

Ausgewählte Kapitel der Theoretischen Informatik

Hoffmann, Ulrich

Publication date:
2007

Document Version
Verlags-PDF (auch: Version of Record)

[Link to publication](#)

Citation for pulished version (APA):

Hoffmann, U. (2007). *Ausgewählte Kapitel der Theoretischen Informatik*. (Final; Band 12, Nr. 3). Universität Lüneburg.

General rights

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

- Users may download and print one copy of any publication from the public portal for the purpose of private study or research.
- You may not further distribute the material or use it for any profit-making activity or commercial gain
- You may freely distribute the URL identifying the publication in the public portal ?

Take down policy

If you believe that this document breaches copyright please contact us providing details, and we will remove access to the work immediately and investigate your claim.



Ausgewählte Kapitel der Theoretischen Informatik

Ulrich Hoffmann

12. Jahrgang, Heft 3, August 2002, ISSN 0939-8821

Technical Reports and Working Papers

Hrsg: Hinrich E. G. Bonin

Volgershall 1, D-21339 Lüneburg

Phone: xx49.4131.677175 Fax: xx49.4131.677140

Inhaltsverzeichnis

1	Einleitung.....	3
1.1	Einige grundlegende Bezeichnungen.....	4
1.2	Problemklassen.....	18
1.3	Ein intuitiver Algorithmusbegriff.....	24
2	Modelle der Berechenbarkeit.....	29
2.1	Deterministische Turingmaschinen.....	29
2.2	Random Access Maschinen.....	50
2.3	Programmiersprachen.....	61
2.4	Universelle Turingmaschinen.....	63
2.5	Nichtdeterminismus.....	71
3	Grenzen der Berechenbarkeit.....	86
3.1	Eigenschaften rekursiv aufzählbarer und rekursiver Mengen.....	90
4	Elemente der Theorie Formaler Sprachen und der Automatentheorie ...	108
4.1	Grammatiken und formale Sprachen.....	108
4.2	Typ-0-Sprachen.....	111
4.3	Typ-1-Sprachen.....	112
4.4	Typ-2-Sprachen.....	118
4.5	Typ-3-Sprachen.....	133
4.6	Eigenschaften im tabellarischen Überblick.....	143
5	Praktische Berechenbarkeit.....	145
5.1	Der Zusammenhang zwischen Optimierungs- und Entscheidungsproblemen.....	157
5.2	Komplexitätsklassen.....	164
5.3	Die Klassen P und NP.....	169
5.4	NP-Vollständigkeit.....	181
5.5	Bemerkungen zur Struktur von NP.....	193
6	Approximation von Optimierungsaufgaben.....	198
6.1	Relativ approximierbare Probleme.....	204
6.2	Polynomiell zeitbeschränkte und asymptotische Approximationsschemata.....	220
7	Weiterführende Konzepte.....	226
7.1	Randomisierte Algorithmen.....	226
7.2	Modelle randomisierter Algorithmen.....	236
8	Übungsaufgaben.....	246
	Literaturauswahl.....	255

1 Einleitung

Der vorliegende Text trifft eine (subjektive) Auswahl aus Themen der Theoretischen Informatik, deren Kenntnis neben vielen anderen Themen als unabdingbar im Rahmen der Informatikausbildung angesehen wird. Auch wenn die direkte praktische Relevanz der dargestellten Inhalte nicht immer sofort sichtbar wird, so zählen diese doch zu den grundlegenden Bildungsinhalten, die jedem Informatiker und Wirtschaftsinformatiker vermittelt werden sollten. Zudem steht für die (dem einen oder anderen auch sehr theoretisch erscheinenden) Ergebnisse gerade der Komplexitätstheorie außer Frage, daß sie einen unmittelbaren Einfluß auf die praktische Umsetzung und algorithmische Realisierung von Problemlösungen haben.

Der Text verzichtet stellenweise auf die Darstellung ausformulierter (mathematisch exakter) Beweise. Diese können in der angegebenen Literatur nachgelesen werden bzw. sie werden in der zugehörigen Vorlesung und den Übungen behandelt.

Die Theoretische Informatik als eine der Säulen der Informatik kann natürlich nicht umfassend im Rahmen einer einsemestrigen Lehrveranstaltung dargestellt werden. Ein Weg, sich den wesentlichen Inhalten zu nähern, zu denen die Theorie der Berechenbarkeit und Entscheidbarkeit, die Automatentheorie, die Theorie der Formalen Sprachen und die Komplexitätstheorie zählen, wird in dem Lehrbuch von Schöning: *Theoretische Informatik kurzgefaßt* aufgezeigt, in dem die angesprochenen Themen in kompakter Form präsentiert werden. Dem vorliegenden Text liegt ein ähnlicher Ansatz zugrunde, wobei hier jedoch eine teilweise unterschiedliche Schwerpunktbildung erfolgt. Er verzichtet auf den Anspruch einer umfassenden oder ansatzweisen Darstellung *aller* wichtigen Teilgebiete der Theoretischen Informatik und konzentriert sich auf einige Grundlagen, die zum Verständnis der prinzipiellen Leistungsfähigkeit von Algorithmen und Computern notwendig erscheinen. Der Blick richtet sich also auf Inhalte, die der Beantwortung von Fragen dienlich sein können wie

- Wie kann man die Begriffe der algorithmischen Berechenbarkeit formal fassen?
- Was können Computer und Algorithmen leisten und gibt es prinzipielle Grenzen der algorithmischen Berechenbarkeit?
- Mit welchen Modellierungswerkzeugen und Beschreibungsmitteln lassen sich berechenbare Menge charakterisieren?
- Mit welchem algorithmischen Aufwand ist im konkreten Fall bei einer Problemlösung zu rechnen?
- Wie kann man die inhärente Komplexität einiger Problemstellungen praktisch nutzen?

Die Behandlung dieser Fragestellungen führt auf die Betrachtung von Modellen der Berechenbarkeit. Hierbei spielt die Turingmaschine eine herausragende Rolle. Ergänzend wird

wegen seiner inhaltlichen Nähe zu realen Computern das Modell der Random Access Maschine behandelt. Auf die Darstellung von Theorien wie die m -rekursiven Funktionen oder das λ -Kalkül soll hier (im wesentlichen aus Platzgründen) verzichtet werden, obwohl sie eine hohe mathematische Eleganz und Relevanz aufweisen.

Innerhalb der Theoretischen Informatik nimmt das Teilgebiet Automatentheorie und Formale Sprachen einen breiten Raum ein. Historisch hat es einen großen Einfluß auf die Entwicklung von Programmiersprachen und Sprachübersetzern, sowie auf die Modellierung komplexer Anwendungssysteme und Betriebssysteme ausgeübt. Im folgenden wird dieses Teilgebiet jedoch nur im Überblick dargestellt, da die Themen Programmiersprachen und Compilerbau in der Wirtschaftsinformatikausbildung höchstens am Rand gestreift werden und auch die Modellierungsmöglichkeiten beispielsweise paralleler Abläufe durch Petrinetze und anderer Automatentypen innerhalb des jeweiligen Anwendungsgebiets behandelt werden können.

Das Thema der praktisch durchführbaren Berechnungen, d.h. der in polynomieller Zeit zu lösenden Probleme, führt auf die Theorie der NP-Vollständigkeit und der Approximation schwerer Optimierungsprobleme. Die damit verbundenen Fragestellungen werden in den beiden letzten Kapiteln aufgezeigt.

Im folgenden Text werden im einzelnen nicht die Literaturquellen angegeben, denen die jeweiligen Inhalte entnommen wurden. Das Literaturverzeichnis führt eine Reihe wichtiger Standardwerke auf, die die Themen in großer Ausführlichkeit behandeln. Dort finden sich auch zusätzliche Übungen und weiterführende Quellenangaben.

1.1 Einige grundlegende Bezeichnungen

Im vorliegenden Kapitel werden grundlegende Definitionen angeführt und einige in den folgenden Kapiteln verwendete (mathematische) Grundlagen zitiert. Dabei wird eine gewisse Vertrautheit mit der grundlegenden Symbolik der Mathematik vorausgesetzt, z.B. mit Schreibweisen wie

$a \in A$, $A \subseteq B$, $A \cup B$, $A \cap B$, $A \setminus B$.

Es seien A und B mathematisch logische Aussagen, die wahr oder falsch sein können. In den nachfolgenden Kapiteln werden häufig daraus Aussagen der Form

„Aus A folgt B “ bzw.

„ A impliziert B “ bzw.

„Wenn A gilt, dann gilt auch B “, gelegentlich auch geschrieben

„ $A \Rightarrow B$ “

gebildet und bewiesen.

Für einen Beweis der Aussage $A \Rightarrow B$ geht man folgendermaßen vor:

Man nimmt an, daß A wahr ist. Durch eine „geeignete“ Argumentation (Anwendung logischer Schlüsse) zeigt man, daß dann auch B wahr ist.

Alternativ kann man auch folgendermaßen argumentieren:

Man nimmt an, daß B *nicht* wahr ist. Durch eine „geeignete“ Argumentation (Anwendung logischer Schlüsse) zeigt man, daß dann auch A nicht wahr ist. Diese Argumentation beruht auf der Tatsache, daß $A \Rightarrow B$ logisch äquivalent (siehe unten) zu $\neg B \Rightarrow \neg A$ ist, d.h. $A \Rightarrow B$ und $\neg B \Rightarrow \neg A$ haben stets denselben Wahrheitswert.

Die Aussage „ $A \Leftrightarrow B$ “ steht für $A \Rightarrow B$ und $B \Rightarrow A$. Um die Aussage $A \Leftrightarrow B$ zu beweisen, sind also zwei „Richtungen“ zu zeigen, d.h. jeweils ein Beweis für $A \Rightarrow B$ und ein Beweis für $B \Rightarrow A$ zu erbringen. Alternativ kann man auch $A \Rightarrow B$ und $\neg A \Rightarrow \neg B$ beweisen. Die Gültigkeit von $A \Leftrightarrow B$ bedeutet, daß A und B beide wahr oder beide falsch sind.

Statt $A \Leftrightarrow B$ sagt man auch

„ A ist (logisch) äquivalent zu B “ bzw.

„ A gilt genau dann, wenn B gilt“.

Die Elemente einer Menge A seien durch eine Eigenschaft E_A beschrieben, die allen Elementen von A zukommt: $A = \{a \mid a \text{ hat die Eigenschaft } E_A\}$. Entsprechend werde die Menge B durch eine Eigenschaft E_B bestimmt: $B = \{b \mid b \text{ hat die Eigenschaft } E_B\}$. Um zu beweisen, daß die Teilmengenbeziehung $A \subseteq B$ gilt, geht man folgendermaßen vor:

Man nehme $x \in A$. Dann weiß man, daß x die Eigenschaft E_A besitzt. Durch eine „geeignete“ Argumentation unter Ausnutzung der Eigenschaft E_A weist man nach, daß x auch die Eigenschaft E_B besitzt, d.h. $x \in B$. Man zeigt also $x \in A \Rightarrow x \in B$.

Logisch äquivalent ist folgende Argumentation:

Man nehme für ein Element x an, daß es *nicht* in B liegt: $x \notin B$, d.h. x hat *nicht* die Eigenschaft E_B . Durch eine „geeignete“ Argumentation unter Ausnutzung der Tatsache, daß für x die Eigenschaft E_B nicht zutrifft, weist man nach, daß x auch nicht die Eigenschaft E_A besitzt, d.h. daß $x \notin A$ ist.

Um die Gleichheit zweier Mengen A und B nachzuweisen, d.h. $A = B$, hat man die beiden Inklusionen $A \subseteq B$ und $B \subseteq A$ nachzuweisen, d.h. $x \in A \Leftrightarrow x \in B$.

Mit $\mathbf{P}(M)$ wird die **Potenzmenge** der Menge M bezeichnet, d.h. $\mathbf{P}(M) = \{L \mid L \subseteq M\}$.

Es gilt für alle Mengen A, B und C :

$$A \cap B \subseteq A, A \cap B \subseteq B, A \subseteq A \cup B, B \subseteq A \cup B$$

$$A \cup B = B \cup A, A \cap B = B \cap A \text{ (Kommutativgesetze)}$$

$$A \cup B = (A \setminus B) \cup (A \cap B) \cup (B \setminus A), \text{ die rechte Seite ist eine disjunkte Zerlegung}^1 \text{ von } A \cup B,$$

$$A \cap (B \cup C) = (A \cap B) \cup (A \cap C), A \cup (B \cap C) = (A \cup B) \cap (A \cup C) \text{ (Distributivgesetze)}.$$

Für Mengen $A_1, A_2, \dots, A_n$ wird das **kartesische Produkt** definiert als

$$A_1 \times A_2 \times \dots \times A_n = \{(a_1, a_2, \dots, a_n) \mid a_1 \in A_1, a_2 \in A_2, \dots, a_n \in A_n\}.$$

Für eine Menge A bezeichnet $|A|$ die **Anzahl der Elemente** (oder die **Mächtigkeit**) von A .

Es seien A und B Mengen. Eine Vorschrift, die jedem Element $a \in A$ *genau* ein Element $b \in B$ zuordnet, heißt **(totale) Funktion** von A nach B , geschrieben:

$$f: \begin{cases} A \rightarrow B \\ a \rightarrow f(a) \end{cases}$$

häufig auch in der Form

$$f: A \rightarrow B, f(a) = \dots$$

Die Angabe $f: A \rightarrow B$ legt fest, daß einem Element vom (Daten-) Typ, der „charakteristisch“ für A ist, jeweils ein Element vom (Daten-) Typ, der „charakteristisch“ für B ist, zugeordnet wird. Beispielsweise könnte die Menge A aus Objekten vom Objekttyp T und die Menge B aus natürlichen Zahlen bestehen. Dann legt die Angabe $f: A \rightarrow B$ fest, daß jedem Objekt vom Objekttyp T in der Menge A durch f eine natürliche Zahl, die beispielsweise als Primärschlüsselwert interpretierbar ist, zugeordnet wird. Die Angabe $f(a) = \dots$ beschreibt, wie diese Zuordnung für jedes Element $a \in A$ geschieht.

Formal ist eine (totale) Funktion $f: A \rightarrow B$ eine Abbildungsvorschrift (genauer: eine zweistellige Relation), die folgenden beiden Bedingungen genügt:

- (i) Sie ist **linkstotal**: für jedes $a \in A$ gibt es $b \in B$ mit $f(a) = b$
- (ii) Sie ist **rechtseindeutig**: $f(a) = b_1$ und $f(a) = b_2$ implizieren $b_1 = b_2$.

Die Menge A heißt **Definitionsbereich** von f , die Menge

¹ Eine Zerlegung $M = M_1 \cup M_2$ der Menge M heißt disjunkt, wenn $M_1 \cap M_2 = \emptyset$ ist.

$$f(A) = \{b \mid b \in B, \text{ und es gibt } a \in A \text{ mit } f(a) = b\}$$

heißt **Wertebereich** von f . Es ist $f(A) \subseteq B$.

Eine Abbildungsvorschrift $f: A \rightarrow B$ heißt **partielle Funktion**, wenn lediglich die obige Bedingung (ii) erfüllt ist. Eine partielle Funktion $f: A \rightarrow B$ ordnet u.U. nicht jedem $a \in A$ ein $b \in B$ zu.

Eine Abbildungsvorschrift $f: A \rightarrow B$ heißt **injektiv**, wenn sie **linkseindeutig** ist, d.h. wenn gilt:

für jedes $a_1 \in A$ und jedes $a_2 \in A$ mit $a_1 \neq a_2$ ist $f(a_1) \neq f(a_2)$.

