

MATHEMATIK III FÜR VERMESSUNGSWESEN

a.o.Univ.Prof. Dr. Michael Drmota

Wintersemester 2000/2001

Vorwort

Die Vorlesung *Mathematik III für Vermessungswesen* wird seit dem Wintersemester 1993/94 in dieser Fassung gelesen. Ihr Ziel ist es insbesondere, *diskrete Strukturen* und *Verfahren* vorzustellen. Außerdem wird besonderen Wert auf den *algorithmischen Aspekt* gelegt, um die Verbindung zu Computeranwendungen herzustellen.

Im ersten Kapitel werden die Grundbegriffe der *Aussagenlogik* bereitgestellt, um einerseits die Sprache der *Mengenlehre* (Kapitel 2) und in der Folge alle weiteren mathematischen Begriffe und Zusammenhänge prägnanter und fundierter formulieren zu können. Kernstück des zweiten Kapitels ist die disjunktive (resp. konjunktive) Normalform von Mengenausdrücken. Diese algorithmisch einfach zu bestimmenden Normalformen treten im zweiten Teil der *elementaren Aussagenlogik* (Kapitel 7) unter einem leicht veränderten Blickwinkel wieder auf und finden Anwendungen in der Schaltalgebra.

Der Begriff der *Relation* ist in der Mathematik von zentraler Bedeutung. Beispielsweise sind *Funktionen* (Kapitel 4) und *Graphen* (Kapitel 5) “nur” Spezialfälle von Relationen. Im dritten Kapitel wird insbesondere mit dem *Warshall-Algorithmus* zur Bestimmung der transitiven Hülle einer Relation die Erreichbarkeitsrelation auf einem Graphen vorbereitet.

Die im fünften Kapitel behandelte *Graphentheorie* ist eines der größten Gebiete der *diskreten Mathematik*. Gerade in diesem Kapitel steht der algorithmische Aspekt besonders im Vordergrund. Es werden der *Markierungsalgorithmus* zur Untersuchung auf *Azyklichkeit*, der *Kruskal-Algorithmus* zur Bestimmung eines minimalen Gerüsts, der *Dijkstra-Algorithmus* zur Ermittlung von kürzesten Wegen und der *Ford-Fulkerson-Algorithmus* zur Berechnung eines maximalen Flusses besprochen.

Das Vordringen des Computers hat insbesondere auch *algebraischen Methoden* neue Anwendungsbereiche erschlossen. Besonders prominent ist hier die hauptsächlich auf Gruppen- und Körpertheorie basierende *algebaische Codierungstheorie*. Andererseits finden sich *zahlentheoretische Methoden* bei *Verschlüsselungsverfahren*, wie z.B. beim RSA-Verfahren. Schließlich bedient sich die *formale Logik*, die Grundlage der *theoretischen Informatik*, der Sprache der *allgemeinen Algebra*.

In der *Aussagenlogik* finden sich vor allem Konzepte der *Booleschen Algebra* wieder, wie z.B. die disjunktive (resp. konjunktive) Normalform sowie der *Algorithmus von Quine-McCluskey*. Letzterer dient zur Optimierung von *logischen Schaltkreisen*.

Dieses Skriptum ist nur ein Skelettskriptum. Beweise werden nicht ausgeführt. Das Lesen

des Skriptums kann daher den Besuch der Vorlesung nicht ersetzen. Es ist dazu gedacht, die Vorlesungsmitschrift zu ergänzen.

Gelegentlich wurden erläuternde Zwischenbemerkungen in kleingedruckter Schrift eingefügt. Diese Teile gehören nicht zum Prüfungsstoff.

Ich möchte mich an dieser Stelle noch bei den Herren Univ.-Prof. Dr. Peter Kirschenhofer und Univ.Prof. Dr. Günter Karigl, die die Vorlesung in den Jahren 1993–1995 gelesen haben, für die Überlassung ihrer Vorlesungsunterlagen, von denen ich besonders bei der Verfassung dieses Skriptums sehr profitiert habe, bedanken.

Wien, im September 2000

Michael Drmota

Inhaltsverzeichnis

Vorwort	I
1 Elementare Aussagenlogik I	1
1.1 Verknüpfungen von Aussagen	1
1.2 Boolesche Variable, Tautologien	3
1.3 Prädikatenlogik und Quantoren	5
2 Mengen	8
2.1 Grundbegriffe	8
2.2 Operationen mit Mengen	9
2.3 Charakteristische Funktion	11
2.4 Elementtabelle	13
2.5 Disjunktive und konjunktive Normalform	14
2.6 Potenzmenge	17
2.7 Stirlingsche Approximationsformel	18
3 Relationen	20
3.1 Grundlegende Begriffe	20
3.2 Äquivalenzrelation	22
3.3 Ordnungsrelation	24
3.4 Transitive Hülle einer Relation	26
3.5 Warshall-Algorithmus	29
4 Funktionen	31
4.1 Injektive, surjektive und bijektiver Funktionen	31

4.2	Permutationen	33
4.3	Abzählbare Mengen	36
5	Graphen	38
5.1	Grundlegende Begriffe	38
5.2	Adjazenz- und Inzidenzmatrix	41
5.3	Azyklische Graphen	43
5.4	Zusammenhang von Graphen	44
5.4.1	Ungerichtete Graphen	44
5.4.2	Gerichtete Graphen	45
5.5	Bäume und Wälder	47
5.6	Planare Graphen	52
5.7	Eulersche und Hamiltonsche Linien	54
5.8	Algorithmen für Netzwerke	54
5.8.1	Minimales Gerüst: Kruskal-Algorithmus	55
5.8.2	Kürzester Weg: Dijkstra-Algorithmus	56
5.8.3	Maximaler Fluß: Ford-Fulkerson-Algorithmus	59
6	Algebraische Strukturen	64
6.1	Gruppen	64
6.2	Ringe und Körper	72
6.3	Codierungstheorie	78
6.3.1	Grundlegende Begriffe	78
6.3.2	Gruppencodes	80
6.3.3	Das RSA-Verfahren	82
6.4	Verbände und Boolesche Algebren	83
6.5	Allgemeine Algebren	86
7	Elementare Aussagenlogik II	90
7.1	Konjunktive und disjunktive Normalformen für Formeln	90
7.2	Boolesche Funktionen, Schaltalgebra	91
7.3	Optimierung von logischen Schaltkreisen	93
7.3.1	Quine-McCluskey-Algorithmus	93

7.3.2 Hauptterme und absolut eliminierbare Terme	95
Literaturverzeichnis	96
Index	97
Übungsbeispiele	101

Kapitel 1

Elementare Aussagenlogik I

1.1 Verknüpfungen von Aussagen

Unter einer (**mathematischen**) **Aussage**¹ versteht man einen sprachlichen Ausdruck, dem eindeutig ein **Wahrheitswert** **wahr** (= **w**) oder **falsch** (= **f**) zugeordnet werden kann. Üblicherweise wird eine *Aussage* durch einen *Aussagesatz* formuliert, wie z.B.:

Es regnet.

Die Straße ist naß.

Es ist Winter.

Es ist kalt.

Man beachte, daß man hier von dem Prinzip ausgeht, daß eine Aussage nur entweder wahr oder falsch sein kann. Dieses Prinzip heißt **Prinzip vom ausgeschlossenen Dritten** (Terium non datur). Man spricht daher auch oft von **zweiwertiger Aussagenlogik**.

Aussagen können auf verschiedene Arten und Weisen miteinander verknüpft werden.

1. **Konjunktion**: Aus zwei Aussagen wird durch Einfügen des Wortes **und** eine neue Aussage gewonnen, die Konjunktion der beiden Aussagen. Z.B. ist

Es ist Winter **und** *es ist kalt.*

die Konjunktion der beiden Aussagen

Es ist Winter. *Es ist kalt.*

Die Konjunktion zweier Aussagen erhält genau dann den Aussagewert **w**, wenn die beiden ursprünglichen Aussagen den Wert **w** haben. (In allen anderen Fällen erhält die Konjunktion den Wert **f**.)

¹proposition

2. **Disjunktion:** Aus zwei Aussagen wird durch Einfügen des Wortes **oder** eine neue Aussage gewonnen, die Disjunktion der beiden Aussagen. Z.B. ist

*Die Straße ist naß **oder** es regnet.*

die Disjunktion der beiden Aussagen

Die Straße ist naß. Es regnet.

Die Disjunktion zweier Aussagen erhält genau dann den Aussagewert **w**, wenn wenigstens eine der ursprünglichen Aussagen den Wert **w** hat. (Nur wenn beide Aussagen den Wert **f** haben, hat auch die Disjunktion den Wert **f**.)

3. **Implikation:** Aus zwei Aussagen wird durch Einfügen der Worte **wenn - dann** eine neue Aussage gewonnen, die Implikation der beiden Aussagen. Z.B. ist

***Wenn** es regnet, **dann** ist die Straße naß.*

die Implikation der beiden Aussagen

Es regnet. Die Straße ist naß.

Die Implikation zweier Aussagen hat genau dann den Wert **f**, wenn die erste Aussage den Wert **w**, aber die zweite den Wert **f** hat. Man beachte, daß bei der Implikation die Reihenfolge wesentlich ist. Weiters hat eine Implikation immer den Wert **w**, wenn die erste Aussage den Wert **f** hat, unabhängig davon, wie die zweite Aussage lautet (ex falso quodlibet).

Statt der Worte *wenn - dann* kann man auch *aus - folgt* oder *impliziert* verwenden.

4. **Äquivalenz:** Aus zwei Aussagen wird durch Einfügen der Worte **genau dann - wenn** eine neue Aussage gewonnen, die Äquivalenz der beiden Aussagen. Z.B. ist

*Die Straße ist **genau dann** naß **wenn** es regnet.*

die Äquivalenz der beiden Aussagen

Die Straße ist naß. Es regnet.

Die Äquivalenz zweier Aussagen hat genau dann den Wert **w**, wenn den beiden ursprünglichen Aussagen dieselben Wahrheitswerte zugeordnet sind.

Statt der Worte *genau dann - wenn* werden auch die Worte *dann und nur dann - wenn* oder *äquivalent zu* verwendet.

5. **Negation:** Fügt man in einer Aussage (an geeigneter Stelle) das Wort **nicht** ein, so entsteht eine neue Aussage, die Negation der ursprünglichen Aussage. Z.B. ist

*Es ist **nicht** Winter.*

die Negation der Aussage

Es ist Winter

Eine negierte Aussage hat genau dann den Wert **w**, wenn die unnegierte (ursprüngliche) Aussage den Wert **f** hat.

Zur Vereinfachung der Notation ist es günstig, Aussagen durch *Symbole* $p_1, p_2, \dots$ zu bezeichnen und anstelle von *und* das Symbol $\wedge$, anstelle von *oder* das Symbol $\vee$, anstelle von *wenn – dann* das Symbol $\rightarrow$, anstelle von *genau dann, wenn* das Symbol $\leftrightarrow$ und anstelle von *nicht* das Symbol $\neg$ zu verwenden. Die Symbole $\wedge, \vee, \rightarrow, \leftrightarrow, \neg$ werden in diesem Zusammenhang als **Junktoren** bezeichnet. Ist z.B.

$p_1 = \text{Die Straße ist naß.}$
 $p_2 = \text{Es regnet.}$
 $p_3 = \text{Die Straße wird gereinigt.}$

so kann die Aussage

Wenn die Straße naß ist, dann regnet es oder die Straße wird gereinigt.

durch

$$p_1 \rightarrow (p_2 \vee p_3)$$

formalisiert werden.

Die Operationssymbole $\wedge, \vee, \rightarrow, \leftrightarrow, \neg$ können sinnvollerweise auch auf die Wahrheitswerte **w** und **f** angewandt werden.

$\wedge$	f	w	$\vee$	f	w	$\rightarrow$	f	w	$\leftrightarrow$	f	w	$\neg$	f	w
f	f	f	f	f	w	f	w	w	f	w	f	w	f	w
w	f	w	w	w	w	w	f	w	w	f	w	w	w	f

Durch diese Wahl ist sichergestellt, daß der Wahrheitswert einer durch Verknüpfungen von Einzelaussagen $p_1, p_2, \dots$ gewonnenen Aussage dadurch bestimmt werden kann, daß man $p_1, p_2, \dots$ durch ihre Wahrheitswerte ersetzt und die Junktoren als Operationssymbole für **w** und **f** interpretiert.

1.2 Boolesche Variable, Tautologien

Um noch ein wenig allgemeiner sein zu können führen wir folgende Definitionen ein:

Definition 1.1 Eine **Boolesche Variable** ist eine Variable, die Werte **w** und **f** annehmen kann.

Eine (**wohlgeformte**) **Formel** F ist ein Ausdruck, der sich in endlich vielen Schritten aus Booleschen Variablen und den Operationen $\wedge, \vee, \rightarrow, \leftrightarrow, \neg$ aufbauen läßt.

Eine wohlgeformte Formel F heißt **gültig** oder **Tautologie**², wenn sie für jede Belegung der Booleschen Variablen mit Werten aus **B** immer den Wert **w** hat.

Eine wohlgeformte Formel F heißt **erfüllbar**, wenn es mindestens eine Belegung der Booleschen Variablen gibt, so daß der Wert **w** ist.

Eine wohlgeformte Formel F heißt **nicht erfüllbar**, wenn ihr Wert für alle möglichen Belegungen der Booleschen Variablen **f** ist.

Um eine Formel zu analysieren, kann man, da in ihr nur endlich viele Boolesche Variable auftreten, eine **Wahrheitstafel** aufstellen. (Das Verfahren ist ganz analog zur Aufstellung einer Elementtabelle für Mengenausdrücke, die im Kapitel 2 besprochen werden.)

Beispiel 1.2 Für die Formel $F = b \leftrightarrow (a \rightarrow (\neg a \vee b))$ erhält man folgende Wahrheitstafel:

a	b	$b \leftrightarrow (a \rightarrow (\neg a \vee b))$
w	w	w
w	f	f
f	w	w
f	f	w

In der zweiten Zeile wird etwa jener Fall diskutiert, wo a den Wert **w** und b den Wert **f** annimmt. $\neg a$ hat dann den Wert **f**, $\neg a \vee b$ den Wert **f**, $a \rightarrow (\neg a \vee b)$ den Wert **f** und schließlich $b \leftrightarrow (a \rightarrow (\neg a \vee b))$ den Wert **w**. Es zeigt sich, daß die betrachtete Formel (in 3 von 4 Fällen) erfüllbar ist, aber keine Tautologie ist.

Zwei i.a. verschiedene Formeln können durchaus immer denselben Wahrheitswert haben. Sie sind von einem gewissen Standpunkt aus ununterscheidbar.

Definition 1.3 Zwei (wohlgeformte) Formeln F_1, F_2 heißen (**semantisch**) **äquivalent**, i.Z. $F_1 \iff F_2$, wenn die Formel $F = (F_1 \leftrightarrow F_2)$ eine Tautologie ist.

Zwei Formeln sind also genau dann äquivalent, wenn sie bei jeder Wahrheitsbelegung der vorkommenden Booleschen Variablen denselben Wahrheitswert liefern.

In ähnlicher Weise wird die (semantische) Implikation definiert.

Definition 1.4 Eine Formel F_1 impliziert eine Formel F_2 , i.Z. $F_1 \implies F_2$, wenn die Formel $F = (F_1 \rightarrow F_2)$ eine Tautologie ist.

Dies bedeutet, daß F_2 immer wahr ist, wenn F_1 wahr ist.

²tautology

Satz 1.5 Für Boolesche Variable a, b, c gelten die folgenden (semantischen) Äquivalenzen und Implikationen.

- | | |
|---|--|
| 1. $a \wedge b \iff b \wedge a,$ | $a \vee b \iff b \vee a,$ |
| 2. $a \wedge (b \wedge c) \iff (a \wedge b) \wedge c,$ | $a \vee (b \vee c) \iff (a \vee b) \vee c,$ |
| 3. $a \wedge (a \vee b) \iff a,$ | $a \vee (a \wedge b) \iff a,$ |
| 4. $a \wedge (b \vee c) \iff (a \wedge b) \vee (a \wedge c),$ | $a \vee (b \wedge c) \iff (a \vee b) \wedge (a \vee c),$ |
| 5. $\neg(a \wedge b) \iff \neg a \vee \neg b,$ | $\neg(a \vee b) \iff \neg a \wedge \neg b,$ |
| 6. $a \wedge b \iff \neg(a \rightarrow \neg b),$ | $a \vee b \iff \neg a \rightarrow b,$ |
| $a \leftrightarrow b \iff (a \rightarrow b) \wedge (b \rightarrow a)$ | |
| $\iff \neg((a \rightarrow b) \rightarrow \neg(b \rightarrow a)),$ | |
| 7. $a \rightarrow b \iff \neg a \vee b,$ | $a \leftrightarrow b \iff (a \wedge b) \vee (\neg a \wedge \neg b),$ |
| 8. $a \wedge b \implies a,$ | $a \implies a \vee b,$ |
| 9. $a \wedge (a \rightarrow b) \implies b.$ | |

Setzt man etwa $1 = a \vee \neg a$ und $0 = a \wedge \neg a$, und interpretiert man $\iff$ als *Gleichheitszeichen*, so erfüllen die wohlgeformten Formeln wegen 1.–5. alle Gesetze einer **Booleschen Algebra**, siehe Kapitel 6.

6. sagt aus, daß jede wohlgeformte Formel semantisch äquivalent zu einer wohlgeformten Formel ist, in der nur die Junktoren $\rightarrow$ und $\neg$ vorkommen. Andererseits kann man wegen 7. auch eine semantisch äquivalente Formel finden, die nur die Junktoren $\wedge$, $\vee$ und $\neg$ verwendet.

9. ist der sogenannte **modus ponens** (Abtrennungsregel).

1.3 Prädikatenlogik und Quantoren

In der Mathematik haben sogenannte *atomare* (d.h. unzerlegbare) Aussagen eine *Subjekt-Prädikatstruktur*, z.B. kann in

Alexander ist groß.

Alexander als Subjekt und *groß* als Prädikat interpretiert werden. Die Aussage

Das Haus ist groß.

unterscheidet sich von der ersten nur im Subjekt. Es ist daher naheliegend, das Prädikat *groß* zu einem Symbol P zu abstrahieren und die **Aussageform** $P(x)$

x ist groß.

zu betrachten, wobei jetzt x als **Gegenstandsvariable** fungiert. Setzt man für x einen Wert ein, so entsteht aus der Aussageform $P(x)$ wieder eine Aussage, z.B. $P(\text{Alexander})$, $P(\text{Haus})$, der wieder ein Wahrheitswert zugeordnet werden kann. Andererseits kann man Aussagen der Form

Für alle ... gilt ... und Es gibt ein ... so daß ...

mit Hilfe von Aussageformen bilden:

Für alle x gilt $P(x)$. Es gibt ein x so daß $P(x)$.

Formalisiert wird dies durch die **Quantoren**, den **Allquantor** $\forall$ und den **Existenzquantor** $\exists$. So bedeutet

$$\forall x P(x),$$

daß alle möglichen x die Eigenschaft P haben und

$$\exists x P(x),$$

daß es wenigstens ein x gibt, das die Eigenschaft P hat.

Der Wahrheitswert von $\forall x P(x)$ ist genau dann **w**, wenn die Aussage $P(x)$ für alle möglichen x der Wert **w** hat, entsprechend hat $\exists x P(x)$ genau dann den Wert **w**, wenn es wenigstens ein x gibt, für das $P(x)$ den Wert **w** hat.

Eine Aussageform $P(x)$ heißt auch **einstelliges Prädikat**. Entsprechend werden auch **mehrstellige Prädikate** $Q(x_1, x_2, \dots, x_n)$ verwendet. Beispielsweise ist

x ist größer als y

ein zweistelliges Prädikat.

Selbstverständlich kann man Prädikate mit Junktoren untereinander bzw. mit gewöhnlichen Aussagen verbinden und, solange noch **freie Gegenstandsvariable** vorhanden sind, Quantoren anwenden. Führt man dies mit Aussagensymbolen $(p_1, p_2, \dots)$ und Prädikatsymbolen $(P_1, P_2, \dots)$ durch, so erhält man eine sogenannte **Formel**.

Beispielsweise ist

$$\forall x_1 ((P_1(x_1) \wedge p_1) \rightarrow (\exists x_2 (P_2(x_2) \wedge p_2) \rightarrow P_3(x_1, x_2)))$$

so eine Formel, in der die beiden auftretenden Gegenstandsvariablen gebunden sind.

Üblicherweise werden Quantoren in mathematischen Aussagen verwendet, um über *Elemente* einer *Menge* (siehe Kapitel 2) eine Aussage zu treffen. Man verwendet

$$\forall x \in E : P(x)$$

als Kurzschreibweise für $\forall x ((x \in E) \rightarrow P(x))$ und

$$\exists x \in E : P(x)$$

als Kurzschreibweise für $\exists x ((x \in E) \wedge P(x))$.

Beispiel 1.6 Sei P_1 das Prädikat *ist ein Mensch* und P_2 das Prädikat *ist sterblich*, so wird die Aussage *Alle Menschen sind sterblich* durch

$$\forall x (P_1(x) \rightarrow P_2(x))$$

formalisiert.

In ähnlicher Weise wie in der Aussagenlogik gibt es auch in der Prädikatenlogik (semantische) Äquivalenzen und Implikationen:

Satz 1.7 *Es gelten die folgenden (semantischen) Äquivalenzen und Implikationen.*

1. $\forall x \forall y P(x, y) \iff \forall y \forall x P(x, y), \quad \exists x \exists y P(x, y) \iff \exists y \exists x P(x, y),$
2. $a \wedge \forall x P(x) \iff \forall x (a \wedge P(x)), \quad a \vee \exists x P(x) \iff \exists x (a \vee P(x)),$
3. $a \wedge \exists x P(x) \iff \exists x (a \wedge P(x)), \quad a \vee \forall x P(x) \iff \forall x (a \vee P(x)),$
4. $\neg(\forall x P(x)) \iff \exists x \neg P(x), \quad \neg(\exists x P(x)) \iff \forall x \neg P(x),$
5. $\forall x P(x) \implies P(x_0), \quad P(x_0) \implies \exists x P(x),$
6. $\forall x (P(x) \rightarrow Q(x)) \wedge P(x_0) \implies Q(x_0),$
7. $\exists x \forall y P(x, y) \implies \forall y \exists x P(x, y).$

Kapitel 2

Mengen

2.1 Grundbegriffe

Die Mengenlehre wurde vor etwa 100 Jahren von Georg Cantor begründet. Er benutzte damals die folgende Definition, die zwar streng formal widersprüchlich ist, sich für unsere Anwendungen aber durchaus als zweckmäßig und ausreichend erweist.

Definition 2.1 (Cantor) *Eine Menge¹ ist eine Zusammenfassung von wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens zu einem Ganzen.*

Beispielsweise stellt sich heraus, daß die *Menge aller Mengen*, die nach der obigen Definition eine Menge sein müßte, ein widersprüchlicher Begriff ist. Formal wurde dieser Widerspruch dadurch gelöst, daß die Mengenlehre streng axiomatisch aufgebaut wurde (Axiomensystem von Zermelo und Fraenkel). Noch einfacher ist es, eine große Menge, ein **Universum** E , vorauszusetzen und nur Teilmengen des Universums zu betrachten. Dadurch können keine Widersprüche wie zuvor entstehen.

Beispiel 2.2 Die Zahlen 1, 2, 3 bilden eine Menge $A = \{1, 2, 3\}$, ebenso wie die ganzen Zahlen $\mathbb{Z}$.

Definition 2.3 *Die Objekte x , die in einer Menge A zusammengefaßt werden, bezeichnet man als **Elemente**² der Menge A . Man sagt auch, daß x in A enthalten ist und schreibt $x \in A$ ³. Ist x in A nicht enthalten, so schreibt man dafür $x \notin A$.*

*Die Menge $\emptyset$, die keine Elemente enthält, heißt **leere Menge**⁴.*

Definition 2.4 *Eine Menge A heißt **Teilmenge**⁵ einer Menge B , i.Z. $A \subseteq B$, wenn jedes Element x aus A auch in B enthalten ist.*

¹set

²element

³ x is contained in A

⁴empty set

⁵subset

Definition 2.5 Zwei Mengen A, B sind gleich, i.Z. $A = B$, wenn sie dieselben Elemente enthalten.

Satz 2.6 Zwei Mengen A, B sind genau dann gleich, wenn sowohl A Teilmenge von B als auch B Teilmenge von A ist, d.h.

$$A = B \iff A \subseteq B \wedge B \subseteq A.$$

Es gibt verschiedene Möglichkeiten, eine Menge anzugeben. Die einfachste Möglichkeit ist die **aufzählende Darstellung**, z.B. $A = \{1, 2, 3\}$, die sich aber nur für endliche (und gelegentlich für abzählbare) Mengen eignet. Die häufigste Form ist die **beschreibende Darstellung**

$$A = \{x \mid P(x)\},$$

wobei $P(x)$ ein *Prädikat* ist in der freien Variablen x ist. Die Menge A enthält nun jene x , für die $P(x)$ gilt. Beispielsweise ist $\{x \mid x \in \mathbf{Z} \wedge 1 \leq x \leq 3\}$ die Menge $\{1, 2, 3\}$. Verlangt man, daß die zu beschreibende Menge A Teilmenge einer Menge E sein soll, d.h. $P(x)$ hat die Form $x \in E$ und $Q(x)$, so wird anstelle $A = \{x \mid x \in E \wedge Q(x)\}$ einfach

$$A = \{x \in E \mid Q(x)\}$$

geschrieben.

Die Definition einer Menge impliziert, daß ein Element x nur *einmal* in eine Menge A aufgenommen werden kann, x ist entweder in A enthalten oder nicht in A enthalten. Es ist aber oft sinnvoll *verallgemeinerte Mengen* zu betrachten, wo die Elemente mit einer gewissen Vielfachheit auftreten, z.B. $A = \{1, 1, 2, 2, 2, 3, 4, 4\}$. Solche Objekte werden als **Multimengen**⁶ bezeichnet.

2.2 Operationen mit Mengen

Definition 2.7 Die **Vereinigung**⁷ $A \cup B$ zweier Mengen A, B enthält genau jene Elemente x , die in A oder in B enthalten sind, d.h.

$$A \cup B = \{x \mid x \in A \vee x \in B\}.$$

Definition 2.8 Der **Durchschnitt**⁸ $A \cap B$ zweier Mengen A, B enthält genau jene Elemente x , die sowohl in A als auch in B enthalten sind, d.h.

$$A \cap B = \{x \mid x \in A \wedge x \in B\}.$$

Es ist oft notwendig, nicht nur zwei oder endlich viele Mengen zu vereinigen, sondern ein ganzes System von Mengen zu vereinigen. Dafür verwendet man die folgende Notation.

⁶multiset

⁷union

⁸intersection

Definition 2.9 Sei I eine Menge (Indexmenge) und für alle $i \in I$ sei A_i eine Menge. Dann ist durch

$$\bigcup_{i \in I} A_i = \{x \mid \exists i \in I : x \in A_i\}$$

die **Vereinigung** aller A_i , $i \in I$, und

$$\bigcap_{i \in I} A_i = \{x \mid \forall i \in I : x \in A_i\}$$

der **Durchschnitt** aller A_i , $i \in I$.

Wie im Abschnitt 2.1 schon angedeutet wurde, ist es oft sinnvoll, nur Teilmengen einer großen Menge, eines Universums E zu betrachten. Die folgende Definition ist nur in einem solchen Kontext anzuwenden.

Definition 2.10 Sei $A \subseteq E$. Dann bezeichnet

$$A' = \{x \in E \mid x \notin A\}$$

das **Komplement**⁹ von A .

Satz 2.11 Seien A, B, C und A_i , $i \in I$, Teilmengen einer Menge E . Dann gelten die folgenden Rechenregeln.

- | | | |
|-----|--|--|
| 1. | $A \cup B = B \cup A,$ | $A \cap B = B \cap A,$ |
| 2. | $A \cup (B \cup C) = (A \cup B) \cup C,$ | $A \cap (B \cap C) = (A \cap B) \cap C,$ |
| 3. | $A \cap (B \cup C) = (A \cap B) \cup (A \cap C),$ | $A \cup (B \cap C) = (A \cup B) \cap (A \cup C),$ |
| 3.' | $A \cap \bigcup_{i \in I} A_i = \bigcup_{i \in I} (A \cap A_i),$ | $A \cup \bigcap_{i \in I} A_i = \bigcap_{i \in I} (A \cup A_i),$ |
| 4. | $A \cup \emptyset = \emptyset \cup A = A,$ | $A \cap E = E \cap A = A$ |
| 5. | $A \cup A' = E,$ | $A \cap A' = \emptyset$ |
| 6. | $A \cup A = A,$ | $A \cap A = A$ |
| 7. | $A \cup E = E,$ | $A \cap \emptyset = \emptyset$ |
| 8. | $A \cup (A \cap B) = A,$ | $A \cap (A \cup B) = A$ |
| 9. | $(A \cup B)' = A' \cap B',$ | $(A \cap B)' = A' \cup B'$ |
| 9.' | $(\bigcup_{i \in I} A_i)' = \bigcap_{i \in I} A_i',$ | $(\bigcap_{i \in I} A_i)' = \bigcup_{i \in I} A_i'$ |

1. nennt man auch **Kommutativgesetz**, 2. **Assoziativgesetz**, 3. **Distributivgesetz**, 8. **Verschmelzungsgesetz** und 9. **DeMorgansche Regel**.

Definition 2.12 Die **Mengendifferenz**¹⁰ $A \setminus B$ zweier Mengen A, B enthält genau jene Elemente x , die in A , aber nicht in B enthalten sind, d.h.

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

⁹complement

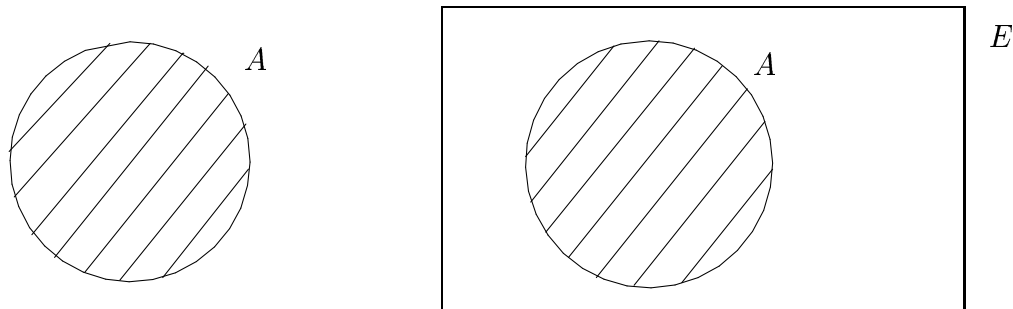
¹⁰set difference

Die **symmetrische Differenz**¹¹ $A \Delta B$ zweier Mengen A, B ist durch

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

gegeben.

Es ist oft nützlich, eine Menge A bildlich durch ein sogenanntes **Venn-Diagramm** darzustellen.



Auf diesem Weg lassen sich die Mengenoperationen Vereinigung, Durchschnitt, Komplement, Mengendifferenz und symmetrische Differenz auf einfache Art graphisch verdeutlichen.

2.3 Charakteristische Funktion

Definition 2.13 Sei A Teilmenge von E (dem Universum). Dann heißt die Funktion $\chi_A : E \rightarrow \{0, 1\}$, die durch

$$\chi_A(x) = \begin{cases} 1 & \text{für } x \in A, \\ 0 & \text{für } x \notin A \end{cases}$$

definiert ist, **charakteristische Funktion**¹² der Menge A .

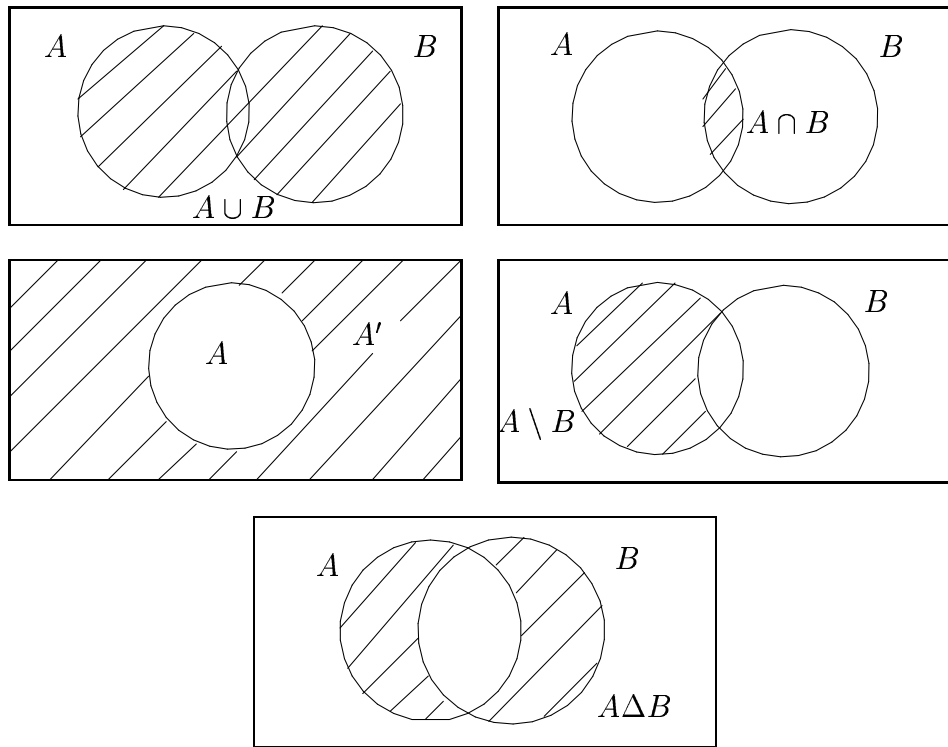
Man könnte anstelle von Mengen auch nur charakteristische Funktionen betrachten. Die charakteristische Funktion einer Menge A beschreibt die Menge A vollständig. Für jedes x kann entschieden werden, ob $x \in A$ gilt oder $x \notin A$:

$$A = \{x \mid \chi_A(x) = 1\}.$$

Es gilt daher auch die folgende Charakterisierung.

¹¹symmetric difference

¹²characteristic function



Satz 2.14 Zwei Teilmengen A, B von E sind genau dann gleich, wenn ihre charakteristischen Funktionen χ_A, χ_B gleich sind, d.h. für alle $x \in E$ gilt $\chi_A(x) = \chi_B(x)$:

$$A = B \iff \chi_A = \chi_B.$$

Entsprechend lassen sich Teilmengen charakterisieren.

Satz 2.15 Es seien A, B Teilmengen einer Menge E . A ist dann genau dann Teilmenge von B , wenn $\chi_A(x) \leq \chi_B(x)$ für alle $x \in E$ gilt:

$$A \subseteq B \iff \chi_A \leq \chi_B.$$

Weiters lassen sich die in Abschnitt 2.2 eingeführten Rechenoperationen auch mit Hilfe charakteristischer Funktionen beschreiben.

Satz 2.16 Es seien A, B Teilmengen einer Menge E . Dann gelten für die charakteristischen Funktionen die folgenden Rechenregeln:

$$\begin{aligned} \chi_{A \cup B}(x) &= \max(\chi_A(x), \chi_B(x)), \\ \chi_{A \cap B}(x) &= \chi_A(x) \cdot \chi_B(x) = \min(\chi_A(x), \chi_B(x)), \\ \chi_{A'}(x) &= 1 - \chi_A(x), \\ \chi_{A \setminus B}(x) &= \chi_A(x) - \chi_{A \cap B}(x) = \chi_A(x)(1 - \chi_B(x)), \\ \chi_{A \Delta B}(x) &= \chi_A(x) + \chi_B(x) \bmod 2. \end{aligned}$$

Charakteristische Funktionen dienen daher auch, mit Mengen in algebraischer Art und Weise umzugehen. Beispielsweise können Mengenidentitäten hergeleitet werden.

Beispiel 2.17 Möchte man etwa zeigen, daß

$$A \cap (B \Delta C) = (A \cap B) \Delta (A \cap C)$$

für alle Mengen A, B, C gilt, so braucht nur nachgewiesen zu werden, daß die entsprechenden charakteristischen Funktionen gleich sind. In diesem Fall ist dies sogar ganz leicht möglich:

$$\begin{aligned} \chi_{A \cap (B \Delta C)} &= \chi_A \cdot \chi_{B \Delta C} \\ &= \chi_A \cdot (\chi_B + \chi_C \bmod 2) \\ &= (\chi_A \cdot \chi_B) + (\chi_A \cdot \chi_C) \bmod 2 \\ &= \chi_{A \cap B} + \chi_{A \cap C} \bmod 2 \\ &= \chi_{(A \cap B) \Delta (A \cap C)}. \end{aligned}$$

Das dritte Gleichheitszeichen ergibt sich übrigens aus dem Distributivgesetz der modulo 2-Rechnung (Siehe Kapitel 6).

2.4 Elementtabelle

Ein nichtalgebraisches Verfahren zum Überprüfen von Mengenidentitäten ist eine **Elementtabelle**. Dieses Verfahren soll an der folgenden Identität

$$A \cap (B \Delta C) = (A \cap B) \Delta (A \cap C)$$

demonstriert werden.

Da zwei Mengen nach Definition 2.5 genau dann gleich sind, wenn jedes (potentielle) Element x dann und nur dann in der einen Menge enthalten ist, wenn es in der anderen enthalten ist, genuegt es, bei drei involvierten Mengen A, B, C acht Fälle zu unterscheiden, entsprechend ob ein x in A resp. B resp. C enthalten ist oder nicht.

