

Mathematik und Logik 2005W

Franz Binder

24. Januar 2006

Inhaltsverzeichnis

1	Logik	3
1.1	Aussagen und Beweise	3
1.2	Aussagenlogik	5
	Implikation, ($\Rightarrow$)	5
	Konjunktion (Logisches Und, $\wedge$, &)	6
	Disjunktion, Logisches Oder	8
1.3	Konstruktionen/Beweise	10
	Implikation	10
	Konjunktion	11
	Disjunktion	12
1.4	Prädikatenlogik	13
	Existenz-Quantor	14
1.5	Boolsche Logik	16
	Negation	16
	Entscheidbare Aussagen	17
	Wahrheitstafeln	18
	Gleichungen der Booleschen Algebra	18
	Disjunktive Normalform	18
	Klassische Logik	20
2	Mengen	22
2.1	Äquivalenzrelationen	22
2.2	Konstruktionen für Mengen	23
	Funktionen	23
	Direktes Produkt	25
	Direkte Summe	25
2.3	Potenzmenge	26
	Teilmengen	26
	Mengenalgebra	27
	Beispiele	29
2.4	Gleichmächtigkeit	30
2.5	Konstruktion der Zahlenmengen	32
	Ganze Zahlen	32
	Rationale Zahlen	35
	Irrationale Zahlen	36
	Reelle Zahlen	36
2.6	Kombinatorik	37
	Permutationen	37

Kombinationen	37
-------------------------	----

3 Rekursion 39

3.1 Natürliche Zahlen	39
Konstruktoren	39
Peano-Induktion	39
Selektoren	40
Gleichheit	40
Vorgänger	42
Addition	44
3.2 Teilbarkeit	45
3.3 Modulare Arithmetik	47
3.4 Primzahlen	49
RSA-Verfahren	51
3.5 Listen	53
Introduktion: Listen-Konstruktoren	53
Elimination: Listen-Induktion	53
Selektor: Rekursion	54
Gleichheit	54
Länge	56
Hintenanfügen und Umkehren	56
Listenverkettung	57
Auswahl	59
Geordnete Listen	60

Kapitel 1

Logik

Die *mathematische Logik* verwendet mathematische Methoden, um das logische Denken formal zu beschreiben. Neben dem klassischen Zweck, die verschiedenen Möglichkeiten des logischen Schließens zu beschreiben, zu präzisieren und darüber zu reflektieren, dient die mathematische Logik in jüngerer Zeit in zunehmendem Maße vor allem dazu, Wissen aller Fachgebiete exakt darzustellen, abzufragen, zu überprüfen und allgemein einer automatischen Verarbeitung zugänglich zu machen.

1.1 Aussagen und Beweise

Zentrales Objekt der Betrachtung in der Logik sind die Aussagen, zusammen mit Kriterien um festzustellen, ob (und warum) diese wahr oder falsch sind. Es liegt somit nahe, zuerst den Begriff *Aussage* näher zu präzisieren.

Nicht jeder Satz einer natürlichen Sprache kommt als Aussage in Frage. Frage- und Befehlssätze kommen selbstverständlich nicht als Aussagen in Betracht; aber auch viele Sätze, die in einem grammatikalischen Sinne Aussagesätze sind, bezeichnen nicht wirklich Aussagen, etwa: zynische oder aus dem Zusammenhang gerissene Sätze, oder allgemein Sätze, deren Bedeutung nicht klar definiert ist.

Da die Gültigkeit von Aussagen ein zentrales Thema ist, wird als populäre Definition gerne verwendet, daß genau diejenigen Sätze als Aussagen zugelassen sind, die entweder wahr oder falsch sind. Dabei tritt das prinzipielle Problem auf, daß eigentlich zuerst definiert werden sollte, wann ein Satz wahr oder falsch ist (was gerade die Aufgabe der Logik wäre). Diese Definition wird daher oft modifiziert zu: Aussagen sind Sätze, von denen es einen *Sinn* macht zu fragen, ob sie wahr oder falsch sind. Damit ist das Problem auf die Sinnfrage reduziert, welche außerhalb von Mathematik und Linguistik behandelt wird, und wohl für immer ungelöst bleiben wird.

Der klassische Ansatz, um dieses Problem zu lösen ist, eine formale Sprache zu definieren, welche genau die mathematischen Aussagen enthält, und von diesen dann implizit festzulegen, welche wahr und welche falsch sein sollten. Dieser Ansatz war tatsächlich sehr erfolgreich, und der Großteil der Kenntnisse in der mathematischen Logik wurde auf diese Weise gewonnen. Ob dadurch tatsächlich das logische Denken beschrieben wird, ist freilich eine andere Frage. Immerhin

kann man die Ergebnisse auf deren praktische Verwertbarkeit überprüfen.

Es stellt sich heraus, daß es für die Praxis gar nicht so wichtig ist, festzulegen, was Aussagen tatsächlich sind, und auch nicht, ob sie wahr oder falsch sind. Ein klassisches Resultat der Logik besagt sogar, daß dies sowieso unmöglich ist. Aus praktischer Sicht viel wichtiger ist vielmehr, ob man die Gültigkeit einer Aussage feststellen (*beweisen*) kann, oder nicht. Es geht also primär darum, welche Beziehungen zwischen einer Aussage und möglichen Beweisen bestehen. Wir präzisieren die beiden Begriffe daher vorerst nur durch eine einfache, sich in der Folge aber als sehr wirkungsvoll herausstellende nicht-formale Annahme.

1.1.1 DEFINITION (AUSSAGE). Eine *Aussage*(*proposition*) ist eine Konstruktion, durch welche festgelegt wird, wie ihre *Beweise*(*proof*) zu konstruieren sind. Eine Aussage, für welche ein Beweis konstruiert werden kann, heißt *wahr*(*true*).

Diese Definition vermeidet bewußt jede unnötige nähere Festlegung von Begriffen wie *Aussage*, *Beweis*, *wahr*, *falsch* *Sinn*, *Bedeutung*. Auch, ob es, neben den beweisbaren Aussagen, noch weitere wahre Aussagen gibt, wird durch diese Definition bewußt nicht festgelegt. Jedenfalls gilt eine Aussage aber nicht unbedingt als falsch, wenn sie nicht bewiesen werden kann.

Der hier verwendete Begriff des Beweises ist freilich noch näher festzulegen. Die allgemeine Vorgangsweise dazu ist die, daß verschiedene logische Operatoren definiert werden, mit der mehrere Aussagen verknüpft werden, um eine neue Aussage zu erhalten. Von jedem Junktore muß dann nur noch festgelegt werden, wie die Beweise der betroffenen Aussagen mit denen der neuen Aussage in Beziehung stehen.

1.1.2 BEISPIEL.

1. Aussagen wie $5 > 3$ sind trivial bzw. leicht nachzurechnen. Wie deren Beweise aussehen, hängt hauptsächlich davon ab, was man unter natürlichen Zahlen und dem Größerbegriff genau versteht, und ist für die Praxis irrelevant. Wichtig ist hier nur, ob es einen Beweis gibt oder nicht.
2. Die Aussage $3 > 5$ hat keinen Beweis.
3. Jede natürliche Zahl ist eine Beweis, dass es eine natürliche Zahl gibt. Daher kann man die Mengen der natürlichen Zahlen mit der Gesamtheit aller Beweise von „es gibt eine natürliche Zahl“ identifizieren.
4. Die Aussage „Zu jeder natürlichen Zahl gibt es eine größere“ erfordert als Beweis unter anderem eine Konstruktion, welche imstande ist, zu jeder beliebigen natürlichen eine größere zu konstruieren. Die „Funktion“ $(+1)$, also $x \mapsto x + 1$, kann das. Ebenso $(+2)$, $(+100)$, $(2*)$, $(2^{\wedge})$, und viele mehr.

1.2 Aussagenlogik

Implikation, ($\Rightarrow$)

Formation

1.2.1 DEFINITION. Sind P und Q Aussagen, dann bezeichnet $P \Rightarrow Q$ ebenfalls eine Aussage, die *Implikation* von P und Q .

Introduktion

1.2.2 DEFINITION. Um $P \Rightarrow Q$ zu beweisen, muß man Q beweisen, wobei man einen Beweis von P voraussetzen darf.

Diese Definition wird formal durch die folgende Regel ausgedrückt:

1.2.3 AXIOM ($\Rightarrow$ -EINFÜHRUNGSREGEL).

$$\frac{\begin{array}{|c|} \hline P \\ \vdots \\ Q \\ \hline \end{array}}{P \Rightarrow Q} \Rightarrow \mathcal{I}$$

Dabei bezeichnet der eingerahmte Teil einen *Beweis-Rahmen* (*proof box*). Dieser bedeutet, daß man, mit Hilfe der oberhalb des Rahmens bezeichneten Annahme (*Hypothese*) P , die Aussage Q am unteren Rand des Rahmens herleiten kann.

In der Praxis werden Beweise verbal wiedergegeben. Die obige Regel liest sich dann etwa folgendermaßen:

Wir möchten $P \Rightarrow Q$ beweisen. Dazu nehmen wir an, es gelte P . Daraus leiten wir nun Q her.

Elimination

Wurde $P \Rightarrow Q$ mittels der Introduktionsregel bewiesen, so gibt es einen Beweis von Q , welcher unter der Annahme P funktioniert. Steht nun zusätzlich ein Beweis von P zur Verfügung, so hat man einen Beweis von Q :

1.2.4 AXIOM ($\Rightarrow$ -ELIMINATIONSREGEL, *Modus Ponens*).

$$\frac{P \Rightarrow Q \quad P}{Q} \Rightarrow \mathcal{E}$$

In verbalen Beweisen spiegelt sich die Eliminationsregel etwa folgendermaßen wieder:

Da wir P bewiesen haben, gilt wegen $P \Rightarrow Q$ auch Q .

Transitivität

1.2.5 SATZ. Die Implikation ist transitiv, d.h. wir haben die hergeleitete Schlussregel

$$\frac{A \Rightarrow B \quad B \Rightarrow C}{A \Rightarrow C} \Rightarrow \text{Transitivität}$$

Beweis.

$$\begin{array}{c}
 A \Rightarrow B \quad B \Rightarrow C \\
 \boxed{
 \begin{array}{c}
 A \\
 \boxed{
 \begin{array}{c}
 B \Rightarrow C \quad \frac{A \Rightarrow B \quad A}{B} \Rightarrow \mathcal{E} \\
 \hline
 C \Rightarrow \mathcal{E}
 \end{array}
 \\
 \hline
 A \Rightarrow C
 \end{array}
 } \Rightarrow \mathcal{I}
 \end{array}$$

Oder verbal:

Annahmen: $A \Rightarrow B$ und $B \Rightarrow C$. Zu zeigen haben wir $A \Rightarrow C$. Gemäß $\Rightarrow \mathcal{I}$ nehmen die zusätzliche Hypothese A an und versuchen damit C herzuleiten. Zweimaliges Anwenden von $\Rightarrow \mathcal{E}$ liefert aus A und $A \Rightarrow B$ zuerst B , und daraus dann, zusammen mit $B \Rightarrow C$, die Behauptung C . $\square$

Konjunktion (Logisches Und, $\wedge$, &)

Formation

1.2.6 DEFINITION. Seien P und Q Aussagen. Dann bezeichnet $P \wedge Q$ ebenfalls eine Aussage, die *Konjunktion* von P und Q .

Introduktion

1.2.7 DEFINITION. Um die Aussage $P \wedge Q$ zu beweisen, muß man sowohl P als auch Q beweisen.

Diese Definition läßt sich mit der folgenden *Einführungsregel* (*introduction rule*) formalisieren:

1.2.8 AXIOM ($\wedge$ -EINFÜHRUNGSREGEL).

$$\frac{P \quad Q}{P \wedge Q} \wedge \mathcal{I} \tag{1.1}$$

Diese Regel ermöglicht es, eine Konjunktion aus anderen Aussagen herzuleiten, also in die Gesamtheit der bewiesenen Aussagen einzuführen. Bei dieser Notation steht die neu eingeführte Aussage unter dem Strich, und die Aussagen, welche vorher zu beweisen sind, stehen über dem Strich.

Durch die Regel $\wedge \mathcal{I}$ wird das umgangssprachliche Und, so wie es in der Mathematik gebraucht wird, vollständig beschrieben, d.h. es gibt keine weiteren Einführungsregeln für $\wedge$. Man beachte aber, daß das Wort „und“, vor allem außerhalb der Mathematik, oft auch in anderen Bedeutungen gebraucht wird, etwa „und daher“, „und trotzdem“, „und dann“, „aber“, usw.

Elimination

Die Bemerkung, daß es nur die eine Einführungsregel für die Konjunktion gibt, läßt sich nicht direkt in Form einer Regel anschreiben. Wir können diese Tatsache aber auch folgendermaßen verstehen: Wurde $P \wedge Q$ bewiesen, dann hat man auch einen Beweis von P und einen von Q , da ja ein Beweis von $P \wedge Q$

gerade aus diesen zwei Beweisen besteht. Auf diese Art läßt sich das Wissen um einen Beweis von $P \wedge Q$ verwerten. Für jede der Prämissen der Einführungsregel erhalten wir eine dazugehörige *Eliminationsregel* (*elimination rule*).

1.2.9 AXIOM ($\wedge$ -ELIMINATIONSREGELN).

$$\frac{P \wedge Q}{P} \wedge \mathcal{E}_0 \quad \frac{P \wedge Q}{Q} \wedge \mathcal{E}_1 \quad (1.2)$$

Die Eliminationsregeln drücken aus, daß wenn $P \wedge Q$ bewiesen worden ist, dann irgendwann einmal auch P und Q bewiesen worden sind, und daher auch deren Beweise zur Verfügung stehen müssen. Bei den Eliminationsregeln steht der betroffene Operator stets oberhalb des Striches.

Kommutativität

1.2.10 SATZ. *Falls $A \wedge B$ bewiesen werden kann, dann auch $B \wedge A$.*

Beweis. Der formale Beweis

$$\frac{A \wedge B}{\frac{\frac{A \wedge B}{B} \wedge \mathcal{E}_1 \quad \frac{A \wedge B}{A} \wedge \mathcal{E}_0}{B \wedge A} \wedge \mathcal{I}}$$

kann folgendermaßen gelesen werden (in (sehr) verbaler Ausführung): Die Annahme besagt $A \wedge B$. Zu zeigen ist $B \wedge A$. Aus der Annahme erhalten wir mit $\wedge \mathcal{E}_1$, daß B gilt, und ebenso mit $\wedge \mathcal{E}_0$, daß A gilt. Damit erhält man mit $\wedge \mathcal{I}$ schließlich $B \wedge A$, wie zu zeigen war. $\square$

Analog gestaltet sich der Beweis für die *Assoziativität* der Konjunktion:

$$\frac{(A \wedge B) \wedge C}{A \wedge (B \wedge C)} \wedge\text{-associative}_0 \quad \frac{A \wedge (B \wedge C)}{(A \wedge B) \wedge C} \wedge\text{-associative}_1$$

Logische Äquivalenz

1.2.11 DEFINITION. Die logische *Äquivalenz* $P \Leftrightarrow Q$ ist definiert als Abkürzung für $(P \Rightarrow Q) \wedge (Q \Rightarrow P)$.

1.2.12 NOTATION. Statt $P \Rightarrow Q$ kann man auch gleichbedeutend $Q \Leftarrow P$ schreiben.

1.2.13 THEOREM. *Die logische Äquivalenz ist eine Äquivalenzrelation, d.h. es gelten*

Reflexivität $A \Leftrightarrow A$;

Symmetrie Wenn $A \Leftrightarrow B$, dann auch $B \Leftrightarrow A$;

Transitivität Wenn $A \Leftrightarrow B$ und $B \Leftrightarrow C$, dann auch $A \Leftrightarrow C$.

Disjunktion, Logisches Oder

Formation

1.2.14 DEFINITION. Seien P und Q Aussagen. Dann bezeichnet $P \vee Q$ ebenfalls eine Aussage, die *Disjunktion* von P und Q .

Introduktion

1.2.15 DEFINITION. Um die *Disjunktion* $A \vee B$ zu beweisen, ist P zu beweisen oder auch Q

Auch diese Definition läßt sich in Form von logischen Schlußregeln formulieren:

1.2.16 AXIOM ($\vee$ -EINFÜHRUNGSREGELN).

$$\frac{P}{P \vee Q} \vee \mathcal{I}_0 \quad \frac{Q}{P \vee Q} \vee \mathcal{I}_1$$

Hier erhalten wir zwei Einführungsregeln, weil es zwei Möglichkeiten gibt, eine Disjunktion zu beweisen.

Elimination

Jede der Einführungsregeln hat nur eine Prämisse, daher gibt es auch nur eine Eliminationsregel. Sie ist aber etwas komplizierter, weil zwei Fälle unterschieden werden müssen.

1.2.17 AXIOM (BEWEIS DURCH FALLUNTERSCHIEDUNG).

$$\frac{P \implies R \quad Q \implies R}{P \vee Q \implies R} \vee \mathcal{E}$$

Distributivität

Wir zeigen eines der beiden Distributivgesetze:

1.2.18 SATZ.

$$A \wedge (B \vee C) \Leftrightarrow (A \wedge B) \vee (A \wedge C).$$

Beweis. Diese Aussage ist eine Konjunktion (Äquivalenz ist ja nur die Abkürzung für die Konjunktion der Implikationen in beide Richtungen). Wir haben also beide Richtungen zu zeigen.

$\implies$: Es gelte $A \wedge (B \vee C)$. Zu zeigen ist $(A \wedge B) \vee (A \wedge C)$. Mittels $\wedge \mathcal{E}$ erhalten wir $B \vee C$. Für jeden dieser Fälle zeigen wir $(A \wedge B) \vee (B \wedge C)$:

Fall: es gelte B . Wir erhalten A aus der Voraussetzung mittels $\wedge \mathcal{E}_0$, und daher mittels $\wedge \mathcal{I}$ auch $A \wedge B$, was mit $\vee \mathcal{I}_0$ zu $(A \wedge B) \vee (B \wedge C)$ wird.

Fall: es gelte C : Analog zum vorigen Fall, erhalten wir ebenfalls $(A \wedge B) \vee (B \wedge C)$.

$$\begin{array}{c}
 A \wedge (B \vee C) \\
 \hline
 \begin{array}{cc}
 B & C \\
 \hline
 \begin{array}{c}
 \frac{A \wedge (B \vee C)}{A} \wedge \mathcal{E} \quad B \\
 \hline
 A \wedge B \quad \wedge \mathcal{I} \\
 \hline
 (A \wedge B) \vee (A \wedge C) \quad \vee \mathcal{I}
 \end{array}
 &
 \begin{array}{c}
 \frac{A \wedge (B \vee C)}{A} \wedge \mathcal{E} \quad C \\
 \hline
 A \wedge C \quad \wedge \mathcal{I} \\
 \hline
 (A \wedge B) \vee (A \wedge C) \quad \vee \mathcal{I}
 \end{array}
 \end{array} \\
 \hline
 \frac{A \wedge (B \vee C)}{B \vee C} \wedge \mathcal{E} \quad (A \wedge B) \vee (A \wedge C) \quad \vee \mathcal{E}
 \end{array}$$

$\Leftarrow$: Wir nehmen wegen $(\Rightarrow \mathcal{I})$ an, daß $(A \wedge B) \vee (A \wedge C)$ und versuchen sowohl A als auch $B \vee C$ herzuleiten ($\wedge \mathcal{I}$). Gemäß ($\vee \mathcal{E}$) betrachten wir zuerst den Fall $(A \wedge B)$. Dann gilt sowohl A (mit $\wedge \mathcal{E}0$) als auch B (mit $\wedge \mathcal{E}1$), und damit $B \vee C$ (mit $\vee \mathcal{I}0$). Analog für den Fall $A \wedge C$. $\square$

Die Umkehrung haben wir hier verbal, aber nicht streng formal gemacht; die andere Variante ist eine gute **Übung!**

1.3 Konstruktionen/Beweise

Objekte unserer Betrachtung seien nun nicht mehr die Aussagen, sondern deren Beweise. Jeder Aussage entspricht eine gewisse Klasse von Beweisen, nämlich die Beweise dieser Aussage.

Bisher haben wir uns nur damit beschäftigt, ob eine Aussage einen Beweis hat oder nicht. Nun sehen wir uns die Beweise selbst etwas genauer an.

1.3.1 NOTATION. Ist x ein Beweis der Aussage A , dann schreiben wir $x: A$.

Umgekehrt kann man jede Zusammenfassung von Objekten mit einer Aussage identifizieren. Damit kann $x: A$ auch bedeuten, daß x den Datentyp A hat, was oft mit $x \in A$ bezeichnet wird.

