R} : x \geq y \vee x < y$ und $\forall y \in \mathbb{R} : \forall x \in \mathbb{R} : x \geq y \vee x < y$ logisch äquivalent.

Beachte, dass dies nicht gilt für unterschiedliche Quantoren:

Beispiel 2.5.14. Die Formel $\forall x \in \mathbb{N} : \exists y \in \mathbb{N} : x^2 = y$ ist wahr (das Quadrat jeder natürlichen Zahl ist wieder eine natürliche Zahl), aber $\exists y \in \mathbb{N} : \forall x \in \mathbb{N} : x^2 = y$ (das Quadrat jeder natürlichen Zahl ist ein und dieselbe natürliche Zahl) ist falsch.

„Es gibt einen Studierenden, der alle Aufgaben in der Klausur lösen kann.“ (in Quantorschreibweise: „ $\exists x \in \text{Studierende} : \forall y \in \text{Aufgaben} : x \text{ löst } y$ “) ist nicht äquivalent zur Aussage, dass jede Aufgabe der Klausur durch irgendeinen Studierenden gelöst werden kann (in Quantorschreibweise: „ $\forall y \in \text{Aufgaben} : \exists x \in \text{Studierende} : x \text{ löst } y$ “).

2.6 Schlussbemerkungen

Zu Formeln mit Quantoren gäbe es noch viel mehr zu sagen, die Rahmen der Veranstaltung aber sprengen würden, und daher in vertiefenden Veranstaltungen behandelt werden. Zu guter Letzt sein noch angemerkt, dass in vielen sprachlichen und mathematischen Sätzen oft implizit allquantifiziert wird, ohne dass der Quantor genannt wird. Z.B. der Satz zur ersten binomischen Formel wird oft als $(a + b)^2 = a^2 + b^2 + 2ab$ geschrieben. Implizit ist gemeint, dass diese Gleichung für alle $a, b \in \mathbb{R}$ gilt.

3 Mengen

3.1 Was sind Mengen?

Eine mathematisch einwandfreie Definition einer Menge ist schwierig. Georg Cantor definierte in (Can95) „Unter einer ‘Menge’ verstehen wir jede Zusammenfassung M von bestimmten wohlunterschiedenen Objekten m unserer Anschauung oder unseres Denkens (welche die ‘Elemente’ von M genannt werden) zu einem Ganzen.“

Verwendet man diese Definition, so kann dies zu Widersprüchen führen, z.B. zur Russellschen Antinomie (Rus37): Bertrand Russel definierte die Menge $R := \{x \mid x \notin x\}$ – die Menge aller Mengen, die sich nicht selbst enthalten. Diese Definition führt zum Widerspruch, denn weder $R \in R$ noch $R \notin R$ kann gelten.

Anstatt nun in die axiomatische Mengenlehre einzusteigen, halten wir es einfach und definieren Mengen nicht genau, sondern verwenden zwei Verfahren, um Mengen zu beschreiben: Mengenbeschreibungen durch Aufzählung und Mengenbeschreibungen durch Eigenschaften. Zunächst führen wir einige Notationen ein.

Die Objekte einer Menge nennt man *Elemente*. Wir können eine Menge aufschreiben, indem wir ihre Elemente aufzählen. Z.B. ist $M_1 = \{1, 2, 3\}$ die Menge der Zahlen 1,2 und 3 und $M_2 = \{\text{Schwarz}, \text{Weiß}\}$ ist die Menge der Farben Schwarz und Weiß.

Definition 3.1.1 (Notation für Mengen). *Die Elemente einer Menge werden mit Kommas getrennt und durch geschweifte Klammern $\{$ und $\}$ umschlossen.*

Wenn ein Objekt m in einer Menge M enthalten ist, so schreiben wir $m \in M$. Wenn das Objekt m kein Element der Menge M ist, so schreiben wir $m \notin M$.

Die Menge, welche keine Elemente enthält, nennt man leere Menge. Wir schreiben $\emptyset$ oder $\{\}$ für die leere Menge.

Beispiel 3.1.2. Für $M = \{2, 4, 6, 9\}$ gilt $4 \in M$ und $5 \notin M$. Ebenso gilt $4 \notin \emptyset$ und $\text{Weiß} \notin \{\}$.

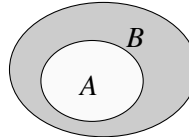
Definition 3.1.3 (Teilmenge und Gleichheit). *Eine Menge M ist eine Teilmenge von N , wenn N alle Elemente von M enthält, d.h. wenn $\forall x \in M : x \in N$ gilt. Wir schreiben dann $M \subseteq N$.*

Die Menge M ist eine echte Teilmenge von N , wenn $M \subseteq N$ gilt und N mindestens ein Element enthält, welches nicht in M enthalten ist. D.h. es gilt $\forall x \in M : x \in N \wedge \exists y \in N : y \notin M$. Wir schreiben dann $M \subset N$.

Zwei Mengen M und N sind gleich (geschrieben $M = N$), wenn $M \subseteq N$ und $N \subseteq M$ gilt.

Unmittelbar aus der Definition der Mengengleichheit folgt, dass die Reihenfolge der Elemente in einer Menge keine Rolle spielt. Z.B. ist $\{1, 2, 3\} = \{3, 2, 1\}$. Ebenso folgt, dass Elemente in einer Menge nicht mehrfach betrachtet werden. Z.B. ist $\{1, 2, 2, 3, 3, 3\} = \{1, 2, 3\}$.

Mengen kann man durch Venn-Diagramme darstellen (benannt nach John Venn): Dabei verwendet man Kreise oder Ellipsen zum Symbolisieren der Mengen und ordnet diese so an, wie sie den Beziehungen zwischen den Mengen entsprechen. Gilt $A \subset B$ so zeichnet man:



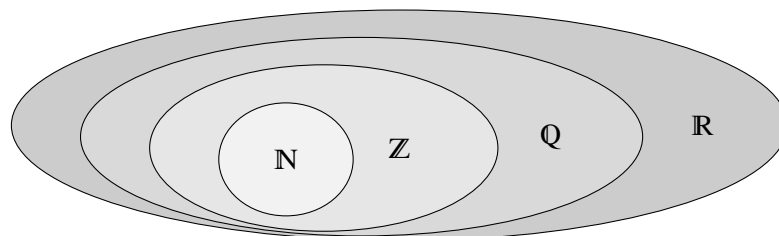
Definition 3.1.4 (Zahlenmengen). Wir definieren Symbole für bekannte Zahlenmengen:

- Wir schreiben $\mathbb{N}$ für die Menge der natürlichen Zahlen ($\mathbb{N} := \{1, 2, 3, \dots\}$) und $\mathbb{N}_0$ für die Menge der natürlichen Zahlen und 0 ($\mathbb{N}_0 = \{0, 1, 2, 3, \dots\}$).
- Wir schreiben $\mathbb{Z}$ für die Menge der ganzen Zahlen ($\mathbb{Z} := \{\dots, -2, -1, 0, 1, 2, \dots\}$).
- Wir schreiben $\mathbb{Q}$ für die Menge der rationalen Zahlen (alle gekürzten Brüche, bzw. endliche oder periodische Dezimalzahlen).
- Wir schreiben $\mathbb{R}$ für die Menge der reellen Zahlen (alle (endliche, periodische, nicht periodische) Dezimalzahlen).

Die Zahlenmengen aus Definition 3.1.4 sind jeweils echte Teilmengen: $\mathbb{N} \subset \mathbb{Z}$, $\mathbb{Z} \subset \mathbb{Q}$, $\mathbb{Q} \subset \mathbb{R}$. Die echten Teilmengenbeziehungen werden z.B. bezeugt durch:

- $-1 \in \mathbb{Z}$ aber $-1 \notin \mathbb{N}$
- $\frac{2}{3} \in \mathbb{Q}$ aber $\frac{2}{3} \notin \mathbb{Z}$
- $\sqrt{2} \in \mathbb{R}$ aber $\sqrt{2} \notin \mathbb{Q}$

Als Venn-Diagramm:



Das Beschreiben einer Menge durch Aufzählen ist oft mühsam oder schwierig oder gar unmöglich. Eine andere Möglichkeit, Mengen zu beschreiben, ist es, aus einer gegebenen Menge nur die Elemente herauszufiltern, die eine bestimmte Eigenschaft haben. Die Eigenschaften beschreibt man mittels einer Aussageform $P(x)$.

Definition 3.1.5 (Mengenbeschreibung durch Eigenschaft). Sei M eine bereits definierte Menge und $P(x)$ eine Aussageform (mit x aus M). Dann beschreibt $\{x \in M \mid P(x)\}$ die Teilmenge von M , die alle Elemente x enthält, die $P(x)$ erfüllen.

Das Konstrukt $\{x \in M \mid P(x)\}$ wird dabei als „mit der Eigenschaft“ gelesen. D.h. die Menge besteht aus „allen Elemente x aus M mit der Eigenschaft $P(x)$ “.

Beispiel 3.1.6. Wir beschreiben einige Mengen mit einer Mengenbeschreibung durch Eigenschaft: Sei $ungerade(x)$ eine Aussageform, die wahr ist, wenn x eine ungerade Zahl ist. Die Menge alle ungeraden natürlichen Zahlen kann dann durch

$$\{x \in \mathbb{N} \mid ungerade(x)\}$$

beschrieben werden.

Die natürlichen Zahlen können als Teilmenge der ganzen Zahlen mittels

$$\{x \in \mathbb{Z} \mid x > 0\}$$

beschrieben werden.

Die Menge der Primzahlen lässt sich durch

$$\{x \in \mathbb{N} \mid x > 1 \text{ und } x \text{ ist nur durch } 1 \text{ und sich selbst teilbar}\}$$

beschreiben.

Als Erweiterung der Mengenbeschreibung durch Eigenschaft erlaubt man auch die folgende Syntax

$$\{h(x) \mid x \in X, P(x)\}$$

Dabei ist h eine Funktion, die auf x angewendet wird und $P(x)$ eine Aussageform. Beachte, dass in diesem Fall, $x \in X$ hinter $\mid$ steht.

Beispiel 3.1.7. Die Menge aller Quadratzahlen größer als 999 lässt sich beschreiben durch

$$\{x^2 \mid x \in \mathbb{N}, x^2 > 999\}.$$

Die Menge aller Kubikzahlen lässt sich beschreiben durch

$$\{x^3 \mid x \in \mathbb{N}\}.$$

Im letzten Beispiel haben wir die Aussageform weggelassen. Damit ist gemeint, dass wir eine Aussageform verwenden, die immer wahr ist.

Übungsaufgabe 3.1.8. Stelle die folgenden Mengen durch Aufzählen der Elemente dar:

- $\{x \in \mathbb{N} \mid x \text{ teilt die Zahl } 42\}$
- $\{x \in \mathbb{Z} \mid x \text{ ist ungerade und } -4 \leq x < 5\}$

Definition 3.1.9 (Intervalle). Für $a < b \in \mathbb{R}$ definieren wir Schreibweisen für Intervalle:

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

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

Linksseitig halboffenes Intervall: $(a, b] := \{x \in \mathbb{R} \mid a < x \leq b\}$

Rechtsseitig halboffenes Intervall: $[a, b) := \{x \in \mathbb{R} \mid a \leq x < b\}$

Manchmal werden bei offenen oder halboffenen Intervallen anstelle der runden Klammer eine eckige benutzt, die „falsch herum“ ist (d.h. $]a, b[$, $]a, b]$ und $[a, b[$).

Für eine Menge M ist die Aussage $x \in \{x \in M \mid P(x)\}$ äquivalent zu $x \in M \wedge P(x)$. Wenn M ein Universum ist, welches aus dem Kontext klar ist, kann man auch direkt $P(x)$ daraus machen.

Analog erlauben wir anstelle von $\{x \in M \mid P(x)\}$ auch die Schreibweise $\{x \mid x \in M \wedge P(x)\}$ oder auch $\{x \mid P(x)\}$, wenn klar ist aus welchem Universum (d.h. aus welcher Menge) die Elemente x stammen.

3.1.1 Operationen auf Mengen

Definition 3.1.10 (Schnitt, Vereinigung, Differenz). *Seien M und N Mengen.*

1. Der Schnitt von M und N , geschrieben $M \cap N$, enthält alle Elemente, die sowohl in M als auch in N sind, d.h.

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

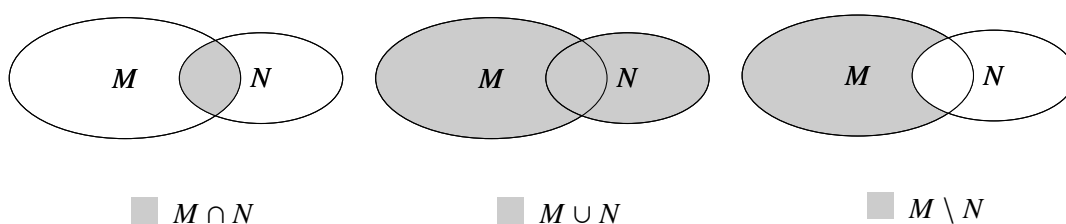
2. Die Vereinigung von M und N , geschrieben $M \cup N$, enthält alle Elemente, die in M oder in N (oder in beiden Mengen) liegen, d.h.

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

3. Die Differenzmenge $M \setminus N$ (M ohne N) enthält alle Elemente, die in M , aber nicht in N liegen:

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

Venn-Diagramme zur Illustration von Vereinigung, Schnitt und Differenz sind:



Beispiel 3.1.11. Wenn $A = \{1, 2, 3\}$ und $B = \{2, 4, 6\}$, dann ist $A \cap B = \{2\}$ und $A \cup B = \{1, 2, 3, 4, 6\}$ und $A \setminus B = \{1, 3\}$

Wenn M die Menge aller Informatikstudierenden und N die Menge aller Erstsemester ist, dann ist $M \cap N$ die Menge aller Informatikstudierenden im ersten Semester. Die Menge $M \cup N$ ist die

Menge aller Studierenden, die im ersten Semester sind oder Informatik studieren (oder beides).
Die Menge $M \setminus N$ sind alle Informatik-Studierende, die nicht im ersten Semester studieren.

Wenn C die Menge aller durch 3 teilbaren Zahlen ist, und D die Menge aller geraden Zahlen ist, dann ist $C \cap D$ die Menge aller durch 6 teilbaren Zahlen. Die Menge $C \cup D$ enthält genau die Zahlen, die durch 2 oder durch 3 teilbar sind. Die Menge $C \setminus D$ enthält genau die Zahlen, die durch 3 aber nicht durch 6 teilbar sind.

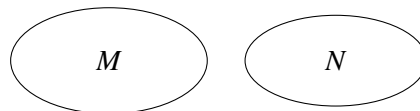
Übungsaufgabe 3.1.12. Bestimme die folgenden Mengen für $A = \{4, 8, 12, 16, 20\}$, $B = \{2, 6, 10, 14, 18\}$ und $C = \{6, 12, 18\}$ jeweils die Ergebnismengen:

- $A \cap B$
- $A \cap C$
- $B \cap C$
- $A \cup B$
- $A \cup B \cup C$
- $A \setminus B$
- $A \setminus C$
- $A \setminus (B \cap C)$

Definition 3.1.13 (Disjunkte Mengen). Zwei Mengen M und N bezeichnet man als disjunkt, wenn sie kein gemeinsames Element haben, d.h. $M \cap N = \emptyset$.

Z.B. sind die Mengen $\{1, 2, 3\}$ und $\{4, 5, 6\}$ disjunkt. Die Menge aller roten Socken und die Menge aller grünen Hemden sind disjunkt. Auch die Mengen $\{x \in U \mid P(x)\}$ und $\{x \in U \mid \neg P(x)\}$ sind für jede Menge U und Aussageform $P(x)$ disjunkt.

Eine Veranschaulichung disjunkter Mengen als Venn-Diagramm ist:



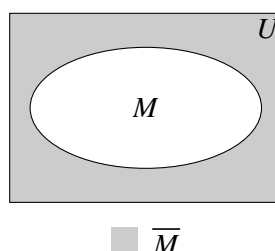
Hat man eine Grundmenge (ein Universum) und eine Teilmenge dieser Grundmenge, so kann man das Komplement als Differenz von U und M bilden:

Definition 3.1.14 (Komplement). Sei U ein Universum, $M \subseteq U$ eine Teilmenge, dann definieren wir das Komplement $\overline{M}$ von M als: $\overline{M} := U \setminus M$.

Manche Bücher verwenden anstelle von $\overline{M}$ auch M^C .

Beispiel 3.1.15. Sei U die Grundmenge aller Menschen und E die Menge aller Volljährigen. Dann ist $\overline{E}$ die Menge aller Kinder und Jugendlichen. Sei $\mathbb{N}$ die Grundmenge, G die Menge aller geraden natürlichen Zahlen. Dann ist $\overline{G}$ die Menge aller positiven ungeraden Zahlen.

Das Komplement kann durch folgendes Venn-Diagramm veranschaulicht werden:



3.1.2 Mengen und Logik

Sämtliche Mengenoperationen werden letztlich auf Logikoperationen zurückgeführt, denn für ein Universum U und Teilmengen M, N davon gilt:

- Schnitt: $M \cap N := \{x \in U \mid x \in M \wedge x \in N\}$
- Vereinigung: $M \cup N := \{x \in U \mid x \in M \vee x \in N\}$
- Differenz: $M \setminus N := \{x \in U \mid x \in M \wedge \neg(x \in N)\}$
- Komplement: $\overline{M} := \{x \in U \mid \neg(x \in M)\}$

Da $\cap, \cup, \neg$ auf die logischen Verknüpfungen $\wedge, \vee, \neg$ zurückgeführt werden können, übertragen sich die Gesetze der Logik ebenfalls auf die Mengenoperationen:

Satz 3.1.16 (Satz von De Morgan). *Seien M und N Teilmengen einer Menge U . Dann gilt:*

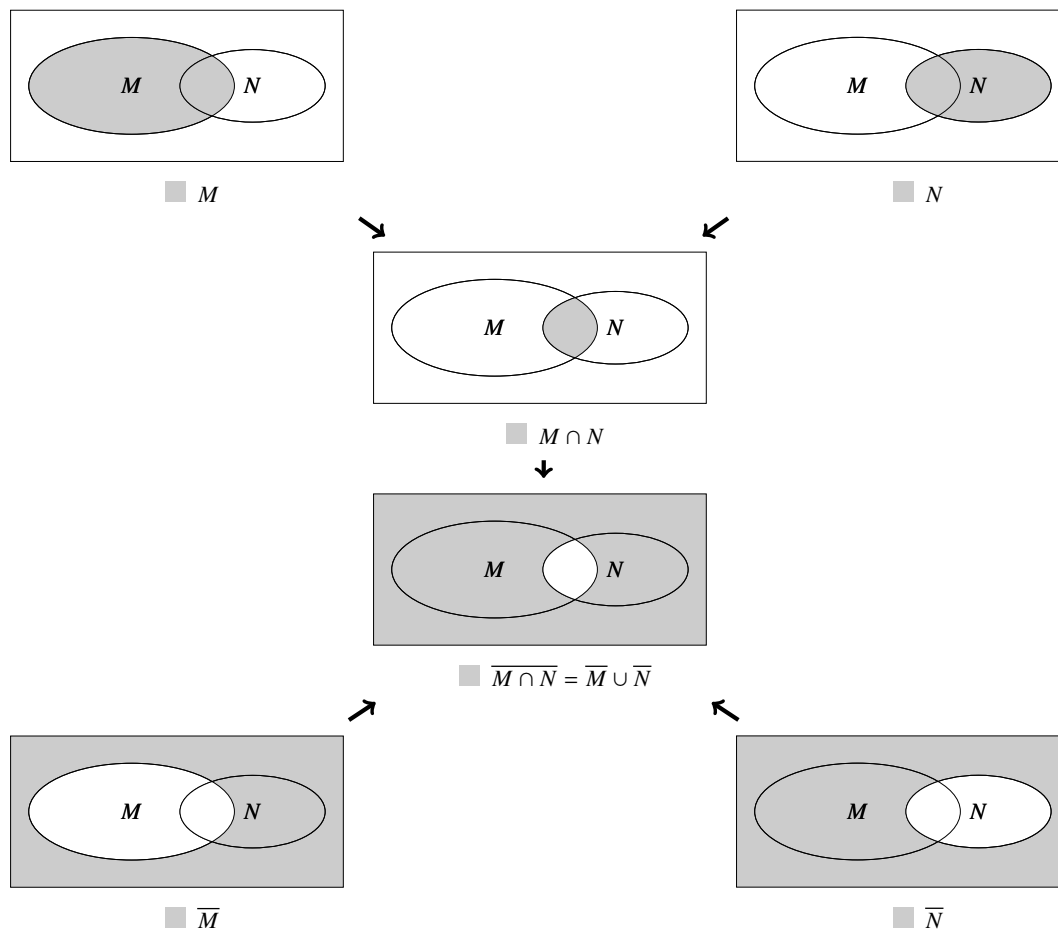
1. $\overline{M \cap N} = \overline{M} \cup \overline{N}$
2. $\overline{M \cup N} = \overline{M} \cap \overline{N}$

Beweis. Wir zeigen den ersten Teil. Der zweite Teil kann analog bewiesen werden:

$$\begin{aligned}
 & \overline{M \cap N} \\
 &= \{x \in U \mid \neg(x \in (M \cap N))\} && \text{(Einsetzen } \neg) \\
 &= \{x \in U \mid \neg(x \in \{x \in U \mid x \in M \wedge x \in N\})\} && \text{(Einsetzen } \cap) \\
 &= \{x \in U \mid \neg(x \in M \wedge x \in N)\} && \text{(Vereinfachung)} \\
 &= \{x \in U \mid \neg(x \in M) \vee \neg(x \in N)\} && \text{(logischer De Morgan)} \\
 &= \{x \in U \mid x \in \{x \in U \mid \neg(x \in M)\} \vee x \in \{x \in U \mid \neg(x \in N)\}\} && \text{(Vereinfachung)} \\
 &= \{x \in U \mid \neg(x \in M)\} \cup \{x \in U \mid \neg(x \in N)\} && \text{(Einsetzen } \cup) \\
 &= \overline{M} \cup \overline{N} && \text{(Einsetzen } \neg)
 \end{aligned}$$

□

Eine Illustration des 1. De Morganschen Gesetz für Mengen mit Venn-Diagrammen ist:

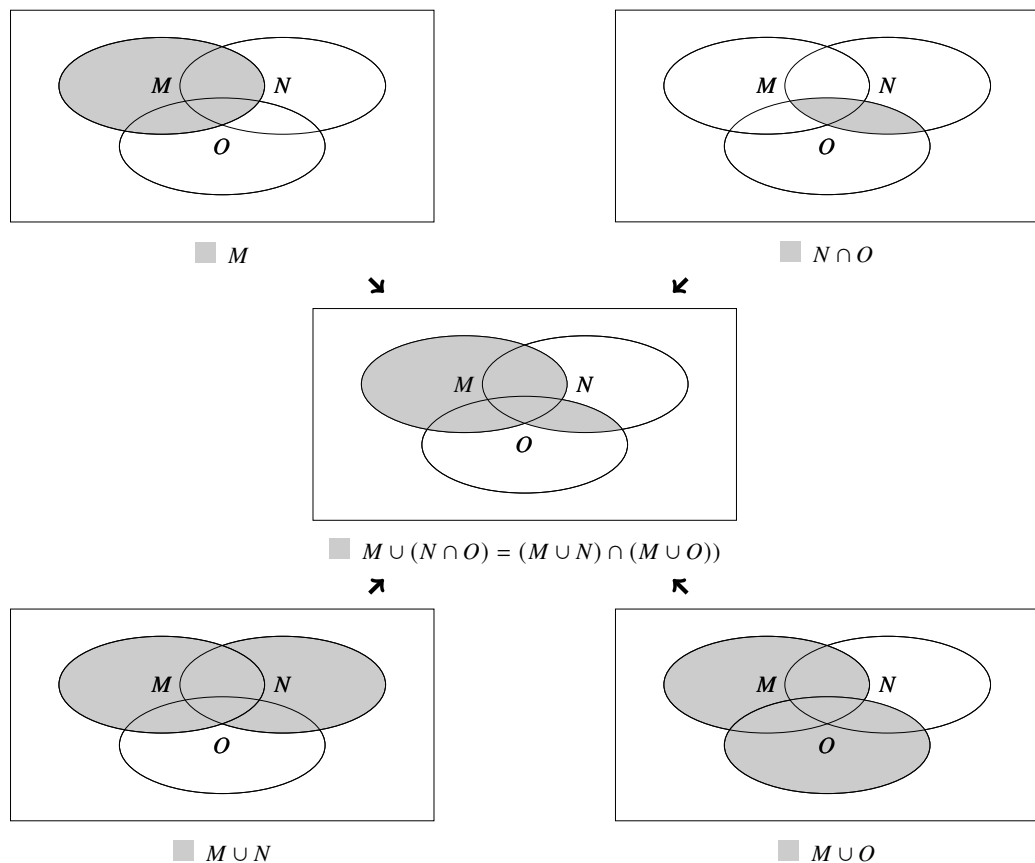


Analog zum Satz von de Morgan lassen sich auch weitere Gesetze direkt von der Logik auf die Mengenoperationen übertragen:

Satz 3.1.17 (Rechengesetze für Mengen). Für alle Teilmengen M, N, O der Grundmenge U gelten die folgenden Gesetze:

1. Kommutativgesetz: $M \cap N = N \cap M$ und $M \cup N = N \cup M$
2. Assoziativgesetz: $((M \cap N) \cap O) = (M \cap (N \cap O))$ und $((M \cup N) \cup O) = (M \cup (N \cup O))$
3. Distributivgesetz: $M \cup (N \cap O) = (M \cup N) \cap (M \cup O)$ und $M \cap (N \cup O) = (M \cap N) \cup (M \cap O)$
4. Existenz neutraler Elemente: $M \cap U = M$ und $M \cup \emptyset = M$
5. Existenz des Komplements: $M \cap \overline{M} = \emptyset$ und $M \cup \overline{M} = U$

Man kann die Gesetze mit Venn-Diagrammen veranschaulichen. Wir machen das exemplarisch für das 1. Distributivgesetz:



Übungsaufgabe 3.1.18.

- Veranschauliche das zweite Distributivgesetz mit Venn-Diagrammen.
- Veranschauliche das Gesetz $(M \setminus N) \cap O = (M \cap O) \setminus (N \cap O)$ mit Venn-Diagrammen.
- Warum ist $(M \setminus (N \setminus O))$ nicht notwendigerweise gleich zu $((M \setminus N) \setminus O)$? Gebe Mengen M, N, O an, sodass die Gleichheit gilt.

Auch die weiteren Gesetze der Logik (siehe Satz 2.2.17) können auf Mengen übertragen werden:

Satz 3.1.19. Seien M, N Teilmengen von U . Dann gilt:

1. Absorptionsgesetze: $M \cap (M \cup N) = M$ und $M \cup (M \cap N) = M$
2. Idempotenzgesetze: $M \cup M = M$ und $M \cap M = M$
3. Involutionsgesetz (doppeltes Komplement): $\overline{\overline{M}} = M$
4. Extremalgesetze: $M \cup U = U$ und $M \cap \emptyset = \emptyset$

3.2 Kartesisches Produkt

Definition 3.2.1 (Paare und Tupel). Seien $M_1, \dots, M_n$ Mengen und $x_1 \in M_1, \dots, x_n \in M_n$. Dann nennt man $(x_1, \dots, x_n)$ ein n -Tupel (oder kurz, ein Tupel).

Für $n = 2$ spricht man auch von einem (geordneten) Paar und 3-Tupel werden auch als Tripel bezeichnet.

Man beachte, dass die Reihenfolge der Elemente in einem n -Tupel relevant ist, d.h. $(1, \text{Grün}, 3) \neq (1, 3, \text{Grün})$. Mehrfaches Auftreten gleicher Elemente ist innerhalb eines Tupels erlaubt und wird unterschieden: Z.B. ist $(1, 1, 2, 2, 3)$ ein 5-Tupel und es gilt $(1, 2, 3) \neq (1, 1, 2, 2, 3)!$

Das kartesische Produkt erzeugt Paare aus zwei Mengen:

Definition 3.2.2 (Kartesisches Produkt). Für zwei Mengen M und N sei das kartesische Produkt $M \times N$ (auch Kreuzprodukt, gesprochen „ M kreuz N “) definiert als die Menge aller Paare (x, y) , sodass x aus M und y aus N stammt:

$$M \times N = \{(x, y) \mid x \in M, y \in N\}$$

Wenn M oder N die leere Menge ist, so ist das Kreuzprodukt leer (d.h. $M \times \emptyset = \emptyset = \emptyset \times M$).

Beispiel 3.2.3. Wir nennen einige Beispiele:

- Für $M = \{1\}$ und $N = \{4\}$ ist $M \times N = \{(1, 4)\}$ und $N \times M = \{(4, 1)\}$.
- Für $M = \{1, 2\}$ und $N = \{4, 5, 6\}$ ist

$$\begin{aligned} M \times N &= \{(1, 4), (1, 5), (1, 6), (2, 4), (2, 5), (2, 6)\} \text{ und} \\ N \times M &= \{(4, 1), (4, 2), (5, 1), (5, 2), (6, 1), (6, 2)\} \end{aligned}$$

- Für $M = \{A, B, C, D, E, F, G, H\}$ und $N = \{1, 2, 3, 4, 5, 6, 7, 8\}$ beschreibt $M \times N$ die Menge aller Felder eines Schachbretts.
- Für $M = \{\spadesuit, \clubsuit, \diamondsuit, \heartsuit\}$ und $N = \{7, 8, 9, 10, B, D, K, A\}$ stellt $M \times N$ eine Repräsentation der Spielkarten eines Skatblatts dar.
- Für $M = \{1, 2\}$, $N = \{A, B, C\}$, $O = \{\text{Rot}, \text{Grün}\}$ ist

$$\begin{aligned} (M \times N) \times O &= \{((1, A), \text{Rot}), ((1, B), \text{Rot}), ((1, C), \text{Rot}), \\ &\quad ((2, A), \text{Rot}), ((2, B), \text{Rot}), ((2, C), \text{Rot}), \\ &\quad ((1, A), \text{Grün}), ((1, B), \text{Grün}), ((1, C), \text{Grün}), \\ &\quad ((2, A), \text{Grün}), ((2, B), \text{Grün}), ((2, C), \text{Grün})\} \text{ und} \\ M \times (N \times O) &= \{(1, (A, \text{Rot})), (1, (A, \text{Grün})), (1, (B, \text{Rot})), \\ &\quad (1, (B, \text{Grün})), (1, (C, \text{Rot})), (1, (C, \text{Grün})), \\ &\quad (2, (A, \text{Rot})), (2, (A, \text{Grün})), (2, (B, \text{Rot})), \\ &\quad (2, (B, \text{Grün})), (2, (C, \text{Rot})), (2, (C, \text{Grün}))\} \end{aligned}$$

- Das Kreuzprodukt $\mathbb{Z} \times \mathbb{N}$ ist eine Repräsentation der rationalen Zahlen, wobei Elemente Paare der Form (Zähler, Nenner) sind.