Eine Abbildungsvorschrift $f: A \rightarrow B$ heißt **surjektiv**, wenn sie rechtstotal ist, d.h. wenn gilt:

$$f(A) = B.$$

Eine Abbildungsvorschrift $f: A \rightarrow B$ heißt **bijektiv**, wenn sie injektiv und surjektiv ist. In diesem Fall gibt es eine eindeutig bestimmte **Umkehrfunktion** $f^{-1}: B \rightarrow A$ mit $f^{-1}(f(a)) = a$ und $f\left(f^{-1}(b)\right) = b$.

Die **Menge der natürlichen Zahlen** wird definiert durch $\mathbf{N} = \{0, 1, 2, 3, 4, \dots\}$.

Die Theoretische Informatik untersucht häufig Eigenschaften von Zeichenketten, die sich aus einzelnen Symbolen (Buchstaben, Zeichen) zusammensetzen.

Es sei Σ eine endliche nichtleere Menge. Die Elemente von Σ heißen **Symbole** oder **Zeichen** oder **Buchstaben**. Σ heißt **Alphabet**. Eine Aneinanderreihung $a_1 a_2 \dots a_n$ von n Symbolen (Zeichen, Buchstaben) $a_1 \in \Sigma$, $a_2 \in \Sigma$, ..., $a_n \in \Sigma$ heißt **Wort (Zeichenkette) über Σ** . Die **Länge** von $a_1 a_2 \dots a_n$ wird durch $|a_1 a_2 \dots a_n| = n$ definiert. Das **leere Wort (leere Zeichenkette)**, das keinen Buchstaben enthält, wird mit e bezeichnet; es gilt $|e| = 0$.

Formal lassen sich die Wörter über Σ wie folgt definieren:

- (i) e ist ein Wort über Σ mit $|e| = 0$
- (ii) Ist x ein Wort über Σ mit $|x| = n - 1$ und $a \in \Sigma$, dann ist xa ein Wort über Σ mit $|xa| = n$
- (iii) y ist ein Wort über Σ genau dann, wenn es mit Hilfe der regeln (i) und (ii) konstruiert wurde.

Die **Menge aller Wörter (beliebiger Länge)** über Σ , einschließlich des leeren Worts, wird mit Σ^* bezeichnet. Zusätzlich wird $\Sigma^+ = \Sigma^* \setminus \{e\}$ gesetzt.

Sind $x = a_1 a_2 \dots a_n$ und $y = b_1 b_2 \dots b_m$ zwei Wörter aus Σ^* , so heißt das Wort $xy = a_1 a_2 \dots a_n b_1 b_2 \dots b_m$ die **Konkation** von x und y . Die Länge von xy ist $|xy| = |x| + |y|$.

Für $x \in \Sigma^*$ definiert man $x^0 = e$ und $x^{n+1} = x^n x$ für $n \geq 0$.

Eine Teilmenge $L \subseteq \Sigma^*$ heißt **(formale) Sprache** über dem Alphabet Σ .

Für Sprachen $L_1 \subseteq \Sigma^*$ und $L_2 \subseteq \Sigma^*$ definiert man das **Produkt** von L_1 und L_2 durch $L_1 \cdot L_2 = \{xy \mid x \in L_1 \text{ und } y \in L_2\}$.

Der **Abschluß** L^* einer Sprache $L \subseteq \Sigma^*$ wird durch folgende Regeln definiert:

- (i) $L^0 = \{e\}$
- (ii) $L^{n+1} = L^n \cdot L$ für $n \geq 0$
- (iii) $L^* = \bigcup_{n \geq 0} L^n$.

Unter dem **positiven Abschluß** L^+ einer Sprache $L \subseteq \Sigma^*$ versteht man $L^+ = \bigcup_{n \geq 1} L^n$.

Offensichtlich ist $L^* = L^+ \cup \{e\}$.

Es sei $w \in L^*$. Dann ist entweder $w = e$ oder es gibt ein $n \in \mathbf{N}$ mit $n \geq 1$ und $w \in L^n$. In diesem Fall gibt es n Wörter $w_1 \in L, \dots, w_n \in L$ mit $w = w_1 \dots w_n$. Jedes Wort $w_1, \dots, w_n$ werde durch seine einzelnen Buchstaben dargestellt: $w_1 = a_{1,1} \dots a_{1,i_1}, \dots, w_n = a_{n,1} \dots a_{n,i_n}$ mit $a_{l,k} \in \Sigma$ für $l = 1, \dots, n$ und $k = 1, \dots, i_l$. Dann ist

$$w = a_{1,1} \dots a_{1,i_1} \dots a_{n,1} \dots a_{n,i_n}. \text{ Für die Länge von } w \text{ gilt } |w| = \sum_{l=1}^n |w_l| = \sum_{l=1}^n i_l.$$

Ein Wort $w \in L^+$. Kann man also in der Form $w = w_1 x = y w_n$ mit $x \in L^{n-1}$ und $y \in L^{n-1}$ schreiben. Wegen $L^{n-1} \subseteq L^*$ sind beide Teilwörter x und y in L^* . Insgesamt ergibt sich $w \in L \cdot L^*$ und $w \in L^* \cdot L$ und damit $L^+ \subseteq L \cdot L^*$ und $L^+ \subseteq L^* \cdot L$. Die umgekehrten Inklusionen $L \cdot L^* \subseteq L^+$ und $L^* \cdot L \subseteq L^+$ ergeben sich entsprechend, so daß damit

$$L^+ = L \cdot L^* = L^* \cdot L$$

gezeigt ist.

Es gilt

$$L \cdot (L_1 \cup L_2) = (L \cdot L_1) \cup (L \cdot L_2),$$

$L \cdot (L_1 \cap L_2) \subseteq (L \cdot L_1) \cap (L \cdot L_2)$, und es gibt Beispiele für Mengen L , L_1 und L_2 , für die diese Inklusion nicht umkehrbar ist (beispielsweise wählt man $L = \{e, 1\}$, $L_1 = \{0\}$ und $L_2 = \{10\}$).

Eine Menge M ist **endlich von der Mächtigkeit (Kardinalität) n** , wenn es eine bijektive Abbildung $f : M \rightarrow \{1, \dots, n\}$ gibt, d.h. man kann die Elemente in M mit den natürlichen Zahlen $1, \dots, n$ durchnummerieren: $M = \{m_1, \dots, m_n\}$.

Eine Menge M ist **(von der Mächtigkeit bzw. Kardinalität) unendlich**, wenn sie eine bijektive Abbildung $f : N \rightarrow M$ zwischen einer echten Teilmenge $N \subset M$ und M gibt.

Zwei Mengen M_1 und M_2 heißen **gleichmächtig** (oder **von gleicher Kardinalität**), wenn es eine bijektive Abbildung zwischen M_1 und M_2 gibt.

Eine Menge M , die gleichmächtig mit der Menge N (der natürlichen Zahlen) ist, heißt **abzählbar unendlich**. Eine Menge heißt **(höchstens) abzählbar**, wenn sie endlich oder abzählbar unendlich ist.

Ist M eine abzählbar unendliche Menge, dann gibt es eine bijektive Abbildung $f : N \rightarrow M$, d.h. $M = \{f(0), f(1), f(2), \dots\}$. Intuitiv: Man kann die Elemente von M mit natürlichen Zahlen durchnummerieren; jedes $m \in M$ hat dabei eine eindeutige Nummer i , d.h. $m = f(i)$.

Für ein endliches Alphabet $\Sigma = \{a_1, \dots, a_n\}$ kann man die Wörter aus Σ^* (in **lexikographischer Reihenfolge**) gemäß folgender Tabelle notieren und durchnummerieren:

Nummer	Wort aus Σ^*	Bemerkung	Nummer	Wort aus Σ^*	Bemerkung
0	e	Wort der Länge 0	$n + n^2 + 1$	$a_1 a_1 a_1$	Wörter der Länge 3
1	a_1	Wörter der Länge 1	$n + n^2 + 2$	$a_1 a_1 a_2$	
2	a_2		...	...	
...	...		$n^2 + 2n$	$a_1 a_1 a_n$	
n	a_n		$n^2 + 2n + 1$	$a_1 a_2 a_1$	
$n + 1$	$a_1 a_1$	Wörter der Länge 2	...	...	
$n + 2$	$a_1 a_2$		$n^2 + 3n$	$a_1 a_2 a_n$	
...	...		...	...	
$2n$	$a_1 a_n$		$n + n^2 + n^3$	$a_n a_n a_n$	
$2n + 1$	$a_2 a_1$		...	...	
$2n + 2$	$a_2 a_2$				
...	...				
$3n$	$a_2 a_n$				
...	...				
$n + n^2$	$a_n a_n$				

Das leere Wort e erhält dabei die Nummer 0; die Wörter der Länge $k > 0$ bekommen die Nummern $\left(\sum_{i=1}^{k-1} n^i\right) + 1$ bis $\sum_{i=1}^k n^i$. Es gilt also:

Satz 1.1-1:

Es sei Σ ein endliche Alphabet. Dann gilt:

1. Σ^* ist abzählbar unendlich.
2. Jede Sprache $L \subseteq \Sigma^*$ ist entweder endlich oder abzählbar unendlich.

Teil 2. des Satzes folgt aus der Tatsache, daß jede Teilmenge einer abzählbar unendlichen Menge abzählbar ist.

Satz 1.1-2:

Die Potenzmenge $\mathbf{P}(M) = \{L \mid L \subseteq M\}$ einer endlichen Menge M mit n Elementen enthält 2^n Elemente (Teilmengen von M).

Beweis:

Es sei $M = \{a_0, \dots, a_{n-1}\}$ eine endliche Menge mit n Elementen. Jede Teilmenge $N \subseteq M$ mit k Elementen, etwa $N = \{a_{i_0}, \dots, a_{i_{k-1}}\}$ kann als 0-1-Vektor der Länge n dargestellt werden. Dabei sind alle Komponenten dieses Vektors gleich 0 bis auf die Komponenten an den Positionen $i_0, \dots, i_{k-1}$; dort steht jeweils eine 1. Umgekehrt kann jeder 0-1-Vektor der Länge n als Teilmenge von M interpretiert werden. ///

Ist M dagegen abzählbar unendlich, so gilt:

Satz 1.1-3:

Die Potenzmenge $\mathbf{P}(M) = \{L \mid L \subseteq M\}$ einer abzählbar unendlichen Menge M ist nicht abzählbar (man sagt: sie ist **überabzählbar**).

Beweis:

Es wird die **Diagonalisierungstechnik** angewendet:

Es sei $f: \mathbf{N} \rightarrow M$ eine Abzählung von M , d.h. $M = \{f(0), f(1), f(2), \dots\}$ mit $f(i) \neq f(j)$ für $i \neq j$. Angenommen, $\mathbf{P}(M)$ sei abzählbar, etwa mit Hilfe der bijektiven Abbildung $g: \mathbf{N} \rightarrow \mathbf{P}(M)$, d.h. $\mathbf{P}(M) = \{g(0), g(1), g(2), \dots\}$. Ein Wert $f(i)$ ist ein Element von M , ein Wert $g(i)$ ist eine Teilmenge von M . Es wird eine Teilmenge D von M durch $D = \{f(i) \mid f(i) \notin g(i)\}$ definiert.

Da D eine Teilmenge von M ist (denn alle Werte $f(i)$ sind Elemente von M), ist D Element von $\mathbf{P}(M)$, hat also eine Nummer $n_D \in \mathbf{N}$, d.h. $D = g(n_D)$. Es wird nun untersucht, ob das Element von M mit der Nummer n_D (das ist das Element $f(n_D)$) in D liegt oder nicht:

Es gilt $f(n_D) \in D$ nach Definition von D genau dann, wenn $f(n_D) \notin g(n_D)$ ist. Wegen $D = g(n_D)$ ist dieses gleichbedeutend mit $f(n_D) \notin D$. Dieser Widerspruch zeigt, daß die Annahme, $\mathbf{P}(M)$ sei abzählbar, falsch ist. ///

Mit $M = \Sigma^*$ sieht man:

Satz 1.1-4:

Die Menge $\{L \mid L \subseteq \Sigma^*\}$ aller Sprachen über einem endlichen Alphabet Σ ist nicht abzählbar.

Im folgenden seien $f : \mathbf{N} \rightarrow \mathbf{R}$ und $g : \mathbf{N} \rightarrow \mathbf{R}$ zwei Funktionen. Die Funktion f ist von der **(Größen-) Ordnung** $O(g(n))$, geschrieben $f(n) \in O(g(n))$, wenn gilt:

es gibt eine Konstante $c > 0$, die von n nicht abhängt, so daß $|f(n)| \leq c \cdot |g(n)|$ ist für jedes $n \in \mathbf{N}$ bis auf höchstens endlich viele Ausnahmen („...“, so daß $|f(n)| \leq c \cdot |g(n)|$ ist für fast alle $n \in \mathbf{N}$ “).

Einige Regeln:

$$f(n) \in O(f(n))$$

Für $d = \text{const.}$ ist $d \cdot f(n) \in O(f(n))$

Es gelte $|f(n)| \leq c \cdot |g(n)|$ für jedes $n \in \mathbf{N}$ bis auf höchstens endlich viele Ausnahmen. Dann ist $O(f(n)) \subseteq O(g(n))$, insbesondere $f(n) \in O(g(n))$.

$$O(O(f(n))) = O(f(n));$$

hierbei ist

$$O(O(f(n))) = \left\{ h \mid \begin{array}{l} h : \mathbf{N} \rightarrow \mathbf{R} \text{ und es gibt eine Funktion } g \in O(f(n)) \text{ und eine Konstante } c > 0 \\ \text{mit } |h(n)| \leq c \cdot |g(n)| \text{ für jedes } n \in \mathbf{N} \text{ bis auf höchstens endlich viele Ausnahmen} \end{array} \right\}$$

Im folgenden seien S_1 und S_2 zwei Mengen, deren Elemente miteinander arithmetisch verknüpft werden können, etwa durch den Operator $\circ$. Dann ist

$$S_1 \circ S_2 = \{ s_1 \circ s_2 \mid s_1 \in S_1 \text{ und } s_2 \in S_2 \}.$$

Mit dieser Notation gilt:

$$O(f(n)) \cdot O(g(n)) \subseteq O(f(n) \cdot g(n)),$$

$$O(f(n) \cdot g(n)) \subseteq |f(n)| \cdot O(g(n)),$$

$$O(f(n)) + O(g(n)) \subseteq O(|f(n)| + |g(n)|).$$

Satz 1.1-5:

Ist p ein Polynom vom Grade m , $p(n) = \sum_{i=0}^m a_i \cdot n^i$ mit $a_i \in \mathbf{R}$ und $a_m \neq 0$, so ist

$$p(n) \in O(n^k) \text{ für } k \geq m,$$

$$n^m \in O(n^k) \text{ für } k \geq m.$$

Konvergiert $f(n) = \sum_{i=0}^{\infty} a_i \cdot n^i$ für $n \leq r$, so ist für jedes $k \geq 0$: $f(n) = \sum_{i=0}^k a_i \cdot n^i + g(n)$ mit $g(n) \in O(n^{k+1})$.

Für $a > 1$ und $b > 1$ ist $\log_a(n) = \frac{1}{\log_b(a)} \cdot \log_b(n)$. Daher gilt

Satz 1.1-6:

Für $a > 1$ und $b > 1$ ist $\log_a(n) \in O(\log_b(n))$.

Wegen $e^{h(n)} = 2^{\log_2(e) \cdot h(n)}$ gilt

Satz 1.1-7:

$e^{h(n)} \in O(2^{O(h(n))})$.

Für jedes Polynom $p(n)$ und jede Exponentialfunktion a^n mit $a > 1$ ist $a^n \notin O(p(n))$, jedoch $p(n) \in O(a^n)$.

Für jede Logarithmusfunktion $\log_a(n)$ mit $a > 1$ und jedes Polynom $p(n)$ ist $p(n) \notin O(\log_a(n))$, jedoch $\log_a(n) \in O(p(n))$.