Implikation

Konstruktor

Wir annotieren die Einführungsregel für die Implikation mit den konkreten Beweisen. Der Beweis von A ist rein hypothetisch und wird mit einer Variablen, die sonst noch keine Bedeutung hat, bezeichnet. Die Notation $t[x]$ bedeutet dann, daß t irgend eine Konstruktion ist, in der das x vorkommen kann.

1.3.2 AXIOM (FUNKTIONSDEFINITION).

$$\frac{\boxed{\begin{array}{c} x: A \\ \vdots \\ t[x]: B \end{array}}}{x \mapsto t[x]: A \Rightarrow B} \Rightarrow \mathcal{I}$$

Ein Beweis von $A \Rightarrow B$ ist also eine „Funktion“ von A nach B und die Einführungsregel entspricht der üblichen Art, wie man Funktionen definiert. Eine klassische Variante dieser Schreibweise ist $\lambda x. t[x]$ (daher λ -Kalkül). Statt einer Definition $f = x \mapsto x^2$ schreibt man üblicherweise $f(x) = x^2$. Dabei sind oft auch Muster erlaubt, etwa in

$$f(x, y) = (y, x).$$

Dies ist insbesondere im Zusammenhang mit Konstruktoren sinnvoll.

Selektor

1.3.3 AXIOM (FUNKTIONSANWENDUNG).

$$\frac{f: A \Rightarrow B \quad x: A}{fx: B} \Rightarrow \mathcal{E}$$

Einem Beweis mit der Eliminationsregel entspricht also Anwendung einer Funktion. Oft schreibt man statt fx auch $f(x)$. Für spezielle Funktionen sind auch andere Konventionen üblich, etwa die Formen xf oder x^f (vgl.: $5!$, $\sin'$).

Nun annotieren wir auch den Beweis der Transitivität:

$$\frac{\begin{array}{c} A \\ \boxed{\begin{array}{c} g: B \Rightarrow C \quad \frac{f: A \Rightarrow B \quad x: A}{fx: B} \Rightarrow \mathcal{E} \\ \hline g(fx): C \end{array}} \Rightarrow \mathcal{I} \\ x \mapsto g(fx): A \Rightarrow C \end{array}}$$

Der Beweis der Transitivität der Implikation entspricht also der Hintereinanderausführung von Funktionen. Für die Konstruktion $x \mapsto g(fx)$ schreiben wir $g \circ f$.

Konjunktion

Konstruktor

Wir annotieren die Einführungsregel der Konjunktion mit den entsprechenden Beweisen:

1.3.4 AXIOM (PAAR-KONSTRUKTOR).

$$\frac{x: A \quad y: B}{(x, y): A \wedge B} \wedge \mathcal{I}$$

Die Konstruktion $(,)$ zum Paare bilden macht aus einem Beweis x von A und einem Beweis y von B den Beweis (x, y) von $A \wedge B$. Man spricht daher von einem *Konstruktor*.

Die Aussage $A \wedge B$ entspricht daher einem Datentyp, dessen Objekte aus zwei Komponenten bestehen, einem vom Typ A , und einem vom Typ B . Die meisten höheren Programmiersprachen bieten die Möglichkeiten, neue Datentypen auf diese Weise zu definieren. Dabei können natürlich auch mehr als zwei Objekte zu einem zusammengefaßt werden. Die Details, wie das gemacht wird, sind recht unterschiedlich, ebenso wie die Namen (Verbund-Datentypen, Records, Strukturen, Datensätze). In Java verwendet man Klassen. Einem Konstruktor entspricht in Java eine Konstruktor-Methode, mit der ein neues Objekt der Klasse angelegt wird.

Selektoren

Wir annotieren nun auch die beiden Eliminationsregeln mit den entsprechenden Beweisen.

$$\frac{(x, y): A \wedge B}{x: A} \wedge \mathcal{E}_0 \quad \frac{(x, y): A \wedge B}{y: B} \wedge \mathcal{E}_1$$

Hier haben wir die Konstruktion des Beweises von $A \wedge B$ explizit angeschrieben (ein Beweis von $A \wedge B$ muß ja so aussehen). Die Konstruktion eines Beweises von A bzw. B besteht dann im Selektieren der passenden Komponente. Eine alternative Notation ist die folgende:

$$\frac{z: A \wedge B}{\text{fst } z} \wedge \mathcal{E}_0 \quad \frac{z: A \wedge B}{\text{snd } z} \wedge \mathcal{E}_1$$

Hier war es nicht notwendig, die Konstruktion des Beweises der Prämisse anzuschreiben, dafür wurden die *Selektoren* explizit bezeichnet.

Auch in Java gibt es in einer Klasse üblicherweise Selektor-Methoden, mit denen auf die Komponenten der Objekte einer Klasse zugegriffen werden kann.

Kommutativität mit Beweisannotation:

$$\boxed{\frac{\frac{A \wedge B}{y: B} \wedge \mathcal{E}_1 \quad \frac{A \wedge B}{x: A} \wedge \mathcal{E}_0}{(y, x): B \wedge A} \wedge \mathcal{I}}$$

Damit haben wir die Regel für die *Kommutativität* der Konjunktion

$$\frac{A \wedge B}{B \wedge A} \wedge\text{-Kommutativität},$$

eine *hergeleitete Schlußregel* (*derived rule*) (kein Axiom!) hergeleitet; und ein Beweis davon ist die Konstruktion, welche die beiden Komponenten eines Paares vertauscht.

Disjunktion

Konstruktoren

1.3.5 AXIOM.

$$\frac{x: A}{\text{Left } x: A \vee B} \vee \mathcal{I}_0 \quad \frac{y: B}{\text{Right } y: A \vee B} \vee \mathcal{I}_1$$

Ein Beweis von $A \vee B$ ist also entweder ein Beweis von A oder ein Beweis von B ; der verwendete Konstruktor gibt an, welcher dieser beiden Fälle eintritt. Es handelt sich also um die „Vereinigung“ der Beweise von A und der Beweise von B . Es liegt somit eine Vereinigung von zwei Datentypen vor.

Selektor

Wir erhalten damit die folgenden ergänzten Schlußregeln

1.3.6 AXIOM (FALLUNTERSCHIEDUNG, (IF, SWITCH, CASE)).

$$\frac{z: A \vee B \quad \boxed{\begin{array}{c} x: A \\ \vdots \\ a[x]: C \end{array}} \quad \boxed{\begin{array}{c} y: B \\ \vdots \\ b[y]: C \end{array}}}{\text{switch } z \begin{array}{l} \text{case Left } x \mapsto a[x] \\ \text{case Right } y \mapsto b[y] \end{array} : C} \vee \mathcal{E}$$

Dabei sind wieder $a[x]$ und $b[y]$ irgendwelche Terme, in denen x bzw. y vorkommen können.

1.4 Prädikatenlogik

Um eine Aussage der Form $\bigwedge_{x: X} A[x]$ zu beweisen, muss man $A[x]$ für ein beliebiges $x: X$ zeigen. Die entsprechende Regeln lautet hiermit:

1.4.1 AXIOM.

$$\frac{\boxed{\begin{array}{c} x: X \\ \vdots \\ A[x] \end{array}}}{\bigwedge_{x: X} A[x]} \forall\mathcal{E}$$

Eine alternative Schreibweise ist $\forall(x: X).A[x]$.

1.4.2 BEMERKUNG. Die Implikation ergibt sich als Spezialfall einer All-Aussage, in der $A[x]$ nicht von x abhängt.

Elimination

Wenn $A[x]$ für ein beliebiges x bewiesen wurde, dann kann man auch für die Variable x einen beliebigen Ausdruck t einsetzen. Die Notation $A[t]$ bezeichnet jene Aussage, die aus $A[x]$ entsteht, wenn alle Vorkommen der Variablen x durch t ersetzt werden.

1.4.3 AXIOM.

$$\frac{\bigwedge_{x: X} A[x] \quad t: X}{A[t]} \forall\mathcal{E}$$

Konstruktor

Die Natur der Beweise einer All-Aussage ergibt sich wieder durch Annotation der Einführungsregel:

1.4.4 AXIOM.

$$\frac{\boxed{\begin{array}{c} x: X \\ \vdots \\ t[x]: A[x] \end{array}}}{x \mapsto t[x] \quad \bigwedge_{x: X} A[x]} \forall\mathcal{E}$$

Damit ergibt sich prinzipiell derselbe Konstruktor wie bei der Implikation. Allerdings haben hier die Funktionswerte nicht alle denselben Datentyp, sondern den Typ $A[x]$, welcher von x abhängen kann.

Selektor

1.4.5 AXIOM.

$$\frac{f: \bigwedge_{x: X} A[x] \quad t: X}{ft: A[t]} \forall \mathcal{E}$$

Auch hier handelt es sich um eine Verallgemeinerung der üblichen Funktionsanwendung, bei der die Datentypen des Ergebnisses vom Input abhängen können.

Existenz-Quantor

Formation

$A[x]$ bezeichne eine Aussage, in der x vorkommen kann.

1.4.6 DEFINITION. Ist $A[x]$ für jedes $x: X$ eine Aussage, dann ist auch $\bigvee_{x: X} A[x]$ eine *Existenz-Aussage*.

Eine alternative Schreibweise ist $\exists(x: X).A[x]$.

Introduktion

1.4.7 DEFINITION. Um eine Aussage der Form $\bigvee_{x: X} A[x]$ zu beweisen, muss man $a: X$ konstruieren, sodaß man $A[a]$ beweisen kann.

1.4.8 AXIOM ($\bigvee$ -EINFÜHRUNGSREGEL).

$$\frac{a: X \quad A[a]}{\bigvee_{x: X} A[x]} \exists \mathcal{I} \tag{1.3}$$

1.4.9 BEMERKUNG. Die Konjunktion ergibt sich als Spezialfall einer Existenz-Aussage, in der $A[x]$ nicht von x abhängt.

Elimination

1.4.10 AXIOM.

$$\frac{\bigvee_{x: X} A[x] \quad \boxed{\begin{array}{c} y: X \quad A[y] \\ \vdots \\ C \end{array}}}{C} \exists \mathcal{E}$$

Diese Eliminationsregel entspricht der Einführung von „einem solchen x so daß $A[x]$ gilt“. Genauer: es gelte $\bigvee_{x: X} A[x]$ und die Aussage C sei zu beweisen; dann darf man ein neues Symbol y einführen und annehmen, daß $A[y]$ gilt. Dieses neue Symbol darf in C nicht vorkommen (sonst wäre es nicht neu), es darf aber mitunter mit dem x übereinstimmen, da das x innerhalb der Existenzaussage nur innerhalb dieser eine Bedeutung hat (es ist durch den Existenzquantor gebunden).

Konstruktor

1.4.11 AXIOM ($\bigvee$ -EINFÜHRUNGSREGEL).

$$\frac{a: X \quad t: A[a]}{(a, t): \bigvee_{x: X} A[x]} \exists\mathcal{I} \quad (1.4)$$

Der Beweis einer Existenzaussage ist somit ein Paar; allerdings darf der Datentyp der zweiten Komponente von der ersten Komponente abhängen.

Selektoren

1.4.12 AXIOM.

$$(a, t): \frac{\bigvee_{x: X} A[x] \quad \boxed{\begin{array}{c} y: X \quad u: A[y] \\ \vdots \\ f[y, a]: C \end{array}}}{\{(y, u) := (a, t)C; f[y, u]: C\}} \exists\mathcal{E}$$

Der Selektor ist daher eine lokale Definition.

1.5 Boolesche Logik

Negation

Sind P und Q Aussagen, die bewiesen wurden, dann gilt jedenfalls $P \iff Q$, d.h. alle bewiesenen Aussagen sind äquivalent. Mit dem Symbol $\top$ (sprich: top) möchten wir irgendeine solche Aussage bezeichnen (andere Bezeichnungen: T, W, true, wahre Aussage). Umgekehrt soll $\perp$ (sprich: bottom) einen Widerspruch (falsche Aussage) bezeichnen (andere Bezeichnungen: F, false, falsum, $\frac{1}{2}$). Wir können diese Begriffe wie gewohnt mit Schlußregeln definieren.

1.5.1 AXIOM. Sei P irgendeine Aussage. Dann gelten die Schlußregeln:

$$\frac{P}{\top} \top \mathcal{I} \quad \frac{\perp}{P} \perp \mathcal{E}$$

Es gibt aber keine $\top \mathcal{E}$ - oder $\perp \mathcal{I}$ -Regeln. Für $\top$ ist alles ein Beweis, dafür kann man aus $\top$ aber auch nichts schließen. Umgekehrt kann man aus $\perp$ alles herleiten, aber es gibt keine Möglichkeit, $\perp$ aus anderen Aussagen zu beweisen.

Der Sinn des Symbols $\perp$ besteht darin, die Negation von Aussagen zu definieren. Die Idee ist: wenn einer Aussage jede andere gefolgt werden kann, dann muß sie wohl falsch sein. Die Aussage „Jede Aussage ist wahr“ wäre eine derartige Aussage; aber auch aus etwas konkretem, wie etwa $0 = 1$, kann man jede Aussage herleiten. Ganz allgemein nennt man ein logisches System widersprüchlich, wenn darin jede Aussage herleitbar ist.

1.5.2 DEFINITION. Die Formel $\neg A$ (nicht A) ist eine Abkürzung für $A \Rightarrow \perp$ und bedeutet, daß A falsch ist.

1.5.3 SATZ. Die aussagenlogische Formel

$$A \wedge \neg A \Rightarrow \perp$$

ist allgemeingültig.

Beweis. Direkt aus der Definition durch einfache Anwendung der Implikations-Elimination. $\square$

1.5.4 SATZ. Die Aussage

$$A \Rightarrow \neg \neg A$$

ist allgemeingültig.

Beweis. Dies ist ein Spezialfall von $A \Rightarrow (A \Rightarrow C) \Rightarrow C$. $\square$

1.5.5 BEMERKUNG. Bemerkenswerterweise läßt sich die Allgemeingültigkeit von $\neg \neg A \Rightarrow A$ mit den bisher besprochenen Regeln nicht herleiten.

1.5.6 SATZ. Die Aussage

$$\neg \neg \neg A \Rightarrow \neg A$$

ist allgemeingültig.

Beweis. Dies ist ein Spezialfall von $((A \implies C) \implies C) \implies A \implies C$. $\square$

1.5.7 BEMERKUNG. Dieser Satz besagt insbesondere, daß die Umkehrung von 1.5.4 zumindest für alle Aussagen der Form $\neg A$ gilt.

1.5.8 SATZ. *Die aussagenlogische Formel*

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

ist allgemeingültig.

Beweis. Dies ist ein Spezialfall von $(A \vee B \implies C) \iff (A \implies C) \wedge (B \implies C)$. $\square$

1.5.9 FOLGERUNG. *Die Aussage*

$$\neg\neg(A \vee \neg A)$$

ist allgemeingültig.

Beweis. In 1.5.8 setzen wir statt $\neg A$ für B ein. Den Rest erledigt 1.5.3. $\square$

Entscheidbare Aussagen

Die doppelte Negation in 1.5.9 ist etwas unbefriedigend. Sicherlich interessant sind Aussagen, für welche dies vermieden werden kann.

1.5.10 DEFINITION. Eine Aussage A heißt *entscheidbar*, wenn $A \vee \neg A$ bewiesen werden kann.

1.5.11 BEMERKUNG. Man beachte, daß der Entscheidbarkeitsbegriff im Kontext der klassischen Logik anders definiert wird.

1.5.12 SATZ. *Die Aussage*

$$A \vee \neg A \implies (\neg\neg A \implies A).$$

ist allgemeingültig.

Beweis. Gemäß zweimaliger Implikations-Introduktion nehmen wir die Aussagen $A \vee \neg A$ und $\neg\neg A$ an und haben A herzuleiten. Wir führen eine Fallunterscheidung $(\vee\mathcal{E})$ nach $A \vee \neg A$ durch. Im ersten Fall ist nichts mehr zu zeigen. Gilt aber $\neg A$, so können wir, zusammen mit $\neg\neg A$, mittels $(\implies\mathcal{E})$ den Widerspruch $\perp$ herleiten, woraus mit $\perp\mathcal{E}$ (unter anderem) A folgt. $\square$

1.5.13 BEMERKUNG. Die Umkehrung von 1.5.4 gilt somit auch für alle entscheidbaren Aussagen.

1.5.14 BEMERKUNG. Entscheidbare Aussagen erfüllen somit in jedem Fall $A \vee \neg A$ und $\neg\neg A \iff A$, sowie, leicht anders formuliert:

$$(A \iff \top) \vee (A \iff \perp).$$

Das macht deutlich, daß die Menge der entscheidbaren Aussagen mit der logischen Äquivalenz als Gleichheitsbegriff nur zwei Elemente hat; man schreibt dafür

$$\mathbb{B} = \{\top, \perp\}.$$

Wahrheitstafeln

Die Gültigkeit von aus entscheidbaren Aussagen zusammengesetzten Formeln läßt sich stets durch Fallunterscheidungen überprüfen. Diese wird üblicherweise in Tabellenform vorgenommen, wie im folgenden Satz.

1.5.15 SATZ. Für alle entscheidbaren Aussagen A, B gilt

$$\neg(A \wedge B) \Leftrightarrow (\neg A \vee \neg B).$$

Beweis. Wir untersuchen tabellarisch alle möglichen Wahrheitswerte für A und B :

A	B	$A \wedge B$	$\neg(A \wedge B)$	$\neg A$	$\neg B$	$\neg A \vee \neg B$	$\neg(A \wedge B) \Leftrightarrow (\neg A \vee \neg B)$
$\top$	$\top$	$\top$	$\perp$	$\perp$	$\perp$	$\perp$	$\top$
$\top$	$\perp$	$\perp$	$\top$	$\perp$	$\top$	$\top$	$\top$
$\perp$	$\top$	$\perp$	$\top$	$\top$	$\perp$	$\top$	$\top$
$\perp$	$\perp$	$\perp$	$\top$	$\top$	$\top$	$\top$	$\top$

Da in der letzten Spalte lauter $\top$ stehen, ist die Aussage $\neg(A \wedge B) \Leftrightarrow (\neg A \vee \neg B)$ somit für alle entscheidbaren Aussagen gültig. □

1.5.16 BEMERKUNG. Satz 1.5.15 ist dual zu Satz 1.5.8, der für allgemeine Aussagen bewiesen wurde. Ohne die Zusatzvoraussetzung gilt im Satz 1.5.15 nur die Implikation von rechts nach links.

1.5.17 BEMERKUNG. Man beachte, daß die Methode mit den Wahrheitstafeln ganz entscheidend davon abhängt, daß alle verwendeten Aussagen entscheidbar sind.

Gleichungen der Boolschen Algebra

1.5.18 THEOREM. Die Operationen $\wedge, \vee, \neg$ auf $\mathbb{B}$ erfüllen für entscheidbare Aussagen a, b, c die folgenden Gleichungen:

$(a \wedge b) \wedge c = a \wedge (b \wedge c)$	$(a \vee b) \vee c = a \vee (b \vee c)$
$a \wedge b = b \wedge a$	$a \vee b = b \vee a$
$a \wedge (b \vee c) = (a \wedge b) \vee (a \wedge c)$	$a \vee (b \wedge c) = (a \vee b) \wedge (a \vee c)$
$a \wedge \neg a = \perp$	$a \vee \neg a = \top$
$\neg(a \wedge b) = \neg a \vee \neg b$	$\neg(a \vee b) = \neg a \wedge \neg b$

1.5.19 DEFINITION. Jede Menge mit Operationen, welche diese Gleichungen erfüllen, heißt eine *Boolsche Algebra*.

Disjunktive Normalform

In diesem Abschnitt beschäftigen wir uns nur mit entscheidbaren Aussagen, bzw. gelte das Prinzip von ausgeschlossenen Dritten.