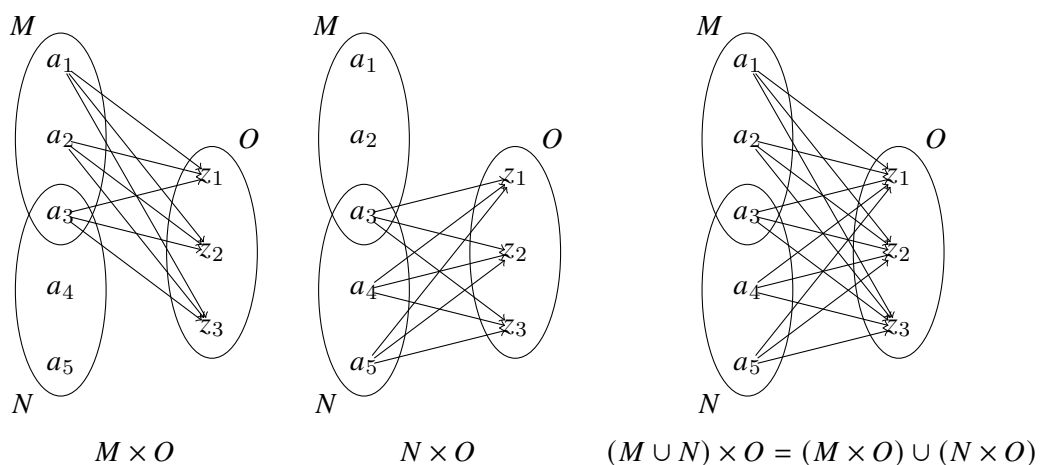
Übungsaufgabe 3.2.4. Berechne $M \times N$ und $N \times M$ für $M = \{1, 2, 3, 5\}$ und $N = \{8, 13\}$.

Beachte, dass das kartesische Produkt *nicht* kommutativ und *nicht* assoziativ ist. Jedoch gelten die folgenden Distributivgesetze für das kartesische Produkt:

Satz 3.2.5. Es gelten die folgenden Gesetze:

- Distributivgesetze für $\cup$ und $\times$:
 $(M \cup N) \times O = (M \times O) \cup (N \times O)$ und $M \times (N \cup O) = (M \times N) \cup (M \times O)$
- Distributivgesetze für $\cap$ und $\times$:
 $(M \cap N) \times O = (M \times O) \cap (N \times O)$ und $M \times (N \cap O) = (M \times N) \cap (M \times O)$
- Distributivgesetze für $\setminus$ und $\times$:
 $(M \setminus N) \times O = (M \times O) \setminus (N \times O)$ und $M \times (N \setminus O) = (M \times N) \setminus (M \times O)$

Eine Veranschaulichung des ersten Distributivgesetzes ist die folgende Illustration wobei Paare des Kreuzprodukts durch Pfeile verbunden sind:



3.2.1 Allgemeines kartesisches Produkt

Das eingeführte Kreuzprodukt ist *zweistellig*. Man kann dies jedoch auf eine beliebige Zahl von Mengen verallgemeinern, sodass anstelle von Paaren n -Tupel erzeugt werden.

Definition 3.2.6 (Allgemeines kartesisches Produkt). Seien $M_1, \dots, M_n$ Mengen, dann ist das kartesische Produkt $M_1 \times \dots \times M_n$ definiert durch

$$M_1 \times \dots \times M_n = \{(x_1, \dots, x_n) \mid x_1 \in M_1, \dots, x_n \in M_n\}$$

Ist eine der Mengen M_i leer, so gilt $M_1 \times \dots \times M_n = \emptyset$.

Beispiel 3.2.7. Wir zeigen einige Beispiele:

- Für $M = \{0, 1\}$ ist

$$M \times M \times M = \{(0, 0, 0), (0, 0, 1), (0, 1, 0), (0, 1, 1), (1, 0, 0), (1, 0, 1), (1, 1, 0), (1, 1, 1)\}.$$

- Für $M = \{a, b\}, N = \{C, D\}, O = \{1, 2\}$ ist

$$\begin{aligned} M \times N \times O &= \{(a, C, 1), (a, C, 2), (a, D, 1), (a, D, 2), \\ &\quad (b, C, 1), (b, C, 2), (b, D, 1), (b, D, 2)\} \\ O \times M \times N &= \{(1, a, C), (1, a, D), (1, b, C), (1, b, D), \\ &\quad (2, a, C), (2, a, D), (2, b, C), (2, b, D)\} \end{aligned}$$

- Das Kreuzprodukt $\mathbb{R} \times \mathbb{R} \times \mathbb{R}$ beschreibt die Menge aller dreidimensionalen Punkte im Raum.
- Das Kreuzprodukt $[-\frac{1}{2}, \frac{1}{2}] \times [-\frac{1}{2}, \frac{1}{2}] \times [-\frac{1}{2}, \frac{1}{2}]$ beschreibt alle dreidimensionalen Punkte eines Würfels mit Seitenlänge 1, dessen Mittelpunkt im Ursprung liegt.

Für das n -fache kartesische Produkt einer Menge M mit sich selbst schreiben wir auch M^n . Z.B. schreiben wir $\mathbb{R}^3$ für $\mathbb{R} \times \mathbb{R} \times \mathbb{R}$.

3.3 Mächtigkeiten und Zählformeln

Definition 3.3.1 (Mächtigkeit einer Menge). Sei M eine Menge. Dann bezeichnen wir mit $|M|$ die Anzahl der Elemente von M und nennen dies die Mächtigkeit (oder Kardinalität) von M . Wenn die Mächtigkeit eine natürliche Zahl oder 0 ist, so nennen wir M endlich, anderenfalls ist M unendlich und schreiben in diesem Fall $|M| = \infty$.

Beispiel 3.3.2. Einige Mächtigkeiten:

- $|\{-1, 0, 1, 2, 3, 4\}| = 6$
- $|\emptyset| = 0$
- $|\{\emptyset, 1, \{2, 3, 4, 5\}\}| = 3$ (denn die Menge hat die 3 Elemente $\emptyset$, 1 und $\{2, 3, 4, 5\}$).
- $|\mathbb{N}| = \infty, |\mathbb{R}| = \infty$
- $|\mathbb{N} \setminus \mathbb{Z}| = 0, |\mathbb{Z} \setminus \mathbb{N}| = \infty$

Übungsaufgabe 3.3.3. Bestimme die Mächtigkeit von $M = \{A, C, F, G, H\}$.

Wie ist die Mächtigkeit von $M = \{x \in \mathbb{N} \mid x \text{ ist Primzahl} \wedge 1 \leq x \leq 20\}$?

Gebe Mengen M_1 und M_2 an, sodass gilt: $|M_1| = 5, |M_2| = 6, |M_1 \cup M_2| = 7$.

Offensichtlich gilt:

Satz 3.3.4. Seien M, N endliche Mengen mit $M \subseteq N$. Dann gilt $|N \setminus M| = |N| - |M|$.

Satz 3.3.5 (Mächtigkeit des Komplements). *Sei U ein endliches Universum und M eine Teilmenge von U . Dann gilt $|\overline{M}| = |U \setminus M| = |U| - |M|$.*

Z.B. kann man die Anzahl aller Volljährigen berechnen, wenn man die Größe der Weltbevölkerung und die Anzahl der Kinder und Jugendlichen kennt $|\text{Volljährige}| = |\text{Kinder und Jugendliche}| = |\text{Menschen} \setminus \text{Kinder und Jugendliche}| = |\text{Menschen}| - |\text{Kinder und Jugendliche}|$

Satz 3.3.6 (Summenformel). *Seien M und N endliche Mengen. Dann gilt:*

$$|M \cup N| = |M| + |N| - |M \cap N|$$

Wenn M und N disjunkt sind, gilt sogar $|M \cup N| = |M| + |N|$.

Beweis. Die wesentliche Idee ist, dass man sowohl in $|M|$ als auch in $|N|$ die gemeinsamen Elemente von $|M|$ und $|N|$ zählt und diese daher einmal abgezogen werden müssen. Formaler setzt sich $M \cup N$ aus drei disjunkten Mengen zusammen: $M \setminus N$, $N \setminus M$ und $M \cap N$. Also gilt $|M \cup N| = |M \setminus N| + |N \setminus M| + |M \cap N| = |M \setminus (M \cap N)| + |N \setminus (M \cap N)| + |M \cap N| = |M| - |M \cap N| + |N| - |M \cap N| + |M \cap N| = |M| + |N| - |M \cap N|$. $\square$

Beispiel 3.3.7. Für $M = \{1, 2, 3\}$ und $N = \{2, 3, 4\}$ gilt $|M \cup N| = |M| + |N| - |M \cap N| = 3 + 3 - 2 = 4$.

Übungsaufgabe 3.3.8. Für die Essensbestellung werden die Teilnehmenden eines Seminars nach ihren Essensvorlieben (Deutsch oder Italienisch) befragt.

20 Personen melden sich für Italienisch, 10 Personen melden sich für Deutsch, 5 Personen haben sich sowohl für Italienisch als auch für Deutsch gemeldet.

Wie viele Personen wurden befragt (alle haben geantwortet)?

Die Summenformel für drei endliche Mengen M , N und O lautet:

$$|M \cup N \cup O| = |M| + |N| + |O| - |M \cap N| - |M \cap O| - |N \cap O| + |M \cap N \cap O|$$

Die Begründung für die Korrektheit der Formel geht analog zum Beweis von Satz 3.3.6:

- Elemente, die nur in einer der Mengen vorkommen, werden in $|M| + |N| + |O|$ genau einmal gezählt.
- Elemente, die in genau zwei der Mengen vorkommen, werden in $|M| + |N| + |O|$ genau zweimal gezählt und in $|M \cap N| + |M \cap O| + |N \cap O|$ genau einmal gezählt.
- Elemente, die in allen drei Mengen vorkommen, werden in $|M| + |N| + |O|$ genau dreimal gezählt und in $|M \cap N| + |M \cap O| + |N \cap O|$ genau dreimal gezählt.

Von der Summe $|M| + |N| + |O|$ werden daher durch $-|M \cap N| - |M \cap O| - |N \cap O|$ die doppelt gezählten Elemente einmal abgezogen, aber die dreifach gezählten Elemente dreimal abgezogen. Daher müssen diese durch $+|M \cap N \cap O|$ noch einmal hinzugefügt werden.

Verallgemeinert man die Summenformel auf die Vereinigung beliebig vieler endlicher Mengen, dann erhält man die Siebformel:

Satz 3.3.9 (Siebformel). *Seien $M_1, \dots, M_n$ endliche Mengen. Dann gilt:*

$$|M_1 \cup \dots \cup M_n| = \alpha_1 - \alpha_2 + \alpha_3 - \alpha_4 \dots$$

wobei α_i die Summe aller Mächtigkeiten aller Schnitte von i Mengen ist, d.h. α_i berechnet sich durch die folgenden Schritte:

- Für je i Mengen der Mengen $M_1, \dots, M_n$ bilde deren Schnitt.
- Bestimme die Mächtigkeiten der Schnittmengen.
- Summiere die Mächtigkeiten der Schnittmengen.

Übungsaufgabe 3.3.10. *Wie lautet die Siebformel für vier Mengen?*

Satz 3.3.11 (Produktformel). *Seien M und N endliche Mengen, dann gilt $|M \times N| = |M| \cdot |N|$. Für k endliche Mengen $M_1, \dots, M_k$ gilt $|M_1 \times \dots \times M_k| = |M_1| \cdot \dots \cdot |M_k|$.*

Beweis. Sei $|M| = m$ und $|N| = n$. Für jedes Paar $(x, y) \in M \times N$ hat man m Möglichkeiten für x . Ist x festgelegt, so hat man n weitere Möglichkeiten um y festzulegen. Also gibt es $m \cdot n$ Möglichkeiten.

Analog kann man für das allgemeine kartesische Produkt argumentieren¹. □

Beispiel 3.3.12. *Für vierstellige PINs am Geldautomat gibt es $|\{0, \dots, 9\}^4| = 10^4 = 10000$ Möglichkeiten.*

Ein Schachbrett hat $|\{A, \dots, H\} \times \{1, \dots, 8\}| = |\{A, \dots, H\}| \cdot |\{1, \dots, 8\}| = 8 \cdot 8 = 64$ Felder.

Übungsaufgabe 3.3.13. *Wie viele Karten hat ein Skatblatt, welches durch $\{\spadesuit, \clubsuit, \diamondsuit, \heartsuit\} \times \{7, 8, 9, 10, B, D, K, A\}$ repräsentiert wird?*

Binäre Tupel sind Tupel, die nur aus 0en und 1en bestehen. Genauer ist ein binäres n -Tupel ein Tupel der Form $(b_1, \dots, b_n)$ mit $b_i \in \{0, 1\}$.

Beispiel 3.3.14. *Die Menge aller binären 3-Tupel ist $\{(0, 0, 0), (0, 0, 1), (0, 1, 0), (0, 1, 1), (1, 0, 0), (1, 0, 1), (1, 1, 0), (1, 1, 1)\}$. Das sind 8 Stück. Wie sieht es allgemein aus?*

Satz 3.3.15. *Die Anzahl der binären n -Tupel ist 2^n .*

¹Es fehlen uns hier zu diesem Zeitpunkt noch die Mittel über alle k zu argumentieren. Das holen wir später nach.

Beweis. Verwende die Produktformel für das n -fache Kreuzprodukt von $\{0, 1\}$: $|\{0, 1\}^n| = |\{0, 1\}|^n = 2^n$ $\square$

Definition 3.3.16 (Potenzmenge). Sei M eine Menge. Die Menge aller Teilmengen von M nennt man Potenzmenge von M . Wir schreiben $\mathcal{P}(M)$ für die Potenzmenge.

$$\mathcal{P}(M) := \{M' \mid M' \subseteq M\}$$

Beispiel 3.3.17. $\mathcal{P}(\{1, 2\}) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$

Satz 3.3.18. Jede n -elementige Menge hat genau 2^n Teilmengen, d.h. für endliche Mengen M gilt $|\mathcal{P}(M)| = 2^{|M|}$.

Beweis. Sei $M = \{x_1, \dots, x_n\}$ eine beliebige Nummerierung der n Elemente von M . Für jede Teilmenge N von M erzeuge ein binäres n -Tupel $B_N = (b_{N,1}, \dots, b_{N,n})$ wobei $b_{N,i} = 0$, wenn $x_i \notin N$ und $b_{N,i} = 1$ wenn $x_i \in N$. Damit lässt sich leicht einsehen, dass jede Teilmenge genau ein binäres n -Tupel repräsentiert und umgekehrt jedes binäre n -Tupel genau eine Teilmenge repräsentiert (wenn zwei Tupel verschieden sind, sind auch die Teilmengen verschieden und umgekehrt.). Kurzum, es gibt genauso viele Teilmengen, wie es binäre n -Tupel gibt. Aus Satz 3.3.15 wissen wir bereits, dass dies 2^n viele sind. $\square$

Übungsaufgabe 3.3.19. Sei $M = \{a, b, c\}$, wobei die Elemente in der Reihenfolge a, b, c nummeriert sind. Welche binären 3-Tupel stellen die Teilmengen $\emptyset$, $\{b, c\}$, $\{a, c\}$ und $\{a, b, c\}$ entsprechend dem letzten Beweis jeweils dar?

Definition 3.3.20 (Binomialkoeffizienten). Die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge wird mit $\binom{n}{k}$ (gesprochen „ n über k “) bezeichnet. Diese Zahlen heißen Binomialkoeffizienten.

Beispiel 3.3.21. Für $M = \{1, 2, 3, 4\}$ gibt es $\binom{4}{2} = 6$ zweielementige Teilmengen (die Teilmengen sind $\{1, 2\}, \{1, 3\}, \{1, 4\}, \{2, 3\}, \{2, 4\}$ und $\{3, 4\}$).

$\binom{n}{0} = 1$ für jedes n , da es genau eine nullelementige Teilmenge gibt – die leere Menge.

$\binom{n}{n} = 1$ für jedes n , da es genau eine n -elementige Teilmenge gibt – die gesamte Menge.

$\binom{0}{0} = 1$ da die leere Menge sich selbst als Teilmenge hat.

$\binom{0}{n} = 0$ für jedes $n > 0$, da die leere Menge nur sich selbst als Teilmenge hat.

$\binom{n}{1} = n$ für jedes n , da es für jedes Element eine einelementige Menge gibt, und diese eine Teilmenge ist (gilt auch für $n = 0!$).

Man kann Binomialkoeffizienten auf verschiedene Weise ausrechnen. Eine Möglichkeit ist die Rekursionsformel zu verwenden:

Satz 3.3.22 (Rekursionsformel für Binomialkoeffizienten). Sei $1 \leq k \leq n$ mit $k, n \in \mathbb{N}$. Dann gilt

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

Beweis. Sei M eine n -elementige Menge und $m \in M$. Sei $M_k \subseteq \mathcal{P}(M)$ die Menge aller k -elementigen Teilmengen von M (d.h. $\binom{n}{k} = |M_k|$). Wir teilen die Mengen in M_k in zwei disjunkte Mengen auf:

- Mengen, die m nicht enthalten. Die Anzahl solcher Teilmengen ist gleich zur Anzahl der k -elementigen Teilmengen von $M \setminus \{m\}$, welche durch den Binomialkoeffizienten $\binom{n-1}{k}$ ausgedrückt werden kann.
- Mengen, die m enthalten. Die Anzahl solcher Teilmengen ist gleich zur Anzahl der $k-1$ -elementigen Teilmengen von $M \setminus \{m\}$, denn hat man diese Mengen, so muss man lediglich m als Element zu jeder Menge hinzufügen. Diese ändert die Anzahl nicht. Diese Anzahl ist der Binomialkoeffizient $\binom{n-1}{k-1}$

Da beide Mengen disjunkt sind, ist die Summe der Elemente beider Mengen gleich zur Anzahl aller k -elementigen Teilmengen. $\square$

Eine weitere Berechnungsmethode für Binomialkoeffizienten ist die folgende explizite Formel, welche die *Fakultät* von natürlichen Zahlen verwendet. Für eine natürliche Zahl n ist die Fakultät von n , geschrieben $n!$, das Produkt $n! := n \cdot (n-1) \cdot \dots \cdot 2 \cdot 1$. Auch für die 0 definieren wir die Fakultät als $0! := 1$.

Beispiel 3.3.23. Z.B. ist $4! = 4 \cdot 3 \cdot 2 \cdot 1 = 24$.

Satz 3.3.24 (Explizite Formel für Binomialkoeffizienten). Seien $k, n \in \mathbb{N}_0$ mit $0 \leq k \leq n$. Dann gilt:

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

Beispiel 3.3.25. Bei Lotto „6 aus 49“ werden 6 Zahlen aus 49 gegebenen Zahlen gezogen. Dafür gibt es $\binom{49}{6} = \frac{49!}{6! \cdot 43!} = \frac{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44}{6!} = 13.983.816$ Möglichkeiten. Die Wahrscheinlichkeit für 6 Richtige liegt daher bei $1/13.983.816 = 0,000000071 \dots$. Die Wahrscheinlichkeit auf 6 Richtige mit Superzahl ist noch mal ein Zehntel davon (da die Superzahl eine Zahl aus 0 bis 9 ist).

3.4 Schlussbemerkungen

Genau wie zur Logik, gibt es zur Mengenlehre und zur Definition von Mengen sehr viel Literatur. Eine knappe Einführung findet sich in (TT13), die Siebformel findet man z.B. in (BZ14). Eine ausführliches Kapitel zu Mengen findet man in (MM24). Ein ganzes Buch zur mathematischen Mengenlehre ist z.B. (Ebb21).

4 Relationen, Funktionen und Abzählbarkeit von Mengen

4.1 Relationen

Wir betrachten als Beispiel eine Menge von Personen $P = \{\text{Anna, Bernd, Claus}\}$ und eine Menge von Tieren $T = \{\text{Hund, Katze, Maus}\}$. Wie wir bereits wissen, enthält das kartesische Produkt $P \times T$ alle Paare bestehend aus einer Person und einem Tier

$$P \times T = \{(\text{Anna, Hund}), (\text{Bernd, Hund}), (\text{Claus, Hund}), \\ (\text{Anna, Katze}), (\text{Bernd, Katze}), (\text{Claus, Katze}), \\ (\text{Anna, Maus}), (\text{Bernd, Maus}), (\text{Claus, Maus})\}$$

Eine Teilmenge von $P \times T$, d.h. eine Menge, die im Allgemeinen nur eine Auswahl dieser Paare enthält, nennt man eine (binäre) *Relation*. Man drückt damit aus, ob Objekte der einen Menge *in Relation* zu den Objekten der anderen Menge stehen. Nehme an, dass Anna Katzen aber keine Hunde und Mäuse mag, Bernd sowohl Katzen als auch Hunde aber keine Mäuse mag, und Claus nur Mäuse mag. Dann kann man die Relation „mag“ als Relation R ausdrücken mit

$$R = \{(\text{Anna, Katze}), (\text{Bernd, Hund}), (\text{Bernd, Katze}), (\text{Claus, Maus})\}.$$

Definition 4.1.1 (Relation). *Eine Relation zwischen zwei Mengen M und N ist eine Teilmenge des kartesischen Produkts $M \times N$.*

Beispiel 4.1.2. Für $T = \{\text{Hund, Katze, Maus}\}$ und $N = \{\text{Fleisch, Getreide, Milch}\}$ ist das kartesische Produkt

$$T \times N = \{(\text{Hund, Fleisch}), (\text{Hund, Getreide}), (\text{Hund, Milch}), \\ (\text{Katze, Fleisch}), (\text{Katze, Getreide}), (\text{Katze, Milch}), \\ (\text{Maus, Fleisch}), (\text{Maus, Getreide}), (\text{Maus, Milch})\}$$

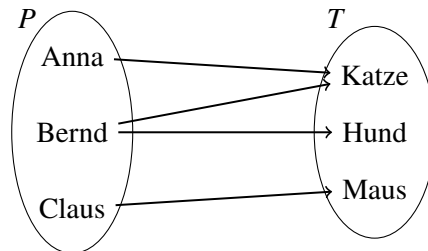
Eine Relation ist z.B.

$$R = \{(x, y) \in (T \times E) \mid x \text{ ernährt sich am liebsten von } y\} \\ = \{(\text{Hund, Fleisch}), (\text{Katze, Milch}), (\text{Maus, Getreide})\}$$

Beispiel 4.1.3. Die „Kleiner-Beziehung“ auf reellen Zahlen ist eine Relation:

$$R := \{(x, y) \in \mathbb{R} \times \mathbb{R} \mid x < y\}$$

Zur Illustration kann man Relationen im Mengendiagramm durch Pfeile darstellen, wobei für jedes Paar (x, y) ein Pfeil $x \rightarrow y$ gezeichnet wird. Für die vorher definierte Relation „mag“ kann man zeichnen:



Übungsaufgabe 4.1.4. Sei $S = \{\text{Istanbul, New York, Tokyo, Wiesbaden}\}$ und $K = \{\text{Europa, Asien, Nordamerika}\}$. Zähle die Elemente der folgenden Relation R auf

$$R = \{(x, y) \in (S \times K) \mid \text{Stadt } x \text{ liegt in Kontinent } y\}.$$

Definition 4.1.5. Sei R eine Relation zwischen den Mengen M und N , d.h. $R \subseteq M \times N$. Wenn $(x, y) \in R$, so sagt man „ x steht in Relation R mit y “ und schreibt dafür auch $x R y$. Für $(x, y) \notin R$ schreibt man auch $x \neg R y$.

Wenn wir eine konkrete Relation R betrachten, bezeichnen wir diese oft mit $\sim_R$ und schreiben $x \sim_R y$ für $x R y$ und $x \not\sim_R y$ für $x \neg R y$. Wenn R aus dem Kontext klar ist, lassen wir den Index R auch manchmal weg und schreiben nur $x \sim y$ bzw. $x \not\sim y$.

Beispiel 4.1.6. Beispiele für Relationen aus dem Alltag:

- $x \sim y$: „ x ist befreundet mit y “
- $x \sim y$: „ x ist verwandt mit y “
- $x \sim y$: „ x und y sind Nachbarn“

Beispiel 4.1.7. Sei $R = \{(x, y) \in (\mathbb{Z} \times \mathbb{Z}) \mid x - y \text{ ist ungerade}\}$. Dann gilt zum Beispiel

- $(3, 2) \in R$, auch geschrieben als $3 R 2$, $3 \sim_R 2$ oder $3 \sim 2$.
- $(-10, 5) \in R$, auch geschrieben als $-10 R 5$, $-10 \sim_R 5$ oder $-10 \sim 5$.
- $(20, 10) \notin R$, auch geschrieben als $20 \neg R 10$, $20 \not\sim_R 10$ oder $20 \not\sim 10$.

Definition 4.1.8 (Umkehrrelation). Sei $R \subseteq M \times N$ eine Relation zwischen M und N . Dann heißt

$$R^{-1} = \{(y, x) \in N \times M \mid (x, y) \in R\}$$

die Umkehrrelation von R (manchmal auch inverse Relation zu R).

Beispiel 4.1.9. Für $R = \{(\text{Hund, Fleisch}), (\text{Katze, Milch}), (\text{Maus, Getreide})\}$ ist die Umkehrrelation $R^{-1} = \{(\text{Fleisch, Hund}), (\text{Getreide, Maus}), (\text{Milch, Katze})\}$

Die Relation „ist Nachfahre von“ ist die Umkehrrelation der Relation „ist Vorfahre von“.

Definition 4.1.10 (Komposition von Relationen). Seien $R \subseteq M \times N$ und $S \subseteq N \times O$. Dann ist die Komposition $R \circ S$ definiert als:

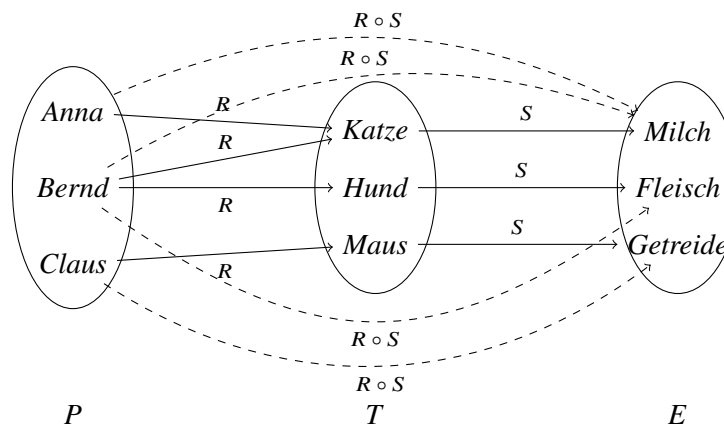
$$R \circ S := \{(x, z) \in (M \times O) \mid \exists y \in N : x R y \wedge y S z\}$$

Graphisch kann man sich die Komposition als das Verschmelzen von zwei Pfeilen zu einem neuen Pfeil vorstellen, wobei der erste Pfeil zu R und der zweite Pfeil zu S gehört.

Beispiel 4.1.11. Für $P = \{\text{Anna, Bernd, Claus}\}$, $T = \{\text{Hund, Katze, Maus}\}$ und $E = \{\text{Fleisch, Getreide, Milch}\}$ und die beiden Relationen $R \subseteq (P \times T)$ und $S \subseteq (T \times E)$ mit $R = \{(\text{Anna, Katze}), (\text{Bernd, Hund}), (\text{Bernd, Katze}), (\text{Claus, Maus})\}$ und $S = \{(\text{Hund, Fleisch}), (\text{Katze, Milch}), (\text{Maus, Getreide})\}$ ist die Komposition $R \circ S \subseteq (P \times E)$:

$$R \circ S = \{(\text{Anna, Milch}), (\text{Bernd, Milch}), (\text{Bernd, Fleisch}), (\text{Claus, Getreide})\}$$

Wir veranschaulichen die Komposition $R \circ S$ graphisch:



Übungsaufgabe 4.1.12. Für die Mengen

$$T = \{\text{Hund, Katze, Maus}\}, E = \{\text{Fleisch, Getreide, Milch}\} \text{ und } L = \{\text{Bauernhof, Molkerei, Supermarkt, Fleischerei}\}$$

und die Relationen $S \subseteq (T \times E)$ und $Q \subseteq (E \times L)$ mit

$$S = \{(\text{Hund, Fleisch}), (\text{Katze, Milch})\}, \\ Q = \{(\text{Fleisch, Fleischerei}), (\text{Fleisch, Supermarkt}), (\text{Getreide, Bauernhof}), \\ (\text{Getreide, Supermarkt}), (\text{Milch, Molkerei}), (\text{Milch, Supermarkt})\}$$

berechne $S \circ Q$. Welche Bedeutung könnten Q und $S \circ Q$ haben?

4.1.1 Mehrstellige Relationen

Genau wie beim Kreuzprodukt kann man statt zwei auch beliebig viele Mengen verwenden, um eine Relation darüber zu bilden.

Definition 4.1.13 (*n*-stellige Relation). Eine *n*-stellige Relation *R* ist eine Teilmenge des allgemeinen kartesischen Produkts von *n* Mengen $M_1, \dots, M_n$:

$$R \subseteq M_1 \times \dots \times M_n$$

Elemente einer *n*-stelligen Relation sind *n*-Tupel.

Beispiel 4.1.14. Mit den Mengen

Veranstaltung = {EI, AM, DS, OOSE}

Raum = {B001, B002}

Tag = {Mo, Di, Mi, Do, Fr}

Beginn = {0815, 1000, 1145, 1415}

ist $R \subseteq (\text{Veranstaltung} \times \text{Raum} \times \text{Tag} \times \text{Beginn})$ mit

$$R = \{(AM, B002, Di, 1000), (EI, B001, Fr, 1000), (DS, B001, Mo, 0815), \\ (OOSE, B002, Mo, 1400), (OOSE, B002, Do, 1145)\}$$

eine 4-stellige Relation, welche einen Erstsemester-Vorlesungsplan darstellen könnte.

4.1.2 Anwendung: Datenbanksysteme

Relationen und die sogenannte *Relationale Algebra* (d.h. Relationen und Operationen auf diesen) bilden die Grundlage für viele Datenbanken und Datenbanksysteme. Diese stellen daher eine wichtige Anwendung von Relationen in der Informatik dar.