Für jede Logarithmusfunktion $\log_a(n)$ mit $a > 1$ und jede Wurzelfunktion $\sqrt[m]{n}$ ist $\sqrt[m]{n} \notin O(\log_a(n))$, jedoch $\log_a(n) \in O(\sqrt[m]{n})$.

Insgesamt ergibt sich damit

Satz 1.1-8:

Es seien $a > 1$, $b > 1$, $m \in \mathbf{N}_{>0}$, $p(n)$ ein Polynom; dann ist

$O(\log_b(n)) \subset O(\sqrt[m]{n}) \subset O(p(n)) \subset O(a^n)$.

Ein **Boolescher Ausdruck (Boolesche Formel, Formel der Aussagenlogik)** ist eine Zeichenkette über dem Alphabet $A = \{ \wedge, \vee, \neg, (,), 0, 1, x \}$, der eine Aussage der Aussagenlogik

darstellt². Innerhalb eines Booleschen Ausdrucks kommen Variablen vor, die mit dem Zeichen x beginnen und von einer Folge von Zeichen aus $\{0, 1\}$ ergänzt werden. Diese ergänzende 0-1-Folge, die an das Zeichen x „gebunden“ ist, wird als Indizierung der Variablen interpretiert. Im folgenden werden daher Variablen als indizierte Variablen geschrieben, wobei die indizierende 0-1-Folge als Binärzahl in Dezimaldarstellung angegeben wird. Kommen die Zeichen 0 bzw. 1 innerhalb eines Booleschen Ausdrucks nicht an das Zeichen x gebunden vor, so werden sie als Konstanten FALSCH (FALSE) bzw. WAHR (TRUE) interpretiert.

Beispielsweise steht $(x_1 \wedge \neg x_5) \vee (\neg x_3 \wedge (x_4 \vee x_5)) \vee 0$ für den Booleschen Ausdruck $(x1 \wedge \neg x101) \vee (\neg x11 \wedge (x100 \vee x101)) \vee 0$.

Die Zeichen $\wedge$, $\vee$ bzw. $\neg$ stehen für die üblichen logischen Junktoren *UND*, *ODER* bzw. *NICHT*. Ein Boolescher Ausdruck wird nur bei korrekter Verwendung der Klammern „(“ und „)“ als syntaktisch korrekt angesehen: zu jeder sich schließenden Klammer „)“ muß es weiter links in der Zeichenkette eine sich öffnende Klammer „(“ geben, und die Anzahlen der sich öffnenden und schließenden Klammern müssen gleich sein; außerdem müssen die Junktoren $\wedge$, $\vee$ bzw. $\neg$ korrekt verwendet werden.

Die Bedeutungen der Junktoren $\wedge$, $\vee$ bzw. $\neg$ ergeben sich aus folgenden Tabellen.

x	$\neg x$
FALSE	TRUE
TRUE	FALSE

$x \vee y$	$y = \text{FALSE}$	$y = \text{TRUE}$
$x = \text{FALSE}$	FALSE	TRUE
$x = \text{TRUE}$	TRUE	TRUE

$x \wedge y$	$y = \text{FALSE}$	$y = \text{TRUE}$
$x = \text{FALSE}$	FALSE	FALSE
$x = \text{TRUE}$	FALSE	TRUE

Unter einem **Literal** innerhalb eines Booleschen Ausdrucks versteht man eine Variable x_i oder eine negierte Variable $\neg x_i$.

Ein Boolescher Ausdruck F heißt **erfüllbar**, wenn es eine Ersetzung der Variablen in F durch Werte FALSE bzw. TRUE gibt, wobei selbstverständlich gleiche Variablen durch die gleichen Werte ersetzt werden (man sagt: die Variablen werden mit FALSE bzw. TRUE **belegt**), so daß sich bei Auswertung der so veränderten Formel F gemäß den Regeln über die Junktoren der Wert TRUE ergibt.

Ein Boolescher Ausdruck F ist in **konjunktiver Normalform**, wenn F die Form

² Die Syntax eines Booleschen Ausdrucks bzw. einer Formel der Aussagenlogik wird hier als bekannt vorausgesetzt.

$$F = F_1 \wedge \dots \wedge F_m$$

und jedes F_i die Form

$$F_i = (y_{i_1} \vee \dots \vee y_{i_k})$$

hat, wobei y_{i_j} ein Literal oder eine Konstante 0 oder 1 ist, d.h. für eine Variable (d.h. $y_{i_j} = x_l$) oder für eine negierte Variable (d.h. $y_{i_j} = \neg x_l$) oder für eine Konstante 0 bzw. 1 steht.

Die Teilformeln F_i von F bezeichnet man als die **Klauseln** (von F). Natürlich ist eine Klausel selbst in konjunktiver Normalform.

Eine Klausel F_i von F ist durch eine Belegung der in F vorkommenden Variablen erfüllt, d.h. besitzt den Wahrheitswert *WAHR*, wenn mindestens ein Literal in F_i erfüllt ist d.h. den Wahrheitswert *WAHR* besitzt. F ist erfüllt, wenn alle in F vorkommenden Klauseln erfüllt sind.

Zu jedem Booleschen Ausdruck F gibt es einen Booleschen Ausdruck in konjunktiver Normalform, der genau dann erfüllbar ist, wenn F erfüllbar ist.

Nicht jeder Boolesche Ausdruck, selbst wenn er syntaktisch korrekt ist, ist erfüllbar. Beispielsweise ist $x \wedge \neg x$ nicht erfüllbar. Mit Hilfe eines systematischen Verfahrens kann man die Erfüllbarkeit eines Booleschen Ausdrucks F mit n Variablen dadurch testen, daß man nacheinander alle 2^n möglichen Belegungen der Variablen in F mit Wahrheitswerten *WAHR* bzw. *FALSCH* erzeugt. Immer, wenn eine neue Belegung generiert worden ist, wird diese in die Variablen eingesetzt und der Boolesche Ausdruck ausgewertet. Die Entscheidung, ob F erfüllbar ist oder nicht, kann u.U. erst nach Überprüfung aller 2^n Belegungen erfolgen.

Ein **gerichteter Graph** $G = (V, E)$ besteht aus einer endlichen Menge $V = \{v_1, \dots, v_n\}$ von **Knoten** (vertices) und einer endlichen Menge $E = \{e_1, \dots, e_k\} \subseteq V \times V$ von **Kanten** (edges).

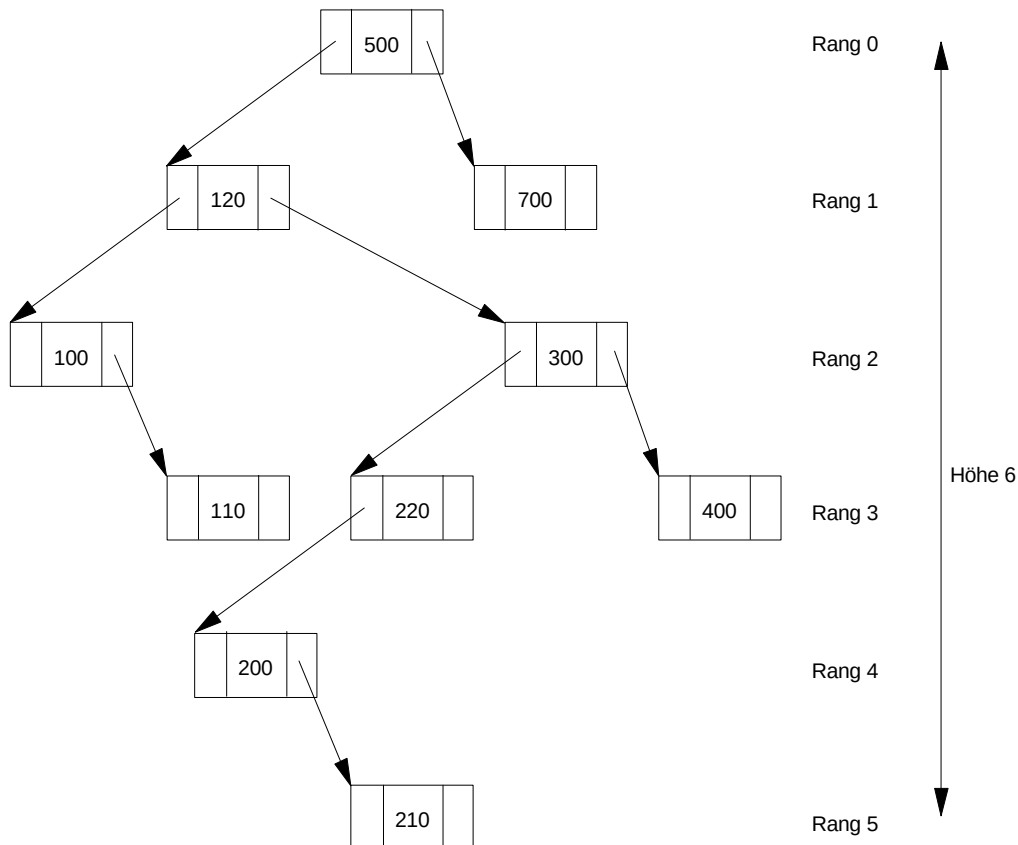
Die Kante $e = (v_i, v_j)$ läuft von v_i nach v_j (verbindet v_i mit v_j). Der Knoten v_i heißt **Anfangsknoten** der Kante $e = (v_i, v_j)$, der Knoten v_j **Endknoten** von $e = (v_i, v_j)$. Zu einem Knoten $v \in V$ heißt $\text{pred}(v) = \{v' \mid (v', v) \in E\}$ die **Menge der direkten Vorgänger** von v , $\text{succ}(v) = \{v' \mid (v, v') \in E\}$ die **Menge der direkten Nachfolger** von v .

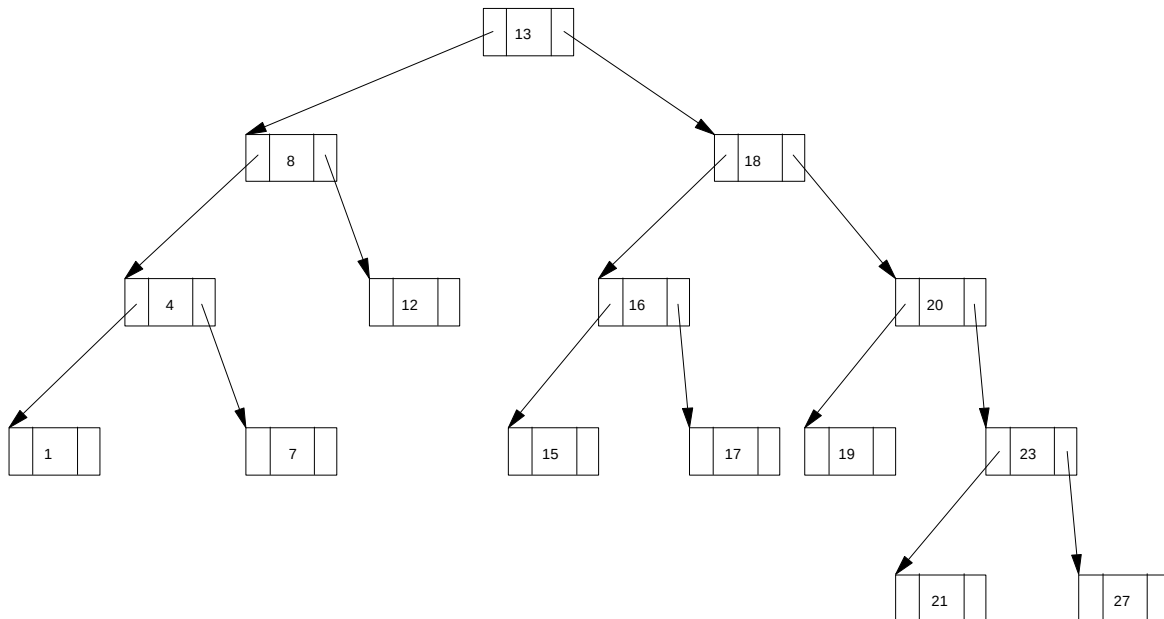
Ein **Binärbaum** $B_n = (V, E)$ mit n Knoten wird durch folgende Eigenschaften charakterisiert:

1. Entweder ist $n \geq 1$ und $|V| = n + 1$ und $|E| = n$,
oder es ist $n = 0$ und $V = E = \emptyset$ (**leerer Baum**).

2. Bei $n \geq 1$ gibt es genau einen Knoten $r \in V$, dessen Menge direkter Vorgänger leer ist; dieser Knoten heißt **Wurzel** von B_n .
3. Bei $n \geq 1$ besteht die Menge der direkten Vorgänger eines jeden Knotens, der nicht die Wurzel ist, aus genau einem Element.
4. Bei $n \geq 1$ besteht die Menge der direkten Nachfolger eines jeden Knotens aus einem Element oder zwei Elementen oder ist leer. Ein Knoten, dessen Menge der direkten Nachfolger leer ist, heißt **Blatt**.

In einem Binärbaum $B = (V, E)$ gibt es für jeden Knoten $v \in V$ genau einen **Pfad** von der Wurzel r zu v , d.h. es gibt eine Folge $((a_0, a_1), (a_1, a_2), \dots, (a_{m-1}, a_m))$ mit $r = a_0$, $v = a_m$ und $(a_{i-1}, a_i) \in E$ für $i = 1, \dots, m$. Der Wert m gibt die **Länge des Pfads** an. Um den Knoten v von der Wurzel aus über die Kanten des Pfads zu erreichen, werden m Kanten durchlaufen. Diese Länge wird auch als **Rang des Knotens** v bezeichnet.





Der Rang eines Knotens lässt sich auch folgendermaßen definieren:

- (i) Die Wurzel hat den Rang 0.
- (ii) Ist v ein Knoten im Baum mit Rang $r-1$ und w ein direkter Nachfolger von v , so hat w den Rang r .

Unter der **Höhe eines Binärbaums** versteht man den maximal vorkommenden Rang eines Blattes + 1. Die Höhe ist gleich der Anzahl der Knoten, die auf einem Pfad maximaler Länge von der Wurzel zu einem Blatt durchlaufen werden.

In einem Binärbaum bilden alle Knoten mit demselben Rang ein **Niveau des Baums**. Das Niveau 0 eines Binärbaums enthält genau einen Knoten, nämlich die Wurzel. Das Niveau 1 enthält mindestens 1 und höchstens 2 Knoten. Das Niveau j enthält höchstens doppelt so viele Knoten wie das Niveau $j-1$. Daher gilt:

Satz 1.1-9:

1. Das Niveau $j \geq 0$ eines Binärbaums enthält mindestens einen und höchstens 2^j Knoten. Die Anzahl der Knoten vom Niveau 0 bis zum Niveau j (einschließlich) beträgt mindestens $j+1$ Knoten und höchstens $\sum_{i=0}^j 2^i = 2^{j+1} - 1$ Knoten.
2. Ein Binärbaum hat maximale Höhe, wenn jedes Niveau genau einen Knoten enthält. Er hat minimale Höhe, wenn jedes Niveau eine maximale Anzahl von Knoten enthält. Also gilt für die Höhe $h(B_n)$ eines Binärbaums mit n Knoten:

$$\lceil \log_2(n+1) \rceil \leq h(B_n) \leq n.$$
3. Für die Anzahl $b(h)$ der Blätter eines Binärbaums der Höhe h gilt

$$1 \leq b(h) \leq 2^{h-1}.$$
4. Für die Mindesthöhe $h(b)$ eines Binärbaums mit $b \geq 1$ vielen Blättern gilt

$$h(b) \geq \lceil \log_2(b) + 1 \rceil.$$
5. Die Anzahl (strukturell) verschiedener Binärbäume mit n Knoten beträgt

$$\frac{1}{n+1} \binom{2n}{n} = \frac{4^n}{n\sqrt{\pi n}} + C \frac{4n}{\sqrt{n^5}} \quad \text{mit einer reellen Konstanten } C > 0.$$
6. Die *mittlere* Anzahl von Knoten, die von der Wurzel aus bis zur Erreichung eines beliebigen Knotens eines Binärbaums mit n Knoten (gemittelt über alle n Knoten) besucht werden muß, d.h. der *mittlere „Abstand“ eines Knotens von der Wurzel*, ist $\sqrt{\pi n} + C$ mit einer reellen Konstanten $C > 0$. Im günstigsten Fall (wenn also alle Niveaus voll besetzt sind) ist der größte Abstand eines Knotens von der Wurzel in einem Binärbaum mit n Knoten gleich $\lceil \log_2(n+1) \rceil \approx \log_2(n)$, im ungünstigsten Fall ist dieser Abstand gleich n .