Eine Aussage wie $\neg(A \Rightarrow (B \wedge C))$ hat viele äquivalente Formen. Mitunter ist es sinnvoll eine *kanonische* Form zu berechnen. „Ausmultiplizieren“ ist eine

Möglichkeit. Dabei wird zuerst die jede Implikation der Form $A \Rightarrow B$ durch $\neg A \vee B$ ersetzt. Dann werden die Gesetze von De Morgan

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

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

immer von links nach rechts angewendet, sodaß alle Negationen immer weiter nach innen wandern, bis sie nur noch unmittelbar vor elementaren Aussagen vorkommen. Doppelte Negationen werden weggelassen. Ebenso werden die Distributivgesetze

$$A \wedge (B \vee C) \Leftrightarrow (A \wedge B) \vee (A \wedge C)$$

$$A \vee (B \wedge C) \Leftrightarrow (A \vee B) \wedge (A \vee C)$$

so angewendet, daß die Konjunktionen immer weiter nach innen wandern. Am Schluß ergibt sich damit eine Disjunktion von Konjunktionen von elementaren Aussagen oder deren Negation. Beispiel: $\neg(A \Rightarrow (B \wedge C)) = \neg(\neg A \vee (B \wedge C)) = \neg\neg A \wedge \neg(B \wedge C) = A \wedge (\neg B \vee \neg C) = (A \wedge \neg B) \vee (A \wedge \neg C)$. Dies ist allerdings noch nicht die disjunktive Normalform. Für diese ist zusätzlich gefordert, daß in jeder der Konjunktionen jede in der Gesamtaussage auftretende elementare Aussage vorkommt (negiert oder nicht). Dies erreicht man mit der allgemeingültigen Regel $A \Leftrightarrow (A \wedge B) \vee (A \wedge \neg B)$. Im Beispiel rechnen wir daher weiter: $= (A \wedge \neg B \wedge C) \vee (A \wedge \neg B \wedge \neg C) \vee (A \wedge B \wedge \neg C) \vee (A \wedge \neg B \wedge \neg C) = (A \wedge \neg B \wedge C) \vee (A \wedge \neg B \wedge \neg C) \vee (A \wedge B \wedge \neg C)$. Diese logische Formel ist jetzt eine disjunktive Normalform.

1.5.20 DEFINITION. Eine logische Formel ist in *disjunktiver Normalform*, wenn sie eine Disjunktion von Konjunktionen aller darin vorkommenden elementaren Aussagen oder deren Negationen ist.

1.5.21 SATZ. Die *disjunktive Normalform* ist, bis auf Umordnungen, eindeutig bestimmt.

Statt durch Ausmultiplizieren kann die disjunktive Normalform auch an der Wahrheitstafel abgelesen werden: Beispiel:

A	B	C	$B \wedge C$	$A \Rightarrow (B \wedge C)$	$\neg(A \Rightarrow (B \wedge C))$
⊤	⊤	⊤	⊤	⊤	⊥
⊤	⊤	⊥	⊥	⊥	⊤
⊤	⊥	⊤	⊥	⊥	⊤
⊤	⊥	⊥	⊥	⊥	⊤
⊥	⊤	⊤	⊤	⊤	⊥
⊥	⊤	⊥	⊥	⊤	⊥
⊥	⊥	⊤	⊥	⊤	⊥
⊥	⊥	⊥	⊥	⊤	⊥

Die drei Zeilen, in denen in der letzten Spalte ein $\top$ steht, entsprechen exakt der disjunktiven Normalform: $(A \wedge B \wedge \neg C) \vee (A \wedge \neg B \wedge C) \vee (A \wedge \neg B \wedge \neg C)$.

1.5.22 SATZ. In der Booleschen Logik sind die Aussagen $A \Rightarrow B$ und $\neg A \vee B$ gleichbedeutend.

Beweis. Übung! □

Jeder Term, der sich aus Variablennamen und den Operationssymbolen $\wedge, \vee, \neg$ syntaktisch korrekt zusammensetzt, heißt ein *Boolscher Ausdruck*.

1.5.23 THEOREM. *Sei $f: \mathbb{B}^n \rightarrow \mathbb{B}$ eine Funktion in beliebig vielen Variablen mit Werten in $\mathbb{B}$. Dann gibt es einen boolschen Ausdruck $t[x_1, \dots, x_n]$, sodaß $f(x_1, \dots, x_n) = t[x_1, \dots, x_n]$.*

Beweis. Ein passender Ausdruck (sogar in disjunktiver Normalform) kann unmittelbar aus der Wertetabelle abgelesen werden. □

1.5.24 BEMERKUNG. Dieser Satz hat die praktische Bedeutung, daß sich jede derartige Funktion (und damit eigentlich jede Funktion zwischen endlichen Mengen) aus einfachen Schaltkreisen für die logischen Operationen (UND-, ODER- und NICHT-Gatter) realisieren läßt. Freilich ist die Verwendung der disjunktiven Normalform für diesen Zweck nicht unbedingt sinnvoll; daher gibt es Verfahren, um die Anzahl bzw. Komplexität der verwendeten Bauelemente zu minimieren

Klassische Logik

Die *klassische Logik* zeichnet sich durch Annahme eines der folgenden Prinzipien aus.

1.5.25 AXIOM (REDUCTIO AD ABSURDUM). *Die aussagenlogische Formel $\neg\neg A \implies A$ ist allgemeingültig.*

1.5.26 AXIOM (TERTIUM NON DATUR). *Die aussagenlogische Formel $A \vee \neg A$ ist allgemeingültig.*

1.5.27 THEOREM. *Die beiden Prinzipien reductio ad absurdum und tertium non datur sind äquivalent.*

Beweis. Satz 1.5.9 beweist die eine und Satz 1.5.12 die andere Richtung. □

1.5.28 SATZ. *Bei Verwendung von klassischer Logik gilt stets*

$$\bigvee_{x: X} A(x) \iff \neg \bigwedge_{x: X} \neg A(x).$$

Beweis. Übung. □

1.5.29 BEMERKUNG. Die hier besprochenen Regeln gehen von der Idee aus, daß es eine absolute Wahrheit gibt, auch wenn sie uns Menschen verborgen bleibt. So ist es zwar möglich, daß eine Aussage weder beweisbar noch widerlegbar ist, aber im Prinzip muß sie doch wahr oder falsch sein (bzw.: wenn sie nicht falsch ist, so muß sie doch wahr sein).

Daneben hat die klassische Logik den Vorteil, daß sich einige Sätze der Mathematik einfacher formulieren oder beweisen lassen (z.B. durch Verzicht auf doppelte Negationen). Außerdem läßt sich die Gültigkeit aussagenlogische Formeln einfach mittels Wahrheitstabellen entscheiden.

Die überwältigende Mehrheit der mathematischen Literatur im 20. Jahrhundert wurde vom Standpunkt der klassischen Logik aus verfaßt, d.h. sie verwendet diese zusätzlichen Prinzipien ohne jeden Kommentar.

In neuerer Zeit gibt es aber immer mehr Anwendungen von nicht-klassischen Logik-Systemen. So lassen sich die in der Informatik so bedeutenden Begriffe der Entscheidbarkeit und der Berechenbarkeit in einer klassischen Logik nicht direkt ausdrücken, sodaß man gezwungen ist, diese Konzepte über die Hintertür (z.B. mittels Maschinenmodell) quasi in einer eigenen Welt künstlich einzuführen. In der Physik spielt die (dort so bedeutende) Frage der Meßbarkeit eine ähnliche Rolle.

Dies führt zu der Betrachtungsweise, daß ein Objekt, welches nicht berechnet (oder gemessen) werden kann, denselben Effekt hat wie eines, das gar nicht existiert.

In diesem Skriptum wird zumeist nicht mit der klassischen Logik, sondern mit der *konstruktiven* Logik (also ohne die soeben besprochenen zusätzlichen Prinzipien) gearbeitet, eben weil sich so einige wichtige Konzepte der Informatik treffender beschreiben lassen. Auf wesentliche Unterschiede zur klassischen Logik wird aber stets hingewiesen.

Kapitel 2

Mengen

2.1 Äquivalenzrelationen

Laut Cantor: *Eine Menge ist eine Zusammenfassung bestimmter wohlunterscheidbarer Objekte unserer Anschauung.*

Um zu bestimmen, welche Objekte in die Zusammenfassung aufgenommen werden sollen, ist festzulegen, wie diese zu konstruieren sind. Dies entspricht exakt unserer Definition des Begriffs *Aussage*.

Objekte, die zu einer Menge zusammengefaßt werden, sollten aber darüberhinaus auch *wohlunterscheidbar* sein. Das heißt, es muß ein vernünftiger Gleichheitsbegriff festgelegt werden. *Vernünftig* heißt hier, daß der Gleichheitsbegriff auf jeden Fall den Regeln für eine Äquivalenzrelation genügen muß.

2.1.1 DEFINITION. Eine binäre Relation $(=): X \rightarrow X \rightarrow \Omega$ heißt *Äquivalenzrelation*, wenn für alle $x, y, z: X$ gilt:

$$\begin{aligned}x &= x && (\text{reflexiv}); \\x = y &\iff y = x && (\text{symmetrisch}); \\x = y \wedge y = z &\implies x = z && (\text{transitiv}).\end{aligned}$$

2.1.2 BEMERKUNG. Um eine Menge vollständig zu definieren, ist also (im Prinzip) dreierlei zu tun:

- Festlegen, welche Objekte dazugehören;
- Eine Relation definieren;
- Beweisen, daß diese Relation tatsächlich die Eigenschaft einer Äquivalenzrelation erfüllt.

Diese Arbeit kann man sich natürlich sparen, wenn einfache Methoden zur Verfügung stehen, um aus bekannten Mengen neue zu konstruieren.

2.1.3 BEMERKUNG. Man beachte, daß je zwei Objekten lediglich eine Aussage zugeordnet werden muß; d.h. es muß klar sein, was bewiesen werden müßte, um zwei Objekte als gleich nachzuweisen. Es ist aber keinesfalls erforderlich, daß festgestellt werden kann, ob es so einen Beweis auch tatsächlich gibt.

2.1.4 NOTATION. Jede Menge kommt zusammen mit einem Gleichheitsbegriff. Mit $x = y : X$ ist gemeint, daß x und y gleiche Objekte von X bezeichnen. Oft ist aus dem Zusammenhang klar, daß x und y als Elemente von X zu betrachten sind, sodaß man stattdessen einfach $x = y$ schreibt. Die Notation $x : X$ kann auch als Abkürzung für $x = x : X$ aufgefaßt werden.

2.1.5 DEFINITION. Eine Menge X heißt *diskret*, wenn für alle $x : X$ gilt:

$$x = x \vee \neg(x = x).$$

2.1.6 BEMERKUNG. Diese Bezeichnung macht keinen Sinn, wenn die Regel vom ausgeschlossenen Dritten (klassische Logik) angenommen wird. Für diesen Fall, muß *diskret* anders (etwa mit einem Maschinenmodell) definiert werden.

2.1.7 BEISPIEL. Die Gesamtheit aller logischen Aussagen, zusammen mit der logischen Äquivalenz, bildet eine Menge (es war einfach, nachzuweisen, daß die logische Äquivalenz tatsächlich ein Äquivalenzrelation ist). Wir bezeichnen sie mit Ω , und nennen sie manchmal auch die Menge aller Wahrheitswerte. In klassischer Logik (d.h. mit *reductio ad absurdum*) gilt natürlich $\Omega = \mathbb{B}$.

2.1.8 BEISPIEL. Die wahrscheinlich fundamentalste aller Mengen ist die Menge der natürlichen Zahlen $\mathbb{N}$. Diese ist diskret.

2.2 Konstruktionen für Mengen

Funktionen

Sind A und B Mengen, so ist auch die Gesamtheit aller *Funktionen* zwischen ihnen eine Menge. Dies ist als Verfeinerung des Begriffs der Implikation von Aussagen aufzufassen und wird daher mit den entsprechenden Introduktions- und Eliminationsregeln beschrieben:

2.2.1 AXIOM ($\rightarrow$ -INTRODUKTIONSREGEL).

$$\frac{\begin{array}{c} x : A \\ \vdots \\ t[x] : B \end{array}}{x \mapsto t[x] : A \rightarrow B} \rightarrow \mathcal{I}$$

Genauer:

$$\frac{\begin{array}{c} x = y : A \\ \vdots \\ s[x] = t[y] : B \end{array}}{(x \mapsto s[x]) = (y \mapsto t[x]) : A \rightarrow B} \rightarrow \mathcal{I}$$

Dabei wird mit $t[x]$ der Term t bezeichnet und gleichzeitig ausgedrückt, daß darin an bestimmten Stellen ein x vorkommen kann.

Für die Konstruktion $x \mapsto t[x]$ ist auch die Notation $\lambda x. t[x]$ gebräuchlich, und sie heißt daher auch λ -Abstraktion. In Java würde man für $f = (x \mapsto t[x]) : A \rightarrow B$ in etwa folgendes schreiben:

```

B f (A x){
    return (t[x]);
}

```

(Für A und B sind natürlich konkrete Java-Datentypen einzusetzen, für $t[x]$ irgendein Ausdruck, in dem typischerweise das x vorkommt.)

2.2.2 AXIOM ($\rightarrow$ -ELIMINATIONSREGEL).

$$\frac{f : A \rightarrow B \quad x : A}{fx : B} \rightarrow \mathcal{E}$$

Genauer:

$$\frac{f = g : A \rightarrow B \quad x = y : A}{fx = gy : B} \rightarrow \mathcal{E}$$

Der zur λ -Abstraktion gehörige Selektor ist damit die Funktionsanwendung (Funktionsapplikation). Statt fx ist die Notation $f(x)$ sehr gebräuchlich, wenn deutlich gemacht werden soll, daß es sich um keine Multiplikation handelt. Für viele spezielle Funktionen sind auch Postfix, Infix, oder andere spezielle Notationen gebräuchlich.

Der Zusammenhang zwischen der λ -Abstraktion und der Funktionsapplikation ergibt sich wieder aus den passenden β und η -Regeln.

2.2.3 AXIOM (β -REGEL). *Ist $t[x] : B$, für jedes $x : A$, und $a : A$, so gilt*

$$(x \mapsto t[x])a = t[a] : B.$$

Dabei bezeichnet $t[a]$ jenen Ausdruck, welcher entsteht, wenn in $t[x]$ jedes Vorkommen von x durch a ersetzt wird.

2.2.4 AXIOM (η -REGEL). *Für $f : A \rightarrow B$ gilt*

$$(x \mapsto fx) = f : A \rightarrow B.$$

2.2.5 SATZ. *Zwei Funktionen $f, g : A \rightarrow B$ sind genau dann gleich, wenn $fx = gx : B$, für alle $x : A$.*

Dies bedeutet, daß es für die Gleichheit von Funktionen nicht entscheidend ist, wie diese konstruiert wurden, sondern ob sie bei jedem Input identische Ergebnisse liefern, also von *außen betrachtet* gleich erscheinen. (Auch eine eventuelle Laufzeit spielt für den Gleichheitsbegriff keine Rolle.)

Von der Implikation haben wir gezeigt, daß sie transitiv ist. Die entsprechende Konstruktion ist eine grundlegende Operation für Funktionen:

2.2.6 DEFINITION. Seien $f : A \rightarrow B$, $g : B \rightarrow C$. Dann heißt

$$g \circ f = (x \mapsto g(fx)) : A \rightarrow C$$

die *Hintereinanderausführung* (*Komposition*) der Funktionen f und g .

2.2.7 SATZ. *Die binäre Operation $(\circ) : (B \rightarrow C) \rightarrow (A \rightarrow B) \rightarrow (A \rightarrow C)$ ist assoziativ, d.h. für alle $f : A \rightarrow B$, $g : B \rightarrow C$, $h : C \rightarrow D$ gilt*

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

Insbesondere bildet daher die Menge der Funktionen $A \rightarrow A$, zusammen mit der Komposition, eine Halbgruppe; die identische Funktion $\text{id } A = (x \mapsto x) : A \rightarrow A$ ist das neutrale Element dieser Halbgruppe.

Direktes Produkt

Sind A und B Mengen, so ist auch die Gesamtheit aller Paare, die man daraus bilden kann, eine Menge. Dies ist als Verfeinerung des Begriffs der Konjunktion von Aussagen aufzufassen und wird daher, so wie diese, mit Introduktions- und Eliminationsregeln beschrieben:

2.2.8 AXIOM ($\times$ -INTRODUKTIONSREGEL).

$$\frac{x:A \quad y:B}{(x,y): A \times B} \times \mathcal{I} \quad (2.1)$$

Genauer:

$$\frac{x_1 = x_2 : A \quad y_1 = y_2 : B}{(x_1, y_1) = (x_2, y_2) : A \times B} \times \mathcal{I} \quad (2.2)$$

2.2.9 AXIOM ($\times$ -ELIMINATIONSREGEL).

$$\frac{z:A \times B}{\text{fst } z:A} \times \mathcal{E} \qquad \frac{z:A \times B}{\text{snd } z:B} \times \mathcal{E} \quad (2.3)$$

Genauer:

$$\frac{z_1 = z_2 : A \times B}{\text{fst } z_1 = \text{fst } z_2 : A} \times \mathcal{E} \qquad \frac{z_1 = z_2 : A \times B}{\text{snd } z_1 = \text{snd } z_2 : B} \times \mathcal{E} \quad (2.4)$$

Die Notation $(,)$ bezeichnet einen *Konstruktor*, während fst und snd *Selektoren* bezeichnen. Der Zusammenhang wird mit den folgenden Regeln beschrieben:

2.2.10 AXIOM (β -REGEL). Für $x:A$ und $y:B$ gilt

$$\text{fst}(x,y) = x : A \qquad \text{snd}(x,y) = y : B$$

2.2.11 AXIOM (η -REGEL). Für $z:A \times B$ gilt

$$(\text{fst } z, \text{snd } z) = z : A \times B$$

Damit ist garantiert, daß Paare genau dann gleich sind, wenn sie auf dieselbe Weise konstruiert wurden, d.h. wenn beide Komponenten übereinstimmen.

Direkte Summe

Sind A und B Mengen, so bildet auch die Gesamtheit derjenigen Objekte, welche zu einer dieser beiden Mengen gehören, ebenfalls eine Menge. Dies ist als eine Verfeinerung des Begriffs der Disjunktion aufzufassen und wird daher mit den entsprechenden Introduktions- und Eliminationsregeln beschrieben:

2.2.12 AXIOM ($+$ -INTRODUKTIONSREGELN).

$$\frac{x:A}{\text{Left } x:A+B} + \mathcal{I} \qquad \frac{x:B}{\text{Right } x:A+B} + \mathcal{I}$$

Genauer:

$$\frac{x_1 = x_2 : A}{\text{Left } x_1 = \text{Left } x_2 : A+B} + \mathcal{I} \qquad \frac{y_1 = y_2 : B}{\text{Right } y_1 = \text{Right } y_2 : A+B} + \mathcal{I}$$

Statt $A + B$ ist auch die Bezeichnung $A \uplus B$ üblich.

Die Verwendung von zwei verschiedenen Konstruktoren stellt sicher, daß von jedem Element von $A + B$ entschieden werden kann, ob es aus A oder aus B stammt, insbesondere auch dann, wenn A und B dieselbe Menge bezeichnen.

2.2.13 AXIOM (+-ELIMINATIONSREGEL).

$$\frac{f: A \rightarrow C \quad g: B \rightarrow C \quad z: A + B}{\text{either } f \text{ } g \text{ } z: C} + \mathcal{E}$$

Genauer:

$$\frac{f_1 = f_2: A \rightarrow C \quad g_1 = g_2: B \rightarrow C \quad z_1 = z_2: A + B}{\text{either } f_1 \text{ } g_1 \text{ } z_1 = \text{either } f_2 \text{ } g_2 \text{ } z_2: C} + \mathcal{E}$$

So wie bei der Disjunktion, haben wir auch hier eine etwas kompliziertere Eliminationsregel. Sie ist aber etwas besser zu verstehen, wenn man folgendes beachtet: Für jedes $z: A+B$ gilt either $f \text{ } g \text{ } z: C$; daher ist either $f \text{ } g: A+B \rightarrow C$. Der Selektor either entspricht damit einer Fallunterscheidung: either $f \text{ } g$ verhält sich auf A so wie f , auf B aber so wie g . Dies verdeutlicht vor allem die folgende Regel:

2.2.14 AXIOM (β -REGEL). Für $f: A \rightarrow C$, $g: B \rightarrow C$, $x: A$, $y: B$ gilt

$$\text{either } f \text{ } g \text{ (Left } x) = f x: C \qquad \text{either } f \text{ } g \text{ (Right } y) = g y: C$$

2.2.15 AXIOM (η -REGEL). Für $z: A + B$ gilt

$$\text{either Left Right } z = z: A + B.$$

2.3 Potenzmenge

Teilmengen

Es sei G eine Menge, die im folgenden Grundmenge genannt wird. Für eine Eigenschaft $P: G \rightarrow \Omega$ sollte $\{g: G \mid P(g)\}$ eine Menge bezeichnen, die all jene Elemente von G umfaßt, welche die Eigenschaft P besitzen.

2.3.1 DEFINITION. Um ein Element der Menge $\{g: G \mid P(g)\}$ zu konstruieren, muß man ein Element $g: G$ konstruieren und dann $P(g)$ beweisen. Es ist also ein Beweis $(g, \varphi): \bigvee_{g: G} P(g)$ zu konstruieren.

Als Gleichheitsbegriff legen wir fest

$$(g_1, \varphi_1) = (g_2, \varphi_2): \{g: G \mid P(g)\} : \iff g_1 = g_2: G,$$

was tatsächlich eine Äquivalenzrelation ist.