Die wesentliche Idee von *relationalen Datenbanken* besteht darin, dass Daten (wie z.B. obiger Erstsemester-Vorlesungsplan) als Relationen in Form von Tabellen gespeichert werden: Jedes *n*-Tupel der *n*-stelligen Relation ist eine Zeile in der Datenbanktabelle und repräsentiert einen Datensatz. Für unser Beispiel könnte eine solche Tabelle „Vorlesungsplan“ wie folgt aussehen:

Veranstaltung	Raum	Tag	Beginn
AM	B002	Di	1000
EI	B001	Fr	1000
DS	B001	Mo	0815
OOSE	B002	Mo	1400
OOSE	B002	Do	1145

Anfragesprachen wie z.B. *SQL* (Structured Query Language) manipulieren oder transformieren diese Tupel. Z.B. kann man sich von den Veranstaltungen, die montags stattfinden, die Kürzel und Uhrzeiten in SQL in etwa wie folgt berechnen lassen:

```
SELECT Veranstaltung, Beginn FROM Vorlesungsplan WHERE Tag='Mo'
```

Dies liefert die Tabelle

Veranstaltung	Beginn
DS	0815
OOSE	1400

welche eine zweistellige Relation $R' \subseteq (\text{Veranstaltung} \times \text{Beginn})$ repräsentiert, d.h. die obige SQL-Anfrage transformiert eine Relation in eine andere.

4.1.3 Äquivalenzrelationen und weitere Eigenschaften von Relationen

Wir betrachten nun binäre Relationen, die Teilmenge des Kreuzprodukts gleicher Mengen sind, d.h. Relationen $R \subseteq M \times M$. Diese nennt man auch *Relationen auf M* oder *homogene Relationen*. Eine sehr wichtige Rolle spielen die Äquivalenzrelationen:

Definition 4.1.15 (Äquivalenzrelation). *Sei $\sim$ eine Relation auf M . Dann ist $\sim$*

- reflexiv, falls für alle $x \in M$ gilt: $x \sim x$.
- symmetrisch, falls für alle $x, y \in M$ gilt: Wenn $x \sim y$ gilt, dann gilt auch $y \sim x$.
- transitiv, falls für alle $x, y, z \in M$ gilt: Wenn $x \sim y$ und $y \sim z$, dann gilt auch $x \sim z$.

Eine Relation, die reflexiv, symmetrisch und transitiv ist, nennt man Äquivalenzrelation.

Beispiel 4.1.16. *Wir betrachten einige Beispiele:*

- Die Relation „hat denselben Nachnamen wie“ ist reflexiv, symmetrisch und transitiv und daher eine Äquivalenzrelation.
- Die Relation „ist ein Vorfahre von“ ist nicht reflexiv (man ist kein Vorfahre von sich selbst), nicht symmetrisch (ein Elternteil ist Vorfahre des Kindes, aber nicht umgekehrt), aber transitiv (wenn x ein Vorfahre von y , und y ein Vorfahre von z , dann ist x auch ein Vorfahre von z). Damit ist „ist ein Vorfahre von“ keine Äquivalenzrelation.
- Die Relation „sind verheiratet“ ist nicht reflexiv, symmetrisch und nicht transitiv. Damit ist „sind verheiratet“ keine Äquivalenzrelation.
- Die Gleichheits-Relation $R \subseteq \mathbb{R}^2$ mit $R = \{(x, y) \in \mathbb{R}^2 \mid x = y\}$ ist reflexiv (denn $x = x$ für alle $x \in \mathbb{R}$), symmetrisch (denn aus $x = y$ folgt stets $y = x$) und transitiv (wenn $x = y$ und $y = z$, dann gilt auch $x = z$) und damit eine Äquivalenzrelation.
- Die Kleiner-Relation $R \subseteq \mathbb{R}^2$ mit $R = \{(x, y) \in \mathbb{R}^2 \mid x < y\}$ ist nicht reflexiv (denn $1 \not< 1$), nicht symmetrisch ($1 < 2$, aber $2 \not< 1$) und transitiv (wenn $x < y$ und $y < z$, dann gilt auch $x < z$). Damit ist R keine Äquivalenzrelation.

Übungsaufgabe 4.1.17. Sind die folgenden Relationen reflexiv, symmetrisch, transitiv? Sind sie Äquivalenzrelationen?

- die Relation „besuchen dieselbe Schule“ auf der Menge aller Schüler:innen
- die Relation „haben ein gemeinsames Hobby“ auf der Menge aller Personen
- die Relation „ist Schwester von“ auf der Menge aller Personen
- die Relation $\{(x, y) \in \mathbb{R}^2 \mid x \geq y\}$
- die Relation $\{(x, y) \in \mathbb{N}^2 \mid x \text{ ist Teiler von } y\}$

Satz 4.1.18. Eine Relation R ist genau dann symmetrisch, wenn $R^{-1} = R$ gilt.

Beweis. Wir zeigen beide Richtungen.

- Sei R symmetrisch. Sei $(a, b) \in R$ beliebig. Aufgrund der Symmetrie gilt $(b, a) \in R$. Damit gilt $(a, b) \in R^{-1}$. Da (a, b) beliebig gewählt wurde, gilt dies für alle (a, b) und damit folgt $R \subseteq R^{-1}$. Nun sei $(a, b) \in R^{-1}$ beliebig. Dann gilt $(b, a) \in R$ und aufgrund der Symmetrie auch $(a, b) \in R$. Damit folgt $R^{-1} \subseteq R$ und insgesamt $R = R^{-1}$.
- Sei $R^{-1} = R$ und $(a, b) \in R$. Dann gilt $(a, b) \in R^{-1}$ (denn $R^{-1} = R$) und damit auch $(b, a) \in R$ (da $(a, b) \in R^{-1}$). Da dies für jedes $(a, b) \in R$ gilt, folgt die Symmetrie von R . $\square$

Äquivalenzrelationen sind wichtig, da sie sich wie Gleichheiten verhalten, aber nicht fordern, dass die Elemente syntaktisch gleich sind. Z.B. ist die logische Äquivalenz $\equiv$ auf aussagenlogischen Formeln eine Äquivalenzrelation. Weitere Beispiele für Äquivalenzrelationen aus der Mathematik und der Informatik sind:

- Parallelität von Geraden (Geraden sind parallel, wenn sie in der selben Ebene liegen, sich aber nie schneiden)
- Kongruenz von Dreiecken (Dreiecke sind kongruent, wenn ihre Winkel und Seitenlängen gleich sind).
- Äquivalenz von Programmen: Hier gibt es verschiedene Möglichkeiten diese zu definieren, z.B. syntaktische Gleichheit, Gleichheit bis auf Umbenennung von Variablen, Semantische Gleichheit z.B. gleiches Ein- / Ausgabeverhalten, u.s.w.

Wir definieren die Gleichheit von Brüchen $\sim$: Sei B die Menge der Brüche, d.h. $\frac{a}{b}$ mit $a \in \mathbb{Z}, b \in \mathbb{Z} \setminus \{0\}$. Dann sei $\sim \subseteq B^2$ definiert als $\sim := \{(\frac{a}{b}, \frac{c}{d}) \mid a \cdot d = b \cdot c\}$.

Satz 4.1.19. Die Gleichheit von Brüchen $\sim$ ist eine Äquivalenzrelation.

Beweis. Wir prüfen alle drei Eigenschaften:

- *Reflexivität:* Für jeden Bruch $\frac{a}{b}$ gilt $\frac{a}{b} \sim \frac{a}{b}$, da $a \cdot b = b \cdot a$.
- *Symmetrie:* Seien $\frac{a}{b}, \frac{c}{d}$ Brüche mit $\frac{a}{b} \sim \frac{c}{d}$. Dann gilt $a \cdot d = b \cdot c$ und damit auch $c \cdot b = d \cdot a$, was zeigt $\frac{c}{d} \sim \frac{a}{b}$.

- **Transitivität:** Seien $\frac{a}{b}, \frac{c}{d}, \frac{e}{f}$ Brüche mit $\frac{a}{b} \sim \frac{c}{d}$ und $\frac{c}{d} \sim \frac{e}{f}$. Dann gilt sowohl $a \cdot d = b \cdot c$ als auch $c \cdot f = d \cdot e$. Umformen der ersten Gleichung ergibt $a = \frac{b \cdot c}{d}$. Multiplizieren mit f ergibt $a \cdot f = \frac{b \cdot c \cdot f}{d}$. Einsetzen von $c \cdot f = d \cdot e$ ergibt $a \cdot f = \frac{b \cdot c \cdot f}{d} = \frac{b \cdot d \cdot e}{d} = b \cdot e$. Damit folgt $\frac{a}{b} \sim \frac{e}{f}$. $\square$

Übungsaufgabe 4.1.20. Welche der folgenden Relationen sind Äquivalenzrelationen?

- $R = \{(x, y) \in \mathbb{N}^2 \mid x \text{ und } y \text{ sind ungerade}\}$
- $R = \{(x, y) \in \mathbb{N}^2 \mid x - y \text{ ist gerade}\}$
- $R = \{(x, y) \in \mathbb{N}^2 \mid x + y \text{ ist gerade}\}$
- $R = \{(x, y) \in \mathbb{N}^2 \mid x - y \text{ ist ungerade}\}$
- $R = \{(x, y) \in \mathbb{N}^2 \mid x + y \text{ ist ungerade}\}$

Definition 4.1.21 (Kongruenz modulo n). Sei $n > 1$ eine natürliche Zahl. Die Relation $\equiv_n$ auf $\mathbb{Z}$ heißt Kongruenz modulo n und ist definiert durch:

$$\begin{aligned}\equiv_n &= \{(x, y) \in \mathbb{Z} \mid y - x \text{ ist ein ganzzahliges Vielfaches von } n\} \\ &= \{(x, y) \in \mathbb{Z} \mid \exists k \in \mathbb{Z} : y - x = k \cdot n\}\end{aligned}$$

Wenn $x \equiv_n y$, dann sagen wir x und y sind kongruent modulo n .

Eine alternative Schreibweise für $x \equiv_n y$ ist $x \equiv y \pmod{n}$.

Beispiel 4.1.22. Zahlen sind kongruent modulo 3, wenn ihre Differenz ein Vielfaches von 3 ist.

- $6 \equiv_3 9$, denn $9 - 6 = 3$
- $5 \equiv_3 11$, denn $11 - 5 = 6$ ist ein Vielfaches von 3
- $42 \equiv_3 30$, denn $30 - 42 = -12$ ist ein Vielfaches von 3

Satz 4.1.23. Die Relation $\equiv_n$ ist eine Äquivalenzrelation.

Beweis. Wir prüfen alle drei Eigenschaften einer Äquivalenzrelation:

- **Reflexivität:** Sei $x \in \mathbb{Z}$. Dann gilt $x \equiv_n x$, denn $x - x = 0 = 0 \cdot n$.
- **Symmetrie:** Sei $x, y \in \mathbb{Z}$ mit $x \equiv_n y$. Dann gibt es ein $k \in \mathbb{Z}$, sodass $y - x = k \cdot n$. Da $x - y = -k \cdot n$ und $-k \in \mathbb{Z}$, bezeugt $-k$, dass auch $y \equiv_n x$ gilt.
- **Transitivität:** Sei $x \equiv_n y$ und $y \equiv_n z$. Dann gibt es $k_1, k_2 \in \mathbb{Z}$ mit $y - x = k_1 \cdot n$ und $z - y = k_2 \cdot n$. Die erste Gleichung kann umgeformt werden zu $y = k_1 \cdot n + x$. Einsetzen der Gleichung in die zweite Gleichung ergibt $z - k_1 \cdot n - x = k_2 \cdot n$ was zu $z - x = (k_1 + k_2)n$ umgeformt werden kann. Die Zahl $k_1 + k_2 \in \mathbb{Z}$ bezeugt daher, dass $x \equiv_n z$ gilt. $\square$

Definition 4.1.24. Bei einer ganzzahligen Division mit Rest von a durch n bezeichnet man den kleinsten nicht-negativen Rest als $a \bmod n$ („ a modulo n “).

Beispiel 4.1.25. $6 \bmod 3 = 0$ (denn $2 \cdot 3 + 0 = 6$), $11 \bmod 3 = 2$ (denn $3 \cdot 3 + 2 = 11$), $-11 \bmod 3 = 1$ (denn $(-4) \cdot 3 + 1 = -11$)

Beachte, dass mod in $10 \text{ mod } 3 = 1$ als binärer Operator auf Ganzzahlen verwendet wird, aber in $10 \equiv 1 \text{ mod } 3$ kein Operator ist, sondern ein syntaktisches Konstrukt bestehend aus $\equiv$ und mod . Wenn $a_1 \equiv_n a_2 \equiv_n \dots \equiv_n a_m$ gilt, dann schreibt man manchmal auch $a_1 \equiv a_2 \equiv \dots \equiv a_m \text{ mod } n$.

Offensichtlich gilt, dass für zwei Zahlen $a \equiv_n b$ (alternativ geschrieben $a \equiv b \text{ mod } n$) genau dann gilt, wenn $a \text{ mod } n = b \text{ mod } n$ gilt.

Übungsaufgabe 4.1.26. Berechne $9 \text{ mod } 5$, $27 \text{ mod } 8$, $1 \text{ mod } 2024$ und $-4 \text{ mod } 6$.

Definition 4.1.27 (Äquivalenzklasse). Sei $\sim$ eine Äquivalenzrelation auf M und $x \in M$. Dann ist

$$[x]_{\sim} := \{y \in M \mid y \sim x\}$$

die Äquivalenzklasse von x .

Man nennt x einen Repräsentanten der Äquivalenzklasse $[x]_{\sim}$.

Durch die Äquivalenzklassen fasst man äquivalente Elemente zu einer Menge zusammen. Die Äquivalenzklasse $[x]_{\sim}$ ist die Menge aller Elemente, die zu x in Relation stehen.

Beispiel 4.1.28. Für die Äquivalenzrelation $\sim =$ „besuchen dieselbe Schule“ und den Schüler Bernd ist $[\text{Bernd}]_{\sim}$ die Menge aller Schüler:innen, die in Bernds Schule gehen.

Für die Kongruenz $\equiv_3$ enthält $[0]_{\equiv_3}$ alle Zahlen, die ohne Rest durch 3 teilbar sind:

$$[0]_{\equiv_3} = \{y \in \mathbb{Z} \mid y \equiv_3 0\} = \{y \in \mathbb{Z} \mid \exists k \in \mathbb{Z} : 0 - y = 3 \cdot k\} = \{\dots, -6, -3, 0, 3, 6, \dots\}$$

Übungsaufgabe 4.1.29. Beschreibe die Elemente der Äquivalenzklasse $[4]_{\equiv_5}$.

Satz 4.1.30. Äquivalente Elemente haben dieselbe Äquivalenzklasse: Sei $\sim$ eine Äquivalenzrelation auf M . Es gilt $x \sim y$ genau dann, wenn $[x]_{\sim} = [y]_{\sim}$.

Beweis. Wir zeigen beide Richtungen:

- Sei $x \sim y$. Wir müssen zeigen $[x]_{\sim} = [y]_{\sim}$. Sei $z \in [x]_{\sim}$, dann gilt $z \sim x$ und mit der Transitivität von $\sim$ und $x \sim y$ folgt $z \sim y$ und daher $z \in [y]_{\sim}$. Daher folgt $[x]_{\sim} \subseteq [y]_{\sim}$. Umgekehrt sei nun $z \in [y]_{\sim}$. Dann gilt $z \sim y$ und mit Symmetrie $y \sim x$. Aus $x \sim y$ und der Transitivität folgt $x \sim z$. Mit Symmetrie folgt $z \sim x$ und damit $z \in [x]_{\sim}$. Damit haben wir auch $[y]_{\sim} \subseteq [x]_{\sim}$ gezeigt.
- Sei $[x]_{\sim} = [y]_{\sim}$. Da $\sim$ reflexiv ist, gilt $x \in [x]_{\sim} = [y]_{\sim}$. Aus $x \in [y]_{\sim}$ folgt $y \sim x$ und aufgrund der Symmetrie auch $x \sim y$. $\square$

Satz 4.1.31. Äquivalenzklassen sind gleich oder disjunkt.

Beweis. Angenommen es gibt $z \in ([x]_{\sim} \cap [y]_{\sim})$. Dann gilt $x \sim z$ und $y \sim z$ und mit Satz 4.1.30 $[x]_{\sim} = [y]_{\sim}$. $\square$

Korollar 4.1.32. Sei $\sim$ eine Äquivalenzrelation auf M . Dann liegt jedes Element von M in genau einer Äquivalenzklasse.

Aus dem Korollar folgt, dass die Äquivalenzklassen die Menge M partitionieren. Eine Partition einer Menge M ist eine Menge von Teilmengen von M , die

- paarweise disjunkt sind und
- deren Vereinigung wieder M ergibt.

Die durch eine Äquivalenzrelation $\sim$ auf M gegebene Partition ist $\{[m]_{\sim} \mid m \in M\}$. Diese bezeichnet man auch als *Quotient von M nach $\sim$* und schreibt $M/\sim$ für sie. Die Mächtigkeit von $M/\sim$ bezeichnet man auch als *Index der Äquivalenzrelation $\sim$* .

Beispiel 4.1.33. Die Äquivalenzklassen von $\equiv_3$ bilden eine Partition von $\mathbb{Z}$. Die Partition ist $\{[0]_{\equiv_3}, [1]_{\equiv_3}, [2]_{\equiv_3}\}$. D.h. der Index von $\equiv_3$ ist 3.

Beispiel 4.1.34. Die Äquivalenzklassen von „besucht dieselbe Schule“ bilden eine Partition aller Schüler:innen. Jede Äquivalenzklasse enthält genau die Schüler:innen einer Schule. Der Index von „besucht dieselbe Schule“ entspricht der Anzahl an Schulen.

Beispiel 4.1.35. Die Äquivalenzklassen von $=$ auf $\mathbb{N}$ sind $[i]_>= = \{i\}$, d.h. jede Äquivalenzklasse enthält nur den Repräsentanten selbst. Die Partition dazu ist $\{[1]_>=, [2]_>=, \dots\} = \{[i]_>= \mid i \in \mathbb{N}\}$. Der Index von $=$ auf $\mathbb{N}$ ist daher unendlich.

Hat man eine homogene Relation, so kann man diese erweitern (durch Hinzufügen neuer Paare), sodass sie transitiv ist. Verwendet man so wenig neue Paare wie möglich, so spricht man vom transitiven Abschluss der Relation transitiven Hülle der Relation.

Definition 4.1.36. Transitiver und reflexiv-transitiver Abschluss einer Relation Sei $R \subseteq M \times M$ eine homogene Relation. Dann sei R^+ die kleinste Erweiterung von R die transitiv ist. Der reflexiv-transitive Abschluss von R wird mit R^* bezeichnet und ist die Vereinigung $R^* := R^+ \cup \{(a, a) \mid a \in M\}$.

Man kann den transitiven Abschluss einer endlichen Relation R wie folgt berechnen:

- Starte mit $S := R$.
- Solange es noch a, b, c gibt mit $(a, b) \in S$, $(b, c) \in S$ und $(a, c) \notin S$, setze $S := S \cup \{(a, c)\}$
- Gebe $R^+ = S$ aus.

Beispiel 4.1.37. Für $R \subseteq \mathbb{N} \times \mathbb{N}$ mit $R = \{(1, 2), (2, 3), (3, 4)\}$ ist $R^+ = \{(1, 2), (1, 3), (1, 4), (2, 3), (2, 4), (3, 4)\}$ und $R^* = \{(1, 1), (2, 2), (3, 3), (4, 4), (1, 2), (1, 3), (1, 4), (2, 3), (2, 4), (3, 4)\}$.

Für $R = \{(i, i+1) \mid i \in \mathbb{N}\}$ ist $R^+ = \{(i, j) \mid i, j \in \mathbb{N} \text{ und } i < j\}$ und $R^* = \{(i, j) \mid i, j \in \mathbb{N} \text{ und } i \leq j\}$.

Zum Abschluss definieren wir noch weitere Eigenschaften homogener Relationen:

Definition 4.1.38 (Quasiordnung, Halbordnung, Ordnung, Striktordnung). Sei $\sim$ eine Relation auf einer Menge M . Dann ist $\sim$

- irreflexiv, falls für alle $x \in M$ gilt: $x \not\sim x$.
- antisymmetrisch, falls für alle $x, y \in M$ gilt: Wenn $x \sim y$ und $y \sim x$, dann gilt $x = y$.
- asymmetrisch, falls für alle $x, y \in M$ gilt: Wenn $x \sim y$, dann $y \not\sim x$.
- total, falls für alle $x, y \in M$ gilt: $x \sim y$ oder $y \sim x$.

Eine Relation,

- die transitiv und reflexiv ist, nennt man Quasiordnung.
- die transitiv, reflexiv und antisymmetrisch ist, nennt man Halbordnung oder partielle Ordnung.
- die transitiv, reflexiv, antisymmetrisch und total ist, nennt man (totale) Ordnung.
- die transitiv, irreflexiv und asymmetrisch ist, nennt man Striktordnung.

Beispiel 4.1.39. Die Relation $\leq$ auf natürlichen Zahlen ist eine totale Ordnung. Die Relation $<$ auf natürlichen Zahlen ist eine Striktordnung.

Die Teilmengenbeziehung $\subseteq$ auf $\mathcal{P}(\mathbb{N})$ ist eine partielle Ordnung: Sie ist reflexiv, da $(M \subseteq M)$, antisymmetrisch (denn aus $M \subseteq N$ und $N \subseteq M$ folgt $M = N$) und transitiv. Sie ist nicht total, da z.B. weder $\{1, 2\} \subseteq \{3, 4\}$ noch $\{3, 4\} \subseteq \{1, 2\}$ gilt.

Sei $\leq_{\text{Obst}}$ die folgende Relation auf Äpfel und Birnen: $a \leq_{\text{Obst}} b$ wenn a und b beides Äpfel oder beides Birnen sind und das Gewicht von b nicht größer als das Gewicht von a ist. $\leq_{\text{Obst}}$ ist eine partielle Ordnung, aber keine totale Ordnung (da $\leq_{\text{Obst}}$ keine Äpfel mit Birnen vergleicht).

Sei $\preceq$ die folgende Relation auf allen Wörtern:

$w \preceq w'$, wenn w und w' mit demselben Buchstaben anfangen
und das Wort w' nicht kürzer ist als das Wort w .

Z.B. gilt „Affe“ $\preceq$ „Ananas“ und „Anna“ $\preceq$ „Affe“. Dann ist $\preceq$ transitiv und reflexiv, und daher eine Quasiordnung, aber $\preceq$ ist keine partielle Ordnung, da $\preceq$ nicht antisymmetrisch ist („Anna“ $\preceq$ „Affe“ und „Affe“ $\preceq$ „Anna“) und keine Äquivalenzrelation, da $\preceq$ auch nicht symmetrisch ist („Affe“ $\preceq$ „Ananas“ aber „Ananas“ $\not\preceq$ „Affe“).

4.2 Funktionen

Wir definieren weitere Eigenschaften von binären Relationen:

Definition 4.2.1 (linkstotal, rechtstotal, linkseindeutig, rechtseindeutig). Sei R eine Relation mit $R \subseteq M \times N$. Dann heißt R

- linkstotal, wenn es zu jedem $x \in M$ mindestens ein $y \in N$ mit $x R y$ gibt.
- rechtstotal, wenn es zu jedem $y \in N$ mindestens ein $x \in M$ mit $x R y$ gibt.

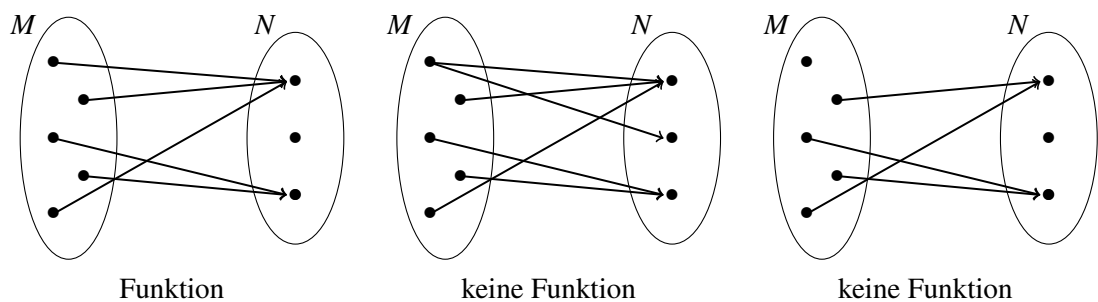
- linkseindeutig, wenn es zu jedem $y \in N$ höchstens ein $x \in M$ mit $x R y$ gibt.
- rechtseindeutig, wenn es zu jedem $x \in M$ höchstens ein $y \in N$ mit $x R y$ gibt.

Aufbauend auf diesen Begriffen können wir definieren, was wir unter einer Funktion verstehen. Diese sind spezielle Relationen: *Jedem* Element aus M muss *genau ein* Element aus N zugeordnet werden.

Sei $P = \{\text{Anna, Bert, Carl}\}$ und $A = \{21, 24, 27\}$. Dann besteht das kartesische Produkt aus allen möglichen 9 Paaren. Die Relation f ordne jeder Person eindeutig ihr Alter zu: $f = \{(\text{Anna}, 27), (\text{Bert}, 21), (\text{Carl}, 21)\}$. Da jeder Person aus P *genau ein* Alter aus A zugeordnet wird ist f eine Funktion.

Definition 4.2.2. Seien M und N Mengen. Eine Funktion f (manchmal auch Abbildung) von M nach N ist eine Relation $f \subseteq M \times N$, so dass es zu jedem $x \in M$ genau ein $y \in N$ gibt mit $(x, y) \in f$, d.h. f ist linkstotal und rechtseindeutig.

Für die folgenden Relationen ist nur die erste eine Funktion, denn in der zweiten werden einem Element aus M zwei Elemente aus N zugeordnet (nicht rechtseindeutig) und in der dritten wird einem Element aus M gar kein Element aus N zugeordnet (nicht linkstotal).



Relationen, die rechtseindeutig aber nicht linkstotal sind, werden auch als *partielle Funktionen* bezeichnet, da sie nicht überall definiert sind.

Definition 4.2.3. Für eine Funktion $f \subseteq M \times N$ schreiben wir auch $f : M \rightarrow N$. Da es für $x \in M$ genau ein Paar gibt mit $(x, y) \in f$, schreiben wir auch $y = f(x)$. Dabei ist M der Definitionsbereich von f , N der Wertebereich von f und $\{f(x) \mid x \in X\} \subseteq N$ ist der Bildbereich von f .

Neben Angabe der Menge von Paaren, gibt es weitere Möglichkeiten, Funktionen darzustellen: als Wertetabelle, textuell (z.B. „ f berechnet das Quadrat der Eingabe“), mit einer Funktionsgleichung (z.B. $f(x) = x^2$), als Funktionsgraph,

Definition 4.2.4 (Injektivität, Surjektivität, Bijektivität). Sei $f : M \rightarrow N$ eine Funktion.

- f heißt injektiv, wenn für alle $x_1, x_2 \in M$ gilt: Wenn $x_1 \neq x_2$, dann $f(x_1) \neq f(x_2)$. D.h. keine zwei verschiedenen Elemente aus M werden durch f auf dasselbe Element aus N abgebildet.
- f heißt surjektiv, wenn es für jedes $y \in N$ ein $x \in M$ gibt, mit $f(x) = y$. D.h. jedes Element aus N wird durch f „getroffen“. In diesem Fall sind Bildbereich und Wertebereich identisch.
- f heißt bijektiv (auch eineindeutig), wenn f injektiv und surjektiv ist.

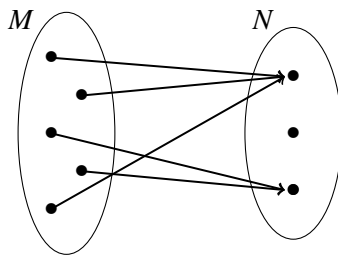
Beispiel 4.2.5. Sei M eine Menge von Mänteln und N eine Menge von Garderobenhaken. Sei $f : M \rightarrow N$ eine Funktion, die Mäntel auf Haken aufhängt.

Injektivität von f bedeutet, dass niemals zwei oder mehr Mäntel an ein und denselben Haken hängen.

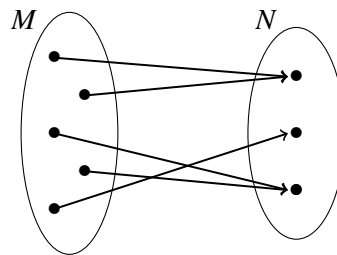
Surjektivität von f bedeutet, dass an jedem Haken mindestens ein Mantel hängt.

Bijektivität von f bedeutet, dass an jedem Haken genau ein Mantel hängt.

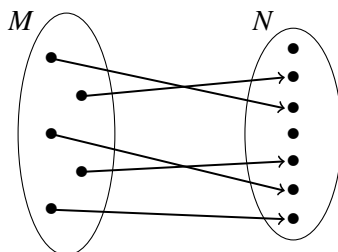
Zur Veranschaulichung zeigen wir Illustrationen für alle Möglichkeiten:



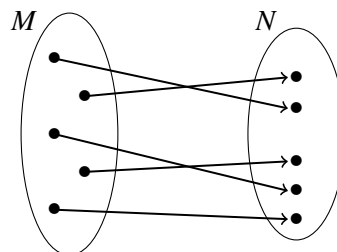
nicht injektiv, nicht surjektiv, nicht bijektiv



nicht injektiv, surjektiv, nicht bijektiv



injektiv, nicht surjektiv, nicht bijektiv



injektiv, surjektiv, bijektiv

Beispiel 4.2.6. Wir nennen einige Beispiele:

- Die Funktion, die jeder Person deren Geburtsdatum zuordnet, ist nicht injektiv, da mehrere Menschen am selben Tag Geburtstag haben. (Wir nehmen an, dass die Personen eindeutig gegeben sind, ansonsten (z.B. bei Verwendung des Namens) wäre die Relation noch nicht einmal eine Funktion).
- Die Funktion, die jedem Datum dessen Wochentag zuordnet, ist nicht injektiv, aber surjektiv.

- Die Funktion, die jeder Fahrgestellnummer eines Kraftfahrzeugs sein KFZ-Kennzeichen zuordnet, ist injektiv. Sie ist nicht surjektiv, wenn man alle möglichen Kennzeichen als Wertebereich betrachtet. Wenn man als Wertebereich nur die vergebenen Kennzeichen verwendet, ist sie auch surjektiv.
- Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist nicht injektiv (da z.B. $f(-2) = f(2) = 4$) und nicht surjektiv (z.B. gibt es kein x mit $f(x) = -1$).
- Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = 2x + 3$ ist bijektiv: Wenn $f(x_1) = 2x_1 + 3 = f(x_2) = 2x_2 + 3$, dann folgt $x_1 = x_2$ und zu jedem $y \in \mathbb{R}$ gibt es ein x mit $2x + 3 = y$, nämlich $x = \frac{1}{2} \cdot (y - 3)$.
- Die Funktion $f : \mathbb{N} \rightarrow \mathbb{Z}$ mit $f(x) = -x$ ist injektiv aber nicht surjektiv.
- Die Funktion $f : \mathbb{Z} \rightarrow \mathbb{Z}$ mit $f(x) = -x$ ist bijektiv.
- Die Betragsfunktion $f : \mathbb{Z} \rightarrow \mathbb{N}_0$ mit $f(x) = |x|$ ist nicht injektiv, aber surjektiv.