1.2 Problemklassen

In der Informatik beschäftigt man sich häufig damit, Problemstellungen mit Hilfe von Rechnerprogrammen zu lösen. Meist wird man dabei das Programm so entwerfen, daß es nicht nur die Lösung für eine einzige konkrete Problemstellung findet, sondern für eine ganze Klasse von Problemen, die natürlich alle „ähnlich“ spezifiziert sind. Beispielsweise wird ein Sortierverfahren in der Lage sein, nicht nur einen Satz von 100 Zahlen zu sortieren, sondern eine beliebige Anzahl von Zahlen, eventuell sogar von feiner strukturierten Objekten.

Ein **Problem** ist eine zu beantwortende Fragestellung, die von **Problemparametern** (Variablenwerten, Eingaben usw.) abhängt, deren genaue Werte in der Problembeschreibung zunächst unspezifiziert sind, deren Typen jedoch in die Problembeschreibung eingeht. Ein Problem wird beschrieben durch:

1. eine allgemeine Beschreibung aller Parameter, von der die Problemlösung abhängt; diese Beschreibung spezifiziert die (Problem-) **Instanz (Eingabeinstanz)**
2. die Eigenschaften, die die Antwort, d.h. die Problemlösung, haben soll.

Eine spezielle Problemstellung erhält man durch Konkretisierung einer Problem Instanz, d.h. durch die Angabe spezieller Parameterwerte in der Problembeschreibung.

Im folgenden werden einige grundlegende **Problemtypen** unterschieden und zunächst an Beispielen erläutert.

Problem des Handlungsreisenden als Optimierungsproblem (minimum traveling salesperson problem)

Instanz: $x = (C, d)$

$C = \{c_1, \dots, c_n\}$ ist eine endliche Menge und $d : C \times C \rightarrow \mathbb{N}$ eine Funktion.

Die Menge C kann beispielsweise als eine Menge von Orten und die Wert $d(c_i, c_j)$ können als Abstand zwischen c_i und c_j interpretiert werden.

Lösung: Eine Permutation (Anordnung) $p : [1:n] \rightarrow [1:n]$, d.h. (implizit) eine Anordnung

$\langle c_{p(1)}, \dots, c_{p(n)} \rangle$ der Werte in C , die den Wert

$$\sum_{i=1}^n d(c_{p(i)}, c_{p(i+1)}) + d(c_{p(n)}, c_{p(1)})$$

minimiert (unter allen möglichen Permutationen von $[1:n]$).

Dieser Ausdruck gibt die Länge einer „kürzesten Tour“ an, die in $c_{p(1)}$ startet, nacheinander alle Orte einmal besucht und direkt vom letzten Ort $c_{p(n)}$ nach $c_{p(1)}$ zurückkehrt. Der Wert $p(i)$ gibt an, wo man sich im i -ten Schritt befindet.

Ein Beispiel einer Instanz ist die Menge $C = \{c_1, \dots, c_4\}$ mit den Werten $d(c_1, c_2) = 10$, $d(c_1, c_3) = 5$, $d(c_1, c_4) = 9$, $d(c_2, c_3) = 6$, $d(c_2, c_4) = 9$, $d(c_3, c_4) = 3$ und $d(c_i, c_j) = d(c_j, c_i)$.

Die Permutation $p = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 1 & 2 & 4 & 3 \end{pmatrix}$, d.h. die Anordnung $\langle c_1, c_2, c_4, c_3 \rangle$ der Werte in C ist eine (optimale) Lösung mit Wert 27.

Eine verallgemeinerte Formulierung des Problems des Handlungsreisenden bezogen auf endliche Graphen lautet:

Problem des Handlungsreisenden auf Graphen als Optimierungsproblem

Instanz: $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

Lösung: Eine **Tour** durch G , d.h. eine geschlossene Kantenfolge

$$\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle \text{ mit } (v_{i_j}, v_{i_{j+1}}) \in E \text{ für } j = 1, \dots, n-1,$$

in der jeder Knoten des Graphen (als Anfangsknoten einer Kante) genau einmal vorkommt, die minimale Kosten (unter allen möglichen Touren durch G) verursacht. Man kann o.B.d.A. $v_{i_1} = v_1$ setzen.

Die Kosten einer Tour $\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$ ist dabei die Summe der Gewichte der Kanten auf der Tour, d.h. der Ausdruck

$$\sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1}).$$

Das Problem des Handlungsreisenden findet vielerlei Anwendungen in den Bereichen

- **Transportoptimierung:**
Ermittlung einer kostenminimalen Tour, die im Depot beginnt, $n-1$ Kunden erreicht und im Depot endet.
- **Fließbandproduktion:**
Ein Roboterarm soll Schrauben an einem am Fließband produzierten Werkstück festdrehen. Der Arm startet in einer Ausgangsposition (über einer Schraube), bewegt sich dann von einer zur nächsten Schraube (insgesamt n Schrauben) und kehrt in die Ausgangsposition zurück.
- **Produktionsumstellung:**
Eine Produktionsstätte stellt verschiedene Artikel mit denselben Maschinen her. Der Herstellungsprozeß verläuft in Zyklen. Pro Zyklus werden n unterschiedliche Artikel produziert. Die Änderungskosten von der Produktion des Artikels v_i auf die des Artikels v_j betragen $w((v_i, v_j))$ (Geldeinheiten). Gesucht wird eine kostenminimale Produktionsfolge. Das Durchlaufen der Kante (v_i, v_j) entspricht dabei der Umstellung von Artikel v_i auf Artikel v_j . Gesucht ist eine Tour (zum Ausgangspunkt zurück), weil die Kosten des nächsten, hier des ersten, Zyklusstarts mit einbezogen werden müssen.

Problem des Handlungsreisenden auf Graphen als Berechnungsproblem

Instanz: $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

Lösung: Der Wert $\sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1})$ einer kostenminimalen Tour $\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$ durch G .

Zu beachten ist hierbei, daß nicht eine kostenminimale Tour selbst gesucht wird, sondern lediglich der Wert einer kostenminimalen Tour. Eventuell ist es möglich, diesen Wert (durch geeignete Argumentationen und Hinweise) zu bestimmen, ohne eine kostenminimale Tour explizit anzugeben. Das Berechnungsproblem scheint daher „einfacher“ zu lösen zu sein als das Optimierungsproblem.

Problem des Handlungsreisenden auf Graphen als Entscheidungsproblem

Instanz: $G = (V, E, w)$ und $K \in \mathbf{R}_{\geq 0}$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

Lösung: Die Antwort auf die Frage:

Gibt es eine kostenminimale Tour $\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$ durch G , deren Wert $\leq K$ ist?

Bei dieser Problemstellung ist nicht einmal der Wert einer kostenoptimalen Tour gesucht, sondern lediglich eine Entscheidung, ob dieser Wert „nicht zu groß“, d.h. kleiner als eine vorgegebene Schranke ist. Das Entscheidungsproblem scheint daher „noch einfacher“ zu lösen zu sein als das Optimierungsproblem.

Im vorliegenden Fall des Problems des Handlungsreisenden befindet man sich jedoch im Irrtum: In einem noch zu präzisierenden Sinne sind bei diesem Problem alle Problemvarianten algorithmisch gleich schwierig zu lösen.

Das Beispiel läßt sich verallgemeinern, wobei im folgenden vom (vermeintlich) einfacheren Problemtyp zum komplexeren Problemtyp übergegangen wird.

Die Instanz x eines Problems Π ist eine endliche Zeichenkette über einem endlichen Alphabet Σ_Π , das dazu geeignet ist, derartige Problemstellungen zu formulieren, d.h. $x \in \Sigma_\Pi^*$.

Es werden folgende **Problemtypen** unterschieden:

Entscheidungsproblem Π :

Instanz: $x \in \Sigma_\Pi^*$

und Spezifikation einer Eigenschaft, die einer Auswahl von Elementen aus Σ_Π^* zukommt, d.h. die Spezifikation einer Menge $L_\Pi \subseteq \Sigma_\Pi^*$ mit

$$L_\Pi = \{ u \in \Sigma^* \mid u \text{ hat die beschriebene Eigenschaft} \}$$

Lösung: Entscheidung „ja“, falls $x \in L_\Pi$ ist,
Entscheidung „nein“, falls $x \notin L_\Pi$ ist.

Bemerkung: In konkreten Beispielen für Entscheidungsprobleme wird gelegentlich die Problemspezifikation nicht in dieser strengen formalen Weise durchgeführt. Insbesondere wird die Spezifikation der die Menge $L_\Pi \subseteq \Sigma_\Pi^*$ beschreibenden Eigenschaft direkt bei der Lösung angegeben.

Bei der Lösung eines Entscheidungsproblems geht es also darum, bei Vorgabe einer Instanz $x \in \Sigma_\Pi^*$ zu entscheiden, ob x zur Menge L_Π gehört, d.h. eine genau spezifizierte Eigenschaft, die genau allen Elementen in L_Π zukommt, besitzt, oder nicht.

Es zeigt sich, daß der hier formulierte Begriff der Entscheidbarkeit sehr eng gefaßt ist. Eine erweiterte Definition eines Entscheidungsproblems verlangt bei der Vorgabe einer Instanz $x \in \Sigma_\Pi^*$ nach endlicher Zeit lediglich eine positive Entscheidung „ja“, wenn $x \in L_\Pi$ ist. Ist $x \notin L_\Pi$, so kann die Entscheidung eventuell nicht in endlicher Zeit getroffen werden. Dieser Begriff der Entscheidbarkeit führt auf die rekursiv aufzählbaren Mengen (im Gegensatz zu den entscheidbaren Mengen) und ist Gegenstand von Kapitel 3.

Berechnungsproblem Π :

Instanz: $x \in \Sigma_{\Pi}^*$

und die Beschreibung einer Funktion $f : \Sigma_{\Pi}^* \rightarrow \Sigma'^*$.

Lösung: Berechnung des Werts $f(x)$.

Optimierungsproblem Π :

Instanz: 1. $x \in \Sigma_{\Pi}^*$

2. Spezifikation einer Funktion SOL_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ **eine Menge zulässiger Lösungen** zuordnet

3. Spezifikation einer **Zielfunktion** m_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ und $y \in \text{SOL}_{\Pi}(x)$ einen Wert $m_{\Pi}(x, y)$, den **Wert einer zulässigen Lösung**, zuordnet

4. $\text{goal}_{\Pi} \in \{\min, \max\}$, je nachdem, ob es sich um ein Minimierungs- oder ein Maximierungsproblem handelt.

Lösung: $y^* \in \text{SOL}_{\Pi}(x)$ mit $m_{\Pi}(x, y^*) = \min\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Minimierungsproblem (d.h. $\text{goal}_{\Pi} = \min$) bzw. $m_{\Pi}(x, y^*) = \max\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Maximierungsproblem (d.h. $\text{goal}_{\Pi} = \max$).

Der Wert $m_{\Pi}(x, y^*)$ einer optimalen Lösung wird auch mit $m_{\Pi}^*(x)$ bezeichnet.

In dieser (formalen) Terminologie wird das Handlungsreisenden-Minimierungsproblem wie folgt formuliert:

Instanz: 1. $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w : E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht

2. $\text{SOL}(G) = \{T \mid T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle \text{ ist eine Tour durch } G\}$; eine Tour durch G ist eine geschlossene Kantenfolge, in der jeder Knoten des Graphen (als Anfangsknoten einer Kante) genau einmal vorkommt

3. für $T \in \text{SOL}(G)$, $T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$, ist die Zielfunktion

definiert durch $m(G, T) = \sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1})$

4. $goal = \min$

Lösung: Eine Tour $T^* \in \text{SOL}(G)$, die minimale Kosten (unter allen möglichen Touren durch G) verursacht, und $m(G, T^*)$.

Im folgenden werden zunächst hauptsächlich Entscheidungsprobleme (kommt einem Wort $x \in \Sigma_{\Pi}^*$ eine Eigenschaft zu, d.h. gilt $x \in L_{\Pi}$ für eine Sprache $L_{\Pi} \subseteq \Sigma_{\Pi}^*$?) und Berechnungsprobleme (man berechne den Wert $f(x)$ für eine Funktion $f: \Sigma_{\Pi}^* \rightarrow \Sigma'^*$) behandelt. Diese Probleme sollen algorithmisch gelöst werden. Dazu müssen die Begriffe der Entscheidbarkeit und der Berechenbarkeit präziser gefaßt und geeignete Modelle Berechnungsmodelle definiert werden.

1.3 Ein intuitiver Algorithmusbegriff

Ein **Algorithmus** ist eine Verfahrensvorschrift (Prozedur, Berechnungsvorschrift), die aus einer endlichen Menge eindeutiger Regeln besteht, die eine endliche Aufeinanderfolge von Operationen spezifiziert, so daß eine Lösung zu einem Problem bzw. einer spezifischen Klasse von Problemen daraus erzielt wird.

Konkret kann man sich einen Algorithmus als ein Computerprogramm vorstellen, das in einer **Pascal-ähnlichen Programmiersprache** formuliert ist. Darunter versteht man Programmiersprachen, die

- Deklarationen von Variablen (mit geeigneten Datentypen) zulassen
- die üblichen arithmetischen Operationen mit Konstanten und Variablen und Wertzuweisungen an Variablen enthalten
- Kontrollstrukturen wie Sequenz (Hintereinanderreihung von Anweisungen, blockstrukturierte Anweisungen, Prozeduren), Alternativen (**IF ... THEN ... ELSE, CASE ... END**) und WHILE-Schleifen (**WHILE ... DO ...**) besitzen.

Von einem Algorithmus erwartet man eine Reihe von **Eigenschaften**, damit er als „effektives Rechenverfahren“ gelten kann:

1. Die Verfahrensvorschrift (das Programm) soll aus einem endlichen Text bestehen.
2. Der Ablauf einer Berechnung soll schrittweise als Folge elementarer Rechenschritte erfolgen.
3. Das Verfahren soll **deterministisch** sein, d.h. in jedem Stadium einer Berechnung soll vollständig und eindeutig bestimmt sein, welcher elementare Rechenschritt als nächster getan wird. Ein Text „Bei Eingabe von x kann man für $f(x)$ einen von endlich vielen Werten aussuchen“ ist als Teil eines Algorithmus nicht zulässig.
4. Das Verfahren soll abgeschlossen sein, d.h. welcher Rechenschritt als nächster getan wird, soll ausschließlich von den Eingabewerten und den vorangegangenen berechneten Zwischenergebnissen abhängen. Ein Text „ $f(x) = x$, wenn es gerade regnet bzw. $= 2x$ sonst“ ist als Teil eines Algorithmus nicht zulässig.
5. Das Verfahren soll im Prinzip beliebig große Zahlen handhaben können.

Typische **Fragestellungen** bei einem gegebenen Algorithmus für eine Problemlösung sind:

- Hält der Algorithmus immer bei einer gültigen Eingabe nach endlich vielen Schritten an?
- Berechnet der Algorithmus bei einer gültigen Eingabe eine korrekte Antwort?