Die Menge $\{g: G \mid P(g)\}$ nennt man eine *Teilmenge* von G .

2.3.2 BEMERKUNG. Ist $(g, \varphi): \{g: G \mid P(g)\}$, so ist wegen dem gewählten Gleichheitsbegriff der Beweis φ irrelevant; relevant ist nur, daß es einen solchen gibt. Dennoch ist es nicht ganz korrekt, wenn man $g: \{g: G \mid P(g)\}$ schreibt. Wir definieren deshalb ein geeignetes 2-stelligen Prädikat.

2.3.3 DEFINITION. Ist $(g, \varphi): \{g: G \mid P(g)\}$, dann sagen wir, daß g in der Teilmenge $\{g: G \mid P(g)\}$ enthalten ist und schreiben für diesen Sachverhalt $g \in \{g: G \mid P(g)\}$. D.h.

$$g \in \{g: G \mid P(g)\} : \Longleftrightarrow P(g).$$

Damit definieren wir eine Ordnungsrelation für Teilmengen.

2.3.4 DEFINITION. Seien A und B Teilmengen von G . Dann definieren wir

$$A \subseteq B : \Longleftrightarrow \bigwedge_{g \in G} (g \in A \implies g \in B),$$

welche reflexiv und transitiv ist, und gemäß Antisymmetrie den folgenden Gleichheitsbegriff

$$A = B \Longleftrightarrow A \subseteq B \wedge B \subseteq A,$$

ergibt, welcher bedeutet, daß zwei Teilmengen genau dann gleich sind, wenn sie dieselben Elemente enthalten. Damit ist auch die Gleichheit von Teilmengen festgelegt; Die (geordnete) Menge aller Teilmengen von G heißt deren *Potenzmenge* und wird mit $\mathcal{P}G$ bezeichnet.

2.3.5 BEMERKUNG. Leicht anders formuliert haben wir

$$A = B : \mathcal{P}(G) \Longleftrightarrow \bigwedge_{g: G} (g \in A \Longleftrightarrow g \in B).$$

Dies bedeutet, daß zwei Teilmengen genau dann als gleich betrachtet werden, wenn sie dieselben Elemente haben. Man beachte aber, daß dies nur dann wirklich Sinn macht, wenn klar ist, auf welche Weise die betroffenen Mengen in einer gemeinsamen Grundmenge enthalten sind.

Mengenalgebra

Sind A, B zwei Mengen, die nichts miteinander zu tun haben, und ist $a: A$, $b: B$, dann macht es im allgemeinen keinen Sinn zu fragen, ob a ein Element von B ist, oder ob $a = b$. Es muß zuerst festgelegt werden, wie die Elemente von A und B zu identifizieren sind. Dies ist insbesondere dann klar, wenn beide Teilmengen einer gemeinsamen Grundmenge sind. Dann lassen sich die üblichen Mengenoperationen definieren:

2.3.6 NOTATION. Statt $\neg(g \in G)$ schreibt man üblicherweise $g \notin G$.

2.3.7 DEFINITION. Seien $A, B: \mathcal{P}(G)$. Mengentheoretische Vereinigung, Durchschnitt und Differenz werden definiert durch

$$A \cup B = \{g: G \mid g \in A \vee g \in B\},$$

$$A \cap B = \{g: G \mid g \in A \wedge g \in B\},$$

$$A \setminus B = \{g: G \mid g \in A \wedge g \notin B\}.$$

Eine spezielle Rolle spielen die *leere Menge* $\emptyset = \{g: G \mid \perp\}$ und die *Gesamtmenge* $\{g: G \mid \top\}$, welche üblicherweise mit der Grundmenge G selbst identifiziert wird. Ist die Grundmenge aus dem Zusammenhang klar, definiert man auch das mengentheoretische Komplement

$$\complement A = G \setminus A \quad (\text{Komplement}).$$

Für das Komplement sind auch die Notationen A' und $\overline{A}$ gebräuchlich.

Ist $G \subseteq H$, so kann jede Teilmenge von G auch als Teilmenge von H aufgefaßt werden. Für die meisten der oben definierten Operationen spielt dies keine Rolle, wohl aber beim Komplement: denn wäre $T: \mathcal{P}H$, dann wäre $\mathcal{C}T = H \setminus T$, und nicht $G \setminus T$.

2.3.8 BEMERKUNG. Identifiziert man G mit $\{g: G \mid \top\}$, so verschwindet formal der Unterschied zwischen $g: G$ und $g \in G$. Ersteres ist eine deklarative Variante, d.h. g wird als Element von G eingeführt. Letzteres dagegen ist eine Aussage, welche gegebenenfalls zu beweisen ist.

Die Rechenregeln für diese Operationen ergeben sich ziemlich direkt aus den entsprechenden Regeln für Aussagen. Im folgenden beschäftigen wir uns aber nur mit einer speziellen Klasse von Teilmengen.

Ohne *reductio ad absurdum* kann es natürlich passieren, daß weder $g \in B$ noch $g \notin B$ bewiesen werden kann.

2.3.9 DEFINITION. Sei G eine Menge. Ein $A: \mathcal{P}(G)$ heißt *entscheidbar* wenn für alle $g: G$ gilt

$$g \in A \vee g \notin A.$$

Diese Definition macht natürlich nur Sinn, wenn *reductio ad absurdum* nicht generell angenommen wird. In klassischer Logik definiert man entscheidbare Teilmengen daher anders (nämlich, als mittels einer abstrakt definierten Maschine entscheidbar).

2.3.10 SATZ. *Die entscheidbaren Teilmengen sind gerade diejenigen, welche sich als*

$$A = \{g: G \mid P(g)\},$$

darstellen lassen, mit $P: G \rightarrow \mathbb{B}$ (nicht nur mit $P: G \rightarrow \Omega$).

2.3.11 BEMERKUNG. Hier ist vor allem gemeint, daß das P tatsächlich berechenbar ist.

2.3.12 SATZ. *Die entscheidbaren Teilmengen bilden mit Vereinigung, Durchschnitt und Komplement eine Boolsche Algebra.*

Beweis. Alle erforderlichen Gleichheiten ergeben sich unmittelbar aus der Definition und der entsprechenden Gleichheit in der Boolschen Logik. $\square$

2.3.13 BEMERKUNG. Die Menge der entscheidbaren Teilmengen ist schon wesentlich besser handhabbar als die ganze Potenzmenge, aber für den praktischen Einsatz, etwa in Datenbanken, immer noch zu groß bzw. kompliziert. Daher beschränkt man sich bei derartigen Anwendungen üblicherweise auf die Menge aller endlichen Teilmengen. Allerdings bilden diese keine Boolsche Algebra mehr, da das Komplement einer endlichen Menge nicht mehr endlich sein muß (wenn die Grundmenge unendlich ist). Man kann aber stattdessen die Mengendifferenz verwenden (die Differenz endlicher Teilmengen einer diskreten Menge ist wieder endlich). Die Gesetze von De Morgan müssen ersetzt werden durch

$$\begin{aligned} A \setminus (B \cup C) &= (A \setminus B) \cap (A \setminus C), \\ A \setminus (B \cap C) &= (A \setminus B) \cup (A \setminus C); \end{aligned}$$

und die Regeln für das Komplement durch

$$A \cap (B \setminus A) = \emptyset$$

$$A \cup (B \setminus A) = B.$$

Eine solche algebraische Struktur heißt *Boolscher Ring*. Mit endlichen Teilmengen einer diskreten Menge können alle relevanten Operationen und Vergleiche effektiv durchgeführt werden (für bloß entscheidbare Teilmengen hat man dagegen z.B. kein allgemeines Verfahren, um festzustellen, ob eine solche leer ist. In der Theorie der formalen Sprachen werden diverse Zwischenstufen diskutiert. So bilden etwa die regulären Sprachen eine Boolsche Algebra, die auch unendliche Mengen umfaßt und in der alles entscheidbar und berechenbar bleibt.

Beispiele

Sei $(-0): \mathbb{N} \rightarrow \mathbb{Z}$ definiert durch $(-0)(x) = x - 0$. Damit wird eine Teilmenge von $\mathbb{Z}$ definiert, nämlich $\{x - 0 \mid x: \mathbb{N}\}: \mathscr{P}\mathbb{Z}$, die positiven ganzen Zahlen. Die Teilmenge der negativen ganzen Zahlen erhält man analog mit $(0-): \mathbb{N} \rightarrow \mathbb{Z}$, $(0-)(x) = 0 - x$ als $\{0 - x \mid x: \mathbb{N}\}: \mathscr{P}\mathbb{Z}$. Die Quadratzahlen erhalten wir mit $\{x^2 \mid x: \mathbb{N}\}$, aber auch z.B. mit $\{x^2 \mid x: \mathbb{Z}\}$. Die Abbildung (-0) ist besonders natürlich, insbesondere ist sie eine injektive Funktion, d.h. es gilt

$$x = y : \mathbb{N} \iff x - 0 = y - 0 : \mathbb{Z};$$

außerdem ist sie mit der Addition verträglich:

$$(x - 0) + (y - 0) = (x + y) - 0,$$

und auch mit der Multiplikation:

$$(x - 0)(y - 0) = xy - 0.$$

Da es nur eine derart natürliche Funktion von $\mathbb{N}$ nach $\mathbb{Z}$ gibt, sind wir geneigt, $\mathbb{N}$ selbst als Teilmenge von $\mathbb{Z}$ anzusehen. Die dabei unterstellte natürliche Funktion gibt dann an, wie die Elemente dabei zu *identifizieren* sind, nämlich als

$$x = x - 0 : \mathbb{Z}, \quad \text{für } x \in \mathbb{N}.$$

Man schreibt daher auch $\mathbb{N} = \{x - 0 \mid x: \mathbb{Z}\}$; oder $x: \mathbb{Z}$, falls $x: \mathbb{N}$, und daher auch $\mathbb{N}: \mathscr{P}\mathbb{Z}$. Zu Mißverständnissen kann es dabei kaum kommen, weil die Funktion (-0) mit allen Operationen verträglich ist.

Auch die Funktion $(0-)$ ermöglicht es, $\mathbb{N}$ mit den negativen ganzen Zahlen zu identifizieren. Dies geht jedoch nur so lange gut, als keine Multiplikation vorkommt: werden $x, y: \mathbb{N}$ als natürliche Zahlen multipliziert und dann mittels $(0-)$ mit einer ganzen Zahl identifiziert, erhält man $0 - xy$; multipliziert man dagegen die entsprechenden ganzen Zahlen, so erhält man $(0 - x)(0 - y) = xy - 0$, was im allgemeinen nicht dasselbe ist. Diese Identifikation ist daher nicht natürlich.

Auf analoge Weise kann man $\mathbb{Z}$ als Teilmenge von $\mathbb{Q}$ auffassen, mit $z = \frac{z}{1} : \mathbb{Q}$. Diese Identifikation mag (so wie die vorherige) recht trivial erscheinen. Man beachte aber, daß in vielen Programmiersprachen sehr wohl zwischen der ganzen Zahl 3 und der rationalen Zahl $\frac{3}{1}$, dargestellt als Paar von ganzen Zahlen oder als Gleitkommazahl, zu unterscheiden ist. Die Identifikation geschieht dann durch implizite oder explizite Typumwandlung.

2.4 Gleichmächtigkeit

Es ist für eine Menge eigentlich völlig irrelevant, welche konkreten Objekte sie enthält. So besteht die Menge der natürlichen Zahlen konzeptionell aus den Objekten $O, SO, S(SO), S(S(SO)), \dots$. Praktischerweise werden sie aber meist als Dezimalzahlen $0, 1, 2, 3, \dots, 756, \dots$ notiert, also durch bestimmte Zeichenketten. In der Informatik wählt man statt dessen gerne das Binärsystem, d.h. man stellt Zahlen durch Zeichenketten über $\{0, 1\}$ dar, welche im Computer wiederum durch Folgen elektrischer oder magnetischer Zustände realisiert werden. Dabei ist stets klar, wie zwischen diesen Darstellungen hin und her gewechselt werden kann.

2.4.1 DEFINITION. Eine Funktion $f: A \rightarrow B$ heißt *bijektiv* (oder eine *Bijektion*), wenn es eine Funktion $g: B \rightarrow A$ gibt, sodaß

$$g \circ f = \text{id}_A \qquad f \circ g = \text{id}_B.$$

Die Funktion g ist, wenn sie existiert, durch diese Eigenschaft eindeutig bestimmt, heißt die *Umkehrfunktion* von f und wird oft mit f^{-1} bezeichnet.

Eine Bijektion bewirkt damit so etwas wie eine andere Darstellung (Umbenennung, Codierung) der Elemente von A durch die Elemente von B . Wegen der Umkehrbarkeit kann man jederzeit wieder zu den ursprünglichen Elementen von A zurückkehren. So ist z.B. die Darstellung von abstrakten oder realen Objekten in einem Computer eine Bijektion zwischen der Menge der Objekte, die man eigentlich betrachtet, und einer bestimmten Menge von (irgendwie realisierten) Bitfolgen.

2.4.2 DEFINITION. Zwei Mengen A, B heißen *gleichmächtig* wenn es eine Bijektion $f: A \rightarrow B$ gibt. Man schreibt dann

$$|A| = |B|.$$

2.4.3 SATZ. *Gleichmächtigkeit ist eine Äquivalenzrelation für Mengen.*

Beweis. Wir müssen überprüfen, ob sie reflexiv, symmetrisch und transitiv ist:

Reflexivität: $\text{id}: A \rightarrow A$ ist bijektiv, denn

$$\text{id}_A^{-1} = \text{id}_A.$$

Daher ist $|A| = |A|$.

Symmetrie: Sei $|A| = |B|$. Dann gibt es laut Definition eine Bijektion $f: A \rightarrow B$, und auch $f^{-1}: B \rightarrow A$ ist ebenfalls bijektiv, denn es gilt

$$(f^{-1})^{-1} = f.$$

Daher ist $|B| = |A|$.

Transitivität: Sei $|A| = |B|$ und $|B| = |C|$. Laut Definition gibt es daher Bijektionen $f: A \rightarrow B$ und $g: B \rightarrow C$. Dann ist aber auch $g \circ f: A \rightarrow C$ eine Bijektion, denn es gilt

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

Daher ist $|A| = |C|$.

Damit ist die Gleichmächtigkeit ein zulässiger Gleichheitsbegriff für Mengen.

2.4.4 SATZ. Sind A und A' gleichmächtig, und ebenso B und B' , also $|A| = |A'|$ und $|B| = |B'|$, so gelten auch

$$\begin{aligned} |A \times B| &= |A' \times B'|; \\ |A + B| &= |A' + B'|; \\ |A \rightarrow B| &= |A' \rightarrow B'|; \\ |\mathcal{P}(A)| &= |\mathcal{P}(A')|. \end{aligned}$$

Damit ist die Gleichmächtigkeit mit den gängigen Operationen für Mengen verträglich und scheint tatsächlich ein sinnvoller Gleichheitsbegriff für Mengen zu sein. Allerdings bezeichnet man die Mengen mit diesem Gleichheitsbegriff üblicherweise als *Kardinalzahlen*, damit es zu keiner Verwechslung mit dem Gleichheitsbegriff für Teilmengen kommt, die nur dann als gleich betrachtet werden, wenn sie dieselben Elemente enthalten (was aber nur dann Sinn macht, wenn klar ist, auf welche Weise sie in einer gemeinsamen Grundmenge eingebettet sind; siehe Abschnitt über Potenzmengen).

2.4.5 DEFINITION. Eine Funktion $f: A \rightarrow B$ heißt

1. *injektiv* falls für alle $x, y: A$ gilt

$$fx = fy : B \implies x = y : A;$$

2. *surjektiv* falls es zu jedem $y: B$ ein $x: A$ gibt sodaß

$$y = fx : B.$$

Man beachte, daß diese beiden Eigenschaften in einem gewissen Sinne dual zu den beiden Eigenschaft sind, die jede Funktion erfüllen muss: die Injektivität ist dual zur Wohldefiniertheit ($x = y : A \implies fx = fy : B$) ist; und die Surjektivität ist dual zur Totalität von Funktionen (zu jedem $x: A$ gibt es ein $y: B$, sodaß $fx = y : B$).

2.4.6 SATZ. Sei $f: A \rightarrow B$. Dann gilt

1. f ist injektiv wenn es ein $g: B \rightarrow A$ gibt, sodaß $g \circ f = \text{id } A$.
2. f ist surjektiv wenn es ein $g: B \rightarrow A$ gibt, sodaß $f \circ g = \text{id } B$.
3. f ist genau dann bijektiv wenn sie injektiv und surjektiv ist.

Man nennt daher eine injektive Funktion auch *links-invertierbar* und eine surjektive Funktion auch *rechts-invertierbar*. Bijektive Funktionen sind hiermit (beidseitig) *invertierbar*.

2.4.7 DEFINITION. Sei A eine Menge. Dann heißt A

1. *endlich* falls es ein $n: \mathbb{N}$ gibt sodaß

$$|A| = |\{0, \dots, n-1\}|;$$

2. *unendlich* falls es eine Injektion $\mathbb{N} \rightarrow A$ gibt (oder, wenn es zu jeder endlichen Teilmenge ein weiteres Element gibt).
3. *abzählbar unendlich* falls

$$|A| = |\mathbb{N}|;$$

4. in B einbettbar ($|A| \leq |B|$) falls es eine Injektion von A nach B gibt.

Statt $|\{0, \dots, n-1\}|$ kann man auch einfach n schreiben. Weiters werden gelegentlich verwendet: $|\mathbb{N}| = \aleph_0$ und $|\mathbb{R}| = c$.

2.4.8 SATZ.

1. Für $n, m: \mathbb{N}$ gilt

$$|\{0, \dots, n-1\}| = |\{0, \dots, m-1\}| \iff n = m;$$

2. Für endliche Mengen A, B gilt

$$|A \times B| = |A| \cdot |B|$$

$$|A + B| = |A| + |B|,$$

$$|A \rightarrow B| = |B|^{|A|};$$

daher verwendet man auch die Notation B^A für $A \rightarrow B$.

3. Ist die Menge A in die unendliche Menge B einbettbar, so gilt

$$|B| = |A \times B|, \quad (\text{falls } |A| \geq 1),$$

$$|B| = |A + B|,$$

$$|B| \leq |A \rightarrow B|;$$

4. Falls B mindestens 2 verschiedene Elemente hat, dann gibt es keine Surjektion von A nach $A \rightarrow B$.

2.4.9 BEISPIEL. $0 < 1 < \dots < n = |\{0, \dots, n-1\}| = |\{1, \dots, n\}| = |\{2, 4, \dots, 2n\}| < |\mathbb{N}| = |\mathbb{Z}| = |\mathbb{Q}| = |\mathbb{Q} \times \mathbb{Q}| = |\mathbb{Q}^n| = |\mathbb{Q}^*| = |\mathbb{A}|$. Die Mengen $\mathbb{R}, \mathbb{Q} \rightarrow \mathbb{B}, \mathbb{R} \rightarrow \mathbb{R}$ sind dagegen nicht abzählbar.

Weiters: $|A^\emptyset| = 1$, sogar wenn $A = \emptyset$, d.h. $|\emptyset^\emptyset| = 1$. Dagegen ist $|\emptyset^A| = 0$, sofern $|A| \geq 1$.

Um zu sehen, daß $\mathbb{N} \rightarrow \mathbb{B}$ nicht abzählbar ist, gehen wir folgendermaßen vor: Sei $f: \mathbb{N} \rightarrow (\mathbb{N} \rightarrow \mathbb{B})$. Dann definieren wir $g: \mathbb{N} \rightarrow \mathbb{B}$ durch

$$g(n) = \neg(f(n)(n)).$$

Wenn f surjektiv ist, dann gibt es ein $n: \mathbb{N}$, sodaß $g = f(n)$. Dann ist aber auch $g(n) = f(n)(n)$, was genau der Definition von g widerspricht. Es kann daher keine Surjektion von $\mathbb{N}$ nach $\mathbb{N} \rightarrow \mathbb{B}$ geben.

2.5 Konstruktion der Zahlenmengen

Ganze Zahlen

Es gibt zwei unterschiedliche Möglichkeiten, die ganzen Zahlen auf die natürlichen Zahlen zurückzuführen.

1. Variante

Die erste funktioniert so ähnlich wie das Geldwesen: alle Geldscheine sind zwar positiv, können aber in beide Richtungen wandern (Einnahmen/Ausgaben).

2.5.1 AXIOM ($\mathbb{Z}$ -INTRODUKTIONSREGEL).

$$\frac{x : \mathbb{N} \quad y : \mathbb{N}}{x - y : \mathbb{Z}} \mathbb{Z}\mathcal{I} \quad (2.5)$$

Der Ausdruck $x - y$ ist hier rein formal zu verstehen (Konstruktor, keine Rechnung) und bedeutet bei der obigen Interpretation: Einnahmen in der Höhe von x Euro stehen Ausgaben in der Höhe von y Euro gegenüber.