Übungsaufgabe 4.2.7. Prüfe für jede der folgenden Funktionen, ob diese injektiv, surjektiv oder bijektiv sind.

- Die Funktion, die jedem Tripel aus (Autor, Buchtitel, Verlag) die IBAN des Buches zuordnet.
- Die Funktion, die jedem Tripel aus (Autor, Buchtitel, Verlag) das Erscheinungsjahr zuordnet.
- Die Funktionen

$$\begin{aligned} f &: \{\spadesuit, \clubsuit, \diamondsuit, \heartsuit\} \rightarrow \{\square, \square, \square, \square\}, \\ g &: \{\spadesuit, \clubsuit, \diamondsuit, \heartsuit\} \rightarrow \{\square, \square, \square, \square\}, \\ h &: \{\spadesuit, \clubsuit, \diamondsuit, \heartsuit\} \rightarrow \{\square, \square, \square, \square\} \end{aligned}$$

gegeben durch die Wertetabelle:

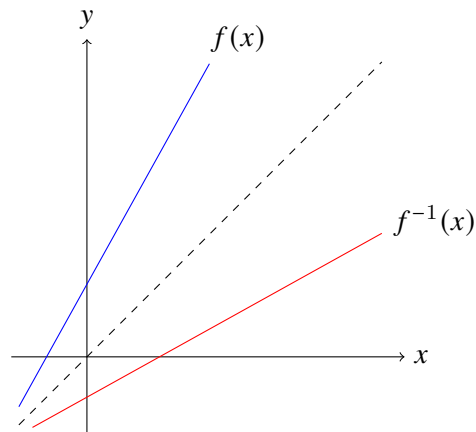
	$\spadesuit$	$\clubsuit$	$\diamondsuit$	$\heartsuit$
f	$\square$	$\square$	$\square$	$\square$
g	$\square$	$\square$	$\square$	$\square$
h	$\square$	$\square$	$\square$	$\square$

Definition 4.2.8 (Umkehrfunktion). Jede bijektive Funktion $f : M \rightarrow N$ ist umkehrbar. D.h. die Umkehrrelation f^{-1} ist wieder eine Funktion, die Umkehrfunktion. Bei der Umkehrfunktion sind Definitionsbereich und Wertebereich vertauscht, d.h. $f^{-1} : N \rightarrow M$.

Satz 4.2.9. Sei $f : M \rightarrow N$ eine Bijektion und $f^{-1} : N \rightarrow M$. Dann gilt $f^{-1}(f(x)) = x$ für alle x aus M und $f(f^{-1}(x)) = x$ für alle x aus N .

Beispiel 4.2.10. Zur Umrechnung von Grad-Celsius in Grad-Fahrenheit kann die Funktion $f(x) = \frac{9}{5}x + 32$ verwendet werden. Wie rechnet man Grad-Fahrenheit in Grad-Celsius um? Benutze die Umkehrfunktion von f !

Graphisch kann man die Umkehrfunktion durch eine Spiegelung an der Geraden $y = x$ erhalten:



Zum Berechnen der Funktionsgleichung für f^{-1} ersetze zunächst $f(x)$ durch y in der Funktionsgleichung von f . Das ergibt $y = \frac{9}{5}x + 32$. Löse danach die Gleichung nach x auf: Das ergibt $x = (y - 32) \frac{5}{9} = \frac{5}{9}y - \frac{160}{9}$. Schließlich vertausche x und y und ersetze y durch $f^{-1}(x)$. Das ergibt $f^{-1}(x) = \frac{5}{9}x - \frac{160}{9}$.

Übungsaufgabe 4.2.11. Zeichne die Funktion $f(x) = 2x + 1$. Zeichne und berechne die Umkehrfunktion von f .

4.3 Abzählbarkeit

In diesem Abschnitt geht es im Wesentlichen um die Frage: Wann haben zwei Mengen die gleiche Anzahl an Elementen?

Bei endlichen Mengen M und N ist diese Frage einfach: Zähle die Elemente und vergleiche, bzw. nutze die Mächtigkeit, d.h. $|M|$ und $|N|$ sind natürliche Zahlen (oder 0) und diese können mit $=$ und $<$ verglichen werden. Wenn eine der beiden Mengen unendlich viele Elemente hat, die andere aber nicht, dann ist auch klar, welche Menge mehr Elemente hat.

Wenn jedoch beide Mengen unendlich viele Elemente haben, so ist zunächst unklar, ob diese gleich viele Elemente haben (gilt „unendlich“ = „unendlich“?). Der Umgang mit der Unendlichkeit ist nicht so einfach, wie man denken mag.

Beispiel 4.3.1. Betrachte das folgende von David Hilbert vorgeschlagene Gedankenexperiment.

Wenn ein Hotel endliche viele Zimmer hat, dann kann man keine Gäste mehr aufnehmen, sobald alle Zimmer belegt sind (klar!). Betrachte nun ein Hotel mit unendlich vielen Zimmern, die mit den natürlichen Zahlen durchnummeriert sind (Zimmer i mit $i \in \mathbb{N}$). Wenn alle Zimmer belegt sind, gibt es eine Zuordnung Gast i zu Zimmer i (für $i \in \mathbb{N}$). Aber kann dieses Hotel keine Gäste mehr aufnehmen? Es kann einen Gast aufnehmen, indem alle vorherigen Gäste ein Zimmer weiter rutschen (Gast i ist in Zimmer $i + 1$). Dann ist Zimmer 1 wieder frei und der neue Gast kann dieses Zimmer beziehen.

Angenommen, es kommt nicht nur ein Gast, sondern ein Bus mit unendlich vielen (die mit natürlichen Zahlen durchnummeriert sind, z.B. Neugast i mit $i \in \mathbb{N}$). Kann das Hotel diese Gäste aufnehmen? Ja, wir können unendlich viele freie Zimmer schaffen, indem jeder alte Gast i in das Zimmer mit Nummer $2i$ zieht. Dann sind die Zimmer mit ungerader Nummer alle wieder frei, dort können die neuen Gäste einziehen (Neugast i wohnt in Zimmer $2i - 1$).

Selbst wenn unendlich viele Busse (durchnummeriert mit natürlichen Zahlen, Bus i mit $i \in \mathbb{N}$) mit jeweils unendlich vielen Gästen (durchnummeriert: Gast j aus Bus i mit $j \in \mathbb{N}$) kommen, kann man diese im Hotel unterbringen: Jeder alte Gast i wechselt auf Zimmer 2^i , und Gast j aus Bus i zieht in das Zimmer mit Nummer p_i^j ein, wobei p_i die $i+1$. Primzahl ist. Man muss sich dabei nur klar machen, dass diese alle auf verschiedene Zimmer verteilt werden: Je zwei Zimmernummern aus unterschiedlichen Bussen (oder Bus und Altgast), haben keine gemeinsamen Teiler und Zimmernummern aus demselben Bus sind aufgrund der unterschiedlichen Potenz verschieden.

Es gibt sogar noch unendlich viele freie Zimmer, z.B. die Zimmer mit den Nummern 6^i mit $i \in \mathbb{N}$.

Das Beispiel suggeriert, dass alle Begriffe von unendlich ein und dieselbe Größe haben. Aber das Beispiel ist wohl gewählt, denn es wurde stets gefordert, dass die Gäste, Busse, Zimmer u.s.w. mit den natürlichen Zahlen durchnummeriert werden können. Das entspricht dem Begriff der Abzählbarkeit, den wir gleich definieren. Zunächst definieren wir, was es nun bedeutet, dass zwei Mengen gleichmächtig sind.

Definition 4.3.2. Zwei Mengen M und N werden genau dann als gleichmächtig bezeichnet, wenn es eine Bijektion $f : M \rightarrow N$ gibt.

Die Definition wird für endliche als auch unendliche Mengen verwendet. Für endliche Mengen passt der Begriff auch gut: Statt die Elemente abzuzählen, ordnen wir diese zu. Das passt sehr gut zum Beispiel der Haken und Mäntel. Das Aufhängen ist die Zuordnung.

Bemerkung 4.3.3. Die Relation „sind gleichmächtig“ auf Mengen ist eine Äquivalenzrelation.

Definition 4.3.4. Eine Menge M heißt abzählbar, wenn M endlich oder gleichmächtig zu $\mathbb{N}$ ist (im zweiten Fall sagt man auch, dass M abzählbar unendlich ist).

Satz 4.3.5. Die Mengen $\mathbb{N}$ und $\mathbb{Z}$ sind gleichmächtig, d.h. $\mathbb{Z}$ ist abzählbar.

Beweis. Sei $f : \mathbb{Z} \rightarrow \mathbb{N}$ mit $f(x) = 2 \cdot x + 1$ für $x \geq 0$ und $f(x) = -2 \cdot x$ für $x < 0$. Dann bildet f die nicht-negativen Zahlen auf die ungeraden natürlichen Zahlen, und die negativen auf die geraden natürlichen Zahlen ab.

Die Funktion f ist injektiv: Sei $f(a) = f(b)$. Dann müssten a und b entweder beide negativ oder beide nicht-negativ sein, ansonsten wäre eine der Zahlen $f(a)$, $f(b)$ gerade und die andere Zahl ungerade. Wenn beide negativ sind, dann muss $-2a = -2b$ und damit $a = b$ gelten. Wenn beide nicht-negativ sind, dann muss $2a + 1 = 2b + 1$ und damit $a = b$ gelten.

Die Funktion f ist surjektiv: Sei $c \in \mathbb{N}$. Wenn c gerade ist, dann ist $-c/2 \in \mathbb{Z}$ und $f(-c/2) = c$. Wenn c ungerade ist, dann ist $(c - 1)/2 \in \mathbb{Z}$ und $f((c - 1)/2) = c$. $\square$

Es sei angemerkt, dass im letzten Beweis auch eine bijektive Funktion $f : \mathbb{N} \rightarrow \mathbb{Z}$ ausreichend wäre.

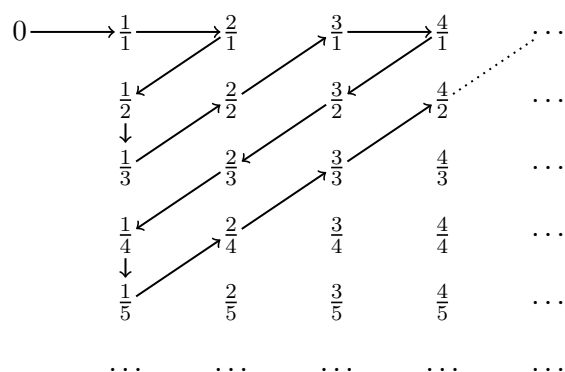
Übungsaufgabe 4.3.6. Zeige, dass die Menge aller Vielfachen von 3 abzählbar ist.

Satz 4.3.7. Die Menge $\mathbb{Q}$ der rationalen Zahlen ist abzählbar, d.h. gleichmächtig zu $\mathbb{N}$.

Beweis. Man muss sich klarmachen, dass man die Brüche abzählen kann. Dafür kann man das Cantorsche Abzählbarschema verwenden. Man trägt die nicht-negativen Brüche zunächst dafür nach folgendem Schema in einer zweidimensionalen Matrix auf:

0	$\frac{1}{1}$	$\frac{2}{1}$	$\frac{3}{1}$	$\frac{4}{1}$	...
	$\frac{1}{2}$	$\frac{2}{2}$	$\frac{3}{2}$	$\frac{4}{2}$	...
	$\frac{1}{3}$	$\frac{2}{3}$	$\frac{3}{3}$	$\frac{4}{3}$	...
	$\frac{1}{4}$	$\frac{2}{4}$	$\frac{3}{4}$	$\frac{4}{4}$	...
	$\frac{1}{5}$	$\frac{2}{5}$	$\frac{3}{5}$	$\frac{4}{5}$	...
	...	...	...	...	...

Nun zähle diese Brüche durch Ablaufen der Diagonalen ab:



In dieser Folge der Bruchzahlen entferne doppelte (nicht gekürzte) Brüche und füge hinter jeden verbleibenden Bruch $\frac{p}{q}$ ($\neq 0$) den negativen Bruch $-\frac{p}{q}$ hinzu.

Das ergibt die Folge $0 \rightarrow \frac{1}{1} \rightarrow -\frac{1}{1} \rightarrow \frac{2}{1} \rightarrow -\frac{2}{1} \rightarrow \frac{1}{2} \rightarrow -\frac{1}{2} \rightarrow \frac{3}{1} \rightarrow -\frac{3}{1} \rightarrow \frac{2}{2} \rightarrow -\frac{2}{2} \rightarrow \frac{4}{1} \rightarrow -\frac{4}{1} \rightarrow \frac{3}{2} \rightarrow -\frac{3}{2} \rightarrow \frac{1}{3} \rightarrow -\frac{1}{3} \rightarrow \frac{4}{2} \rightarrow -\frac{4}{2} \rightarrow \frac{1}{4} \rightarrow -\frac{1}{4} \rightarrow \frac{5}{1} \rightarrow -\frac{5}{1} \rightarrow \frac{2}{3} \rightarrow -\frac{2}{3} \rightarrow \frac{3}{3} \rightarrow -\frac{3}{3} \rightarrow \frac{1}{5} \rightarrow -\frac{1}{5} \rightarrow \frac{4}{3} \rightarrow -\frac{4}{3} \rightarrow \frac{3}{4} \rightarrow -\frac{3}{4} \rightarrow \frac{1}{5} \rightarrow -\frac{1}{5} \rightarrow \frac{5}{2} \rightarrow -\frac{5}{2} \rightarrow \frac{2}{5} \rightarrow -\frac{2}{5} \rightarrow \frac{3}{5} \rightarrow -\frac{3}{5} \rightarrow \frac{4}{5} \rightarrow -\frac{4}{5} \rightarrow \dots$
 Verwende diese Folge zum Abzählen, d.h. $f : \mathbb{N} \rightarrow \mathbb{Q}$ mit $f(1) = 0$, $f(2) = \frac{1}{1}$, $f(3) = -\frac{1}{1}$, ...
 Dann ist f eine Bijektion. $\square$

Nicht alle unendlichen Mengen sind abzählbar:

Satz 4.3.8. Die Potenzmenge der natürlichen Zahlen $\mathcal{P}(\mathbb{N})$ ist nicht abzählbar. Es gilt sogar, der folgende Satz von Cantor: Die Potenzmenge $\mathcal{P}(M)$ jeder Menge M hat eine größere Mächtigkeit als M .

Beweis. Für endliche Mengen ist $|\mathcal{P}(M)| = 2^{|M|} > |M|$. Für unendliche Mengen: Nehme an, es gibt eine bijektive Abbildung $f : M \rightarrow \mathcal{P}(M)$. Dann definiere

$$A := \{x \in M \mid x \notin f(x)\}$$

D.h. A enthält alle Elemente x aus M , die durch f auf Teilmengen von M abgebildet werden, die x nicht selbst enthalten.

Da $A \subseteq M$ und f surjektiv ist, gibt es ein $a \in M$ mit $f(a) = A$.

Nun betrachte zwei Fälle:

- $a \in f(a) = A$: Das ist ein Widerspruch, da A nur Elemente enthält, für die $x \notin f(x)$ gilt, also hier insbesondere $a \notin f(a)$ gilt.
- $a \notin f(a) = A$: Das ist ein Widerspruch, da $a \in A$ gelten müsste, denn da $a \notin f(a)$ gilt, ist die Aussageform in der Definition von A für $x = a$ erfüllt.

In beiden Fällen erhält man einen Widerspruch. Daher war unsere Annahme falsch: Die Bijektion f existiert nicht. $\square$

Satz 4.3.9. Die Menge der reellen Zahlen ist nicht abzählbar. Es gilt sogar: Das reelle Intervall $[0, 1)$ ist nicht abzählbar.

Beweis. Der Beweis folgt durch Widerspruch. Nehme also an, dass sich die reellen Zahlen aus $[0, 1)$ abzählen lassen. D.h. es gibt eine Bijektion $f : \mathbb{N} \rightarrow [0, 1)$. Wir schreiben die Dezimalbrüche $f(1), f(2), f(3), \dots$ in dieser Reihenfolge untereinander auf:

$$\begin{aligned} f(1) &= 0, \quad d_{1,1} \ d_{1,2} \ d_{1,3} \ d_{1,4} \ d_{1,5} \cdots \\ f(2) &= 0, \quad d_{2,1} \ d_{2,2} \ d_{2,3} \ d_{2,4} \ d_{2,5} \cdots \\ f(3) &= 0, \quad d_{3,1} \ d_{3,2} \ d_{3,3} \ d_{3,4} \ d_{3,5} \cdots \\ f(4) &= 0, \quad d_{4,1} \ d_{4,2} \ d_{4,3} \ d_{4,4} \ d_{4,5} \cdots \\ f(5) &= 0, \quad d_{5,1} \ d_{5,2} \ d_{5,3} \ d_{5,4} \ d_{5,5} \cdots \\ &\dots \end{aligned}$$

Dabei sind $d_{i,j}$ die einzelnen Ziffern der Nachkommastellen.

Nun konstruiere einen neuen Dezimalbruch $b = 0, b_1 b_2 \dots$: Wähle für jede der Ziffern b_i immer so, dass gilt $b_i \neq d_{i,i}$ (das geht immer, wir haben für jedes b_i sogar $|\{0, \dots, 9\} \setminus \{d_{i,i}\}| = 9$ Möglichkeiten).

$$\begin{aligned} b &= 0, \quad b_1 \quad b_2 \quad b_3 \quad b_4 \quad b_5 \dots \\ f(1) &= 0, \quad \mathbf{d_{1,1}} \quad d_{1,2} \quad d_{1,3} \quad d_{1,4} \quad d_{1,5} \dots \\ f(2) &= 0, \quad d_{2,1} \quad \mathbf{d_{2,2}} \quad d_{2,3} \quad d_{2,4} \quad d_{2,5} \dots \\ f(3) &= 0, \quad d_{3,1} \quad d_{3,2} \quad \mathbf{d_{3,3}} \quad d_{3,4} \quad d_{3,5} \dots \\ f(4) &= 0, \quad d_{4,1} \quad d_{4,2} \quad d_{4,3} \quad \mathbf{d_{4,4}} \quad d_{4,5} \dots \\ f(5) &= 0, \quad d_{5,1} \quad d_{5,2} \quad d_{5,3} \quad d_{5,4} \quad \mathbf{d_{5,5}} \dots \\ &\dots \end{aligned}$$

Es gibt keine Zahl $n \in \mathbb{N}$ mit $f(n) = b$, denn b kommt offensichtlich in der Tabelle nicht vor, denn b ist an der i . Nachkommastelle verschieden von $f(i)$. Dies ist ein Widerspruch dazu, dass f surjektiv ist. Daher kann f nicht existieren.

□

Definition 4.3.10. Eine Menge die nicht abzählbar ist, nennt man überabzählbar.

Satz 4.3.11. Die reellen Zahlen, das reelle Intervall $[0, 1)$ als auch die Potenzmenge der natürlichen Zahlen, die Potenzmenge der ganzen Zahlen, die Potenzmenge der rationalen Zahlen sind allesamt überabzählbar.

Beachte, dass Satz 4.3.8 auch zeigt, dass die überabzählbaren Mengen nicht alle gleichmächtig sind. So ist z.B. die Potenzmenge der reellen Zahlen wieder mächtiger als die reellen Zahlen selbst. Daraus kann man eine Theorie entwickeln, die sich mit diesen verschiedene Unendlichkeiten beschäftigt und daraus wieder Zahlen macht (diese heißen Kardinalzahlen). Interessanterweise kann man auch mit diesen Zahlen wieder rechnen.

4.4 Schlussbemerkungen

Einführungen zu Relationen und Funktionen sind z.B. in (TT13; Ber24; Ebb21; IL21) zu finden. Zur Unendlichkeit, Hilberts Hotel und weiterführenden Konstruktionen sei (OE17) empfohlen.

5 Beweise und Beweisen

5.1 Einführendes

Mathematische Texte bestehen oft aus Definitionen, Sätzen und deren Beweise. Wir gehen zunächst auf diese und weitere Begriffe ein, um klar zu machen, wann was verwendet wird, und wo die Beweise ins Spiel kommen.

5.1.1 Aufbau mathematischer Texte

Ein *Axiom* ist eine Grundaussage, die nicht bewiesen, sondern als gültig angenommen wird. Eine *Definition* führt neue Begriffe und Notationen ein. Diese enthalten keine Aussagen und werden daher auch nicht bewiesen. Ein *Satz* formuliert Aussagen (oder genauer: fast immer Wenn-Dann-Aussagen), diese müssen bewiesen werden. Je nach Wichtigkeit sind Sätze *Theoreme* (sehr wichtig), *Sätze* (wichtig), *Lemmas* (Hilfssätze, die zum Beweis anderer Sätze oder Theoreme verwendet werden), *Korollare* (Sätze, die direkt aus anderen Sätzen folgen, und daher ohne Beweis gegeben werden). *Beweise* zeigen die Korrektheit von Sätzen. *Bemerkungen* enthalten Erläuterungen, Motivationen, manchmal interessante Beispiele u.s.w. *Vermutungen* sind Aussagen, die nicht bewiesen sind und deren Wahrheitswert daher ungeklärt ist. Vermutung stellt man meistens dann auf, wenn man (noch) nicht weiß, wie man die Aussage beweisen soll, aber es einen (vernünftigen) Grund gibt, warum die Aussage stimmen könnte. Ein Beispiel ist die berühmte Goldbachsche Vermutung, dass sich jede gerade natürliche Zahl größer als 2 also Summe zweier Primzahlen darstellen lässt. Die Vermutung stammt aus dem Jahr 1742 und ist bis heute nicht bewiesen (aber auch nicht widerlegt). *Beispiele* illustrieren definierte Begriffe und Aussagen, manchmal auch Konstruktionen, die innerhalb von Beweisen verwendet werden. Sie tragen meistens sehr zum Verständnis bei. Es ist sehr sinnvoll alle diese Kategorien zu nummerieren und zu referenzieren, damit eindeutig klar ist, was warum folgt, welche Definition genau gemeint ist, u.s.w.

5.1.2 Warum will man beweisen?

Formale Beweise klären Sachverhalte ein für allemal, die Aussage stimmt oder sie stimmt nicht. Sie helfen Kritik und Irrtümer auszuräumen. Schließlich dienen sie auch der Erschaffung von neuem Wissen. In der Informatik haben wir es oft mit Problemen zu tun, für welche Programme (Algorithmen) entworfen werden sollen, um diese zu lösen. Bevor man diese Algorithmen einsetzt, wäre man sich gerne sicher, dass der Algorithmus korrekt ist oder dass er eine gewisse Garantie bezüglich Laufzeit oder Platzverhalten hat. Beweise können diese Frage ein für allemal

beantworten und sie geben oft einen tieferen Einblick in den Algorithmus – wie und warum er funktioniert. Aber auch negative Aussagen der Form „Für dieses Problem gibt es keinen Algorithmus“ kann man beweisen. Manchmal hilft ein solcher Beweis, da man keine Zeit vergeudet mit der Suche nach dem Unmöglichen.

5.1.3 Was kennzeichnet Sätze und Beweise?

Beispiele genügen nicht, um als Beweis zu dienen. Z.B. kann man die Behauptung „Für alle $n \in \mathbb{N}$: $n^2 + n + 41$ ist eine Primzahl“ für die ersten 30 Zahlen prüfen, und wird keinen Ausreißer finden. Trotzdem ist dies kein Beweis! Insbesondere betrachte $n = 41$: Das ergibt $41^2 + 41 + 41 = 41 \cdot 43$, was keine Primzahl ist (schon für $n = 40$ ergibt sich keine Primzahl, denn $40^2 + 40 + 41 = 40 \cdot (40 + 1) + 41 = 41 \cdot (40 + 1) = 41 \cdot 41$).

Mathematische Sätze müssen daher *allgemein bewiesen* oder durch ein *Gegenbeispiel widerlegt* werden. Die reine Betrachtung von Beispielen ist nicht ausreichend.

Mathematische Sätze sind Wenn-Dann-Aussagen der Form $F \rightarrow G$. Unter den Aussagen F (der *Voraussetzung*) folgt eine andere Aussage G (die *Behauptung*).

Beispiel 5.1.1. Der Satz des Pythagoras lautet nicht nur $a^2 + b^2 = c^2$, sondern:

Wenn a, b, c die Seitenlängen eines rechtwinkligen Dreiecks sind, dann gilt $a^2 + b^2 = c^2$.

Die Behauptung ergibt sich aus der Voraussetzung durch die Gesetze der Logik und *nur* durch diese. Diese Ableitung nennt man einen *Beweis*. Man beendet einen mathematischen Beweis mit der Box $\square$ oder mit „Q.E.D.“ (was für „quod erat demonstrandum“¹ steht). Der Sinn dabei ist, das Ende des Beweises eindeutig zu kennzeichnen, damit Lesende wissen, wann der Beweis vorbei ist.

5.2 Beweisarten

Wir erläutern verschiedene Arten einen Beweis zu führen.

5.2.1 Direkter Beweis

Beim direkten Beweis, wird aus der Voraussetzung F „direkt“ die Behauptung G bewiesen. Dies geschieht nach dem Muster:

Beweis. Sei F erfüllt.
... (Diese Lücke ist zu füllen)
Also gilt G .

$\square$

¹lat. für „was zu beweisen war“

Beispiel 5.2.1. Wir beweisen den folgenden Satz mit einem direkten Beweis

Satz. Wenn eine natürliche Zahl n durch 10 teilbar ist, dann ist n gerade.

Beweis. Sei n durch 10 teilbar. Dann gibt es $k \in \mathbb{N}$ mit $n = 10 \cdot k$. D.h. $n = 2 \cdot (5 \cdot k)$. Daher ist 2 ein Teiler von n . Also ist n gerade. $\square$

Übungsaufgabe 5.2.2. Beweise den folgenden Satz mit einem direkten Beweis:

Satz. Seien $a, b \in \mathbb{Q}$ rationale Zahlen mit $a < b$. Dann gibt es $c \in \mathbb{Q}$ mit $a < c < b$.

5.2.2 Kontraposition

Beim Beweis durch Kontraposition verwendet man die logische Äquivalenz $(F \rightarrow G) \equiv (\neg G \rightarrow \neg F)$. D.h. anstelle von F auf G zu schließen, schließt man von $\neg G$ auf $\neg F$. Das Muster für einen Beweis durch Kontraposition ist daher:

Beweis. Sei $\neg G$ erfüllt.
... (Diese Lücke ist zu füllen)
Also gilt $\neg F$. $\square$

Beispiel 5.2.3. Wir beweisen den folgenden Satz:

Satz. Wenn eine natürliche Zahl n durch 10 teilbar ist, dann ist n gerade.

Beweis. Wir verwenden Kontraposition und zeigen die äquivalente Aussage: Wenn n eine ungerade Zahl ist, dann ist n nicht durch 10 teilbar. Sei n eine ungerade natürliche Zahl. Dann ist 2 kein Teiler von n . Dann ist $n \neq 2 \cdot k$ für alle $k \in \mathbb{N}$. Dann ist $n \neq 10 \cdot k'$ für alle $k' \in \mathbb{N}$. Dann ist n nicht durch 10 teilbar. $\square$

Beispiel 5.2.4. Wir beweisen den folgenden Satz mit Kontraposition.

Satz. Wenn n^2 eine ungerade Zahl ist, dann ist n eine ungerade Zahl.

Beweis. Wir verwenden Kontraposition und zeigen die äquivalente Aussage: Wenn n keine ungerade Zahl ist, dann ist n^2 keine ungerade Zahl. Sei n keine ungerade natürliche Zahl. Dann ist n gerade. Dann ist 2 ein Teiler von n , d.h. $n = 2k$ für ein $k \in \mathbb{N}$. Dann ist $n^2 = 4k^2 = 2 \cdot (2k^2)$. Dann ist n^2 eine gerade Zahl. Also ist n^2 keine ungerade Zahl. $\square$

Übungsaufgabe 5.2.5. Beweise durch Kontraposition: Wenn $n, m \in \mathbb{N}$ mit $n \cdot m$ ist nicht durch 3 teilbar, dann sind auch n und m nicht durch 3 teilbar.

5.2.3 Beweis durch Widerspruch

Beim Beweis durch Widerspruch nimmt man an, dass die zu zeigende Aussage falsch ist, d.h. F gilt, aber G gilt nicht. Dann zeigt, man dass dies ein Widerspruch ist, was letztendlich zeigt, dass die Annahme falsch war und G gelten muss. Logisch begründbar ist dieses Vorgehen wie folgt. Statt $F \rightarrow G$ als wahr (eine Tautologie) zu zeigen, zeigt man $\neg(F \rightarrow G)$ ist falsch (ein Widerspruch). Logisch umgeformt ergibt dies $\neg(F \rightarrow G) \equiv \neg(\neg F \vee G) \equiv F \wedge \neg G$.

Das Muster für einen Beweis durch Widerspruch ist:

Beweis. Sei F erfüllt. Angenommen G wäre falsch.

... (Diese Lücke ist zu füllen)

Also ergibt sich ein Widerspruch. Daher war die Annahme falsch und G gilt. $\square$

Beispiel 5.2.6. Wir haben bereits beim Beweis der Überabzählbarkeit der reellen Zahlen einen Widerspruchsbeweis geführt.

Beispiel 5.2.7. Wir beweisen den folgenden Satz mit einem Widerspruchsbeweis.

Satz. Wenn n^2 eine ungerade Zahl ist, dann ist n eine ungerade Zahl.

Beweis. Sei n^2 eine ungerade Zahl. Wir nehmen an n ist eine gerade Zahl. Dann ist $n = 2k$ für ein $k \in \mathbb{N}$. Dann ist $n^2 = 4k^2 = 2(2k^2)$. Daher ist n^2 eine gerade Zahl – ein Widerspruch. Daher war die Annahme falsch und n muss eine ungerade Zahl sein. $\square$

Beispiel 5.2.8. Wir beweisen den folgenden Satz mit einem Widerspruchsbeweis.

Satz. Die reelle Zahl $\sqrt{2}$ ist keine rationale Zahl.