Die positive Beantwortung beider Fragen erfordert einen mathematischen **Korrektheitsbeweis** des Algorithmus. Bei positiver Beantwortung nur der zweiten Frage spricht man von **partieller Korrektheit**. Für die partielle Korrektheit ist lediglich nachzuweisen, daß der Algorithmus bei einer gültigen Eingabe, bei der er nach endlich vielen Schritten anhält, ein korrektes Ergebnis liefert.

- Wieviele Schritte benötigt der Algorithmus bei einer gültigen Eingabe **höchstens (worst case analysis)** bzw. **im Mittel (average case analysis)**, d.h. welche **(Zeit-) Komplexität** hat er im schlechtesten Fall bzw. im Mittel? Dabei ist es natürlich besonders interessant nachzuweisen, daß die Komplexität des Algorithmus von der jeweiligen Formulierungsgrundlage (Programmiersprache, Maschinenmodell) weitgehend unabhängig ist.

Entsprechend kann man nach dem benötigten **Speicherplatzbedarf (Platzkomplexität)** eines Algorithmus fragen.

Die Beantwortung dieser Fragen für den schlechtesten Fall gibt **obere Komplexitätsschranken** (Garantie für das Laufzeitverhalten bzw. den Speicherplatzbedarf) an.

- Gibt es zu einer Problemlösung eventuell ein „noch besseres“ Verfahren (mit weniger Rechenschritten, weniger Speicherplatzbedarf)? Wieviele Schritte wird jeder Algorithmus mindestens durchführen, der das vorgelegte Problem löst?

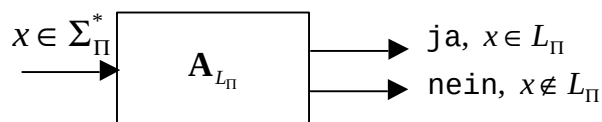
Die Beantwortung dieser Frage liefert untere Komplexitätsschranken.

Entsprechend der verschiedenen Problemtypen (Entscheidungsproblem, Berechnungsproblem, insbesondere Optimierungsproblem) gibt es auch unterschiedliche **Typen von Algorithmen**:

Ein **Algorithmus A_{L_Π} für ein Entscheidungsproblem Π** mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_\Pi^*$

Ausgabe: ja (accept), falls $x \in L_\Pi$ gilt
nein (reject), falls $x \notin L_\Pi$ gilt.



Bei Eingabe von $x \in \Sigma_\Pi^*$ wird mit $A_{L_\Pi}(x)$, $A_{L_\Pi}(x) \in \{ \text{ja}, \text{nein} \}$, die Entscheidung von A_{L_Π} bezeichnet.

Es gilt also für $x \in \Sigma_\Pi^*$:

Es ist $x \in L_\Pi$ genau dann, wenn A_{L_Π} bei Eingabe von $x \in \Sigma_\Pi^*$ mit $A_{L_\Pi}(x) = \text{ja}$ stoppt.

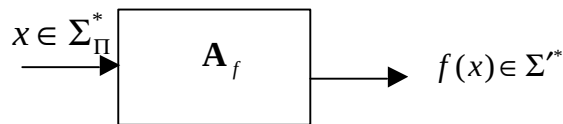
Wie bereits in Kapitel 1.2 bemerkt, erweist sich der definierte Begriff des Entscheidungsalgorithmus als zu eng gefaßt, wenn man fordert, daß der Algorithmus bei jeder Eingabe $x \in \Sigma_\Pi^*$ stoppt (entweder mit Antwort ja oder nein). Es werden daher auch Entscheidungsprobleme algorithmisch behandelt, die Verfahren zulassen, die nicht bei jeder Eingabe anhal-

ten. Bei Eingabe von $x \in \Sigma_{\Pi}^*$ stoppt der Algorithmus mit Antwort ja, falls $x \in L_{\Pi}$ ist. Für $x \notin L_{\Pi}$ wird zugelassen, daß der Algorithmus entweder mit Antwort nein stoppt oder nicht terminiert. Dieser Algorithmusbegriff führt auf den Begriff der **rekursiv aufzählbaren Mengen**.

Ein **Algorithmus A_f für ein Berechnungsproblem Π** (Berechnung einer Funktion $f : \Sigma_{\Pi}^* \rightarrow \Sigma'^*$) hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_{\Pi}^*$

Ausgabe: $f(x) \in \Sigma'^*$



Mit $A_f(x)$ wird das Berechnungsergebnis von A_f bei Eingabe von $x \in \Sigma_{\Pi}^*$ bezeichnet, d.h. $A_f(x) = f(x)$. Falls A_f bei einer Eingabe $x \in \Sigma_{\Pi}^*$ nicht anhält, ist $f(x)$ nicht definiert. In diesem Fall berechnet A_f eine partielle Funktion.

Mit diesen sehr „vagen“ Begriff der Berechenbarkeit läßt sich bereits zeigen, daß es Funktionen gibt, die nicht berechenbar sind. Genauer:

Satz 1.3-1:

Es gibt nichtberechenbare Funktionen der Form $g : \mathbf{N} \rightarrow \{0, 1\}$.

Beweis:

Wieder wird die Diagonalisierungstechnik verwendet (vgl. den Beweis von Satz 1.1-3): Jedes Berechnungsverfahren ist ein Algorithmus, der mit Hilfe eines endlichen Alphabets Σ formuliert werden kann. Die Menge B aller Berechnungsverfahren für Funktionen der Form $g : \mathbf{N} \rightarrow \{0, 1\}$ ist daher eine Teilmenge von Σ^* und damit abzählbar (unendlich). Es gibt daher eine bijektive Funktion $f : \mathbf{N} \rightarrow B$, d.h. $B = \{f(0), f(1), f(2), \dots\}$. Die folgende Argumentation wird klarer, wenn man B in der Form $B = \{f_0, f_1, f_2, \dots\}$ darstellt, d.h. die Nummer einer Funktion in B als Index angibt. Es wird eine Funktion $g_0 : \mathbf{N} \rightarrow \{0, 1\}$ wie folgt definiert:

$$g_0 : \begin{cases} \mathbf{N} & \rightarrow \{0, 1\} \\ n & \rightarrow \begin{cases} 0 & \text{falls } f_n(n) = 1 \text{ ist.} \\ 1 & \text{falls } f_n(n) = 0 \end{cases} \end{cases}$$

Angenommen, die Funktion g_0 komme in B vor. Dann hat g_0 eine Nummer $i_g \in \mathbf{N}$ in der Aufzählung von B , also $g_0 = f_{i_g}$. Nun gibt es zwei Möglichkeiten: entweder $g_0(i_g) = 0$ oder $g_0(i_g) = 1$.

Ist $g_0(i_g) = 0$, dann ist (nach Definition von g_0) $f_{i_g}(i_g) = 1$. Wegen $g_0 = f_{i_g}$ entsteht der Widerspruch $0 = g_0(i_g) = f_{i_g}(i_g) = 1$.

Ist $g_0(i_g) = 1$, dann ist (nach Definition von g_0) $f_{i_g}(i_g) = 0$. Wieder ergibt sich ein Widerspruch, nämlich $1 = g_0(i_g) = f_{i_g}(i_g) = 0$.

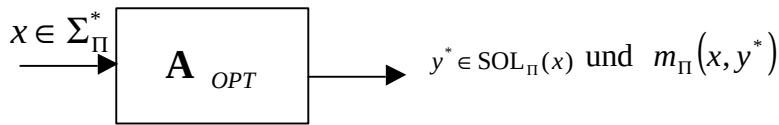
Daher ist $g_0 \notin B$ und damit eine nichtberechenbare Funktion der Form $g : \mathbf{N} \rightarrow \{0, 1\}$. ///

Der Vollständigkeit halber soll noch die Form eines Algorithmus zur Lösung eines Optimierungsproblems, wie es in Kapitel 1.2 definiert wurde, angeführt werden:

Ein **Algorithmus** $\mathbf{A}_{OPT}$ **für ein Optimierungsproblem** Π mit Zielfunktion m_Π und Optimierungsziel $goal_\Pi \in \{\min, \max\}$ hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_\Pi^*$

Ausgabe: $y^* \in \text{SOL}_\Pi(x)$ und $m_\Pi(x, y^*) = goal_\Pi \{m_\Pi(x, y) \mid y \in \text{SOL}_\Pi(x)\}$



Mit $\mathbf{A}_{OPT}(x)$ wird die von $\mathbf{A}_{OPT}$ bei Eingabe von $x \in \Sigma_\Pi^*$ ermittelte Lösung bezeichnet, d.h.

$\mathbf{A}_{OPT}(x) = (y^*, m_\Pi(x, y^*))$. Hier wird implizit vorausgesetzt, daß für $x \in \Sigma_\Pi^*$ die Menge $\text{SOL}_\Pi(x)$ nichtleer ist.

2 Modelle der Berechenbarkeit

In diesem Kapitel werden zwei unterschiedliche Ansätze beschrieben, die den Begriff der Berechenbarkeit mathematisch exakt formulieren: das Modell der Turingmaschine und das Modell der Random Access Maschine. Beide Modelle sind „theoretische“ Modelle, da man weder eine Turingmaschine noch eine Random Access Maschine (wegen der Modellierung des eingesetzten Speichers) physisch bauen kann. Sie ähneln jedoch sehr der Architektur heutiger Rechner.

Beide Modelle haben ihre Stärken. Die Turingmaschine ist Basismodell in der Automaten-theorie, die ihrerseits eng verknüpft mit der Theorie der Formalen Sprachen ist. Diese wiederum hat erst die Grundlage dazu gelegt, daß wir heute über Programmiersprachen und schnelle Compiler und über mächtige Modellierungswerkzeuge in der Anwendungsentwicklung verfügen. Die Random Access Maschine ist ein Modell eines Rechners einschließlich einer sehr einfachen Assemblersprache, so daß es eher einen mechanischen Zugang zum Begriff der Berechenbarkeit liefert. Es stellt sich heraus, daß die Modelle äquivalent in dem Sinne sind, daß sie in der Lage sind, die gleiche Menge von Funktionen zu berechnen bzw. die gleichen Mengen zu entscheiden (in einem noch zu präzisierenden Sinn).

Ein weiterer Ansatz, der jedoch auch hier nicht vertieft wird, beschäftigt sich damit, eine Programmiersprache auf ihre minimale Anzahl möglicher Anweisungstypen zu reduzieren. Auch hier stellt sich heraus, daß die Berechnungsfähigkeit dieses Modells mit der einer Turingmaschine äquivalent ist.

Auch weitere hier nicht behandelte Ansätze wie die eher mathematisch ausgerichteten Theorien der m -rekursiven Funktionen oder des Churchsche λ -Kalküls haben bisher keinen formalen umfassenderen Berechenbarkeitsbegriff erzeugt. Daher wird die als **Churchsche These** bekannte Aussage als gültig angesehen, nach der das im folgenden behandelte Modell der Turingmaschine die Formalisierung des Begriffs „Algorithmus“ darstellt.

2.1 Deterministische Turingmaschinen

Das hier beschriebene Modell zur Berechenbarkeit wurde 1936 von Alan Turing (1912 – 1954) vorgestellt, und zwar noch vor der Entwicklung der Konzepte moderner Computer. Grundlage ist dabei die Vorstellung, daß eine Berechnung mit Hilfe eines mechanischen Verfahrens durchgeführt wird, wobei Zwischenergebnisse auf einem Rechenblatt notiert und später wieder verwendet werden können. Die seinem Initiator zu Ehren genannte Turingmaschine stellt ein *theoretisches Modell* zur Beschreibung der Berechenbarkeit dar und ist trotz seiner Ähnlichkeit mit heutigen Computern hardwaremäßig *nicht* realisierbar, da es über ei-

nen abzählbar unendlich großen Speicher verfügt. Es hat wesentlichen Einfluß sowohl auf die mathematische Logik als auch auf die Entwicklung heutiger Rechner genommen.

Eine **deterministische k -Band-Turingmaschine (k -DTM)** TM ist definiert durch

$$TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$$

mit:

1. Q ist eine endliche nichtleere Menge: die **Zustandsmenge**
2. Σ ist eine endliche nichtleere Menge: das **Arbeitsalphabet**; Σ enthält alle Zeichen, die in Feldern der Bänder (siehe unten) stehen können
3. $I \subseteq \Sigma$ ist eine endliche nichtleere Menge: das **Eingabealphabet**; mit den Zeichen aus I werden die Wörter der Eingabe gebildet
4. $b \in \Sigma \setminus I$ ist das **Leerzeichen**
5. $q_0 \in Q$ ist der **Anfangszustand** (oder **Startzustand**)
6. $q_{accept} \in Q$ ist der **akzeptierende Zustand** (**Endzustand**)
7. $d : (Q \setminus \{q_{accept}\}) \times \Sigma^k \rightarrow Q \times (\Sigma \times \{L, R, S\})^k$ ist eine partielle Funktion, die **Überföhrungs-funktion**; insbesondere ist $d(q, a_1, \dots, a_k)$ für $q = q_{accept}$ nicht definiert; zu beachten ist, daß d für einige weitere Argumente eventuell nicht definiert ist.

Anschaulich kann man sich die **Arbeitsweise einer k -DTM** wie folgt vorstellen:

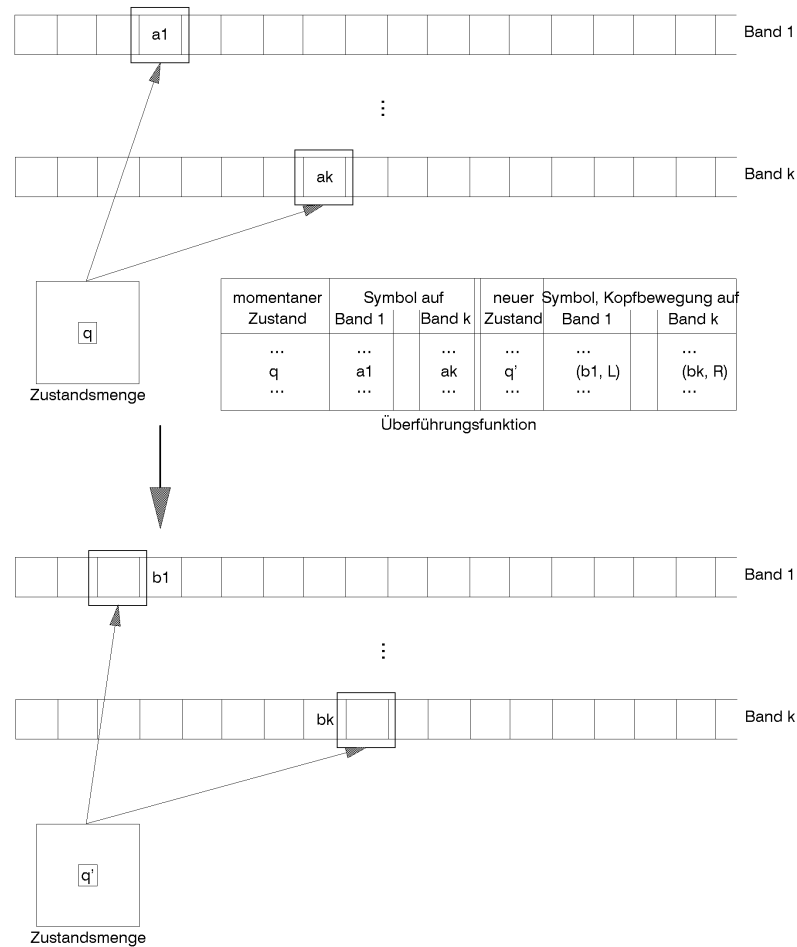
Es gibt k (Speicher-) **Bänder**, die jeweils in Zellen eingeteilt sind und nach rechts unendlich lang sind. Jede Zelle eines jeden Bands kann einen Buchstaben des Arbeitsalphabets aufnehmen. Eine Zelle, die das Leerzeichen enthält, wird als **leere Zelle** bezeichnet. Für jedes Band gibt es genau einen **Schreib/Lesekopf**, der jeweils genau über einer Zelle steht. Dieser kann den Zellinhalt lesen, neu beschreiben und zur linken oder rechten Nachbarzelle übergelien oder auf der Zelle stehenbleiben. Welche Aktion der jeweilige Schreib/Lesekopf unternimmt, wird in der Überföhrungsfunktion d in Abhängigkeit vom Zustand (des Steuerwerks) und der auf den Bändern gelesenen Zeichen angegeben.