Diese Regel hat dieselbe Form wie die Introduktionsregel für das direkte Produkt. Eine ganze Zahl wird durch zwei natürliche Zahlen dargestellt. Entsprechend der Tatsache, daß höhere Einnahmen durch gleichviel höhere Ausgaben ausgeglichen werden (den selben Gewinn ergeben), verwenden wir allerdings als Gleichheitsbegriff nicht den für Paare, sondern den folgenden.

2.5.2 AXIOM. Für $x_1, y_1, x_2, y_2 : \mathbb{N}$ definieren wir

$$x_1 - y_1 = x_2 - y_2 : \mathbb{Z} \iff x_1 + y_2 = x_2 + y_1 : \mathbb{N}. \quad (2.6)$$

Damit wurde der Gleichheitsbegriff für ganze Zahlen auf den für die natürlichen Zahlen zurückgeführt. Freilich muß nachgewiesen werden, daß dieser Gleichheitsbegriff tatsächlich eine Äquivalenzrelation ist, was aber nicht schwer ist.

Die Addition läßt sich ebenfalls leicht auf $\mathbb{Z}$ definieren (wie bei einer Zusammenlegung von 2 Konten):

2.5.3 DEFINITION. Für $x_1 - y_1 : \mathbb{Z}$ und $x_2 - y_2 : \mathbb{Z}$ definieren wir

$$(x_1 - y_1) + (x_2 - y_2) = (x_1 + x_2) - (y_1 + y_2) : \mathbb{Z}.$$

Damit wird $(\mathbb{Z}, +)$ nicht nur zu einer Halbgruppe oder einem Monoid (mit $0 - 0$ als neutralem Element), sondern auch zu einer Gruppe, mit inversem Element:

$$-(x - y) = y - x.$$

Und mit der üblichen Multiplikation

$$(m_1 - n_1)(m_2 - n_2) = (m_1 m_2 - n_1 n_2) - (m_1 n_2 + m_2 n_1) : \mathbb{Z}$$

erhält man sogar einen Ring.

Eine echte Rechnung ergibt sich, wenn man einen Kontoausgleich vornimmt:

2.5.4 SATZ. Zu jedem $z : \mathbb{Z}$ gibt es ein eindeutiges $n : \mathbb{N}$ sodaß

$$z = n - 0 : \mathbb{Z} \vee z = 0 - n : \mathbb{Z}.$$

Damit erhält jede ganze Zahl eine kanonische Darstellung. Das hier konstruierte n wird üblicherweise mit $|z|$ bezeichnet. Den Ausdruck $0 - n$ kürzt man üblicherweise mit $-n$ ab, während $n - 0$ als $+n$ oder, noch einfacher, als n geschrieben wird. Die ganze Zahl $0 - 0$ kann auf beide Arten interpretiert werden,

daher ist $-0 = +0 : \mathbb{Z}$ Das Vorzeichen gibt an, welcher der Fälle für ein $z : \mathbb{Z}$ eintritt:

$$\text{sgn } z = \begin{cases} +1 & \text{falls } z = n - 0, \text{ aber nicht } z = 0 - n; \\ -1 & \text{falls } z = 0 - n, \text{ aber nicht } z = n - 0; \\ 0 & \text{falls } z = n - 0 \text{ und } z = 0 - n. \end{cases}$$

Es gilt stets $z = |z| \cdot \text{sgn } z$.

2. Variante

Die andere Art der Konstruktion der ganzen Zahlen aus den natürlichen geht von der zuletzt erwähnten Notation aus und verwendet eine Vorgangsweise wie bei der direkten Summe:

2.5.5 AXIOM ($\mathbb{Z}$ -INTRODUKTIONSREGELN).

$$\frac{n : \mathbb{N}}{+n : \mathbb{Z}} \mathbb{Z}\mathcal{I} \qquad \frac{n : \mathbb{N}}{-n : \mathbb{Z}} +\mathcal{I}$$

Die Gleichheit wird dann ähnlich wie bei der direkten Summe festgelegt

$$\begin{aligned} +m = +n : \mathbb{Z} &\iff m = n : \mathbb{N} \\ -m = -n : \mathbb{Z} &\iff m = n : \mathbb{N} \end{aligned}$$

aber

$$+m = -n : \mathbb{Z} \iff m = 0 = n : \mathbb{N}.$$

(im Gegensatz zu $\text{Left } x = \text{Right } y : A + B \iff \perp$).

Die üblichen Operationen lassen sich dann durch Fallunterscheidung (either) definieren. Dabei ist zu beachten, daß für $f, g : \mathbb{N} \rightarrow C$ die Konstruktion either $f g$ nur dann eine Funktion $\mathbb{Z} \rightarrow C$ ergibt, wenn $f(0) = g(0)$. Beispiel:

$$|.| = \text{either id id},$$

wobei $\text{id} : \mathbb{N} \rightarrow \mathbb{N}$ die identische Funktion bezeichnet. Oder, ausgeschrieben:

$$\begin{aligned} |+n| &= n : \mathbb{N}, \\ |-n| &= n : \mathbb{N}, \end{aligned}$$

für alle $n : \mathbb{N}$.

Ähnlich ergibt sich die Negation

$$(-) = \text{either}(-)(+) : \mathbb{Z} \times \mathbb{Z}$$

Dabei bezeichnet das $(-)$ auf der linken Seite die unäre Operation, während die Zeichen rechts die Konstruktoren bezeichnen. Ausgeschrieben wird es ein wenig deutlicher:

$$\begin{aligned} -(+n) &= -n : \mathbb{Z} \\ -(-n) &= +n : \mathbb{Z} \end{aligned}$$

Bei der Addition hat man vier Fälle zu unterscheiden. Die binäre Subtraktion ergibt sich dann als $a - b = a + (-b) : \mathbb{Z}$.

Der Minusoperator ist hier schon ziemlich überladen. Er hat mehrere Bedeutungen, abhängig von den Datentypen der Objekte, auf die er angewandt wird:

- $(-): \mathbb{N} \rightarrow \mathbb{N} \rightarrow \mathbb{Z}$: binärer Konstruktor, gemäß 1. Konstruktionsmethode;
- $(-): \mathbb{Z} \rightarrow \mathbb{Z}$: unäre Operation;
- $(-): \mathbb{Z} \rightarrow \mathbb{Z} \rightarrow \mathbb{Z}$: binäre Operation;
- $(-): \mathbb{N} \rightarrow \mathbb{Z}$: unäre Funktion bzw. Konstruktor nach der zweiten Konstruktionsmethode.

Die Überladung ist hier kein Problem, weil sie mit der natürlichen Einbettung der von $\mathbb{N}$ in $\mathbb{Z}$ stets kompatibel ist.

Rationale Zahlen

Die Konstruktion der rationalen Zahlen aus den ganzen erfolgt fast genauso wie die der ganzen aus den natürlichen (erste Variante). Der einzige wesentliche Unterschied besteht darin, daß 0 nicht als Nenner auftreten darf.

2.5.6 AXIOM ($\mathbb{Q}$ -INTRODUKTIONSREGEL).

$$\frac{p: \mathbb{Z} \quad q: \mathbb{Z} \quad q \neq 0: \mathbb{Z}}{\frac{p}{q}: \mathbb{Z}} \mathbb{QI} \quad (2.7)$$

Die Gleichheit wird dann genau analog definiert.

2.5.7 DEFINITION. Für $\frac{p_1}{q_1}: \mathbb{Q}$ und $\frac{p_2}{q_2}: \mathbb{Q}$ definieren wir

$$\frac{p_1}{q_1} = \frac{p_2}{q_2}: \mathbb{Q} \iff p_1 q_2 = p_2 q_1: \mathbb{Z}.$$

Die Definition der Multiplikation entspricht dann exakt der Definition der Addition für ganze Zahlen, während für die Addition die etwas kompliziertere Version mit dem gemeinsamen Nenner erforderlich ist.

Dem Kontoausgleich bei den ganzen Zahlen entspricht bei den rationalen Zahlen das Kürzen. Entsprechend erhält man für jede rationale Zahl durch Kürzen eine kanonische Form,

2.5.8 SATZ. Zu jeder rationalen Zahl $r: \mathbb{Q}$ gibt es genau ein Paar von ganzen Zahlen $p, q: \mathbb{Z}$ sodaß

$$r = \frac{p}{q}: \mathbb{Q} \quad q > 0 \quad p, q \text{ sind teilerfremd.}$$

Auch eine der 2. Konstruktionsmethode für ganze Zahlen entsprechende Variante gibt es hier: dabei rechnet man ausschließlich mit gekürzten Brüchen.

Ähnlich wie das Minuszeichen bei den ganzen Zahlen, ist hier der Bruchstrich überladen: Als $\mathbb{Z} \rightarrow \mathbb{Z} \rightarrow \mathbb{Q}$ ist er ein Konstruktor, als $\mathbb{Q} \rightarrow \mathbb{Q} \rightarrow \mathbb{Q}$ ist er eine binäre Operation (Doppelbruch).

Die rationalen Zahlen sind das Standardbeispiel für einen Körper; d.h. alle vier Grundrechnungsarten können wie gewohnt ausgeführt werden.

Irrationale Zahlen

Schon seit der Antike ist bekannt: es läßt sich kein $r: \mathbb{Q}$ finden, sodaß $r^2 = 2$, was geometrisch bedeutet, daß sich kein Zahlenverhältnis gibt, welches die Diagonale eines Quadrats mit dessen Seitenlänge in Verbindung bringt. Dies ist leicht zu erkennen. Sei etwa $r = \frac{p}{q}$, dann ist $r^2 = 2$ gleichbedeutend zu $p^2 = 2q^2$, $q \neq 0$. Die höchste Zweierpotenz, welche p^2 teilt, ist gerade; dasselbe gilt für q^2 , weshalb sie für $2q^2$ ungerade sein muß. Weil aber $2q^2 = p^2$ ist, folgt ein Widerspruch.

Man kann das Problem lösen, indem man $\sqrt{2}$ als weitere (irrationale) Zahl zuläßt, unter Fortführung der üblichen Rechengesetze, was durchaus möglich ist. Allerdings fehlen dann immer noch die Zahlen $\sqrt{3}$, $\sqrt[3]{2}$. Man kann auch als Zahlen alle Lösungen von Gleichungen beliebig hohen Grades mit rationalen Koeffizienten zulassen. Die Gesamtheit aller Lösung solcher Gleichungen bildet den Körper $\mathbb{A}$ der algebraischen Zahlen. Auch in diesem können alle Grundrechnungsarten durchgeführt werden, allerdings nur mit großem Aufwand. Aber selbst unter diesen gibt es z. B. keine Zahl, welche das Verhältnis des Umfangs eines Kreises zu dessen Diagonale beschreibt. D.h., π ist keine algebraische Zahl, sondern *transzendent* (konkret: $a_n \pi^n + \dots + a_0 = 0$ ist nur möglich, wenn alle Koeffizienten $a_i = 0$). Dasselbe gilt für die Eulersche Zahl e und viele weitere Zahlen, welche bei der Bildung von Grenzwerten auftreten. Es gibt also immer noch „Lücken“ auf der Zahlengeraden.

Reelle Zahlen

Alle Lücken auf der Zahlengeraden lassen sich schließen, wenn man die Lücken selbst als Zahlen auffaßt. Dies führt zur Konstruktion der reellen Zahlen $\mathbb{R}$. Auf die Details sei hier verzichtet. Wesentlich ist, daß man jede reelle Zahl mit rationalen Zahlen beliebig genau approximieren kann. D.h.: Sei $x: \mathbb{R}$; dann gibt es zu jedem $\varepsilon > 0$ ein $q: \mathbb{Q}$, sodaß

$$|x - q| \leq \varepsilon$$

Ist etwa $\varepsilon = 10^{-5}$, so hat man eine Approximation auf 5 Nachkomma-Dezimalstellen. Es bietet sich daher an, reelle Zahlen als Folge immer besser werdender Approximationen darzustellen. Auf diese Weise entsteht erneut ein Körper, d.h. auch innerhalb der reellen Zahlen können alle Grundrechnungsarten uneingeschränkt und mit den gewohnten Rechengesetzen ausgeführt werden. Ebenso die diversen Grenzwertübergänge, die in der Infinitesimalrechnung unerlässlich sind. Freilich muß man sich diesen Fortschritt teuer erkaufen: für reelle Zahlen ist die Gleichheit unentscheidbar, d.h. es gibt kein allgemeines Verfahren, welche von zwei reellen Zahlen feststellt, ob sie gleich sind oder nicht; oder: die Aussage

$$x = y : \mathbb{R} \text{ und } x \neq y : \mathbb{R}$$

ist ohne Prinzip des ausgeschlossenen Dritten nicht mehr allgemeingültig (d.h. die Gleichheit von reellen Zahlen ist nicht entscheidbar), und das exakte Rechnen mit reellen Zahlen ist sehr rechenaufwendig (aber nicht unmöglich!!!). Computer verwenden daher heutzutage immer noch Gleitkommazahlen (floating point) statt denn echten reellen Zahlen. Diese bestehen nur aus einer einzigen Approximation (etwa auf 15 Stellen). Auf diese Weise kann schnell gerechnet werden. Allerdings können sich die Rundungsfehler sehr übel auswirken, was zu

bösen Überraschungen führen kann. Insbesondere erfüllen die Gleitkommazahlen *nicht* die üblichen Rechengesetze, vor allem nicht das Assoziativgesetz. Die Nicht-Entscheidbarkeit der Gleichheit für reelle Zahlen wirkt sich so aus, daß Gleitkommazahlen, die sehr nahe beisammen liegen (innerhalb des Rundungsfehlers, mit dem zu rechnen ist), überhaupt nicht verglichen werden können (bzw. dürften). Jeder Test der Form $x = y$ für Gleitkommazahlen ist sinnlos; stattdessen ist stets festzustellen, ob $|x - y| < \varepsilon$, was bedeutet, daß x und y gleich sein *könnten* (wobei ε dem zu erwartenden Rundungsfehler entspricht). Um die Gleichheit definitiv feststellen zu können, müßten unendlich viele Stellen verglichen werden.

2.6 Kombinatorik

Permutationen

2.6.1 DEFINITION. Sei L eine Liste der Länge n , deren Elemente alle verschieden sind. Eine Liste die daraus durch Umordnen der Elemente entsteht, heißt eine *Permutation* von L . Die Anzahl aller möglichen Permutationen von L hängt nur von n ab (nicht von den konkreten Elementen in der Liste). Sie heißt *Fakultät* von n und wird meist mit $n!$ bezeichnet.

In der leeren Liste kann nicht viel umgeordnet werden: es gibt genau eine Permutation davon. Von einer Liste mit $n + 1$ Elementen bilden wir zuerst alle Permutationen der ersten n Elemente. Für das letzte Element bleiben dann genau $n + 1$ Stellen, an denen es eingefügt werden kann. Damit erhalten wir die Beziehungen

$$\begin{aligned} 0! &= 1, \\ (n + 1)! &= (n + 1) \cdot n!, \end{aligned}$$

welche als rekursive Definition verwendet werden können.

Kombinationen

2.6.2 DEFINITION. Sei A eine Menge mit n Elementen. Die Anzahl der k -elementigen Teilmengen (*Kombinationen*) von A hängt nur von n ab (nicht von A), heißt *Binomialkoeffizient* und wird meist mit $\binom{n}{k}$ bezeichnet.

Jede Menge hat genau eine 0-elementige Teilmenge. Die leere Menge hat sonst keine Teilmenge. Sei nun eine Menge mit $n + 1$ Elementen gegeben. Wir zählen zuerst die $(k + 1)$ -elementigen Teilmengen, welche das letzte Element nicht enthalten; dies sind $\binom{n}{k+1}$. Dann zählen wir die $(k + 1)$ -elementigen Teilmengen, welche das letzte Element sehr wohl enthalten. Diese entsprechen genau den k -elementigen Teilmengen, welche das letzte Element nicht enthalten. Daraus ergibt sich

$$\binom{n}{0} = 1 \qquad \binom{0}{k+1} = 0 \qquad \binom{n+1}{k+1} = \binom{n}{k+1} + \binom{n}{k}$$

2.6.3 SATZ. Seien $n, k: \mathbb{N}$ mit $0 \leq k \leq n$. Dann gilt

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

Beweis. Induktion nach n . □

2.6.4 THEOREM (BINOMISCHER LEHRSATZ). *Sei $n \in \mathbb{N}$ und a, b aus irgendeiner Menge, in der die üblichen Rechenregeln für $+$ und $*$ gelten. Dann gilt*

$$(a + b)^n = \sum_{i=0}^n \binom{n}{i} a^i b^{n-i}$$

Beweis. Induktion nach n . □

Kapitel 3

Rekursion

3.1 Natürliche Zahlen

Konstruktoren

Eine grundlegende Eigenschaft der *natürlichen Zahlen* ist, daß es zu jeder natürlichen Zahl eine weitere (größere) gibt, deren *Nachfolger*. Ebenso grundlegend ist, daß es überhaupt eine natürliche Zahl gibt, etwa die *Null*. Das bedeutet:

Um eine natürliche Zahl zu konstruieren, kann man sie entweder als Nachfolger einer bereits konstruierten natürlichen Zahl konstruieren, oder man konstruiert die spezielle natürliche Zahl 0. Dies läßt sich formal durch die folgenden Einführungsregeln beschreiben.

3.1.1 AXIOM.

$$\frac{k : \mathbb{N}}{Sk : \mathbb{N}} \text{NI}_S \qquad \frac{}{0 : \mathbb{N}} \text{NI}_0$$

So wie bei der Disjunktion bzw. direkten Summe gibt es zwei Einführungsregeln und entsprechend die zwei Konstruktoren 0 und S .

Peano-Induktion

Eine weitere grundlegende Eigenschaft der natürlichen Zahlen ist, daß man mit den oben eingeführten Konstruktionsmethoden alle natürlichen Zahlen konstruieren kann. Dies läßt sich wieder durch eine entsprechende Eliminationsregel beschreiben, allerdings, so wie bei der Disjunktion, in einer etwas komplizierten Form mit Fallunterscheidung. Dazu sei $A[n]$ irgendeine Aussage, in der die natürliche Zahl n vorkommen kann. Um zu zeigen, daß sie für jedes beliebige $n : \mathbb{N}$ gilt, kann man folgendermaßen vorgehen: Wir zeigen, daß, wenn die Aussage $A[k]$ für irgendein $k : \mathbb{N}$ gilt, dann gilt sie auch für dessen Nachfolger, d.h. es muß dann auch $A[Sk]$ gelten. Außerdem zeigen wir $A[0]$. Damit haben wir die Aussage für jedes n gezeigt, welches gemäß Einführungsregeln konstruiert wurde. Wenn es keine weiteren Möglichkeiten geben soll, natürliche Zahlen zu konstruieren, dann ist damit tatsächlich $A[n]$ für alle n bewiesen.

3.1.2 AXIOM (PEANO-INDUKTION).

$$\frac{\bigwedge_{k: \mathbb{N}} A[k] \implies A[Sk] \quad A[0]}{\bigwedge_{n: \mathbb{N}} A[n]} \text{--}\mathcal{NE}$$

Die erste Voraussetzung heit *Induktionsschlu*, die zweite *Induktionsanfang*.

Man kann das Induktionsprinzip auch folgendermaen begrnden: Wenn $A[n]$ fr alle $n: \mathbb{N}$ gezeigt werden sollte, und wenn es auer $0, S0, SS0, SSS0, SSSS0, \dots$ keine weiteren natrlichen Zahlen gibt, dann reicht es, jede der Aussagen $A[0], A[S0], A[SS0], A[SSS0], A[SSSS0]$, usw. beweisen. Haben wir aber $A[0]$ bereits bewiesen, so kann man diese Tatsache im Beweis von $A[S0]$ ausnutzen, d.h. man mu nur noch zeigen, da man $A[S0]$ aus $A[0]$ herleiten kann. Ebenso kann man dann beim Beweis von $A[SS0]$ verwenden, da man bereits $A[S0]$ gezeigt hat. Wir gehen also folgendermaen vor:

$$A[0] \quad \boxed{\begin{array}{c} A[0] \\ \vdots \\ A[S0] \end{array}} \quad \boxed{\begin{array}{c} A[S0] \\ \vdots \\ A[SS0] \end{array}} \quad \boxed{\begin{array}{c} A[SS0] \\ \vdots \\ A[SSS0] \end{array}} \quad \boxed{\begin{array}{c} A[SSS0] \\ \vdots \\ A[SSSS0] \end{array}} \quad \dots$$

Dabei sind immer noch unendlich viele Beweise zu fhren. Diese sind aber alles spezielle Flle des Induktionsschlusses.