Beweis. Wir verwenden einen Beweis durch Widerspruch. Sei $\sqrt{2}$ eine rationale Zahl. Dann gibt es $p \in \mathbb{N}_0, q \in \mathbb{N}$ mit $\frac{p}{q} = \sqrt{2}$, wobei p und q teilerfremd sind (der Bruch ist gekürzt).

Dann ist $\left(\frac{p}{q}\right)^2 = \frac{p^2}{q^2} = 2$ oder anders ausgedrückt $p^2 = 2q^2$. Damit ist p^2 durch 2 teilbar, was nur sein kann, wenn p durch 2 teilbar ist. Dann muss es eine Zahl m geben mit $p = 2m$. Dann gilt $4m^2 = 2q^2$ oder anders ausgedrückt $q^2 = 2m^2$. Daher ist q^2 durch 2 teilbar, was nur sein kann, wenn q durch 2 teilbar ist, d.h. es muss n geben mit $2n = q$. Dann ist aber $\frac{p}{q} = \frac{2m}{2n}$ noch kürzbar. Ein Widerspruch! Daher war unsere Annahme falsch und $\sqrt{2}$ ist keine rationale Zahl. $\square$

5.2.4 Äquivalenzen

Äquivalenzen sind Sätze der Form $F \leftrightarrow G$. Da dies logisch äquivalent zu $(F \rightarrow G) \wedge (G \rightarrow F)$ ist, können wir den Beweis einer solchen Äquivalenz in zwei Teilen führen nach dem Muster:

Beweis. Wir zeigen beide Richtungen.

- $F \rightarrow G: \dots$
- $G \rightarrow F: \dots$

□

Beispiel 5.2.9. Wir beweisen den folgenden Satz mit einem Äquivalenzbeweis:

Satz. Für jede ganze Zahl $z \in \mathbb{Z}$ gilt: $z > 0 \leftrightarrow -z < 0$.

Beweis. Wir zeigen beide Richtungen.

- Sei $z > 0$. Dann gibt es eine natürliche Zahl k mit $z - k = 0$. Umformen ergibt $k = z$. Dann gilt $-z = -k$, d.h. $-z = 0 - k$. Damit folgt $-z < 0$.
- Sei $-z < 0$. Dann gibt es eine natürliche Zahl k mit $0 - k = z$ d.h. $-z = -(0 - k) = -0 + k = 0 + k$. D.h. $-z > 0$. □

5.3 Fallunterscheidung

Ein oft verwendetes Mittel beim Beweisen ist das Unterscheiden in verschiedene Fälle. Da $(F \rightarrow G) \wedge (\neg F \rightarrow G) \equiv G$ kann man die Aussage G beweisen, indem man die Fälle F gilt und F gilt nicht jeweils einmal annimmt und zeigt, dass in beiden Fällen auch G gelten muss. Man unterscheidet also die Fälle

- F gilt
- F gilt nicht

Das ganze kann man natürlich wiederholt machen (z.B. noch zusätzlich unterscheiden, ob Formel Q gilt, oder Q nicht gilt (was immer Q genau aussagt). Wichtig beim Fallunterscheiden ist, dass man keinen Fall vergisst („unter den Tisch fallen lässt“).

Wir betrachten als ganz einfaches Beispiel

Satz 5.3.1. Für jede ganze Zahl z gilt $z^2 \geq 0$.

Beweis. Wir unterscheiden zwei Fälle: $z < 0$ und $z \geq 0$.

1. $z < 0$: Dann ist $z^2 = z \cdot z$ das Produkt zweier negativer Zahlen, was positiv ist. Also gilt die Aussage.
2. $z \geq 0$. Wir unterscheiden nochmal: $z = 0$ oder $z > 0$.
 - a) $z = 0$: $0^2 = 0 \cdot 0 = 0 \geq 0$
 - b) $z > 0$: Dann ist $z^2 = z \cdot z$ das Produkt zweier positiver Zahlen, welches positiv ist.

□

Noch ein Beispiel:

Satz 5.3.2. Für jede natürliche Zahl $n > 1$ ist $n(n^2 - 1)$ ein Vielfaches von 3.

Beweis. Wir unterscheiden drei Fälle: $n \bmod 3 = 0$, $n \bmod 3 = 1$, $n \bmod 3 = 2$

1. $n \bmod 3 = 0$, dann gibt es $k \in \mathbb{N}$ mit $n = 3 \cdot k$ und $3 \cdot k \cdot (9k^2 - 1)$ ist offensichtlich ein Vielfaches von 3.
2. $n \bmod 3 = 1$, dann gibt es $k \in \mathbb{N}$ mit $n = 3 \cdot k + 1$ und $(3k + 1) \cdot ((3k + 1)^2 - 1) = (3k + 1) \cdot (9k^2 + 6k + 1 - 1) = (3k + 1) \cdot (9k^2 + 6k) = (3k + 1) \cdot (3k^2 + 2k) \cdot 3$, was ein Vielfaches von 3 ist.
3. $n \bmod 3 = 2$, dann gibt es $k \in \mathbb{N}$ mit $n = 3 \cdot k + 2$ und $(3k + 2) \cdot ((3k + 2)^2 - 1) = (3k + 2) \cdot (9k^2 + 12k + 4 - 1) = (3k + 2) \cdot (9k^2 + 12k + 3) = (3k + 2) \cdot (3k^2 + 4k + 1) \cdot 3$, was ein Vielfaches von 3 ist.

□

Beachte, dass wir hier aus logischer Sicht zweimal unterschieden haben. Z.B. erst in $n \bmod 3 = 0$ und $n \bmod 3 \neq 0$ und dann für den zweiten Fall nochmal in $n \bmod 3 = 1$ und $n \bmod 3 \neq 1$ (da bleibt dann nur $n \bmod 3 = 2$ übrig)

5.4 Schubfachprinzip

Das Schubfachprinzip ist ein einfaches Mittel, um Aussagen über endliche Mengen zu beweisen. Man teilt dabei die Objekte einer Menge in unterschiedliche Kategorien (d.h. paarweise disjunkte Teilmengen):

Satz 5.4.1 (Schubfachprinzip, auch Taubenschlagprinzip). *Seien m Objekte in n Kategorien eingeteilt. Wenn $m > n$ ist, so gibt es mindestens eine Kategorie, die mindestens zwei Objekte enthält.*

Beweis. Wir führen einen Beweis durch Widerspruch und nehmen an, dass wir $m > n$ Objekte in n Kategorien verteilt haben, aber keine Kategorie mehr als ein Objekt enthält.

Dann ist die Gesamtzahl aller Objekte höchstens so groß wie die Anzahl an Kategorien, d.h. höchstens n . Dies ist ein Widerspruch dazu, dass wir $m > n$ Objekte gegeben haben. Daher war unsere Annahme falsch und das Schubfachprinzip gilt. □

Der Name Taubenschlagprinzip kommt von Dirichlets Veranschaulichung des Prinzips: Wenn man viele Tauben auf wenige Taubenschläge verteilt, dann sitzen in einem Taubenschlag mindestens zwei Tauben.

Beispiel 5.4.2. *Weitere einfache Beispiele, die aus dem Schubfachprinzip folgen, sind:*

- *Unter 13 Personen gibt, es mindestens zwei, die im selben Monat Geburtstag haben, unter 367 gibt es mindestens zwei die am selben Tag und Monat Geburtstag haben.*
- *Unter 4 Studierenden aus den Studiengängen AI,TS,WI gibt es mindestens zwei aus demselben Studiengang.*

Übungsaufgabe 5.4.3. *Wie viele Personen sind mindestens notwendig, damit zwei am gleichen Wochentag Geburtstag haben?*

Beispiel 5.4.4. *In der Sockenkiste von Emil befinden sich 8 graue und 8 braune Socken. Wie viele muss er herausnehmen, um*

- *garantiert zwei gleichfarbige Socken zu erhalten?*
- *garantiert zwei graue Socken zu erhalten?*

Für die erste Frage, kann man Satz 5.4.1 mit zwei Kategorien anwenden (graue und braune Socken). Das zeigt sofort, dass man nur 3 Socken nehmen muss, um garantiert zwei Socken einer Kategorie (d.h. gleichfarbige Socken) zu erhalten. Für die zweite Frage, muss man vom schlimmsten Fall ausgehen: Emil nimmt erst alle braunen Socken heraus. Danach bekommt er sicher zwei graue. D.h. er muss 10 Socken herausnehmen.

Als größeres Beispiel zeigen wir die folgende Aussage:

Satz 5.4.5. *Unter je sechs natürlichen Zahlen gibt es stets zwei, deren Differenz durch 5 teilbar ist.*

Beweis. Wir verwenden das Schubfachprinzip. Die Objekte sind die 6 natürlichen Zahlen $k_1, \dots, k_6$. Die Kategorien sind: durch 5 teilbar ohne Rest, durch 5 teilbar mit Rest 1, durch 5 teilbar mit Rest 2, durch 5 teilbar mit Rest 3, durch 5 teilbar mit Rest 4

Das Schubfachprinzip sagt, dass mindestens eine der Kategorien auf zwei Zahlen zutrifft. Seien dies die Zahlen k_i und k_j (mit $1 \leq i < j \leq 6$) und die Kategorie sei „durch 5 teilbar mit Rest m “. Dann existieren $x, y \in \mathbb{Z}$ mit $k_i = 5 \cdot x + m$ und $k_j = 5 \cdot y + m$. Die Differenz $k_i - k_j = 5x - m - 5y - m = 5(x - y)$ ist offensichtlich durch 5 teilbar. $\square$

Übungsaufgabe 5.4.6. *Zeige mit dem Schubfachprinzip: Unter 4 natürlichen Zahlen gibt es stets zwei, deren Differenz durch 3 teilbar ist.*

Unter 13 Personen haben mindestens 2 im selben Monat Geburtstag, aber es gilt auch: Unter 25 Personen haben mindestens 3 im selben Monat Geburtstag. Man kann das Schubfachprinzip verallgemeinern:

Satz 5.4.7 (Verallgemeinertes Schubfachprinzip). *Seien m Objekte in n Kategorien eingeteilt. Wenn $m > r \cdot n$, dann gibt es mindestens eine Kategorie, die mindestens $r + 1$ Objekte enthält.*

Beweis. Beweis durch Widerspruch. Hätte jede Kategorie höchstens r Objekte, so gäbe es höchstens $r \cdot n$ Objekte, was ein Widerspruch zu $m > r \cdot n$ ist. $\square$

Übungsaufgabe 5.4.8. *Wie viele Tauben muss man auf 5 Taubenschläge mindestens verteilen, damit es sicher mindestens einen Taubenschlag mit mindestens 4 Tauben gibt?*

5.5 Vollständige Induktion

Wir verwenden im folgenden das *Summenzeichen*: Sei $k \leq n$, $h : \mathbb{N} \rightarrow \mathbb{R}$ eine reellwertige Funktion mit Definitionsbereich $\mathbb{N}$. Dann bezeichnet $\sum_{i=k}^n h(i)$ die Summe

$$h(k) + h(k+1) + \cdots + h(n)$$

. Für $k > n$ setzen wir per Definition (leere Summe) $\sum_{i=k}^n h(i) := 0$. D.h. das Summenzeichen dient dazu, dass man die Punkte $\dots$ nicht mehr benötigt. Z.B. ist

$$\sum_{i=1}^5 i^2 = 1^2 + 2^2 + 3^2 + 4^2 + 5^2 = 1 + 4 + 9 + 16 + 25 = 55 \text{ und } \sum_{i=3}^4 2^i = 2^3 + 2^4 = 8 + 16 = 24$$

Analog zum Summenzeichen gibt es auch ein Produktzeichen

$$\prod_{i=k}^n h(i) = h(k) \cdot h(k+1) \cdots h(n),$$

wobei das *leere Produkt* als $\prod_{i=k}^n := 1$ (falls $k > n$) definiert ist. Z.B. ist $\prod_{i=1}^n i = n!$

Oft möchte man Aussagen der Form „Für alle natürlichen Zahlen n gilt $A(n)$ “ nachweisen. Dabei ist $A(n)$ eine Aussage, die von n abhängt.

Beispiele aus der Mathematik und Informatik sind:

- Für alle Folgen von n Elementen berechnet der Algorithmus die sortierte Folge.
- Der Sortieralgorithmus benötigt bei n Eingaben nicht mehr als $n \log_2 n$ Vergleiche.
- Die Anzahl der möglichen Sitzordnungen für die Klausur für n Studierende auf n Stühlen ist $n!$.
- ...

Da es unendlich viele natürliche Zahlen gibt, kann man den Beweis einer solchen Aussage nicht einzeln für alle Fälle nachweisen. Die vollständige Induktion ist ein Beweisprinzip, welches genau für solche Fälle hilft:

Definition 5.5.1 (Beweisprinzip der Vollständigen Induktion). *Um zu zeigen, dass eine Aussage $A(n)$ für jede natürliche Zahl $n \in \mathbb{N}$ gilt, genügt es, die folgenden beiden Aussagen zu zeigen:*

1. (*Induktionsanfang/Induktionsbasis*): $A(1)$ gilt.
2. (*Induktionsschritt*): Für jede beliebige Zahl $n \in \mathbb{N}$ gilt: Nehme an, dass $A(n)$ gilt (*Induktionsannahme / Induktionshypothese*). Zeige, dass dann auch $A(n+1)$ gilt.

Für die Korrektheit des Beweisprinzips muss man sich klar machen, warum die Aussage $A(m)$ für beliebiges $m \in \mathbb{N}$ gelten muss: Fange bei $A(1)$ an: Die Induktionsbasis zeigt, dass die Aussage gilt. Jetzt wende $m - 1$ mal den Induktionsschritt an:

Da $A(1)$ gilt, gilt auch $A(2)$.

Da $A(2)$ gilt, gilt auch $A(3)$.

...

Da $A(m - 1)$ gilt, gilt auch $A(m)$.

5.5.1 Beispiele

Wir demonstrieren die vollständige Induktion anhand der Gaußschen Summenformel. Die Geschichte dazu ist, dass Carl Friedrich Gauß die Schulaufgabe bekam die Zahlen von 1 bis 100 zu summieren. Gauß löste die Aufgabe rasend schnell, indem er nicht die einzelnen Zahlen summierte, sondern beobachtete, dass die Summe der jeweils i . und $101-i$. Zahl genau 101 ergibt. Wenn man dann noch zählt, dass es genau $100/2=50$ von diesen Zahlen gibt, so erhält man die Summe als $50 \cdot 101 = 5050$. Die Gaußsche Summenformel verallgemeinert die Formel für die ersten n (statt 100) Zahlen.

Satz 5.5.2. Für alle $n \in \mathbb{N}$ gilt:

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

Beweis. Wir zeigen die Aussage durch vollständige Induktion über n , d.h. die Aussage $A(n)$ ist

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

- Induktionsbasis: Offensichtlich gilt Aussage $A(1) = \sum_{i=1}^1 i = 1$.
- Induktionsschritt: Sei $n \in \mathbb{N}$ eine beliebige Zahl. Induktionsannahme: $A(n)$ gilt, d.h. $\sum_{i=1}^n i = \frac{(n+1)n}{2}$. Wir müssen $A(n+1)$ zeigen, d.h. $\sum_{i=1}^{n+1} i = \frac{(n+2)(n+1)}{2}$ ist zu zeigen. Es gilt $\sum_{i=1}^{n+1} i = \left(\sum_{i=1}^n i \right) + n + 1$ nach Definition. Die Induktionsannahme zeigt: Letzteres ist gleich zu $\frac{(n+1)n}{2} + n + 1$. Durch Ausrechnen erhält man $\frac{(n+1)n}{2} + n + 1 = \frac{(n+1)n + 2n + 2}{2} = \frac{(n+2)(n+1)}{2}$, also insgesamt die zu zeigende Eigenschaft. $\square$

Satz 5.5.3. Die Summe der ersten n ungeraden Zahlen ist gleich zu n^2 (als Summenformel geschrieben $\sum_{i=1}^n 2i - 1 = n^2$).

Beweis. Mit vollständiger Induktion über n :

- Induktionsbasis $n = 1$: Da $1 = 1^2$ stimmt die Aussage.
- Induktionsschritt: Sei $n \geq 1$ beliebig und es gelte, dass die Summe der ersten n ungeraden Zahlen gleich zu n^2 ist (Induktionsannahme). Betrachte die Summe der ersten $n + 1$ ungeraden Zahlen. Die $n + 1$. ungerade Zahl ist $2(n + 1) - 1$. Daher ist die Summe der ersten $n + 1$ ungeraden Zahlen $\sum_{i=1}^{n+1} 2i - 1 = (\sum_{i=1}^n 2i - 1) + 2(n + 1) - 1 = (\sum_{i=1}^n 2i - 1) + 2n + 1$. Anwenden der Induktionshypothese zeigt, dass diese Summe gleich zu $n^2 + 2n + 1$ ist. Die erste binomische Formel liefert nun, dass dies wiederum gleich zu $(n + 1)^2$ ist. D.h. wir haben gezeigt $\sum_{i=1}^{n+1} 2i - 1 = (n + 1)^2$. $\square$

Übungsaufgabe 5.5.4. Beweise mit vollständiger Induktion nach n : $\sum_{i=0}^n 2^i = 2^{n+1} - 1$

5.5.2 Vollständige Induktion mit anderem Startwert

Manche Aussagen $A(n)$ gelten nicht für alle $n \in \mathbb{N}$, sondern erst ab einem bestimmten Wert, d.h. für alle n mit $n \geq n_0$. Betrachte z.B. die Aussage $n! > 2^n$. Diese gilt nicht für $n \in \{1, 2, 3\}$, aber für alle $n \geq 4$.

Das Schema aus Definition 5.5.1 passt also nicht, für einen Beweis. Ein verallgemeinerte Variante ist (in der wir auch 0 als Startwert zulassen)

Definition 5.5.5 (Beweisprinzip der Vollständigen Induktion mit anderem Startwert). Sei $n_0 \in \mathbb{N}_0$. Um zu zeigen, dass eine Aussage $A(n)$ für jede Zahl $n \geq n_0 \in \mathbb{N}_0$ gilt, genügt es, die folgenden beiden Aussagen zu zeigen:

1. (Induktionsanfang/Induktionsbasis): $A(n_0)$ gilt.
2. (Induktionsschritt): Für jede beliebige Zahl $n \geq n_0 \in \mathbb{N}_0$ gilt: Nehme an, dass $A(n)$ gilt (Induktionsannahme / Induktionshypothese). Zeige, dass dann auch $A(n + 1)$ gilt.

Das Schema aus Definition 5.5.1 ergibt sich aus Definition 5.5.5, indem man $n_0 = 1$ setzt. Die Korrektheit des verallgemeinerten Schemas, lässt sich mit dem normalen Schema zeigen, indem man definiert $A'(n) := A(n_0 + n)$ und anschließend mit dem Schema aus Definition 5.5.1 zeigt, dass $A'(n)$ für alle $n \in \mathbb{N}$ gilt.

Wir verwenden die vollständige Induktion mit Startwert:

Satz 5.5.6. Für jede natürliche Zahl $n \geq 4$ gilt $n! > 2^n$.

Beweis. Wir zeigen die Aussage mit vollständiger Induktion mit Startwert $n = 4$:

- Induktionsbasis: Für $n = 4$ gilt $n! = 4 \cdot 3 \cdot 2 \cdot 1 = 24 > 16 = 2^4$.
- Induktionsschritt. Sei $n \geq 4$. Wir nehmen an $n! > 2^n$ (Induktionsannahme) und zeigen $(n + 1)! > 2^{n+1}$: $(n + 1)! = (n + 1)n! > (n + 1)2^n > 2 \cdot 2^n = 2^{n+1}$. Hierbei nutzen wir im zweiten Schritt die Induktionsannahme und im dritten Schritt, dass für $n \geq 4$, die Ungleichung $n + 1 > 2$ gilt.

□

Bemerkung 5.5.7. Man sollte sich klar machen, dass die Aussage

Für alle $n \in \mathbb{N} : A(n)$, wobei $A(n) := n \geq 4 \rightarrow n! > 2^n$.

äquivalent zu Satz 5.5.6 ist, aber sich nicht so einfach mit vollständiger Induktion mit Startwert 1 zeigen lässt. Die Basis $A(1)$ gilt sofort, da $1 \geq 4$ falsch ist. Im Induktionsschritt müsste man aber aus $A(n) = n \geq 4 \rightarrow n! > 2^n$ die Aussage $A(n+1) := n+1 \geq 4 \rightarrow (n+1)! > 2^{(n+1)}$ folgern, was erst funktioniert, wenn sichergestellt ist, dass $n \geq 4$ ist.

Eine Alternative wäre es, die Aussagen $A(1), A(2), A(3), A(4)$ als Basis zu zeigen, um dann im Induktionsschritt erst auf die Induktionsannahme für $n \geq 4$ zurückzugreifen.

Zum Abschluss des Abschnitts beweisen Satz 3.3.18 erneut, diesmal mit einem Induktionsbeweis:

Satz 5.5.8. Jede n -elementige Menge hat genau 2^n Teilmengen, d.h. für endliche Mengen gilt $|\mathcal{P}(M)| = 2^{|M|}$.

Beweis. Hier müssen wir mit $n = 0$ anfangen, da wir die Aussage auch für leere Mengen verwenden wollen.

- Induktionsbasis $n = 0$: Die leere Menge $\emptyset$ hat nur sich als Teilmenge, d.h. $|\mathcal{P}(\emptyset)| = |\{\emptyset\}| = 1 = 2^0$
- Induktionsschritt: Sei $n \geq 0$ beliebig. Wir nehmen an, dass jede n -elementige Menge 2^n Teilmengen hat. Sei M eine $n+1$ -elementige Menge. Sei m ein Element aus M (das muss existieren, da M mindestens ein Element hat). Dann ist $M \setminus \{m\}$ eine n -elementige Menge. Die Induktionsannahme, zeigt, dass $M \setminus \{m\}$ 2^n Teilmengen hat. Zu jeder dieser Teilmengen können wir m hinzufügen (und erhalten eine Teilmenge von M) oder nicht hinzufügen (und erhalten trotzdem eine Teilmenge von M). D.h. M hat doppelt so viele Teilmengen wie $M \setminus \{m\}$, was $2 \cdot 2^n = 2^{n+1}$ ergibt.

□

5.5.3 Starke vollständige Induktion, Fibonacci-Zahlen und der Goldene Schnitt

Manchmal genügt es nicht, im Induktionsschritt $A(n) \rightarrow A(n+1)$ nur auf die Gültigkeit von $A(n)$ zurückzugreifen, sondern man benötigt z.B. auch die Gültigkeit von $A(n-1)$.

Das Induktionsprinzip der *starken Induktion* ermöglicht dies:

Definition 5.5.9 (Beweisprinzip der starken vollständigen Induktion). Sei $n_0 \in \mathbb{N}$. Um zu zeigen, dass eine Aussage $A(n)$ für jede natürliche Zahl $n \geq n_0 \in \mathbb{N}$ gilt, genügt es, die folgende Aussage zu zeigen:

(Induktionsschritt): Für jede beliebige Zahl $n \geq n_0 \in \mathbb{N}$ gilt: Nehme an, dass für alle i mit $n_0 \leq i \leq n-1$ die Aussage $A(i)$ gilt (Induktionsannahme / Induktionshypothese). Zeige, dass dann auch $A(n)$ gilt.

Diese Schema kommt ohne expliziten Basisfall aus, aber dieser versteckt sich im Induktionsschritt: Die Gültigkeit von $A(n_0)$ ist durch den Induktionsschritt abgedeckt ($n = n_0$), aber die Voraussetzung „für alle $i : n_0 \leq i \leq n-1$ gilt $A(i)$ “ ist leer, denn es gibt keine i mit $n_0 \leq i \leq n_0-1$. D.h. in diesem Fall muss $A(n_0)$ ohne Induktionsannahme gezeigt werden.

Das Zurückgreifen auf vorherige Annahmen im Induktionsschritt erfordert im Allgemeinen, dass man mehr Basisfälle zeigen muss, genauer gilt: Wenn der Induktionsschritt zum Folgern der Gültigkeit von $A(n)$ die Gültigkeit von $A(n-k)$ benötigt, dann kann man dieses Argument nicht für $A(n_0), A(n_0+1), \dots, A(n_0+k-1)$ verwenden, da jedes mal „hinter“ $A(n_0)$ geschaut werden müsste. Daher muss man all diese Fälle als Basis zeigen.

Betrachte als Beispiel den folgenden **falschen Beweis**:

Aussage: Alle natürlichen Zahlen ab 3 sind ungerade.

Falscher Induktionsbeweis:

Induktionsbasis: Für $n = 3$ gilt die Aussage.

Induktionsschritt: Sei $n \geq 3$ beliebig. Induktionsannahme $A(i)$ gilt für $1 \leq i \leq n-1$, d.h. die Zahlen $1 \leq i \leq n-1$ sind ungerade. Da damit insbesondere $n-2$ ungerade ist, ist auch $n-2+2$ ungerade. Damit folgt n ist ungerade.

Der Fehler im Induktionsschritt ist z.B. für die Zahl $n = 4$. Diese greift auf $n-2 = 2$ zurück, aber die Induktionsannahme gibt $A(2)$ nicht her. Daher hätten wir $A(4)$ auch direkt als Basis zeigen müssen (was wir nicht können, da 4 keine ungerade Zahl ist)

5.5.4 Fibonacci-Zahlen

Leonardo Fibonacci beschrieb (1202) das Wachstum einer Kaninchenpopulation nach den folgenden Regeln:

- Zu Beginn gibt es ein Kaninchenpaar.
- Jedes Kaninchenpaar braucht 2 Monate nach der Geburt, bis es geschlechtsreif ist.
- Von da an gebiert es in jedem Monat ein neues Paar.
- Alle Kaninchen leben ewig.

Fragt man sich nun, wie viele Kaninchenpaare es zu Beginn des n . Monat gibt, kann man zunächst eine Tabelle erstellen:

Monat	Paare
1	1
2	1
3	2 (1 altes Paar + 1 neues Paar (Nachkommen aller Paare aus Monat 1))
4	3 (2 alte Paare + 1 neues Paar (Nachkommen aller Paare aus Monat 2))
5	5 (3 alte Paare + 2 neue Paare (Nachkommen aller Paare aus Monat 3))
6	8 (5 alte Paare + 3 neue Paare (Nachkommen aller Paare aus Monat 4))
7	13 (8 alte Paare + 5 neue Paare (Nachkommen aller Paare aus Monat 5))

Allgemein kann man die Kaninchenzahl zu Beginn des n Monats berechnen als die Anzahl an Kaninchen des $n - 1$. Monats plus die neuen Paare, welche genauso viele sind, wie die Kaninchenpaare zu Beginn des $n - 2$. Monats.

Die Folge $1, 1, 2, 3, 5, 8, 13, 21, \dots$ nennt man die *Fibonacci-Zahlen*. Die n . Fibonacci-Zahl kann durch die Funktion $\text{fib} : \mathbb{N} \rightarrow \mathbb{N}$ rekursiv definiert werden:

$$\text{fib}(n) := \begin{cases} 1 & \text{für } n \in \{1, 2\} \\ \text{fib}(n-1) + \text{fib}(n-2) & \text{für } n > 2 \end{cases}$$

Satz 5.5.10 (Binet-Formel). Für alle $n \in \mathbb{N}$ gilt

$$\text{fib}(n) = \frac{\left(\frac{1+\sqrt{5}}{2}\right)^n - \left(\frac{1-\sqrt{5}}{2}\right)^n}{\sqrt{5}}$$

Beweis. Wir verwenden starke vollständige Induktion. Da wir im Induktionsschritt die Binet-Formel für $\text{fib}(n-1)$ und $\text{fib}(n-2)$ verwenden, zeigen wir die Induktionsbasis für $n = 1$ und $n = 2$.

- Induktionsbasis für $n = 1$ und $n = 2$:

$$\text{fib}(1) = 1 = \frac{\left(\frac{1+\sqrt{5}}{2}\right)^1 - \left(\frac{1-\sqrt{5}}{2}\right)^1}{\sqrt{5}} = \frac{\frac{1}{2} + \frac{\sqrt{5}}{2} - \frac{1}{2} + \frac{\sqrt{5}}{2}}{\sqrt{5}} = \frac{\sqrt{5}}{\sqrt{5}} = 1$$

$$\text{fib}(2) = 1 = \frac{\left(\frac{1+\sqrt{5}}{2}\right)^2 - \left(\frac{1-\sqrt{5}}{2}\right)^2}{\sqrt{5}} = \frac{\left(\frac{1+2\sqrt{5}+5}{4}\right) - \left(\frac{1-2\sqrt{5}+5}{4}\right)}{\sqrt{5}} = \frac{\frac{\sqrt{5}}{4}}{\frac{\sqrt{5}}{4}} = 1$$

- Induktionsschritt: Sei $n \geq 3$ beliebig. Wir nehmen an, dass die Binet-Formel für $n-1$ und $n-2$ gilt, und zeigen, dass sie dann auch für n gilt:

$$\begin{aligned} \text{fib}(n) &= \text{fib}(n-1) + \text{fib}(n-2) \\ &= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-1} - \left(\frac{1-\sqrt{5}}{2}\right)^{n-1}}{\sqrt{5}} + \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-2} - \left(\frac{1-\sqrt{5}}{2}\right)^{n-2}}{\sqrt{5}} \\ &= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-1} - \left(\frac{1-\sqrt{5}}{2}\right)^{n-1} + \left(\frac{1+\sqrt{5}}{2}\right)^{n-2} - \left(\frac{1-\sqrt{5}}{2}\right)^{n-2}}{\sqrt{5}} \\ &= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-1} + \left(\frac{1+\sqrt{5}}{2}\right)^{n-2} - \left(\left(\frac{1-\sqrt{5}}{2}\right)^{n-1} + \left(\frac{1-\sqrt{5}}{2}\right)^{n-2}\right)}{\sqrt{5}} \\ &= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-2} \left(\left(\frac{1+\sqrt{5}}{2}\right) + 1\right) - \left(\left(\frac{1-\sqrt{5}}{2}\right)^{n-2} \left(\left(\frac{1-\sqrt{5}}{2}\right) + 1\right)\right)}{\sqrt{5}} \end{aligned}$$

$$\begin{aligned}
&= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{n-2} \left(\frac{1+\sqrt{5}}{2}\right)^2 - \left(\frac{1-\sqrt{5}}{2}\right)^{n-2} \left(\frac{1-\sqrt{5}}{2}\right)^2}{\sqrt{5}} \\
&= \frac{\left(\frac{1+\sqrt{5}}{2}\right)^n - \left(\frac{1-\sqrt{5}}{2}\right)^n}{\sqrt{5}}
\end{aligned}$$

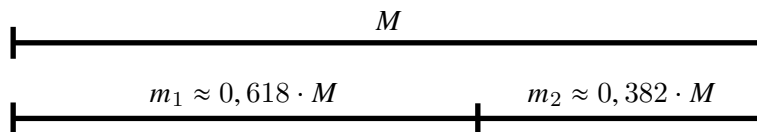
Dabei wurden die Gleichheiten $\left(\frac{1+\sqrt{5}}{2}\right) + 1 = \left(\frac{1+\sqrt{5}}{2}\right)^2$ und $\left(\frac{1-\sqrt{5}}{2}\right) + 1 = \left(\frac{1-\sqrt{5}}{2}\right)^2$ verwendet. Beide kann man nachrechnen:

$$\begin{aligned}
\left(\frac{1+\sqrt{5}}{2}\right)^2 &= \left(\frac{(1+\sqrt{5})^2}{4}\right) = \left(\frac{1+2\sqrt{5}+5}{4}\right) = \left(\frac{6+2\sqrt{5}}{4}\right) = \left(\frac{3+\sqrt{5}}{2}\right) = \left(\frac{1+\sqrt{5}}{2}\right) + 1 \\
\left(\frac{1-\sqrt{5}}{2}\right)^2 &= \left(\frac{(1-\sqrt{5})^2}{4}\right) = \left(\frac{1-2\sqrt{5}+5}{4}\right) = \left(\frac{6-2\sqrt{5}}{4}\right) = \left(\frac{3-\sqrt{5}}{2}\right) = \left(\frac{1-\sqrt{5}}{2}\right) + 1
\end{aligned}$$

□

Die Zahl $\phi = \frac{1+\sqrt{5}}{2} \approx 1,61$ wird auch als der *goldene Schnitt* bezeichnet. Man kann zeigen, dass $\text{fib}(n+1)/\text{fib}(n)$ mit wachsendem n sich immer mehr an ϕ nähert (man sagt die Folge konvergiert gegen ϕ).