Die Arbeitsweise der k -DTM wird durch ein **endliches Steuerwerk** festgelegt, dessen Zustandsüberföhrungen **getaktet** ablaufen und durch d beschrieben werden:

Liest der Kopf des i -ten Bandes den Buchstaben $a_i \in \Sigma$ ($i = 1, \dots, k$), ist das Steuerwerk im Zustand q und ist

$$d(q, a_1, \dots, a_k) = (q', (b_1, d_1), \dots, (b_k, d_k)),$$

dann geht das Steuerwerk in den Zustand q' über, der i -te Kopf schreibt b_i und geht zur linken Nachbarzelle für $d_i = L$ (falls dieses möglich ist), zur rechten Nachbarzelle für $d_i = R$ oder bleibt für $d_i = S$ über der Zelle stehen.



Ist $d(q, a_1, \dots, a_k)$ nicht definiert, so **halt** die k -DTM **im Zustand q an (stoppt im Zustand q)**. Sobald also der akzeptierende Zustand q_{accept} erreicht wird, halt die Turingmaschine an, da $d(q, a_1, \dots, a_k)$ fur $q = q_{accept}$ nicht definiert ist. Sie kann aber auch eventuell vorher in einem anderen Zustand anhalten (namlich dann, wenn $d(q, a_1, \dots, a_k)$ nicht definiert ist), oder sie **kann beliebig lange weiterlaufen (sie halt nicht an)**. Diese Situation tritt beispielsweise dann ein, wenn die Turingmaschine in einen Zustand $q \neq q_{accept}$ kommt und $d(q, a_1, \dots, a_k) = (q, (a_1, S), \dots, (a_k, S))$ ist. Eine andere Moglichkeit eines endlosen Weiterlaufens ergibt sich dann, wenn die Turingmaschine in einen Zustand $q \neq q_{accept}$ kommt, alle Kopfe uber Zellen stehen, die das Leerzeichen b enthalten, rechts dieser Zellen auf allen Bandern nur noch Leerzeichen stehen und $d(q, b, \dots, b) = (q, (b, R), \dots, (b, R))$ ist.

Die Werte $d(q, a_1, \dots, a_k) = (q', (b_1, d_1), \dots, (b_k, d_k))$ der Uberfohrungsfunktion werden hufig in Form einer endlichen Tabelle angegeben:

momentaner Zustand	Symbol unter dem Schreib/Lesekopf auf			neuer Zustand	neues Symbol, Kopfbewegung auf		
	Band 1	...	Band k		Band 1	...	Band k
q	a_1	...	a_k	q'	b_1, d_1	...	b_k, d_k

Eine **Konfiguration** K einer k -DTM TM beschreibt den gegenwärtigen Gesamtzustand von TM , d.h. den gegenwärtigen Zustand zusammen mit den Bandinhalten und den Positionen der Schreib/Leseköpfe:

$$K = (q, (a_1, i_1), \dots, (a_k, i_k)) \text{ mit } q \in Q, a_j \in \Sigma^*, i_j \geq 1 \text{ für } j = 1, \dots, k.$$

Diese Konfiguration K wird folgendermaßen interpretiert:

TM ist im Zustand q , das j -te Band (für $j = 1, \dots, k$) enthält linksbündig die endliche Zeichenkette a_j (jeder Buchstabe von a_j belegt eine Zelle), gefolgt von Zellen, die das Leerzeichen enthalten (leere Zellen), der Schreib/Lesekopf des j -ten Bands steht über der i_j -ten Zelle; für $i_j \in [1: |a_j|]$ ist dieses der i_j -te Buchstabe von a_j , für $i_j \geq |a_j| + 1$ steht der Schreib/Lesekopf hinter a_j über einer Zelle, die das Leerzeichen enthält.

Ist a_j der i_j -te Buchstabe von a_j bzw. $a_j = b$ für $i_j \geq |a_j| + 1$, und ist

$$d(q, a_1, \dots, a_k) = (q', (b_1, d_1), \dots, (b_k, d_k)),$$

dann geht die TM in die **Folgekonfiguration** K' über, die durch

$$K' = (q', (b_1, i'_1), \dots, (b_k, i'_k))$$

definiert wird. Dabei entsteht b_j aus a_j durch Ersetzen von a_j durch b_j . Die Positionen i'_j der Schreib/Leseköpfe der Folgekonfiguration K' lauten

$$i'_j = \begin{cases} i_j + 1 & \text{für } d_j = R \\ i_j & \text{für } d_j = S \\ i_j - 1 & \text{für } d_j = L \text{ und } i_j \geq 2. \end{cases}$$

Man schreibt in diesem Fall

$$K \Rightarrow K'.$$

Die Bezeichnung $K \Rightarrow^* K'$ besagt, daß entweder keine Konfigurationsänderung stattgefunden hat (es ist dann $K = K'$) oder daß es eine Konfiguration K_1 gibt mit $K \Rightarrow K_1$ und $K_1 \Rightarrow^* K'$ (auch geschrieben als $K \Rightarrow K_1 \Rightarrow^* K'$).

Man schreibt $K \Rightarrow^m K'$ mit $m \in \mathbb{N}$, wenn K' aus K durch m Konfigurationsänderungen hervorgegangen ist, d.h. wenn es m Konfigurationen $K_1, \dots, K_m$ gibt mit

$$K \Rightarrow K_1 \Rightarrow \dots \Rightarrow K_m = K'.$$

Für $m = 0$ ist dabei $K = K'$.

Eine Konfiguration K_0 , die den Anfangszustand q_0 und auf dem 1. Band ein Wort $w \in I^*$ enthält, wobei sich auf allen anderen Bändern nur Leerzeichen befinden und die Köpfe über den am weitesten links stehenden Zellen stehen, d.h. eine Konfiguration der Form

$$K_0 = (q_0, (w, 1), (e, 1), \dots, (e, 1)),$$

heißt **Anfangskonfiguration mit Eingabewort** w . Da $w \in I^*$ ist und das Leerzeichen nicht zu I gehört, kann TM (bei entsprechender Definition der Überföhrungsfunktion) das Ende von w , nämlich das erste Leerzeichen im Anschluß an w , erkennen. Im folgenden wird gelegentlich für eine Eingabe $w \in I^*$ auch $w \in \Sigma^*$ geschrieben und dabei implizit angenommen, daß w nur Buchstaben aus I enthält.

Eine Konfiguration K_{accept} , die den akzeptierenden Zustand q_{accept} enthält, d.h. eine Konfiguration der Form

$$K_{accept} = (q_{accept}, (a_1, i_1), \dots, (a_k, i_k)),$$

heißt **akzeptierende Konfiguration (Endkonfiguration)**.

Ein Wort w über dem Eingabealphabet wird von TM **akzeptiert**, wenn gilt:

$$(q_0, (w, 1), (e, 1), \dots, (e, 1)) \Rightarrow^* K_{accept}$$

mit einer Endkonfiguration K_{accept} .

Die von einer k -DTM TM **akzeptierte Sprache** ist die Menge

$$\begin{aligned} L(TM) &= \{ w \mid w \in I^* \text{ und } w \text{ wird von } TM \text{ akzeptiert} \} \\ &= \{ w \mid w \in I^* \text{ und } (q_0, (w, 1), (e, 1), \dots, (e, 1)) \Rightarrow^* (q_{accept}, (a_1, i_1), \dots, (a_k, i_k)) \}. \end{aligned}$$

Zu beachten ist, daß TM eventuell auch dann bereits eine Endkonfiguration erreicht, wenn das Eingabewort w noch gar nicht komplett gelesen ist. Auch in diesem Fall gehört w zu $L(TM)$.

Für $w \notin L(TM)$ hält TM entweder nicht im Zustand q_{accept} , oder TM läuft unendlich lange weiter, d.h. die Überföhrungsfolge $K_0 \Rightarrow K'$ läßt sich unendlich lang fortsetzen, ohne daß eine Konfiguration erreicht wird, die den Endzustand enthält.

Eine 2-DTM Turingmaschine zur Akzeptanz von

$$L(TM) = \{ w \mid w \in \{0,1\}^+ \text{ und } w \text{ ist die Binärdarstellung einer geraden Zahl} \} \cup \{e\}$$

Die 2-DTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ wird gegeben durch

$Q = \{q_0, q_1, q_2\}$, $\Sigma = \{0, 1, b\}$, $I = \{0, 1\}$, $q_{accept} = q_1$ mit der Überföhrungsfunktion d , die durch folgende Tabelle definiert ist:

momentaner Zustand	Symbol unter dem Schreib/Lesekopf auf		neuer Zustand	neues Symbol, Kopfbewegung auf	
	Band 1	Band 2		Band 1	Band 2
q_0	b	b	q_1	b, S	$1, S$
	0	b	q_0	$0, R$	b, S
	1	b	q_2	$1, R$	b, S
q_2	b	b	q_2	b, R	b, S
	0	b	q_0	$0, R$	b, S
	1	b	q_2	$1, R$	b, S

TM stoppt bei Eingabe von w nicht, falls w die Binärdarstellung einer ungeraden Zahl ist.

Eine 2-DTM Turingmaschine zur Akzeptanz der Palindrome über $\{0, 1\}$

Die folgende Turingmaschine 2-DTM TM akzeptiert genau die Palindrome über $\{0, 1\}$:

$$L(TM) = \{ w \mid w \in \{0,1\}^+ \text{ und } w = a_1 \dots a_{n-1} a_n = a_n a_{n-1} \dots a_1 \} \cup \{e\}.$$

$$TM = (Q, \Sigma, I, d, b, q_0, q_{accept}) \text{ mit } Q = \{q_0, \dots, q_5\}, \Sigma = \{0, 1, b, \#\}, I = \{0, 1\}, q_{accept} = q_5.$$

TM arbeitet wie folgt:

1. Die 1. Zelle des 2. Bandes wird mit $\#$ markiert. Dann wird das Eingabewort auf das 2. Band kopiert; der Kopf des 1. Bandes steht jetzt unmittelbar rechts des Eingabeworts.
2. Der Kopf des 2. Bandes wird bis zum Zeichen $\#$ zurückgesetzt.
3. Der Kopf des 1. Bandes wird jeweils um 1 Zelle nach links und der Kopf des 2. Bandes um 1 Zelle nach rechts verschoben. Wenn die von den Köpfen jeweils gelesenen Symbole sämtlich übereinstimmen, ist das Eingabewort ein Palindrom, und die Maschine geht in den Zustand $q_f = q_5$ und stoppt. Sonst stoppt die Maschine in einem von q_f verschiedenen Zustand.

Die Überföhrungsfunktion wird durch folgende Tabelle gegeben:

momentaner Zustand	Symbol unter dem Schreib/Lesekopf auf		neuer Zustand	neues Symbol, Kopfbewegung auf	
	Band 1	Band 2		Band 1	Band 2
q_0	0	b	q_1	0, S	$\#, R$
	1	b	q_1	1, S	$\#, R$
	b	b	q_{accept}	b, S	b, S
q_1	0	b	q_1	0, R	0, R
	1	b	q_1	1, R	1, R
	b	b	q_2	b, S	b, L
q_2	b	0	q_2	b, S	0, L
	b	1	q_2	b, S	1, L
	b	$\#$	q_3	b, L	$\#, R$
q_3	0	0	q_4	0, S	0, R
	1	1	q_4	1, S	1, R
q_4	0	0	q_3	0, L	0, S
	0	1	q_3	0, L	1, S
	1	0	q_3	1, L	0, S
	1	1	q_3	1, L	1, S
	0	b	q_{accept}	0, S	b, S
	1	b	q_{accept}	1, S	b, S

Beobachtet man die Arbeitsweise einer Turingmaschine TM bei einem Eingabewort w , so liegt nach endlich vielen Überführungen eine der folgenden Situationen vor:

1. Fall: TM ist bereits in einem Zustand $q \neq q_{accept}$ stehengeblieben (d.h. $d(q, \dots)$ ist nicht definiert). Dann ist $w \notin L(TM)$.
2. Fall: TM ist bereits im Zustand q_{accept} stehengeblieben. Dann ist $w \in L(TM)$.
3. Fall: TM ist noch nicht stehengeblieben, d.h. TM befindet sich in einem Zustand $q \neq q_{accept}$, für den $d(q, \dots)$ definiert ist. Dann ist noch nicht entschieden, ob $w \in L(TM)$ oder $w \notin L(TM)$ gilt. TM ist eventuell noch nicht lange genug beobachtet worden. Es ist nicht „vorhersagbar“, wie lange TM beobachtet werden muß, um eine Entscheidung zu treffen (es ist **algorithmisch unentscheidbar**).

Eine Turingmaschine TM kann als **Berechnungsvorschrift** definiert werden: Das 1. Band wird als **Eingabeband** und ein Band, etwa das k -te Band, als **Ausgabeband** ausgezeichnet. Die Turingmaschine TM **berechnet eine partielle Funktion** $f_{TM} : I^* \rightarrow \Sigma^*$, wenn gilt:

Startet TM im Anfangszustand mit w , d.h. startet TM mit w auf dem Eingabeband im Zustand q_0 (alle anderen Bänder sind leer), und stoppt TM nach endlich vielen Schritten im akzeptierenden Zustand q_{accept} , dann wird die Bandinschrift y des Ausgabebands von TM als Funktionswert $y = f_{TM}(w)$ interpretiert. Falls TM überhaupt nicht oder nicht im Endzustand stehenbleibt, dann ist $f_{TM}(w)$ nicht definiert. Daher ist f_{TM} eine partielle Funktion.

Eine 2-DTM zur Berechnung der (totalen) Funktion $f : \begin{cases} \mathbb{N} & \rightarrow & \mathbb{N} \\ n & \rightarrow & n+1 \end{cases}$

Die folgende 2-DTM berechnet die (totale) Funktion $f : \begin{cases} \mathbb{N} & \rightarrow & \mathbb{N} \\ n & \rightarrow & n+1 \end{cases}$:

Hierbei wird ein $n \in \mathbb{N}$ als Zeichenkette in seiner Binärdarstellung auf das Eingabeband geschrieben; die höchstwertige Stelle steht dabei ganz links. Der Wert 0 wird als Zeichenkette 0 eingegeben, alle anderen Werte $n \in \mathbb{N}$ ohne führende Nullen. Entsprechend steht abschließend $n+1$ in seiner Binärdarstellung ohne führende Nullen auf dem Ausgabeband (2. Band).

$TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ mit $Q = \{q_0, \dots, q_{19}\}$, $\Sigma = \{0, 1, b, \#\}$, $I = \{0, 1\}$, $q_{accept} = q_{19}$.

TM arbeitet wie folgt:

1. Auf das 2. Band wird zunächst das Zeichen # geschrieben, das das linke Ende des 2. Bandes markieren soll. Dann wird auf das 2. Band eine 0 geschrieben und das Eingabewort w auf das 2. Band kopiert (auf dem 2. Band steht jetzt $\#0w$, d.h. das Eingabewort w ist durch eine führende 0 ergänzt worden).
2. Auf dem 2. Band werden von rechts alle 1'en in 0'en invertiert, bis die erste 0 erreicht ist; diese wird durch 1 ersetzt (Addition $n := n + 1$).
3. Der Kopf des 2. Bandes wird auf die Anfangsmarkierung # zurückgesetzt und geprüft, ob rechts dieses Zeichens eine 0 steht (das bedeutet, daß die anfangs an w angefügte führende 0 wieder entfernt werden muß).
4. In diesem Fall wird der Inhalt des 2. Bandes komplett um zwei Position nach links verschoben. Dadurch werden die führende 0 und das Zeichen # auf dem 2. Band entfernt. Es ist zu beachten, daß am rechten Ende des 2. Bandes 2 Leerzeichen gesetzt werden.
5. Andernfalls wird der Inhalt des 2. Bandes um 1 Position nach links verschoben und dadurch das Zeichen # entfernt.