Selektoren

Wir fhren eine Bezeichnung fr den Beweis der Konklusion in der Eliminationsregel ein.

3.1.3 DEFINITION.

$$\frac{\varphi: \bigwedge_{k: \mathbb{N}} (A[k] \implies A[Sk]) \quad a: A[0]}{\text{rec } \varphi x: \bigwedge_{n: \mathbb{N}} A[n]} \text{--}\mathcal{NE}$$

Das φ ordnet jedem $k: \mathbb{N}$ und jedem $x: A[k]$ ein $\varphi_k(x): A[Sk]$ zu. Der erste Parameter wird der bersichtlichkeit halber als Index geschrieben. Man kann das auch als $\varphi_k: A[k] \implies A[Sk]$ lesen.

Gleichheit

Wie bei den allgemeinen Konstruktionen fr Mengen, ergnzen wir auch hier die Regeln um den Gleichheitsbegriff.

3.1.4 AXIOM ($\mathbb{N}^=$ -ELIMINATION).

$$\frac{\varphi = \psi: \prod_{k: \mathbb{N}} (A[k] \rightarrow A[Sk]) \quad a = b: A[0]}{\text{rec } \varphi a = \text{rec } \psi a: \prod_{n: \mathbb{N}} A[n]} \text{--}\mathcal{NE}^=$$

Dies bedeutet insbesondere, daß der binäre Operator

$$\text{rec}: \prod_{k: \mathbb{N}} (A[k] \rightarrow A[Sk]) \rightarrow A[0] \rightarrow \prod_{n: \mathbb{N}} A[n]$$

wohldefiniert ist.

Durch verschiedenartige Anwendung dieser Regeln ergeben sich wieder die zugehörigen β - und η -Regeln:

3.1.5 AXIOM.

$$\frac{\varphi: \prod_{k: \mathbb{N}} (A[k] \rightarrow A[Sk]) \quad a: A[0] \quad n: \mathbb{N}}{\text{rec } \varphi a (Sn) = \varphi_n(\text{rec } \varphi a n) : A[Sn]} \mathbb{N}\beta_S \quad \frac{a: A[0]}{\text{rec } \varphi a 0 = a : A[0]} \mathbb{N}\beta_0$$

3.1.6 SATZ. Sei $f = \text{rec } \varphi a$, mit Datentypen wie oben. Dann gilt

$$\begin{aligned} f(0) &= a, \\ f(Sn) &= \varphi_n(f(n)), \end{aligned}$$

und f ist die einzige Funktion diese beiden Gleichungen erfüllt.

Beweis. Wir setzen die Definition ein und vereinfachen mit β -Regeln:

$$\begin{aligned} f(0) &= \text{rec } \varphi a 0 && \text{(Vorraussetzung)} \\ &= a && (\mathbb{N}\beta_0); \\ f(Sn) &= \text{rec } \varphi a (Sn) && \text{(Vorraussetzung)} \\ &= \varphi_n(\text{rec } \varphi a n) && (\mathbb{N}\beta_S) \\ &= \varphi_n(f(n)) && \text{(Voraussetzung)}. \end{aligned}$$

Damit ist der erste Teil bewiesen.

Um die Eindeutigkeit zu zeigen, sei f irgendeine Funktion, welche diese Gleichungen erfüllt. Dazu zeigen wir, daß $f = \text{rec } \varphi a$. Gemäß Extensionalität heißt das, daß wir für jedes $n: \mathbb{N}$ zu zeigen haben: $f(n) = \text{rec } \varphi a n$. Dies erledigen wir mit Induktion nach n .

Induktionsanfang: zu zeigen ist: $f(0) = \text{rec } \varphi a 0$, was sofort folgt.

Induktionsschritt: Wir nehmen an, daß $f(n) = \text{rec } \varphi a n$, und haben zu zeigen: $f(Sn) = \text{rec } \varphi a (Sn)$. Auch dieser Schritt besteht in einfachem Nachrechnen:

$$\begin{aligned} f(Sn) &= \varphi_n(f(n)) && \text{(Annahme)} \\ &= \varphi_n(\text{rec } \varphi a n) && \text{(Induktionshypothese)} \\ &= \text{rec } \varphi a (Sn) && (\mathbb{N}\mathcal{I}\beta_S). \end{aligned}$$

□

Der erste Parameter des Rekursionsoperators (φ) hat einen *abhängigen Datentyp*: Der Typ von φ_k hängt von k ab. So wichtig diese Form beim Beweisen mit Induktion ist, so wenig unterstützen heutige Programmiersprachen diesen Datentyp. Für den Fall, daß in der Aussage A gar kein k vorkommt, erhalten wir eine deutlich einfachere Variante der Rekursion, insbesondere wenn wir uns

auch noch auf den Fall beschränken, daß auch φ nicht von k abhängt. Dann kann man den Rekursionsoperator auch als

$$\text{rec}: (A \rightarrow A) \rightarrow A \rightarrow \mathbb{N} \rightarrow A$$

betrachten. Die β - und η -Regeln lauten dann:

$$\frac{\varphi: A \rightarrow A \quad a: A \quad n: \mathbb{N}}{\text{rec } \varphi a (Sn) = \varphi(\text{rec } \varphi a n) : A} \mathbb{N}\beta_S \qquad \frac{a: A}{\text{rec } \varphi a 0 = a : A} \mathbb{N}\beta_0$$

$$\frac{n: \mathbb{N}}{\text{rec } S 0 n = n : \mathbb{N}} \mathbb{N}\eta$$

Intuitiv bedeutet $\text{rec } \varphi a n$ damit einfach, daß die Funktion φ auf a n -mal angewendet wird, d.h.

$$\text{rec } \varphi a n = \underbrace{\varphi(\dots \varphi(\varphi(a)) \dots)}_{n\text{-mal}}$$

Mit dieser Betrachtungsweise ist auch die allgemeinere Variante besser zu verstehen:

$$\text{rec } \varphi a n = \varphi_{n-1}(\dots \varphi_1(\varphi_0(a)) \dots)$$

Vorgänger

3.1.7 LEMMA. Für alle $m, n: \mathbb{N}$ gilt

$$Sm = Sn \implies m = n.$$

Beweis. Wir definieren die *Vorgänger-Funktion* (*Predecessor*). Sie sollte die Gleichungen

$$\begin{aligned} P(Sn) &= n \\ P \ 0 &= 0 \end{aligned}$$

erfüllen. Tatsächlich läßt sich eine solche Funktion darstellen als

$$P = \text{rec}(n \mapsto x \mapsto n) 0.$$

Wir nehmen nun $Sm = Sn$ an und wenden auf beide Seiten P an. Damit haben wir $P(Sm) = P(Sn)$ hergeleitet. Wegen $P(Sm) = m$ und $P(Sn) = n$ erhalten wir daher $m = n$, womit die Behauptung $Sm = Sn \implies m = n$ bewiesen ist. $\square$

3.1.8 SATZ. Für alle $m: \mathbb{N}$ gilt

$$0 = Sm \implies 0 = S0.$$

Beweis. Wir beweisen mit Induktion nach m .

Induktionsanfang: Zu zeigen ist: $0 = S0 \implies 0 = S0$, was trivialerweise erfüllt ist.

Induktionshypothese: $0 = Sm \implies 0 = S0$. Zu zeigen: $0 = S(Sm) \implies 0 = S0$.

Wir nehmen also $0 = S(Sm)$ an und haben $0 = S0$ zu zeigen. Dazu wenden wir die Vorgänger-Funktion P auf beide Seiten der Voraussetzung an, und erhalten

$$P0 = P(S(Sm)).$$

Da $P0 = 0$ und $P(S(Sm)) = Sm$ ist, erhalten wir somit $0 = Sm$, was gerade die Voraussetzung in der Induktionshypothese ist. Daher gilt auch $0 = S0$. $\square$

Notation: Wir verwenden die Abkürzung $1 := S0$.

3.1.9 SATZ. Für alle $m, n: \mathbb{N}$ gilt

$$m = n \vee (m = n \implies 0 = 1).$$

Beweis. Wir führen eine Induktion nach m .

Induktionsanfang: Wir haben zu zeigen: $0 = n \vee (0 = n \implies 0 = 1)$. Dies wird im Lemma 3.1.10 unten gezeigt.

Induktionshypothese: $m = n \vee (m = n \implies 0 = 1)$

Zu zeigen: $Sm = n \vee (Sm = n \implies 0 = 1)$.

Wir zeigen wieder mit Induktion, diesmal nach n : Induktionsanfang: Zu zeigen: $Sm = 0 \vee (Sm = 0 \implies 0 = 1)$. Der zweite Teil dieser Aussage wurde im Satz 3.1.8 gezeigt.

Induktionshypothese: $Sm = n \vee (Sm = n \implies 0 = 1)$.

Zu zeigen: $Sm = Sn \vee (Sm = Sn \implies 0 = 1)$.

Die neue Induktionshypothese brauchen wir nicht, wohl aber die vorige. Da diese eine Disjunktion ist, gehen wir mit Fallunterscheidung vor.

Fall $m = n$: Dann ist auch $Sm = Sn$, und damit der erste Teil der Behauptung bewiesen.

Fall $m = n \implies 0 = 1$: Dann versuchen wir den zweiten Teil zu beweisen. Da dieser eine Implikation ist, nehmen wir $Sm = Sn$ an und versuchen damit $0 = 1$ herzuleiten. Wir wenden beiderseits die Vorgänger-Funktion an $P(Sm) = P(Sn)$, woraus sich durch deren Definition $m = n$ ergibt, und die Voraussetzung für diesen Fall angewandt werden kann. $\square$

3.1.10 LEMMA. Sei $n: \mathbb{N}$; dann gilt

$$0 = n \vee (0 = n \implies 0 = 1).$$

Beweis. Induktion nach n .

Induktionsanfang:

Zu zeigen: $0 = 0 \vee (0 = 0 \implies 0 = 1)$.

Offensichtlich gilt der linke Teil.

Induktionshypothese: $m = 0 \vee (m = 0 \implies 0 = 1)$

Zu zeigen: $Sm = 0 \vee (Sm = 0 \implies 0 = 1)$.

Wegen Satz 3.1.8 gilt der rechte Teil. (Die Induktionshypothese wurde nicht gebraucht.) $\square$

3.1.11 SATZ. Seien $P, Q: \Omega$ beliebige Aussagen. Dann gilt

$$0 = 1 \implies (P \iff Q)$$

Beweis. Wenn $0 = 1$, so gilt wegen der Wohldefiniertheit von rec auch

$$\text{rec } \varphi P 0 = \text{rec } \varphi P 1 : \Omega,$$

wobei $\varphi: \Omega \rightarrow \Omega$ eine konstante Funktion ist: $\varphi(R) = Q$. Da aber $\text{rec } \varphi P 0 = P : \Omega$ und $\text{rec } \varphi P 1 = Q : \Omega$ ist, muß auch $P = Q : \Omega$ sein, d.h. $P \iff Q$. $\square$

Addition

Mittels Rekursion kann man Funktionen auf den natürlichen Zahlen definieren. Die einfachste ist $\text{rec } S$. Dies ist eine Funktion in zwei Variablen. Bei $\text{rec } S m n$ wird von n ausgehend m mal der Nachfolger gebildet, was offensichtlich der Summe von m und n entspricht. Die β -Reduktionen Regeln ergeben damit

$$\begin{aligned}\text{rec } S m 0 &= n, \\ \text{rec } S m (Sn) &= S(\text{rec } S m n).\end{aligned}$$

Dies läßt sich leichter lesen, wenn statt $\text{rec } S$ der binäre Operator $(+)$ verwendet wird:

$$\begin{aligned}m + 0 &= n, \\ m + Sn &= S(m + n).\end{aligned}$$

3.1.12 DEFINITION. Die lineare Ordnung für natürliche Zahlen $m, n: \mathbb{N}$ wird festgelegt durch

$$m \leq n \iff \bigvee_{p: \mathbb{N}} m + p = n.$$

3.1.13 SATZ. Die lineare Ordnung $\leq$ auf den natürlichen Zahlen ist eine Ordnungsrelation.

3.1.14 SATZ. Seien $a, b, c: \mathbb{N}$; dann gilt

$$a \leq b \implies a + c \leq b + c$$

3.1.15 SATZ. Für $a, b: \mathbb{N}$ gilt

$$\begin{aligned}0 &\leq b \\ \neg(Sa \leq 0) \\ Sa \leq Sb &\iff a \leq b\end{aligned}$$

3.1.16 SATZ. Für alle $a, b: \mathbb{N}$ gilt

$$a \leq b \vee \neg(a \leq b)$$

3.1.17 SATZ. Für alle $a, b: \mathbb{N}$ gilt

$$a \leq b \vee b \leq a$$

3.2 Teilbarkeit

3.2.1 DEFINITION. Eine Zahl $d: \mathbb{Z}$ heißt ein *Teiler* von $n: \mathbb{Z}$ wenn es ein $q: \mathbb{Z}$ gibt sodaß $n = qd$. Man sagt dann auch d *teilt* n , und schreibt $d \mid n$. Es gilt also

$$d \mid n \iff \bigvee_{q: \mathbb{Z}} n = qd. \quad (3.1)$$

Man beachte die Analogie zur Definition von $m \leq n$.

3.2.2 SATZ. *Teilbarkeit ist transitiv, d.h. für alle $a, b, c: \mathbb{Z}$ gilt*

$$a \mid b \wedge b \mid c \implies a \mid c$$

Beweis. Seien $a \mid b$ und $b \mid c$. Gemäß der Definition der Teilbarkeit gibt es daher $x, y: \mathbb{Z}$, sodaß $b = x \cdot a$ und $c = y \cdot b$. Einsetzen der ersten Gleichheit in die zweite ergibt damit $c = y \cdot (x \cdot a) = (y \cdot x) \cdot a$ (Assoziativgesetz), was wieder genau der Definition von $a \mid c$ entspricht. $\square$

3.2.3 SATZ. *Seien $d, n, m, z: \mathbb{Z}$ und $d \mid n$, $d \mid m$. Dann gelten auch $d \mid n + m$, $d \mid n - m$, und $d \mid zn$.*

3.2.4 BEMERKUNG. Die Teilbarkeitsrelation in $\mathbb{Z}$ ist nicht definit: für jede ganze Zahl $a: \mathbb{Z}$ gelten zwar $a \mid (-a)$ und $(-a) \mid a$, aber nicht $a = -a$ (außer wenn $a = 0$).

3.2.5 SATZ. *Seine $a, b: \mathbb{Z}$; dann gilt*

$$a \mid b \wedge b \mid a \iff |a| = |b|.$$

Eigentlich kann man sich bei der Teilbarkeit auf die von positiven Zahlen beschränken, den es gilt:

3.2.6 SATZ. *Seien $a, b: \mathbb{Z}$. Dann gilt*

$$a \mid b \iff |a| \mid |b|$$

Schränkt man die Teilbarkeit auf die natürlichen Zahlen ein, so erhält man tatsächlich eine Ordnungsrelation.

3.2.7 BEMERKUNG. Teilbarkeit für natürliche Zahlen $m, n: \mathbb{N}$ wird am natürlichsten als

$$d \mid n \iff \bigvee_{q: \mathbb{N}} n = q \cdot d. \quad (3.2)$$

Identifiziert man aber $m, n: \mathbb{N}$ mit den ganzen Zahlen $m, n: \mathbb{Z}$, erhält man jedoch aus der Definition der Teilbarkeit für ganze Zahlen

$$d \mid n \iff \bigvee_{q: \mathbb{Z}} n = q \cdot d,$$

was ein wenig anders aussieht (man beachte, daß hier $q: \mathbb{Z}$ statt $q: \mathbb{N}$). Erst wenn nachgewiesen ist, daß beide Definition übereinstimmen (was in diesem Fall leicht ist) kann man natürliche Zahlen auch im Zusammenhang mit Teilbarkeit in der üblichen Weise mit ganzen Zahlen identifizieren.

3.2.8 DEFINITION. Man nennt d einen *gemeinsamen Teiler* von $a, b: \mathbb{Z}$, wenn $d \mid n$ und $d \mid m$ gilt.

3.2.9 BEMERKUNG. Man beachte: Jede Zahl teilt 0, denn es gilt stets $0 = x \cdot 0$, woraus unmittelbar $x \mid 0$ folgt. Es gilt sogar $0 \mid 0$. Andere Zahlen teilt 0 jedoch nicht, denn $0 \mid a$ bedeutet ja gerade daß $a = x \cdot 0$, woraus aber $a = 0$ folgt.

0 ist damit die größte Zahl, wenn diese nach Teilbarkeit geordnet werden. Die Zahl 1 dagegen teilt jede andere Zahl, und ist damit die kleinste.

3.2.10 SATZ. Ist d ein gemeinsamer Teiler von m und n , und t ein Teiler von d , dann ist auch t ein gemeinsamer Teiler von m und n .

Beweis. Gemäß Voraussetzung gibt es Zahlen $a, b, c: \mathbb{Z}$, sodaß

$$m = adn = bdd = ct;$$

einfaches Einsetzen ergibt dann mit der Assoziativität

$$m = (ac)tn = (bc)t,$$

womit die Behauptung unmittelbar folgt. □

3.2.11 DEFINITION. Ein gemeinsamer Teiler d von a und b heißt ein *größter gemeinsamer Teiler*, wenn jeder andere gemeinsame Teiler t von a und b bereits ein Teiler von d ist.

Man beachte, daß sich das *größer* in dieser Definition auf die Teilbarkeitsrelation bezieht, und vorerst nichts mit der normalen linearen Größenrelation zu tun hat. Es ist damit auch nicht von vorneherein klar, ob es überhaupt größte gemeinsame Teiler gibt, und ob diese eindeutig bestimmt sind. Letzteres ist aber einfach:

3.2.12 SATZ. Seien $a, b: \mathbb{Z}$. Zwei größte gemeinsame Teiler von a und b stimmen bis auf das Vorzeichen überein.

Als den größten gemeinsamen Teiler bezeichnet man dann meistens die positive Variante.

Manchmal ist der größte gemeinsame Teiler leicht zu bestimmen.

3.2.13 SATZ. Ist d ein Teiler von a , so ist d ein größter gemeinsamer Teiler von d und a .

Beweis. Weil $d \mid d$ (Teilbarkeit ist reflexiv) ist d klarerweise ein gemeinsamer Teiler von d und a . Ist nun t ein weiterer gemeinsamer Teiler von d und a , dann ist insbesondere t ein Teiler von d , weshalb d tatsächlich ein größter gemeinsamer Teiler ist. □

3.2.14 LEMMA. Seien $m, n: \mathbb{Z}$ und $m = qn + r$, mit $q, r: \mathbb{Z}$. Dann ist ein $d: \mathbb{Z}$ genau dann ein gemeinsamer Teiler von m und n wenn d ein gemeinsamer Teiler von n und r ist.

Beweis. Wir nehmen an, daß $d \mid m$ und $d \mid n$, und zeigen damit daß $d \mid n$ und $d \mid r$. Der erste Teil steht schon in der Voraussetzung. Wegen $r = m - qn$, und weil $d \mid m$ und $d \mid n \mid qn$, gilt wegen Satz 3.2.3 auch $d \mid r$.

Die Umkehrung funktioniert ganz ähnlich. □

3.2.15 THEOREM (EUCLIDISCHER ALGORITHMUS). *Seien $m, n: \mathbb{Z}$. Dann gibt es genau einen (bis auf das Vorzeichen) größten gemeinsamen Teiler $d: \mathbb{N}$. Ferner gibt es ganzzahlige Koeffizienten $x, y: \mathbb{Z}$ sodaß*

$$d = xm + yn.$$

Beweis. Wir beweisen mit Induktion nach $|n|$.

Ist $n = 0$, so ist m ein größter gemeinsamer Teiler von m und n . Wegen $d = m = 1m + 0n$ sind auch geeignete Koeffizienten leicht gefunden.