Die schöne Eigenschaft von ϕ ist die folgende: Wenn man eine Strecke M in zwei Teilstücke m_1 und m_2 teilt, wobei $m_1 > m_2$ sodass gilt $m_1/m_2 = M/m_1$, dann ist dieses Verhältnis $= \phi$. Graphisch veranschaulicht



Dieses Verhältnis wird als besonders ästhetisch wahrgenommen und in der Kunst oder im Design zur Gestaltung oft verwendet.

5.6 Induktive Definitionen

Genau wie Beweise per Induktion kann man auch Strukturen per Induktion definieren. Dabei hat man einen Basisfall und einen induktiven Schritt, um aus bestehenden (kleineren Dingen) die nächsten größeren zu erzeugen.

Z.B. kann man die natürlichen Zahlen induktiv definieren durch:

Basis: 1 ist eine natürliche Zahl.

Schritt: Wenn n eine natürliche Zahl ist, dann ist auch der Nachfolger $n+1$ eine natürliche Zahl.

Arithmetische Ausdrücke (Ausdrücke die aus Zahlen und $+$, $-$, $\cdot$, $/$ aufgebaut sind) kann man induktiv definieren. Die kleinsten unteilbaren arithmetischen Ausdrücke sind Zahlen, daher bilden sie die Basis. Mithilfe der Operatoren kann man dann größere Ausdrücke bauen:

- Basis: Ganze Zahlen sind arithmetische Ausdrücke.
- Schritt: Wenn a, b arithmetische Ausdrücke sind, dann sind auch $(a + b)$, $(a \cdot b)$, $(a - b)$ und (a/b) arithmetische Ausdrücke.

Die aussagenlogischen Formeln (Definition 2.1.4) als auch die prädikatenlogischen Formeln (Definition 2.5.6) waren ebenfalls induktiv definiert.

Hat man eine solche induktive Definition einer Struktur, so kann man Eigenschaften der Instanzen mittels *struktureller Induktion* führen. Dabei prüft man die Fälle der induktiven Definitionen: Man zeigt, dass die Eigenschaft für die Basis gilt. Im Induktionsschritt nimmt man an, dass die Eigenschaft für die Teilstrukturen gilt und zeigt, dass sie dann auch für jede zusammengesetzte Struktur gilt.

Z.B. kann man zeigen:

Satz 5.6.1. *Für jeden arithmetische Ausdruck a gilt: Wenn n -Operatoren in a vorkommen, dann enthält a mindestens $n + 1$ Zahlen.*

Beweis. Induktion über die Struktur von a .

- Induktionsbasis: a ist eine Zahl. Dann enthält a keine Operatoren und $1 = 0 + 1$ Zahlen. Die Aussage stimmt also
- Induktionsschritt: Seien a, b arithmetische Ausdrücke mit n_a und n_b Operatoren. Wir nehmen als Induktionsannahme an, dass a $n_a + 1$ Zahlen und b $n_b + 1$ Zahlen enthält. Betrachte nun $(a \text{ op } b)$ mit $\text{op} \in \{+, -, \cdot, /\}$: Der Ausdruck enthält $n_a + n_b + 1$ Operatoren und (gemäß Annahme) $n_a + 1 + n_b + 1 = (n_a + n_b + 1) + 1$ Zahlen.

□

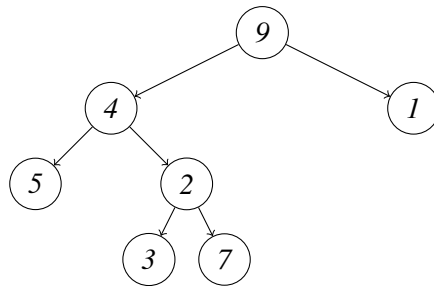
Als weiteres Beispiel betrachten wir (gerichtete) Binärbäume (als Vorgriff auf das nächste Kapitel, denn grundsätzlich sind Binärbäume spezielle Graphen, die wir erst im nächsten Kapitel kennenlernen). Im folgenden ist ein *Knoten* v ein Element aus einer beliebigen (endlichen) Menge V (der *Knotenmenge*). Ein Binärbaum besteht genau aus den Knoten in V (d.h. jeder Knoten $v \in V$ kommt genau einmal im Baum vor).

Definition 5.6.2. *Ein Binärbaum ist induktiv definiert durch:*

- Basis: Ein einzelner Knoten $v \in V$ ist ein Binärbaum (wir bezeichnen diesen Baum als Blatt).
- Schritt: Seien B_1 und B_2 Binärbäume, sodass die Knoten in B_1 und die Knoten in B_2 disjunkt sind und sei $v \in V$ ein Knoten, der nicht in B_1 oder B_2 vorkommt. Dann ist auch die Struktur bestehend aus v (genannt Wurzel) und B_1 (linker Teilbaum) und B_2 (rechter Teilbaum) auch ein Binärbaum.

Solche Binärbäume kann man leicht zeichnen, indem man die Wurzel oben zeichnet und linken und rechten Teilbaum dann links und rechts unter die Wurzel schreibt und mit Pfeilen verbindet (später werden wir diese Pfeile als *Kanten* eines gerichteten Graphen bezeichnen). Einzelne Knoten werden mit (beschrifteten) Kreisen gezeichnet.

Beispiel 5.6.3. Die folgende Zeichnung zeigt einen Binärbaum mit Knotenmenge $V = \{1, 2, 3, 5, 7, 9\}$:



Die *Höhe* eines Binärbaums ist induktiv definiert durch:

- Basis: Ein Blatt hat Höhe 1.
- Schritt: Ein Binärbaum mit Wurzel v und linkem Teilbaum B_1 und rechtem Teilbaum B_2 hat Höhe $1 + H$, wobei H das Maximum der Höhen von B_1 und B_2 ist.

Beispiel 5.6.4. Der Baum aus Beispiel 5.6.3 hat Höhe 4, denn:

- Das Blatt 3 und das Blatt 7 haben jeweils Höhe 1. Daher hat der Teilbaum mit Wurzel 2, die Höhe $1 + \max(1, 1) = 1 + 1 = 2$.
- Das Blatt 5 hat Höhe 1. Daher hat der Teilbaum mit Wurzel 4 die Höhe $1 + \max(1, 2) = 1 + 2 = 3$.
- Das Blatt 1 hat Höhe 1. Daher hat der Baum mit Wurzel 9 die Höhe $1 + \max(3, 1) = 1 + 3 = 4$.

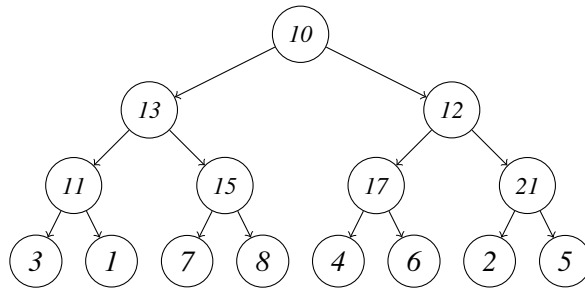
Die *Tiefe* eines Knotens v in einem Binärbaum B ist die Anzahl der Pfeile (d.h. Kanten) von der Wurzel bis zum Knoten v .

Beispiel 5.6.5. Die Tiefen aller Knoten des Baums aus Beispiel 5.6.3 sind

Knoten	Tiefe
1	1
2	2
3	3
5	2
7	3
9	0

Ein Binärbaum ist *vollständig*, wenn alle Blätter in derselben Tiefe sind.

Beispiel 5.6.6. Der Baum aus Beispiel 5.6.3 ist nicht vollständig, da z.B. Blätter 1 und 3 nicht in derselben Tiefe sind. Der folgende Baum ist ein vollständiger Binärbaum:



Satz 5.6.7. Ein vollständiger Binärbaum B mit Höhe h hat $2^h - 1$ Knoten.

Beweis. Beweis durch strukturelle Induktion über B .

- Basis: B ist ein einzelnes Blatt. Dann ist $h = 1$. Da $2^1 - 1 = 1$, gilt die Behauptung.
- Induktionsschritt: Nehme als Induktionshypothese an, dass die Aussage für linken und rechten Teilbaum von B gilt. Seien B_1 und B_2 die beiden Teilbäume von B mit Höhen h_1 und h_2 . Die Induktionshypothese zeigt daher, dass B_1 $2^{h_1} - 1$ und B_2 $2^{h_2} - 1$ Knoten haben. Damit hat B $1 + 2^{h_1} - 1 + 2^{h_2} - 1 = 2^{h_1} + 2^{h_2} - 1$ Knoten. Da die Blätter von B_1 und B_2 in gleicher Tiefe liegen müssen (sonst wäre B kein vollständiger Binärbaum), gilt $h_1 = h_2$. Damit folgt $h = 1 + \max(h_1, h_2) = 1 + h_1 = 1 + h_2$ und B hat $2^{h_1} + 2^{h_2} - 1 = 2^{h-1} + 2^{h-1} - 1 = 2 \cdot 2^{h-1} - 1 = 2^h - 1$ Knoten. $\square$

5.7 Schlussbemerkungen

Beweistechniken insbesondere verschiedene Schemata des Induktionsbeweises werden gut in (OE17) erläutert. Auch das englische Buch (Cum21) ist eine sehr empfehlenswerte Lektüre, um Beweistechniken, induktive Beweise und häufige Fehler beim Beweisen kennenzulernen. Kapitel über das Beweisen, die Teile unseres Kapitels abdecken (und andere nicht behandelte Themen) sind z.B. in (Ber24; MM24; BZ14) zu finden.

Kapitel zu induktiven Definitionen und dazugehörigen induktiven Beweisen finden sich z.B. in (KS07; MM24).

In (BZ14) wird das Schubfachprinzip in einem Kapitel behandelt.

6 Grundlagen der Graphentheorie

6.1 Einführendes

In diesem Kapitel lernen wir Graphen kennen. Durch Graphen können viele Probleme *modelliert* werden: Dabei *abstrahiert* man vom konkreten Problem und stellt dieses als Graph dar. Anschließend löst man das Problem als Fragestellung auf den Graphen und schließlich transformiert man die Lösung zurück, um das ursprüngliche Problem zu lösen.

Das wesentliche Lernziel des Kapitels ist es, die Grundbegriffe für Graphen und deren Eigenschaften zu studieren.

6.2 Graphen und grundlegende Begriffe

Definition 6.2.1 (Graph). Ein (ungerichteter) Graph G ist ein Paar $G = (V, E)$, wobei

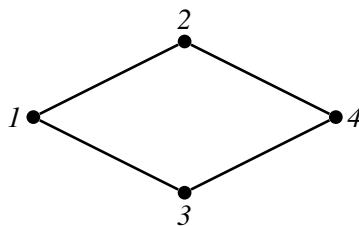
- V eine Menge von Knoten (engl. „vertexes“) ist, manchmal auch Ecken genannt, und
- $E \subseteq \{\{u, v\} \mid u \in V, v \in V\}$ eine Menge von Kanten (engl. „edges“) ist.

Wenn V endlich ist, spricht man von einem endlichen Graph.

Wir behandeln im wesentlichen nur endliche Graphen. Eine Kante drückt eine Verbindung zwischen zwei Knoten aus. Wenn zwei Knoten durch eine Kante verbunden sind, so nennt man sie *adjazent* oder *benachbart*.

Graphisch werden Knoten als Punkte (oder kleine Kreise) und Kanten als Verbindungen zwischen den Punkten gezeichnet. Manchmal schreibt man die Knotennamen mit an die Knoten.

Beispiel 6.2.2. Den Graph $G = (\{1, 2, 3, 4\}, \{\{1, 2\}, \{1, 3\}, \{2, 4\}, \{3, 4\}\})$ kann man zeichnen als



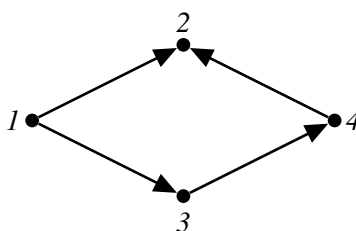
Neben ungerichteten Graphen gibt es auch gerichtete Graphen. Dabei haben die Kanten eine Richtung und sie sind daher Relationen auf der Knotenmenge E .

Definition 6.2.3 (Gerichteter Graph). Ein gerichteter Graph G ist ein Paar $G = (V, E)$, wobei

- V eine Menge von Knoten ist und
- $E \subseteq V \times V$ eine Menge von gerichteten Kanten ist.

Bei gerichteten Graphen werden Kanten als Pfeile gezeichnet.

Beispiel 6.2.4. Den gerichteten Graph $G = (\{1, 2, 3, 4\}, \{(1, 2), (1, 3), (2, 4), (3, 4)\})$ kann man zeichnen als



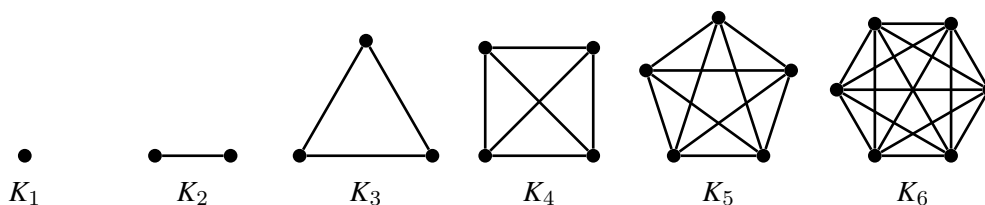
Man kann ungerichtete Graphen durch gerichtete Graphen repräsentieren, indem man Kanten in beiden Richtungen einfügt, bzw. fordert, dass die Kantenrelation symmetrisch ist. Wir bleiben jedoch bei der Definition mit zweielementigen Mengen.

Wir gehen auf einige weitere Varianten von Graphen ein. Ein *Multigraph*, erlaubt mehrere Kanten zwischen denselben Knoten (man spricht dann auch von Mehrfachkanten), die Kanten sind in diesem Fall keine Menge, sondern eine sogenannte *Multimenge*, diese kann man sich als Menge von Paaren (e, i) vorstellen, wobei $i \in \mathbb{N}$, angibt, wie oft das Element e in der Multimenge vorkommt. Graphen mit *Schlingen* sind Graphen, die auch Schlingen erlauben: Eine *Schlinge* ist eine Kante von und zum selben Knoten: Unsere Definition von gerichteten Graphen erlaubt Schlingen, während in unserer Definition von ungerichteten Graphen Schlingen nicht erlaubt sind (man kann die Definition erweitern, um Schlingen auch dort zu erlauben, wir benötigen diese jedoch nicht.) Einen Graph ohne Schlingen nennt man auch *schlingenfrei*. Ein Graph ohne Schlingen und ohne Mehrfachkanten nennt man auch *schlicht*.

Wir betrachten im wesentlichen schlichte, ungerichtete Graphen und machen deutlich, wenn wir andere Graphen betrachten.

Für einen Graph mit n Knoten, ist der *Nullgraph* der Graph mit leerer Kantenmenge und der *vollständige Graph*, der Graph mit maximaler Kantenanzahl, d.h. die Kantenmenge ist $E = \{\{u, v\} \mid u, v \in V\}$. Der vollständige Graph mit n Knoten wird aus als K_n bezeichnet.

Beispiel 6.2.5. Die Graphen K_1, K_2, K_3, K_4, K_5 und K_6 kann man zeichnen als:



Satz 6.2.6. Der vollständige Graph K_n hat $n \cdot (n - 1)/2$ Kanten.

Beweis. Wir verwenden vollständige Induktion über die Anzahl der Knoten n .

- Induktionsbasis: K_1 hat einen Knoten und 0 Kanten. Daher stimmt die Aussage denn $1 \cdot 0/2 = 0$
- Induktionsschritt: Wir nehmen an, dass K_n genau $n \cdot (n - 1)/2$ Kanten hat und betrachten K_{n+1} . Betrachte einen beliebigen Knoten v aus K_{n+1} . Dieser hat n benachbarte Knoten. Entferne Knoten v und alle Kanten zwischen v und seinen Nachbarn. Der entstehende Graph ist K_n und hat nach Induktionshypothese $n \cdot (n - 1)/2$ Kanten. Addieren wir die n Kanten, so zeigt dies, dass K_{n+1} $n + n \cdot (n - 1)/2 = (2n + n \cdot (n - 1))/2 = ((2 + n - 1) \cdot n)/2 = (n + 1) \cdot n/2$ Kanten hat. $\square$

Definition 6.2.7. Seien $G_1 = (V_1, E_1)$ und $G_2 = (V_2, E_2)$ Graphen.

- G_2 ist ein Teilgraph von G_1 , wenn gilt $V_2 \subseteq V_1$ und $E_2 \subseteq E_1$
- $G_1 \cup G_2 = (V_1 \cup V_2, E_1 \cup E_2)$ ist die Vereinigung von G_1 und G_2 .
- $G_1 \cap G_2 = (V_1 \cap V_2, E_1 \cap E_2)$ ist der Schnitt von G_1 und G_2 .

Für Graph $G = (V, E)$ ist der Komplementgraph $\bar{G}$ von G definiert als $\bar{G} = (V, \bar{E})$ mit

$$\bar{E} = \{\{u, v\} \mid u, v \in V, u \neq v, \{u, v\} \notin E\}.$$

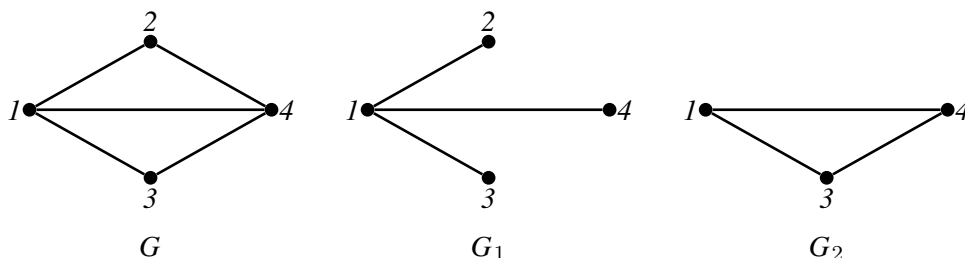
Für Graph $G = (V, E)$ und Knotenmenge $V' \subseteq V$ ist

$$G' = (V', \{\{u, v\} \in E \mid u \in V', v \in V'\})$$

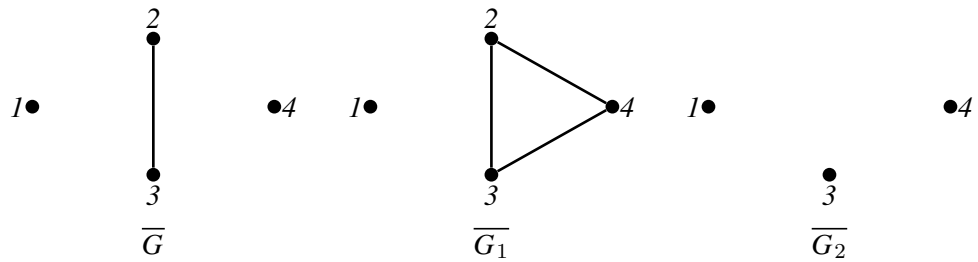
der von V' induzierte Teilgraph von G

Der Komplementgraph von G entsteht, indem man alle fehlenden Kanten hinzufügt und die alten Kanten entfernt.

Beispiel 6.2.8. Für $G = (\{1, 2, 3, 4\}, \{\{1, 2\}, \{1, 3\}, \{1, 4\}, \{2, 4\}, \{3, 4\}\})$ ist der Graph $G_1 = (\{1, 2, 3, 4\}, \{\{1, 2\}, \{1, 3\}, \{1, 4\}\})$ ein Teilgraph von G und der von $\{1, 3, 4\}$ induzierte Teilgraph von G ist $G_2 = (\{1, 3, 4\}, \{\{1, 3\}, \{1, 4\}, \{3, 4\}\})$.



Die Komplementgraphen von G , G_1 , G_2 sind:

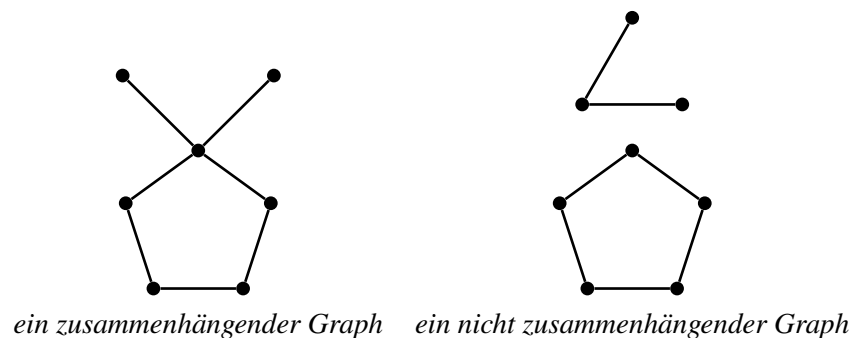


Definition 6.2.9. In einem Graph $G = (V, E)$ ist ein Weg (oder Pfad) von Knoten v_1 nach Knoten v_k eine Menge von Kanten $\{\{v_1, v_2\}, \{v_2, v_3\}, \dots, \{v_{k-1}, v_k\}\} \subseteq E$. Wenn $v_1 = v_k$, dann ist der Weg ein Kreis. Wenn G einen Kreis enthält, so heißt G zyklisch, anderenfalls nennen wir G kreisfrei oder azyklisch.

Graph G ist zusammenhängend, wenn es für alle $u, v \in V$ mit $u \neq v$ einen Weg von u nach v gibt.

Ist ein Graph nicht zusammenhängend, so zerfällt er in mehrere Teile (diese nennt man auch Zusammenhangskomponenten).

Beispiel 6.2.10. Wir zeigen zwei Beispiele:



ein zusammenhängender Graph ein nicht zusammenhängender Graph

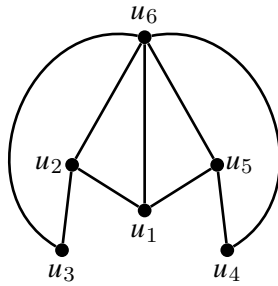
Für einen Graph $G = (V, E)$ ist der Grad eines Knotens $u \in V$, geschrieben $\deg_G(u)$, die Anzahl der Kanten, die von diesem Knoten ausgehen:

$$\deg_G(u) := |\{v \mid \{u, v\} \in E\}|$$

Wenn $\deg_G(u) = 0$, so nennt man u isoliert. Für gerichtete Graphen kann man den Ausgrad (Anzahl ausgehender Kanten) und den Ingrad (Anzahl eingehender Kanten) eines Knotens unterscheiden.

Beispiel 6.2.11. In einem vollständigen Graphen K_n ist der Grad jedes Knotens $\deg_{K_n}(u) = n-1$, da jeder Knoten u mit allen anderen $n-1$ Knoten verbunden ist.

Beispiel 6.2.12. Für den links gezeigten Graphen G , ist rechts der Grad jedes Knotens rechts angegeben:



$$\begin{aligned}\deg_G(u_1) &= 3 \\ \deg_G(u_2) &= 3 \\ \deg_G(u_3) &= 2 \\ \deg_G(u_4) &= 2 \\ \deg_G(u_5) &= 3 \\ \deg_G(u_6) &= 5\end{aligned}$$

6.2.1 Gleichheit und Isomorphie

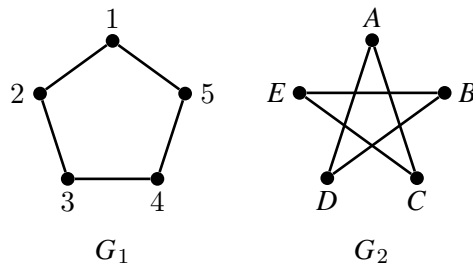
Bisher haben wir oft Graphen gezeichnet, ohne die Knoten und Kantenmenge explizit anzugeben. Wenn man die Knoten jedoch durch unterscheidbare Objekte darstellt (und nicht durch schwarze Kreise), dann stellt man fest, dass es auf die Namen der Objekte ankommt. Daher unterscheidet man Gleichheit und Isomorphie von Graphen:

Definition 6.2.13. Zwei Graphen $G_1 = (V_1, E_1)$ und $G_2 = (V_2, E_2)$ sind gleich ($G_1 = G_2$), wenn gilt $V_1 = V_2$ und $E_1 = E_2$. Sie sind isomorph (geschrieben $G_1 \cong G_2$), wenn es eine Bijektion $f : V_1 \rightarrow V_2$ gibt, sodass für alle Knoten $u, v \in V_1$ gilt: $\{u, v\} \in E_1$ genau dann, wenn $\{f(u), f(v)\} \in E_2$.

Beispiel 6.2.14. Die beiden Graphen

$$\begin{aligned}G_1 &= (V_1, E_1) = (\{1, 2, 3, 4, 5\}, \{\{1, 2\}, \{2, 3\}, \{3, 4\}, \{4, 5\}, \{5, 1\}\}) \text{ und} \\ G_2 &= (V_2, E_2) = (\{A, B, C, D, E\}, \{\{A, C\}, \{C, E\}, \{B, E\}, \{B, D\}, \{A, D\}\})\end{aligned}$$

sind isomorph.



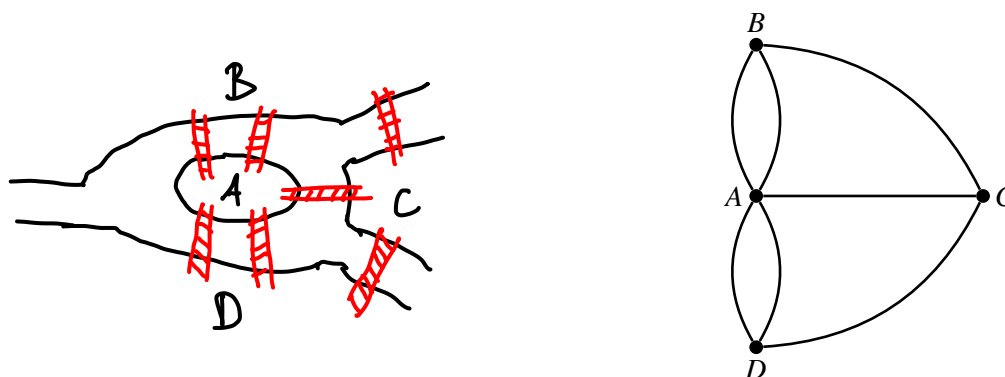
Für die Bijektion $f = \{1 \mapsto A, 2 \mapsto D, 3 \mapsto B, 4 \mapsto E, 5 \mapsto C\}$ gilt $\{u, v\} \in E_1$ genau dann, wenn $\{f(u), f(v)\} \in E_2$.

Übungsaufgabe 6.2.15. Zeige: Die Relation $\cong$ auf Graphen ist eine Äquivalenzrelation.

6.3 Eulergraphen

Das von Leonhard Euler im 18. Jahrhundert untersuchte *Königsberger Brückenproblem* bestand darin: Im damalig Königsberg vereinigten sich Alter Pregel und Neuer Pregel zum Fluss Pregel.

Insgesamt gab es nur sieben Brücken über die Flüsse. Euler fragte sich, ob es einen Rundweg gibt, der alle Brücken genau einmal überquert und dort endet, wo man begonnen hat. Ein Skizze zur Lage der Brücken ist links gezeigt. Übersetzt man die Landkarte in einen Graphen, indem man die Landabschnitte A, B, C, D als Knoten und die Brücken als Kanten auffasst, so erhält man einen Multigraphen, der rechts gezeigt ist:



Eine *Eulertour* (auch Eulerkreis) in einem Graph (oder Multigraph) $G = (V, E)$ ist Kreis von einem Knoten u zum selben Knoten u , der jede Kante aus E genau einmal enthält. Ein Graph heißt *eulersch*, wenn er eine Eulertour enthält.

Das Königsberger Brückenproblem wäre genau dann positiv lösbar, wenn der oben gezeigte Graph eulersch ist.

Es gilt jedoch der Satz von Euler:

Satz 6.3.1. *Wenn ein zusammenhängender Graph (oder Multigraph) eulersch ist, dann ist der Grad jedes Knoten gerade.*

Beweis. Sei $G = (V, E)$ ein eulerscher Graph und $v \in V$: Da G eine Eulertour hat und G zusammenhängend ist, durchläuft die Tour auch den Knoten v . Wenn v nicht der Anfangs- und Endknoten der Tour ist, dann werden bei jedem Durchlaufen von v zwei Kanten verwendet:

eine zum Hinein- und eine zum Herauslaufen: $\overset{v}{\text{-----}} \rightarrow$ oder $\leftarrow \text{-----} \overset{v}{}$. Daher folgt $\deg_G(v) = 2 \cdot k$, wobei k die Anzahl an Durchläufen durch v ist.

Für den Anfangs- und Endknoten folgt dasselbe für die Durchläufe, wobei noch für den Beginn eine Kante und für das Ende eine Kante besucht wird, was aber auch eine gerade Anzahl ergibt. $\square$

Satz 6.3.1 zeigt, dass es in Königsberg den gesuchten Rundweg nicht gibt, da es mehrere Knoten mit ungeradem Grad gibt.