Ist z.B. ein x in A und B , aber nicht in C enthalten, so ist es in A und in $B \Delta C$ und damit auch in $A \cap (B \Delta C)$ enthalten. Andererseits ist es in $A \cap B$, aber nicht in $A \cap C$ enthalten, womit es allerdings in $(A \cap B) \Delta (A \cap C)$ enthalten ist. Ein x , das in A und B , aber nicht in C enthalten ist, ist daher sowohl in $A \cap (B \Delta C)$ als auch in $(A \cap B) \Delta (A \cap C)$ enthalten. Die anderen sieben Fälle können ähnlich behandelt werden, und in jedem Fall ist x entweder (wie im eben behandelten Fall) Element von beiden Teilen oder kein Element. Da mit den acht Fällen alle möglichen Situationen abgedeckt sind, müssen $A \cap (B \Delta C)$ und $(A \cap B) \Delta (A \cap C)$ nach Definition 2.5 gleich sein. In einer **Elementtabelle** können alle Fälle übersichtlich dargestellt und so der Nachweise von Mengenidentitäten erbracht werden. Die oben angeführte Überlegung entspricht übrigens der zweiten Zeile.

A	B	C	$B\Delta C$	$A \cap (B\Delta C)$	$A \cap B$	$A \cap C$	$(A \cap B)\Delta(A \cap C)$
$\in$	$\in$	$\in$	$\notin$	$\notin$	$\in$	$\in$	$\notin$
$\in$	$\in$	$\notin$	$\in$	$\in$	$\in$	$\notin$	$\in$
$\in$	$\notin$	$\in$	$\in$	$\in$	$\notin$	$\in$	$\in$
$\in$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$
$\notin$	$\in$	$\in$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$
$\notin$	$\in$	$\notin$	$\in$	$\notin$	$\notin$	$\notin$	$\notin$
$\notin$	$\notin$	$\in$	$\in$	$\notin$	$\notin$	$\notin$	$\notin$
$\notin$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$	$\notin$

Man kann auf einem Blick erkennen, ob Mengengleichheit besteht, die fünfte und die achte Spalte müssen gleich sein. Entsteht in wenigstens einem Fall ein unterschiedliches Bild (d.h. x ist in einem Teil enthalten, im anderen aber nicht), so sind die betrachteten Mengenausdrücke nicht gleich. Es gibt dann Mengen $A, B, C, \dots$, für die die Identität nicht gilt. (Solche können auch leicht konstruiert werden.)

Durch eine einfache Modifikation können auch Enthaltenseinsrelationen ($\subseteq$) von Mengenausdrücken überprüft werden. (Ist links ein $\in$ -Zeichen, so muß auch rechts eines sein.)

2.5 Disjunktive und konjunktive Normalform

Im Abschnitt 2.4 wurden, um zwei Mengenausdrücke auf Gleichheit zu untersuchen, eine Fallunterscheidung getroffen, je nachdem ob ein x in einer der involvierenden Mengen enthalten ist oder nicht. Sind beispielsweise drei Mengen A, B, C involviert, so werden $2^3 = 8$ Fälle unterschieden. (Sind n Mengen $A_1, A_2, \dots, A_n$ beteiligt, so benötigt man 2^n Unterfälle.) Diese Fälle wurden in der Elementtabelle systematisch erfaßt. Beispielsweise behandelt die zweite Zeile den Fall, wo x in A und B , aber nicht in C enthalten ist, d.h. man behandelt die Menge $A \cap B \cap C'$. Es stellt sich heraus, daß diese Menge sowohl in $A \cap (B\Delta C)$ als auch in $(A \cap B)\Delta(A \cap C)$ enthalten ist. Die Menge $A \cap B' \cap C'$, die der vierten Zeile entspricht, ist hingegen ganz im Komplement dieser Mengen enthalten.

Nach Analyse aller $2^3 = 8$ Fälle erhält man, daß $M = A \cap (B\Delta C) = (A \cap B)\Delta(A \cap C)$ auch durch

$$M = (A \cap B \cap C') \cup (A \cap B' \cap C)$$

gegeben ist. Einen Ausdruck dieser Form nennt man *disjunktive Normalform*.

Für das Komplement M' ergibt sich übrigens

$$\begin{aligned} M' &= (A \cap B \cap C) \cup (A \cap B' \cap C') \\ &\quad \cup (A' \cap B \cap C) \cup (A' \cap B \cap C') \\ &\quad \cup (A' \cap B' \cap C) \cup (A' \cap B' \cap C'). \end{aligned}$$

Wendet man auf diese Darstellung die DeMorganschen Regeln an, so erhält man eine

alternative Darstellung von M :

$$\begin{aligned} M &= (A' \cup B' \cup C') \cap (A' \cup B \cup C) \\ &\quad \cap (A \cup B' \cup C') \cap (A \cup B' \cup C) \\ &\quad \cap (A \cup B \cup C') \cap (A \cup B \cup C) \end{aligned}$$

Eine Darstellung dieser Form heißt auch *konjunktive Normalform*.

Im folgenden werden wir die Notation

$$A^\delta = \begin{cases} A & \text{für } \delta = 1, \\ A' & \text{für } \delta = 0 \end{cases}$$

verwenden.

Definition 2.18 Seien $A_1, A_2, \dots, A_n$ Teilmengen einer Menge E . Für jede $m \times n$ -Matrix $J = (\delta_{kl})_{1 \leq k \leq m, 1 \leq l \leq n}$ ($0 \leq m \leq 2^n$) mit Elementen $\delta_{kl} \in \{0, 1\}$ und paarweise verschiedenen Zeilen bezeichnet

$$\bigcup_{k=1}^m \left(\bigcap_{l=1}^n A_l^{\delta_{kl}} \right)$$

eine **disjunktive Normalform**. Entsprechend heißt

$$\bigcap_{k=1}^m \left(\bigcup_{l=1}^n A_l^{\delta_{kl}} \right)$$

konjunktive Normalform.

Demnach gibt es genau 2^{2^n} disjunktive resp. konjunktive Normalformen von n Mengen.

Satz 2.19 Seien $A_1, A_2, \dots, A_n$ Teilmengen einer Menge E , und sei M ein Mengenausdruck in $A_1, A_2, \dots, A_n$, der nur die Operationen Durchschnitt ($\cap$), Vereinigung ($\cup$), Komplement ($'$), Mengendifferenz ($\setminus$) und symmetrische Differenz (Δ) verwendet, dann läßt sich M (bis auf die Reihenfolge der Terme) eindeutig als disjunktive resp. konjunktive Normalform darstellen.

Die entsprechenden Darstellungen können aus einer Elementtabelle für M sofort abgelesen werden. Ersetzt man die Elementsymbole $\in$ durch 1 und die Nicht-Elementsymbole durch 0, so setzt sich die Matrix J für die disjunktive Normalform aus jenen Zeilen der ersten n Spalten zusammen, wo in der resultierenden Spalte 1 steht. Vertauscht man alle 0 und 1, so erhält man auf dieselbe Art und Weise die Matrix J für die konjunktive Normalform. Beinhaltet die disjunktive Normalform von M genau k Terme, so setzt sich die konjunktive Normalform aus $2^n - k$ Termen zusammen.

Es gibt eine weitere Methode, die disjunktive resp. die konjunktive Normalform eines Mengenausdrucks M zu erhalten, ohne die recht aufwendige Elementtabelle aufstellen zu müssen. Der folgende Algorithmus liefert die disjunktive Normalform:

1. Eliminieren der Symbole Mengendifferenz ($\setminus$) und symmetrische Differenz (Δ), d.h. $A \setminus B$ wird durch $A \cap B'$ und $A \Delta B$ wird durch $(A \cap B') \cup (A' \cap B)$ ersetzt.
2. Eliminieren von Komplementbildungen außerhalb von Klammern, d.h. $(A \cap B)'$ wird durch $A' \cup B'$ und $(A \cup B)'$ wird durch $A' \cap B'$ ersetzt.
3. Eliminieren von Durchschnitten ($\cap$) bei Klammern mit Hilfe des Distributivgesetzes, d.h. $A \cap (B \cup C)$ wird durch $(A \cap B) \cup (A \cap C)$. (Durchschnitte werden in Klammern *hineingezogen*.) Die bisherigen Umformungen ergeben daher für M eine Darstellung der Form

$$M = \bigcup_i M_i,$$

wobei M_i gewisse Durchschnitte von $A_1, A_2, \dots, A_n$ bzw. deren Komplementen $A'_1, A'_2, \dots, A'_n$ sind. Da $A \cap A = A$ und $A \cap A' = \emptyset$ gilt, können gewisse M_i weiter vereinfacht werden bzw. $M_i = \emptyset$ und es braucht nicht weiter berücksichtigt werden.

4. Vervollständigung der M_i . Kommt in einem M_i weder A_j noch A'_j vor so ersetzt man M_i durch

$$M_i = M_i \cap (A_j \cup A'_j) = (M_i \cap A_j) \cup (M_i \cap A'_j).$$

Auf diese Art und Weise erhält man schließlich

$$M = \bigcup_j \tilde{M}_j,$$

wobei jedes $\tilde{M}_j$ die Form

$$\tilde{M}_j = A_1^{\delta_1} \cap A_2^{\delta_2} \cap \dots \cap A_n^{\delta_n}$$

mit $\delta_1, \delta_2, \dots, \delta_n \in \{0, 1\}$ hat. Sind gewisse $\tilde{M}_j$ gleich so brauchen sie wegen der Beziehung $A \cup A = A$ nur einmal aufgenommen werden. Das Resultat ist nun die *disjunktive Normalform* von M .

Beispiel 2.20 Die disjunktive Normalform von $A \cap (B \cap C)'$ erhält man aus dem eben angeführten Algorithmus zu

$$\begin{aligned} A \cap (B \cap C)' &= A \cap (B' \cup C') \\ &= (A \cap B') \cup (A \cap C') \\ &= (A \cap B' \cap C) \cup (A \cap B' \cap C') \\ &\cup (A \cap B \cap C') \cup (A \cap B' \cap C') \\ &= (A \cap B' \cap C) \cup (A \cap B' \cap C') \cup (A \cap B \cap C'). \end{aligned}$$

Auf ähnliche Art und Weise erhält man die konjunktive Normalform:

$$A \cap (B \cap C)' = A \cap (B' \cup C')$$

$$\begin{aligned}
&= (A \cup B \cup C) \cap (A \cup B \cup C') \\
&\cap (A \cup B' \cup C) \cap (A \cup B' \cup C') \\
&\cap (A \cup B' \cup C') \cap (A' \cup B' \cup C') \\
&= (A \cup B \cup C) \cap (A \cup B \cup C') \cap (A \cup B' \cup C) \\
&\cap (A \cup B' \cup C') \cap (A' \cup B' \cup C').
\end{aligned}$$

Da zwei disjunktive (resp. konjunktive) Normalformen nur dann dieselbe Menge representieren, wenn in der Normalform dieselben Terme vorkommen, kann man mit Hilfe von diesen Normalformen auch Mengenidentitäten überprüfen, z.B. haben $A \cap (B \Delta C)$ und $(A \cap B) \Delta (A \cap C)$ dieselbe disjunktive Normalform

$$M = (A \cap B \cap C') \cup (A \cap B' \cap C).$$

2.6 Potenzmenge

Definition 2.21 Die **Potenzmenge** $\mathbf{P}(A)$ einer Menge A ist die Menge aller Teilmengen von A , d.h.

$$\mathbf{P}(A) = \{C \mid C \subseteq A\}.$$

Beispiel 2.22

$$\begin{aligned}
\mathbf{P}(\{1, 2\}) &= \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}, \\
\mathbf{P}(\emptyset) &= \{\emptyset\}
\end{aligned}$$

Manchmal bezeichnet man die Potenzmenge einer Menge A auch als 2^A . Ein Grund dafür ist der folgende Satz. ($|A|$ bezeichnet die Anzahl der Elemente von einer endlichen Menge A .)

Satz 2.23 Für eine endliche Menge A gilt

$$|\mathbf{P}(A)| = 2^{|A|},$$

d.h. eine Menge mit n Elementen hat genau 2^n Teilmengen.

Definition 2.24 Seien $0 \leq k \leq n$ ganze Zahlen. Der **Binomialkoeffizient**¹³ $\binom{n}{k}$ (sprich: n über k) ist definiert durch

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

($0! = 1$ und $n! = 1 \cdot 2 \cdots (n-1) \cdot n$ für $n > 0$.)

¹³binomial coefficient

Für ganze Zahlen n, k mit $n \geq 0$ und $k < 0$ oder $k > n$ setzt man auch $\binom{n}{k} = 0$. Mit Hilfe dieser Zusatzdefinition gilt der folgende Satz uneingeschränkt und ist Grundlage des **Pascalschen Dreiecks**.

Satz 2.25 *Die Binomialkoeffizienten erfüllen die Rekursion*

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

Satz 2.26 *Eine endliche Menge mit n Elementen hat genau $\binom{n}{k}$ Teilmengen mit k Elementen.*

Man beachte, daß daraus folgt, daß $\binom{n}{k}$ immer eine natürliche Zahl ist, was aus der Definition nicht unmittelbar ersichtlich ist.

Beispiel 2.27 In einer Liga von n Mannschaften müssen genau $\binom{n}{2} = \frac{1}{2}n(n-1)$ Spiele ausgetragen werden, damit jeder gegen jeden gespielt hat.

Korollar 2.28 *Für $n \geq 0$ gilt*

$$\sum_{k=0}^n \binom{n}{k} = 2^n.$$

Eine Erweiterung dieses Korollars ist der sogenannte **Binomische Lehrsatz**

Satz 2.29 *Für $n \geq 0$ und beliebige $x, y \in \mathbb{C}$ gilt*

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

Korollar 2.28 ergibt sich aus der Spezialfall $x = y = 1$.

2.7 Stirlingsche Approximationsformel

Im Abschnitt 2.6 wurde für die Definition des Binomalkoeffizienten bereits die Bezeichnung $n! = 1 \cdot 2 \cdot \dots \cdot (n-1) \cdot n$ (sprich: n Fakultät¹⁴ oder n faktorielle) verwendet. Aufgrund seiner Definition läßt sich $n!$ nur rekursiv berechnen. Außerdem wird $n!$ mit wachsendem n außerordentlich groß. In vielen Anwendungen, etwa in der Statistik, ist es nicht notwendig, den genauen Wert von $n!$ zu kennen. Ein guter Näherungswert ist ausreichend. Im folgenden soll eine asymptotische Formel für $n!$ angegeben werden.

¹⁴ n factorial

Definition 2.30 Zwei Folgen $(a_n)_{n \geq 0}$, $(b_n)_{n \geq 0}$ heißen *asymptotisch gleich*, i.Z. $a_n \sim b_n$ ($n \rightarrow \infty$), wenn

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

Satz 2.31 (Stirlingsche Approximationsformel) $n!$ läßt sich folgendermaßen asymptotisch approximieren:

$$n! \sim n^n e^{-n} \sqrt{2\pi n} \quad (n \rightarrow \infty).$$

Beispiel 2.32 Aus der obigen Approximationsformel ergibt sich

$$\begin{aligned} \binom{2k}{k} &\sim \frac{(2k)^{2k} e^{-2k} \sqrt{4k\pi}}{k^{2k} e^{-2k} 2\pi k} \\ &= 4^k \frac{1}{\sqrt{\pi k}}. \end{aligned}$$

Da es genau $2^{2k} = 4^k$ Teilmengen einer $2k$ -elementigen Menge gibt, ist die *Wahrscheinlichkeit*, bei einer zufälligen Auswahl einer Teilmenge genau k Elemente zu erhalten, ungefähr $1/\sqrt{\pi k}$.

Kapitel 3

Relationen

3.1 Grundlegende Begriffe

Um den mengentheoretischen Begriff einer **Relation** zu motivieren, sollen zunächst einige Beispiele angegeben werden.

Beispiel 3.1 Man betrachte eine Gruppe von Personen. Sind a und b zwei Personen dieser Gruppe, so können folgende Situationen eintreten:

- a und b kennen einander,
- a kennt b , aber b kennt nicht a ,
- b kennt a , aber a kennt nicht b ,
- a und b kennen einander nicht.

Beispiel 3.2 Gewisse Städte können durch Direktflüge voneinander erreicht werden, andere nicht.

Beispiel 3.3 Eine Schachtel (alter) Schrauben kann so sortiert werde, daß jeweils Schrauben gleicher Länge in ein eigenes Fach kommen.

Abstrahiert man von den angegebenen Beispielen, so steht man vor folgender Situation. Zwei Elemente a, b einer Menge stehen miteinander in einer gewissen *Relation* (wobei die Reihenfolge eine Rolle spielen kann) oder eben nicht. Um diesen *Relationsbegriff* mathematisch streng zu fassen, benötigt man noch den Begriff des kartesischen Produkts.

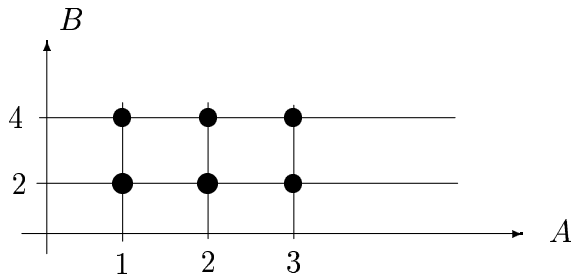
Definition 3.4 Das **kartesische Produkt**¹ $A \times B$ zweier Mengen A, B ist die Menge aller geordneten Paare (a, b) mit $a \in A$ und $b \in B$, d.h.

$$A \times B = \{(a, b) \mid a \in A \wedge b \in B\}$$

¹(cartesian) product

In derselben Weise definiert man auch das Produkt $A_1 \times A_2 \times \cdots \times A_n$ von endlich vielen Mengen $A_1, A_2, \dots, A_n$ als die Menge aller n -tupel $(a_1, a_2, \dots, a_n)$ mit $a_j \in A_j$ ($1 \leq j \leq n$). Sind alle Mengen A_j gleich einer Menge A , so schreibt man statt $A \times A \times \cdots \times A$ auch A^n .

Beispiel 3.5 Sei $A = \{1, 2, 3\}$ und $B = \{2, 4\}$. Dann ist $A \times B = \{(1, 2), (1, 4), (2, 2), (2, 4), (3, 2), (3, 4)\}$. Dies lässt sich auch folgendermaßen *kartesisch* darstellen:



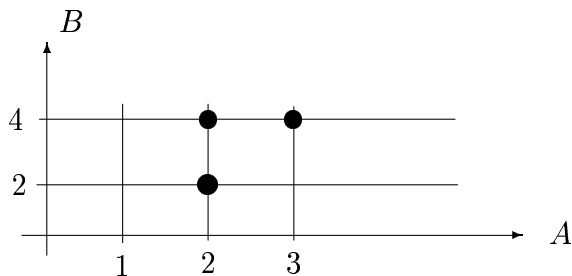
Definition 3.6 Eine **Relation**² R zwischen zwei Mengen A und B ist eine Teilmenge des kartesischen Produkts $A \times B$.

Ist $A = B$ so spricht man von einer **binären Relation**.³

Anstelle von $(a, b) \in R$ schreibt man auch aRb , anstelle von $(a, b) \notin R$ auch $a \not R b$.

Die drei einleitenden Beispiele sind übrigens alle binäre Relationen.

Beispiel 3.7 Sei $A = \{1, 2, 3\}$ und $B = \{2, 4\}$. Dann ist $R = \{(2, 2), (2, 4), (3, 4)\}$ eine Relation zwischen A und B , d.h. $2R2$, $2R4$ und $3R4$, aber $1 \not R 2$, $1 \not R 4$ und $3 \not R 2$. Dies kann natürlich auch graphisch ausgedrückt werden:

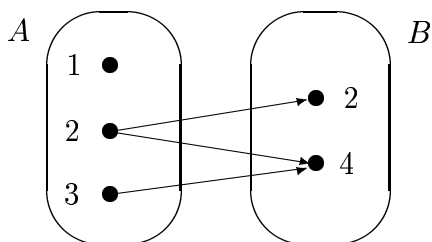


Es gibt verschiedene Möglichkeiten, eine Relation $R \subseteq A \times B$ bildlich darzustellen:

²relation

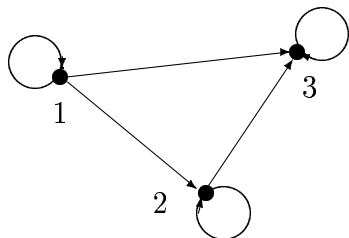
³binary relation

1. **Kartesische Darstellung:** In der kartesischen Darstellung von $A \times B$ (siehe Beispiel 3.5) werden nur jene Elemente (= Punkte) von $A \times B$ *markiert*, die der Relation $R \subseteq A \times B$ angehören (siehe Beispiel 3.7).
2. **Pfeildiagramm:** Die Mengen A, B werden durch Venndiagramme dargestellt und Paare $(a, b) \in R$ durch einen Pfeil verbunden:



3. **Graph einer binären Relation:** Da bei einer binären Relation $A = B$ gilt (d.h. $R \subseteq A \times A = A^2$), reicht es, (im Gegensatz zum Pfeildiagramm) die Menge A nur einmal zu repräsentieren. Paare $(a, b) \in A \times A$, die in Relation stehen, werden nun ähnlich wie beim Pfeildiagramm miteinander verbunden. Der **Graph** $G(R)$ besteht daher aus einer Menge von *Punkten (Knoten)* entsprechend den Elementen aus A und einer Mengen von *gerichteten Kanten*, die genau jene Punkte miteinander verbinden, die miteinander in Relation stehen.

Ist beispielsweise $A = \{1, 2, 3\}$ und $R = \{(1, 1), (1, 2), (1, 3), (2, 2), (2, 3), (3, 3)\}$, so entsteht folgendes Bild:



Man beachte, daß in einem Graphen einer Relation auch sogenannte *Schlingen* auftreten, das sind Kanten, die von $a \in A$ wieder auf a zeigen (a steht mit sich selbst in Relation: aRa).

3.2 Äquivalenzrelation

Definition 3.8 Eine binäre Relation R auf einer Menge A heißt **Äquivalenzrelation**⁴, wenn folgende drei Eigenschaften erfüllt sind:

1. $\forall a \in A : aRa$ (**Reflexivität**),

⁴equivalence relation

2. $\forall a, b \in A : aRb \rightarrow bRa$ (**Symmetrie**),
 3. $\forall a, b, c \in A : (aRb \wedge bRc) \rightarrow aRc$ (**Transitivität**).

Eine Relation mit der Eigenschaft 1. heißt **reflexiv**⁵, eine mit der Eigenschaft 2. **symmetrisch**⁶ und eine mit der Eigenschaft 3. **transitiv**⁷. Eine Äquivalenzrelation ist also reflexiv, symmetrisch und transitiv.

Beispiel 3.9 Gleichheitsrelation: A sei eine beliebige Menge und $R = "="$, d.h. jedes Element $a \in A$ steht nur mit sich selbst in Relation.

Beispiel 3.10 Allrelation: A sei eine beliebige Menge und $R = A^2$, d.h. jedes Element $a \in A$ steht mit allen anderen Elementen $b \in A$ in Relation.

Beispiel 3.11 Sei $A = \mathbf{Z}$ und $aRb : \Longleftrightarrow a \equiv b \pmod{2}$.

In der *kartesischen Darstellung* einer Relation R äußert sich die Reflexivität dadurch, daß die 1. Hauptdiagonale (1. Mediane) in der Relation enthalten ist. Ist die Relation symmetrisch, so ist die kartesische Darstellung symmetrisch zur 1. Hauptdiagonale, d.h. R geht durch Spiegelung an der 1. Hauptdiagonale in sich über. Die Transitivität hat in der kartesischen Darstellung keine offensichtliche Entsprechung.

Stellt man eine Relation R als Graph $G(R)$ dar, so entspricht einer reflexiven Relation ein Graph, bei dem jeder Punkt (Knoten) mit einer Schlinge ausgestattet ist. Bei einer symmetrischen Relation treten die Kanten (mit Ausnahme der Schlingen) gepaart auf. Zu einer Kante von a nach b gibt es immer auch eine Gegenkante von b nach a . Ebenso läßt sich die Transitivität sofort übersetzen.

Im Beispiel 3.11 fällt auf, daß durch die Äquivalenzrelation die ganzen Zahlen $\mathbf{Z}$ in zwei Teilmengen zerlegt wird, in die geraden Zahlen und in die ungeraden Zahlen. Alle geraden Zahlen stehen miteinander in Relation, und entsprechend alle ungeraden Zahlen. Sie sind jeweils kongruent modulo 2. Aber eine gerade Zahl steht mit keiner ungeraden Zahl in Relation.

Ein entsprechender Sachverhalt gilt ganz allgemein. Äquivalenzrelationen zerlegen die Grundmenge in sogenannte Äquivalenzklassen. Dies soll nun präzisiert werden.

Definition 3.12 Ein System von nichtleeren Teilmengen A_i ($i \in I$) einer Menge A heißt **Partition**⁸ oder **Zerlegung** von A , wenn die A_i ($i \in I$) paarweise disjunkt sind, d.h.

$$A_i \cap A_j = \emptyset \quad \text{für } i \neq j,$$

⁵reflexive

⁶symmetric

⁷transitive

⁸partition

und A die Vereinigung

$$A = \bigcup_{i \in I} A_i$$

ist.

Definition 3.13 Sei R eine Äquivalenzrelation auf A . Für $a \in A$ heißt die Menge

$$K(a) = \{b \in A \mid aRb\}$$

die von a erzeugte **Äquivalenzklasse**.⁹

Man beachte, daß wegen der Reflexivität immer $a \in K(a)$ gilt. Weiters gilt die folgende Eigenschaft.

Lemma 3.14 Sei R eine Äquivalenzrelation auf A . Dann gilt

$$aRb \iff K(a) = K(b).$$

Daraus ergibt sich leicht der folgende Zusammenhang zwischen Äquivalenzrelationen und Partitionen.

Satz 3.15 Sei R eine Äquivalenzrelation auf A . Dann bilden die (verschiedenen) Äquivalenzklassen der Elemente von A eine Partition von A .

Sei umgekehrt A_i ($i \in I$) eine Partition von A und bezeichne $C(a)$ ($a \in A$) jene Teilmenge A_j der Partition mit $a \in A_j$. Definiert man aRb genau für jene $a, b \in A$, für die $C(a) = C(b)$ gilt, so ist R eine Äquivalenzrelation.

3.3 Ordnungsrelation

Definition 3.16 Eine binäre Relation R auf einer Menge A heißt **Halbordnung** oder **partielle Ordnung**¹⁰, wenn folgende drei Eigenschaften erfüllt sind:

1. $\forall a \in A : aRa$ (**Reflexivität**),
2. $\forall a, b \in A : (aRb \wedge bRa) \rightarrow a = b$ (**Antisymmetrie** oder **Identität**),
3. $\forall a, b, c \in A : (aRb \wedge bRc) \rightarrow aRc$ (**Transitivität**).

Eine Relation mit der Eigenschaft 2. heißt antisymmetrisch¹¹. Eine Halbordnung ist daher eine reflexive, antisymmetrische und transitive Relation

⁹equivalence class

¹⁰partial order

¹¹antisymmetric

Definition 3.17 Eine Halbordnung R auf einer Menge A heißt **Totalordnung**¹², wenn für je zwei Elemente $a, b \in A$ entweder aRb oder bRa gilt, d.h. je zwei Elemente sind vergleichbar.

Beispiel 3.18 $A = \mathbf{R}$ mit $R = “\leq”$ bildet eine Totalordnung (Ordnung der reellen Zahlen).

Beispiel 3.19 $A = \mathbf{N}$ mit $mRn :\iff “m \text{ teilt } n”$ ist eine Halbordnung, aber keine Totalordnung.

$A = \mathbf{Z}$ mit $mRn :\iff “m \text{ teilt } n”$ ist keine Halbordnung, da R auf $\mathbf{Z}$ nicht mehr antisymmetrisch ist (z.B. -2 teilt 2 und umgekehrt, aber $-2 \neq 2$).

Beispiel 3.20 $A = \mathbf{P}(M)$ (Potenzmenge einer Menge M) mit $BRC \iff B \subseteq C$ bildet eine Halbordnung, aber für $|M| > 1$ keine Totalordnung.

Insbesondere sind damit alle Relationen $R \subseteq A \times B$ zwischen zwei (festen) Mengen A, B in natürlicher Weise geordnet.

Wie jede binäre Relation kann man natürlich auch Halbordnungen durch einen Graphen darstellen. Viele der darzustellenden Kanten sind allerdings redundant, sie lassen sich aus den definierenden Eigenschaften leicht wieder rekonstruieren. Führt man die folgenden drei Schritte durch, so erhält man aus dem Graphen $G(R)$ einer Halbordnung R das **Hassediagramm**¹³ von R :

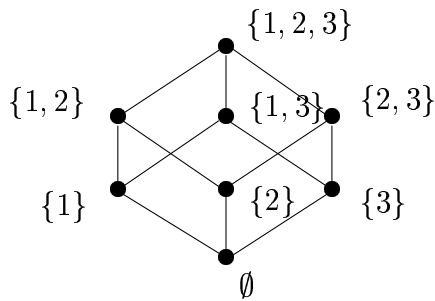
- Weglassen aller Schlingen.
- Weglassen aller Kanten, die sich aufgrund der Transitivitätsbedingung rekonstruieren lassen, d.h. ist aRb , aber gibt es kein c mit aRc und cRb , so bleibt die Kante von a nach b erhalten, allen anderen Kanten werden *gestrichen*. Mit anderen Worten: nur die *unmittelbaren Nachfolger* von a werden von a mit einer Kante verbunden.
- Weglassen aller Orientierungen. Wegen der Antisymmetrie kann für $a \neq b$ entweder aRb oder bRa gelten aber nie beides zugleich. Zur Übersicht zeichnet man bei aRb ($a \neq b$) b oberhalb von a und kann die Orientierung der Kante weglassen.

Oft wird eine Halbordnung auf einer endlichen Menge nur durch Angabe des Hassediagramms definiert.

Beispiel 3.21 Sei $A = \mathbf{P}(M) = \{\emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\}\}$ die Potenzmenge von $M = \{1, 2, 3\}$ mit der Inklusion ($\subseteq$) als Relation (siehe Beispiel 3.20). Dann hat das Hassediagramm dieser Halbordnung die folgende Gestalt:

¹²linear order

¹³Hasse diagramm



3.4 Transitive Hülle einer Relation

Definition 3.22 Sei $R \subseteq A \times B$ eine Relation. Dann ist die **inverse Relation** $R^{-1} \subseteq B \times A$ durch

$$R^{-1} = \{(b, a) \in B \times A \mid aRb\},$$

d.h. $bR^{-1}a \iff aRb$.

Definition 3.23 Seien $R \subseteq A \times B$ und $S \subseteq B \times C$ zwei Relationen. Dann ist die **Komposition**¹⁴ $R \circ S$ eine Relation zwischen A und C (d.h. $R \circ S \subseteq A \times C$), so daß $a(R \circ S)c$ genau dann gilt, wenn es ein $b \in B$ mit aRb und bSc gibt.

Ist $R \subseteq A \times A = A^2$ eine binäre Relation, so bezeichnet man $R \circ R$ durch R^2 , $R \circ R \circ R$ durch R^3 usw. Aus Vollständigkeitsgründen definiert man auch $R^1 := R$ und $R^0 := \Delta = \{(a, a) \mid a \in A\}$.

Lemma 3.24 Sei $R \subseteq A \times A$ eine binäre Relation. Dann gelten die folgenden Charakterisierungen:

- R ist genau dann reflexiv, wenn $R^0 \subseteq R$.
- R ist genau dann symmetrisch, wenn $R^{-1} = R$.
- R ist genau dann transitiv, wenn $R^2 \subseteq R$.

Definition 3.25 Seien $A = \{a_1, a_2, \dots, a_m\}$ und $B = \{b_1, b_2, \dots, b_n\}$ zwei endliche Mengen und $R \subseteq A \times B$ eine Relation zwischen A und B . Die **Matrix** M_R **der Relation** R ist eine $m \times n$ -Matrix $(m_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ mit

$$m_{ij} = \begin{cases} 1 & \text{für } a_i R b_j, \\ 0 & \text{für } a_i \not R b_j, \end{cases}$$

d.h. das Element m_{ij} im Schnittpunkt der i -ten Zeile und der j -ten Spalte ist genau dann 1, wenn a_i mit b_j in Relation stehen.

$\boxplus$	0	1
0	0	1
1	1	1

Im folgenden wird untersucht, wie man mit Hilfe der Matrix M_R einer Relation R rechnen kann. Entscheidend dafür wird der Gebrauch der binäre Operation $\boxplus$ sein, die durch die oben angegebene Operationstafel definiert ist. Die Operation $\boxplus$ unterscheidet sich von der gewöhnlichen Addition $+$, dadurch, daß Resultate, die positiv sind, durch 1 ersetzt werden.

Definition 3.26 Seien $M = (m_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ und $M' = (m'_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ zwei $m \times n$ -Matrizen mit Eintragungen $m_{ij}, m'_{ij} \in \{0, 1\}$. Dann bezeichnet $M \boxplus M' = (m_{ij} \boxplus m'_{ij})$ und $M \boxdot M' = (m_{ij} \cdot m'_{ij})$.

Ist weiters $P = (p_{jl})_{1 \leq j \leq n, 1 \leq l \leq k}$ eine $n \times k$ -Matrix mit Eintragungen $p_{jl} \in \{0, 1\}$. Dann bezeichne $M \circ P = (q_{il})_{1 \leq i \leq m, 1 \leq l \leq k}$ jene Matrix, die durch

$$q_{il} = m_{i1}p_{1l} \boxplus m_{i2}p_{2l} \boxplus \cdots \boxplus m_{in}p_{nl}$$

gegeben ist.

Satz 3.27 Seien $R \subseteq A \times B$, $R' \subseteq A \times B$, $S \subseteq B \times C$ Relation. Dann gilt für die entsprechenden Matrizen:

- $M_{R^{-1}} = (M_R)^T$.
- $M_{R \cup R'} = M_R \boxplus M_{R'}$.
- $M_{R \cap R'} = M_R \boxdot M_{R'}$.
- $M_{R \circ S} = M_R \circ M_S$.

Mit Hilfe dieser Eigenschaften läßt sich auch Lemma 3.24 für Matrizen umformulieren.

Satz 3.28 Sei $R \subseteq A \times A$ eine binäre Relation mit Matrix $M_R = (m_{ij})_{1 \leq i, j \leq |A|}$. Dann gelten die folgenden Charakterisierungen:

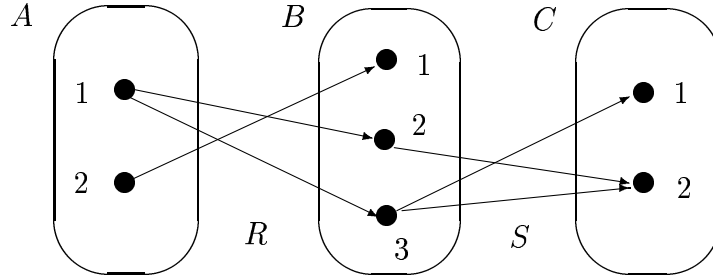
- R ist genau dann reflexiv, wenn $m_{ii} = 1$ für $1 \leq i \leq |A|$.
- R ist genau dann symmetrisch, wenn M_R symmetrisch ist.
- R ist genau dann transitiv, wenn alle Elemente von $M_R \circ M_R$ kleiner oder gleich den entsprechenden Elementen von M_R sind.

¹⁴composition

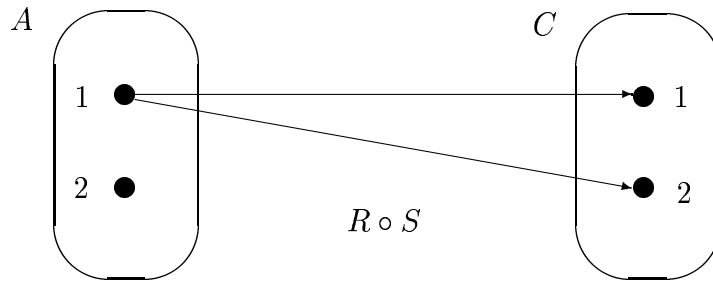
Beispiel 3.29 Sei $A = \{1, 2\}$, $B = \{1, 2, 3\}$ und $C = \{1, 2\}$. Weiters seien die Relationen $R \subseteq A \times B$ und $S \subseteq B \times C$ durch die Matrizen

$$M_R = \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 0 \end{pmatrix}, \quad M_S = \begin{pmatrix} 0 & 0 \\ 0 & 1 \\ 1 & 1 \end{pmatrix}$$

gegeben. Aus den entsprechenden *Pfeildiagrammen*



ergibt sich daraus unschwer das Pfeildiagramm der Komposition $R \circ S$:



Die dazugehörige Matrix ist

$$M_{R \circ S} = \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix},$$

die sich auch mit Hilfe des letzten Punktes von Satz 3.27 berechnen läßt:

$$\begin{aligned} \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 0 \end{pmatrix} \circ \begin{pmatrix} 0 & 0 \\ 0 & 1 \\ 1 & 1 \end{pmatrix} &= \begin{pmatrix} 0 \boxplus 0 \boxplus 1 & 0 \boxplus 1 \boxplus 1 \\ 0 \boxplus 0 \boxplus 0 & 0 \boxplus 0 \boxplus 0 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix}. \end{aligned}$$

Definition 3.30 Die **transitive Hülle**¹⁵ oder **Erreichbarkeitsrelation** R^+ einer binären Relation R ist durch

$$R^+ = \bigcup_{k=1}^{\infty} R^k$$

¹⁵transitive closure

gegeben, d.h. aR^+b genau dann, wenn es ein $k \geq 1$ und $c_1, c_2, \dots, c_{k-1} \in A$ gibt, die $aRc_1, c_1Rc_2, \dots, c_{k-2}Rc_{k-1}$ und $c_{k-1}Rb$ erfüllen.