Die Überföhrungsfunktion wird durch folgende Tabelle gegeben:

momentaner Zustand	Symbol unter dem Schreib/Lesekopf auf		neuer Zustand	neues Symbol, Kopfbewegung auf		Bemerkung
	Band 1	Band 2		Band 1	Band 2	
q_0	0	b	q_1	0, S	#, R	linkes Ende für $n+1$ markieren
	1	b	q_1	1, S	#, R	
q_1	0	b	q_2	0, S	0, R	führende 0 erzeugen
	1	b	q_2	1, S	0, R	
q_2	0	b	q_2	0, R	0, R	Inhalt des Eingabebandes auf 2. Band kopieren
	1	b	q_2	1, R	1, R	
	b	b	q_3	b , S	b , L	
q_3	b	0	q_4	b , S	1, S	1 addieren: anhängende 1'en invertieren, erste 0 von rechts invertieren
	b	1	q_3	b , S	0, L	
q_4	b	0	q_4	b , S	0, L	auf # zurückgehen
	b	1	q_4	b , S	1, L	
	b	#	q_5	b , S	#, R	

../..

q_5	b	0	q_6	b, S	0, R	muß führende 0 entfernt werden?
	b	1	q_{15}	b, S	1, S	
q_6	b	0	q_7	b, S	0, L	gelesenes Zeichen im Zustand merken
	b	1	q_8	b, S	1, L	
	b	b	q_{13}	b, S	b, L	
q_7	b	0	q_9	b, S	0, L	2. Kopf um 1 Positi- on nach links setzen
	b	1	q_9	b, S	1, L	
q_8	b	0	q_{10}	b, S	0, L	2. Kopf um 1 Positi- on nach links setzen
	b	1	q_{10}	b, S	1, L	
q_9	b	0	q_{11}	b, S	0, R	0 schreiben
	b	1	q_{11}	b, S	0, R	
	b	#	q_{11}	b, S	0, R	
q_{10}	b	0	q_{11}	b, S	1, R	1 schreiben
	b	1	q_{11}	b, S	1, R	
	b	#	q_{11}	b, S	1, R	
q_{11}	b	0	q_{12}	b, S	0, R	2. Kopf um 1 Positi- on nach rechts set- zen
	b	1	q_{12}	b, S	1, R	
q_{12}	b	0	q_6	b, S	0, R	2. Kopf 1 weitere Position nach rechts setzen
	b	1	q_6	b, S	1, R	
q_{13}	b	0	q_{14}	b, S	b, L	letztes Zeichen löschen und nach links gehen
	b	1	q_{14}	b, S	b, L	
q_{14}	b	0	$q_{19} = q_f$	b, S	b, S	letztes Zeichen löschen und STOP
	b	1	$q_{19} = q_f$	b, S	b, S	
q_{15}	b	0	q_{16}	b, S	0, L	gelesenes Zeichen im Zustand merken
	b	1	q_{17}	b, S	1, L	
	b	b	q_{14}	b, S	b, L	
q_{16}	b	0	q_{18}	b, S	0, R	0 schreiben
	b	1	q_{18}	b, S	0, R	
	b	#	q_{18}	b, S	0, R	

../..

q_{17}	b	0	q_{18}	b, S	1, R	1 schreiben
	b	1	q_{18}	b, S	1, R	
	b	#	q_{18}	b, S	1, R	
q_{18}	b	0	q_{15}	b, S	0, R	2. Kopf um 1 Posi-
	b	1	q_{15}	b, S	1, R	tion nach rechts set-
						zen

Die Konzepte der Berechnung partieller Funktionen und das Akzeptieren von Sprachen mittels Turingmaschinen sind äquivalent:

Satz 2.1-1:

Berechnet die Turingmaschine TM die partielle Funktion $f_{TM} : I^* \rightarrow \Sigma^*$, dann kann man eine Turingmaschine TM' konstruieren mit $L(TM') = \{ w \# y \mid y = f_{TM}(w) \}$.

Akzeptiert die Turingmaschine TM die Sprache $L(TM)$, dann läßt sich TM so modifizieren, daß sie die partielle Funktion $f_{TM} : I^* \rightarrow \Sigma^*$ berechnet, die definiert ist durch

$f_{TM}(w) = y$ genau dann, wenn gilt:

$$(q_0, (w, 1), (e, 1), \dots, (e, 1), (e, 1)) \Rightarrow^* (q_{accept}, (a_1, i_1), (a_2, i_2), \dots, (a_{k-1}, i_{k-1}), (y, |y| + 1))$$

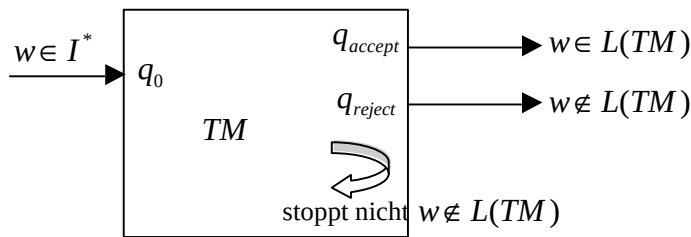
Die Aussage läßt die Interpretation zu, daß die Verifikation eines Funktionsergebnisses genauso komplex ist wie die Berechnung des Funktionsergebnisses selbst.

Die Überföhrungsfunktion d einer Turingmaschine TM werde folgendermaßen modifiziert: Die Zustandsmenge Q wird um einen neuen Zustand q_{reject} erweitert. Für alle bisherigen Zustände $q \in Q$ mit $q \neq q_{accept}$, für die $d(q, \dots)$ nicht definiert ist, werden in d die Zeilen

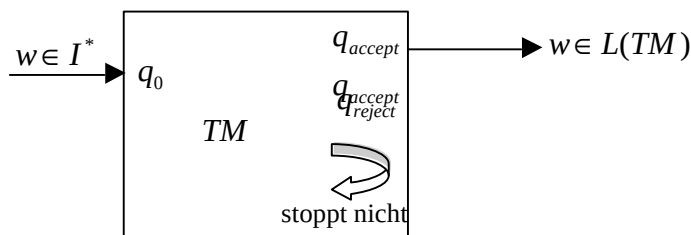
$$d(q, a_1, \dots, a_k) = (q_{reject}, (a_1, S), \dots, (a_k, S)) \text{ mit } a_i \in \Sigma \text{ für } i = 1, \dots, k$$

aufgenommen. Für q_{reject} ist d nicht definiert. Kommt die so modifizierte Turingmaschine bei Eingabe eines Wortes $w \in I^*$ in einen Zustand $q \in Q$ mit $q \neq q_{accept}$, für die $d(q, \dots)$ bisher nicht definiert war (es ist dann $w \notin L(TM)$), so macht die modifizierte Turingmaschine noch eine Überföhrung, ohne die Bandinhalte zu ändern, und stoppt im Zustand q_{reject} , ohne das Wort w zu akzeptieren. Man kann daher annehmen, daß eine Turingmaschine TM so definierte Zustände q_{accept} und q_{reject} enthält. Der Zustand q_{accept} heißt **akzeptierender Zustand**, der Zustand q_{reject} heißt **nicht-akzeptierender (verwerfender) Zustand**. Startet TM bei Ein-

gabe eines Wortes $w \in I^*$ im Anfangszustand q_0 , so zeigt sie folgendes Verhalten: Entweder kommt TM in den Zustand q_{accept} , dann ist $w \in L(TM)$ und **die Eingabe w wird akzeptiert**; oder TM kommt in den Zustand q_{reject} , dann ist $w \notin L(TM)$ und **die Eingabe w wird verworfen**; oder TM läuft unendlich lange weiter, dann ist ebenfalls $w \notin L(TM)$. Man kann TM daher als Blackbox darstellen, die eine Eingangsschnittstelle besitzt, über die im Anfangszustand q_0 ein Wort $w \in I^*$ eingegeben wird und der Start der Berechnung gemäß der Überföhrungsfunktion d erfolgt. Sie besitzt zwei „aktivierbare“ Ausgangsschnittstellen: Die erste Ausgangsschnittstelle wird von TM dann aktiviert, wenn TM bei der Berechnung den Zustand q_{accept} erreicht und stoppt; es ist dann $w \in L(TM)$. Die zweite Ausgangsschnittstelle wird von TM aktiviert, wenn TM bei der Berechnung den Zustand q_{reject} erreicht und stoppt; es ist dann $w \notin L(TM)$. Wird keine Ausgangsschnittstelle aktiviert, weil TM nicht anhält, dann ist ebenfalls $w \notin L(TM)$. Eine graphische Repräsentation von TM sieht wie folgt aus.



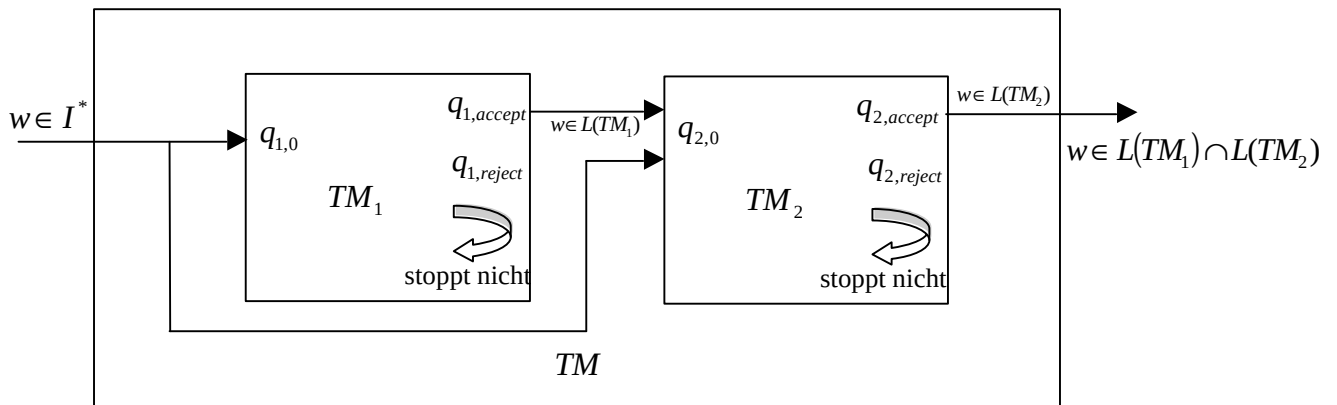
Da die Turingmaschine TM dazu verwendet wird, um festzustellen, ob ein Eingabewort $w \in I^*$ von ihr akzeptiert wird, kann man TM vereinfacht auch folgendermaßen darstellen:



Turingmaschinen können zu neuen Turingmaschinen zusammengesetzt werden: Es seien zwei Turingmaschinen $TM_1 = (Q_1, \Sigma_1, I, d_1, b, q_{1,0}, q_{1,accept})$ und $TM_2 = (Q_2, \Sigma_2, I, d_2, b, q_{2,0}, q_{2,accept})$ mit disjunkten Zustandsmengen ($Q_1 \cap Q_2 = \emptyset$) und demselben Eingabealphabet $I \subseteq \Sigma_1$ und $I \subseteq \Sigma_2$ und jeweils getrennten Bändern gegeben. Die Anzahl der Bänder von TM_1 sei k_1 , die

von TM_2 sei k_2 . Jeweils das 1. Band ist das Eingabeband. Beispiele für die Möglichkeiten der **Zusammensetzung** sind:

- Man kann eine neue Turingmaschine TM durch **Hintereinanderschaltung** von TM_1 und TM_2 konstruieren, die die Bänder von TM_1 und TM_2 umfaßt, d.h. $k_1 + k_2$ viele Bänder besitzt, und folgendermaßen arbeitet: Eine Eingabe $w \in I^*$ für die zusammengesetzte Turingmaschine TM wird auf das 1. Band von TM_1 gegeben und auf das 1. Band von TM_2 kopiert. Jetzt wird zunächst das Verhalten von TM_1 auf w simuliert. Falls TM_1 das Wort w akzeptiert, d.h. in den Zustand $q_{1,accept}$ gelangt, wird das Verhalten von TM_2 auf w simuliert. TM akzeptiert das Wort w , falls TM_2 das Wort w akzeptiert. Es gilt:
 $L(TM) = L(TM_1) \cap L(TM_2)$. Graphisch läßt sich TM wie folgt darstellen.



Die formale Notation von TM mit Hilfe der Überföhrungsfunktion ist etwas mühsam:

Es ist $TM = (Q, \Sigma, I, d, b, q_{1,0}, q_{2,accept})$. Dabei ist

$Q = Q_1 \cup Q_2 \cup \{q_{copy}, q_{skip}\}$ die Zustandsmenge von TM mit zwei neuen Zuständen q_{copy} und q_{skip} , die nicht in $Q_1 \cup Q_2$ enthalten sind,

$\Sigma = \Sigma_1 \cup \Sigma_2 \cup \{\#\}$ das Arbeitsalphabet von TM mit einem neuen Symbol $\#$, das nicht in $\Sigma_1 \cup \Sigma_2$ enthalten ist,

$d : (Q \setminus \{q_{2,accept}\}) \times \Sigma^{k_1+k_2} \rightarrow Q \times (\Sigma \times \{L, R, S\})^{k_1+k_2}$ die Überföhrungsfunktion, die sich aus den Überföhrungsfunktionen d_1 und d_2 wie folgt ergibt:

Die Bänder von TM werden von 1 bis $k_1 + k_2$ numeriert; die ersten k_1 Bänder sind die ursprünglichen Bänder von TM_1 . Insbesondere ist das Eingabeband von TM das ursprüngliche Eingabeband von TM_1 . Das Band mit der Nummer $k_1 + 1$ ist das ursprüngliche Eingabeband von TM_2 . Ein Eingabewort $w \in I^*$ wird auf das erste Band von TM gegeben.

Es sei $n = |w|$. In die erste Zelle des Bands mit der Nummer $k_1 + 1$ wird das Zeichen # geschrieben und w auf dieses Band kopiert. Das Zeichen # dient dazu, das linke Ende des Bands mit der Nummer $k_1 + 1$ zu erkennen. Nach diesem Kopiervorgang stehen die Köpfe des ersten und des $(k_1 + 1)$ -ten Bands jeweils über der $(n+1)$ -ten Zelle, alle übrigen Köpfe wurden nicht bewegt. Das Ende des Kopiervorgangs wurde dadurch erkannt, daß auf dem ersten Band das Leerzeichen gelesen wurde. Nun werden die Köpfe des ersten und des $(k_1 + 1)$ -ten Bands beide zurück auf das erste Zeichen von w gesetzt.