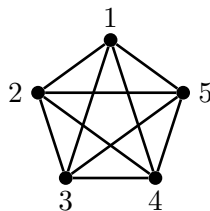
Auch die Umkehrung des Satzes gilt:

Satz 6.3.2. Wenn in einem zusammenhängenden Graph (oder Multigraph) alle Knoten geraden Grad haben, dann ist G eulersch.

Wir verzichten auf den genauen Beweis und skizzieren ihn nur: Man startet mit beliebigem Knoten und besucht von dort aus seine benachbarten Knoten u.s.w. Dabei werden die verwendeten Kanten als besucht markiert. Auf diese Weise findet man Kreise, die nicht notwendigerweise alle Kanten benutzen. Ist dies der Fall, so setzt man an den nicht verwendeten Kanten erneut an und sucht den nächsten Kreis, u.s.w. Danach werden die einzelnen Kreise zu einem Kreis verbunden, was immer möglich ist, da der Graph zusammenhängend ist.

Daraus folgt, dass die Graphen K_i mit ungeradem i eulersch sind.

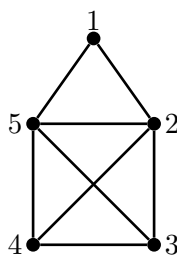
Beispiel 6.3.3. Eine Eulertour im K_5 besucht die Knoten in der Reihenfolge 1, 2, 3, 4, 5, 1, 3, 5, 2, 4, 1



Eine Verallgemeinerung der Eulertour ist ein *Eulerweg*: Dieser ist ein Weg von einem Knoten u zu einem Knoten v , der jede Kante genau einmal besucht. D.h. im Unterschied zur Eulertour, muss Anfangs- und Endknoten nicht derselbe sein.

Ein offener Eulerweg ist ein Eulerweg, der keine Eulertour ist, d.h. der Anfangs- und Endknoten sind im offenen Weg echt verschieden.

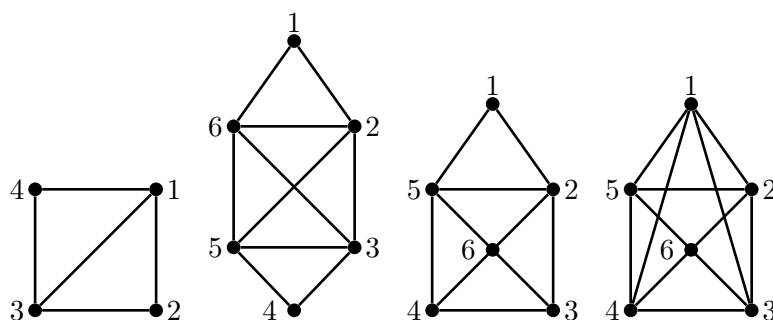
Beispiel 6.3.4. Das „Haus vom Nikolaus“ ist ein Graph mit einem offenen Eulerweg, z.B. wenn man die Knoten in der Reihenfolge 3, 4, 5, 1, 2, 5, 3, 2, 4 besucht:



Satz 6.3.5. Ein zusammenhängender Graph (bzw. Multigraph) hat genau dann einen offenen Eulerweg, wenn zwei Knoten ungeraden Grad und alle anderen Knoten geraden Grad haben. Der Eulerweg beginnt und endet in den Knoten mit ungeradem Grad.

- Beweis.*
- Sei $G = (V, E)$ ein Graph (oder Multigraph) mit $u \neq v \in G$ und $\deg_G(u)$, $\deg_G(v)$ sind ungerade, $\deg_G(w)$ ist gerade für alle $w \notin \{u, v\}$. Dann hat $G' = (V, E \cup \{u, v\})$ nur Knoten mit geradem Grad und daher eine Eulertour. Da dies ein Kreis ist, kann man ihn auch mit der Kante $\{u, v\}$ beginnen. Entferne diese Kante aus der Tour. Dann ergibt dies einen offenen Eulerweg, der in v beginnt und in u endet, oder umgekehrt.
 - Sei G ein Graph, der einen offenen Eulerweg hat. Dann müssen für jeden Knoten, der nicht Anfangs- oder Endknoten ist, beim Durchlaufen zwei Kanten verwendet werden – eine zum Hinein- und eine zum Herauslaufen. Daher haben diese Knoten alle geraden Grad. Für Anfangs- und Endknoten gilt dasselbe beim Durchlaufen, wobei jeweils noch eine Kante zum Herauslaufen (Startknoten) und zum Hineinlaufen (Endknoten) hinzukommt. Daher haben diese beiden Knoten einen ungeraden Grad. Mit der gleichen Argumentation ist klar, dass nur Knoten mit ungeradem Grad als Start- oder Endknoten verwendet werden können. $\square$

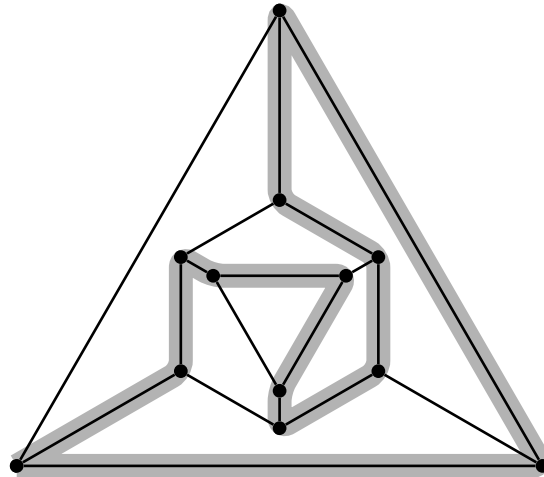
Übungsaufgabe 6.3.6. Welche der folgenden Graphen sind eulersch? Welche der Graphen haben einen offenen Eulerweg? Gebe die Knoten an, wie sie von einer Eulertour oder einem Eulerweg besucht werden.



Während Eulertouren alle Kanten eines Graphs besuchen, besucht ein Hamiltonkreis jeden Knoten:

Definition 6.3.7. Ein Hamiltonkreis in einem Graph $G = (V, E)$ ist ein Kreis von einem Knoten u zum selben Knoten u , der alle Knoten aus $V \setminus \{u\}$ genau einmal und u zweimal (am Anfang und am Ende) besucht.

Beispiel 6.3.8. Der grau hinterlegte Kreis ist ein Hamiltonkreis im folgenden Graphen:



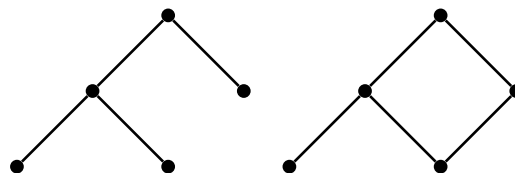
Die Frage, ob ein Graph einen Hamiltonkreis enthält, ist vermutlich viel schwieriger zu beantworten als die Frage nach einer Eulertour. Es existiert bis heute kein Programm, welches die Frage in annehmbarer Laufzeit für jeden Graph entscheiden kann. Das Problem ist eines der *NP-vollständigen* Probleme. Für solche Probleme gibt es bis heute keine effizienten Algorithmen, mit denen man alle Instanzen in annehmbarer Laufzeit entscheiden könnte. Vielmehr geht man davon aus, dass es ein solches Verfahren nicht gibt (das ist die berühmte $P = NP$ -Frage der Informatik, die bis heute unbeantwortet ist.).

Eine Anwendung für das Hamiltonkreis-Problem ist das Traveling-Salesman-Problem: Dabei fragt man nach der kürzesten Rundreise durch alle Städte einer gegebenen Landkarte, die jede Stadt einmal besucht. Dies ist eine Variante eines Hamiltonkreise, wobei Kanten zusätzlich Entfernungen (dies nennt man ein Kantengewicht in der Graphentheorie) haben.

6.4 Bäume

Ein Graph enthält einen Kreis, wenn es einen Weg von v_1 nach v_1 gibt. Ein zusammenhängender, azyklischer Graph wird *Baum* genannt. Ein azyklischer Graph, der nicht notwendigerweise zusammenhängend ist, wird auch *Wald* genannt, da er in Bäume zerfällt.

Beispiel 6.4.1. Der links gezeichnete Graph ist ein Baum, der rechts gezeichnete Graph ist zyklisch und daher kein Baum.



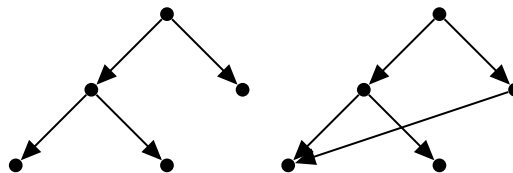
Satz 6.4.2. Jeder Baum mit $n \in \mathbb{N}$ Knoten hat $n - 1$ Kanten.

Beweis. Mit vollständiger Induktion über die Anzahl an Knoten.

- Induktionsbasis $n = 1$: Ein Baum mit 1 Knoten hat keine (und daher $n - 1 = 0$) Kanten.
- Induktionsschritt: Sei $n \in \mathbb{N}$. Wir nehmen an, dass jeder Baum mit n Knoten genau $n - 1$ Kanten hat. Betrachte nun einen Baum mit $n+1$ Knoten: Wähle einen beliebigen Knoten mit Grad 1. Diesen muss es geben, da der Graph zusammenhängend und kreisfrei ist. Entferne diesen Knoten mitsamt seiner Kante. Dann ist der Graph immer noch zusammenhängend und kreisfrei, d.h. ein Baum. Dieser Baum hat n Knoten und nach Induktionshypothese $n - 1$ Kanten. Da wir nur eine Kante entfernt haben, hatte der gegebene Baum $n = (n+1) - 1$ Kanten. $\square$

Für *gerichtete* Graphen gibt es auch eine Definition für Bäume: Ein (gerichteter) Baum ist ein zusammenhängender gerichteter Graph, der azyklisch ist (auch DAG genannt, von engl. *directed acyclic graph*), sodass der ungerichtete Graph, der durch Weglassen der Richtung entsteht, ein Baum ist.

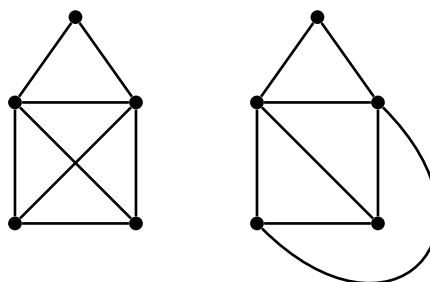
Beispiel 6.4.3. Der links gezeichnete gerichtete Graph ist ein Baum, der rechts gezeichnete gerichtete Graph ist ein DAG, aber kein Baum, da er nach Weglassen der Richtung ein zyklischer ungerichteter Graph ist.



6.5 Planare Graphen

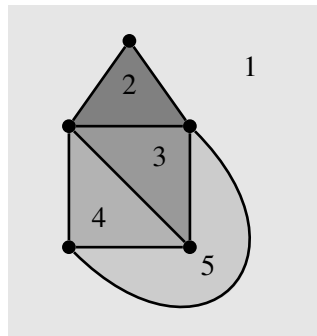
Ein Graph ist *planar*, falls er in der Ebene gezeichnet werden kann, ohne dass sich Kanten überschneiden.

Beispiel 6.5.1. Das Haus vom Nikolaus ist ein planarer Graph, da er ohne Überschneidungen gezeichnet werden kann:



Eine Anwendung von planaren Graphen in der Informatik ist der digitale Schaltungsentwurf: Dort identifiziert man die Bauteile mit Knoten und die Kanten stellen die Leiterbahnen zwischen den Verbindungen dar. Man möchte die Kontakte verbinden, ohne dass sich die Leitungen kreuzen. D.h. man sucht eine planare Einbettung des Graphen in die Ebene.

Ein planar gezeichneter Graph zerlegt die Ebene in Gebiete. Jedes Gebiet wird durch einen Kreis umschlossen, außer das (immer existierende) Außengebiet. Obige planare Zeichnung des Haus vom Nikolaus zerlegt die Ebene in 5 Gebiete:



Für planare Graphen gilt die folgende Eulersche Polyederformel:

Satz 6.5.2 (Eulersche Polyederformel). *Sei G ein zusammenhängender planarer Graph mit n Knoten, m Kanten und g Gebieten. Dann gilt $g = m - n + 2$.*

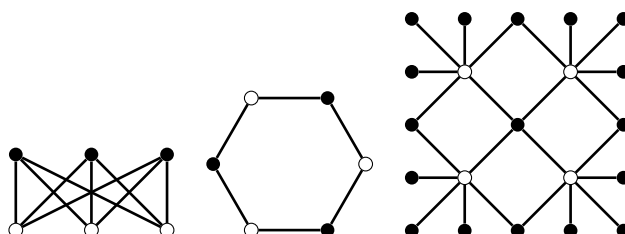
Beweis. Wir verwenden vollständige Induktion nach g .

- Induktionsbasis $g = 1$: Wenn $g = 1$, dann hat der Graph G keine Kreise und da er zusammenhängend ist, ist er ein Baum. Dann gilt mit Satz 6.4.2 $m = n - 1$ und daher $m - n + 2 = n - 1 - n + 2 = 1$.
- Induktionsschritt: Wir nehmen an, dass die Behauptung für alle Graphen mit g Gebieten gilt. Sei G ein Graph mit $g + 1$ Gebieten ($g \geq 1$), m Kanten und n Knoten. Dann muss G einen Kreis enthalten. Entferne eine Kante eines Kreises aus G . Sei dies der Graph G' . Graph G' hat g Gebiete, genauso viele Knoten und eine Kante weniger als G . Nach Induktionshypothese gilt für diesen Graph $g = (m - 1) - n + 2$. Damit gilt für G : $g + 1 = m - n + 1$ und damit folgt die Aussage. $\square$

6.6 Bipartite Graphen

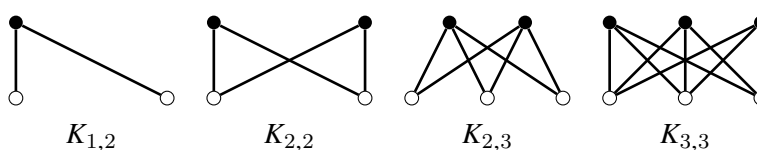
Ein Graph heißt *bipartit*, wenn man allen Knoten eine Farbe aus {schwarz, weiß} so zuordnen kann, dass Kanten nur schwarze mit weißen Knoten verbinden. Diese Zuordnung partitioniert die Knoten des Graphs eine Partition bestehend aus zwei Knotenmengen (die schwarzen und die weißen Knoten).

Beispiel 6.6.1. Beispiele für bipartite Graphen sind:



Ein bipartiter Graph mit Bipartition $\{V_1, V_2\}$ ist *vollständig bipartit*, wenn jeder Knoten in V_1 mit jedem Knoten in V_2 durch eine Kante verbunden ist. Wenn $|V_1| = m$ und $|V_2| = n$, dann bezeichnet man den vollständig bipartiten Graph mit $K_{m,n}$:

Beispiele sind:



Lemma 6.6.2. Jeder Kreis in einem bipartiten Graph hat eine gerade Anzahl an Kanten.

Beweis. Sei $\{v_1, v_2\}, \{v_2, v_3\} \dots, \{v_n, v_1\}$ ein Kreis. Die Knoten müssen abwechselnd unterschiedlich gefärbt sein, was sofort zeigt $\{v_1, v_3, \dots\}$ sind gleich gefärbt und $\{v_2, v_4, \dots\}$ sind gleich gefärbt. Wenn n ungerade wäre, dann hätten v_1 und v_n die gleiche Farbe, was nicht möglich ist, da sie über eine Kante verbunden sind. $\square$

Beispiel 6.6.3. Drei Häuser möchten alle eine Verbindung mit dem Gaswerk, Elektrizitätswerk und ein und demselben Internetanbieter. Kann man dies so machen, dass sich die Leitungen nicht kreuzen? Die Frage ist äquivalent zur Frage, ob der Graph $K_{3,3}$ planar ist. Der nächste Satz zeigt, dass dies nicht zutrifft.

Satz 6.6.4. $K_{3,3}$ ist nicht planar.

Beweis. Beweis durch Widerspruch. Nehme an, $K_{3,3}$ ist planar. Betrachte die Einbettung in die Ebene. Da es keine Kreise mit ungerader Anzahl an Kanten geben kann (siehe Lemma 6.6.2), muss jedes Gebiet mit mindestens 4 Kanten umgeben sein. Wenn wir alle diese Kanten zusammenzählen, ergibt dies $4g$. Dabei können wir jede Kante höchstens zweimal gezählt haben, denn sie gehört zu maximal zwei Gebieten (entweder trennt sie zwei Gebiete, oder sie liegt im selben Gebiet). Das ergibt die Ungleichung $4g \leq 2m$, oder umgeformt $2g \leq m$. Da $K_{3,3}$ 9 Kanten hat, kann daher nur $g \leq 4$ gelten. Die Eulersche Polyederformel liefert hingegen für $n = 6$ und $m = 9$: $g = m - n + 2 = 9 - 6 + 2 = 5$. Das ist ein Widerspruch! Die Annahme war falsch, $K_{3,3}$ ist nicht planar. $\square$

Auch der vollständige Graph K_5 ist nicht planar (wir verzichten auf den Beweis):

Satz 6.6.5. *Der vollständige Graph K_5 und der vollständig bipartite Graph $K_{3,3}$ sind nicht planar.*

Die Graphen K_5 und $K_{3,3}$ spielen für die Planarität eine besondere Rolle, denn nach Sätzen von Kuratowski und von Wagner folgt, dass jeder nichtplanare Graph K_5 oder $K_{3,3}$ als Minor hat. Ein Minor entsteht, indem man Knoten und Kanten entfernt und durch Kantenkontraktion verschmilzt (dabei entfernt man eine Kante und verschmilzt die dazugehörigen Knoten).

6.7 Färbungen

Wir haben bei bipartiten Graphen schon eine Knotenfärbung (mit 2 Farben) gesehen. Man kann dieses Problem auf k Farben verallgemeinern und sich z.B. Fragen, wie viele Farben man benötigt, um einen Graphen zu färben, ohne adjazente Knoten gleich zu färben. Allgemein ist das Färbbarkeitsproblem ein schwieriges (NP-vollständiges) Problem.

Verwandt (aber einfacher) ist das Vierfarbenproblem: Die Frage dabei ist, ob man die Länder einer Landkarte stets mit nicht mehr als 4 Farben färben kann, ohne dass benachbarte Länder die gleiche Farbe haben. Übersetzt in die Graphentheorie kann man die Länder als Knoten auffassen und die Kanten stellen Nachbarschaftsbeziehungen der Länder dar. D.h. die Frage ist, ob man jeden planaren Graphen mit 4 Farben färben kann.

Der berühmte Vierfarbensatz besagt, dass tatsächlich vier Farben ausreichend sind. Der Satz erlangte einerseits Berühmtheit, da die Zeit von Problemformulierung bis zum Beweis gut 100 Jahre dauerte und andererseits war er das erste große mathematische Problem, das mithilfe von Computern gelöst wurde: Der Beweis reduziert die zu betrachteten Fälle auf unter 2000, die dann mit dem Computer überprüft wurden.

Anwendungen der Färbbarkeit sind z.B. Konfliktgraphen: Eine Kante repräsentiert einen Konflikt. Z.B. kann man damit eine Sitzordnung erstellen: Kanten repräsentieren, wer nicht mit wem am selben Tisch sitzen darf. Die Frage nach der Färbung ist dann: Wer sitzt an welchem Tisch, und wie viele Tische werden benötigt?

6.8 Repräsentation von Graphen

Man kann Graphen auf verschiedene Art und Weise repräsentieren, um dann Algorithmen auf der Repräsentation zu definieren oder auszuführen. Je nach Algorithmus kann die eine oder die andere Repräsentationsform vorteilhaft sein. Wir stellen einige gebräuchliche Repräsentationen vor.

6.8.1 Adjazenzmatrix

Bei der Repräsentation durch eine Adjazenzmatrix stellt man für einen Graph mit n Knoten $\{1, \dots, n\}$ die Kantenrelation als zweidimensionale Matrix M der Dimension $n \times n$ mit Einträgen

aus $\{0, 1\}$ dar. Ist der entsprechende Eintrag $M_{ij} = 1$, so existiert eine Kante zwischen Knoten i und j . Die Darstellung funktioniert für gerichtete und ungerichtete Graphen. Bei ungerichteten Graphen gilt $M_{ij} = M_{ji}$, d.h. die Matrix ist an der Diagonalen gespiegelt (die Matrix ist symmetrisch). Für Multigraphen kann man Einträge aus $\mathbb{N}_0$ anstelle von $\{0, 1\}$ verwenden. Im Rechner kann die Adjazenzmatrix leicht durch ein zweidimensionales Feld (Array) implementiert werden. Der Zugriff auf einzelne Inhalte ist normalerweise schnell (in konstanter Laufzeit) möglich. Der Platzbedarf für einen Graph ist jedoch stets quadratisch in der Anzahl der Knoten. Der Grad eines Knotens kann z.B. berechnet werden, indem man die entsprechende Zeile des Knotens in der Matrix aufsummiert.

6.8.2 Adjazenzlisten

Bei der Darstellung mit Adjazenzlisten speichert man für jeden Knoten eine Liste seiner benachbarten Knoten. Will man prüfen, ob eine Kante $\{i, j\}$ existiert, so sucht man die Liste einer der beiden Knoten (z.B. die Liste für i) und schließlich durchsucht man die Liste nach dem anderen Knoten j . Der Nachteil ist, dass dies länger dauern kann, als der gleiche Test mit einer Adjazenzmatrix. Der Grad eines Knotens i kann berechnet werden, indem man die Länge der Liste für Knoten i ausrechnet. Ein Vorteil der Darstellung durch Adjazenzlisten ist, dass man nur soviel Platz braucht, wie es Kanten im Graphen gibt.

6.8.3 Implizite Repräsentation

Hat man unendliche Graphen, so kann man diese nicht komplett im Rechner repräsentieren. Gleiches gilt, wenn die Graphen sehr groß sind (man denke z.B. an Landkarten und Routensuche auf diesen). Hier kann man eine implizite Repräsentation verwenden: Statt aller Knoten, speichert man nur eine Menge von Startknoten (das sind diejenigen, die für das jeweilige Problem interessant sind) und man stellt eine Nachfolgerfunktion zur Verfügung, die für jeden Knoten dessen Nachfolger berechnen kann. Diese kann z.B. durch eine Funktionsgleichung repräsentiert werden oder auch auf eine Datenbank zugreifen und dort nachschauen. Algorithmen auf solchen Repräsentationen schauen niemals den gesamten Graph auf einmal an, sondern untersuchen den Graph nur soweit, wie er benötigt wird.

6.9 Schlussbemerkungen

Die meisten Ausführungen in diesem Kapitel richten sich nach (BZ14). Weitere Einführungen zur Graphentheorie finden sich z.B. in Kapiteln in (TT13; Ber24; MM24; BZ14). Eine Monographie zur Graphentheorie ist (Die17), ältere Standardwerke zum Thema sind z.B. (Wil10; Tru93).

7 Algebraische Grundstrukturen

7.1 Einführendes

Die Algebra ist ein Teilgebiet der Mathematik. Sprachlich stammt das Wort Algebra vom arabischen Wort „al-ğabr“ ab, was als „Zusammenfügen gebrochener Teile“ übersetzt werden kann. Geschichtlich stammt die Algebra vom Lösen von Gleichungen mit Unbekannten ab. Wir beschäftigen uns jedoch mit der abstrakten Algebra, die Mengen mit Operationen und deren Eigenschaften betrachtet. Die Idee dabei ist Mengen und Operationen als Struktur aufzufassen, und diese zu benennen, wenn sie gemeinsame Eigenschaften haben. Wir werden dies für die sogenannten Gruppen und die sogenannten Körper in diesem Kapitel ausführen.

Hat man definiert, welche Eigenschaften (auch Axiome genannt) z.B. eine Gruppe ausmachen, kann man allgemein Eigenschaften für alle Gruppen nachweisen und sie für Instanzen der Gruppen verwenden.

Neben bekannten unendlichen Zahlenmengen und Operationen darauf, gehen wir insbesondere auf die modulare Arithmetik ein, was man als das Rechnen mit Resten bezeichnen könnte. Diese Strukturen sind interessant, da die zugrundeliegenden Mengen der Strukturen endlich sind.

Als Anwendung aus der Kryptographie werden wir die RSA-Verschlüsselung erläutern.

7.2 Gruppen

Sei M eine Menge. Eine binäre Verknüpfung $\circ : M \times M \rightarrow M$ bildet zwei Elemente aus M auf ein drittes Element ab.

Definition 7.2.1. Eine Gruppe $(G, \circ)$ besteht aus einer Menge G und einer Verknüpfung $\circ$, sodass die folgenden Gesetze (genannt Gruppenaxiome) erfüllt sind:

(G0) *Abgeschlossenheit:* Wenn $x, y \in G$, dann ist auch $x \circ y \in G$.

(G1) *Assoziativität:* Für alle $x, y, z \in G$: $(x \circ y) \circ z = x \circ (y \circ z)$.

(G2) *Existenz eines neutralen Elements:* Es gibt ein Element $e \in G$ für das gilt $e \circ x = x \circ e = x$ für alle $x \in G$.

(G3) *Existenz inverser Elemente:* Für jedes Element $x \in G$ gibt es ein Element $x^{-1} \in G$ mit $x \circ x^{-1} = x^{-1} \circ x = e$.

Wenn zusätzlich gilt:

(G4) *Kommutativität:* Für alle $x, y \in G$: $x \circ y = y \circ x$.

dann ist G eine abelsche Gruppe (oder auch kommutative Gruppe).

Manchmal nennt man das neutrale Element auch 1, manchmal auch 0 (dann wird das Inverse auch als $-x$ anstelle von x^{-1} bezeichnet.)

Satz 7.2.2. *Das Paar $(\mathbb{Z}, +)$, wobei $+$ die übliche Addition ist, ist eine abelsche Gruppe.*

Beweis. Wir prüfen die Gruppenaxiome.

(G0): Offensichtlich ist $\mathbb{Z}$ abgeschlossen bezüglich $+$.

(G1): $(x + y) + z = x + (y + z)$ gilt für die Addition auf ganzen Zahlen x, y und z .

(G2): Das neutrale Element $e = 0$ erfüllt das Gruppenaxiom, denn $0 + x = x = x + 0$ für alle $x \in \mathbb{Z}$.

(G3): Das inverse Element zu $x \in \mathbb{Z}$ ist: $x^{-1} = -x$, denn es gilt $x + (-x) = (-x) + x = 0$ für alle $x \in \mathbb{Z}$.

(G4): Für alle $x, y \in \mathbb{Z}$ gilt: $x + y = y + x$. □

Bemerkung 7.2.3. *Das Paar $(\mathbb{N}, +)$ ist keine Gruppe, da es z.B. kein neutrales Element gibt. Das Paar $(\mathbb{N}_0, +)$ ist ebenfalls keine Gruppe, da es z.B. kein inverses Element zu 1 gibt. Das Paar $(\mathbb{Z} \setminus \{0\}, \cdot)$ ist keine Gruppe, da es z.B. zu $3 \in \mathbb{Z}$ kein inverses Element in $\mathbb{Z}$ gibt. Das Paar $(\{-2, -1, 0, 1, 2\}, +)$ (mit $+$ ist normale Addition auf $\mathbb{Z}$) ist keine Gruppe, da $+$ nicht abgeschlossen für die Menge $\{-2, -1, 0, 1, 2\}$ ist.*

Satz 7.2.4. *Das Paar $(\mathbb{R} \setminus \{0\}, \cdot)$ wobei $\cdot$ die übliche Multiplikation ist, ist eine abelsche Gruppe.*

Beweis. Wir prüfen die Gruppenaxiome.

(G0): Offensichtlich ist $\mathbb{R} \setminus \{0\}$ abgeschlossen bezüglich der Multiplikation.

(G1): $(x \cdot y) \cdot z = x \cdot (y \cdot z)$ gilt für die Multiplikation auf reellen Zahlen.

(G2): Das neutrale Element $e = 1 \in \mathbb{R}$ erfüllt das Gruppenaxiom, denn $1 \cdot x = x = x \cdot 1$ für alle $x \in \mathbb{R} \setminus \{0\}$.

(G3): Das inverse Element zu $x \in \mathbb{R} \setminus \{0\}$ ist: $x^{-1} = \frac{1}{x}$, denn es gilt $x \cdot \frac{1}{x} = \frac{1}{x} \cdot x = 1$ für alle $x \in \mathbb{R} \setminus \{0\}$.

(G4:) Für alle $x, y \in \mathbb{R} \setminus \{0\}$ gilt: $x \cdot y = y \cdot x$. □

Bemerkung 7.2.5. *$(\mathbb{R}, \cdot)$ ist keine Gruppe, da es zu 0 kein inverses Element gibt (es gibt keine reelle Zahl x mit $0 \cdot x = 1$!).*

Wir untersuchen nun *endliche* Strukturen, d.h. endliche Mengen mit Verknüpfung darauf. Wie wir schon gesehen haben, taugen in diesem Fall für endliche Mengen von Ganzzahlen die normale Addition oder Multiplikation nicht, da sie nicht abgeschlossen bezüglich der Verknüpfung sind.

Definition 7.2.6. *Wir definieren die Menge $\mathbb{Z}_n$ als Menge aller natürlichen Zahlen mit 0, die kleiner n sind:*

$$\mathbb{Z}_n := \{0, 1, \dots, n-1\}$$

Man nennt die Menge $\mathbb{Z}_n$ auch Restklassenmenge modulo n .

Die Addition $+$ auf $\mathbb{Z}_n$ sei die Addition modulo n , d.h. für $x, y \in \mathbb{Z}_n$ berechnen wir $(x + y) \bmod n$.

Zur Erinnerung $(x + y) \bmod n$ kann berechnet werden, indem man erst $x + y$ berechnet und anschließend mit Rest durch n teilt. Das Ergebnis von $(x + y) \bmod n$ ist dann der Rest. Beim Umformen verwenden wir auch die Kongruenz modulo n und schreiben z.B. $x \equiv y \bmod n$. Die gute Eigenschaft beim Rechnen modulo n ist, dass wir den Rest auch zwischen den Rechenschritten bilden dürfen, und damit die Zahlen „klein“ halten können.

Beispiel 7.2.7. Wir betrachten $(\mathbb{Z}_n, +)$ für $n = 3, 4, 5$. Die jeweilige Addition modulo n kann man tabellarisch auftragen:

+	0	1	2
0	0	1	2
1	1	2	0
2	2	0	1

Addition in $\mathbb{Z}_3$

+	0	1	2	3
0	0	1	2	3
1	1	2	3	0
2	2	3	0	1
3	3	0	1	2

Addition in $\mathbb{Z}_4$

+	0	1	2	3	4
0	0	1	2	3	4
1	1	2	3	4	0
2	2	3	4	0	1
3	3	4	0	1	2
4	4	0	1	2	3

Addition in $\mathbb{Z}_5$

Satz 7.2.8. Für alle $n \in \mathbb{N}$ gilt: $(\mathbb{Z}_n, +)$ ist eine Gruppe.