Ist $|n| > 0$, dann gibt es Zahlen q, r mit $m = qn + r$ und $0 \leq r < |n|$. Laut Induktionsvoraussetzung haben n und r einen größten gemeinsamen Teiler d , und dieser läßt sich als $d = xn + yr$ darstellen, mit passenden Koeffizienten $x, y: \mathbb{Z}$. Wegen Lemma 3.2.14 ist aber d auch ein größter gemeinsamer Teiler von m und n . Weiters gilt $d = xn + yr = xn + y(m - qn) = ym + (x - yq)n$. Damit sind auch für n und m passende Koeffizienten gefunden. $\square$

Dieser Beweis liefert ein einfaches und trotzdem recht effizientes Verfahren, um den größten gemeinsamen Teiler, und auch passende Koeffizienten zu finden. Die einfachere Variante berechnet nur den ggT:

$$\text{gcd } m \ 0 = m;$$

$$\text{gcd } m \ n = \text{gcd } n \ r \quad \text{wobei } m = qn + r, \ 0 \leq r < |n|.$$

Die erweiterte Variante berechnet gleichzeitig auch passende Koeffizienten:

$$\text{xgcd } m \ 0 = m;$$

$$\begin{aligned} \text{xgcd } m \ n &= (d, y, x - yq) \quad \text{wobei } m = qn + r, \ 0 \leq r < |n| \\ &\text{und } (d, x, y) = \text{xgcd } n \ r. \end{aligned}$$

3.2.16 BEMERKUNG. Die Einschränkung für r bei diesem Algorithmus ist nicht die einzig mögliche. Eine sinnvolle Alternative ist es zum Beispiel, den betragsmäßig kleinsten Rest zu verwenden, d.h. $r \leq \lfloor \frac{n}{2} \rfloor$.

3.2.17 SATZ. *Sei $m: \mathbb{N}$. Dann gibt es zu jedem $e: \mathbb{Z}$ mit $\text{gcd } m \ e = 1$ einen Kehrwert modulo m , d.h. ein $d: \mathbb{Z}$, sodaß $de \equiv 1 \pmod{m}$.*

Beweis. Mit dem Euklidischen Algorithmus (3.2.15, erweiterte Version) erhalten wir ganzzahlige Koeffizienten x, y sodaß $1 = xm + yn$ gilt, was unmittelbar $1 \equiv yn \pmod{m}$ zur Folge hat. Der Koeffizient y ist also eine geeignete Wahl für das geforderte d . $\square$

3.3 Modulare Arithmetik

3.3.1 DEFINITION. Sei $m: \mathbb{N}$; dann heißen ganze Zahlen $a, b: \mathbb{Z}$ *kongruent modulo m* (Schreibweise: $a \equiv_m b$) falls m ein Teiler von deren Differenz $a - b$ ist.

Neben der erwähnten sind noch die Schreibweisen: $a \equiv b \pmod{m}$, oft auch

abgekürzt zu $a \equiv b \pmod{m}$, und auch $a = b \pmod{m}$ bzw $a = b \pmod{m}$ üblich. Also:

$$\begin{aligned} m \mid (a - b) &\iff a \equiv_m b \\ a &\equiv b \pmod{m} \\ a &\equiv b \pmod{m} \\ a &= b \pmod{m} \\ a &= b \pmod{m}. \end{aligned}$$

Wenn das m aus dem Zusammenhang klar ist, schreibt man oft auch nur $a \equiv b$ oder sogar $a = b$.

Es ist leicht nachzurechnen, daß die Kongruenz modulo m stets eine Äquivalenzrelation ist. Man kann also den Gleichheitsbegriff auf $\mathbb{Z}$ durch die Kongruenz modulo m ersetzen.

3.3.2 DEFINITION. Sei $m: \mathbb{N}$. Dann bezeichnet $\mathbb{Z}_m = \mathbb{Z} / \equiv_m$ die Menge der Restklassen modulo m .

Damit erhalten wir eine weitere Notation für $a \equiv_m b$, nämlich: $a = b : \mathbb{Z}_m$.

3.3.3 SATZ. Die Kongruenz modulo m ist mit der Addition, der Subtraktion und der Multiplikation verträglich, d.h. sind $a \equiv_m b$ und $c \equiv_m d$, so gelten auch

$$a + c \equiv_m b + d \qquad a - c \equiv_m b - d \qquad a \cdot c \equiv_m b \cdot d.$$

Dieser Satz garantiert, daß diese drei Grundrechnungsarten auch für Restklassen wohldefiniert sind. Man beachte, daß dies für die Division nicht funktioniert: z.B.: $40 \equiv_5 30$ und $10 \equiv_5 15$, aber es gilt nicht $40/10 \equiv_5 30/15$. Dafür kann man auch Potenzieren (mit ganzzahligen Exponenten):

3.3.4 SATZ. Sei $m: \mathbb{Z}$, $a = b : \mathbb{Z}_m$, $n: \mathbb{N}$; dann gilt

$$a^n = b^n : \mathbb{Z}_m$$

Somit kann man modulo m ganz normal rechnen, falls keine Divisionen dabei sind.

3.3.5 SATZ. Sei $m: \mathbb{N}$ und $a: \mathbb{Z}_m$. Dann gibt es genau dann ein $b: \mathbb{Z}_m$ mit $ab = 1 : \mathbb{Z}_m$ falls $\text{ggT}(a, m) = 1$.

Beweis. Gilt $\text{ggT}(a, m) = 1$, so liefert der erweiterte Euklidische Algorithmus Lösungen $x, y: \mathbb{Z}$, sodaß $ax + my = 1$. Diese Gleichung gilt natürlich auch modulo m , und vereinfacht sich in diesem Fall zu $ax = 1 : \mathbb{Z}_m$, womit jedes $b = x : \mathbb{Z}_m$ eine passende Lösung ist.

Für die Umkehrung sei $ab = 1 : \mathbb{Z}_m$. Das heißt aber, daß es ein y gibt, sodaß $ab - 1 = my$, bzw $ab - my = 1$. Jeder gemeinsame Teiler von a und m ist damit auch ein Teiler von 1, d.h. $\text{ggT}(a, m) = 1$. $\square$

3.3.6 SATZ. Die Berechnung von $a^n : \mathbb{Z}_m$ ist durch sukzessives Quadrieren (und sofortiges Reduzieren modulo m) effizient möglich.

Beweis. Wir verwenden die Gleichungen;

$$\begin{aligned} a^0 &= 1; \\ a^{2n} &= (a^n)^2; \\ a^{2n+1} &= a(a^n)^2. \end{aligned}$$

Für ganze Zahlen bringt das nichts, weil sie bei jedem Quadrieren doppelt so lange werden. Wenn wir aber in jedem Schritt modulo m reduzieren können, bleiben alle Zwischenergebnisse durch m beschränkt. Ist etwa n in der Größenordnung von 2^{1000} (also etwa 300 Dezimalstellen), so muß man damit nur etwa 1000 mal Quadrieren und etwa ebensooft multiplizieren. Mit dem normalen Verfahren würde man dagegen $2^{1000} \approx 10^{300}$ Multiplikationen benötigen, was selbst mit Computern in astronomischer Größe unmöglich wäre. $\square$

3.4 Primzahlen

3.4.1 DEFINITION. Eine natürliche Zahl $n > 1$ heißt *Primzahl* wenn sie keine echten Teiler hat (also nur 1 und n selbst. 1 ist per Definition keine Primzahl).

3.4.2 LEMMA. Eine natürliche Zahl $p > 1$ ist genau dann eine Primzahl, wenn gilt:

$$\bigwedge_{n,m \in \mathbb{N}} (p \mid nm \implies p \mid n \vee p \mid m).$$

Ein bekanntes Verfahren zur Bestimmung aller Primzahlen ist das **Sieb des Eratosthenes**: Man schreibt alle natürlichen Zahlen ab 2 in einer Liste an. (Freilich ist die Liste eigentlich unendlich lange, aber man kann ja nach Bedarf weiterschreiben.) Die erste Zahl (also 2) wird als Primzahl gekennzeichnet und all ihre nicht-trivialen Vielfachen werden aus der Liste gestrichen (denn diese sind keine Primzahlen). Die nächste (noch nicht gestrichene) Zahl ist 3. Auch diese wird als Primzahl gekennzeichnet und ihre Vielfachen gestrichen. Die nächste Zahl nicht gestrichene Zahl ist daher 5, welche ebenfalls als Primzahl gekennzeichnet wird, und deren Vielfache gestrichen werden. Auf diese Weise fährt man unendlich lange fort und findet weiters die Primzahlen 7, 11, 13, 17, 19, 23, ...

Man beachte: dieses Verfahren läuft unendlich lange. Dies ist aber kein Problem, da dennoch jede einzelne Primzahl nach endlicher Zeit gefunden wird.

Um zu Testen, ob eine Zahl n prim ist, kann man nachsehen, ob sie in der vom Sieb des Eratosthenes erzeugten Liste aufscheint (nie gestrichen wird). Oder man testet alle Zahlen bis zu ca. $\sqrt{n}$, ob sie n teilen (wenn nicht, ist n eine Primzahl). Beide Verfahren sind für große Zahlen (z.B. 300-stellige Zahlen) allerdings nicht brauchbar. In der Praxis verwendet man probabilistische Tests, welche bedeutend schneller sind, und für die ausreichend zuverlässig arbeiten (siehe unten).

3.4.3 THEOREM (HAUPTSATZ DER ARITHMETIK). Jede natürliche Zahl n hat eine eindeutige Zerlegung in Primfaktoren

$$n = p_1^{k_1} \cdots p_m^{k_m},$$

wobei alle p_i Primzahlen sind, mit $p_i < p_{i+1}$, und $k_i > 1$.

Leider ist das Auffinden von Primfaktor-Zerlegungen für große Zahlen im allgemeinen noch schwieriger als ein Primzahltest, und es gibt auch kein probabilistisches Verfahren. „Leider“ sollte man hier allerdings nicht sagen, denn gerade diese Tatsache, daß sich für große Zahlen die Primzahlzerlegung praktisch nicht bestimmen läßt, ist die Grundlage der gängigsten Verfahren zur Verschlüsselung von Daten (siehe unten) und daher von größter Bedeutung in der Praxis.

3.4.4 SATZ. Seien p, q verschiedene Primzahlen, und für $x, y: \mathbb{Z}$ gelte

$$x \equiv y \pmod{p} \quad \text{und} \quad x \equiv y \pmod{q}.$$

Dann gilt auch

$$x \equiv y \pmod{pq}.$$

Beweis. Laut Voraussetzung gilt $p \mid (x - y)$ und $q \mid (x - y)$, und daher auch $pq \mid (x - y)$. $\square$

3.4.5 DEFINITION. Sein $n: \mathbb{N}$. Dann bezeichnet $\varphi(n)$ die Anzahl der Zahlen $k: \mathbb{N}$, $k < n$, welche zu n teilerfremd sind. Die Funktion φ heißt die *Eulersche φ -Funktion*.

Ist p eine Primzahl, so gilt $\varphi(p) = p - 1$, und allgemeiner, $\varphi(p^k) = p^{k-1}(p - 1)$. Ferner gilt $\varphi(mn) = \varphi(m)\varphi(n)$, falls m und n teilerfremd sind. Damit kann man φ für alle Zahlen leicht berechnen, für die man die Primfaktorzerlegung kennt.

3.4.6 THEOREM. Für $n: \mathbb{N}$ und alle $a: \mathbb{Z}$ teilerfremd zu n gilt

$$a^{\varphi(n)} \equiv 1 \pmod{n}.$$

3.4.7 FOLGERUNG. Seien p, q verschiedene Primzahlen, dann gilt für beliebige $a: \mathbb{Z}$

$$a^p \equiv a \pmod{p} \quad \text{und} \quad a^{\varphi(pq)+1} \equiv a \pmod{pq}.$$

Beweis. Der erste Teil ist trivial, falls $a \equiv 0 \pmod{p}$; ansonsten folgt er aus dem vorigen Satz mittels Multiplikation mit a .

Für den zweiten Teil zeigen wir zunächst daß die Kongruenz modulo p gilt: Wieder ist der Fall $a \equiv 0 \pmod{p}$ trivial; ansonsten folgt wieder mit dem vorigen Satz:

$$a^{\varphi(pq)+1} = a^{(p-1)(q-1)+1} = (a^{p-1})^{q-1} a \equiv 1 \pmod{p}.$$

Analog zeigt man die Kongruenz modulo q , woraus mit Satz 3.4.4 auch die Kongruenz modulo pq folgt. $\square$

Die Eigenschaft $a^{p-1} \equiv 1 \pmod{p}$ kann als Primzahltest gebraucht werden: ist sie für ein beliebiges $a < p$ verletzt, so weiß man, daß p keine Primzahl sein kann. Ist sie dagegen für ein a erfüllt, so kann man vermuten, daß sie eine Primzahl ist, gilt sie für mehrere a , dann ist es schon ziemlich wahrscheinlich. Damit hat man einen einfachen brauchbaren Test; allerdings gibt es Nicht-Primzahlen, die nach diesem Test wie Primzahlen aussehen (Pseudoprime). Besser ist der Miller-Rabin-Test: Dieser funktioniert nach einem ähnlichen Prinzip, ermöglicht es aber die Fehlerwahrscheinlichkeit tatsächlich beliebig klein zu wählen.

RSA-Verfahren

Das RSA-Verfahren ist das bekannteste asymmetrische Verschlüsselungsverfahren. Der Grundgedanke dabei ist, daß jeder Teilnehmer an einem Kommunikationssystem zwei Schlüssel bekommt: einen öffentlichen Schlüssel P , der sowie eine Telefonnummer oder Email-Adresse allen anderen Teilnehmern bekannt gemacht wird, und einen geheimen Schlüssel S , welchen jeder Teilnehmer für sich behält. Ein Schlüssel ist dabei eine Funktion, welche eine Nachricht in eine andere Nachricht umwandelt; dabei sollte aber gelten

$$P \circ S = \text{id} = S \circ P.$$

Wir nehmen jetzt an, daß Alice eine geheime Nachricht a an Bob senden will. Dazu verschlüsselt Alice die Nachricht a mit Bob's öffentlichem Schlüssel (der ist ja allgemein bekannt); dabei entsteht eine unleserliche Nachricht $P(a)$. Diese kann bedenkenlos z.B. via Internet verschickt werden. Bob kann sie aber lesen, denn er (und nur er) hat ja den dazupassenden geheimen Schlüssel S , und es gilt ja $S(P(a)) = a$.

Nun möchte Bob eine Nachricht b verschicken (egal an wen) und dabei sicherstellen, daß jeder garantiert weiß, daß diese tatsächlich von ihm stammt. Dazu verschlüsselt er mit seinem geheimen Schlüssel, d.h. er verschickt $S(b)$. Jeder kann überprüfen, daß diese Nachricht tatsächlich von Bob stammt, da die Anwendung seines öffentlichen Schlüssels auf diese Nachricht $P(S(b)) = b$ ergibt, und nur Bob eine solche Nachricht erzeugen kann.

Das ganze System funktioniert natürlich nur dann, wenn es praktisch unmöglich ist, den geheimen Schlüssel zu finden, selbst wenn man den öffentlichen Schlüssel und vielleicht auch noch zahlreiche mit dem geheimen Schlüssel verschlüsselte Nachrichten kennt.

Der Grundgedanke beim RSA-Verfahren ist es nun, sich zu Nutze zu machen, daß es zwar relativ leicht ist, zwei Zahlen miteinander zu multiplizieren (selbst wenn sie hunderte Stellen haben), es aber umgekehrt im allgemeinen sehr schwierig ist, die Primfaktoren einer solchen Zahl zu bestimmen, obwohl diese ja auch eindeutig bestimmt sind.

Dabei geht man (leicht vereinfacht) folgendermaßen vor:

Man wählt zwei große Primzahlen p und q (512 bit große Zahlen gelten heutzutage als sicher, das sind etwa 150 Dezimalstellen, Primzahltest!), und bilde deren Produkt $n = pq$, sowie $\varphi = (p - 1)(q - 1)$. Weiters wähle man eine zu φ teilerfremde Zahl e . Damit erhält man den öffentlichen Schlüssel

$$P(a) = a^e \bmod m.$$

Bekanntgegeben werden also m und e , nicht aber die Primzahlen p, q und auch nicht φ .

Um den dazupassenden geheimen Schlüssel zu erhalten, bestimmen wir ein d mit $de \equiv 1 \pmod{\varphi}$, was möglich ist, weil e zu φ teilerfremd gewählt wurde (Erweiterter Euklidischer Algorithmus). Der geheime Schlüssel ist dann

$$S(a) = a^d \bmod m.$$

Tatsächlich gilt wegen $de \equiv 1 \pmod{\varphi}$ und $a^\varphi \equiv 1 \pmod{m}$,

$$S(P(a)) \equiv (a^e)^d \equiv a^{ed} \equiv a^1 \equiv 1 \pmod{m}.$$

Ist also $a < m$, so kann a tatsächlich rekonstruiert werden.

Man beachte, daß wir, um d zu bestimmen, die Zahl φ gebraucht haben, welche wiederum leicht aus p und q berechnet werden kann. Es ist aber nicht ersichtlich, wie es direkt aus dem Produkt $m = pq$ bestimmt werden kann. Auf der Annahme, daß dies tatsächlich praktisch unmöglich ist, beruht das RSA-Verfahren. Bisher ist es nicht geknackt worden.

3.5 Listen

Sei im folgenden A ein beliebiger Datentyp.

Introduktion: Listen-Konstruktoren

3.5.1 DEFINITION. Der Datentyp aller *Listen* über A wird mit $\text{List } A$ oder A^* bezeichnet und folgendermaßen definiert: Um ein Element von A^* zu konstruieren, muß man entweder die leere Liste $[]$ konstruieren oder $\text{cons } h \ t$, d.h. ein Element $h : A$ zusammen mit einer Liste $t : A^*$.

3.5.2 NOTATION. Der Ausdruck $\text{cons } h \ t$ sollte bedeuten, daß h bei der Liste t vorne eingefügt wird. Wir verwenden dafür daher die suggestivere Notation $h \triangleleft t$; und ein Ausdruck der Form $a \triangleleft b \triangleleft \ell$ ist so wie $a \triangleleft (b \triangleleft \ell)$ zu verstehen (andersherum würden die Datentypen gar nicht passen). Weiters verwenden wir oft $[a, b, c, \dots, z]$ als übersichtlichere Darstellung statt $a \triangleleft b \triangleleft c \triangleleft \dots \triangleleft z \triangleleft []$; also insbesondere $[a] = a \triangleleft []$ und $[a, b] = a \triangleleft b \triangleleft []$.

Ein Beispiel für eine Liste über $\mathbb{N}$ ist $2 \triangleleft 6 \triangleleft 1 \triangleleft 17 \triangleleft []$ (oder einfacher notiert: $[2, 6, 1, 17]$).

Die Definition können wir formal durch die folgenden Einführungsregeln notieren:

3.5.3 AXIOM.

$$\frac{}{[] : A^*} \text{List } [] \quad \frac{h : A \quad t : A^*}{h \triangleleft t : A^*} \text{List } \triangleleft$$

Oder wir legen einfach die Datentypen der betroffenen Konstruktoren fest:

$$\begin{aligned} [] &: A^*; \\ \triangleleft &: A \rightarrow A^* \rightarrow A^*. \end{aligned}$$

Elimination: Listen-Induktion

Sei $P : A^* \rightarrow \Omega$ eine beliebige Eigenschaft auf A^* , d.h. ein Prädikat, welches jeder Liste eine Aussage zuordnet. Wenn es gelingt, jede der Aussagen

$$\begin{aligned} &P([]), \\ &P([]) \implies P([17 \triangleleft []]), \\ &P(17 \triangleleft []) \implies P(1 \triangleleft 17 \triangleleft []), \\ &P(1 \triangleleft 17 \triangleleft []) \implies P([6 \triangleleft 1 \triangleleft 17 \triangleleft []]) \end{aligned}$$

zu beweisen, dann ist klarerweise (mehrmalige Anwendung der $\Rightarrow \mathcal{E}$) auch die Aussage $P([6 \triangleleft 1 \triangleleft 17 \triangleleft []])$ bewiesen.

Gelingt es nun, für alle $h : A$ und $t : A^*$ die Aussage

$$P(t) \implies P(h \triangleleft t)$$

nebst $P([])$ zu beweisen, so sind auch alle Aussagen der Form $P(a \triangleleft b \triangleleft \dots \triangleleft z \triangleleft [])$ (bzw. $P[a, b, \dots, z]$) für beliebige $a, b, \dots, z : A$ bewiesen, d.h. für alle Listen, die sich auch $[]$ und $\triangleleft$ aufbauen lassen. Die folgende Eliminationsregel besagt daher gerade, daß tatsächlich alle Listen derart entstehen.

3.5.4 AXIOM (PRIMITIVE INDUKTION ÜBER LISTEN).

$$\frac{\frac{\bigwedge_{t:A^*} (P(t) \implies \bigwedge_{h:A} P(h \triangleleft t)) \quad P[\Box]}{\bigwedge_{\ell:A^*} P[\ell]} \text{List-}\mathcal{E}}$$

Das bedeutet: Kann man zeigen, daß eine Aussage für die leere Liste gilt und, wenn sie für eine Liste t gilt, dann auch für jede Liste der Form $h \triangleleft t$, dann gilt sie für jede Liste.