Entsprechend wird

$$R^* = \bigcup_{k=0}^{\infty} R^k = R^+ \cup R^0$$

als **reflexiv-transitive Hülle** von R bezeichnet.

Es ist unmittelbar klar, daß R^+ die kleinste transitive Relation ist, die R enthält. Entsprechend ist R^* die kleinste reflexive und transitive Relation, die R enthält.

Satz 3.31 *Es sei A eine endliche Menge und $R \subseteq A \times A$ eine binäre Relation auf A . Dann gilt*

$$R^+ = \bigcup_{k=1}^{|A|-1} R^k.$$

Satz 3.31 eröffnet eine Möglichkeit, die transitive Hülle R^+ einer Relation R mit Hilfe der Matrix M_R effektiv zu berechnen. Offensichtlich gilt

$$M_{R^+} = M_R \boxplus M_R^2 \boxplus \dots \boxplus M_R^{m-1},$$

wobei $m = |A|$ bezeichnet. Man überlegt sich sofort, daß man dafür im allgemeinen $O(m^4)$ Rechenschritte benötigt.

3.5 Warshall-Algorithmus

Die Berechnung der transitiven Hülle R^+ einer Relation R mit Hilfe von Matrizen ist mit ihren $O(m^4)$ Rechenschritten sehr aufwendig. Der folgende Algorithmus von Warshall erreicht eine signifikante Verbesserung, er benötigt (wie man sofort erkennt) nur $O(m^3)$ Schritte.

Warshall-Algorithmus zur Bestimmung von R^+

1. **Input:** Es sei $M_R = (m_{ij})_{1 \leq i, j \leq m}$ die Matrix einer binären Relation R auf $A = \{a_1, a_2, \dots, a_m\}$.

2. **Algorithmus:**

```

for  $i := 1$  to  $m$  do
  for  $j := 1$  to  $m$  do
    if  $m_{ji} = 1$  then
      for  $k := 1$  to  $m$  do
        if  $m_{ik} = 1$  then

```

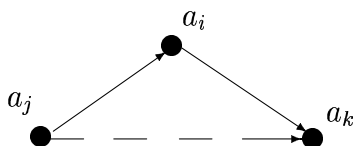
```

                 $m_{jk} := 1;$ 
            end;
        end;
    end;
end;

```

3. **Output:** Die resultierende Matrix $(m_{ij})_{1 \leq i, j \leq m}$ ist nun die Matrix M_{R^+} der transitiven Hülle von R .

In diesem Algorithmus werden also alle Elemente $a_1, a_2, \dots, a_m$ durchlaufen. Im ersten Schritt betrachtet man a_1 und sucht nun a_j und a_k , für die $a_j R a_1$ und $a_1 R a_k$ gilt. Hat man solche gefunden, so erweitert man die Relation um das Paar (a_j, a_k) , das ja sicherlich in R^+ enthalten sein muß. Dies geschieht, indem man $m_{jk} := 1$ setzt.



Wichtig ist, daß sich die Matrix $(m_{ij})_{1 \leq i, j \leq m}$ in diesem Algorithmus *dynamisch* ändert. Wurde im ersten Durchlauf der ersten Schleife ($i = 1$) ein Eintrag m_{2k} auf 1 gesetzt, so ist dieser im zweiten Durchlauf der ersten Schleife ($i = 2$) schon verfügbar, und es werden möglicherweise im zweiten Durchlauf mehr Einträge auf 1 gesetzt als wenn nur die ursprüngliche Matrix M_R zur Verfügung stünde.

Ist also $a, b \in A$ mit $a R^+ b$ und $a R b$, so gibt es eine natürliche Zahl $k \leq m - 2$ und Elemente $c_1, c_2, \dots, c_k \in A$ mit $a R c_1, c_1 R c_2, \dots, c_{k-1} R c_k$ und $c_k R b$, so daß $a, c_1, c_2, \dots, c_k, b$ voneinander verschieden sind. Ist nun etwa $c_j = a_1$ für ein c_j in dieser Kette, so nimmt der Warshall-Algorithmus im ersten Durchlauf der ersten Schleife das Paar (c_{j-1}, c_{j+1}) in die Relation auf, d.h. der *Weg* von a nach b ist in der vergrößerten Relation kürzer geworden, $c_j = a_1$ wird sozusagen überbrückt. Bei den weiteren Durchläufen der ersten Schleife wird der *Weg* von a nach b jeweils dann verkürzt, wenn a_i einem der verbleibenden c_j gleich ist. Da alle c_j voneinander verschieden sind, wird schließlich a mit b *verbunden*. Daher berechnet der Warshall-Algorithmus tatsächlich M_{R^+} .

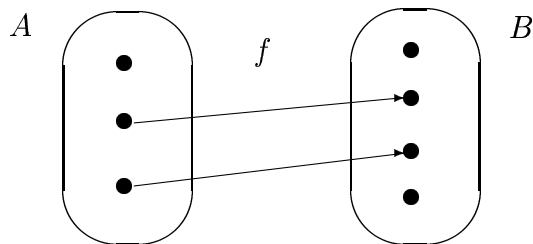
Kapitel 4

Funktionen

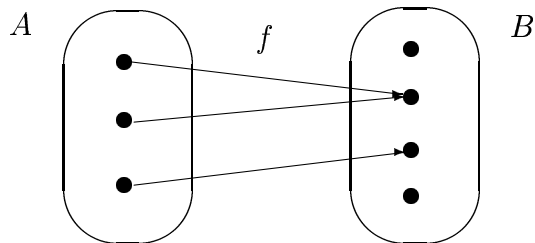
4.1 Injektive, surjektive und bijektiver Funktionen

Definition 4.1 Seien A, B zwei nichtleere Mengen.

Eine **partielle Funktion**¹ $f : A \rightarrow B$ ist eine Relation $R_f \subseteq A \times B$ mit der Eigenschaft, daß zu jedem $a \in A$ höchstens ein $b \in B$ mit aR_fb existiert. In diesem Fall schreibt man auch $b = f(a)$.



Eine **Funktion**² oder **Abbildung**³ $f : A \rightarrow B$ ist eine Relation $R_f \subseteq A \times B$ mit der Eigenschaft, daß zu jedem $a \in A$ genau ein $b \in B$ mit aR_fb existiert. Auch hier schreibt man $b = f(a)$.



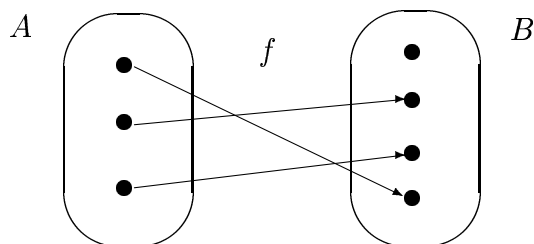
¹partial function

²function

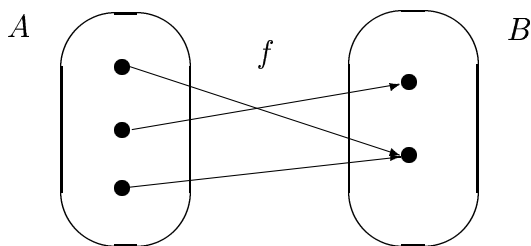
³mapping

Man kann eine Funktion $f : A \rightarrow B$ auch als *Zuordnung* oder als *Automat* interpretieren. Einem $a \in A$ wird ein (und nur ein) $b = f(a) \in B$ zugeordnet bzw. bei Vorgabe von $a \in A$ wird ein $b = f(a) \in B$ ausgegeben.

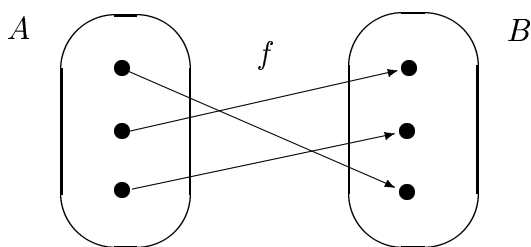
Definition 4.2 Eine Funktion $f : A \rightarrow B$ heißt **injektiv**⁴ oder **Injektion**, wenn es zu jedem $b \in B$ höchstens ein $a \in A$ mit $b = f(a)$ gibt.



Eine Funktion $f : A \rightarrow B$ heißt **surjektiv**⁵ oder **Surjektion**, wenn es zu jedem $b \in B$ mindestens ein $a \in A$ mit $b = f(a)$ gibt.



Eine Funktion $f : A \rightarrow B$ heißt **bijektiv**⁶ oder **Bijektion**, wenn es zu jedem $b \in B$ genau ein $a \in A$ mit $b = f(a)$ gibt.



Satz 4.3 Eine Funktion $f : A \rightarrow B$ ist genau dann bijektiv, wenn sie injektiv und surjektiv ist.

Satz 4.4 Haben zwei endliche Mengen A, B gleich viele Elemente, d.h. $|A| = |B|$, dann ist eine injektive Funktion $f : A \rightarrow B$ auch surjektiv und damit bijektiv. Entsprechend ist eine surjektive Funktion $f : A \rightarrow B$ auch injektiv und ebenfalls bijektiv.

⁴injective

⁵surjective

⁶bijjective

Definition 4.5 Sind $f : A \rightarrow B$ und $g : B \rightarrow C$ Funktionen, so wird die Zusammensetzung $g \circ f : A \rightarrow C$ durch $(g \circ f)(a) = g(f(a))$ definiert. $g \circ f$ ist dann wieder eine Funktion.

Man beachte, daß $g \circ f$ der Relation $R_f \circ R_g$ entspricht, wobei R_f und R_g die entsprechenden Relationen von f und g bezeichnen.

Satz 4.6 Sind die Funktionen $f : A \rightarrow B$ und $g : B \rightarrow C$ beide injektiv (bzw. surjektiv bzw. bijektiv), so ist auch $g \circ f : A \rightarrow C$ injektiv (bzw. surjektiv bzw. bijektiv).

Definition 4.7 Eine Funktion $f^{-1} : B \rightarrow A$ heißt die zu der Funktion $f : A \rightarrow B$ **inverse Funktion**⁷, wenn $f^{-1} \circ f = \text{id}_A$ und $f \circ f^{-1} = \text{id}_B$ sind.

Dabei bezeichnet id_A die **identische Funktion** auf einer Menge A , d.h. $\text{id}_A(a) = a$ für alle $a \in A$.

Satz 4.8 Eine Funktion $f : A \rightarrow B$ besitzt genau dann eine inverse Funktion $f^{-1} : B \rightarrow A$, wenn f bijektiv ist. Die inverse Funktion f^{-1} ist dann auch bijektiv.

4.2 Permutationen

Definition 4.9 Eine bijektive Abbildung $\pi : A \rightarrow A$ auf einer Menge A heißt **Permutation**⁸ auf A .

Die Menge aller Permutationen $\mathbf{S}(A)$ auf einer Menge A heißt auch **symmetrische Gruppe**⁹ auf A , da die Permutationen mit der Zusammensetzung $\circ$ ein Gruppe bilden (siehe Abschnitt 6.1). $\pi\sigma = \sigma \circ \pi$ bezeichnet man als das **Produkt** der Permutationen π und σ , d.h. zuerst wird π ausgeführt und danach σ . Weiters bezeichnet π^{-1} die zu π **inverse Permutation**.

Ist $A = \{a_1, a_2, \dots, a_n\}$ endlich, so identifiziert man die Elemente $a_1, a_2, \dots, a_n$ mit den Zahlen $1, 2, \dots, n$ und schreibt für $\mathbf{S}(A)$ auch $\mathbf{S}_n$.

Satz 4.10 Es gibt genau $n!$ verschieden Permutationen auf den Zahlen $\{1, 2, \dots, n\}$, d.h. $|\mathbf{S}_n| = n!$.

Permutationen auf endlichen Mengen besitzen verschiedene **Darstellungsarten**:

1. **Zweizeilige Darstellung:** $\pi : \{1, 2, \dots, n\} \rightarrow \{1, 2, \dots, n\}$ wird durch

$$\pi = \begin{pmatrix} 1 & 2 & 3 & \cdots & n-1 & n \\ \pi(1) & \pi(2) & \pi(3) & \cdots & \pi(n-1) & \pi(n) \end{pmatrix}$$

⁷inverse function

⁸permutation

⁹symmetric group

dargestellt, z.B.

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

Man kann daher eine Permutation auch als **Umordnung** interpretieren.

Die inverse Permutation π^{-1} erhält man am einfachsten dadurch, daß man die beiden Zeilen (in der zweizeiligen Darstellung von π) vertauscht und dann spaltenweise ordnet:

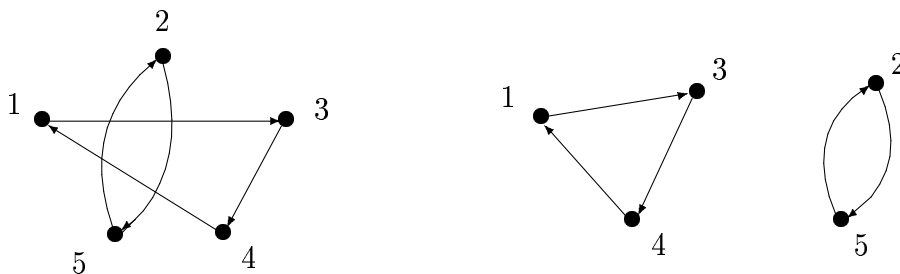
$$\pi^{-1} = \begin{pmatrix} 3 & 5 & 4 & 1 & 2 \\ 1 & 2 & 3 & 4 & 5 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 5 & 1 & 3 & 2 \end{pmatrix}.$$

Ist $\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 2 & 5 & 1 & 3 \end{pmatrix}$ eine weitere Permutation, so kann auch das Produkt $\pi\sigma$ in einfacher Weise bestimmt werden:

$$\pi\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 5 & 4 & 1 & 2 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 2 & 5 & 1 & 3 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 5 & 3 & 1 & 4 & 2 \end{pmatrix},$$

z.B. ist $\pi(1) = 3$ und $\sigma(3) = 5$, also $\pi\sigma(1) = 5$.

2. **Graphische Darstellung:** Die Zahlen $\{1, 2, \dots, n\}$ werden als Punkte (Knoten) dargestellt, und ist $j = \pi(i)$, so verläuft eine Kante von i nach j . (Es gibt daher genau n Kanten.)



Man beachte, daß von jedem Punkt genau eine Kante wegführt und zu jedem Punkt genau eine Kante hinführt. Der Graph muß daher in **Zyklen**¹⁰ zerfallen. Im obigen Beispiel sind dies die Zyklen $1, \pi(1) = 3, \pi(3) = 4, \pi(4) = 1$ und $2, \pi(2) = 5, \pi(5) = 2$. Die Graphen von Permutationen haben daher eine sehr einfache Struktur. Sie bilden eine Menge von Zyklen, wobei natürlich auch **Schlingen**¹¹ auftreten können, und zwar genau dann, wenn ein $j \in \{1, 2, \dots, n\}$ auf sich selbst abgebildet wird, d.h. $\pi(j) = j$. Solche Punkte heißen auch **Fixpunkte**¹².

3. **Zyklendarstellung:** Da jede Permutation $\pi \in \mathbf{S}_n$ in eine Menge von Zyklen zerfällt genügt es, einfach diese anzugeben. Ist etwa $1, \pi(1), \pi(\pi(1)), \pi(\pi(\pi(1))) = \pi^3(1), \dots, \pi^{k-1}(1), \pi^k(1) = 1$ der erste Zyklus, so stellt man diesen durch

$$(1 \ \pi(1) \ \pi^2(1) \ \dots \ \pi^{k-1}(1))$$

¹⁰cycle

¹¹loop

¹²fixed point

dar. Im obigen Beispiel gibt es also die Zyklen $(1\ 3\ 4)$ und $(2\ 5)$. Schreibt man nun alle Zyklen von π hintereinander an, so erhält man die Zyklen­darstellung von π , z.B.

$$\pi = (1\ 3\ 4)(2\ 5).$$

Schreibt man die Zyklen in einer anderen Reihenfolge an, bzw. vertauscht man innerhalb eines Zyklus die Elemente zyklisch, so erhält man auch eine Zyklen­darstellung dieser Permutation, z.B.

$$\pi = (5\ 2)(3\ 4\ 1),$$

die Menge der Zyklen wird dadurch ja nicht verändert.

Indem man in jedem Zyklus mit dem kleinsten Element (*Zyklenführer*) beginnt und die Anordnung der Zyklen so wählt, daß die Zyklenführer absteigend sortiert sind, kann man in dieser Zyklen­darstellung sogar die Klammern weglassen, und die Permutation läßt sich trotzdem eindeutig rekonstruieren.

Sei $\pi = b_1 b_1 \cdots b_n$ diese Darstellung. Ist b_{k+1} das erste Element $< b_1$, so ist $(b_1 b_1 \cdots b_k)$ der erste Zyklus von π . Ist weiters b_{m+1} das erste Element $< b_{k+1}$, so ist $(b_{k+1} b_{k+2} \cdots b_m)$ der zweite Zyklus von π , usw. Im obigen Beispiel ist

$$\pi = 2\ 5\ 1\ 3\ 4$$

diese klammernlose Zyklen­darstellung. Offensichtlich ist $(2\ 5)$ der erste Zyklus, da in jedem Zyklus das kleinste Element an erster Stelle steht. Der zweite Zyklus ist dann $(1\ 3\ 4)$.

Schlingen $\pi(j) = j$ können in der Zyklen­darstellung durch (j) angeschrieben werden (womit alle Zahlen $1, 2, \dots, n$ genau einmal aufgenommen werden) oder einfach weggelassen werden, d.h. die Permutation $\sigma = (1\ 4)(2)(3\ 5)$ kann auch durch $\sigma = (1\ 4)(3\ 5)$ dargestellt werden.

Definition 4.11 Eine Permutation $\pi \in \mathbf{S}_n$ heißt **Transposition**¹³, wenn die Zyklen­darstellung (in der Fixpunkte nicht aufgenommen werden) nur aus einem Zweierzyklus besteht. Ist $j \neq k$, so bezeichnet man die Transposition $(j\ k)$ durch π_{jk} .

Bei einer Transposition π_{jk} werden also nur die Zahlen j und k vertauscht.

Satz 4.12 Ein Zyklus $(b_1 b_2 \cdots b_k) \in \mathbf{S}_n$ der Länge k kann als Produkt von $k - 1$ Transpositionen

$$(b_1 b_2 \cdots b_k) = (b_1 b_2)(b_1 b_3) \cdots (b_1 b_k)$$

dargestellt werden.

Daher ist auch jede Permutation $\pi \in \mathbf{S}_n$ als Produkt von Transpositionen darstellbar.

¹³transposition

Definition 4.13 $\pi \in \mathbf{S}_n$ heißt **gerade Permutation**¹⁴, wenn sie als Produkt einer geraden Anzahl von Transposition darstellbar ist. Man schreibt dafür auch $\text{sgn}(\pi) = 1$. Die Menge $\mathbf{A}_n$ aller geraden Permutation $\pi \in \mathbf{S}_n$ wird als **alternierende Gruppe** bezeichnet. Entsprechend heißt $\pi \in \mathbf{S}_n$ heißt **ungerade Permutation**¹⁵, wenn sie als Produkt einer ungeraden Anzahl von Transposition darstellbar ist. In diesem Fall schreibt man $\text{sgn}(\pi) = -1$.

Beipielsweise sind Zyklen gerader Länge ungerade und Zyklen ungerader Länge gerade.

Satz 4.14 Das Produkt zweier gerader Permutation ist gerade. Ebenso ist das Produkt zweier ungerader Permutationen gerade. Andererseits ist das Produkt einer geraden mit einer ungeraden Permutation ungerade. In jedem Fall gilt daher

$$\text{sgn}(\pi \cdot \sigma) = \text{sgn}(\pi) \cdot \text{sgn}(\sigma).$$

Die Funktion sgn ist daher ein *Gruppenhomomorphismus* (siehe Abschnitt 6.1) von $\mathbf{S}_n$ nach $\{-1, 1\}$ und $\mathbf{A}_n$ ist der Kern von sgn .

Satz 4.15 Für $n \geq 2$ gibt es genauso viele gerade wie ungerade Permutationen in $\mathbf{S}_n$, d.h. $|\mathbf{A}_n| = n!/2$.

4.3 Abzählbare Mengen

Definition 4.16 Zwei Mengen A, B heißen **gleichmächtig**, wenn es eine bijektive Abbildung $f : A \rightarrow B$ gibt. Man schreibt dafür auch $|A| = |B|$ und bezeichnet $|A|$ als die **Kardinalität**¹⁶ von A .

Bei endlichen Mengen bedeutet, daß A und B gleichmächtig sind, daß sie dieselbe Anzahl von Elementen haben. Deshalb wird $|A|$ mit dieser Anzahl gleichgesetzt. Für unendliche Mengen A kann $|A|$ als unendlich große Anzahl (**Kardinalzahl**¹⁷) interpretiert werden.

Definition 4.17 Eine Menge A heißt **abzählbar**¹⁸, wenn sie gleichmächtig zu den natürlichen Zahlen $\mathbf{N}$ ist.

Die Kardinalität einer abzählbaren Menge wird mit $\aleph_0$ bezeichnet. ($\aleph$ – sprich “aleph” – ist ein hebräischer Buchstabe.)

¹⁴even permutation

¹⁵odd permutation

¹⁶cardinality

¹⁷cardinal

¹⁸countable

Für abzählbare Mengen A gibt es daher eine (bijektive) Funktion $f : \mathbf{N} \rightarrow A$, mit anderen Worten, man kann die Elemente von A durch

$$a_1 = f(1), a_2 = f(2), a_3 = f(3), \dots$$

tatsächlich *abzählen*. Jedes Element aus A wird in dieser Liste genau einmal aufgenommen.

Manchmal werden endliche Mengen auch als abzählbar bezeichnet. In diesem Fall bezeichnet man unendlichen Mengen, die abzählbar sind, auch als **abzählbar unendlich**.

Satz 4.18 *Jede unendliche Teilmenge einer abzählbaren Menge ist abzählbar.*

Satz 4.19 *Die Mengen $\mathbf{Z}, \mathbf{Q}$ sind abzählbar.*

Definition 4.20 *Eine unendlichen Menge B , die nicht abzählbar ist, heißt **überabzählbar**¹⁹, d.h. jede injektive Funktion $f : \mathbf{N} \rightarrow B$ ist nicht surjektiv.*

Satz 4.21 (Georg Cantor) *Die Menge der reellen Zahlen $\mathbf{R}$ ist überabzählbar.*

Georg Cantor ist, wie schon erwähnt wurde, der Schöpfer der Mengenlehre. Satz 4.21 ist einer seiner bemerkenswertesten Sätze, so sagt er doch aus, daß es eine Hierarchie des Unendlichen gibt. Es läßt sich sogar mit ganz einfachen Mitteln zeigen, daß es unendlich viele Unendlichkeitsstufen gibt.

Satz 4.22 *Sei A eine beliebige Menge. Dann sind A und deren Potenzmenge $\mathbf{P}(A)$ nicht gleichmächtig.*

Startet man beispielsweise mit den natürlichen Zahlen $\mathbf{N}$, so folgt aus Satz 4.22, daß $\mathbf{P}(\mathbf{N})$ (die Menge aller Teilmengen von $\mathbf{N}$) überabzählbar ist. (Übrigens sind $\mathbf{P}(\mathbf{N})$ und $\mathbf{R}$ gleichmächtig.) Weiters ist $\mathbf{P}(\mathbf{P}(\mathbf{N}))$ mächtiger als $\mathbf{P}(\mathbf{N})$ usw.

¹⁹uncountable

Kapitel 5

Graphen

5.1 Grundlegende Begriffe

Eine binäre Relation R auf einer Menge M kann (wie schon ausgeführt wurde) graphisch dargestellt werden, indem zwischen zwei Punkten $a, b \in M$ ein verbindender *Pfeil* gezeichnet wird, wenn a in Relation R zu b steht, d.h. aRb . Die dabei entstehenden *Graphen* können auch als selbständige (mathematische) Objekte betrachtet werden. Insbesondere in den Anwendungen erweist sich die Sprache der Graphen als außerordentlich nützlich.

Definition 5.1 Ein **Graph**¹ $G = (V, E)$ besteht aus einer *Knotenmenge*² $V = V(G)$ und einer *Kantenmenge*³ $E = E(G)$. Dabei ist eine Kante $e \in E(G)$ entweder **gerichtet**, d.h. ein geordnetes Paar $e = \langle v_1, v_2 \rangle$ von zwei Knoten $v_1, v_2 \in V(G)$, oder **ungerichtet**, d.h. ein ungeordnetes Paar $e = (v_1, v_2)$ von zwei Knoten $v_1, v_2 \in V(G)$. Im gerichtete Fall heißt v_1 **Anfangsknoten** und v_2 **Endknoten** von e .

Sind alle Kanten $e \in E(G)$ gerichtet, so spricht man von einem **gerichteten Graphen**⁴, sind hingegen alle Kanten $e \in E(G)$ ungerichtet, so heißt G **ungerichteter Graph**.

In der obigen Definition sind auch Kanten der Form $\langle v, v \rangle$ (bzw. (v, v)) zugelassen. So eine Kante heißt auch **Schlinge**⁵.

Sind zwei Knoten v, w eines Graphen durch eine Kante verbunden, so heißen v und w auch **adjazent**⁶. Weiters **inzidieren** die Knoten v, w mit einer Kante, die sie verbindet.

Die Anzahl der Knoten eines Graphen G wird auch mit $\alpha_0(G) = |V(G)|$ und die Anzahl der Kanten mit $\alpha_1(G) = |E(G)|$ bezeichnet.

¹graph

²vertex set

³edge set

⁴digraph

⁵loop

⁶adjacent

Es ist auch möglich, Graphen mit **Mehrfachkanten** zu betrachten. Hier kann man aber Kanten nicht mehr mit Paaren von Knoten identifizieren. Einer allgemeinen Kante e werden zwei Knoten v_1, v_2 zugeordnet, die als Anfangs- bzw. Endknoten interpretiert werden können. Zwei verschiedene Kanten können daher (in diesem Rahmen) durchaus dieselben Anfangs- und Endknoten besitzen.

Definition 5.2 *Ein gerichteter oder ungerichteter Graph $G = (V, E)$ heißt **schlicht**⁷ oder **einfach**, wenn G keine Schlingen (und keine Mehrfachkanten) enthält.*

Definition 5.3 *In einem ungerichteten schlichten Graphen G heißen die zu $v \in V(G)$ adjazenten Knoten*

$$\Gamma(v) = \{w \in V(G) \mid (v, w) \in E(G)\}$$

Nachbarn⁸ von v .

Die Anzahl

$$d(v) = |\Gamma(v)| = |\{w \in V(G) \mid (v, w) \in E(G)\}|$$

*der Nachbarn von $v \in V(G)$ ist der **Knotengrad**⁹ von $v \in V(G)$*

In einem gerichteten Graphen G heißen die Elemente von

$$\Gamma^+(v) = \{w \in V(G) \mid \langle v, w \rangle \in E(G)\}$$

Nachfolger¹⁰ des Knoten $v \in V(G)$ und

$$\Gamma^-(v) = \{w \in V(G) \mid \langle w, v \rangle \in E(G)\}$$

Vorgänger¹¹ von $v \in V(G)$. $\Gamma(v) = \Gamma^+(v) \cup \Gamma^-(v)$ bilden die Nachbarn von $v \in V(G)$.

Die Anzahl

$$d^+(v) = |\Gamma^+(v)| = |\{w \in V(G) \mid \langle v, w \rangle \in E(G)\}|.$$

*der Nachfolger von $v \in V(G)$ ist der **Weggrad**¹² von $v \in V(G)$.*

Die Anzahl

$$d^-(v) = |\{w \in V(G) \mid \langle w, v \rangle \in E(G)\}|.$$

*der Vorgänger von $v \in V(G)$ ist der **Hingrad**¹³ von $v \in V(G)$*

Man beachte, daß bei gerichteten Graphen in dieser Definition Schlingen zugelassen sind. Ließe man auch im ungerichteten Fall Schlingen zu, so müßten diese bei der Gradberechnung doppelt gezählt werden.

⁷simple graph

⁸neighbor

⁹vertex degree

¹⁰successor

¹¹predecessor

¹²outdegree

¹³indegree

Satz 5.4 In einem ungerichteten schlichten Graphen G gilt

$$\sum_{v \in V(G)} d(v) = 2|E(G)|.$$

In einem gerichteten Graphen G gilt hingegen

$$\sum_{v \in V(G)} d^+(v) = \sum_{v \in V(G)} d^-(v) = |E(G)|.$$

Definition 5.5 Ein Graph $G' = (V', E')$ heißt **Teilgraph**¹⁴ eines Graphen $G = (V, E)$ wenn $V' \subseteq V$ und $E' \subseteq E$ gelten.

Definition 5.6 Eine Folge von Kanten $e_1, e_2, \dots, e_k \in E(G)$ eines ungerichteten Graphen G heißt **Kantenfolge**¹⁵, wenn es Knoten $v, v_1, v_2, \dots, v_{k-1}, w \in V(G)$ mit

$$e_1 = (v, v_1), e_2 = (v_1, v_2), \dots, e_{k-1} = (v_{k-2}, v_{k-1}), e_k = (v_{k-1}, w)$$

gibt, d.h. man kann die Kanten $e_1, e_2, \dots, e_k$ "ohne Absetzen" durchlaufen. Man sagt auch daß die Kantenfolge e_j ($1 \leq j \leq k$) die Knoten v und w verbindet. Die Anzahl k der Kantenfolge ist die **Länge** der Kantenfolge. Eine Kantenfolge der Länge 0 besteht aus keiner Kante und wird als **leere Kantenfolge** bezeichnet. Sie verbindet jeden Knoten mit sich selbst.

Eine Folge von Kanten $e_1, e_2, \dots, e_k \in E(G)$ eines gerichteten Graphen G heißt **Kantenfolge**, wenn für je zwei aufeinanderfolgende Kanten e_j, e_{j+1} ($1 \leq j < k$) der Endknoten von e_j mit dem Anfangsknoten von e_{j+1} übereinstimmt, d.h. es gibt Knoten $v, v_1, v_2, \dots, v_{k-1}, w$ mit

$$e_1 = \langle v, v_1 \rangle, e_2 = \langle v_1, v_2 \rangle, \dots, e_{k-1} = \langle v_{k-2}, v_{k-1} \rangle, e_k = \langle v_{k-1}, w \rangle.$$

Wiederum spricht man von einer Kantenfolge, die die Knoten v und w (gerichtet) verbindet, v ist der Anfangsknoten und w der Endknoten. Ebenso ist k die Länge der Kantenfolge.

Eine Kantenfolge $e_1, e_2, \dots, e_k \in E(G)$ in einem Graphen G heißt **Kantenzug**, wenn alle Kanten e_j ($1 \leq j \leq k$) voneinander verschieden sind.

Eine Kantenfolge $e_1, e_2, \dots, e_k \in E(G)$ in einem Graphen G heißt **Weg**¹⁶, wenn alle Knoten, die mit den Kanten e_j ($1 \leq j \leq k$) inzidieren, voneinander verschieden sind.

Verbindet eine Kantenfolge (resp. ein Kantenzug resp. ein Weg) die Knoten v und w , so bezeichnet man dies auch durch $KF(v, w)$ (resp. durch $KZ(v, w)$ resp. durch $W(v, w)$).

Eine Kantenfolge, die einen Knoten $v \in V(G)$ mit sich selbst verbindet, heißt **geschlossen**¹⁷.

¹⁴subgraph

¹⁵path

¹⁶simple path

¹⁷closed path

Eine geschlossene Kantenfolge $K(v, v)$ in einem ungerichteten Graphen heißt **Kreis**¹⁸, wenn alle Knoten dieser Kantenfolge mit Ausnahme von v voneinander verschieden sind und keine Kante mehrfach vorkommt.

Eine geschlossene Kantenfolge $K(v, v)$ in einem gerichteten Graphen heißt **Zyklus**¹⁹, wenn alle Knoten dieser Kantenfolge mit Ausnahme von v voneinander verschieden sind.

Satz 5.7 Werden in einem Graphen G zwei Knoten v, w durch eine Kantenfolge verbunden, so gibt es auch einen Weg, der v mit w verbindet.

Gibt es in einem Graphen G eine geschlossene Kantenfolge positiver Länge, so gibt es auch einen Kreis (resp. Zyklus) positiver Länge.

Definition 5.8 Zwei Graphen G_1, G_2 heißen **isomorph**²⁰, wenn es eine bijektive Abbildung

$$\phi: V(G_1) \rightarrow V(G_2)$$

gibt, so daß eine Kante (v, w) (bzw. $\langle v, w \rangle$) genau dann in $E(G_1)$ enthalten ist, wenn die Kante $(\phi(v), \phi(w))$ (bzw. $\langle \phi(v), \phi(w) \rangle$) in $E(G_2)$ enthalten ist.

5.2 Adjazenz- und Inzidenzmatrix

Definition 5.9 Sei G ein Graph mit Knotenmenge $V(G) = \{v_1, v_2, \dots, v_n\}$. Die **Adjazenzmatrix**²¹ $A(G) = (a_{ij})_{1 \leq i, j \leq n}$ ist durch

$$a_{ij} = \begin{cases} 1 & \text{für } (v_i, v_j) \in E(G) \text{ bzw. } \langle v_i, v_j \rangle \in E(G), \\ 0 & \text{sonst} \end{cases}$$

gegeben, d.h. $A(G)$ ist die Matrix M_R der zu G gehörigen Relation R .

Man beachte, daß die Adjazenzmatrix eines ungerichteten Graphen immer symmetrisch ist.

Schlingen äußern sich in der Adjazenzmatrix durch Einträge 1 in der Diagonale. Schlichte Graphen haben daher eine symmetrische Adjazenzmatrix mit verschwindender Diagonale.

Es ist auch sinnvoll, Graphen mit Mehrfachkanten eine Adjazenzmatrix zuzuordnen, wo a_{ij} die Anzahl der Kanten von v_i nach v_j bezeichnet.

Mit Hilfe der Adjazenzmatrix $A(G)$ eines Graphen lassen sich direkt die Knotengrade ablesen.

¹⁸circle, simple closed path

¹⁹cycle

²⁰isomorphic

²¹adjacency matrix

Satz 5.10 Sei $A(G) = (a_{ij})_{1 \leq i, j \leq n}$ die Adjazenzmatrix eines Graphen G mit Knotenmenge $V(G) = \{v_1, v_2, \dots, v_n\}$. Ist G ungerichtet und schlicht, so gilt

$$d(v_i) = \sum_{j=1}^n a_{ij} = \sum_{j=1}^n a_{ji}.$$

Ist G gerichtet, so gilt entsprechend

$$d^+(v_i) = \sum_{j=1}^n a_{ij} \quad \text{und} \quad d^-(v_i) = \sum_{j=1}^n a_{ji}.$$

In einem Graphen G heißt ein Knoten w von einem Knoten v aus **erreichbar**, wenn es eine Kantenfolge gibt, die v mit w verbindet. Die **Erreichbarkeitsrelation** ist ja nichts anderes als die reflexiv-transitive Hülle R^* der zu G gehörigen Relation R . $A(G)$ ist die Matrix M_R dieser Relation R . Die Erreichbarkeitsrelation kann daher (z.B. mit Hilfe des Warshallalgorithmus) effektiv berechnet werden.

Mit Hilfe des folgenden Satzes lassen sich sogar quantitative Resultate erzielen.

Satz 5.11 Sei $A(G) = (a_{ij})_{1 \leq i, j \leq n}$ die Adjazenzmatrix eines Graphen G mit Knotenmenge $V(G) = \{v_1, v_2, \dots, v_n\}$. Dann ist der Eintrag $a_{ij}^{[k]}$ der k -ten Potenz der Adjazenzmatrix

$$A(G)^k = (a_{ij}^{[k]})_{1 \leq i, j \leq n}$$

die Anzahl der Kantenfolgen der Länge k von v_i nach v_j .

Definition 5.12 Sei G ein Graph mit Knotenmenge $V(G) = \{v_1, v_2, \dots, v_n\}$ und Kantenmenge $E(G) = \{e_1, e_2, \dots, e_m\}$.

Die **Inzidenzmatrix** $B(G) = (b_{ij})_{1 \leq i \leq n, 1 \leq j \leq m}$ eines ungerichteten schlichten Graphen G ist durch

$$b_{ij} = \begin{cases} 1 & \text{wenn } v_i \text{ mit } e_j \text{ inzidiert,} \\ 0 & \text{sonst} \end{cases}$$

gegeben.

Die **Inzidenzmatrix** $B(G) = (b_{ij})_{1 \leq i \leq n, 1 \leq j \leq m}$ eines gerichteten Graphen G ist durch

$$b_{ij} = \begin{cases} +1 & \text{wenn } v_i \text{ Anfangsknoten von } e_j \text{ ist,} \\ -1 & \text{wenn } v_i \text{ Endknoten von } e_j \text{ ist,} \\ 0 & \text{sonst} \end{cases}$$

gegeben.