Beweis. Wir prüfen die Gruppenaxiome.

(G0): Da wir modulo n rechnen, ist der Rest zwischen 0 und $n - 1$ und damit in $\mathbb{Z}_n$.

(G1): Es gilt $(x + y) + z \equiv x + (y + z) \bmod n$, da $((a \bmod n) + b) \bmod n = (a + b) \bmod n$ für alle a, b, n .

(G2): Das Element $0 \in \mathbb{Z}_n$ erfüllt das Axiom, denn $x + 0 \equiv x \equiv 0 + x \bmod n$ für alle $x \in \mathbb{Z}_n$.

(G3): Das inverse Element zu $x \in \mathbb{Z}_n$ ist $x^{-1} = n - x$, denn es gilt

$$x + n - x \equiv n \equiv 0 \equiv n - x + x \bmod n. \quad \square$$

Bemerkung 7.2.9. Das Paar $(\mathbb{Z}_n, \cdot)$, wobei $\cdot$ die Multiplikation modulo n ist, ist keine Gruppe: Das Element $0 \in \mathbb{Z}_n$ hat kein multiplikativ inverses Element. Auch andere Elemente haben nicht immer ein multiplikatives Inverses: In $\mathbb{Z}_4$ hat z.B. 2 kein multiplikativ inverses Element, denn es gibt kein 2^{-1} mit $2 \cdot 2^{-1} \equiv 1 \bmod 4$: Wir probieren alle Zahlen aus $\mathbb{Z}_4$ durch: $2 \cdot 0 \equiv 0 \bmod 4$, $2 \cdot 1 \equiv 2 \bmod 4$, $2 \cdot 2 \equiv 0 \bmod 4$ und $2 \cdot 3 \equiv 2 \bmod 4$.

Definition 7.2.10. Der größte gemeinsame Teiler zweier Ganzzahlen x, y wird notiert als $\text{ggT}(x, y)$. Es ist die größte Zahl $a \in \mathbb{N}_0$, sodass es Zahlen $x', y' \in \mathbb{Z}$ gibt mit $x = a \cdot x'$ und $y = a \cdot y'$. Für $\text{ggT}(0, 0)$ definieren wir $\text{ggT}(0, 0) = 0$.

Zwei natürliche Zahlen $x, y \in \mathbb{N}_0$ sind teilerfremd, wenn gilt: $\text{ggT}(x, y) = 1$.

Beachte, dass gilt $\text{ggT}(0, x) = |x| = \text{ggT}(x, 0)$ für alle $x \in \mathbb{Z}$.

Definition 7.2.11. Für $n \in \mathbb{N}$ definieren wir die Menge

$$\mathbb{Z}_n^* = \{x \in \mathbb{Z}_n \mid x \text{ und } n \text{ sind teilerfremd}\}.$$

Beachte, dass $0 \notin \mathbb{Z}_n^*$ für $n > 1$ (da jede Zahl die 0 teilt). Z.B. $\mathbb{Z}_4^* = \{1, 3\}$ und $\mathbb{Z}_{15}^* = \{1, 2, 4, 7, 8, 11, 13, 14\}$. Wenn p eine Primzahl ist, dann ist $\mathbb{Z}_p^* = \mathbb{Z}_p \setminus \{0\}$.

Übungsaufgabe 7.2.12. Zähle jeweils die Elemente der Menge auf:

1. $\mathbb{Z}_7^*$
2. $\mathbb{Z}_{16}^*$
3. $\mathbb{Z}_{21}^*$.

Bemerkung 7.2.13. Die Struktur $(\mathbb{Z}_n^*, +)$ ist offensichtlich keine Gruppe, da das neutrale Element $0 \notin \mathbb{Z}_n^*$ fehlt. Die Frage, ob $(\mathbb{Z}_n^*, \circ)$ eine Gruppe ist, klären wir im folgenden.

Es gilt das folgende Lemma, dessen genauer Beweis in der Literatur gefunden werden kann. Wir werden später erläutern, wie man die Zahlen a und b für gegebenes x und y findet.

Lemma 7.2.14 (Lemma von Bézout). Seien x und y ganze Zahlen, und $d = \text{ggT}(x, y)$. Dann gibt es ganze Zahlen a, b mit $d = ax + by$.

Z.B. ist für $x = 15$ und $y = 6$ der $\text{ggT}(x, y) = \text{ggT}(15, 6) = 3$ und mit $a = 1$ und $b = -2$ ist das Lemma von Bézout für x und y erfüllt (denn $15 - 12 = 3$).

Für $x = 132$ und $y = 123$ ist $\text{ggT}(x, y) = \text{ggT}(132, 123) = 3$ und mit $a = 14$ und $b = -15$ ist das Lemma von Bézout für x und y erfüllt (denn $14 \cdot 132 - 15 \cdot 123 = 1848 - 1845 = 3$).

Satz 7.2.15. Das Paar $(\mathbb{Z}_n^*, \cdot)$, wobei $\cdot$ die Multiplikation modulo n ist, ist eine abelsche Gruppe.

Beweis. Wir prüfen die Gruppenaxiome:

- (G0): Die Multiplikation modulo n in $\mathbb{Z}_n^*$ ist abgeschlossen: Seien $x, y \in \mathbb{Z}_n^*$. Dann sind x und y teilerfremd zu n , damit ist auch das Produkt $x \cdot y$ teilerfremd zu n (ein Teiler von $x \cdot y$ würde auch x oder y teilen).
- (G1): Die Multiplikation ist assoziativ, da die Multiplikation assoziativ in $\mathbb{Z}_n$ ist.
- (G2): Das neutrale Element ist $1 \in \mathbb{Z}_n^*$, denn $1 \cdot x = x \cdot 1 = x$ für alle $x \in \mathbb{Z}_n$.
- (G3): Sei $x \in \mathbb{Z}_n^*$. Da x und n teilerfremd sind, gilt $\text{ggT}(x, n) = 1$ und es gibt Zahlen a, b mit $ax + bn = 1$ (Lemma 7.2.14). Damit gilt auch $(ax + bn) \bmod n = 1$ und auch $((a \bmod n)x + (bn \bmod n)) \bmod n = 1$ und (da bn mit Rest 0 durch n teilbar ist) $((a \bmod n)x) \bmod n = 1$. Setze $x^{-1} = a \bmod n$. Dann gilt $x \cdot x^{-1} \equiv 1 \bmod n$. Es bleibt zu zeigen $x^{-1} = (a \bmod n) \in \mathbb{Z}_n^*$: Aus $((a \bmod n)x) \bmod n = 1$ folgt: Es gibt ein k mit $((a \bmod n)x) - kn = 1$. D.h. $\text{ggT}(((a \bmod n)x), kn) = 1$, denn jeder Teiler von $((a \bmod n)x)$ und kn muss die 1 teilen. Damit folgt sofort, dass $a \bmod n$ und n teilerfremd sind.

(G4): Die Multiplikation ist kommutativ, da sie kommutativ in $\mathbb{Z}_n$ ist. $\square$

Es bleibt zu zeigen, wie man die Zahl a berechnet. Hierfür kann der erweiterte euklidische Algorithmus verwendet werden, den wir im folgenden erläutern.

7.2.1 Der erweiterte euklidische Algorithmus

Bevor wir den erweiterten Euklidischen Algorithmus erläutern, erklären wir den normalen Euklidischen Algorithmus. Dieser dient dazu den größten gemeinsamen Teiler zweier natürlicher Zahlen zu berechnen¹

Satz 7.2.16 (Euklidischer Algorithmus). *Seien $x, y \in \mathbb{N}_0$ und $y > 0$. Der $\text{ggT}(x, y)$ kann wie folgt rekursiv berechnet werden:*

1. Setze $X := x$ und $Y := y$.
2. Teile X durch Y mit einer Division mit Rest, d.h. berechne q und $0 \leq r < Y$ mit:

$$X = q \cdot Y + r$$

3. Wenn der Rest r gleich 0 ist, dann stoppe und liefere Y als Ergebnis für den $\text{ggT}(x, y)$. Anderenfalls, setze $X := Y$ und $Y := r$ und mache weiter mit Schritt (2).

Beweis. Wenn der Algorithmus sofort in Schritt 3 stoppt, ist $\text{ggT}(x, y) = y$ offensichtlich, denn x ist ohne Rest durch y teilbar! Anderenfalls muss man sich klar machen, dass aus $X = q \cdot Y + r$ folgt, dass der $\text{ggT}(X, Y)$ auch r teilen muss, um ein Teiler von X zu sein. Daher kann man mit $\text{ggT}(X, r)$ weiter machen.

Der Algorithmus stoppt nach endlichen vielen Iterationen, da Y in jedem Schritt kleiner wird², und Y niemals auf 0 (oder eine negative Zahl) gesetzt wird. $\square$

Beispiel 7.2.17. *Wir berechnen $\text{ggT}(75, 48)$ mit dem Euklidischen Algorithmus*

- Schritt 1: $X = 75, Y = 48$
- Schritt 2 ergibt $75 = 1 \cdot 48 + 27$, d.h. $q = 1$ und $r = 27$.
- Schritt 3 setzt $X = 48, Y = 27$ und mache weiter mit Schritt 2.
- Schritt 2 ergibt $48 = 1 \cdot 27 + 21$, d.h. $q = 1$ und $r = 21$.
- Schritt 3 setzt $X = 27, Y = 21$ und mache weiter mit Schritt 2.
- Schritt 2 ergibt $27 = 1 \cdot 21 + 6$, d.h. $q = 1$ und $r = 6$.
- Schritt 3 setzt $X = 21, Y = 6$ und mache weiter mit Schritt 2.
- Schritt 2 ergibt $21 = 3 \cdot 6 + 3$, d.h. $q = 3$ und $r = 3$.
- Schritt 3 setzt $X = 6, Y = 3$ und mache weiter mit Schritt 2.

¹Für ganze Zahlen kann man mit den Beträgen rechnen und später das Vorzeichen anpassen.

²Dies gilt auch im Fall $X < Y$: In diesem Fall, dreht der Algorithmus X und Y im nächsten Schritt um.

- Schritt 2 ergibt $6 = 2 \cdot 3 + 0$, d.h. $q = 2$ und $r = 0$.
- Schritt 3 liefert 3 als Ergebnis für $\text{ggT}(75, 48)$.

Um nun die Zahlen a und b aus dem Lemma von Bézout zu gewinnen, kann man die durch den Euklidischen Algorithmus gemachte Rechnung rückwärts betrachten und die erhaltenen Werte wieder einsetzen.

Die Idee ist folgende:

- Verwende die Gleichungen des Euklidischen Algorithmus $X = q \cdot Y + r$
- Nummeriere die Werte X, Y, q, r entsprechend der Iteration zu $X_i = q_i \cdot Y_i + r_i$, wobei mit $i = 2$ gestartet wird, d.h. $X_2 = q_2 \cdot Y_2 + r_2$ mit $X_2 = x, Y_2 = y$
- Löse die Gleichungen nach r_i auf. Das ergibt $r_i = X_i - q_i \cdot Y_i$.
- Das Setzen der X - und Y -Werte in Schritt 3 zeigt, dass im allgemeinen gilt:

$$Y_i = r_{i-1} \text{ und } X_i = Y_{i-1} = r_{i-2}.$$

Wenn wir noch definieren $r_0 = x$ und $r_1 = y$, dann lässt sich $r_i = X_i - q_i \cdot Y_i$ daher schreiben als

$$r_i = r_{i-2} - q_i \cdot r_{i-1}$$

- Sei $r_n = r_{n-2} - q_n \cdot r_{n-1}$ die vorletzte Gleichung (deren Rest r_n der gesuchte $\text{ggT}(x, y)$ ist, d.h.

$$\text{ggT}(x, y) = r_n = r_{n-2} - q_n \cdot r_{n-1}$$

- Setze sukzessive die Gleichungen für Reste r_i für $i = n - 1, \dots, 1$ ein, bis nur noch q_i 's und x und y vorkommen.
- Vereinfache die Gleichung bis sie in der richtigen Form $r_n = ax + by$ ist.

Wir betrachten Beispiel 7.2.17 und wenden die Nummerierung und Benennung an:

Gleichungen aus dem Euklidischen Algorithmus	nach r_i aufgelöst	entspricht
		$r_0 = x = 75$
		$r_1 = y = 48$
$75 = 1 \cdot 48 + 27$, d.h. $q = 1$ und $r = 27$	$27 = 75 - 1 \cdot 48$	$r_2 = r_0 - 1 \cdot r_1$
$48 = 1 \cdot 27 + 21$, d.h. $q = 1$ und $r = 21$	$21 = 48 - 1 \cdot 27$	$r_3 = r_1 - 1 \cdot r_2$
$27 = 1 \cdot 21 + 6$, d.h. $q = 1$ und $r = 6$	$6 = 27 - 1 \cdot 21$	$r_4 = r_2 - 1 \cdot r_3$
$21 = 3 \cdot 6 + 3$, d.h. $q = 3$ und $r = 3$	$3 = 21 - 3 \cdot 6$	$r_5 = r_3 - 3 \cdot r_4$

Wenn man nun in $3 = r_5 = r_3 - 3 \cdot r_4$ sukzessive einsetzt und vereinfacht, so erhält man:

$$\begin{aligned}
 3 &= r_5 = r_3 - 3 \cdot r_4 \\
 &= r_3 - 3 \cdot (r_2 - 1 \cdot r_3) \\
 &= (r_1 - 1 \cdot (r_0 - 1 \cdot r_1)) - 3 \cdot ((r_0 - 1 \cdot r_1) - 1 \cdot (r_1 - 1 \cdot (r_0 - 1 \cdot r_1))) \\
 &= (48 - 1 \cdot (75 - 1 \cdot 48)) - 3 \cdot ((75 - 1 \cdot 48) - 1 \cdot (48 - 1 \cdot (75 - 1 \cdot 48))) \\
 &= (-7 \cdot 75) + (11 \cdot 48)
 \end{aligned}$$

Der erweiterte Euklidische Algorithmus macht genau diese Berechnung, nur er merkt sich die Faktoren direkt, sodass man nicht rückwärts rechnen muss.

Satz 7.2.18 (Erweiterter Euklidischer Algorithmus). *Seien $x, y \in \mathbb{N}$, dann können $a, b, \text{ggT}(x, y)$ mit $\text{ggT}(x, y) = a \cdot x + b \cdot y$ wie folgt berechnet werden:*

1. Setze

$$\begin{aligned} r_0 &= x, & r_1 &= y \\ s_0 &= 1, & s_1 &= 0 \\ t_0 &= 0, & t_1 &= 1 \end{aligned}$$

2. Setze $i = 2$

3. Berechne q_i als den ganzzahligen Anteil der Division von r_{i-2} durch r_{i-1}

4. Berechne

$$\begin{aligned} r_i &= r_{i-2} - q_i \cdot r_{i-1} \\ s_i &= s_{i-2} - q_i \cdot s_{i-1} \\ t_i &= t_{i-2} - q_i \cdot t_{i-1} \end{aligned}$$

5. Falls $r_i = 0$, dann stoppe mit $\text{ggT}(x, y) = r_{i-1} = s_{i-1} \cdot x + t_{i-1} \cdot y$, d.h. die Ausgabe ist $a = s_{i-1}$, $b = t_{i-1}$, $\text{ggT}(x, y) = r_{i-1}$

Anderenfalls, setze $i = i + 1$ und gehe zu 3.

Beispiel 7.2.19. Wir wenden den erweiterten Euklidischen Algorithmus für 75 und 48 an. Dabei gehen wir tabellarisch vor und schreiben $\lfloor m/n \rfloor$ für den ganzzahligen Anteil der Division von m durch n .

i		r_i	s_i	t_i
0		75	1	0
1		48	0	1
i	$q_i = \lfloor r_{i-2}/r_{i-1} \rfloor$	$r_i = r_{i-2} - q_i \cdot r_{i-1}$	$s_i = s_{i-2} - q_i \cdot s_{i-1}$	$t_i = t_{i-2} - q_i \cdot t_{i-1}$
2	$1 = \lfloor 75/48 \rfloor$	$27 = 75 - 1 \cdot 48$	$1 = 1 - 1 \cdot 0$	$-1 = 0 - 1 \cdot 1$
3	$1 = \lfloor 48/27 \rfloor$	$21 = 48 - 1 \cdot 27$	$-1 = 0 - 1 \cdot 1$	$2 = 1 - 1 \cdot (-1)$
4	$1 = \lfloor 27/21 \rfloor$	$6 = 27 - 1 \cdot 21$	$2 = 1 - 1 \cdot (-1)$	$-3 = -1 - 1 \cdot 2$
5	$3 = \lfloor 21/6 \rfloor$	$3 = 21 - 3 \cdot 6$	$-7 = -1 - 3 \cdot 2$	$11 = 2 - 3 \cdot (-3)$
6	$2 = \lfloor 6/3 \rfloor$	$0 = 6 - 2 \cdot 3$	Stop!	

Die Ausgabe ist daher:

$$\text{ggT}(75, 48) = 3 = -7 \cdot 75 + 11 \cdot 48$$

Als Schlussbemerkung führen wir die *Eulersche Funktion* ein, da wir sie im nächsten Abschnitt verwenden werden.

Definition 7.2.20. Die Eulersche Funktion $\phi : \mathbb{N} \rightarrow \mathbb{N}$ ist definiert als $\phi(n) = |\mathbb{Z}_n^*|$.

Für alle Primzahlen $p \neq q$ gilt $\phi(p \cdot q) = (p - 1) \cdot (q - 1)$.

Es gilt der folgende Satz (der Beweis kann in leicht in der Literatur gefunden werden)

Satz 7.2.21 (Satz von Euler). *Für alle $x \in \mathbb{Z}_n^*$ gilt:*

$$x^{\phi(n)} \equiv 1 \pmod{n}$$

7.2.2 Eine Anwendung aus der Kryptographie: Die RSA-Verschlüsselung

Bei der RSA-Verschlüsselung (benannt nach Ron Rivest, Adi Shamir und Leonard Adleman, (RSA78)) handelt es sich um ein asymmetrisches Verschlüsselungsverfahren. Dabei gibt es im Gegensatz zu symmetrischen Verfahren keinen gemeinsamen geheimen Schlüssel, den Sender und Empfänger einer Nachricht kennen müssen, sondern es gibt öffentliche und private Schlüssel. Der Empfänger gibt seinen öffentlichen Schlüssel öffentlich bekannt, der Sender verwendet diesen Schlüssel zum Verschlüsseln. Nur der Empfänger kennt seinen privaten Schlüssel, mit diesem ist es ihm möglich, die verschlüsselte Nachricht zu entschlüsseln.

Wir gehen im folgenden davon aus, dass die Nachricht als kleine Zahl kodiert ist (in Realität würde man eine Folge von Zahlen schicken, wobei jede Zahl einen Teil der Nachricht kodiert). Insbesondere sei die Nachricht x kleiner als die Zahl n im folgenden.

Sei Alice die Empfängerin und Bob der Sender der Nachricht. Die Nachricht sei x , die verschlüsselte Nachricht sei y .

Das RSA-Verfahren geht wie folgt vor:

- Alice wählt zwei große Primzahlen p und q und bildet $n = p \cdot q$ und Alice wählt eine Zahl a die teilerfremd zu $\phi(n) = (p - 1) \cdot (q - 1)$ ist. Sei b das multiplikative Inverse zu a modulo $\phi(n)$. Der öffentliche Schlüssel von Alice ist (n, a) der private Schlüssel ist $\phi(n)$ (und b , aber den kann Alice auch aus a und $\phi(n)$ berechnen)
- Zum *Verschlüsseln* berechnet Bob $y = x^a \pmod{n}$.
- Zum *Entschlüsseln* berechnet Alice $y^b \pmod{n}$.

Satz 7.2.22. *Das RSA-Verfahren ist korrekt, d.h. $(x^a)^b \equiv x \pmod{n}$.*

Beweis. Da b das multiplikative Inverse zu a ist (in $\mathbb{Z}_{\phi(n)}^*$), folgt $a \cdot b \equiv 1 \pmod{\phi(n)}$ und damit gibt es k mit $a \cdot b = k \cdot \phi(n) + 1$.

Jetzt lässt sich $y^b \pmod{n} = x$ durch Umformen zeigen:

$$\begin{aligned} y^b \pmod{n} &= (x^a \pmod{n})^b \pmod{n} = x^{a \cdot b} \pmod{n} = x^{k \cdot \phi(n) + 1} \pmod{n} \\ &= (((x^{\phi(n)}) \pmod{n})^k \pmod{n} \cdot x) \pmod{n} \stackrel{(1)}{=} ((1^k \pmod{n}) \cdot x) \pmod{n} = x \pmod{n} \end{aligned}$$

Dabei folgt Schritt (1) in fast allen Fällen mit dem Satz von Euler Satz 7.2.21 (nämlich für alle Fälle, dass n und x teilerfremd sind). Für alle anderen Fälle beweisen wir

$$x^{ab} \equiv x \pmod{n}$$

direkt: Es bleibt der Fall dass x und n nicht teilerfremd sind. Wir unterscheiden 3 Fälle:

1. x ist durch n teilbar: Dann gibt es k'' mit $x = n \cdot k''$ und daher

$$x^{ab} \equiv (n \cdot k'')^{ab} \equiv 0 \equiv x \pmod{n}$$

folgt direkt.

2. x ist durch q , aber nicht durch p teilbar. Da p prim, gilt $\phi(p) = p - 1$ und mit dem Satz von Euler gilt $x^{p-1} \equiv 1 \pmod{p}$. Da

$$1 = x^{p-1} \pmod{p} = (x^{(p-1)} \pmod{p})^{(q-1)} \pmod{p} = x^{(p-1) \cdot (q-1)} \pmod{p} = x^{\phi(n)} \pmod{p},$$

folgt $x^{\phi(n)} \equiv 1 \pmod{p}$.

D.h. es gibt ℓ mit $x^{p-1} = 1 + \ell \cdot p$. Damit können wir schließen

$$\begin{aligned} x^{a \cdot b} &\equiv x^{k \cdot \phi(n) + 1} \equiv x \cdot x^{(p-1)(q-1)k} \equiv x(1 + \ell \cdot p)^{(q-1)k} \\ &\stackrel{(2)}{\equiv} x(1 + p \cdot t) \equiv x + xpt \stackrel{(3)}{\equiv} x \pmod{n} \end{aligned}$$

wobei (2) folgt, da Ausmultiplizieren von $(1 + \ell \cdot p)^{(q-1)k}$ zeigt, dass jeder Summand $\ell \cdot p$ als Faktor enthält außer der erste (der nur Einsen multipliziert). Daher muss es ein t geben mit $(1 + \ell \cdot p)^{(q-1)k} = 1 + p \cdot t$. Die Umformung (3) folgt, da x nach Annahme durch q teilbar ist und damit xpt durch $n = p \cdot q$ teilbar ist.

3. x ist durch p aber nicht durch q teilbar: Der Beweis geht analog zum vorherigen, indem man p und q vertauscht. $\square$

Schließlich ist noch zu klären, ob das Verfahren sicher ist, oder der private Schlüssel leicht berechnet werden kann. Die Sicherheit ergibt sich daher, dass es bisher kein schnelles Verfahren gibt, um auf gebräuchlichen Rechnern aus der Zahl n die Zahl $\phi(n)$ bzw. die Primzahlen p und q zu berechnen (dies nennt man Faktorisierung von n). Da ein solches Verfahren bisher nicht entdeckt wurde und man p und q mit steigender Rechenkraft immer größer wählen kann, ist es schwierig den privaten Schlüssel zu knacken. Mit Quantencomputern ist die schnelle Faktorisierung theoretisch möglich, daher könnte es sein, dass RSA mit der Skalierung von Quantencomputern geknackt wird. Spätestens dann muss man auf andere Verschlüsselungsverfahren umstellen.

7.3 Körper

Ein Körper ist eine Struktur, die Verknüpfungen für Addition und Multiplikation hat und Gesetze garantiert, sodass man „normal“ in dieser Struktur rechnen kann.

Definition 7.3.1. Eine Menge K mit Verknüpfungen $+$ und $\circ$ ist ein Körper, wenn

1. $+$ und $\circ$ sind abgeschlossen auf K .
2. $+$ und $\circ$ sind assoziativ und kommutativ.

3. Es gibt ein neutrales Element 0 bezüglich $+$ und ein neutrales Element $1 \neq 0$ bezüglich $\circ$.
4. Jedes Element $x \in K$ hat ein additives Inverses $-x \in K$ und jedes Element $x \neq 0$ hat ein multiplikatives Inverses $x^{-1} \in K$.
5. Das Distributivgesetz gilt: $x \circ (y + z) = x \circ y + x \circ z$ und $(x + y) \circ z = (x \circ z) + (y \circ z)$ für alle $x, y, z \in K$

Alternativ kann man auch sagen: $(K, +, \circ)$ ist ein Körper, wenn $(K, +)$ eine abelsche Gruppe mit neutralem Element 0 und $(K \setminus \{0\}, \cdot)$ eine abelsche Gruppe ist und das Distributivgesetz gilt.

Beispiel 7.3.2. $(\mathbb{Q}, +, \cdot)$ ist ein Körper, $(\mathbb{R}, +, \cdot)$ ist ein Körper, die komplexen Zahlen sind ein Körper.

Es gibt auch endliche Körper, d.h. Körper $(K, +, \circ)$, wobei K eine endliche Menge ist. $\mathbb{Z}_2 = \{0, 1\}$ mit Addition und Multiplikation modulo 2 ist ein Körper. $\mathbb{Z}_3 = \{0, 1, 2\}$ mit Addition und Multiplikation modulo 3 ist ein Körper.

Allgemein gilt:

Satz 7.3.3. $\mathbb{Z}_p$ mit der Addition und Multiplikation modulo p ist genau dann ein Körper, wenn p eine Primzahl ist.

Eine Richtung ist einsichtig, da $\mathbb{Z}_p^*$ eine multiplikative Gruppe ist und $\mathbb{Z}_p^* = \mathbb{Z}_p \setminus \{0\}$. Ist p keine Primzahl, dann ist $\mathbb{Z}_p$ nicht null-teilerfrei, d.h. wenn $p = q_1 \cdot q_2$, dann ist $q_1 \cdot q_2 \equiv 0 \pmod{p}$.

Trotzdem gibt es Körper mit $n - 1$ Elementen, wobei n keine Primzahl ist: Man muss Multiplikation und Addition anders definieren. Allgemein gilt der folgende Satz:

Satz 7.3.4. Es gibt genau dann einen endlichen Körper mit q Elementen wenn $q = p^n$ mit p eine Primzahl und $n \in \mathbb{N}$.

7.4 Schlussbemerkungen

Eine Einführung in die Algebra ist (TT13) zu finden, die auch das RSA-Verfahren genauer beschreiben. Ein ausführliches Kapitel zu algebraischen Strukturen enthält (Ste07). Eine mathematische Einführung in die Algebra ist z.B. in (Wol10) zu finden.

Literatur

- Berghammer, Rudolf: *Mathematik für die Informatik: Grundlegende Begriffe, Strukturen und Anwendungen*. 5. Auflage. Springer Vieweg, 2024. – ISBN 9783658441487
- Beutelspacher, Albrecht ; Zschiegner, Marc-Alexander: *Diskrete Mathematik für Einsteiger: Bachelor und Lehramt*. Springer Fachmedien Wiesbaden, 2014. – ISBN 9783658057800
- Cantor, Georg: Beiträge zur Begründung der transfiniten Mengenlehre. In: *Mathematische Annalen* 46 (1895), Nov, Nr. 4, S. 481–512. – ISSN 1432–1807
- Cummings, Jay: *Proofs: A Long-Form Mathematics Textbook*. 2021. – ISBN 9798595265973
- Diestel, Reinhard: *Graphentheorie*. 5. Auflage. Berlin : Springer Spektrum, 2017. – ISBN 3–662–53633–1
- Ebbinghaus, Heinz-Dieter: *Einführung in die Mengenlehre*. 5. Auflage. Springer Spektrum, 2021. – ISBN 978–3–662–63865–1
- Ebbinghaus, Heinz-Dieter ; Flum, Jörg ; Thomas, Wolfgang: *Einführung in die mathematische Logik*. 6. Auflage. Springer Spektrum, 2018. – ISBN 978–3–662–58028–8
- Iwanowski, Sebastian ; Lang, Rainer: *Diskrete Mathematik mit Grundlagen: Lehrbuch für Studierende von MINT-Fächern*. 2. Auflage. Springer Fachmedien Wiesbaden, 2021. – ISBN 978–3–658–32760–6
- Kreuzer, Martin ; Kühling, Stefan: *Logik für Informatiker*. Pearson Studium, 2006. – ISBN 978–3–8273–7215–4
- Klaeren, Herbert ; Sperber, Michael: *Die Macht der Abstraktion: Einführung in die Programmierung*. Wiesbaden : Teubner, 2007. – ISBN 978–3–8351–9079–5
- Meinel, Christoph ; Mundhenk, Martin: *Mathematische Grundlagen der Informatik: Mathematisches Denken und Beweisen. Eine Einführung*. 7. Auflage. Springer Vieweg, 2024. – ISBN 9783658431358
- Ohlbach, Hans J. ; Eisinger, Norbert: *Design Patterns für mathematische Beweise: Ein Leitfaden insbesondere für Informatiker*. Springer Vieweg, 2017. – ISBN 978–3–6625–5651–1
- Rivest, Ronald L. ; Shamir, Adi ; Adleman, Leonard: A method for obtaining digital signatures and public-key cryptosystems. In: *Communications of the ACM* 21 (1978), Nr. 2, S. 120–126
- Russell, Bertrand: *The principles of mathematics*. London : G. Allen and Unwin, Ltd, 1937
- Steger, Angelika: *Diskrete Strukturen: Band 1 Kombinatorik, Graphentheorie, Alge-*

- bra*. 2. Aufl. Berlin : Springer, 2007. – ISBN 9783540466604
- Trudeau, Richard J.: *Introduction to Graph Theory*. New York : Dover, 1993. – ISBN 0486678709
- Teschl, Gerald ; Teschl, Susanne: *eXamen.press*. Bd. 1: *Mathematik für Informatiker – Diskrete Mathematik und Lineare Algebra*. 4. Auflage. Berlin, Heidelberg : Springer Vieweg, 2013. – ISBN 978-3-642-37971-0
- Wilson, Robin J.: *Introduction to graph theory*. 5. Auflage. Pearson, 2010. – ISBN 9780273728894
- Wolfart, Jürgen: *Einführung in die Zahlentheorie und Algebra*. 2. Aufl. Vieweg+Teubner, 2010. – ISBN 9783834814616