Selektor: Rekursion

Wir ergänzen die Listen-Eliminationsregel durch den passenden Selektor:

3.5.5 AXIOM (PRIMITIVE REKUSION ÜBER LISTEN).

$$\frac{\varphi: \bigwedge_{t:A^*} (P(t) \implies \bigwedge_{h:A} P(h \triangleleft t)) \quad a: P(\Box)}{\text{listrec } \varphi \ a: \bigwedge_{\ell:A^*} P[\ell]} \text{List-}\mathcal{E}$$

$P(t)$ kann dabei durchaus ein beliebiger Datentyp sein (nicht nur eine Aussage). Dieser Selektor hat somit den Typ:

$$\text{listrec}: \prod_{t:A^*} (P(t) \rightarrow \prod_{h:A} P(h \triangleleft t)) \rightarrow A \rightarrow \prod_{\ell:A^*} P(\ell),$$

wobei P als beliebiger durch A^* parametrisierter Datentyp zu verstehen ist. Wenn P nicht parametrisiert ist, dann verschwinden die abhängigen Typen (Allaussagen werden zu Implikationen) und alles vereinfacht sich zu

$$\text{listrec}: (A^* \rightarrow P \rightarrow A \rightarrow P) \rightarrow A \rightarrow A^* \rightarrow P.$$

Dies ist die übliche Form, wie Funktionen auf Listen rekursiv definiert werden.

Gleichheit

Wenn es für A einen vernünftigen Gleichheitsbegriff gibt, dann sollte dieser auch auf A^* übertragen werden.

Zwei Listen sollten genau dann gleich sein, wenn sie auf dieselbe Weise konstruiert wurden. Aus den Introduktions- und Eliminationsregeln ergibt sich damit:

3.5.6 AXIOM.

$$\frac{\frac{\Box = \Box: A^*}{\frac{\ell_1 = \ell_2: A^* \quad a_1 = a_2: A}{a_1 \triangleleft \ell_1 = a_2 \triangleleft \ell_2: A^*}}{\frac{\varphi_1 = \varphi_2: \bigwedge_{t:A^*} (P(t) \implies \bigwedge_{h:A} P(h \triangleleft t)) \quad a_1 = a_2: P(\Box)}{\text{listrec } \varphi_1 \ a_1, \text{listrec } \varphi_2 \ a_2: \bigwedge_{\ell:A^*} P[\ell]} \text{List-}\mathcal{E}}$$

Wir eruieren nun, wie die Konstruktoren $\sqcup$ und $\triangleleft$ mit dem Selektor listrec zusammenhängen. Wir verwenden dazu die Notation wie in der Regel für die Listen-Rekursion. Wurde $\bigwedge_{\ell: A^*} P(\ell)$ mit Listenrekursion bewiesen, dann bezeichnet der Ausdruck $\text{listrec } \phi p \ell$ stets einen Beweis von $P(\ell)$. Für die leere Liste $\square$ haben wir zwei Beweise: $p: P \square$ und $\text{listrec } \phi p \square$. Ebenso für eine Liste der Form $h \triangleleft t$: $\text{listrec } \phi p t$ und $\phi t q a$, wenn $q: P(t)$ ist. Tatsächlich entspricht es gerade der Intention, daß diese jeweils gleich sind.

3.5.7 AXIOM (β -REGEL FÜR LISTEN-REKURSION).

$$\begin{aligned}\text{listrec } \phi p \square &= p \\ \text{listrec } \phi p (h \triangleleft t) &= \phi h t (\text{listrec } \phi p t)\end{aligned}$$

Ist nun $f = \text{listrec } \phi p$, dann gilt somit

$$\begin{aligned}f \square &= p \\ f (h \triangleleft t) &= \phi h t (f t)\end{aligned}$$

Dies ist im Prinzip eine Definition mit Fallunterscheidung; allerdings kommt das zu definierende f in der zweiten Gleichung auf beiden Seiten vor. Es ist daher nicht selbstverständlich, daß es eine solche Funktion überhaupt gibt. Genau das besagt aber die Regel für die Listen-Rekursion. Wichtig dabei ist, daß das f auf der rechten Seite nicht an beliebiger Stelle vorkommen darf, sondern nur in der Kombination $f t$, also ein bestimmter (leichter zu bestimmender) Funktionswert (Rekursionsprinzip).

3.5.8 NOTATION. Das ϕ ist eine dreistellige Operation; dessen Anwendung daher etwas unübersichtlich. Aber für jedes $\ell: A^*$ ist $\phi \ell$ eine binäre Operation, welche man als Infixoperator schreiben kann. Das ℓ wird dabei sinnvollerweise als Index geschrieben. Wir verwenden also statt $\phi \ell a q$ die Notation:

$$a \varphi_{\ell} q,$$

wobei die Klammern in einem Ausdruck wie $a \varphi_{\ell} (b \varphi_{\ell} q)$ weggelassen werden dürfen.

3.5.9 BEISPIEL. Sei $f = \text{listrec } \phi p$ wir berechnen den Funktionswert an der Stelle $[1, 2, 3]$:

$$\begin{aligned}f[1, 2, 3] &= f(1 \triangleleft 2 \triangleleft 3 \triangleleft \square) \\ &= 1 \varphi_{[2,3]} f(2 \triangleleft 3 \triangleleft \square) \\ &= 1 \varphi_{[2,3]} 2 \varphi_{[3]} f(3 \triangleleft \square) \\ &= 1 \varphi_{[2,3]} 2 \varphi_{[3]} 3 \varphi_{\square} f(\square) \\ &= 1 \varphi_{[2,3]} 2 \varphi_{[3]} 3 \varphi_{\square} p\end{aligned}$$

3.5.10 SATZ. Für $h, h': A$ und $t, t': A^*$ sowie jede beliebige Aussage P gelten

$$\begin{aligned}h \triangleleft t = h' \triangleleft t' &\implies h = h' \wedge t = t'; \\ h \triangleleft t = \square &\implies P.\end{aligned}$$

Länge

3.5.11 BEISPIEL. Sei A ein beliebiger Datentyp. Die Länge einer Liste, $\text{length} : A^* \rightarrow \mathbb{N}$, wird definiert als $\text{length} = \text{listrec } S \, 0$, d.h.

$$\begin{aligned} \text{length} [] &= 0; \\ \text{length}(h \triangleleft t) &= S(\text{length } t). \end{aligned}$$

3.5.12 NOTATION. Statt $\text{length } \ell$ schreiben wir einfacher $|\ell|$. Damit lesen sich die obigen Gleichungen als

$$\begin{aligned} |[]| &= 0; \\ |h \triangleleft t| &= S |t|. \end{aligned}$$

Hintenanfügen und Umkehren

Die Funktion cons wird als Einfügen eines Elements am Beginn der Liste interpretiert (z.B. $5 \triangleleft [1, 2, 3] = [5, 1, 2, 3]$). Wir definieren nun eine Funktion, die dementsprechend am hinteren Ende einer Liste ein Element anhängt, z.B. $[1, 2, 3] \triangleright 5 = [1, 2, 3, 5]$.

3.5.13 DEFINITION. Sei A ein beliebiger Datentyp. Die Operation $\triangleright : A^* \rightarrow A \rightarrow A^*$ wird definiert durch

$$\begin{aligned} [] \triangleright a &= a \triangleleft [] \\ (h \triangleleft t) \triangleright a &= h \triangleleft (t \triangleright a) \end{aligned}$$

Auch dies ist tatsächlich eine gültige rekursive Definition. Man kann sie, wenn auch wegen dem zweiten Parameter etwas umständlicher, auch mit dem Rekursionsoperator schreiben:

$$\triangleleft = \text{listrec}(t \mapsto q \mapsto h \mapsto a \mapsto h \triangleleft (t \triangleright a)) (a \mapsto a \triangleleft [])$$

3.5.14 BEMERKUNG. Man beachte, daß durch Vertauschen der Seiten auch die Definition von $\triangleleft$ durch $\triangleright$ entsteht:

$$\begin{aligned} h \triangleleft &= [] \triangleright h \\ h \triangleleft (t \triangleright a) &= (h \triangleleft t) \triangleright a \end{aligned}$$

3.5.15 SATZ. Für $a : A$ und $\ell : A^*$ gilt

$$|\ell \triangleright a| = S |\ell|.$$

Beweis. Wir beweisen mit Listeninduktion nach ℓ .

Induktionsanfang: Zu zeigen ist

$$|[] \triangleright a| = S |[]|.$$

Tatsächlich gilt

$$\begin{aligned} |[] \triangleright a| &= |a \triangleleft []| && \text{(Definition von } \triangleright \text{)} \\ &= S |[]| && \text{(Definition von length).} \end{aligned}$$

Induktionsschritt: Für $t: A^*$ verwenden wir die Induktionshypothese

$$|t \triangleright a| = S |t|$$

und zeigen, daß dann für alle $h: A$ auch

$$|(h \triangleleft t) \triangleright a| = S |h \triangleleft t|$$

gilt. Tatsächlich ergibt sich

$$\begin{aligned} |(h \triangleleft t) \triangleright a| &= |h \triangleleft (t \triangleright a)| && \text{(Definition von } \triangleright \text{)} \\ &= S |t \triangleright a| && \text{(Definition von length)} \\ &= S (S |t|) && \text{(Induktionshypothese)} \\ &= S |h \triangleleft t| && \text{(Def. von length, von rechts nach links).} \end{aligned}$$

□

3.5.16 DEFINITION. Wir definieren eine Funktion *reverse*: $A^* \rightarrow A^*$ die die Reihenfolge der Elemente in einer Liste umkehrt, z.B. $\text{reverse}[1, 2, 3, 4] = [4, 3, 2, 1]$. Sei A ein beliebiger Datentyp. Die rekursiven Gleichungen sind leicht hinzuschreiben:

$$\begin{aligned} \text{reverse}[] &= [], \\ \text{reverse}(h \triangleleft t) &= \text{reverse } t \triangleright h. \end{aligned}$$

Auch diese Definition läßt sich mittels Rekursionsoperator ausdrücken (Übung).

3.5.17 SATZ. Für alle $\ell: A^*$ gilt

$$|\text{reverse } \ell| = |\ell|.$$

Beweis. Einfache Übung für einen Beweis mit Listeninduktion. □

Listenverkettung

3.5.18 BEISPIEL. Wir definieren die *Listenverkettung* als eine binäre Operation $\diamond: A^* \rightarrow A^* \rightarrow A^*$, welche zwei Listen zusammenhängt, z.B. $[1, 2, 3] \diamond [5, 6] = [1, 2, 3, 5, 6]$. Sie erfüllt die Gleichungen

$$\begin{aligned} [] \diamond \ell &= \ell \\ (h \triangleleft t) \diamond \ell &= h \triangleleft (t \diamond \ell) \end{aligned}$$

und ist damit durch eine primitive Rekursion über Listen wohldefiniert.

3.5.19 SATZ. Für alle $u, v, w: A^*$ gilt

$$\begin{aligned} (u \diamond v) \diamond w &= u \diamond (v \diamond w); \\ [] \diamond u &= u = u \diamond []. \end{aligned}$$

Beweis. Übung (Induktion nach u). □

Die Listenverkettung ist somit eine assoziative Operation und die leere Liste ist ein neutrales Element. Man sagt auch, sie bildet ein Monoid.

3.5.20 SATZ. Sei A eine beliebige Menge. Dann gilt stets

$$|\ell \diamond r| = |\ell| + |r|.$$

Beweis. Wir verwenden Induktion nach ℓ .

Induktionsanfang: Zu zeigen ist $|\square \diamond r| = |\square| + |r|$, was wegen der Definitionen klar ist.

Induktionshypothese: $|\ell \diamond r| = |\ell| + |r|$. Zu zeigen: $|(a \triangleleft \ell) \diamond r| = |a \triangleleft \ell| + |r|$. Wir rechnen:

$$\begin{aligned} |(a \triangleleft \ell) \diamond r| &= |a \triangleleft (\ell \diamond r)| \\ &= S \, |\ell \diamond r| \\ &= S(|\ell| + |r|) \\ &= (S \, |\ell|) + |r| \\ &= |a \triangleleft \ell| + |r|. \end{aligned}$$

□

3.5.21 DEFINITION. Seien A und B beliebige Datentypen und $f: A \rightarrow B$. Dann läßt sich f durch „elementweises“ Anwenden auf eine Funktion $f^*: A^* \rightarrow B^*$ hochziehen:

$$\begin{aligned} f^* \square &= \square \\ f^* (x \triangleleft \ell) &= (fx) \triangleleft f^* \ell. \end{aligned}$$

Die Funktion wird oft auch mit map bezeichnet und hat den Typ $\text{map}: (A \rightarrow B) \rightarrow (A^* \rightarrow B^*)$.

3.5.22 NOTATION. Eine weitere Notation für map ist

$$[fx \mid x \leftarrow \ell] := \text{map } f \ell = f^* \ell,$$

welche besonders dann praktisch ist, wenn f durch einen Ausdruck wie z.B. in $f = (x \mapsto x^2)$ gegeben ist, weil man dann einfach $[x^2 \mid x \leftarrow [2, 3, 4]]$ schreiben kann. Diese Notation (list comprehension) ist auf vielfältige Weise erweiterbar.

3.5.23 SATZ (FUNKTOR-EIGENSCHAFT). Es seien A, B, C beliebige Mengen. Dann gilt für alle $f: A \rightarrow B$ und $g: B \rightarrow C$:

$$\begin{aligned} (\text{id}_A)^* &= \text{id}_{A^*} \\ (g \circ f)^* &= g^* \circ f^* \end{aligned}$$

Beweis. Es geht um die Gleichheit von Funktionen. Wir setzen auf beiden Seiten ein beliebiges Element ein (vgl. Extensionalität) und rechnen nach.

Für die erste Gleichung haben wir zu zeigen:

$$\bigwedge_{\ell: A^*} (\text{id}_A)^* \ell = \text{id}_{A^*} \ell$$

Wir beweisen mit Induktion nach ℓ . Für den Induktionsanfang ist zu zeigen: $(\text{id}_A)^* \square = \text{id}_{A^*} \square$; die linke Seite ist gemäß der Definition von map gleich $\square$, die rechte Seite gemäß der Definition von id . Für den Induktionsschritt verwenden

wir die Induktionshypothese $\text{id}_{A^*} t = \text{id}_A^* t$ und zeigen, daß dann auch, für beliebiges $h: A$, gilt $(\text{id}_A)^*(h \triangleleft t) = \text{id}_{A^*}(h \triangleleft t)$:

$$\begin{aligned} (\text{id}_A)^*(h \triangleleft t) &= \text{id}_A h \triangleleft (\text{id}_A)^* t && \text{(Definition von map)} \\ &= \text{id}_A h \triangleleft \text{id}_{A^*} t && \text{(Induktionshypothese)} \\ &= h \triangleleft t && \text{(Definition von id)} \\ &= \text{id}_{A^*}(h \triangleleft t) && \text{(Definition von id)} \end{aligned}$$

Für die zweite Gleichung haben wir zu zeigen:

$$\bigwedge_{\ell: A^*} (g \circ f)^* \ell = (g^* \circ f^*) \ell,$$

was wir wiederum mit Induktion nach ℓ zeigen. (**Übung!**) □

Aufgrund der obigen Eigenschaft nennt man map einen *Funktor*.

3.5.24 SATZ. *Ein paar weitere nette Eigenschaften von map:*

$$\begin{aligned} \text{length} \circ f^* &= \text{length}; \\ f^*(\ell \diamond r) &= f^* \ell \diamond f^* r; \\ f^* \circ \text{concat} &= \text{concat} \circ (f^*)^*. \end{aligned}$$

Beweis. **Übung!** □

Auswahl

3.5.25 DEFINITION. Sei $b: A \rightarrow \mathbb{B}$. Dann definieren wir

$$\begin{aligned} \text{filter } b \ [] &= [] \\ \text{filter } b (a \triangleleft \ell) &= \text{if } ba \text{ then } a \triangleleft \text{filter } b \ell \text{ else filter } b \ell \end{aligned}$$

Der Funktor map verändert normalerweise die Elemente einer Liste, läßt aber deren Länge unverändert. filter dagegen verändert normalerweise die Länge einer Liste, ändert aber nichts an den Elementen selbst.

3.5.26 NOTATION.

$$[a \mid a \leftarrow \ell, ba] := \text{filter } b \ell$$

Diese Notation kann mit der ähnlichen für map nach Belieben kombiniert werden, insbesondere:

$$[fa \mid a \leftarrow \ell, ba] := \text{map } f (\text{filter } b \ell)$$

und noch deutlich allgemeinere Varianten, mit stets offensichtlicher Bedeutung.

Beispiel: $[x^2 \mid x \leftarrow [1, 2, 3, 4, 5, 6, 7, 8, 9], x \text{ ist Primzahl}] = [4, 9, 25, 49]$.

3.5.27 DEFINITION. Sei A eine Menge auf der eine Addition definiert ist. Dann wird die *Summe* aller Elemente einer Liste definiert durch:

$$\begin{aligned} \sum [] &= 0 \\ \sum (a \triangleleft \ell) &= a + \sum \ell. \end{aligned}$$

Anders geschrieben: $\sum = \text{listrec} + 0$. Die Notation wird üblicherweise folgendermaßen modifiziert:

$$\sum_{k=1}^n a_k = \sum [a_k \mid k \leftarrow [1..n]]$$

Wir haben daher insbesondere:

$$\begin{aligned} \sum_{k=1}^0 &= 0 \\ \sum_{k=1}^{Sn} a_k &= \sum_{k=1}^n a_k + a_{Sn} \end{aligned}$$

3.5.28 SATZ. *Einige elementare Beziehungen für Summen:*

$$\begin{aligned} \sum_{k=1}^n (a_k + b_k) &= \sum_{k=1}^n a_k + \sum_{k=1}^n b_k \\ \sum_{k=1}^n a_k \cdot \sum_{\ell=1}^m b_k &= \sum_{k=1}^n \sum_{\ell=1}^m a_k b_\ell \\ \sum_{k=1}^n a_k &= \sum_{k=0}^{n-1} a_{k+1} \end{aligned}$$

Beweis. Alle derartigen Aussagen lassen sich recht einfach mittels Induktion nach n aus den elementaren Eigenschaften der Grundrechnungsarten gewinnen. $\square$

Geordnete Listen

Sei A eine beliebige Menge, auf der eine Ordnung $\preceq: A^* \rightarrow A^* \rightarrow \Omega$ definiert ist.

3.5.29 DEFINITION. Das Prädikat $\text{issorted}: A^* \rightarrow \Omega$ sollte ausdrücken, daß eine Liste sortiert ist. Es sei induktiv durch die folgenden Regeln definiert:

$$\begin{aligned} &\text{issorted}[] \\ &\text{issorted}(a \triangleleft []) \\ &\text{issorted}(h \triangleleft t) \implies a \preceq h \implies \text{issorted}(a \triangleleft h \triangleleft t) \end{aligned}$$

Dies ist so zu verstehen, daß umgekehrt eine Liste nur dann sortiert sein sollte, wenn dies aus einer dieser Regeln folgt. Daraus ergibt sich auf natürliche Weise wieder ein Induktionsprinzip:

$$P([]) \wedge \bigwedge_{a: A} \wedge \bigwedge_{t: A^*} \bigwedge_{h \triangleleft t} (\text{issorted}(h \triangleleft t) \implies \bigwedge_{a: A} (a \preceq h \implies P(a \triangleleft h \triangleleft t))) \implies \bigwedge_{\ell: A^*} (\text{isorted}(\ell) \implies P(\ell))$$

3.5.30 DEFINITION. Die Relation $A^* \rightarrow A^* \rightarrow \Omega$ sollte ausdrücken, daß zwei Listen bis auf die Reihenfolge ihrer Elemente übereinstimmen. Sie sei induktiv

durch die folgenden Regeln definiert. (bezeichnet mit $\sim$), definieren wir drei Regeln:

$$\begin{aligned} \square &\sim \square \\ t_1 \sim t_2 &\implies z \triangleleft t_1 \sim z \triangleleft t_2 \\ t_1 \sim t_2 &\implies n \triangleleft p \triangleleft t_1 \sim p \triangleleft n \triangleleft t_2 \end{aligned}$$

3.5.31 SATZ. *Sei A eine beliebige Menge, auf der eine entscheidbare lineare Ordnung $\preceq$ definiert ist, und $\ell: A^*$. Dann gibt es genau eine Liste ℓ' sodaß $\ell \sim \ell'$ und issorted ℓ' .*

Quicksort. Die folgende Funktion liefert einen konstruktiven Beweis:

$$\begin{aligned} \text{sort } \square &= \square \\ \text{sort}(a \triangleleft \ell) &= \text{sort}[x \mid x \leftarrow \ell, x \preceq a] \diamond [a] \diamond \text{sort}[x \mid x \leftarrow \ell, x \succ a] \end{aligned}$$

Die Rekursion führt hier über die Länge der Liste. □