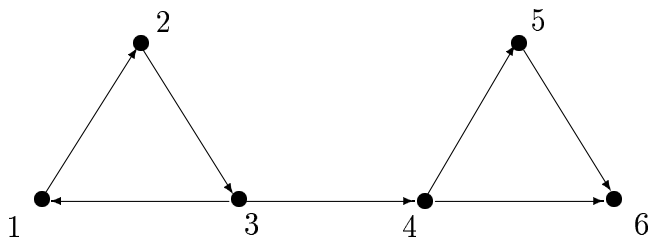
5.3 Azyklische Graphen

Definition 5.13 Ein gerichteter Graph $G = (V, E)$ heißt **azyklisch**²² wenn er keine Zyklen (positiver Länge) enthält.

Der folgende **Markierungsalgorithmus** entscheidet, ob ein Graph azyklisch ist oder nicht.

1. (a) Bestimme alle Knoten v mit Weggrad $d^+(v) = 0$ und markiere diese mit $\oplus$.
 (b) Wurde in 1.(a) kein solcher Knoten gefunden, so ist G **nicht azyklisch** $\rightarrow$ **ENDE**.
2. (a) Sind bereits alle Knoten von G mit $\oplus$ markiert, so ist G **azyklisch** $\rightarrow$ **ENDE**.
 (b) Suche (unmarkierte) Knoten, von denen nur Kanten zu schon markierten führen und markiere diese mit $\oplus$.
3. (a) Wurde in 2.(b) mindestens ein Knoten markiert, so wiederhole 2.
 (b) Wurde in 2.(b) kein Knoten markiert, so ist G **nicht azyklisch** $\rightarrow$ **ENDE**.

Beispiel 5.14 Wendet man den Markierungsalgorithmus auf das Beispiel



an, so durchläuft man folgende Schritte und markiert die angegebenen Punkte:

- 1.(a) $6 \rightarrow \oplus$
- 1.(b)
- 2.(a)
- 2.(b) $5 \rightarrow \oplus$
- 3.(a)
- 2.(a)
- 2.(b) $4 \rightarrow \oplus$
- 2.(a)
- 2.(b)
- 3.(a)
- 3.(b) G ist nicht azyklisch

²²cyclefree graph

Dem Markierungsalgorithmus liegt der folgende Satz zugrunde.

Satz 5.15 *Jeder azyklische Graph G besitzt einen Knoten $v \in V(G)$ mit Weggrad $d^+(v) = 0$ (und entsprechend auch einen Knoten $w \in V(G)$ mit Eingrad $d^-(w) = 0$).*

Dieser Satz begründet zunächst einmal den Schritt 1.(b).

Weiters beachte man, daß die Knoten, die im Schritt 1.(a) (mit $\oplus$) markiert werden, sicherlich nicht in einem Zyklus liegen können, d.h. *streicht* man alle in diesem Schritt markierten Knoten und alle zu ihnen hinführenden Kanten, so entsteht ein kleinerer Graph G' , der genau dann azyklisch ist, wenn G azyklisch ist.

Im Schritt 2.(b) werden alle Knoten markiert, von denen nur Kanten zu bereits markierten führen. In G' sind dies aber genau jene Knoten $v' \in V(G')$, die in G' Weggrad $d_{G'}^+(v') = 0$ haben, d.h. 2.(b) ist nichts anderes als 1.(a) angewandt auf G' . Das sukzessive Markieren stellt daher ein *Abarbeiten* jener Knoten dar, die sicherlich in keinem Zyklus von G liegen.

Schließlich trifft einer der beiden folgenden Fälle zu:

- A. Es können alle Knoten markiert werden (Fall 2.(a)), d.h. kein Knoten von G liegt in einem Zyklus, der Graph G ist daher azyklisch.
- B. Es können nicht alle Knoten markiert werden (Fall 3.(b)), d.h. nach dem Entfernen von gewissen Knoten (die in keinem Zyklus von G liegen können) bleibt ein Graph G'' über, der keine Knoten $v'' \in V(G'')$ mit Weggrad $d_{G''}^+(v'') = 0$ besitzt. Nach Satz 5.15 ist G'' und damit auch G nicht azyklisch.

5.4 Zusammenhang von Graphen

5.4.1 Ungerichtete Graphen

Definition 5.16 *Ein ungerichteter Graph G heißt **zusammenhängend**²³, wenn es zwischen je zwei Knoten $v, w \in V(G)$ einen Weg $W(v, w)$ gibt.*

*Die maximalen zusammenhängenden Teilgraphen eines ungerichteten Graphen G heißen **(Zusammenhangs-) Komponenten**²⁴ von G .*

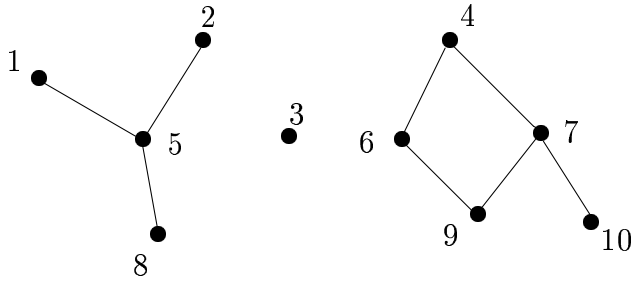
Ein Graph G ist daher genau dann zusammenhängend, wenn er nur aus einer Zusammenhangskomponente besteht.

Beispiel 5.17 Die Zusammenhangskomponenten des folgenden Graphen

sind $K_1 = \{1, 2, 5, 8\}$, $K_2 = \{3\}$ und $K_3 = \{4, 6, 7, 9, 10\}$.

²³connected

²⁴component



Die Zusammenhangskomponenten eines ungerichteten Graphen bilden eine Partition der Knotenmenge. Dieser Partition entspricht in natürlicher Weise eine Äquivalenzrelation R^* (siehe Abschnitt 3.3), es gilt vR^*w genau dann, wenn es einen Weg $W(v, w)$ gibt. Diese Relation ist daher nichts anderes als die reflexiv-transitive Hülle jener (symmetrischen) Relation R , die dem (ungerichteten) Graphen G entspricht.

Die reflexiv-transitive Hülle kann in einfacher Weise (etwa mit dem Warshallalgorithmus) algorithmisch bestimmt werden. Die Komponenten von G können dann sofort aus der Matrix $M_{R^*} = (m_{ij}^*)_{1 \leq i, j \leq n}$ abgelesen werden, so ist etwa

$$K_j = \{v_i \in V(G) \mid m_{ij}^* = 1\}$$

jene Knotenmenge jener Komponente von G , die v_j enthält.

5.4.2 Gerichtete Graphen

Definition 5.18 Ein gerichteter Graph G heißt **stark zusammenhängend**²⁵, wenn für je zwei Knoten $v, w \in V(G)$ gerichtete Wege $W_1(v, w)$ von v nach w und $W_2(w, v)$ von w nach v existieren.

Die maximalen stark zusammenhängenden Teilgraphen eines gerichteten Graphen G heißen **starke Zusammenhangskomponenten** oder **Komponenten des starken Zusammenhangs** von G .

Wieder bilden die Komponenten des starken Zusammenhangs eine Partition der Knotenmenge. Zwei Knoten $v, w \in V(G)$ stehen in der dazugehörigen Äquivalenzrelation R^{**} in Relation, i.Z. $vR^{**}w$, wenn es einen Weg $W(v, w)$ und einen Weg $W(w, v)$ gibt, d.h. $R^{**} = R^* \cap (R^{-1})^*$. Die Matrix $M_{R^{**}} = (m_{ij}^{**})_{1 \leq i, j \leq n} = M_{R^*} \boxtimes M_{R^*}^T$ kann daher etwa mit dem Warshallalgorithmus ermittelt werden. Die Menge

$$K_j = \{v_i \in V(G) \mid m_{ij}^{**} = 1\}$$

ist die Komponente des starken Zusammenhangs, die v_j enthält.

²⁵strictly connected

Definition 5.19 Der Schatten G^u eines gerichteten Graphen G ist ein ungerichteter Graph mit der selben Knotenmenge wie G ($V(G^u) = V(G)$) und (v, w) ist eine (ungerichtete) Kante von G^u , wenn $\langle v, w \rangle$ oder $\langle w, v \rangle$ in $E(G)$ enthalten sind, d.h. gerichtete Kanten werden einfach durch ungerichtete ersetzt.

Ein gerichteter Graph G heißt **schwach zusammenhängend**²⁶, wenn der Schatten G^u zusammenhängend ist. Entsprechend sind die **schwachen Zusammenhangskomponenten** die Komponenten von G^u .

Die Komponenten des schwachen Zusammenhangs können daher wie im ungerichteten Fall bestimmt werden.

Definition 5.20 Sei G ein gerichteter Graph, und bezeichnen $K_1, K_2, \dots, K_r$ die Komponenten des starken Zusammenhangs. Die **Reduktion**²⁷ $G_R = (V(G_R), E(G_R))$ ist jener (gerichtete) Graph mit Knotenmenge

$$V(G_R) = \{K_1, K_2, \dots, K_r\},$$

wobei eine (gerichtete) Kante $\langle K_i, K_j \rangle$ ($i \neq j$) genau dann in $E(G_R)$ enthalten ist, wenn es Knoten $v \in K_i$ und $w \in K_j$ mit $\langle v, w \rangle \in E(G)$ gibt.

Satz 5.21 Die Reduktion G_R eines gerichteten Graphen G ist ein azyklischer Graph.

Definition 5.22 Eine **Knotenbasis** B eines gerichteten Graphen G ist eine Teilmenge $B \subseteq V(G)$ mit den folgenden Eigenschaften:

1. Zu jedem Knoten $v \in V(G)$ gibt es einen Knoten $b \in B$ und einen Weg $W(b, v)$.
2. B ist minimal bezüglich der Eigenschaft 1., d.h. jede echte Teilmenge $B' \subseteq B$ (mit $B' \neq B$) erfüllt 1. nicht.

Für azyklische Graphen läßt sich eine (die) Knotenbasis sehr einfach ermitteln.

Satz 5.23 $B = \{v \in V(G) \mid d^-(v) = 0\}$ ist die einzige Knotenbasis eines azyklischen Graphen G .

Man beachte, daß wegen Satz 5.15 die so definierte Menge B in einem azyklischen Graphen nichtleer ist.

²⁶weakly connected

²⁷reduction

Satz 5.24 Sei G ein gerichteter Graph, und bezeichnen $K_1, K_2, \dots, K_l$ jene Komponenten des starken Zusammenhangs von G mit Hingrad

$$d_{G_R}^-(K_j) = 0 \quad (1 \leq j \leq l),$$

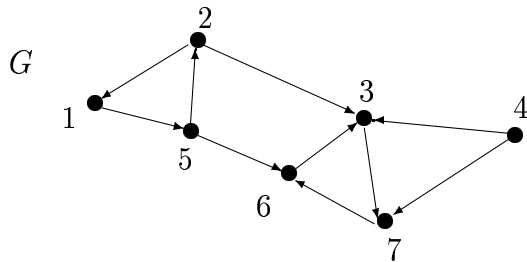
d.h. $B_R = \{K_1, K_2, \dots, K_l\}$ bilden die Knotenbasis von G_R . Ist $v_1 \in K_1, v_2 \in K_2, \dots, v_l \in K_l$ eine beliebige Auswahl von Knoten aus diesen Komponenten, so ist

$$B = \{v_1, v_2, \dots, v_l\}$$

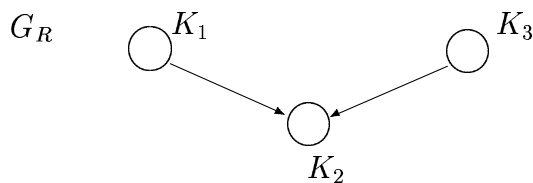
eine Knotenbasis von G . Weiters erhält man auf diese Art und Weise alle Knotenbasen von G .

Korollar 5.25 Alle Knotenbasen eines gerichteten Graphen haben dieselbe Anzahl von Elementen.

Beispiel 5.26 Im folgenden Graphen G sind die Komponenten des starken Zusammenhangs $K_1 = \{1, 2, 5\}$, $K_2 = \{3, 6, 7\}$ und $K_3 = \{4\}$



Die Reduktion G_R von G hat nun die Gestalt

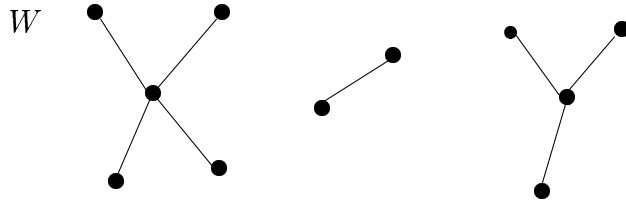


$B_R = \{K_1, K_3\}$ ist die Knotenbasis von G_R . Damit sind $B_1 = \{1, 4\}$, $B_2 = \{2, 4\}$ und $B_3 = \{5, 4\}$ alle Knotenbasen von G .

5.5 Bäume und Wälder

Definition 5.27 Ein schlichter Graph W , der keine Kreise enthält, heißt **Wald**²⁸.

Ein Wald T , der auch zusammenhängend ist, heißt **Baum**²⁹.



Die Zusammenhangskomponenten von einem Wald sind Bäume.

Man beachte, daß in einem Baum T zu je zwei Knoten v, w genau einen Weg $W(v, w)$ gibt. (Da T zusammenhängend ist, muß es es einen Weg $W(v, w)$ geben. Gäbe es aber einen weiteren von $W(v, w)$ verschiedenen Weg, $\tilde{W}(v, w)$, so müßte es auch eine geschlossene Kantenfolge positiver Länge und damit einen Kreis geben, was aber definitionsgemäß ausgeschlossen ist.) Die Länge dieses Weges $W(v, w)$ bezeichnet man als den **Abstand**³⁰ $d_T(v, w)$.

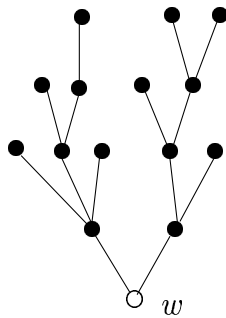
Satz 5.28 Für einen Baum T gilt

$$\alpha_0(T) = \alpha_1(T) + 1.$$

Entsprechend gilt für einen Wald W mit k Komponenten

$$\alpha_0(W) = \alpha_1(W) + k.$$

Zeichnet man in einem Baum einen Knoten $w \in E(T)$ (**Wurzel**)³¹ aus, so kann man sich die Struktur eines Baumes sehr einfach verdeutlichen. Zeichnet man in einer graphischen Darstellung die Nachbarn von w oberhalb von w und deren Nachbarn (mit der Ausnahme w) wieder darüber usw., so entsteht tatsächlich ein Bild, das einem Baum ähnelt.



Man beachte, daß hier tatsächlich von jedem Knoten in neue Knoten verzweigt wird, da ein Baum definitionsgemäß kreisfrei ist.

²⁸forest

²⁹tree

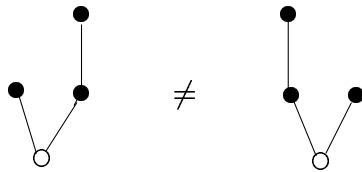
³⁰distance

³¹root

Definition 5.29 Ein Baum T , bei dem ein Knoten w (**Wurzel**) ausgezeichnet ist, heißt **Wurzelbaum**³².

Knoten $v \in V(T)$, $v \neq w$, eines Wurzelbaumes T mit Knotengrad $d(v) = 1$ heißen **Endknoten**, **externe Knoten**³³ oder **Blätter**³⁴. Alle anderen Knoten $v \in V(T)$ heißen **interne Knoten**³⁵.

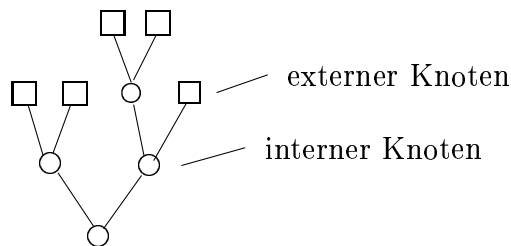
Ein **ebener Wurzelbaum**³⁶ ist ein Wurzelbaum, bei dem die Links-Rechts-Reihenfolge der Nachfolgeäste der Wurzel wesentlich ist, z.B. sind werden die beiden Wurzelbäume



als verschieden betrachtet, obwohl sie natürlich denselben Graphen repräsentieren.

Der Vorteil von Wurzelbäumen ist, daß sie eine *rekursive Struktur* haben. Ist w' ein Nachbar von der Wurzel w , so bilden alle Knoten $v \in V(T)$, deren Verbindungsweg $W(v, w)$ zu w über w' führt, gemeinsam mit w' einen kleineren Wurzelbaum usw.

Definition 5.30 Ein **Binärbaum**³⁷ T ist ein ebener Wurzelbaum, bei dem jeder interne Knoten $v \in V(T)$ mit Ausnahme der Wurzel w Knotengrad $d(v) = 3$ und die Wurzel Knotengrad $d(w) = 2$ hat, d.h. jeder interne Knoten hat genau zwei Nachfolgeknoten.



Satz 5.31 Jeder Binärbaum mit n internen Knoten hat genau $n+1$ externe Knoten, und die Anzahl der verschiedenen Binärbäumen mit n internen Knoten beträgt

$$\frac{1}{n+1} \binom{2n}{n} \sim 4^n n^{-3/2} \frac{1}{\sqrt{\pi}} \quad (n \rightarrow \infty).$$

³²rooted tree

³³external node

³⁴leaf

³⁵internal node

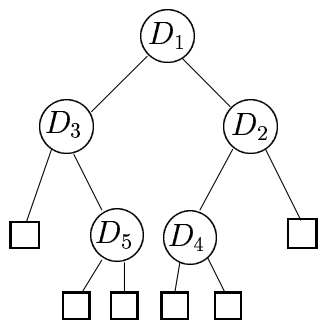
³⁶planar tree

³⁷binary tree

Binärbäume werden wegen ihres einfachen rekursiven Aufbaus als **Datenstrukturen**³⁸ verwendet:

1. **Binärer Suchbaum:**³⁹ Es sollen n Daten $D_1, D_2, \dots, D_n$, die mit n paarweise vergleichbaren *Schlüsseln* $k_1, k_2, \dots, k_n$ versehen sind, in einem (noch aufzubauenden) Binärbaum gespeichert werden. Zunächst wird D_1 als *Wurzel* verwendet. Nun nehme man an, die ersten j Daten $D_1, D_2, \dots, D_j$ seien schon eingetragen, dann wird der Platz für D_{j+1} so bestimmt, daß man bei der Wurzel D_1 beginnend die Schlüssel miteinander vergleicht. Ist etwa $k_{j+1} < k_1$ so geht man (im Binärbaum) nach links weiter, und ist $k_{j+1} > k_1$ so geht man nach rechts weiter. Erreicht man jetzt einen weiteren Knoten D_i , so vergleicht man die Schlüssel k_{j+1} und k_i in derselben Art und Weise und setzt dieses Verfahren so lange fort bis man zu einem unbesetzten Knoten kommt. Dort trägt dort D_{j+1} ein.

Sind etwa die Schlüssel der Datensätze D_1, D_2, D_3, D_4, D_5 durch $k_1 = 3, k_2 = 5, k_3 = 1, k_4 = 4, k_5 = 2$ gegeben, so erhält man folgenden Datenaufbau:



Die Daten $D_1, D_2, \dots, D_n$ sind daher in den internen Knoten gespeichert.

2. **Digitaler Suchbaum:**⁴⁰ Wieder sollen n Daten $D_1, D_2, \dots, D_n$ gespeichert werden, wobei die Schlüssel k_j ($1 \leq j \leq n$) 0-1-Folgen sind. D_1 ist wie im binären Suchbaum die Wurzel. Die Platz der weiteren Datensätze D_j wird aber jetzt nicht durch Vergleich der Schlüssel gefunden, sondern man orientiert sich nur am Schlüssel k_j . Ist das erste bit 0, so geht man (bei der Wurzel beginnend) nach links weiter, ist das erste bit 1, so geht man nach rechts weiter. Im nächsten Schritt verwendet man das zweite bit in derselben Art und Weise. Dieses Verfahren führt man solange durch bis man einen freien Platz gefunden hat. Dort trägt man D_j ein.

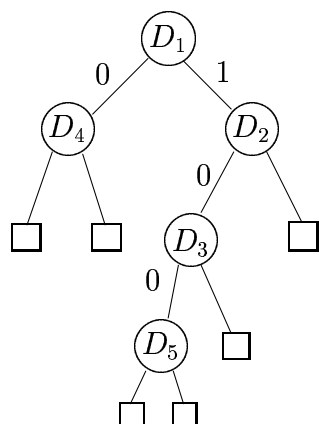
Haben die Datensätze D_1, D_2, D_3, D_4, D_5 die Schlüssel $k_1 = 01101\dots, k_2 = 10110\dots, k_3 = 10100\dots, k_4 = 00101\dots, k_5 = 10010\dots$ so ergibt sich folgender auf der nächsten Seite dargestellter Datenaufbau.

Man beachte, daß die Struktur eines binären bzw. digitalen Suchbaums von der Reihenfolge der Dateneinträge abhängt, der erste Datensatz wird z.B. immer als Wurzel verwendet.

³⁸data structure

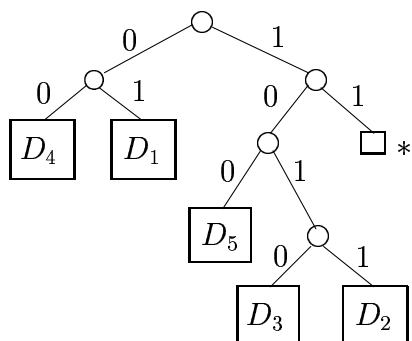
³⁹binary search tree

⁴⁰digital search tree



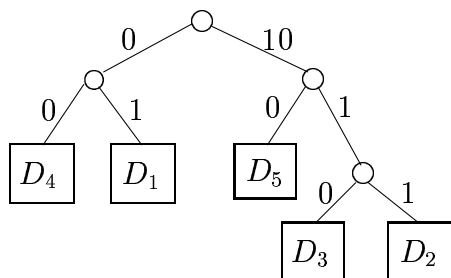
3. **Tries:**⁴¹ Hier werden die Daten $D_1, D_2, \dots, D_n$ nicht als interne Knoten eines Binärbaumes gespeichert, sondern in den externen Knoten. Die Schlüssel k_j ($1 \leq j \leq n$) sind wieder 0-1-Folgen. Die Position von D_j wird durch den kürzesten eindeutigen Anfangsabschnitt von k_j bestimmt, wobei 0 wieder *gehe nach links* und 1 *gehe nach rechts* bedeutet.

Verwendet man dieselben Schlüssel $k_1 = \underline{0}1101\dots$, $k_2 = \underline{1}0110\dots$, $k_3 = \underline{1}0100\dots$, $k_4 = \underline{0}0101\dots$, $k_5 = \underline{1}0010\dots$ wie im letzten Beispiel (nur daß jetzt der kürzeste eindeutige Anfangsabschnitt der Übersicht halber unterstrichen wurde), so hat der zugehörige Trie die folgende Gestalt:



Der Baum ist bei Tries nicht mehr von der Reihenfolge der Daten abhängig. (Es ist bei einem Trie übrigens sehr einfach, neue Datensätze einzutragen. Die jeweiligen kürzesten eindeutigen Anfangsabschnitte müssen dabei nicht immer neu berechnet werden. Im Gegenteil, sie lassen sich aus dem entstandenen Trie ablesen.) Ein Nachteil des Trie ist, daß unbesetzte Endknoten (*) auftreten können. Dieser Nachteil kann dadurch behoben werden, daß *unnötige* Kanten *elminiert* werden, und man erhält den sogenannten **Patricia Trie**:

⁴¹information retrieval



5.6 Planare Graphen

Definition 5.32 Ein Graph G heißt **planar**⁴² oder **eben**, wenn G kreuzungsfrei in der Ebene $\mathbf{R}^2$ dargestellt werden kann, d.h. die Kurven, die die Kanten repräsentieren, haben außer in jenen Punkten, die die Knoten repräsentieren, keine weiteren Schnittpunkte.

Die kreuzungsfreie Darstellung eines Graphen G zerlegt die Ebene in eine endliche Anzahl von Gebieten. Es stellt sich heraus, daß diese Anzahl unabhängig davon ist, wie man G in der Ebene kreuzungsfrei repräsentiert. Man bezeichnet diese Anzahl durch $\alpha_2(G)$.

Satz 5.33 (Eulersche Polyederformel) Für einen zusammenhängenden planaren Graphen G gilt

$$\alpha_0(G) - \alpha_1(G) + \alpha_2(G) = 2.$$

Durch Projektion eines konvexen Polyeders P (Durchschnitt von endlich vielen Halbräumen im $\mathbf{R}^3$) auf eine im Inneren des Polyeders liegende Kugel erhält man auf der Kugeloberfläche eine kreuzungsfreie Darstellung eines Graphen G , der dieselbe Anzahl von Knoten, Kanten und Gebieten hat wie das Polyeder Eckpunkte, Kanten und Flächen. Projiziert man nun die Kugeloberfläche mittels stereographischer Projektion auf die Ebene (wobei als *Nordpol* weder ein Knotenpunkt noch ein Punkt einer Kante des Graphen auf der Kugeloberfläche verwendet werden darf) so wird die kreuzungsfreie Darstellung des Graphen G auf der Kugeloberfläche auf eine kreuzungsfreie Darstellung von G in der Ebene projiziert. Natürlich bleibt auch bei dieser Projektion die Anzahl der Knoten, Kanten und Gebiete gleich. Daher gilt auch

$$\alpha_0(P) - \alpha_1(P) + \alpha_2(P) = 2,$$

wobei $\alpha_0(P)$ die Anzahl der Eckpunkte, $\alpha_1(P)$ die Anzahl der Kanten und $\alpha_2(P)$ die Anzahl der Flächen des Polyeders P bezeichnen.

Aus der eben geführten Überlegung ergibt sich auch die folgende Eigenschaft.

Satz 5.34 Ein Graph läßt sich genau dann kreuzungsfrei in der Ebene darstellen, wenn er kreuzungsfrei auf einer Kugeloberfläche darstellbar ist.

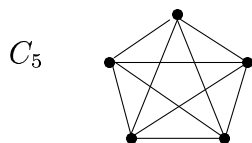
⁴²planar

Definition 5.35 Der **vollständige Graph**⁴³ C_n ist ein ungerichteter schlichter Graph mit n Knoten, bei dem jeder Knoten mit jedem anderen durch eine Kante verbunden ist.

Lemma 5.36 Für einen ungerichteten, schlichten, zusammenhängenden, planaren Graphen G gilt

$$\alpha_1(G) \leq 3\alpha_0(G) - 6.$$

Satz 5.37 Der vollständige Graph C_5 ist nicht planar.



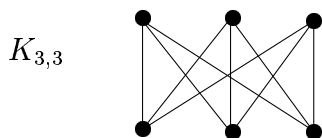
Definition 5.38 Ein (ungerichteter) Graph G heißt **paarer Graph**⁴⁴, wenn dessen Knotenmenge $V(G)$ in zwei disjunkte nichtleere Teilmengen V_1, V_2 zerlegt werden kann, so daß alle Kanten nur zwischen V_1 und V_2 verlaufen.

Der **vollständige paare Graph**⁴⁵ $K_{m,n}$ ($m, n \geq 1$) ist ein schlichter paarer Graph mit $m + n$ Knoten $V(G) = V_1 \cup V_2$, wobei V_1 m Knoten und V_2 n Knoten enthält und jeder Knoten aus V_1 mit allen Knoten aus V_2 durch eine Kante verbunden ist.

Lemma 5.39 Für einen ungerichteten, schlichten, planaren Graphen G mit der Eigenschaft, daß jeder Kreis Länge ≥ 4 hat gilt

$$\alpha_1(G) \leq 2\alpha_0(G) - 4.$$

Satz 5.40 Der vollständige paare Graph $K_{3,3}$ ist nicht planar.



Definition 5.41 Ein Graph H heißt **Unterteilung** eines Graphen G , wenn H aus G dadurch hervorgeht, wenn in einer oder in mehreren Kanten von G zusätzliche Knoten eingefügt werden.

Satz 5.42 (Satz von Kuratowski) Eine Graph G ist genau dann nicht planar, wenn er einen Teilgraphen enthält, der aus C_5 oder $K_{3,3}$ durch eventuelle Unterteilung (von Kanten) entsteht.

⁴³complete graph

⁴⁴bipartite graph

⁴⁵complete bipartite graph

5.7 Eulersche und Hamiltonsche Linien

Definition 5.43 Eine Kantenfolge in einem (gerichteten oder ungerichteten) Graphen G heißt **Eulersche Linie**, wenn sie jeden Knoten und jede Kante enthält, und zwar jede Kante genau einmal.

Bei einer **geschlossenen Eulerschen Linie** stimmen Anfangs- und Endknoten überein, bei einer **offenen Eulerschen Linie** sind sie verschieden.

Satz 5.44 Ein ungerichteter Graph G besitzt genau dann eine geschlossene Eulersche Linie, wenn G zusammenhängend ist und alle Knotengrade $d(v)$ ($v \in V(G)$) gerade sind.

Ein ungerichteter Graph G besitzt genau dann eine offene Eulersche Linie, wenn G zusammenhängend ist und mit der Ausnahme von zwei Knoten $v_1, v_2 \in V(G)$ (mit ungeradem Knotengrad) alle Knotengrade $d(v)$ ($v \in V(G) \setminus \{v_1, v_2\}$) gerade sind.

Satz 5.45 Ein gerichteter Graph G besitzt genau dann eine geschlossene Eulersche Linie, wenn G schwach zusammenhängend ist und alle Knoten $v \in V(G)$ Hin- und Weggrad gleich sind: $d^+(v) = d^-(v)$.

Ein gerichteter Graph G besitzt genau dann eine offene Eulersche Linie, wenn G schwach zusammenhängend ist und mit der Ausnahme von zwei Knoten $v_1, v_2 \in V(G)$, für die $d^+(v_1) = d^-(v_1) + 1$ und $d^+(v_2) = d^-(v_2) - 1$ gilt, bei allen Knoten $v \in V(G) \setminus \{v_1, v_2\}$ Hin- und Weggrad gleich sind: $d^+(v) = d^-(v)$.

Definition 5.46 Eine Kantenfolge in einem (gerichteten oder ungerichteten) Graphen G heißt **Hamiltonsche Linie**, wenn sie jeden Knoten (mit der möglichen Ausnahme, daß Anfangs- und Endpunkt übereinstimmen) genau einmal enthält.

Bei einer **geschlossenen Hamiltonschen Linie** stimmen Anfangs- und Endknoten überein, bei einer **offenen Hamiltonschen Linie** sind sie verschieden.

Im Gegensatz zu den Eulerschen Linien gibt es (bis jetzt noch) kein einfaches Kriterium für die Existenz von Hamiltonschen Linien.

5.8 Algorithmen für Netzwerke

Definition 5.47 Ein **Netzwerk** bzw. ein **bewerteter Graph**⁴⁶ ist ein gerichteter oder ungerichteter Graph $G = (V, E)$, wo jeder Kante $e \in E$ ein Wert (Gewicht, Länge, Kapazität) $w(e) \in \mathbf{R}$ zugeordnet wird, d.h. es gibt noch eine Funktion $w : E \rightarrow \mathbf{R}$.

Üblicherweise wird auch vorausgesetzt, daß für alle Kanten $w(e) \geq 0$ gilt. In manchen Anwendungen ist es auch sinnvoll, die Bewertung $w(e) = \infty$ zuzulassen.

⁴⁶labeled graph

Im folgenden sollen drei ganz unterschiedliche Algorithmen für Netzwerke beschrieben werden, einer zur Bestimmung eines minimalen spannenden Baumen, einen zur Ermittlung des kürzesten Weges zwischen zwei Knoten und einen zur Berechnung eines maximalen Flusses.

5.8.1 Minimales Gerüst: Kruskal-Algorithmus

Definition 5.48 Ein spannender Baum T eines schlichten zusammenhängenden Graphen G ist ein Baum mit $V(T) = V(G)$ und $E(T) \subseteq E(G)$, d.h. er enthält dieselben Knoten wie G und gewisse Kanten von G .

Ein **Gerüst** oder **spannender Wald** W eines schlichten Graphen G ist ein Wald mit $V(W) = V(G)$ und $E(W) \subseteq E(G)$ und denselben Zusammenhangskomponenten wie G , d.h. schränkt man W auf eine Zusammenhangskomponente K von G ein, so ist diese Einschränkung T ein spannender Baum von K .

Definition 5.49 Sei $G = (V, E)$ ein Netzwerk mit Bewertung $w : E \rightarrow \mathbf{R}_0^+$. Für jede Teilmenge $F \subseteq E$ von Kanten heißt

$$w(F) = \sum_{e \in F} w(e)$$

Gewicht⁴⁷ von F .

Definition 5.50 Ein **minimaler spannender Baum** T bzw. ein **minimales Gerüst** W eines Netzwerkes $G = (V, E)$ mit nichtnegativer Bewertung $w : E \rightarrow \mathbf{R}_0^+$ ist ein spannender Baum bzw. ein Gerüst mit kleinstmöglichem Gewicht $w(E(T))$ bzw. $w(E(W))$.

Ein **maximaler spannender Baum** bzw. ein **maximales Gerüst** ist ein spannender Baum bzw. ein Gerüst mit größtmöglichem Gewicht.

Mit Hilfe des **Kruskal-Algorithmus** gelingt es in sehr einfacher Weise ein minimales bzw. maximales Gerüst eines Netzwerkes $G = (V, E)$ mit nichtnegativem Gewicht $w : E \rightarrow \mathbf{R}_0^+$ zu bestimmen.

1. Man nummeriere die Kanten $E = \{e_1, e_2, \dots, e_m\}$ nach steigendem Gewicht:

$$w(e_1) \leq w(e_2) \leq \dots \leq w(e_m).$$

2. Setze $E' := \emptyset$.

```

for  $j := 1$  to  $m$  do
  if  $(V, E' \cup \{e_j\})$  kreisfrei then
     $E' := E' \cup \{e_j\}$ 
  end;
end;
```

⁴⁷ weight

$W = (V, E')$ ist nun ein minimales Gerüst von G .

Um ein maximales Gerüst zu erhalten, muß man in 1. die Kanten nach fallendem Gewicht ordnen.

Hat G n Knoten und k Zusammenhangskomponenten, so hat jedes Gerüst von G genau $n - k$ Kanten. Daher kann 2. abgebrochen werden, wenn E' bereits $n - k$ Kanten enthält.

Im allgemeinen gibt es eine sehr große Anzahl von Gerüsten in einem Graphen, und es ist außerordentlich schwer, sich einen Überblick über alle möglichen Gerüste zu verschaffen. (Allein ein minimales Gerüst ist mit Hilfe des Kruskal-Algorithmus sehr leicht zu bestimmen.) Immerhin gibt es eine Formel für die Anzahl der Gerüste.

Satz 5.51 (Matrix-Baum-Theorem von Kirchhoff) *Sei G ein ungerichteter, schlichter, zusammenhängender Graph mit Knotenmenge $V(G) = \{v_1, v_2, \dots, v_n\}$ und Adjazenzmatrix $A(G)$. Es bezeichne weiters $D(G)$ die Diagonalmatrix der Knotengrade $\text{diag}(d(v_1), d(v_2), \dots, d(v_n))$, so ist eine beliebige $(n - 1) \times (n - 1)$ -Unterdeterminante der Matrix*

$$D(G) - A(G)$$

die Anzahl der spannenden Bäume von G .

Ist G nicht zusammenhängend, so wendet man das Matrix-Baum-Theorem für jede Komponente an und multipliziert die Ergebnisse, um die Anzahl der Gerüste von G zu erhalten.

5.8.2 Kürzester Weg: Dijkstra-Algorithmus

Definition 5.52 *Sei $G = (V, E)$ ein Netzwerk mit Bewertung $w : E \rightarrow \mathbf{R}_0^+$. Die Länge einer Kantenfolge $e_1, e_2, \dots, e_k$ ist durch die Summe der Gewichte*

$$w(\{e_1, e_2, \dots, e_k\}) = \sum_{j=1}^k w(e_j)$$

*gegeben. Die **Distanz**⁴⁸ $d(v, w)$ zwischen zwei Knoten $v_1, v_2 \in V$ ist die kleinstmögliche Länge einer $KF(v_1, v_2)$. Gibt es keine $KF(v_1, v_2)$, so setzt man $d(v_1, v_2) = \infty$.*

Mit Hilfe des **Dijkstra-Algorithmus** wird von einem Knoten $v_0 \in V$ eines Netzwerkes mit nichtnegativem Gewicht die Distanz $d(v_0, v)$ für alle Knoten $v \in V$ bestimmt.

1. Setze $W := \emptyset$, $U := V$, $l(v_0) := 0$, $l(v) := \infty$ für alle $v \in V \setminus \{v_0\}$ und $p(v) = *$ für alle $v \in V$.
2. Man bestimme $m = \min_{v \in U} l(v)$, wähle einen Knoten $z \in U$ mit $l(z) = m$.

⁴⁸distance

3. Man setze $W := W \cup \{z\}$, $U := U \setminus \{z\}$ und für $v \in U$, die $l(v) > l(z) + w(z, v)$ erfüllen, setze man $p(v) := z$ und

$$l(v) := l(z) + w(z, v).$$

4. Ist $W = V \rightarrow$ **ENDE**.

Ist $l(v) = \infty$ für alle $v \in U \rightarrow$ **ENDE**.

Gehe zu 2.

Die Menge $W \subseteq V$ umfaßt in jedem Zeitpunkt des Algorithmus genau jene $v \in V$, für die schon bekannt ist, daß $l(v) = d(v_0, v)$ ist. Für $v \in U$ ist hingegen $l(v)$ die minimale Länge einer Kantenfolge $KF(v_0, v)$, die mit Ausnahme von v nur Knoten aus W enthält. $p(v)$ ist jeweils der Vorgängerknoten von v auf einer minimalen Kantenfolge $KF(v_0, v)$.

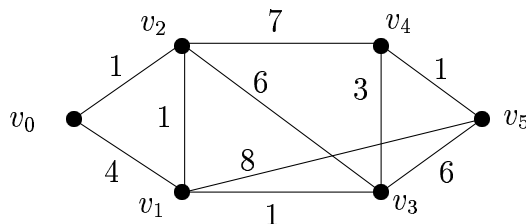
Aus diesen Bedingungen ist klar, daß der in 2. ausgewählte Knoten $z \in U$ in W aufgenommen werden kann. Für $v \in U \setminus \{z\}$ wird in 3. $l(v)$ neu berechnet, da W um z erweitert wurde.

Endet der Algorithmus nicht mit $W = V$, sondern mit der Abbruchbedingung $l(v) = \infty$ für alle $v \in U$, so ist G nicht zusammenhängend, und W ist die Zusammenhangskomponente, die v_0 enthält.

Sucht man übrigens zu einem bestimmten Punkt $v_1 \in V$ die Distanz $d(v_0, v_1)$, so kann der Algorithmus auch dann abgebrochen werden, wenn v_1 in W aufgenommen wird.

Für $v \in W$ kann durch $v, p(v), p(p(v)), \dots, v_0$ eine $KF(v_0, v)$ von kleinstmöglicher Länge rückverfolgt werden. Möchte man alle Kantenfolgen kleinstmöglicher Länge finden, so muß in 3. für alle Knoten v mit $l(v) = l(z) + w(z, v)$ der Knoten z als *alternatives* $p(y)$ vermerkt werden.

Beispiel 5.53 Im folgenden Graphen sollen alle Distanzen $d(v_0, v)$ bestimmt werden.



1. Durchlauf

$$m = 0$$

$$z = v_0$$

$$W = \{v_0\}, U = \{v_1, v_2, v_3, v_4, v_5\}$$

$$l(v_1) = 4, l(v_2) = 1, l(v_3) = l(v_4) = l(v_5) = \infty$$

$$p(v_1) = v_0, p(v_2) = v_0, p(v_3) = p(v_4) = p(v_5) = *$$

2. Durchlauf

$$\begin{aligned}
m &= 1 \\
z &= v_2 \\
W &= \{v_0, v_2\}, U = \{v_1, v_3, v_4, v_5\} \\
l(v_1) &= 2, l(v_3) = 7, l(v_4) = 8, l(v_5) = \infty \\
p(v_1) &= v_2, p(v_3) = v_2, p(v_4) = v_2, p(v_5) = *
\end{aligned}$$

3. Durchlauf

$$\begin{aligned}
m &= 2 \\
z &= v_1 \\
W &= \{v_0, v_1, v_2\}, U = \{v_3, v_4, v_5\} \\
l(v_3) &= 3, l(v_4) = 8, l(v_5) = 10 \\
p(v_3) &= v_1, p(v_4) = v_2, p(v_5) = v_1
\end{aligned}$$

4. Durchlauf

$$\begin{aligned}
m &= 3 \\
z &= v_3 \\
W &= \{v_0, v_1, v_2, v_3\}, U = \{v_4, v_5\} \\
l(v_4) &= 6, l(v_5) = 9 \\
p(v_4) &= v_3, p(v_5) = v_3
\end{aligned}$$

5. Durchlauf

$$\begin{aligned}
m &= 6 \\
z &= v_4 \\
W &= \{v_0, v_1, v_2, v_3, v_4\}, U = \{v_5\} \\
l(v_5) &= 7 \\
p(v_5) &= v_4
\end{aligned}$$

6. Durchlauf

$$\begin{aligned}
m &= 7 \\
z &= v_3 \\
W &= \{v_0, v_1, v_2, v_3, v_4, v_5\}, U = \emptyset \\
\mathbf{ENDE}
\end{aligned}$$

Man erhält damit $d(v_0, v_1) = 2$, $d(v_0, v_2) = 1$, $d(v_0, v_3) = 3$, $d(v_0, v_4) = 6$ und $d(v_0, v_5) = 7$.

5.8.3 Maximaler Fluß: Ford-Fulkerson-Algorithmus

Im folgenden sei $G = (V, E)$ ein gerichteter, bewerteter Graph, auf dem zwei Knoten $s, t \in V$ ausgezeichnet sind. s habe Hingrad $d^-(s) = 0$ und wird als **Quelle**⁴⁹ bezeichnet und t habe Weggrad $d^+(t) = 0$ und wird als **Senke**⁵⁰ bezeichnet. Weiters seien

$$\Gamma^+(x) = \{y \in V \mid \langle x, y \rangle \in E\}$$

die **Nachfolger**⁵¹ von $x \in V$ und

$$\Gamma^-(x) = \{y \in V \mid \langle y, x \rangle \in E\}$$

die **Vorgänger**⁵² von $x \in G$.

Definition 5.54 Sei $G = (V, E)$ ein gerichteter, bewerteter Graph mit Quelle s , Senke t und Bewertung $w : E \rightarrow \mathbf{R}_0^+$. Eine Funktion $\phi : E \rightarrow \mathbf{R}$ heißt **Fluß**⁵³ auf G , wenn

(i) für alle Kanten $e \in E$

$$0 \leq \phi(e) \leq w(e)$$

gilt und

(ii) für alle Knoten $x \in V \setminus \{s, t\}$ die Flußbedingung

$$\sum_{y \in \Gamma^-(x)} \phi(\langle y, x \rangle) = \sum_{y \in \Gamma^+(x)} \phi(\langle x, y \rangle)$$

erfüllt ist, d.h. es fließt genauso viel ab wie zu.

Die **Größe eines Flusses** ϕ ist durch

$$v(\phi) = \sum_{y \in \Gamma^+(s)} \phi(\langle s, y \rangle)$$

gegeben.

Die Größe eines Flusses ϕ kann auch auf andere Art bestimmt werden. Wegen der Flußbedingung (ii) gilt z.B.

$$v(\phi) = \sum_{x \in \Gamma^-(t)} \phi(\langle x, t \rangle),$$

⁴⁹spring

⁵⁰sink

⁵¹successor

⁵²predecessor

⁵³flow

was aus der Quelle s herausfließt, muß in die Senke t wieder zurückfließen. Die Größe jedes Flusses ϕ ist daher sicherlich durch

$$v(\phi) \leq \min \left(\sum_{y \in \Gamma^+(s)} w(\langle s, y \rangle), \sum_{y \in \Gamma^+(s)} w(\langle s, y \rangle) \right)$$

beschränkt. Um noch bessere Abschätzungen zu erhalten, benötigt man den Begriff eines Schnittes.

Definition 5.55 Sei $V(G) = V_1 \cup V_2$ eine Zerlegung der Knotenmenge $V(G)$ in zwei nichtleere, disjunkte Teilmengen, d.h. $V_1 \neq \emptyset$, $V_2 \neq \emptyset$, $V_1 \cup V_2 = V(G)$ und $V_1 \cap V_2 = \emptyset$. Der **Schnitt**⁵⁴ $S = (V_1, V_2)$ auf G besteht dann aus allen Kanten $e \in E(G)$, die Knoten aus V_1 mit Knoten aus V_2 verbinden oder umgekehrt, d.h.

$$S = (V_1, V_2) = (V_1 \rightarrow V_2) \cup (V_1 \leftarrow V_2),$$

wobei

$$(V_1 \rightarrow V_2) = \{\langle x, y \rangle \in E(G) \mid x \in V_1, y \in V_2\}$$

und

$$(V_1 \leftarrow V_2) = \{\langle y, x \rangle \in E(G) \mid x \in V_1, y \in V_2\}$$

bezeichnen.

Definition 5.56 Sei $G = (V, E)$ ein gerichteter, bewerteter Graph und $S = (V_1, V_2)$ ein Schnitt auf G , der die Quelle s und die Senke t trennt, d.h. $s \in V_1$ und $t \in V_2$. Dann bezeichnet

$$c(S) = \sum_{e \in (V_1 \rightarrow V_2)} w(e)$$

die **Kapazität**⁵⁵ des Schnittes S .

Satz 5.57 Sei $G = (V, E)$ ein gerichteter, bewerteter Graph. Dann gilt für jeden Fluß ϕ auf G

$$v(\phi) \leq \min_{S=(V_1, V_2), s \in V_1, t \in V_2} c(S).$$

Die Größe eines Flusses ϕ kann also nie größer sein als die Kapazität irgendeines Schnittes $S = (V_1, V_2)$, der s und t trennt. Interessant ist, daß es auch immer einen Fluß gibt, dessen Größe gleich der kleinsten Kapazität eines solchen Schnittes ist.

Satz 5.58 (Ford-Fulkerson) Sei $G = (V, E)$ ein gerichteter, bewerteter Graph. Dann gilt

$$\max_{\phi \text{ Fluß auf } G} v(\phi) = \min_{S=(V_1, V_2), s \in V_1, t \in V_2} c(S).$$

⁵⁴cut

⁵⁵capacity

Der Satz von Ford und Fulkerson ist in dem Sinn auch konstruktiv, daß man bei ganzwertigen (bzw. rationalwertigen) Bewertungen mit Hilfe des **Algorithmus von Ford und Fulkerson** einen maximalen Fluß immer auch berechnen kann.

Im folgenden sei nun $G = (V, E)$ ein gerichteter, bewerteter Graph mit Quelle s und Senke t , so daß die Gewichtsfunktion $w : E \rightarrow \mathbb{N}_0$ nur nichtnegative ganzzahlige Werte annimmt.

Ähnlich wie beim Dijkstra-Algorithmus wird jeder Knoten $x \in V$ mit einem *signierten Vorgänger* $p(x)$ und einem *Wert* $\delta(x)$ markiert.

1. Setze $\phi(e) := 0$ für alle $e \in E$.
2. Setze $p(s) := +s, \delta(s) := \infty, V_1 := \{s\}$ und $V_2 := V \setminus \{s\}$.
3. (a) Für alle $x \in V_1$ und $y \in \Gamma^+(x) \cap V_2$ führe folgendes aus:

```

    if  $\phi(\langle x, y \rangle) < w(\langle x, y \rangle)$  then
       $p(y) := +x;$ 
       $\delta(y) := \min(\delta(x), w(\langle x, y \rangle) - \phi(\langle x, y \rangle));$ 
       $V_1 := V_1 \cup \{y\}; V_2 := V_2 \setminus \{y\};$ 
    end;

```

- (b) Für alle $x \in V_1$ und $y \in \Gamma^-(x) \cap V_2$ führe folgendes aus:

```

    if  $\Phi(\langle y, x \rangle) > 0$  then
       $p(y) := -x;$ 
       $\delta(y) := \min(\delta(x), \Phi(\langle y, x \rangle));$ 
       $V_1 := V_1 \cup \{y\}; V_2 := V_2 \setminus \{y\};$ 
    end;

```

4. Falls im letzten Durchlauf von 3. V_1 vergrößert wurde oder $t \in V_2$ ist, gehe zu 3.
5. (a) Ist $t \in V_2 \rightarrow$ **ENDE**.

```

    (b)    $x := t;$ 
          while  $x \neq s$  do
            if  $p(x) = +z$  then
               $\phi(\langle z, x \rangle) := \phi(\langle z, x \rangle) + \delta(t);$ 
               $x := z;$ 
            end;
            if  $p(x) = -z$  then
               $\phi(\langle x, z \rangle) := \phi(\langle x, z \rangle) - \delta(t);$ 
               $x := z;$ 
            end;
          end;

```

6. Gehe zu 2.

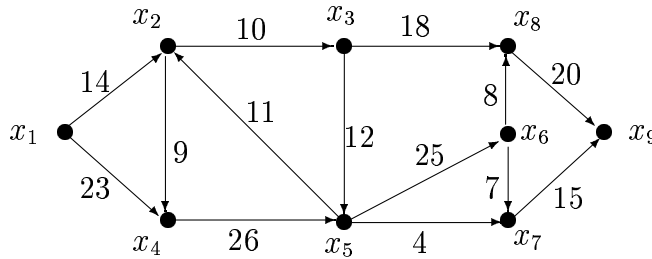
In 1. wird der Fluß zu 0 gesetzt, da der Nullfluß sicherlich ein gültiger Fluß ist. Man könnte hier auch mit einem anderen bekannten Fluß beginnen.

In 2. werden der Markierungen $p(s)$ und $\delta(s)$ initialisiert, und in 3. wird versucht, möglichst alle Knoten zu markieren. Dabei ist $\delta(x)$ jeweils den maximalen Wert, um den man den Fluß entlang eines Weges $W(s, x)$ erhöhen könnte, ohne die Flußbeschränkung $0 \leq \phi \leq w$ zu verletzen. (Wenn man entlang eines Weges den Fluß um einen gewissen Wert erhöht, bleibt die Flußbedingung sicherlich richtig und muß daher an dieser Stelle des Algorithmus nicht berücksichtigt werden.) $p(x) = \pm z$ ist der jeweilig Vorgänger von x auf diesem Weg. $p(x) = +z$, wenn $e = \langle z, x \rangle \in E$, d.h. e ist gleichorientiert wie der Weg $W(s, x)$. Um den Fluß auf $W(s, x)$ zu erhöhen, muß $\phi(e)$ erhöht werden. $p(x) = -z$, wenn $e = \langle x, z \rangle \in E$, d.h. e ist zum Weg $W(s, x)$ gegenorientiert. Um den Fluß auf $W(s, x)$ zu erhöhen, muß $\phi(e)$ daher erniedrigt werden. Alle markierten Knoten werden in V_1 aufgenommen.

Kann man in 3. auch die Senke t markieren, so gibt es einen Weg von s nach t , auf dem der Fluß um den Wert $\delta(t)$ erhöht werden kann. Dies wird in 5. durchgeführt.

2.–6. wird so lange wiederholt, bis es in 3. nicht mehr gelingt, auch die Senke t zu markieren. In diesem Fall ist $S = (V_1, V_2)$, der s und t trennt, und, wie man sofort sieht, $c(S) = v(\phi)$. (3. bricht ab, wenn alle Kanten $e \in (V_1 \rightarrow V_2)$ saturiert sind und wenn auf allen Kanten $e \in (V_1 \leftarrow V_2)$ kein Fluß fließt.) Daher hat man einen maximalen Fluß gefunden.

Beispiel 5.59 Es sei folgendes Netzwerk mit Quelle $s = x_1$ und Senke $t = x_9$ vorgegeben:



Im 1. Durchlauf erhält man folgende Markierungen:

$$\begin{aligned}
 p(x_1) &= +x_1, \delta(x_1) = \infty, \\
 p(x_2) &= +x_1, \delta(x_2) = 14, \\
 p(x_3) &= +x_2, \delta(x_3) = 10, \\
 p(x_4) &= +x_1, \delta(x_4) = 23, \\
 p(x_5) &= +x_3, \delta(x_5) = 10, \\
 p(x_6) &= +x_5, \delta(x_6) = 10, \\
 p(x_7) &= +x_5, \delta(x_7) = 4, \\
 p(x_8) &= +x_3, \delta(x_8) = 10, \\
 p(x_9) &= +x_7, \delta(x_9) = 4.
 \end{aligned}$$

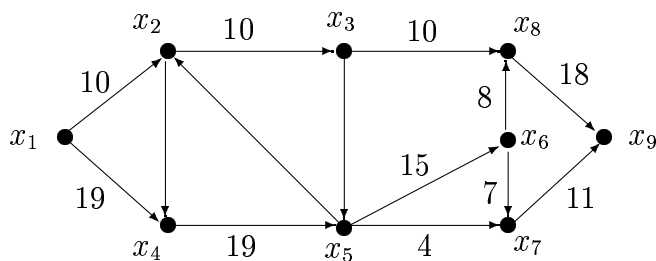
Es gibt also einen Weg von x_1 nach x_9 , auf dem der Fluß um $\delta(x_9) = 4$ erhöht werden kann. Mit Hilfe der Funktion p kann dieser Weg rückverfolgt werden ($x_9 \leftarrow x_7 \leftarrow x_5 \leftarrow x_3 \leftarrow x_2 \leftarrow x_1$). Der Fluß wird daher auf den Kanten $\langle x_1, x_2 \rangle$, $\langle x_2, x_3 \rangle$, $\langle x_3, x_5 \rangle$, $\langle x_5, x_7 \rangle$ und $\langle x_7, x_9 \rangle$ von 0 auf 4 erhöht werden.

Im zweiten Durchlauf kann man den Fluß auf den Kanten $\langle x_1, x_2 \rangle$, $\langle x_2, x_3 \rangle$, $\langle x_3, x_5 \rangle$, $\langle x_5, x_6 \rangle$, $\langle x_6, x_7 \rangle$ und $\langle x_7, x_9 \rangle$ um 6 erhöhen.

Im dritten Durchlauf erhöht man den Fluß auf den Kanten $\langle x_1, x_4 \rangle$, $\langle x_4, x_5 \rangle$, $\langle x_5, x_6 \rangle$, $\langle x_6, x_7 \rangle$ und $\langle x_7, x_9 \rangle$ um 1.

Im nächsten Durchlauf wird der Fluß auf den Kanten $\langle x_1, x_4 \rangle$, $\langle x_4, x_5 \rangle$, $\langle x_3, x_8 \rangle$ und $\langle x_8, x_9 \rangle$ um 10 erhöht und auf der Kante $\langle x_3, x_5 \rangle$ um 10 verkleinert.

Schließlich vergrößert man noch im fünften Durchlauf den Fluß auf den Kanten $\langle x_1, x_4 \rangle$, $\langle x_4, x_6 \rangle$, $\langle x_5, x_6 \rangle$, $\langle x_6, x_8 \rangle$ und $\langle x_8, x_9 \rangle$ um 8 und erhält folgendes Bild, wobei die Eintragungen an den Kanten nun den Fluß representieren.



Im letzten Durchlauf gelingt es nicht mehr, alle Knoten zu markieren. Man erhält den Schnitt $S = (V_1, V_2)$ mit $V_1 = \{x_1, x_2, x_4, x_5, x_6\}$ und $V_2 = \{x_3, x_7, x_8, x_9\}$. Es ist leicht zu verifizieren, daß die Kapazität dieses Schnittes 29 mit dem Wert des momentanen Flusses übereinstimmt.

Kapitel 6

Algebraische Strukturen

6.1 Gruppen

Definition 6.1 Sei A eine nichtleere Menge. Eine **binäre Operation**¹ $\circ$ auf A ist eine Abbildung $A \times A \rightarrow A$, d.h. je zwei Elementen $a, b \in A$ wird ein Element $a \circ b$ zugeordnet.

Ein Paar $\langle A, \circ \rangle$ heißt **algebraische Struktur**², wenn A eine nichtleere Menge ist und $\circ$ eine binäre Operation auf A ist.

Es tritt öfters die Situation ein, daß eine binäre Operation $\circ$ auf einer Menge A nur für Elemente a, b in einer Teilmenge $B \subseteq A$ betrachtet wird. Dafür muß aber gewährleistet sein, daß für alle $a, b \in B$ das Element $a \circ b$ wieder in B liegt. In diesem Fall spricht man von **Abgeschlossenheit** der Operation $\circ$ auf B . Um dieser Situation Rechnung zu tragen, wird die Eigenschaft *Abgeschlossenheit* in der folgenden Liste mitaufgenommen, obwohl dies (wegen der Definition einer binären Operation) nicht unbedingt notwendig wäre.

Definition 6.2 Sei $\langle A, \circ \rangle$ eine algebraische Struktur. Dabei werden folgende Gesetzmäßigkeiten von $\langle A, \circ \rangle$ definiert.

(1) **Abgeschlossenheit**: Für alle $a, b \in A$ ist auch $a \circ b \in A$, (d.h. $\circ$ ist binäre Operation auf A).

(2) **Assoziativgesetz**³: Für alle $a, b, c \in A$ gilt

$$(a \circ b) \circ c = a \circ (b \circ c).$$

¹binary operation

²algebraic structure

³associative law

- (3) **Existenz eines neutralen Elements**⁴: Es gibt ein $e \in A$ mit der Eigenschaft, daß für alle $a \in A$

$$e \circ a = a \circ e = a$$

gilt.

- (4) **Existenz inverser Elemente**⁵: Für alle $a \in A$ gibt es $a' \in A$ mit

$$a \circ a' = a' \circ a = e,$$

wobei e das neutrale Element aus (3) bezeichnet.

- (5) **Kommutativgesetz**⁶: Für alle $a, b \in A$ gilt

$$a \circ b = b \circ a.$$

Benützt man für das binäre Operationssymbol das *Malzeichen* $\cdot$, so schreibt man für das inverse Element von a auch a^{-1} , verwendet man hingegen das *Pluszeichen* $+$, so bezeichnet man das inverse Element von a durch $-a$.

Satz 6.3 *In einer algebraischen Struktur $\langle A, \circ \rangle$ gibt es höchstens ein neutrales Element, und zu jedem $a \in A$ gibt es höchstens ein inverses Element.*

Es wird daher im folgenden nur mehr vom neutralen Element bzw. von inversen Element gesprochen werden, sofern diese existieren.

Definition 6.4 *Eine algebraische Struktur $\langle A, \circ \rangle$ heißt*

- **Gruppoid**, wenn sie (1) erfüllt,
- **Halbgruppe**⁷, wenn sie (1) und (2) erfüllt,
- **Monoid**⁸, wenn sie (1), (2) und (3) erfüllt, und
- **Gruppe**⁹, wenn sie (1), (2), (3) und (4) erfüllt.

*Erfüllt eine der Strukturen Gruppoid, Halbgruppe, Monoid bzw. Gruppe auch (5), so heißen sie auch **kommutative(s)** Gruppoid, Halbgruppe, Monoid bzw. Gruppe.*

Kommutative Gruppen werden auch als **abelsche Gruppen**¹⁰ (im Andenken an den Mathematiker Niels Henrik Abel) bezeichnet

⁴neutral element

⁵inverse

⁶commutative law

⁷semigroup

⁸monoid

⁹group

¹⁰abelian group

Beispiel 6.5 $A = \mathbf{N}$ mit $a \circ b = a^b$ ist nur ein Gruppoid.

Beispiel 6.6 $\langle \mathbf{N}, + \rangle$ ist eine Halbgruppe, $\langle \mathbf{N}_0, + \rangle$ und $\langle \mathbf{N}, \cdot \rangle$ sogar Monoide, aber keine Gruppen.

Beispiel 6.7 Sei Σ eine endliche Menge, das *Alphabet* und bezeichne Σ^* die Menge aller endlichen *Wörter* über Σ , das sind alle endlichen Folgen $x_1 x_2 \dots x_k$ mit $x_j \in \Sigma$ ($1 \leq j \leq k$) ergänzt um das *leere Wort* ϵ . Sind $w_1 = x_{11} x_{12} \dots x_{1k}$ und $w_2 = x_{21} x_{22} \dots x_{2l}$ zwei Wörter in Σ^* so definiert man

$$w_1 \circ w_2 = x_{11} x_{12} \dots x_{1k} x_{21} x_{22} \dots x_{2l} \in \Sigma^*.$$

$\langle \Sigma^*, \circ \rangle$ ist damit ein Monoid mit neutralem Element ϵ . Man bezeichnet Σ^* auch als **freies Monoid**¹¹ über dem Alphabet Σ .

Beispiel 6.8 $\langle \mathbf{Z}, + \rangle$, $\langle \mathbf{Q}, + \rangle$, $\langle \mathbf{Q} \setminus \{0\}, \cdot \rangle$, $\langle \mathbf{R}, + \rangle$, $\langle \mathbf{R} \setminus \{0\}, \cdot \rangle$ etc. sind abelsche Gruppen.

Beispiel 6.9 Alle $n \times n$ -Matrizen $L_{n,n}(\mathbf{R})$ mit Koeffizienten aus $\mathbf{R}$ bilden mit der Matrizenaddition eine Gruppe. Außerdem bilden jede $n \times n$ -Matrizen A mit $\det(A) \neq 0$ bezüglich der Matrizenmultiplikation eine Gruppe. Diese ist für $n \geq 2$ nicht kommutativ.

Beispiel 6.10 Sei M eine beliebige Menge. Dann bildet $\langle \mathbf{P}(M), \Delta \rangle$, d.h. die Teilmengen von M mit der symmetrischen Mengendifferenz, eine Gruppe. Das neutrale Element ist $\emptyset$ und jedes Element ist selbstinvers.

Beispiel 6.11 Die Menge der Permutationen $\mathbf{S}_n$ der Zahlen $\{1, 2, \dots, n\}$ mit der Hintereinanderausführung bilden die sogenannte **symmetrische Gruppe**.

Beispiel 6.12 Sei G ein Graph. Ein Isomorphismus $\phi : G \rightarrow G$ heißt **Automorphismus**. Die Menge aller Automorphismen des Graphen G bildet mit der Hintereinanderausführung eine Gruppe, die sogenannte **Automorphismengruppe** von G .

Beispiel 6.13 Die **Symmetriegruppe** eines gleichseitigen Dreiecks besteht aus allen Isometrien der Ebene, die ein gleichseitiges Dreieck auf sich selbst abbilden. Da ein Dreieck durch seine Eckpunkte eindeutig gegeben ist, reicht es aus, die Auswirkung solcher Isometrien auf die Eckpunkte zu betrachten. Es entstehen gewisse Permutationen der Eckpunkte $\{1, 2, 3\}$. Bei den Drehungen um 0° , 120° und 240° werden die Eckpunkte zyklisch vertauscht, und bei den Spiegelungen an den drei Höhen werden jeweils zwei Eckpunkte miteinander vertauscht. Insgesamt erhält man also sechs verschiedene Symmetrien, die bezüglich Hintereinanderausführung eine Gruppe bilden. In diesem speziellen Fall eines gleichseitigen Dreiecks ist die Symmetriegruppe nichts anderes als die symmetrische Gruppe auf den drei Eckpunkten.

$\circ$	e	$\circ$	e	a	$\circ$	e	a	b
e	e	e	e	a	e	e	a	b
		a	a	e	a	a	b	e
					b	b	e	a

$\circ$	e	a	b	c	$\circ$	e	a	b	c	$\circ$	e	a	b	c	d
e	e	a	b	c	e	e	a	b	c	e	e	a	b	c	d
a	a	b	c	e	a	a	e	c	b	a	a	b	c	d	e
b	b	c	e	a	b	b	c	e	a	b	b	c	d	e	a
c	c	e	a	b	c	c	b	a	e	c	c	d	e	a	b
						d	d	e	a	b	d	e	a	b	c

Beispiel 6.14 Kleine algebraische Strukturen kann man auch durch sogenannte **Operationstabellen** definieren. Um dies zu demonstrieren, werden alle Möglichkeiten von *kleinen Gruppen* mit ≤ 5 Elementen aufgelistet. (e bezeichnet immer das neutrale Element.)

Definition 6.15 Eine (nichtleere) Teilmenge $U \subseteq G$ einer Gruppe $\langle G, \circ \rangle$ heißt **Untergruppe**¹² von G , wenn $\langle U, \circ \rangle$ selbst ein Gruppe ist, i.Z. $\langle U, \circ \rangle \leq \langle G, \circ \rangle$ oder nur $U \leq G$.

Satz 6.16 Sei $\langle G, \circ \rangle$ Gruppe und U nichtleere Teilmenge von G . Dann sind die folgenden drei Bedingungen äquivalent:

- (i) U ist Untergruppe von G ($U \leq G$).
- (ii) Mit $a, b \in U$ sind auch $a \circ b \in U$ und $a' \in U$.
- (iii) Mit $a, b \in U$ ist auch $a \circ b' \in U$.

Beispiel 6.17 Für $m \in \mathbf{N}$ bilden die Mengen $m\mathbf{Z} = \{0, \pm m, \pm 2m, \pm 3m, \dots\}$ Untergruppen bezüglich der Addition.

Definition 6.18 Sei $\langle G, \circ \rangle$ Gruppe, U Untergruppe von G und $a \in G$. Dann heißt

$$a \circ U = \{a \circ u \mid u \in U\}$$

Linksnebenklasse¹³ von U in G und

$$U \circ a = \{u \circ a \mid u \in U\}$$

Rechtsnebenklasse¹⁴ von U in G .

¹¹free monoid

¹²subgroup

¹³left coset

¹⁴right coset

Lemma 6.19 Sei $\langle G, \circ \rangle$ Gruppe und $U \leq G$. Dann gilt

$$b \in a \circ U \iff b \circ U = a \circ U.$$

Satz 6.20 Sei $\langle G, \circ \rangle$ Gruppe und $U \leq G$. Dann bildet die Menge der Linksnebenklassen $a \circ U$ ($a \in G$) eine Partition von G . Die Relation $a \sim b \iff a \circ U = b \circ U$ ist die entsprechende Äquivalenzrelation.

Eine entsprechende Aussage gilt für die Rechtsnebenklassen $U \circ a$.

Satz 6.21 Sei $\langle G, \circ \rangle$ Gruppe und $U \leq G$. Dann sind alle Links- und Rechtsnebenklassen gleichmächtig, d.h. für alle $a \in G$ gilt $|a \circ U| = |U \circ a| = |U|$

Korollar 6.22 Sei $\langle G, \circ \rangle$ endliche Gruppe und $U \leq G$. Dann stimmt die Anzahl der Linksnebenklassen von U in G mit der Anzahl der Rechtsnebenklassen überein. Diese Anzahl ist durch $|G|/|U|$ gegeben.

Definition 6.23 Sei $\langle G, \circ \rangle$ endliche Gruppe und $U \leq G$. Die Anzahl der Links- bzw. Rechtsnebenklassen von U wird als **Index**¹⁵ $|G : U| = |G|/|U|$ von G nach U bezeichnet.

Die Anzahl der Elemente $|G|$ einer Gruppe wird als **Ordnung**¹⁶ von G bzw. als **Gruppenordnung** bezeichnet.

Satz 6.24 (Satz Lagrange) Ist $\langle G, \circ \rangle$ endliche Gruppe so ist die Ordnung $|U|$ einer Untergruppe $U \leq G$ stets Teiler der Gruppenordnung $|G|$.

Definition 6.25 Sei $\langle G, \circ \rangle$ Gruppe mit neutralem Element e . Für $a \in G$ werden die **Potenzen**¹⁷ a^n von a mit $n \in \mathbf{Z}$ folgendermaßen definiert:

$$a^n = \begin{cases} e & \text{für } n = 0, \\ a & \text{für } n = 1, \\ a \circ a^{n-1} & \text{rekursiv für } n > 1, \\ (a')^{-n} & \text{für } n < 0. \end{cases}$$

Ist das Operationssymbol $+$, so schreibt man statt a^n auch na , z.B. $3a$ für $a + a + a$.

Lemma 6.26 Sei $\langle G, \circ \rangle$ Gruppe und $a \in G$. Dann gilt für alle $n, m \in \mathbf{Z}$

$$a^{n+m} = a^n \circ a^m.$$

Weiters sind entweder alle Potenzen a^n ($n \in \mathbf{Z}$) voneinander verschieden oder es gibt ein $n \in \mathbf{N}$ mit $a^n = e$.

¹⁵index

¹⁶order

¹⁷power

Definition 6.27 Sei $\langle G, \circ \rangle$ Gruppe und $a \in G$. Sind alle Potenzen a^n ($n \in \mathbf{Z}$) voneinander verschieden, so hat a **unendliche Ordnung**¹⁸ $\text{ord}_G(a) = \infty$. Andernfalls bezeichnet man

$$\text{ord}_G(a) = \min\{n \in \mathbf{N} \mid a^n = e\}$$

als **Ordnung**¹⁹ von a . a hat dann **endliche Ordnung**²⁰.

Hat $a \in G$ unendliche Ordnung, so ist

$$\langle a \rangle = \{a^n \mid n \in \mathbf{Z}\}$$

die von a erzeugte Untergruppe, bei endlicher Ordnung $\text{ord}_G(a)$ ist

$$\langle a \rangle = \{a^n \mid 0 \leq n < \text{ord}_G(a)\}$$

die von a erzeugt Untergruppe.

Man beachte, daß $\langle a \rangle$ in jedem Fall eine Untergruppe von G bildet $|\langle a \rangle| = \text{ord}_G(a)$ gilt.

Satz 6.28 Sei $\langle G, \circ \rangle$ endliche Gruppe und $a \in G$. Dann hat a endliche Ordnung und es $\text{ord}_G(a)$ ist ein Teiler der Gruppenordnung $|G|$.

Korollar 6.29 Für jedes Element $a \in G$ einer endlichen Gruppe $\langle G, \circ \rangle$ gilt $a^{|G|} = e$.

Satz 6.30 Ist H_i ($i \in I$) ein System von Untergruppen einer Gruppe G , so ist

$$H = \bigcap_{i \in I} H_i$$

wieder eine Untergruppe von G .

Definition 6.31 Sei G eine Gruppe und $K \subseteq G$ eine nichtleere Teilmenge von G . Die von K **erzeugte Untergruppe**²¹ $\langle K \rangle$ ist der Durchschnitt aller Untergruppen $H \leq G$, die K enthalten.

$\langle K \rangle$ ist wegen Satz 6.30 immer eine Untergruppe von G . Man beachte weiters, daß bei einelementigen Mengen $K = \{a\}$ die Definitionen 6.27 und 6.27 übereinstimmen.

Sind G_1, G_2 zwei Gruppen so kann auch das kartesische Produkt $G_1 \times G_2$ in natürlicher Weise zu einer Gruppe gemacht werden.

Definition 6.32 Seien $\langle G_1, \circ \rangle$ und $\langle G_2, \star \rangle$ zwei Gruppen. Dann ist $\langle G_1 \times G_2, \cdot \rangle$ mit der Operation

$$(a_1, b_1) \cdot (a_2, b_2) = (a_1 \circ a_2, b_1 \star b_2) \quad (a_1, a_2 \in G_1, b_1, b_2 \in G_2)$$

das **Produkt**²² der Gruppen G_1, G_2 .

¹⁸infinite order

¹⁹order

²⁰finite order

²¹generated subgroup

²²product of groups

Es ist leicht nachzurechnen, daß $G_1 \times G_2$ tatsächlich die Gruppenaxiome (1)–(4) erfüllt.

Definition 6.33 Eine Untergruppe H einer Gruppe G heißt **Normalteiler**²³, i.Z. $H \trianglelefteq G$, wenn die Links- und Rechtsnebenklassen übereinstimmen.

Offensichtlich ist jede Untergruppe H einer kommutativen Gruppe G ein Normalteiler.

Weiters ist jede Untergruppe H mit Index $|G : H| = 2$ Normalteiler, da es in diesem Fall nur zwei Links- bzw. Rechtsnebenklassen gibt. Die eine ist $e \circ H = H \circ e = H$ und die andere $G \setminus H$.

Satz 6.34 Für eine Untergruppe H einer Gruppe G sind folgende drei Eigenschaften äquivalent:

- (i) $H \trianglelefteq G$.
- (ii) $a \circ H = H \circ a$ für alle $a \in G$.
- (iii) $a \circ H \circ a' \subseteq H$ für alle $a \in G$.

Lemma 6.35 Sei H Normalteiler einer Gruppe G . Dann folgt aus $a_1 \circ H = a_2 \circ H$ und $b_1 \circ H = b_2 \circ H$ auch $(a_1 \circ b_1) \circ H = (a_2 \circ b_2) \circ H$.

Mit Hilfe dieser Eigenschaft von Normalteilern kann auch auf der Menge der Nebenklassen eine Gruppenoperation definiert werden.

Definition 6.36 Sei H Normalteiler einer Gruppe G und bezeichne G/H die Menge der Nebenklassen von G nach H . Dann wird durch die Operation

$$(a \circ H) \circ (b \circ H) := (a \circ b) \circ H$$

eine Gruppenoperation auf G/H definiert. Die Gruppe $\langle G/H, \circ \rangle$ heißt **Faktorgruppe**²⁴ von G nach H .

Es ist bei Faktorgruppen G/H üblich, dasselbe Operationszeichen (hier $\circ$) zu verwenden wie bei der ursprünglichen Gruppe G , da bei der Deutung als *Komplexprodukt*

$$(a \circ H) \circ (b \circ H) = \{c \circ d \mid c \in a \circ H, d \in b \circ H\}$$

tatsächlich (wegen der Normalteilereigenschaft)

$$\begin{aligned} (a \circ H) \circ (b \circ H) &= (H \circ a) \circ (b \circ H) \\ &= (H \circ (a \circ b)) \circ H \\ &= (a \circ b) \circ (H \circ H) \\ &= (a \circ b) \circ H \end{aligned}$$

²³normal subgroup

²⁴factor group

dasselbe Ergebnis erhalten wird. Die Gruppenstruktur von G/H ist daher natürlich. Ist H kein Normalteiler von G , so ist das Komplexprodukt von $a \circ H$ und $b \circ H$ i.a. keine Linksnebenklasse.

Beispiel 6.37 Sei $G = \mathbf{Z}$ (mit der Addition $+$) und $H = m\mathbf{Z}$ (mit $m \in \mathbf{N}$). Dann besteht $\mathbf{Z}/m\mathbf{Z} = \mathbf{Z}_m$ aus m Nebenklassen $\bar{0} = 0 + m\mathbf{Z} = m\mathbf{Z}$, $\bar{1} = 1 + m\mathbf{Z}$, $\dots$, $\overline{m-1} = (m-1) + m\mathbf{Z}$, den sogenannten **Restklassen modulo m** ²⁵. $\mathbf{Z}_m$ ist übrigens eine zyklische Gruppe, sie wird etwa von $\bar{1} = 1 + m\mathbf{Z}$ erzeugt.

Definition 6.38 Eine Abbildung $\varphi : G \rightarrow H$ zwischen zwei Gruppen $\langle G, \circ \rangle$ und $\langle H, \star \rangle$ heißt **Homomorphismus** oder **Gruppenhomomorphismus**²⁶, wenn für alle $a, b \in G$

$$\varphi(a \circ b) = \varphi(a) \star \varphi(b)$$

gilt.

Ist ein Gruppenhomomorphismus φ injektiv, so heißt φ auch **Monomorphismus**²⁷, und ist φ surjektiv, so nennt man ihn **Epimorphismus**²⁸.

Ist φ bijektiv, so heißt er **Isomorphismus**²⁹. Die inverse Abbildung $\varphi^{-1} : H \rightarrow G$ ist dann auch ein Isomorphismus. Existiert zwischen zwei Gruppen G, H ein Isomorphismus, so heißen G und H **isomorph** und man schreibt dafür $G \cong H$.

Lemma 6.39 Ist $\varphi : G \rightarrow H$ ein Gruppenhomomorphismus, so wird das neutrale Element e_G von G auf das neutrale Element e_H von H abgebildet, d.h. $\varphi(e_G) = e_H$. Weiters gilt $\varphi(a') = \varphi(a)'$ für alle $a \in G$.

Definition 6.40 Sei $\varphi : G \rightarrow H$ ein Gruppenhomomorphismus. Das Urbild $\varphi^{-1}(\{e_H\})$ des neutralen e_H wird als **Kern**³⁰ von φ

$$\text{kern}(\varphi) = \{a \in G \mid \varphi(a) = e_H\}$$

bezeichnet.

Weiters nennt man

$$\varphi(G) = \{b \in H \mid \text{es existiert } a \in G \text{ mit } \varphi(a) = b\}$$

als **Bild**³¹ von G unter φ bezeichnet.

²⁵congruence class

²⁶group homomorphism

²⁷monomorphism

²⁸epimorphism

²⁹isomorphism

³⁰kernel

³¹image

Satz 6.41 Sei $\varphi : G \rightarrow H$ ein Gruppenhomomorphismus. Dann ist $\text{kern}(\varphi)$ ein Normalteiler von G und $\varphi(G)$ eine Untergruppe von H .

Einer der wichtigsten Sätze der Gruppentheorie ist der **Homomorphiesatz**

Satz 6.42 Sei $\varphi : G \rightarrow H$ ein Gruppenhomomorphismus. Dann ist die Faktorgruppe $G/\text{kern}(\varphi)$ mit $\varphi(G)$ isomorph:

$$G/\text{kern}(\varphi) \cong \varphi(G)$$

Die Nebenklasse $a \circ \text{kern}(\varphi) \in G/\text{kern}(\varphi)$ entspricht dem Element $\varphi(a) \in \varphi(G)$.

6.2 Ringe und Körper

In den ganzen Zahlen $\mathbf{Z}$ verwendet man (wenigstens) zwei verschiedene binäre Operationen, die *Addition* und die *Multiplikation*. Bezüglich der Addition ist $\mathbf{Z}$ eine Gruppe und bezüglich der Multiplikation ein Monoid. Man erfaßt aber die Struktur der ganzen Zahlen $\mathbf{Z}$ (bezüglich Addition und Multiplikation) nicht vollständig, wenn man nur die Strukturen $\langle \mathbf{Z}, + \rangle$ und $\langle \mathbf{Z}, \cdot \rangle$ betrachtet. Es gelten auch Rechenregeln, wie das *Distributivgesetz*

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

wo Addition und Multiplikation gemeinsam auftreten.

Im folgenden werden daher algebraische Strukturen $\langle A, +, \cdot \rangle$ mit zwei binären Operationen behandelt, die notationstechnischen Gründen mit $+$ (plus) und $\cdot$ (mal) bezeichnet werden (auch wenn sie mit der *gewöhnlichen* Addition und Multiplikation nichts zu tun haben). Entsprechend bezeichnet man das neutrale Element von $+$, sofern eines existiert, mit 0 (Null) und das von $\cdot$ mit 1 (Eins). Das additiv inverse Element von a ist dann $-a$ und das multiplikative a^{-1} . Schließlich wird, um Klammern zu sparen, wie üblich die Multiplikation vor der Addition ausgeführt.

Definition 6.43 Eine algebraische Struktur $\langle R, *, \cdot \rangle$ (mit zwei binären Operationen) heißt **Halbring**³², wenn die folgenden vier Eigenschaften erfüllt sind:

1. $\langle R, + \rangle$ ist ein kommutatives Monoid mit neutralem Element 0.
2. $\langle R, \cdot \rangle$ ist ein Monoid mit neutralem Element 1, das von 0 verschieden ist.
3. Es gelten die **Distributivgesetze**³³:

$$\begin{aligned} a \cdot (b + c) &= a \cdot b + a \cdot c \\ (a + b) \cdot c &= a \cdot c + b \cdot c \end{aligned} \quad \text{für alle } a, b, c \in R.$$

³²semiring

³³distributive law

4. $a \cdot 0 = 0 \cdot a = 0$ für alle $a \in R$.

Beispiel 6.44 $\langle \mathbb{N}_0, +, \cdot \rangle$ ist ein Halbring.

Beispiel 6.45 $R = \{0, 1\}$ mit den Operationen

$+$	0	1	$\cdot$	0	1
0	0	1	0	0	0
1	1	1	1	0	1

ist der sogenannte **Boolesche Halbring**.

Definition 6.46 Eine algebraische Struktur $\langle R, *, \cdot \rangle$ (mit zwei binären Operation) heißt **Ring**³⁴, wenn die folgenden drei Eigenschaften erfüllt sind:

1. $\langle R, + \rangle$ ist eine kommutative Gruppe (mit neutralem Element 0).
2. $\langle R, \cdot \rangle$ ist eine Halbgruppe.
3. Es gelten die Distributivgesetze.

Besitzt R bezüglich $\cdot$ ein neutrales Element (1), so nennt man R **Ring mit Einselement**, und ist R bezüglich $\cdot$ kommutativ, so nennt man R **kommutativen Ring**.

Insbesondere ist jeder Ring mit Einselement auch ein Halbring.

Beispiel 6.47 $\langle \mathbb{Z}, +, \cdot \rangle$ und $\langle \mathbb{Z}_m, +, \cdot \rangle$ ($m \geq 1$) sind kommutative Ringe mit Einselement. (Ähnlich wie die Addition auf $\mathbb{Z}_m = \mathbb{Z}/m\mathbb{Z}$ definiert man auch die Multiplikation zweier Nebenklassen durch $\overline{a} \cdot \overline{b} = (a + m\mathbb{Z}) \cdot (b + m\mathbb{Z}) := (a \cdot b) + m\mathbb{Z} = \overline{a \cdot b}$.)

Beispiel 6.48 Die $n \times n$ -Matrizen $L_{n,n}(R)$ mit Eintragungen aus einem Ring R bilden mit der Matrizenaddition und -multiplikation wieder einen Ring.

Beispiel 6.49 Die Menge $R[x]$ der **Polynome**³⁵ mit Koeffizienten aus R bildet mit der üblichen Polynomaddition und -multiplikation wieder einen Ring, den **Polynomring über R** ³⁶.

³⁴ring

³⁵polynomial

³⁶ring of polynomials

Sind $p(x) = \sum_{k=0}^{\infty} a_k x^k$, $q(x) = \sum_{k=0}^{\infty} b_k x^k$ zwei Polynome über R , d.h. $a_k, b_k \in R$ und nur endliche viele a_k bzw. b_k sind von 0 verschieden, so berechnen sich Summe $p(x) + q(x)$ und Produkt $p(x) \cdot q(x)$ durch

$$\begin{aligned} p(x) + q(x) &= \sum_{k=0}^{\infty} (a_k + b_k) x^k, \\ p(x) \cdot q(x) &= \sum_{k=0}^{\infty} \left(\sum_{j=0}^k a_j \cdot b_{k-j} \right) x^k. \end{aligned}$$

Man bezeichnet das maximale k mit $a_k \neq 0$ eines Polynoms $p(x) = \sum_{k=0}^{\infty} a_k x^k$ als den **Grad**³⁷ von $p(x)$, i.Z. $\text{grad}(p(x))$. Man beachte, daß immer

$$\begin{aligned} \text{grad}(p(x) + q(x)) &\leq \max(\text{grad}(p(x)), \text{grad}(q(x))), \\ \text{grad}(p(x) \cdot q(x)) &\leq \text{grad}(p(x)) + \text{grad}(q(x)) \end{aligned}$$

gelten.

Beispiel 6.50 Die Menge $R[[x]]$ der **(formalen) Potenzreihen**³⁸ $\sum_{k=0}^{\infty} a_k x^k$ mit Koeffizienten aus $a_k \in R$ bildet mit formal denselben Operation $+, \cdot$ wie bei den Polynomen aus Beispiel 6.49 wieder eine Ring, den **Ring der formalen Potenzreihen über R** .

Lemma 6.51 *In jedem Ring R gilt $a \cdot 0 = 0 \cdot a = 0$.*

In einem Halbring kann diese Eigenschaft nicht aus den anderen Axiomen abgeleitet werden.

In einem Ring kann das Produkt zweier Elemente von 0 verschiedener Element 0 sein, in $\mathbf{Z}_6$ gilt z.B. $\bar{2} \cdot \bar{3} = \bar{6} = \bar{0}$.

Definition 6.52 *Eine Element $a \neq 0$ eines Ringes R heißt **Nullteiler**³⁹, wenn es ein $b \neq 0$ aus R gibt, so daß entweder $a \cdot b = 0$ ist oder $b \cdot a = 0$ ist.*

Dieses b ist damit natürlich auch ein Nullteiler.

Definition 6.53 *Ein kommutativer Ring mit Einselement ohne Nullteiler heißt **Integritätsbereich**⁴⁰ oder **Integritätsring**.*

Beispiel 6.54 $\langle \mathbf{Z}, +, \cdot \rangle$ ist ein Integritätsbereich.

³⁷degree

³⁸formal power series

³⁹zero divisor

⁴⁰integral domain

Beispiel 6.55 $\langle \mathbf{Z}_m, +, \cdot \rangle$ ($m \geq 1$) ist nur dann ein Integritätsbereich, wenn m eine Primzahl ist. Ist m zusammengesetzt, so besitzt $\mathbf{Z}_m$ sicherlich Nullteiler.

Satz 6.56 Der Polynomring $R[x]$ und der Ring der formalen Potentreihen $R[[x]]$ über einem Integritätsbereich R sind wieder Integritätsbereiche.

Definition 6.57 Ein Element a eines Ringes (mit Einselement 1) heißt **Einheit**⁴¹, wenn es bezüglich $\cdot$ ein inverses Element a^{-1} gibt.

Die Menge aller Einheiten R^* von R bildet eine Gruppe, die sogenannte **Einheitengruppe**.

Beispiel 6.58 $\mathbf{Z}^* = \{-1, 1\}$.

Beispiel 6.59 $\mathbf{Z}_m^* = \{\bar{a} = a + m\mathbf{Z} \mid \text{ggT}(a, m) = 1\}$. Diese Einheitengruppe heißt auch **Gruppe der primen Restklassen modulo m** . Ihre Ordnung $|\mathbf{Z}_m^*|$ wird als **Eulersche Phifunktion**⁴² $\varphi(m)$ bezeichnet. Kennt man die Primfaktorenzerlegung von $m = p_1^{e_1} p_2^{e_2} \cdots p_n^{e_n}$, so berechnet man $\varphi(m)$ durch

$$\varphi(m) = \prod_{j=1}^n p_j^{e_j-1} (p_j - 1) = m \prod_{j=1}^n \left(1 - \frac{1}{p_j}\right).$$

Wendet man die allgemeine Beziehung $a^{|G|} = e$ auf die Einheitengruppe $\mathbf{Z}_m^*$ an, so erhält man den **kleinen Fermatschen Satz**

$$a^{\varphi(m)} \equiv 1 \pmod{m} \quad \text{für } a \text{ mit } \text{ggT}(a, m) = 1.$$

Insbesondere gilt für eine Primzahl p die Formel $\varphi(p) = p - 1$, d.h. alle Elemente aus $\mathbf{Z}_p$ außer $\bar{0}$ sind Einheiten.

Definition 6.60 Ein kommutativer Ring $\langle K, +, \cdot \rangle$ mit Einselement $1 \neq 0$, in dem jedes Element $a \neq 0$ eine Einheit ist, heißt ein **Körper**⁴³.

Eine algebraische Struktur $\langle K, +, \cdot \rangle$ ist also genau dann ein Körper, wenn $\langle K, + \rangle$ und $\langle K \setminus \{0\}, \cdot \rangle$ kommutative Gruppen sind und die Distributivgesetze gelten.

Beispiel 6.61 $\langle \mathbf{Q}, +, \cdot \rangle$, $\langle \mathbf{R}, +, \cdot \rangle$ und $\langle \mathbf{C}, +, \cdot \rangle$ sind Körper.

Satz 6.62 Jeder Körper ist ein Integritätsbereich.

⁴¹unit

⁴²Euler's totient function

⁴³field

Die Umkehrung dieses Satzes gilt im allgemeinen nicht, wie das Beispiel $\mathbf{Z}$ zeigt. Es gilt allerdings:

Satz 6.63 *Jeder endliche Integritätsbereich ist ein Körper.*

Satz 6.64 $\langle \mathbf{Z}_m, +, \cdot \rangle$ ist genau dann ein Körper, wenn $m = p$ eine Primzahl ist. $\langle \mathbf{Z}_p, +, \cdot \rangle$ ist dann ein **endlicher Körper**⁴⁴ der Ordnung p .

Allgemein läßt sich zeigen, daß es nur für Primzahlpotenzen p^m ($m \geq 1$) endliche Körper mit p^m Elementen gibt. Es gibt daher keinen Körper mit 6 Elementen, aber einen mit 8 und einen mit 9.

Im Beispiel 6.59 wird die Einheitengruppe $\mathbf{Z}_m^* = \{\bar{a} = a + m\mathbf{Z} \mid \text{ggT}(a, m) = 1\}$ betrachtet. Ein Element $\bar{a}$ besitzt daher ein inverses Element $\bar{a}'$ genau dann, wenn $\text{ggT}(a, m) = 1$. Abschließend soll nun der **Euklidische Algorithmus** vorgestellt werden, mit dessen Hilfe man nicht nur den größten gemeinsamen Teiler zweier ganzer Zahlen, sondern auch das inverse Element $\bar{a}'$ sehr effektiv berechnet werden kann. Er beruht auf folgender einfachen Eigenschaft.

Lemma 6.65 (Division mit Rest) *Zu je zwei ganzen Zahlen a, b mit $b \neq 0$ gibt es ganze Zahlen q, r mit*

$$a = bq + r \quad \text{und} \quad 0 \leq r < b.$$

q heißt Quotient und r Rest.

Satz 6.66 *Führt man zu zwei ganzen Zahlen a, b mit $b \neq 0$ die Divisionskette*

$$\begin{aligned} a &= bq_0 + r_0, & 0 < r_0 < b \\ b &= r_0q_1 + r_1 & 0 < r_1 < r_0 \\ r_0 &= r_1q_2 + r_2 & 0 < r_2 < r_1 \\ &\vdots \\ r_{k-2} &= r_{k-1}q_k + r_k & 0 < r_k < r_{k-1} \\ r_{k-1} &= r_kq_{k+1} + 0 & , \end{aligned}$$

durch, so muß diese wegen $b > r_0 > r_1 > r_2 > \cdots \geq 0$ einmal abbrechen, d.h. es gibt irgendeinmal einen verschwindenden Rest. Der letzte Rest $r_k \neq 0$ ist dann der größte gemeinsame Teiler $\text{ggT}(a, b)$.

Beispiel 6.67 Zur Bestimmung des $\text{ggT}(59, 11)$ ermittelt man die Divisionskette

$$\begin{aligned} 59 &= 11 \cdot 5 + 4 \\ 11 &= 4 \cdot 2 + 3 \\ 4 &= 3 \cdot 1 + 1 \\ 3 &= 1 \cdot 3 + 0. \end{aligned}$$

⁴⁴finite field

Es ist also $\text{ggT}(59, 11) = 1$. Umgekehrt kann man mit Hilfe dieser Divisionskette auch den ggT zweier Zahlen als ganzzahlige Linearkombination von a, b darstellen, indem man ausgehend von der Gleichung $\text{ggT}(59, 11) = 4 - 3 \cdot 1$ sukzessive die weiteren Reste 3 und 4 rückeinsetzt:

$$\begin{aligned} 1 &= 4 - 3 \cdot 1 \\ &= 4 - (11 - 4 \cdot 2) \cdot 1 \\ &= 3 \cdot 4 - 1 \cdot 11 \\ &= 3 \cdot (59 - 5 \cdot 11) - 1 \cdot 11 \\ &= 3 \cdot 59 - 15 \cdot 11 - 1 \cdot 11 \\ &= 3 \cdot 59 - 16 \cdot 11. \end{aligned}$$

Das ergibt z.B. $11^{-1} = -16 \equiv 43 \pmod{59}$.

Satz 6.68 *Ist d der größte gemeinsame Teiler der von Null verschiedenen ganzen Zahlen a, b , so gibt es ganze Zahlen k, l mit*

$$ak + bl = d,$$

die mit Hilfe der Divisionskette von a und b effektiv berechnet werden können.

Die Division mit Rest (Lemma 6.65) funktioniert nicht nur für den Bereich der ganzen Zahlen.

Lemma 6.69 *Seien $a(x)$ und $q(x)$ zwei nichtverschwindende Polynome mit Koeffizienten aus einem Körper K . Dann gibt es Polynome $q(x), r(x) \in K[x]$ mit*

$$a(x) = b(x)q(x) + r(x),$$

wobei $r(x)$ entweder das Nullpolynom ist oder $\text{grad}(r(x)) < \text{grad}(b(x))$.

Genauso wie in den ganzen Zahlen kann jetzt mit einer Divisionskette der größte gemeinsame Teiler zweier Polynome über einem Körper bestimmt werden. Dabei heißt ein Polynom $d(x) \in K[x]$ größter gemeinsamer Teiler der Polynome $a(x), b(x) \in K[x]$, wenn $d(x)$ ein gemeinsamer Teiler von $a(x)$ und $b(x)$ ist und jeder gemeinsame Teiler $t(x) \in K[x]$ von $a(x)$ und $b(x)$ ein Teiler von $d(x)$ ist.

Satz 6.70 *Zu je zwei nichtverschwindenden Polynomen $a(x), b(x)$ mit Koeffizienten aus einem Körper K gibt es immer einen größte gemeinsamen Teiler $d(x) \in K[x]$. Weiters gibt es Polynome $k(x), l(x) \in K[x]$ mit*

$$a(x)k(x) + b(x)l(x) = d(x).$$

Ähnlich wie den Ring $\mathbf{Z}_m = \mathbf{Z}/m\mathbf{Z}$ kann man auch einen Faktorpolynomring über einem Körper K aufbauen. Sei $p(x) \in K[x]$ ein festes Polynom (ungleich dem Nullpolynom). Die polynomiellen Vielfachen $p(x)K[x]$ dieses Polynoms bilden eine (additiven) Normalteiler von $K[x]$. Auf der Faktorgruppe $K[x]/p(x)K[x]$ kann aber auch (wie in $\mathbf{Z}/m\mathbf{Z}$) in natürlicher Weise eine Multiplikation definiert werden, und $K[x]/p(x)K[x]$ wird dadurch wieder zu einem Ring.

Satz 6.71 *Sei K ein Körper und $p(x) \in K[x]$ ein Polynom mit Koeffizienten aus K . Dann ist der Faktorring $\langle K[x]/p(x)K[x], +, \cdot \rangle$ genau dann ein Körper, wenn das Polynom $p(x)$ **irreduzibel**⁴⁵ über K ist, d.h. $p(x)$ kann nicht als Produkt zweier Polynome $a(x), b(x) \in K[x]$ mit kleinerem Grad (als $p(x)$) dargestellt werden.*

Beispiel 6.72 Das Polynom $p(x) = x^2 + 1$ ist irreduzibel über $\mathbf{R}$. Daher ist der Faktorring $\mathbf{R}[x]/(x^2 + 1)\mathbf{R}[x]$ ein Körper. Die Nebenklasse $x + (x^2 + 1)\mathbf{R}[x]$ hat die Eigenschaft $(x + (x^2 + 1)\mathbf{R}[x])^2 = -1 + (x^2 + 1)\mathbf{R}[x]$. Es ist leicht einzusehen, daß $\mathbf{R}[x]/(x^2 + 1)\mathbf{R}[x]$ nichts anderes als die komplexen Zahlen $\mathbf{C}$ repräsentiert. Die imaginäre Einheit i muß nur mit $x + (x^2 + 1)\mathbf{R}[x]$ identifiziert werden.

Beispiel 6.73 Ist $q(x)$ ein irreduzibles Polynom vom Grad k über $\mathbf{Z}_p$ (mit einer Primzahl p), so ist $\mathbf{Z}_p[x]/q(x)\mathbf{Z}_p[x]$ ein endlicher Körper mit p^k Elementen.

6.3 Codierungstheorie

6.3.1 Grundlegende Begriffe

Übermittelt man Daten über eine Leitung auf elektronischem Wege oder speichert man Daten ab und liest sie zu einem späteren Zeitpunkt wieder, so können durch verschiedene Einflüsse *Übertragungsfehler* auftreten, d.h. daß die empfangene Nachricht nicht mehr mit der gesendeten übereinstimmt. Aus diesem Grund ist es sinnvoll, die ursprüngliche Nachricht in einer Art und Weise umzukodieren, so daß gewisse Fehler erkannt und möglicherweise sogar korrigiert werden können.

Beispiel 6.74 Liegt eine Nachricht in 0-1-Blöcken gleicher Länge vor, so wird jeder Block um ein 0-1-bit **parity bit** verlängert und zwar so, daß dann in jedem Block eine gerade Anzahl von 1 steht.

Z.B. wird die Nachricht 00101, 11010, 10110, 01011 zu 001010, 110101, 101101, 010111 umkodiert.

Beim Empfänger wird nun getestet, ob jeder empfangene Block eine gerade Anzahl von 1 besitzt. Ist dies einmal nicht so, so muß bei der Übertragung dieses Blocks ein Übertragungsfehler passiert sein. Dieser Fehler kann aber (ohne Rückfrage beim Sender) nicht behoben werden.

⁴⁵irreducible polynomial

Beispiel 6.75 Beim **Wiederholungcode** wird jeder 0-1-Block m -mal wiederholt. Ist z.B. $m = 3$, so können Fehler an einer Stelle korrigiert werden und Fehler an höchstens zwei Stellen erkannt werden.

Definition 6.76 Sei A eine endliche Menge, die im folgenden **Alphabet** bezeichnet werden wird.

Ein (n, k) -**Code**⁴⁶ $C = f(A^k) \subseteq A^n$ ist das Bild einer injektiven Abbildung

$$f : A^k \rightarrow A^n.$$

Die Elemente aus A^k nennt man **Nachrichtenworte** und die aus C **Codeworte**.

Beispiel 6.77 Der Code im Beispiel 6.74 wird mit dem Alphabet $A = \{0, 1\}$ und der Codierungsfunktion $f_1 : A^k \rightarrow A^{k+1}$ mit

$$f_1(a_1, a_2, \dots, a_k) = (a_1, a_2, \dots, a_k, a_1 + a_2 + \dots + a_k \bmod 2).$$

Entsprechend ist die Codierungsfunktion im Beispiel 6.75 $f_2 : A^k \rightarrow A^{mk}$ mit

$$f_2(a_1, \dots, a_k) = (a_1, \dots, a_k, a_1, \dots, a_k, \dots, a_1, \dots, a_k).$$

Entsteht bei der Übertragung von Nachrichten ein Übertragungsfehler, so wird (üblicherweise) vom Empfänger ein Wort w aus A^n empfangen, das kein Codewort ist. Um diese Abweichung zu quantifizieren, verwendet man den Begriff der *Hammingdistanz*.

Definition 6.78 Die **Hammingdistanz** $d(v, w)$ zweier Worte $v = a_1 a_2 \dots a_n$, $w = b_1 b_2 \dots b_n \in A^n$ ist die Anzahl der verschiedenen Stellen $a_j \neq b_j$ ($1 \leq j \leq n$).

Tritt bei der Übertragung eines Wortes an zwei Stellen ein Fehler auf, so ist die Hammingdistanz 2.

Die Hammingdistanz $d(v, w)$ ist übrigens eine **Metrik** auf A^n .

Definition 6.79 Sei $C \subseteq A^n$ ein Code, dann ist die **Minimaldistanz**

$$\delta(C) = \min_{v, w \in C, v \neq w} d(v, w)$$

die minimale Hammingdistanz zweier verschiedener Codeworte.

Satz 6.80 Ist $\delta(C)$ die Minimaldistanz eines Codes $C \subseteq A^n$, so können (Übertragungs-) Fehler an höchstens $k < \delta(C)$ Stellen erkannt und (Übertragungs-) Fehler an höchstens $k < \frac{\delta(C)}{2}$ Stellen korrigiert werden.

⁴⁶code

Bei einem Code ist es immer ein Ziel, ein **Korrekturschema** $K(C)$ aufzustellen. Dies ist eine Tabelle, wo in der ersten Zeile die Codeworte $C = \{c_1, c_2, \dots\}$ aufgelistet sind, und in der jeweiligen Spalte unter einem Codewort c_j stehen alle Worte aus A^n , die zu c_j korrigiert werden. Satz 6.80 sagt zum Beispiel aus, daß jene Worte $w \in A^n$, deren Distanz zu C kleiner als $\delta(C)/2$ sicherlich in jedes Korrekturschema aufgenommen werden. Es kann aber auch noch weitere Worte in A^n geben, die eindeutig korrigiertbar sind. Jedenfalls kann man bei der Erstellung eines Korrekturschemas $K(C)$ zwischen 2 Prinzipien unterscheiden:

1. Es werden in $K(C)$ nur jene Worte eingetragen, für die die Korrektur eindeutig möglich ist. Es können dann möglicherweise nicht alle Worte aus A^n korrigiert werden.
2. Es werden alle Worte aus A^n in $K(C)$ eingetragen. Gibt es zu einem Wort $w \in A^n$ zwei oder mehr verschiedene Codeworte mit gleichem minimalem Abstand, so muß man sich entscheiden, zu welchem Codewort korrigiert werden soll. Dabei können Zusatzinformationen über die Fehlerwahrscheinlichkeit im Übertragungskanal berücksichtigt werden. (Oft sind *Fehlerbündel* wahrscheinlicher als gleichverteilte Fehler.)

6.3.2 Gruppencodes

Definition 6.81 Sei A eine abelsche Gruppe. Ein **Gruppencode** $C \subseteq A^n$ ($n \in \mathbb{N}$) ist ein Code, der gleichzeitig Untergruppe von A^n ist.

Beispiel 6.82 Jede Untergruppe C von $\langle \mathbb{Z}_2^n, + \rangle$ oder $\langle \mathbb{Z}_2^n, + \rangle$ ist ein Gruppencode.

Satz 6.83 Sei $C \subseteq A^n$ ein Gruppencode, dann gilt

$$\delta(C) = \min_{c \in C \setminus \{0\}} w(c),$$

wobei $w(c) = d(c, 0)$ das **Gewicht** eines Wortes bezeichnet.

Die Minimaldistanz eines Gruppencodes kann daher sehr schnell berechnet werden. Weiters besteht ein direkter Zusammenhang zwischen der Nebenklassenzerlegen von C in A^n und dem Korrekturschema.

Satz 6.84 Sei $v + C$ eine Nebenklasse von C mit der Eigenschaft, daß v das einzige Element von $v + C$ mit minimalem Gewicht ist. Dann wird das Element $v + c \in v + C$ ($c \in C$) in jedem Korrekturschema $K(C)$ zu $c \in C$ korrigiert.

Definition 6.85 Sei $v + C$ eine Nebenklasse von C . Ein Element $v_0 \in v + C$ heißt **Nebenklassenführer**, wenn es minimales Gewicht in $v + C$ hat.

Ein Nebenklassenführer muß nicht eindeutig sein (wie wir im Beispiel 6.86 noch sehen werden). Ist er jedoch eindeutig, so macht die Korrektur keine Probleme. Empfängt man ein Wort v aus einer Nebenklasse mit eindeutigem Nebenklassenführer v_0 , so wird v einfach zu $v - v_0 \in C$ korrigiert. Man kann daher das Korrekturschema so aufbauen, daß man zunächst alle Nebenklassen von C bestimmt und jede Nebenklasse mit eindeutigem Nebenklassenführer v_0 in eine *Zeile* des tabellarischen Korrekturschemas einträgt. Unter dem Wort $c \in C$ schreibt man $v_0 + c$. v_0 ist natürlich jenes Wort, das zu $00 \cdots 0 \in C$ korrigiert wird.

Bei jenen Nebenklassen, die keinen eindeutigen Nebenklassenführer besitzen, geht man jetzt prinzipiell genauso vor. Man muß nur unter den möglichen Kandidaten mit minimalem Gewicht einen aussuchen. Dieser wird dann zu $00 \cdots 0$ korrigiert. Die Korrektur der anderen Elemente einer solchen Nebenklasse ist, wenn man sich für einen Kandidaten entschieden hat, dann schon festgelegt. Man muß nur den ausgewählten Nebenklassenführer subtrahieren.

Beispiel 6.86 Es sei $C = \{00000, 11010, 00111, 11101\}$ Untergruppe von $\mathbf{Z}_2^5$. C hat $2^5/4 = 8$ Nebenklassen und Gewicht $w(C) = 3$. Man kann daher Fehler an einer Stelle eindeutig korrigieren. Daher sind die Nebenklassen $10000 + C$, $01000 + C$, $00100 + C$, $00010 + C$, $00001 + C$ sicher paarweise verschieden und die Nebenklassenführer sind $10000, 01000, 00100, 00010, 00001$. Da C auch eine Nebenklasse ist, fehlen nur mehr zwei, nämlich

$$\{01100, 10110, 01011, 10001\}$$

und

$$\{01001, 10011, 01110, 10100\}.$$

In beiden Nebenklassen gibt es zwei Elemente mit minimalem Gewicht 2. Es werden hier 01100 und 01001 als Führer gewählt. Daher ergibt sich das folgende Korrekturschema $K(C)$:

C	00000	11010	00111	11101
$10000 + C$	10000	01010	10111	01101
$01000 + C$	01000	10010	01111	10101
$00100 + C$	00100	11110	00011	11001
$00010 + C$	00010	11000	00101	11111
$00001 + C$	00001	11011	00110	11100
$01100 + C$	01100	10110	01011	10001
$01001 + C$	01001	10011	01110	10100

Ein wichtiger Spezialfall von Gruppencodes sind **Linearcodes**. Hier betrachtet man eine injektive, lineare Abbildung $f : K^k \rightarrow K^n$ zwischen dem k -dimensionalen Vektorraum K^k und dem n -dimensionalen Vektorraum K^n über einem endlichen Körper K . Das Bild $C = f(K^k)$ ist dann ein k -dimensionaler Teilraum von K^n , also eine Untergruppe bezüglich der Addition. f ist die Codierungsabbildung, die auch durch eine **Generatormatrix** G (mit k Zeilen und n Spalten mit vollem Rang k) durch

$$f(v_1, v_2, \dots, v_k) = (v_1 \ v_2 \ \cdots \ v_k)G$$

realisiert werden. Die Zeilen von G bilden auch eine Basis von C .

Der große Vorteil von Linearcodes ist, daß die Fehlerkorrektur sehr einfach durchgeführt werden kann. Dazu benötigt man eine **Kontrollmatrix** H (mit $n - k$ Zeilen und n Spalten mit vollem Rang $n - k$) mit der Eigenschaft, daß GH^T die Nullmatrix ergibt. Es stellt sich heraus, daß zwei Vektoren $(w_1, w_2, \dots, w_n)$, $(w'_1, w'_2, \dots, w'_n)$ genau dann in derselben Nebenklasse von C liegen, wenn

$$(w_1, w_2, \dots, w_n)H^T = (w'_1, w'_2, \dots, w'_n)H^T$$

gilt. Die möglichen Bilder $(w_1, w_2, \dots, w_n)H^T$ liegen in K^{n-k} . Jedes Element aus K^{n-k} entspricht daher genau einer Nebenklassen von C . Zur Korrektur benötigt man daher nur eine Kontrollmatrix H und eine Liste, die jedem Element aus K^{n-k} einen Nebenklassenführer der entsprechenden Nebenklasse von C zuordnet.

6.3.3 Das RSA-Verfahren

Im Unterschied zur Codierungstheorie, wo die Fehlererkennung und -korrektur bei Fehlübertragungen im Vordergrund steht, ist bei der **Verschlüsselung**⁴⁷ das Ziel, sie *abhörsicher* zu machen, d.h. ein potentieller *Lauscher* soll nicht in der Lage sein, die gesendete verschlüsselte Nachricht zu entziffern.

Verschlüsselungsverfahren finden nicht nur militärische Anwendungen, sie werden natürlich auch im zivilen Bereich, etwa im Bankenwesen, eingesetzt.

Das momentan beste Verschlüsselungsverfahren ist das nach Rivest, Schamir und Adleman benannte **RSA-Verfahren**, das sich auch dadurch auszeichnet, daß es mit **öffentlichen Schlüsseln**⁴⁸ arbeitet.

Das RSA-Verfahren beruht auf dem folgenden Satz:

Satz 6.87 Seien p, q zwei verschiedene ungerade Primzahlen, $n = pq$ und $v = \text{kgV}(p - 1, q - 1)$. Dann gilt für alle $a, m \in \mathbf{Z}$

$$a^{mv+1} \equiv a \pmod{n}.$$

Sind also $e, d \in \mathbf{Z}$ zwei Zahlen mit $ed \equiv 1 \pmod{v}$, so gilt für alle $a \in \mathbf{Z}$

$$(a^e)^d \equiv a \pmod{n}.$$

Möchte also eine Person A verschlüsselte Nachrichten empfangen können, so multipliziert A zwei (i.a. mindestens 100-stellige) Primzahlen p, q miteinander und veröffentlicht das Produkt $n = pq$ und eine Zahl e , die zu $v = \text{kgV}(p - 1, q - 1)$ teilerfremd ist.

Ist nun eine weitere Person B daran interessiert, der Person A eine Nachricht zu senden, die nur A lesen kann, so unterteilt er die Nachricht in Blöcke $a_1, a_2, \dots$, so daß jeder Block

⁴⁷enciphering

⁴⁸public key crypto system

durch eine nichtnegative Zahl $< n$ repräsentiert werden kann. Anschließend verschlüsselt er $a_1, a_2, \dots$ durch

$$b_j = a_j^e \bmod n \quad (j \geq 1)$$

und läßt A die Zahlen $b_1, b_2, \dots$ ohne weitere Geheimhaltung zukommen.

A kann nun die ursprüngliche Nachricht $a_1, a_2, \dots$ durch

$$a_j = b_j^d \bmod n \quad (j \geq 1)$$

zurückgewinnen, also entschlüsseln.

Auf dem ersten Blick erscheint das RSA-Verfahren nicht sehr sicher zu sein. Da n und e bekannt sind, sollte man doch $d = e^{-1} \bmod v$ berechnen können. Der Grund, warum dies nicht funktioniert, ist die Tatsache, daß $v = \text{kgV}(p-1, q-1)$ nur mit Hilfe der Primfaktorenzerlegung von $n = pq$ errechnet werden kann, und diese wird nicht bekanntgegeben.

Es hat sich gezeigt, daß die Ermittlung der Primfaktorenzerlegung einer natürlichen Zahl bei wenigen *großen* Primfaktoren am schwierigsten ist. Weiters steigt der Aufwand, die Primfaktorenzerlegung einer Zahl zu bestimmen sehr schnell mit der Größe der vorgegebenen Zahl. Beispielsweise liegt die erwartete Rechenzeit zur Faktorisierung einer 200-stelligen Zahl momentan außerhalb unserer Lebenserwartung.

6.4 Verbände und Boolesche Algebren

Definition 6.88 Eine algebraische Struktur $\langle M, \wedge, \vee \rangle$ heißt **Verband**⁴⁹, wenn folgende Eigenschaften erfüllt sind:

1. $\langle M, \wedge \rangle$ ist eine kommutative Halbgruppe.
2. $\langle M, \vee \rangle$ ist eine kommutative Halbgruppe.
3. Es gelten die **Verschmelzungsgesetze**

$$\begin{aligned} a &= a \wedge (a \vee b), \\ a &= a \vee (a \wedge b) \end{aligned}$$

für alle $a, b \in M$.

Beispiel 6.89 $\langle \mathbf{P}(A), \cap, \cup \rangle$ ist für jede Menge A ein Verband.

Beispiel 6.90 $M = \mathbf{N}$ mit $a \wedge b = \min(a, b)$ und $a \vee b = \max(a, b)$ ist ein Verband.

Beispiel 6.91 $M = \mathbf{N}$ mit $a \wedge b = \text{ggT}(a, b)$ und $a \vee b = \text{kgV}(a, b)$ ist ein Verband.

$\wedge$	0	1
0	0	0
1	0	1

$\vee$	0	1
0	0	1
1	1	1

Beispiel 6.92 $B = \{0, 1\}$ mit den Operationen $\wedge, \vee$ ist ein Verband.

Ebenso ist $B^n = \{(x_1, x_2, \dots, x_n) \mid x_j \in B \ (1 \leq j \leq n)\}$, wobei $\cap$ und $\vee$ *komponentenweise* durchgeführt werden, ein Verband.

$\langle B^n, \cap, \vee \rangle$ ist übrigens zu $\langle \mathbf{P}(A), \cap, \cup \rangle$ *isomorph*, wobei A eine endliche Menge mit n Elementen $A = \{a_1, a_2, \dots, a_n\}$ bezeichnet. Einer Teilmenge $B \subseteq A$ kann mit $(\chi_B(a_1), \chi_B(a_2), \dots, \chi_B(a_n)) \in B^n$ identifiziert werden.

Definition 6.93 Sei $\leq$ eine Halbordnung auf einer Menge A .

Eine Element $c \in A$ heißt **Infimum**⁵⁰ $\inf(a, b)$ zweier Elemente $a, b \in A$, wenn $c \leq a$ und $c \leq b$ ist und für jedes Element $d \in A$ mit $d \leq a$ und $d \leq b$ auch $d \leq c$ gilt.

Eine Element $c \in A$ heißt **Supremum**⁵¹ $\sup(a, b)$ zweier Elemente $a, b \in A$, wenn $a \leq c$ und $b \leq c$ ist und für jedes Element $d \in A$ mit $a \leq d$ und $b \leq d$ auch $c \leq d$ gilt.

Das Infimum wird auch als *größte untere Schranke* und das Supremum als *kleinste obere Schranke* bezeichnet.

Satz 6.94 Sei $\langle M, \wedge, \vee \rangle$ ein Verband. Dann wird durch

$$a \leq b : \Longleftrightarrow a = a \wedge b$$

auf M eine Halbordnung definiert. In dieser Halbordnung gibt es zu je zwei Elementen ein Infimum, nämlich $\inf(a, b) = a \wedge b$, und ein Supremum $\sup(a, b) = a \vee b$.

Ist umgekehrt $\leq$ eine Halbordnung auf M mit der Eigenschaft, daß es zu je zwei Elementen ein Infimum und ein Supremum gibt, so ist M mit den Operationen $a \wedge b = \inf(a, b)$ und $a \vee b = \sup(a, b)$ ein Verband.

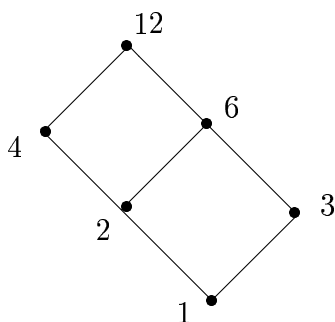
Da jede (endliche) Halbordnung durch ein **Hassediagramm** dargestellt werden kann, ist es auch möglich, einen Verband durch das Hassediagramm der entsprechenden Halbordnung zu repräsentieren.

Beispiel 6.95 $M = \{1, 2, 3, 4, 6, 12\}$ mit $a \wedge b = \text{ggT}(a, b)$ und $a \vee b = \text{kgV}(a, b)$ ist ein Verband. Sein Hassediagramm hat die folgende Gestalt:

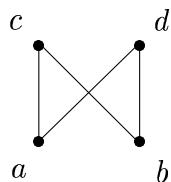
⁴⁹lattice

⁵⁰infimum

⁵¹supremum



Beispiel 6.96 Das folgende Beispiel einer Halbordnung zeigt, daß es auch Halbordnungen auf sehr kleinen Mengen gibt, wo nicht für alle Paare ein Infimum bzw. Supremum existiert.



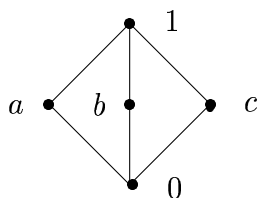
Diese Halbordnung repräsentiert daher keinen Verband.

Definition 6.97 Ein Verband $\langle M, \wedge, \vee \rangle$ heißt **distributiver Verband**, wenn die **Distributivgesetze**

$$\begin{aligned} a \wedge (b \vee c) &= (a \wedge b) \vee (a \wedge c), \\ a \vee (b \wedge c) &= (a \vee b) \wedge (a \vee c) \end{aligned}$$

für alle $a, b, c \in M$ gelten.

Beispiel 6.98 Der durch das folgende Hassediagramm gegebene Verband



ist nicht distributiv. Es gilt

$$a \wedge (b \vee c) = a \wedge 1 = a \quad \text{aber} \quad (a \wedge b) \vee (a \wedge c) = 0 \vee 0 = 0.$$

Gewisse Verbände haben außer den Distributivgesetzen noch weitere Eigenschaften. Beispielsweise hat der Verband $\langle \mathbf{P}(A), \cap, \cup \rangle$ aus Beispiel 6.89 die ganze Menge A als *gemeinsame obere Schranke* und die leere Menge $\emptyset$ als *gemeinsame untere Schranke*. Diese Elemente sind dann natürlich neutrale Elemente für $\cap$ und $\cup$. Weiters gibt es zu jeder Menge $B \in \mathbf{P}(A)$ das Komplement $B' = A \setminus B$ mit den Eigenschaften $\inf(B, B') = B \cap B' = \emptyset$ und $\sup(B, B') = B \cup B' = A$. $\langle \mathbf{P}(A), \cap, \cup \rangle$ bildet eine sogenannte *Boolesche Algebra*.

Definition 6.99 Ein distributiver Verband $\langle M, \wedge, \vee \rangle$ heißt **Boolesche Algebra**, wenn er die folgenden beiden zusätzlichen Eigenschaften besitzt:

1. Es gibt ein neutrales Element $1 \in M$ bezüglich $\wedge$, und es gibt ein neutrales Element $0 \in M$ ($0 \neq 1$) bezüglich $\vee$, d.h. $\langle M, \wedge \rangle$ und $\langle M, \vee \rangle$ sind Monoide.
2. Zu jedem $a \in M$ gibt es ein Komplement $a' \in M$ mit

$$a \vee a' = 1 \quad \text{und} \quad a \wedge a' = 0.$$

Satz 6.100 Eine Boolesche Algebra $\langle M, \wedge, \vee \rangle$ hat die folgenden Eigenschaften:

- Für alle $a \in M$ gilt

$$a \vee 1 = 1 \quad \text{und} \quad a \wedge 0 = 0.$$

- Gelten für ein $b \in M$ die Beziehungen $a \vee b = 1$ und $a \wedge b = 0$, so ist $a' = b$.

- Für alle $a \in M$ gilt

$$(a')' = a.$$

- Für alle $a, b \in M$ gelten die **DeMorganschen Regeln**

$$\begin{aligned} (a \wedge b)' &= a' \vee b', \\ (a \vee b)' &= a' \wedge b'. \end{aligned}$$

Beispiel 6.101 $\langle B^n, \cap, \cup \rangle$ ist für jedes $n \in \mathbb{N}$ eine Boolesche Algebra.

Satz 6.102 Jede endliche Boolesche Algebra ist zu $\langle B^n, \cap, \cup \rangle$ für ein $n \in \mathbb{N}$ isomorph, d.h. es gibt nur endliche Boolesche Algebren der Ordnung 2^n ($n \in \mathbb{N}$), und diese sind bis auf Isomorphie durch $\langle B^n, \cap, \cup \rangle$ gegeben.

6.5 Allgemeine Algebren

Definition 6.103 Sei A eine nichtleere Menge und $n \in \mathbb{N}$. Eine **n -stellige Operation** ω auf A ist eine Abbildung

$$\omega : A^n = A \times A \times \cdots \times A \rightarrow A,$$

d.h. jedem n -Tupel $(a_1, a_2, \dots, a_n) \in A^n$ wird ein Element $\omega(a_1, a_2, \dots, a_n) \in A$ zugeordnet.

Eine **0-stellige Operation** ω auf A ist ein bestimmtes Element aus A .

Beispiel 6.104 In einer Gruppe $\langle G, \circ \rangle$ gibt es neben der zweistelligen (bzw. binären) Operation $\circ$ noch eine 0-stellige, die das neutrale Element e darstellt und eine einstellige, die jedem $a \in G$ das inverse Element a' zuordnet.

Definition 6.105 Ist A eine nichtleere Menge und sind $\omega_1, \omega_2, \dots$ n_1 -stellige, n_2 -stellige, ... Operationen auf A , so heißt das System

$$\langle A, \omega_1, \omega_2, \dots \rangle$$

allgemeine Algebra⁵² und

$$\langle n_1, n_2, \dots \rangle$$

Typ von $\langle A, \omega_1, \omega_2, \dots \rangle$.

Die folgenden (bisher besprochenen) algebraischen Strukturen sind allgemeine Algebren und haben den angegebenen Typ:

Halbgruppe	$\langle A, \circ \rangle$	$\langle 2 \rangle$
Monoid	$\langle A, \circ, e \rangle$	$\langle 2, 0 \rangle$
Gruppe	$\langle A, \circ, e, ' \rangle$	$\langle 2, 0, 1 \rangle$
Ring	$\langle A, +, 0, -, \cdot \rangle$	$\langle 2, 0, 1, 2 \rangle$
Ring mit 1	$\langle A, +, 0, -, \cdot, 1 \rangle$	$\langle 2, 0, 1, 2, 0 \rangle$
Verband	$\langle A, \wedge, \vee \rangle$	$\langle 2, 2 \rangle$
Boolesche Algebra	$\langle A, \wedge, 1, \vee, 0, ' \rangle$	$\langle 2, 0, 2, 0, 1 \rangle$

Körper sind beispielsweise keine allgemeine Algebren, da die Operation, die einem Element a sein multiplikativ inverses Element a^{-1} zuordnet, für $a = 0$ nicht definiert werden kann. Ebenso sind Integritätsbereiche keine allgemeinen Algebren.

Für allgemeine Algebren $\langle A, \omega_1, \omega_2, \dots \rangle$ gibt es universelle Konstruktionen, zu weiteren Algebren deselben Typs überzugehen, nämlich **Unteralgebren**, **direkte Produkte** und **Faktoralgebren**.

Definition 6.106 Sei $\langle A, \omega_1, \omega_2, \dots \rangle$ eine allgemeine Algebra und U eine nichtleere Teilmenge von A . Ist U abgeschlossen bezüglich aller Operationen $\omega_1, \omega_2, \dots$, d.h. für alle $a_1, a_2, \dots, a_{n_j} \in U$ gilt $\omega_j(a_1, a_2, \dots, a_{n_j}) \in U$ ($j = 1, 2, \dots$), so heißt $\langle U, \omega_1, \omega_2, \dots \rangle$ **Unteralgebra**⁵³ von A .

Definition 6.107 Sind $\langle A', \omega'_1, \omega'_2, \dots \rangle$ und $\langle A'', \omega''_1, \omega''_2, \dots \rangle$ zwei Algebren vom selben Typ, so bezeichnet man $\langle A' \times A'', \omega_1, \omega_2, \dots \rangle$ mit den Operationen

$$\omega_j((a'_1, a''_1), (a'_2, a''_2), \dots, (a'_{n_j}, a''_{n_j})) = (\omega'_j(a'_1, a'_2, \dots, a'_{n_j}), \omega''_j(a''_1, a''_2, \dots, a''_{n_j}))$$

als **direktes Produkt**⁵⁴ der Algebren A' und A'' .

⁵²universal algebra

⁵³subalgebra

⁵⁴product

Definition 6.108 Sei $\langle A, \omega_1, \omega_2, \dots \rangle$ eine allgemeine Algebra. Eine Äquivalenzrelation R auf A heißt **Kongruenzrelation**⁵⁵ wenn für alle $a_1, a_2, \dots, a_{n_j}$ und $b_1, b_2, \dots, b_{n_j}$, die $a_1 R b_1, a_2 R b_2, \dots, a_{n_j} R b_{n_j}$ erfüllen, auch

$$\omega_j(a_1, a_2, \dots, a_{n_j}) R \omega_j(b_1, b_2, \dots, b_{n_j}) \quad (j = 1, 2, \dots)$$

gilt.

Für eine Kongruenzrelation R auf A und $a \in A$ bezeichne

$$K(a) = \{b \in A \mid a R b\}$$

jene Äquivalenzklasse von R , die a enthält. Definiert man auf der Menge $A/R = \{K(a) \mid a \in A\}$ die Operationen

$$\bar{\omega}_j(K(a_1), K(a_2), \dots, K(a_{n_j})) = K(\omega_j(a_1, a_2, \dots, a_{n_j})) \quad (j = 1, 2, \dots),$$

so bezeichnet $\langle A/R, \bar{\omega}_1, \bar{\omega}_2, \dots \rangle$ die **Faktoralgebra**⁵⁶ von A nach (der Kongruenzrelation) R .

Beispiel 6.109 Für Gruppen G entsprechen die Kongruenzrelationen R auf G genau den Normalteilern H von G , d.h. die von einer Kongruenzrelation R erzeugten Partition von G ist die Nebenklassenzerlegung eines Normalteilers H von G .

Definition 6.110 Sind $\langle A, \omega_1, \omega_2, \dots \rangle$ und $\langle B, \bar{\omega}_1, \bar{\omega}_2, \dots \rangle$ zwei allgemeine Algebren desselben Typs und gilt für eine Abbildung $\varphi : A \rightarrow B$, so daß für alle $a_1, a_2, \dots, a_{n_j} \in A$

$$\varphi(\omega_j(a_1, a_2, \dots, a_{n_j})) = \bar{\omega}_j(\varphi(a_1), \varphi(a_2), \dots, \varphi(a_{n_j})) \quad (j = 1, 2, \dots),$$

so bezeichnet man $\varphi : A \rightarrow B$ als **Homomorphismus**⁵⁷.

Ist ein Homomorphismus $\varphi : A \rightarrow B$ bijektiv, so heißt φ **Isomorphismus**⁵⁸ und die beiden Algebren A, B heißen **isomorph**⁵⁹.

Beispiel 6.111 Ein Homomorphismus $\varphi : G \rightarrow H$ zwischen zwei Gruppen G, H (aufgefaßt als allgemeine Algebren vom Typ $\langle 2, 0, 1 \rangle$) muß die Bedingungen

$$\begin{aligned} \varphi(a \circ b) &= \varphi(a) \star \varphi(b), \\ \varphi(e_G) &= e_H, \\ \varphi(a') &= \varphi(a)' \end{aligned}$$

erfüllen. Das sind aber genau die Eigenschaften, die auch ein *Gruppenhomomorphismus* erfüllt.

⁵⁵congruence

⁵⁶factor algebra

⁵⁷homomorphism

⁵⁸isomorphisms

⁵⁹isomorphic

Satz 6.112 (Allgemeiner Homomorphiesatz) *Sei $\langle A, \omega_1, \omega_2 \dots \rangle$ eine allgemeine Algebra und R eine Kongruenzrelation auf A . Dann ist $\overline{\varphi} : A \rightarrow A/R, a \mapsto K(a)$, ein Homomorphismus.*

Ist andererseits $\varphi : A \rightarrow B$ ein surjektiver Homomorphismus zwischen zwei allgemeinen Algebren A, B desselben Typs, so gibt es auf A eine Kongruenzrelation R , so daß B und A/R isomorph sind.

Homomorphismen und Kongruenzrelationen leisten daher dasselbe, um aus einer allgemeinen Algebra weitere Algebren (desselben Typs) zu konstruieren.

Kapitel 7

Elementare Aussagenlogik II

7.1 Konjunktive und disjunktive Normalformen für Formeln

Die in Kapitel 1 eingeführten aussagenlogischen Begriffe können nun präzisiert werden. Dazu bezeichne $\mathbf{B} = \{\mathbf{w}, \mathbf{f}\}$ die Menge aus den Symbolen $\mathbf{w}$ und $\mathbf{f}$.

Definition 7.1 Auf der Menge $\mathbf{B} = \{\mathbf{w}, \mathbf{f}\}$ werden die zweistelligen Operation **Konjunktion**¹ ($\wedge$), **Disjunktion**² ($\vee$), **Implikation**³ ($\rightarrow$), **Äquivalenz** ($\leftrightarrow$) und die einstellige Operation **Negation** ($\neg$), wie folgt, definiert:

$\wedge$	$\mathbf{f}$	$\mathbf{w}$	$\vee$	$\mathbf{f}$	$\mathbf{w}$	$\rightarrow$	$\mathbf{f}$	$\mathbf{w}$	$\leftrightarrow$	$\mathbf{f}$	$\mathbf{w}$
$\mathbf{f}$	$\mathbf{f}$	$\mathbf{f}$	$\mathbf{f}$	$\mathbf{f}$	$\mathbf{w}$	$\mathbf{f}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{f}$	$\mathbf{w}$	$\mathbf{f}$
$\mathbf{w}$	$\mathbf{f}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{f}$	$\mathbf{w}$	$\mathbf{w}$	$\mathbf{f}$	$\mathbf{w}$
$\neg$	$\mathbf{f}$	$\mathbf{w}$									
	$\mathbf{w}$	$\mathbf{f}$									

$\langle \mathbf{B}, \wedge, \vee, \rightarrow, \leftrightarrow, \neg \rangle$ ist daher eine allgemeine Algebra vom Typ $\langle 2, 2, 2, 2, 1 \rangle$.

Definition 7.2 Eine **Boolesche Variable** ist eine Variable, die Werte aus $\mathbf{B} = \{\mathbf{w}, \mathbf{f}\}$ annehmen kann.

Die Definitionen für (**wohlgeformte**) **Formel**, **Tautologie** etc. bleiben natürlich gleich. Das erste Ziel ist es, zu Formeln in Analogie zu Mengenausdrücken eine Normalform anzugeben.

¹conjunction

²disjunction

³implication

Genauso wie bei den Mengenausdrücken läßt sich mit Hilfe der Wahrheitstafel (die ja der Elementartabelle entspricht) zu jeder wohlgeformten Formel eine disjunktive resp. konjunktive Normalform bestimmen.

Für Boolesche Variable a verwenden wir die folgende Notation:

$$a^\delta = \begin{cases} a & \text{für } \delta = 1, \\ \neg a & \text{für } \delta = 0 \end{cases}$$

verwenden.

Definition 7.3 Seien $a_1, a_2, \dots, a_n$ Boolesche Variable. Für jede $m \times n$ -Matrix $J = (\delta_{kl})_{1 \leq k \leq m, 1 \leq l \leq n}$ ($0 \leq m \leq 2^n$) mit Elementen $\delta_{kl} \in \{0, 1\}$ mit paarweise verschiedenen Zeilen bezeichnet

$$\bigvee_{k=1}^m \left(\bigwedge_{l=1}^n a_l^{\delta_{kl}} \right)$$

eine **disjunktive Normalform**. Entsprechend heißt

$$\bigwedge_{k=1}^m \left(\bigvee_{l=1}^n a_l^{\delta_{kl}} \right)$$

konjunktive Normalform.

Satz 7.4 Sei F eine wohlgeformte Formel, in der nur die Booleschen Variablen $a_1, a_2, \dots, a_n$ vorkommen. Dann gibt es eine (bis auf die Reihenfolge der Terme) eindeutige zu F äquivalente disjunktive resp. konjunktive Normalform der Booleschen Variablen $a_1, a_2, \dots, a_n$.

Die entsprechenden Darstellungen können wie im Abschnitt 2.5 aus der Wahrheitstafel sofort abgelesen werden. Ersetzt man die Elementsymbole **w** durch 1 und **f** durch 0, so setzt sich die Matrix J für die disjunktive Normalform aus jenen Zeilen der ersten n Spalten zusammen, so in der resultieren Spalte 1 steht. Vertauscht man alle 0 und 1, so erhält man auf dieselbe Art und Weise die Matrix J für die konjunktive Normalform.

Weiters läßt sich der im Abschnitt 2.5 vorgestellte Algorithmus zur Bestimmung der disjunktiven Normalform direkt auf wohlgeformte Formeln anwenden

7.2 Boolesche Funktionen, Schaltalgebra

Definition 7.5 Sei $\mathbf{B} = \{\mathbf{w}, \mathbf{f}\}$ und n eine natürliche Zahl. Eine Funktion $f : \mathbf{B}^n \rightarrow \mathbf{B}$ heißt **Boolesche Funktion**⁴.

⁴Boolean function

Beispiel 7.6 Die Funktionen $f_1(a, b) = a \wedge b$ und $f_2(a, b) = a \vee b$ sind Boolesche Funktionen.

Man kann also mit Hilfe wohlgeformter Formeln sofort Boolesche Funktionen erzeugen. Interessanterweise gilt dies auch umgekehrt.

Satz 7.7 Jede Boolesche Funktion $f : \mathbf{B}^n \rightarrow \mathbf{B}$ ist durch eine wohlgeformte Formel darstellbar.

Eine solche wohlgeformte Formel kann direkt aus der Wertetabelle von f als disjunktive Normalform angegeben werden.

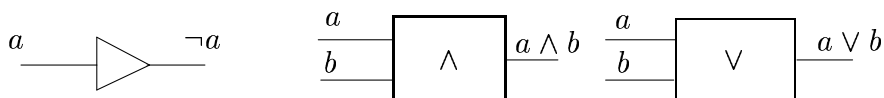
Beispiel 7.8 Ist $f : \mathbf{B}^2 \rightarrow \mathbf{B}$ durch

a	b	$f(a, b)$
w	w	f
w	f	w
f	w	w
f	f	f

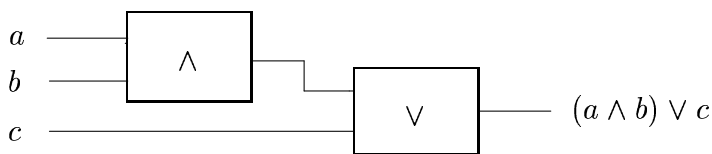
gegeben, so erhält man direkt

$$f(a, b) = (a \wedge \neg b) \vee (\neg a \wedge b).$$

Boolesche Funktionen treten in der Praxis oft bei Steuerungen auf und können durch **logische Schaltkreise** realisiert werden. Dabei werden die beiden logischen Werte **w**, **f** (wie in Digitalrechnern üblich) durch unterschiedliche Spannungszustände realisiert. Graphisch werden wir folgende Symbole verwenden:



Durch Kombination solcher **Gatter** erhält man ein logisches **Modul**, wie das z.B. das dargestellte Modul $(a \wedge b) \vee c$.



Im nächsten Abschnitt wird ein Algorithmus zur Optimierung der Anzahl der Gatter angegeben.

7.3 Optimierung von logischen Schaltkreisen

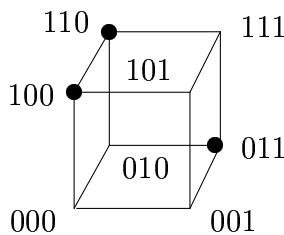
7.3.1 Quine-McCluskey-Algorithmus

Der Algorithmus von Quine-McClysky ist der erste von zwei Teilen, für eine vorgegebene Boolesche Funktion $f : \mathbf{B}^n \rightarrow \mathbf{B}$ eine Darstellung anzugeben, die mit möglichst wenigen Gattern als logische Schaltkreis realisiert werden kann.

Ausgangspunkt des Algorithmus ist die disjunktive Normalform von f . Der erste Schritt, ist die Zeilen der zugehörigen Matrix J , aufgefaßt als 0-1-Folgen der Länge n , aufzulisten. Beispielsweise gehören zu

$$w = (a \wedge b \wedge \neg c) \vee (a \wedge \neg b \wedge \neg c) \vee (\neg a \wedge b \wedge c)$$

die 3 0-1-Folgen 110, 100, 011. Die Idee ist, die 0-1-Folgen als Eckpunkte des Einheitswürfel zu interpretieren:



Die Terme der disjunktiven Normalform entsprechen dann gewissen Eckpunkten dieses Würfels. Die Eckpunkte werden im folgenden auch als **0-Würfel** bezeichnet werden. Weiters wird eine Kante zwischen zwei benachbarten Eckpunkten als **1-Würfel** bezeichnet. Beispielsweise wird der 1-Würfel zwischen den Eckpunkten 110, 100 durch $1y0$ dargestellt, y variiert sozusagen zwischen 0 und 1. Man sagt auch, daß der 1-Würfel $1y0$ die 0-Würfel 110, 100 überdeckt. Ebenso kann man die Seitenfläche der Eckpunkte 110, 100, 000, 010 durch eine sogenannten **2-Würfel** $yy0$ darstellen. Dieser überdeckt dann die 1-Würfel $1y0, y00, 0y0, y10$ und die 0-Würfel 110, 100, 000, 010. Der ganze Würfel wäre dann ein **3-Würfel** yyy .

In ganz analoger Weise definiert man **k-Würfel** ($0 \leq k \leq n$) im Zusammenhang mit 0-1-Folgen der Länge n .

Im obigen Beispiel enthält w also genau 3 0-Würfel 110, 100, 011 und einen 1-Würfel $1y0$. Klareweise kann man, wenn in der disjunktiven Normalform ein k -Würfel enthalten ist, die entsprechenden 2^k Terme durch einen einzigen ersetzen, wo nur mehr jene $n - k$ Boolesche Variable aufgenommen werden, die den 2^k Termen gemeinsam sind. Beispielsweise gilt

$$(a \wedge b \wedge \neg c) \vee (a \wedge \neg b \wedge \neg c) = a \wedge \neg c.$$

Dies motiviert den folgenden **Algorithmus von Quine-McClysky**. Es sei also w die disjunktive Normalform einer Booleschen Funktion in n Booleschen Variablen.

1. Man schreibe alle Zeilen der Matrix J (der disjunktiven Normalform) - also alle 0-Würfel von w - in eine nach der Anzahl der Einsen geordnete Liste und setze $k := 1$.
2. Man kombiniere alle Paare von $(k - 1)$ -Würfel. Falls bei einer Kombination ein k -Würfel erhalten wird, schreibt man ihn in eine Liste der k -Würfel. Danach wird jedes Paar von $(k - 1)$ -Würfel, das einen k -Würfel geliefert hat, gestrichen.
3. Ist $k < n$, so setze $k := k + 1$ und gehe zu 2.
4. Ist $k = n$, so entsprechen die nichtgestrichenen i -Würfel ($1 \leq i \leq n$) Termen p_j , deren Disjunktion äquivalent zu w ist.

Beispiel 7.9 Es sei

$$\begin{aligned} w = & (\neg a \wedge \neg b \wedge \neg c \wedge \neg d) \vee (\neg a \wedge b \wedge \neg c \wedge \neg d) \vee (\neg a \wedge b \wedge \neg c \wedge d) \\ & \vee (\neg a \wedge \neg b \wedge \neg c \wedge d) \vee (a \wedge \neg b \wedge \neg c \wedge d) \vee (a \wedge \neg b \wedge c \wedge d) \end{aligned}$$

die disjunktive Normalform einer wohlgeformten Formel mit 4 Booleschen Variablen a, b, c, d .

Im ersten Schritt werden die Zeilen von J , d.h. die 0-Würfel von w , aufgelistet:

0000
0100
0001
0101
1001
1011.

Als nächstes werden alle Paare dieser 0-Würfel betrachtet und es wird überprüft, ob daraus 1-Würfel gebildet werden können. Es können insgesamt 6 1-Würfel

0y00
000y
010y
0y01
y001
10y1

gebildet werden, die alle 0-Würfel von w überdecken, d.h. alle 6 0-Würfel werden gestrichen.

Daraufhin werden alle Paare dieser 1-Würfel untersucht, um festzustellen, ob daraus 2-Würfel gebildet werden können. Die ersten 4 1-Würfel 0y00, 000y, 010y, 0y01 werden vom 2-Würfel 0y0y überdeckt.

Insgesamt bleiben daher die drei Würfel y001, 10y1, 0y0y übrig, die den Termen

$$\begin{aligned} p_1 &= \neg b \wedge \neg c \wedge d, \\ p_2 &= a \wedge \neg b \wedge d, \\ p_3 &= \neg a \wedge \neg c \end{aligned}$$

entsprechen. w ist daher äquivalent zu

$$p_1 \vee p_2 \vee p_3 = (\neg b \wedge \neg c \wedge d) \vee (a \wedge \neg b \wedge d) \vee (\neg a \wedge \neg c).$$

Es wird sich herausstellen, daß diese Darstellung noch nicht optimal ist. w ist nämlich auch zu

$$p_2 \vee p_3 = (a \wedge \neg b \wedge d) \vee (\neg a \wedge \neg c)$$

äquivalent.

7.3.2 Hauptterme und absolut eliminierbare Terme

Zur weiteren Optimierung wird nun eine Matrix aufgestellt, wobei die Zeilen mit den p_j und die Spalten mit den 0–Würfeln von w indiziert werden. Dabei wird in die Matrix ein x gesetzt, wenn der 0–Würfel der Spalte vom Würfel der Zeile überdeckt wird.

Im oben besprochenen Beispiel ergibt sich folgendes Bild:

	0000	0100	0001	0101	1001	1011
0y0y	x	x	x	x		
10y1					x	x
y001				x	x	

Wenn es in einer Spalte nur ein x gibt, so muß jenes p_j der entsprechenden Zeile sicherlich in eine *Disjunktion*, die zu w äquivalent sein soll, aufgenommen werden. Ein Term mit dieser Eigenschaft heißt **Hauptterm**. Streicht man die Zeilen der Hauptterme und die Spalten der davon überdeckten 0–Würfel, so bleiben jede Spalten übrig, deren 0–Würfel von keinem Hauptterm überdeckt werden. Falls ein restlicher Term p_j keinen 0–Würfel überdeckt, so heißt er **absolut eliminierbar** und kann gestrichen werden.

Im obigen Beispiel sind 0y0y, 10y1 Hauptterme und y001 ist absolut eliminierbar. w ist daher zu

$$p_2 \vee p_3 = (a \wedge \neg b \wedge d) \vee (\neg a \wedge \neg c)$$

äquivalent.

Im allgemeinen sind nicht alle Terme p_j Hauptterme oder absolut eliminierbare Terme. Von den restlichen Termen p_j muß eine geeignete Konjunktion $p_{j_1} \wedge p_{j_2} \wedge \dots$ gebildet werden, die gemeinsam alle restlichen 0–Würfel überdecken und durch möglichst wenige Gatter realisiert werden kann.

Literaturverzeichnis

- [1] A. Gill, *Applied Algebra for the Computer Sciences*, Prentice Hall, Englewood Cliffs, New York, 1976.
- [2] M. Piff, *Discrete Mathematics, An Introduction for software engineers*, Cambridge University Press, Cambridge, 1991.

Index

- Abbildung, 31
- abelsche Gruppe, 65
- Abgeschlossenheit, 64
- Abstand zweier Knoten, 48
- abzählbare Mengen, 36
- adjazente Knoten, 38
- Adjazenzmatrix, 41
- algebraische Struktur, 64
- Algorithmus von Dijkstra, 56
- Algorithmus von Ford und Fulkerson, 59, 61
- Algorithmus von Kruskal, 55
- Algorithmus von Quine-McClyiskey, 93
- Algorithmus von Warshall, 29
- allgemeine Algebra, 87
- Allquantor, 6
- Allrelation, 23
- Alphabet, 79
- alternierende Gruppe, 36
- Anfangsknoten, 38
- antisymmetrische Relation, 24
- Äquivalenz, 2, 90
- Äquivalenzklasse, 24
- Äquivalenzrelation, 22
- äquivalente Formeln, 4
- Assoziativgesetz, 10, 64
- Aussage, 1
- Aussageform, 5
- Automorphismengruppe, 66
- azyklischer Graph, 43
- Baum, 47
- bewerteter Graph, 54
- Bijektion, 32
- bijektive Funktion, 32
- Bild eines Homomorphismus, 71
- Binärbaum, 49
- binäre Operation, 64
- binäre Relation, 21
- binärer Suchbaum, 50
- Binomialkoeffizient, 17
- Binomischer Lehrsatz, 18
- Blatt eines Baumes, 49
- Boolesche Algebra, 5, 86
- Boolesche Funktion, 91
- Boolesche Variable, 3, 90
- Boolescher Halbring, 73
- charakteristische Funktion, 11
- Code, 79
- Codewort, 79
- Datenstruktur, 50
- DeMorgansche Regel, 10, 86
- digitaler Suchbaum, 50
- direktes Produkt, 87
- Disjunktion, 2, 90
- disjunktive Normalform, 15, 91
- Distanz, 56
- distributiver Verband, 85
- Distributivgesetz, 10, 72, 85
- Division mit Rest, 76
- Durchschnitt von Mengen, 9
- ebener Graph, 52
- ebener Wurzelbaum, 49
- einfacher Graph, 39
- Einheit, 75
- Einheitengruppe, 75
- einstelliges Prädikat, 6
- Element, 8
- Elementtabelle, 13
- Endknoten, 38
- Endknoten eines Baumes, 49
- endliche Ordnung, 69

- endlicher Körper, 76
- Epimorphismus, 71
- equivalence, 90
- erfüllbare Formel, 4
- Erreichbarkeitsrelation, 28
- Euklidischer Algorithmus, 76
- Eulersche Linie, 54
- Eulersche Phifunktion, 75
- Eulersche Polyederformel, 52
- Existenzquantor, 6
- externer Knoten eines Baumes, 49

- Faktoralgebra, 88
- Faktorgruppe, 70
- Fakultät, 18
- Fixpunkt einer Permutation, 34
- Fluß, 59
- Formel, 4, 6
- Funktion, 31

- gültige Formel, 4
- Gegenstandsvariable, 5
- Generatormatrix, 81
- Gerüst, 55
- gerade Permutation, 36
- gerichtete Kante, 38
- gerichteter Graph, 38
- geschlossene Eulersche Linie, 54
- geschlossene Hamiltonsche Linie, 54
- geschlossene Kantenfolge, 40
- Gewicht eines Wortes, 80
- Gleichheitsrelation, 23
- Grad eines Polynoms, 74
- Graph, 38
- Graph einer Relation, 22
- Gruppe, 65
- Gruppencode, 80
- Gruppenhomomorphismus, 71
- Gruppenisomorphismus, 71
- Gruppenordnung, 68
- Gruppoid, 65

- Halbgruppe, 65
- Halbordnung, 24
- Halbring, 72

- Hamiltonsche Linie, 54
- Hammingdistanz, 79
- Hassediagramm, 25, 84
- Hauptterm, 95
- Hingrad, 39
- Homomorphiesatz, 72, 89
- Homomorphismus, 71, 88

- identische Funktion, 33
- Implikation, 2, 90
- Index einer Untergruppe, 68
- Indexmenge, 10
- Infimum, 84
- Injektion, 32
- injektive Funktion, 32
- Integritätsbereich, 74
- interner Knoten eines Baumes, 49
- inverse Funktion, 33
- inverse Permutation, 33
- inverses Element, 65
- Inzidenzmatrix, 42
- irreduzibles Polynom, 78
- isomorphe Algebren, 88
- isomorphe Graphen, 41
- Isomorphismus, 71, 88

- Junktor, 3

- Körper, 75
- Kante, 38
- Kantenfolge, 40
- Kantenzug, 40
- Kapazität eines Schnittes, 60
- Kardinalität, 36
- Kardinalzahl, 36
- Kartesische Darstellung einer Relation, 22
- kartesisches Produkt von Mengen, 20
- Kern eines Homomorphismus, 71
- kleiner Fermatscher Satz, 75
- Knoten, 38
- Knotenbasis, 46
- Knotengrad, 39
- kommutativer Ring, 73
- kommutative Gruppe, 65

- kommutatives Monoid, 65
- Kommutativgesetz, 10, 65
- Komplement einer Menge, 10
- Komplexprodukt, 70
- Komponente, 44
- Komposition von Relationen, 26
- Kongruenzrelation, 88
- Konjunktion, 1, 90
- konjunktive Normalform, 15, 91
- Korrekturschema, 80
- Kreis, 41

- Länge einer Kantenfolge, 40
- leere Kantenfolge, 40
- leere Menge, 8
- Linearcode, 81
- Linksnebenklasse, 67
- logischer Schaltkreis, 92

- Mächtigkeit von Mengen, 36
- Markierungsalgorithmus, 43
- Matrix einer Relation, 26
- Matrix-Baum-Theorem, 56
- maximaler spannender Baum, 55
- maximales Gerüst, 55
- Mehrfachkante, 39
- mehrstelliges Prädikat, 6
- Menge, 8
- Mengendifferenz, 10
- Mengenlehre, 8
- Minimaldistanz, 79
- minimaler spannender Baum, 55
- minimales Gerüst, 55
- Monoid, 65
- Monomorphismus, 71
- Multimenge, 9

- Nachbar eines Knotens, 39
- Nachfolger, 59
- Nachfolger eines Knoten, 39
- Nachrichtenwort, 79
- Nebenklasse, 67
- Nebenklassenanführer, 80
- Negation, 2
- Netzwerk, 54

- neutrales Element, 65
- nicht erfüllbare Formel, 4
- Normalteiler, 70
- Nullteiler, 74

- offene Eulersche Linie, 54
- offene Hamiltonsche Linie, 54
- öffentlicher Schlüssel, 82
- Operationstafel, 67
- Ordnung einer Gruppe, 68
- Ordnung eines Elements, 69

- paarer Graph, 53
- parity bit, 78
- partielle Funktion, 31
- Partition, 23
- Pascalsches Dreieck, 18
- Patricia Trie, 51
- Permutation, 33
- Pfeildiagramm einer Relation, 22
- planarer Graph, 52
- Polynom über einem Ring, 73
- Polynomring, 73
- Potenzmenge, 17
- Potenzreihe, 74
- Pozent eines Elements, 68
- Prädikatenlogik, 5
- Produkt von Gruppen, 69
- Produkt von Permutationen, 33

- Quantor, 6
- Quelle, 59

- Rechenregeln für Mengen, 10
- Rechtsnebenklasse, 67
- Reduktion eines Graphen, 46
- reflexiv-transitive Hülle, 29
- reflexive Relation, 23
- Relation, 21
- Restklasse, 71
- Ring, 73
- RSA-Verfahren, 82

- Satz von Kuratowski, 53
- Satz von Lagrange, 68
- Schaltalgebra, 92

- schlichter Graph, 39
- Schlinge, 34, 38
- Schnitt, 60
- schwach zusammenhängender Graph, 46
- schwache Zusammenhangskomponente, 46
- Senke, 59
- spannender Baum, 55
- spannender Wald, 55
- stark zusammenhängender Graph, 45
- starke Zusammenhangskomponente, 45
- Supremum, 84
- Surjektion, 32
- surjektive Funktion, 32
- Symmetriegruppe, 66
- symmetrische Differenz, 11
- symmetrische Gruppe, 33, 66
- symmetrische Relation, 23

- Tautologie, 4
- Teilgraph, 40
- Teilmenge, 8
- Totalordnung, 25
- transitive Hülle, 28
- transitive Relation, 23
- Transposition, 35
- Trie, 51
- Typ einer Algebra, 87

- überabzählbare Mengen, 37
- ungerade Permutation, 36
- ungerichtete Kante, 38
- ungerichteter Graph, 38
- Universum, 8
- Unteralgebra, 87
- Untergruppe, 67
- Unterteilung, 53

- Venn diagramm, 11
- Verband, 83
- Vereinigung von Mengen, 9
- Verschlüsselung, 82
- Verschmelzungsgesetz, 10, 83
- vollständiger Graph, 53
- vollständiger paarer Graph, 53

- Vorgänger, 39, 59

- Wahrheitstafel, 4
- Wahrheitswert, 1
- Wald, 47
- Weg, 40
- Weggrad, 39
- Wiederholungscode, 79
- wohlgeformte Formel, 4
- Wurzel eines Baumes, 48, 49
- Wurzelbaum, 49

- Zerlegung, 23
- zusammenhängender Graph, 44
- Zusammenhangskomponente, 44
- zweiwertige Aussagenlogik, 1
- Zyklen einer Permutation, 34
- Zyklus, 41

Übungsbeispiele

1) Sei a die Aussage *Es gibt eine größte natürliche Zahl.* und b die Aussage *0 ist die größte natürliche Zahl.* Man entscheide, ob die Aussagen $a \rightarrow b$ bzw. $b \rightarrow a$ wahr oder falsch sind.

2–5) Man beweise die folgenden Äquivalenzen mittels Wahrheitstafeln:

2) $a \wedge (b \vee c) \iff (a \wedge b) \vee (a \wedge c).$

3) $a \vee (a \wedge b) \iff a.$

4) $a \leftrightarrow b \iff (a \rightarrow b) \rightarrow \neg(b \rightarrow a).$

5) $\neg(a \rightarrow b) \iff a \wedge \neg b.$

6–9) Beweisen Sie die folgenden Beziehungen mit Hilfe von Elementtafeln oder geben Sie ein konkretes Gegenbeispiel an.

6) $A \cap (B \cap C) = (A \cap B) \cap C.$

7) $(A \setminus B) \setminus C = A \setminus (B \setminus C).$

8) $M \setminus (A \cup B) = (M \setminus A) \cap (M \setminus B).$

9) $(A \cup B) \cap (B \cup C)' \subseteq A \cap B'$

10–11) Beweisen Sie die folgenden Beziehungen durch Rückführung auf die Normalform oder geben Sie ein konkretes Gegenbeispiel an.

10) $A \cap (A \cup B) = A.$

11) $A \triangle (B \cap C) = (A \triangle B) \cap (A \triangle C).$

12–13) Beweisen Sie die folgenden Beziehungen mit Hilfe der entsprechenden charakteristischen Funktionen oder geben Sie ein konkretes Gegenbeispiel an.

12) $A \cup (B \cap C) = (A \cup B) \cap (A \cup C).$

13) $A \cap (B \triangle C) = (A \cap B) \triangle (A \cap C)$.

14) Sei M eine nichtleere endliche Menge. Zeigen Sie, daß M gleich viele Teilmengen mit gerader Elementanzahl wie solche mit ungerader Elementanzahl besitzt, indem Sie ein Verfahren angeben, das aus den Teilmengen der einen Art umkehrbar eindeutig die der anderen Art erzeugt.

15) Man bestimme zu $(A \setminus B)' \cap (A \cup C)$ die disjunktive und konjunktive Normalform

16) Beweisen oder widerlegen Sie die Beziehung

$$(A \times B) \cap (B \times A) = (A \cap B) \times (A \cap B).$$

17) Wie verhält sich der Ausdruck $\binom{3n}{n}$ für $n \rightarrow \infty$?

18) Man beweise die Beziehung $\binom{n+1}{k+1} = \binom{n}{k+1} + \binom{n}{k}$ durch Interpretation von $\binom{n}{k}$ als Anzahl der k -elementigen Teilmengen einer n -elementigen Menge.

19) Man beweise die Beziehung $\binom{n+1}{k+1} = \binom{n}{k+1} + \binom{n}{k}$ mit Hilfe der Formel $\binom{n}{k} = \frac{n!}{k!(n-k)!}$.

20) Die Relation R sei für $m, n \in \{2, 3, 4, 5, 7, 8, 11, 13\}$ definiert durch $mRn : \iff m + n$ ungerade oder $m = n$. Man stelle R durch einen gerichteten Graphen und auch im cartesischen Koordinatensystem dar.

21) Untersuchen Sie, ob die Relation aus Aufgabe 20 reflexiv, symmetrisch bzw. transitiv ist.

22) Wie 21), jedoch für die Relation $mRn : \iff \text{ggT}(m, n) = 1, m, n \in \{1, 2, 3, \dots\}$.

23) Wie 21), jedoch für die Relation $xRy : \iff x - y \in \mathbb{Z}, x, y \in \mathbb{R}$.

24) Sei A eine Menge. Man zeige, daß die Relation $BRC \iff B \subseteq C$ auf der Potenzmenge $\mathfrak{P}(A)$ eine Halbordnung ist.

25) Sei $\preceq$ eine Halbordnung auf einer Menge A . Man zeige, daß dann durch

$$(a, b) \preceq_2 (c, d) \iff (a \preceq c \text{ und } a \neq c) \text{ oder } (a = c \text{ und } b \preceq d)$$

eine Halbordnung auf $A \times A$ definiert wird (lexikographische Ordnung).

26) Gegeben sei die Menge $\{1, 2, \dots, 8\}$ und die Relation $mRn \iff m|n$ (m teilt n). Zeigen Sie, daß R eine Halbordnung ist und zeichnen Sie das Hasse-Diagramm.

27) Bestimmen Sie die Matrizen zur Relation R aus 26).

28) Zeigen Sie: Eine Relation $R \subseteq A \times A$ ist symmetrisch genau dann, wenn $R^{-1} \subseteq R$ gilt und transitiv genau dann, wenn $R^2 \subseteq R$ gilt.

29) Seien S bzw. R die folgenden binären Relationen auf $\mathbb{Z}$:

$$S = \{(a^2, a), a \in \mathbb{Z}\}, R = \{(a^3, a), a \in \mathbb{Z}\}.$$

Untersuchen Sie die beiden Relationen auf Reflexivität, Symmetrie und Transitivität und bestimmen Sie die Relationen $S \circ R$ sowie $R \circ S$.

30) Geben Sie ein Beispiel für zwei binäre Relationen R bzw. S auf der Menge $\{1, 2, 3\}$ an, für die gilt: $S \circ R \neq R \circ S$.

31) Bestimmen Sie die Relation R^+ zur folgenden Relation R auf $\{a, b, c\}$ mit Hilfe der Matrix von R :

$$aRb, cRa, bRb, cRc.$$

32) Gegeben sei die Menge $\{a_1, a_2, \dots, a_n\}$, sowie die folgende Relation R :

$$a_1Ra_2, a_2Ra_3, \dots, a_{n-1}Ra_n.$$

Bestimmen Sie die transitive Hülle R^+ mit Hilfe der Matrix zu R .

33) Bestimmen Sie alle binären Relation R auf einer Menge A , die reflexiv, symmetrisch, transitiv und antisymmetrisch sind.

34) Seien R_1 und R_2 Äquivalenzrelationen auf einer Menge A . Man zeige, daß der Durchschnitt $R_1 \cap R_2$ wieder eine Äquivalenzrelation ist. Gilt dasselbe auch für die Vereinigung?

35) Man bestimme alle Äquivalenzrelationen auf der Menge $A = \{1, 2, 3, 4\}$. Sei $E(A)$ die Menge dieser Äquivalenzrelationen. Man zeige weiters daß für $R_1, R_2 \in E(A)$ durch $R_1 \subseteq R_2$ eine Halbordnung definiert wird und bestimme das Hassediagramm der Halbordnung $(E(A), \subseteq)$.

36) Bestimmen Sie die Relation R^+ zur folgenden Relation R auf $\{a, b, c, d\}$ mit Hilfe des Algorithmus von Warshall:

$$aRb, cRa, bRb, cRc, dRc.$$

37) Bestimmen Sie in Aufgabe 27) R^+ mit dem Algorithmus von Warshall.

38–40) Untersuchen Sie, ob es sich bei den folgenden Relationen $R \subseteq A \times B$ um partielle Funktionen, Funktionen, injektive Funktionen, surjektive Funktionen bzw. bijektive Funktionen handelt.

38) $R = \{(x^2, \frac{1}{x^2}) \mid x \in \mathbb{R}_+\}$, $A = B = \mathbb{R}$.

39) Wie 29) jedoch $A = B = \mathbb{R}_+$.

40) $R = \{(\log_2 x, x) \mid x \in \mathbb{R}_+\}$, $A = B = \mathbb{R}$.

41) Untersuchen Sie, ob π eine Permutation festlegt und geben Sie gegebenenfalls den Graphen, die Zykldarstellung, sowie die Zykldarstellung ohne Klammern an:

$$\pi(k) = 4k + 2 \bmod 10, \quad 0 \leq k \leq 9.$$

42) Schreiben Sie π als Produkt von Zweierzyklen:

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

43) Sei eine Permutation π in zweizeiliger Darstellung gegeben. Unter der Inversionstafel von π versteht man die Folge $(b_1, b_2, \dots, b_n)$, wobei $b_k \geq 0$ angibt, wieviele größere Zahlen in der zweiten Zeile links vom Element k stehen. Wählen wir z.B. π aus 42), so lautet die Inversionstafel $(7, 7, 0, 5, 0, 2, 1, 0, 0)$.

Wie kann man bei Kenntnis der Inversionstafel die Permutation rekonstruieren? Demonstrieren Sie ein geeignetes Verfahren am obigen Beispiel.

44) Man untersuche, ob die Funktionen $f(x) = x^2 \bmod 10$ bzw. $g(x) = x^3 \bmod 10$ auf der Menge $\{0, 1, \dots, 9\}$ bijektiv sind, d.h. Permutationen festlegen!

45) Zeigen Sie, daß $\mathbb{N} \times \mathbb{N} \times \mathbb{N}$ abzählbar ist.

46) Zeigen Sie, daß die Menge aller unendlichen 0-1-Folgen überabzählbar ist

47) Man zeige mittels eines geeigneten graphentheoretischen Modells, daß es in jeder Stadt mindestens zwei Bewohner mit der gleichen Anzahl von Nachbarn gibt.

48) 7 Freunde vereinbaren, daß jeder an 3 andere Ansichtskarten schickt. Man überlege an einem geeigneten graphentheoretischen Modell, ob das so geschehen kann, daß jeder nur denjenigen schreibt, die ihm auch geschrieben haben.

49) Unter n Mannschaften wird ein Turnier ausgetragen und es haben insgesamt schon mindestens $n + 1$ Spiele stattgefunden. Zeigen Sie mit Hilfe eines geeigneten graphentheoretischen Modells, daß mindestens eine Mannschaft bereits an mindestens 3 Spielen teilgenommen hat.

50) Konstruieren Sie, wenn möglich einen ungerichteten Graphen mit den Graden

a) 2, 2, 3, 3, 4, 4

b) 2, 3, 3, 4, 4, 4

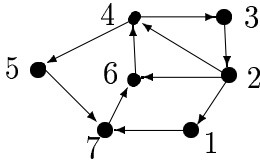
c) 2, 3, 3, 3, 4, 4

51) Ein schlichter Graph $G = (V, E)$ heißt kubisch, wenn jeder Knoten $v \in V$ Knotengrad $d(v) = 3$ hat.

- Geben Sie ein Beispiel für einen kubischen Graphen mit $\alpha_0(G) = 6$ an!
- Gibt es einen kubischen Graphen mit ungerader Knotenanzahl $\alpha_0(G)$?
- Zeigen Sie, daß es zu jedem $n \geq 2$ einen kubischen Graphen mit $\alpha_0(G) = 2n$ gibt!

52) Sei $G = (V, E)$ ein Graph mit $0 < \alpha_1(G) < \alpha_0(G)$. Zeigen Sie, daß mindestens ein Knoten $v \in V$ mit $d(v) \leq 1$ existiert!

53–54) Sei G der folgende gerichtete Graph:

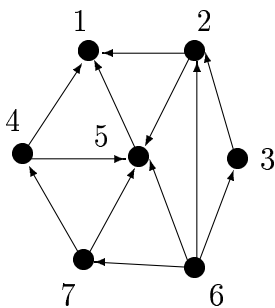
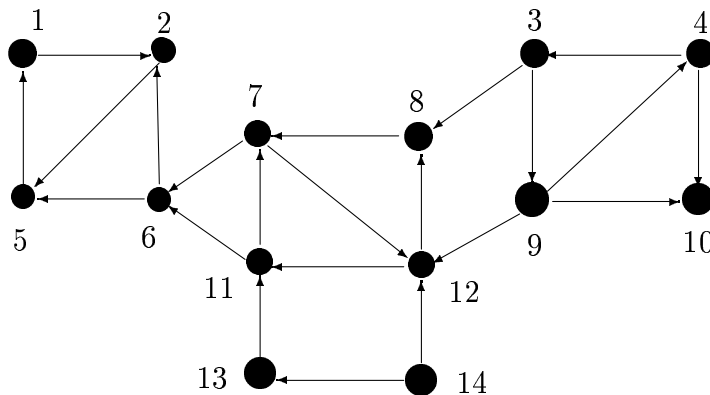


53) Bestimmen Sie die Adjazenzmatrix $\mathfrak{A}(G)$, sowie (mit deren Hilfe) die Anzahl der gerichteten Kantenfolgen der Länge 3 von 4 nach 6.

54) Bestimmen Sie im Graphen G die Anzahl der Zyklen der Länge 3, auf denen der Knoten 4 liegt. (Ein Zyklus ist eine gerichtete Kantenfolge, bei der Anfangs- und Endknoten übereinstimmen, alle anderen auftretenden Knoten aber nur einmal vorkommen.)

55) Zeigen Sie mit Hilfe des Algorithmus von Warshall, daß der Graph G aus 53–54) stark zusammenhängend ist

56) Untersuchen Sie den Graphen H_1 bzw. H_2 mit Hilfe des Markierungsalgorithmus auf Azyklizität.

 H_1  H_2 

57) Man bestimme im Graphen H_2 aus 56) alle Knotenbasen.

58) Sei H_3 der “inverse” Graph von H_2 , d.h. alle Kantenrichtungen werden umgedreht. Bestimmen Sie alle Knotenbasen von H_3 .

59) Zeigen Sie, daß es in jedem Baum T mit $\alpha_0(T) > 1$ wenigstens zwei Endknoten gibt. (Ein Endknoten $v \in V(T)$ hat Knotengrad $d(v) = 1$.)

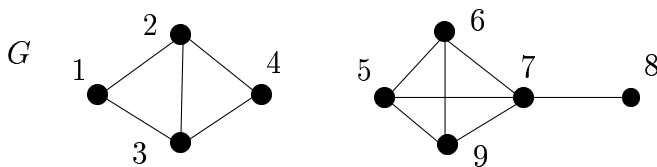
60) Zeigen Sie: Hat ein Baum T einen Knoten $v \in V(T)$ mit Knotengrad $d(v) = d$, so hat T wenigstens d Endknoten.

61) Sei W ein Wald mit k Komponenten. Welcher Zusammenhang besteht dann zwischen $\alpha_0(W)$ und $\alpha_1(W)$?

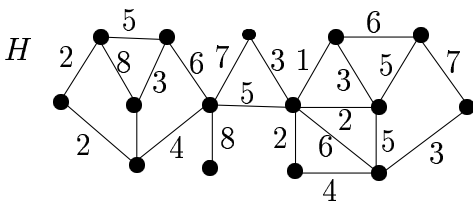
62) Man zeige, daß in einem Baum T jede Kante eine Brücke ist. (Eine Kante $e \in E$ heißt Brücke, wenn der Graph $G \setminus e = (V(G), E(G) \setminus \{e\})$ mehr Zusammenhangskomponenten hat als G .)

63) Sei G ein schlichter Graph mit $\alpha_0(G) > 4$. Man zeige, daß dann entweder G oder G^c einen Kreis enthält. (G^c ist der komplementäre Graph zu G , d.h. G^c enthält die selben Knoten wie G und alle Kanten $(v, w) \in V(G) \times V(G)$, $v \neq w$, die nicht in $E(G)$ enthalten sind.)

64) Man bestimme im folgenden Graphen G mit Hilfe des Matrix-Baum-Theorems die Anzahl der Gerüste auf G .



65) Man bestimme im folgenden Graphen H mit Hilfe des Kruskalalgorithmus einen minimalen und einen maximalen spannenden Baum.



66) Gegeben seien die 6 Daten $A, \dots, F$ durch die Schlüsselfolgen $A = 010010 \dots$, $B = 011101 \dots$, $C = 101100 \dots$, $D = 101011 \dots$, $E = 100110 \dots$, $F = 001010 \dots$. Konstruieren Sie den zugehörigen Trie, Patricia Trie und Digitalen Suchbaum.

67) Die *externe Pfadlänge* eines Tries ist die Summe der Abstände von der Wurzel zu allen *besetzten* Endknoten des Baumes, wobei die Abstände in der Anzahl der Kanten

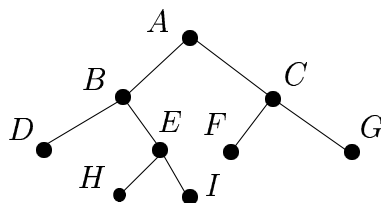
auf dem entsprechenden Weg gemessen werden. Wie groß ist die externe Pfadlänge eines Tries, der N Daten enthält, mindestens? (Hinweis: Welche Gestalt des Binärbaums führt zu kleiner Pfadlänge?)

68) Ein t -ärer Baum ($t \in \mathbb{N}$, $t \geq 2$) ist ein ebener Wurzelbaum, bei dem jeder Knoten entweder 0 Nachfolger (Endknoten) oder genau t Nachfolger (interner Knoten) hat. Für $t = 2$ ergeben sich also genau die Binärbäume. Wieviele Endknoten hat ein t -ärer Baum mit n internen Knoten?

69) Zum *Abarbeiten* der Knoten eines Binärbaumes verwendet man gerne rekursive Algorithmen, die in wohldefinierter Reihenfolge die folgenden Schritte ausführen:

- (1) Bearbeite den aktuellen Knoten,
- (2) Gehe zur Wurzel des linken Nachfolgebaums des aktuellen Knotens,
- (3) Gehe zur Wurzel des rechten Nachfolgebaums des aktuellen Knotens.

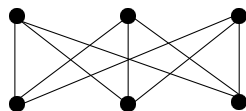
Am Beginn steht man bei der Wurzel des Gesamtbaumes. Führt man die genannten Schritte (1) bis (3) rekursiv in der angegebenen Reihenfolge aus, so spricht man von *Präordertraversierung*. Beim untenstehenden Baum werden die Knoten also in folgender Reihenfolge bearbeitet: $A, B, D, E, H, I, C, F, G$. Wie ändert sich diese Reihenfolge, wenn man im Algorithmus jeweils die Abfolge (2)(1)(3) nimmt (*Inordertraversierung*), wie wenn man die Abfolge (2)(3)(1) wählt (*Postordertraversierung*)?



70) Zeigen Sie: Für einen einfachen, zusammenhängenden, planaren, Graphen G , der keine Kreise mit Länge ≤ 3 besitzt gilt: $\alpha_1(G) \leq 2\alpha_0(G) - 4$.

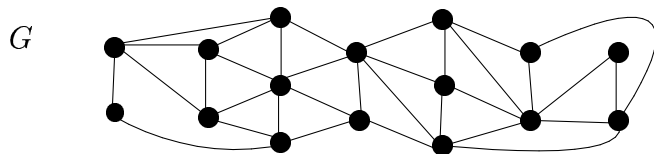
71) Zeigen Sie mit Hilfe von 70): Der vollständige paare Graph $K_{3,3}$ ist nicht planar.

$K_{3,3}$



72) Man untersuche, ob der folgende Graph G (siehe nächste Seite) eine Eulersche Linie besitzt und bestimme gegebenenfalls eine!

73) Sei G ein einfacher, ungerichteter Graph. Dann wird der *line graph* G^* zu G folgendermaßen definiert: $V(G^*) = E(G)$, und (e, f) ist in $E(G^*)$, genau dann, wenn im Graphen G die Kanten e und f einen gemeinsamen Knoten haben. Zeichnen Sie für einen selbstgewählten Graphen G mit 6 Knoten und 10 Kanten den line graph.



74) Zeigen Sie: Ist G ein einfacher, ungerichteter Graph und geschlossen Eulersch, so ist der line graph G^* Hamiltonsch und Eulersch.

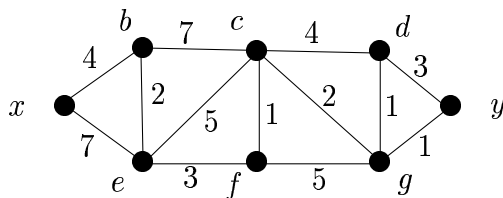
75) Für welche m, n besitzt der vollständige paare Graph $K_{m,n}$ eine geschlossene Hamiltonsche Linie? (Die Knotenmenge V eines vollständigen paaren Graphen $K_{m,n}$ besteht aus 2 disjunkten Teilmengen V_1, V_2 mit $|V_1| = m$ und $|V_2| = n$ und die Kantenmenge E besteht aus allen ungerichteten Kanten (v_1, v_2) mit $v_1 \in V_1$ und $v_2 \in V_2$.)

76) Zeigen Sie mit Hilfe eines graphentheoretischen Modells: Auf jeder Party mit 6 Leuten gibt es 3, die miteinander bekannt sind, oder 3, die sich gegenseitig nicht kennen.

77) Zeigen Sie: Jeder Baum ist ein paarer Graph. (Ein ungerichteter Graph G ist ein *paarer Graph*, wenn die Knotenmenge $V(G)$ in zwei disjunkte, nichtleere Teilmengen V_1, V_2 zerlegt werden kann, so daß es nur Kanten $(v_1, v_2) \in E(G)$ mit $v_1 \in V_1$ und $v_2 \in V_2$ gibt.)

78) Zeigen Sie: Ein ungerichteter Graph G ist genau dann ein paarer Graph, wenn jeder Kreis in G gerade Länge hat.

79) Bestimmen Sie mit dem Algorithmus von Dijkstra einen kürzesten Weg zwischen den Knoten x und y im folgenden Graphen:

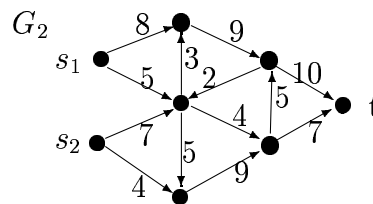
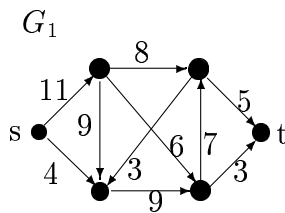


80) Ein Automorphismus φ eines Graphen G ist ein Isomorphismus von G mit sich selbst, d.h. eine bijektive Abbildung $\varphi : V(G) \rightarrow V(G)$, mit der Eigenschaft $(x, y) \in E(G) \Leftrightarrow (\varphi(x), \varphi(y)) \in E(G)$.

Beschreiben Sie alle Automorphismen eines vollständigen paaren Graphen $K_{m,n}$ mit $m \neq n$.

81) Bestimmen Sie im Netzwerk G_1 mit Hilfe des Algorithmus von Ford und Fulkerson einen maximalen Fluß!

82) Bestimmen Sie im Netzwerk G_2 mit zwei Quellen s_1, s_2 mit Hilfe des Algorithmus von Ford und Fulkerson einen maximalen Fluß, indem sie zunächst G_2 durch eine *virtuelle Quelle* s , die vor den Quellen s_1, s_2 liegt, ergänzen.



83) Sei G ein gerichtetes Netzwerk, bei dem auch die Knoten eine *Kapazität* $\bar{w}$ besitzen, d.h. derjenige Teil eines gültigen Flusses ϕ , der über einen Knoten $x \in V(G)$ fließt, ist durch $\bar{w}(x)$ beschränkt. Wie muß man ein Netzwerk G modifizieren, so daß man den Algorithmus von Ford und Fulkerson anwenden kann, um einen maximalen Fluß zu finden, der auch die Knotenkapazitäten respektiert?

84–85) Untersuchen Sie, welche der Eigenschaften (1)–(5) die folgenden algebraischen Strukturen $\langle M, \circ \rangle$ haben:

84) $\Sigma = \{a, b\}$, M die Menge $\{\epsilon\}$ vereinigt mit der Menge aller derjenigen Wörter in Σ^* , die mit a beginnen, und bei denen auf den Buchstaben b nur a folgen darf, während auf a sowohl a als auch b folgen darf, $\circ$ das Aneinanderhängen der Wörter.

85) M die Menge aller Permutationen der Menge $\{1, 2, \dots, n\}$, die sich als Produkt einer geraden Anzahl von Transpositionen (Zweierzyklen) schreiben lassen, $\circ$ die Hintereinanderausführung der Permutationen.

86) M die Menge aller komplexen Zahlen vom Betrag 1, $\circ$ die Multiplikation.

87) $M = \mathbb{R}$, $a \circ b = \frac{a+b}{1-ab}$.

88) Man zeige, daß $\langle \mathbb{Z}, \circ \rangle$ mit der Operation $a \circ b = a + b - ab$ eine Halbgruppe ist. Gibt es ein neutrales Element? Wenn ja, welche Elemente haben Inverse?

89) Sei G die Menge aller $n \times n$ -Matrizen A mit reellen Einträgen und $\det A \neq 0$. Man zeige, daß G mit der gewöhnlichen Matrizenmultiplikation eine Gruppe bildet.

90) Sei U_1 die Menge jener $n \times n$ -Matrizen A aus Beispiel 89) mit $\det A > 0$ und U_2 die Menge jener Matrizen A mit $\det A = 1$. Man zeige daß U_1 und U_2 Untergruppen der Gruppe G (aus Beispiel 89) sind.

91) Man zeige, daß eine nichtleere Teilmenge U einer endlichen Gruppe G genau dann Untergruppe von G ist, wenn für alle $a, b \in U$ auch $a \circ b \in U$ gilt.

92) Sei G die Menge der Permutationen id , (13) , (24) , $(12)(34)$, $(13)(24)$, $(14)(23)$, (1234) , (1432) . Man veranschauliche G , indem man die Permutationen auf die vier Eckpunkte eines Quadrats wirken lasse und als geometrische Operationen interpretiere. Man zeige

mit Hilfe dieser Interpretation, daß G eine Untergruppe der symmetrischen Gruppe S_4 ist. (Symmetriegruppe des Quadrats!)

93) Man bestimme alle Untergruppen der Gruppe G aus Beispiel 92).

94) Sei $U_1 = \langle (1234) \rangle$ die von (der Drehung) (1234) erzeugte Untergruppe von G (aus Beispiel 89) und $U_2 = \langle (13) \rangle$ die von (der Spiegelung) (13) erzeugte Untergruppe von G . Man bestimme die Links- und die Rechtsnebenklassenzerlegungen von U_1 und die von U_2 .

95) Seien $\langle G, \circ \rangle$ sowie $\langle H, * \rangle$ Gruppen. Dann ist das *direkte Produkt* die Menge $G \times H$ mit der Operation $(g_1, h_1) \cdot (g_2, h_2) = (g_1 \circ g_2, h_1 * h_2)$. Zeigen Sie, daß das direkte Produkt wieder eine Gruppe ist.

96) Seien U_1, U_2 Untergruppen von einer Gruppe G . Zeigen Sie, daß dann auch $U_1 \cap U_2$ eine Untergruppe ist.

97) Seien U_1, U_2 Untergruppen von einer Gruppe G . Zeigen Sie, daß $U_1 \cup U_2$ nur dann wieder eine Untergruppe von G ist, wenn $U_1 \subseteq U_2$ oder $U_2 \subseteq U_1$ gilt.

98) Bestimmen Sie alle Untergruppen der Gruppe der Restklassen modulo 4 mit der Addition. Welche sind Normalteiler?

99) Sei $G = \mathbb{Z}_4^3$ die Menge aller Tripel von Elementen aus $\mathbb{Z}_4$ und $+$ die komponentenweise Addition modulo 4. Zeigen Sie, daß G eine kommutative Gruppe ist.

100) Zeigen Sie, daß die Menge $U = \{000, 123, 202, 321\}$ eine Untergruppe der Gruppe G aus Beispiel 99) bildet. Bestimmen Sie weiters die Nebenklassen von U in G .

101) Man zeige, daß die symmetrische Gruppe S_n (das sind alle Permutationen der Zahlen $1, 2, \dots, n$ mit der Hintereinanderausführung) für $n \geq 3$ nicht kommutativ ist, indem man Permutationen $\pi, \sigma \in S_n$ mit $\pi \circ \sigma \neq \sigma \circ \pi$ angebe.

102) Sei A eine nichtleere Menge. Man zeige, daß $\langle \mathfrak{P}(A), \Delta \rangle$ eine Gruppe ist.

103) Sei $\varphi : G \rightarrow H$ ein Gruppenhomomorphismus. Man zeige, daß das Bild

$$\varphi(G) = \{b \in H \mid \text{es existiert } a \in G \text{ mit } \varphi(a) = b\}$$

eine Untergruppe von H ist.

104) Sei G eine endliche Gruppe. Man zeige, daß für alle $a \in G$ die Eigenschaft $a^{|G|} = 1$ gilt.

(Anleitung: Man beachte, daß die Ordnung der von a erzeugten Untergruppe $\langle a \rangle$ ein Teiler von $|G|$ ist.)

105) Sei $\langle G, \circ \rangle$ eine Gruppe und

$$Z = \{x \in G \mid \text{für alle } a \in G \text{ gilt } a \circ x = x \circ a\}$$

das *Zentrum* von G . Man zeige, daß Z ein Normalteiler von G ist.

106) Lösen Sie die folgenden Kongruenzen (dh. Gleichungen in Restklassen) bzw. beweisen Sie die Unlösbarkeit:

a) $7x \equiv 3 \pmod{9}$

b) $6x \equiv 4 \pmod{9}$

c) $6x \equiv 3 \pmod{9}$.

107) Bestimmen Sie alle invertierbaren Elemente in der Menge der Restklassen modulo 12 bezüglich der Multiplikation und stellen Sie die Gruppentafel auf!

108) Bestimmen Sie mit dem Euklidischen Algorithmus $\text{ggT}(209, 77)$.

109) Bestimmen Sie mit Hilfe der Umkehrung des Euklidischen Algorithmus $41^{-1} \pmod{61}$, d.h. $41 \cdot 41^{-1} \equiv 1 \pmod{61}$.

110) Man verschlüssele und entschlüssele die Nachricht *GEOD* gemäß des RSA-Verfahrens mit $n = 47 \cdot 59 = 2773$ und einem geeignet gewählten d , das zu $\text{kgV}(46, 58) = 2 \cdot 23 \cdot 29 = 1334$ teilerfremd ist. (Es sollen immer zwei Buchstaben zusammengefaßt und die Buchstabencodierung $A = 00, B = 01, \dots, Z = 25$ verwendet werden.)

111) Es seien p, q zwei verschiedene Primzahlen, $n = pq$ und $v = \text{kgV}(p-1, q-1)$. Man zeige, daß für alle ganzen Zahlen a die Beziehung $a^{v+1} \equiv a \pmod{n}$ gilt. (Man beachte insbesondere jene Zahlen a , die nicht zu n teilerfremd sind!)

112) Der Code C sei gegeben durch $C = \{010110, 100011, 111000, 001101, 011011\}$. Berechnen Sie die Minimaldistanz $\delta(C)$ und bestimmen Sie daraus an wievielen Stellen man alle Fehler erkennen bzw. korrigieren kann.

113) Stellen Sie für den Code aus 112) ein Korrekturschema auf, indem Sie zu jedem Codewort alle diejenigen Wörter mit Distanz 1 angeben, die eindeutig zu diesem Codewort korrigiert werden können. Können auch die Wörter 010001, 001010 bzw. 111111 eindeutig korrigiert werden?

114) Wir interpretieren die Menge U aus 100) als Gruppencode. Bestimmen Sie die Minimaldistanz, zu jeder Nebenklasse die möglichen Nebenklassenführer mit minimalem Gewicht, sowie für eine Auswahl von Nebenklassenführern ein Korrekturschema.

115) Es sei $f : \mathbb{Z}_2^2 \rightarrow \mathbb{Z}_2^5$ jene Codierungsabbildung, die durch die *Matrizenmultiplikation*

$$f(v_1, v_2) = (v_1, v_2) \cdot G \quad \text{mit} \quad G = \begin{pmatrix} 1 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 1 & 1 \end{pmatrix}$$

gegeben ist. Man zeige, daß $C = f(\mathbb{Z}_2^2)$ einen Gruppencode bildet. Weiters bestimme man C , die Minimaldistanz $\delta(C)$ und wie in 114) ein Korrekturschema.

116) Man zeige, daß in einem Ring R immer $(-a) \cdot b = -(a \cdot b)$ gilt.

117) Es sei $R = \mathbb{Z}[\sqrt{-5}] = \{a + b\sqrt{-5} \mid a, b \in \mathbb{Z}\} \subseteq \mathbb{C}$. Man zeige, daß R mit der aus den komplexen Zahlen vererbten Addition $+$ und Multiplikation $\cdot$ einen Ring bildet.

118) Man zeige, daß die Menge R^* aller invertierbarer Elemente in einem Ring mit 1 bezüglich der Multiplikation eine Gruppe bildet. (Einheiten sind Elemente, die ein multiplikativ inverses Element besitzen.)

119) Sei $R[[x]]$ der Ring der formalen Potenzreihen über dem Ring R . Man zeige:

$$(1 + (-1)x)^{-1} = 1 + x + x^2 + x^3 + \cdots = \sum_{k=0}^{\infty} x^k \in R[[x]].$$

120) Sei R ein Integritätsbereich. Man zeige, daß die Einheiten $R[x]^*$ des Polynomrings $R[x]$ genau die Polynome $a_0 + 0x + 0x^2 + \cdots$ vom Grad 0 mit $a_0 \in R^*$ sind, d.h.

$$R[x]^* = R^*.$$

121) Man bestimme mit Hilfe des Euklidischen Algorithmus alle größten gemeinsamen Teiler der Polynome

$$a(x) = x^4 - 4x^2 - x - 2, b(x) = x^3 + 2x^2 - x - 2 \in \mathbb{R}[x].$$

122) Man zeige, daß das Polynome $x^2 + x + 1$ in $\mathbb{Z}_2$ irreduzibel ist, d.h. es gibt keine Polynome $ax + b, cx + d \in \mathbb{Z}_2[x]$ vom Grad 1 mit $(ax + b)(cx + d) = x^2 + x + 1$.

123) Man ermittle die Operationstabellen für die Addition und die Multiplikation im Ring

$$\tilde{R} = \mathbb{Z}_2[x]/(x^2 + x + 1)\mathbb{Z}_2[x] = \{\overline{a + bx} \mid a, b \in \mathbb{Z}_2\}$$

und zeige, daß $\tilde{R}$ einen Körper mit 4 Elementen bildet.

124) Man suche möglichst einfache Schaltungen, die folgende Boolesche Funktionen realisieren.

x_1	x_2	x_3	f_1	f_2	f_3	f_4
1	1	1	1	1	0	0
1	1	0	1	0	1	0
1	0	1	0	1	0	1
1	0	0	0	1	0	0
0	1	1	1	1	0	1
0	1	0	0	0	1	0
0	0	1	0	1	0	1
0	0	0	0	0	1	1