Für den Kopier- und den Rücksetzvorgang der Köpfe werden folgende Zeilen in die Überföhrungsfunktion d von TM aufgenommen:

$$d \left(q_{1,0}, \underbrace{a, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2} \right) = \left(q_{copy}, \underbrace{(a, S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(\#, R), (b, S), \dots, (b, S)}_{k_2} \right) \quad \text{für alle}$$

$a \in \Sigma_1$ (b ist das Leerzeichen),

$$d \left(q_{copy}, \underbrace{a, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2} \right) = \left(q_{copy}, \underbrace{(a, R), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(a, R), (b, S), \dots, (b, S)}_{k_2} \right) \quad \text{für alle}$$

$a \in I$ (man beachte, daß $a = \text{Leerzeichen}$ hierbei nicht vorkommt),

$$d \left(q_{copy}, \underbrace{b, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2} \right) = \left(q_{skip}, \underbrace{(b, S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(b, L), (b, S), \dots, (b, S)}_{k_2} \right),$$

$$d \left(q_{skip}, \underbrace{a', b, \dots, b}_{k_1}, \underbrace{a, b, \dots, b}_{k_2} \right) = \left(q_{skip}, \underbrace{(a', L), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(a, L), (b, S), \dots, (b, S)}_{k_2} \right) \quad \text{für alle}$$

$a' \in I \cup \{b\}$ und $a \in I$,

$$d \left(q_{skip}, \underbrace{a', b, \dots, b}_{k_1}, \underbrace{\#, b, \dots, b}_{k_2} \right) = \left(q_{1,0}, \underbrace{(a', S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(\#, R), (b, S), \dots, (b, S)}_{k_2} \right) \quad \text{für alle}$$

$a' \in I \cup \{b\}$.

Nun wird das Verhalten von TM_1 simuliert, wobei die Inhalte der Bänder mit den Nummern $k_1 + 1$ bis $k_1 + k_2$ (entsprechend den ursprünglichen Bändern von TM_2) und die Positionen der Köpfe auf diesen Bändern nicht verändert werden. Die dafür erforderlichen Zeilen in der Überföhrungsfunktion d lauten:

$$d \left(q_1, \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{a, b, \dots, b}_{k_2} \right) = \left(q'_1, \underbrace{(b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1})}_{k_1}, \underbrace{(a, S), (b, S), \dots, (b, S)}_{k_2} \right) \quad \text{für}$$

jeden Eintrag $d_1(q_1, a_1, a_2, \dots, a_{k_1}) = (q'_1, (b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1}))$ in der Überföhrungsfunktion d_1 von TM_1 und $a \in \Sigma_1$.

Falls bei der Simulation der akzeptierende Zustand $q_{1,accept}$ von TM_1 erreicht wird, d.h. $w \in L(TM_1)$, wird die Simulation des Verhaltens von TM_2 auf w gestartet. Auf den ersten k_1 ändert sich dabei nichts mehr. Folgende Zeilen werden in die Überföhrungsfunktion d aufgenommen:

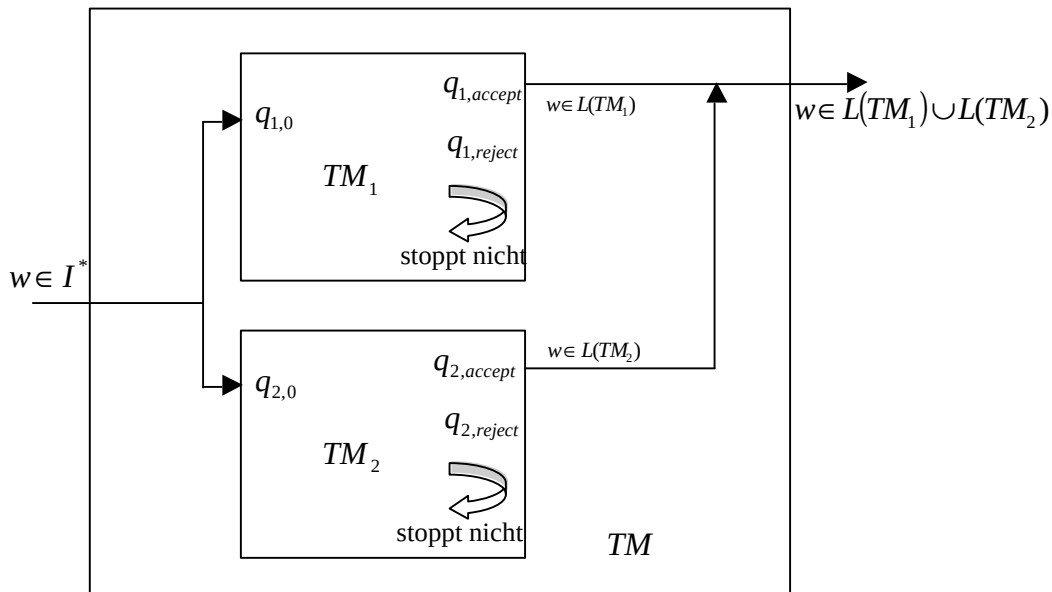
$$d \left(q_{1,accept}, \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{a, b, \dots, b}_{k_2} \right) = \left(q_{2,0}, \underbrace{(a_1, S), (a_2, S), \dots, (a_{k_1}, S)}_{k_1}, \underbrace{(a, S), (b, S), \dots, (b, S)}_{k_2} \right)$$

mit $a_1 \in \Sigma_1, \dots, a_{k_1} \in \Sigma_1$ und $a \in \Sigma_1$ und

$$d \left(q_2, \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2} \right) = \left(q'_2, \underbrace{(a_1, S), (a_2, S), \dots, (a_{k_1}, S)}_{k_1}, \underbrace{(h_1, d_1), (h_2, d_2), \dots, (h_{k_2}, d_{k_2})}_{k_2} \right)$$

für jeden Eintrag $d_2(q_2, g_1, g_2, \dots, g_{k_2}) = (q'_2, (h_1, d_1), (h_2, d_2), \dots, (h_{k_2}, d_{k_2}))$ in der Überföhrungsfunktion d_2 von TM_2 und $a_1 \in \Sigma_1, \dots, a_{k_1} \in \Sigma_1$.

- Durch **Parallelschaltung** kann man aus TM_1 und TM_2 eine neue Turingmaschine TM bilden: ein Wort w wird sowohl auf das 1. Band von TM_1 als auch auf das 1. Band von TM_2 gegeben und das Verhalten beider Turingmaschinen simultan simuliert. Sobald eine der beiden Turingmaschinen w akzeptiert, wird w von TM akzeptiert. Es gilt:
 $L(TM) = L(TM_1) \cup L(TM_2)$.



Auch hier ist die formale Notation von TM mit Hilfe der Überföhrungsfunktion etwas mühsam:

Es ist $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$. Dabei ist

$Q = (Q_1 \times Q_2) \cup \{(q_{copy}, q_{2,0}), (q_{skip}, q_{2,0}), (q_{acc}, q_{acc})\}$ die Zustandsmenge von TM mit zwei neuen Zuständen q_{copy} und q_{skip} , die nicht in Q_1 enthalten sind, und einem neuen Zustand q_{acc} , der nicht in $Q_1 \cup Q_2$ enthalten ist (die Zustände bestehen jetzt aus Paaren von Zuständen der Turingmaschinen TM_1 und TM_2 und drei neuen Zuständen),

$\Sigma = \Sigma_1 \cup \Sigma_2 \cup \{\#\}$ das Arbeitsalphabet von TM mit einem neuen Symbol $\#$, das nicht in Σ_2 enthalten ist,

$d : (Q \setminus \{(q_{acc}, q_{acc})\}) \times \Sigma^{k_1+k_2} \rightarrow Q \times (\Sigma \times \{L, R, S\})^{k_1+k_2}$ die Überföhrungsfunktion, die sich aus den Überföhrungsfunktionen d_1 und d_2 ergibt (siehe unten),

$q_0 = (q_{1,0}, q_{2,0})$ der Anfangszustand und

$q_{accept} = (q_{acc}, q_{acc})$ der akzeptierende Zustand von TM .

Die Bänder von TM werden wieder von 1 bis $k_1 + k_2$ numeriert; die ersten k_1 Bänder sind die ursprünglichen Bänder von TM_1 . Insbesondere ist das Eingabeband von TM das ursprüngliche Eingabeband von TM_1 . Das Band mit der Nummer $k_1 + 1$ ist das ursprüngliche Eingabeband von TM_2 . Ein Eingabewort $w \in I^*$ wird auf das erste Band von TM gegeben. Wieder wird in die erste Zelle des Bands mit der Nummer $k_1 + 1$ das Zeichen $\#$ geschrieben und w auf dieses Band kopiert. Anschließend werden der Kopf des ersten Bands auf den Bandanfang und der Kopf des $(k_1 + 1)$ -ten Bands auf die zweite Zelle (rechte Nachbarzelle der Zelle, die das Zeichen $\#$ enthält) zurückgesetzt. Der Vorgang verläuft ähnlich dem Vorgang, der bei der Hintereinanderschaltung von TM_1 und TM_2 beschrieben wurde. Die erforderlichen Einträge in der Überföhrungsfunktion lauten jetzt:

$$d \left((q_{1,0}, q_{2,0}), \underbrace{a, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2} \right) = \left((q_{copy}, q_{2,0}), \underbrace{(a, S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(\#, R), (b, S), \dots, (b, S)}_{k_2} \right)$$

für alle $a \in \Sigma_1$ (b ist das Leerzeichen),

$$d \left((q_{copy}, q_{2,0}), \underbrace{a, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2} \right) = \left((q_{copy}, q_{2,0}), \underbrace{(a, R), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(a, R), (b, S), \dots, (b, S)}_{k_2} \right)$$

für alle $a \in I$ (man beachte, daß $a =$ Leerzeichen hierbei nicht vorkommt),

$$\begin{aligned}
d\left((q_{copy}, q_{2,0}), \underbrace{b, b, \dots, b}_{k_1}, \underbrace{b, \dots, b}_{k_2}\right) &= \left((q_{skip}, q_{2,0}), \underbrace{(b, S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(b, L), (b, S), \dots, (b, S)}_{k_2}\right) \\
d\left((q_{skip}, q_{2,0}), \underbrace{a', b, \dots, b}_{k_1}, \underbrace{a, b, \dots, b}_{k_2}\right) &= \left((q_{skip}, q_{2,0}), \underbrace{(a', L), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(a, L), (b, S), \dots, (b, S)}_{k_2}\right) \\
&\text{für alle } a' \in I \cup \{b\} \text{ und } a \in I, \\
d\left((q_{skip}, q_{2,0}), \underbrace{a', b, \dots, b}_{k_1}, \underbrace{\#, b, \dots, b}_{k_2}\right) &= \left((q_{1,0}, q_{2,0}), \underbrace{(a', S), (b, S), \dots, (b, S)}_{k_1}, \underbrace{(\#, R), (b, S), \dots, (b, S)}_{k_2}\right) \\
&\text{für alle } a' \in I \cup \{b\}.
\end{aligned}$$

Nun wird parallel das Verhalten von TM_1 und von TM_2 simuliert, wobei für die Simulation von TM_1 nur auf die Inhalte der Bänder mit den Nummern 1 bis k_1 (entsprechend den ursprünglichen Bändern von TM_1) und für die Simulation von TM_2 nur auf die Inhalte der Bänder mit den Nummern $k_1 + 1$ bis $k_1 + k_2$ (entsprechend den ursprünglichen Bändern von TM_2) zugegriffen wird. Die dafür erforderlichen Zeilen in der Überföhrungsfunktion d lauten:

$$\begin{aligned}
&d\left((q_1, q_2), \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2}\right) \\
&= \left((q'_1, q'_2), \underbrace{(b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1})}_{k_1}, \underbrace{(h_1, e_1), (h_2, e_2), \dots, (h_{k_2}, e_{k_2})}_{k_2}\right), \\
&\text{falls } d_1(q_1, a_1, a_2, \dots, a_{k_1}) = (q'_1, (b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1})) \text{ in der Überföhrungsfunktion} \\
&d_1 \text{ von } TM_1 \text{ und } d_2(q_2, g_1, g_2, \dots, g_{k_2}) = (q'_2, (h_1, e_1), (h_2, e_2), \dots, (h_{k_2}, e_{k_2})) \text{ in der Überföhrungsfunktion} \\
&d_2 \text{ von } TM_2 \text{ ist.}
\end{aligned}$$

Falls bei der Simulation ein Zustand der Form $(q_{1, reject}, q_2)$ mit $q_2 \in Q_2$ bzw. der Form $(q_1, q_{2, reject})$ mit $q_1 \in Q_1$ erreicht wird, muß sichergestellt werden, daß die Simulation von TM_2 bzw. von TM_1 weiterläuft. Daher werden auch folgende Einträge in die Überföhrungsfunktion d aufgenommen:

$$\begin{aligned}
&d\left((q_{1, reject}, q_2), \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2}\right) \\
&= \left((q_{1, reject}, q'_2), \underbrace{(a_1, S), (a_2, S), \dots, (a_{k_1}, S)}_{k_1}, \underbrace{(h_1, e_1), (h_2, e_2), \dots, (h_{k_2}, e_{k_2})}_{k_2}\right)
\end{aligned}$$

mit $a_1 \in \Sigma_1, \dots, a_{k_1} \in \Sigma_1$ und falls $d_2(q_2, g_1, g_2, \dots, g_{k_2}) = (q'_2, (h_1, e_1), (h_2, e_2), \dots, (h_{k_2}, e_{k_2}))$ in der Überföhrungsfunktion d_2 von TM_2 ist und

$$d \left((q_1, q_{2, \text{reject}}), \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2} \right) = \left((q'_1, q_{2, \text{reject}}), \underbrace{(b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1})}_{k_1}, \underbrace{(g_1, S), (g_2, S), \dots, (g_{k_2}, S)}_{k_2} \right)$$

für jeden Eintrag $d_1(q_1, a_1, a_2, \dots, a_{k_1}) = (q'_1, (b_1, d_1), (b_2, d_2), \dots, (b_{k_1}, d_{k_1}))$ in der Überföhrungsfunktion d_1 von TM_1 und $g_1 \in \Sigma_2, \dots, g_{k_2} \in \Sigma_2$.

Schließölich soll TM die Eingabe w akzeptieren, wenn TM_1 oder TM_2 die Eingabe akzeptiert. Folgende Einträge werden daher in die Überföhrungsfunktion von TM aufgenommen:

$$d \left((q_1, q_{2, \text{accept}}), \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2} \right) = \left((q_{\text{acc}}, q_{\text{acc}}), \underbrace{(a_1, S), (a_2, S), \dots, (a_{k_1}, S)}_{k_1}, \underbrace{(g_1, S), (g_2, S), \dots, (g_{k_2}, S)}_{k_2} \right) \text{ und}$$

$$d \left((q_1, q_{2, \text{accept}}), \underbrace{a_1, a_2, \dots, a_{k_1}}_{k_1}, \underbrace{g_1, g_2, \dots, g_{k_2}}_{k_2} \right) = \left((q_{\text{acc}}, q_{\text{acc}}), \underbrace{(a_1, S), (a_2, S), \dots, (a_{k_1}, S)}_{k_1}, \underbrace{(g_1, S), (g_2, S), \dots, (g_{k_2}, S)}_{k_2} \right)$$

für alle $q_1 \in Q_1, q_2 \in Q_2, a_1 \in \Sigma_1, \dots, a_{k_1} \in \Sigma_1$ und $g_1 \in \Sigma_2, \dots, g_{k_2} \in \Sigma_2$.

- Dadurch, daß man die Turingmaschine TM_2 als Unterprogramm in die Berechnung der Turingmaschine TM_1 einsetzt, erhält man eine neue Turingmaschine TM , die prinzipiell folgendermaßen abläuft: Eine Eingabe $w \in I^*$ wird in TM_1 eingegeben. Ein Band von TM_1 wird als „Parameterübergabeband“ für TM_2 ausgezeichnet. Sobald TM_1 in einen als Parameterübergabezustand ausgezeichneten Zustand q_2 gelangt, wird der Inhalt des Parameterübergabebands auf das erste Band von TM_2 kopiert. Wird diese Eingabe von TM_2 akzeptiert, wobei bei Akzeptanz der Inhalt des k_2 -ten Bands y lautet, wird das Parameterübergabeband gelöscht, y darauf kopiert, das erste Band von TM_2 gelöscht und die Be-

rechnung von TM_1 fortgesetzt. Die Ausformulierung der Details dieser Konstruktion werden dem Leser als Übung überlassen.

Obige Überlegungen und Ausführungen in Kapitel 3.1 zeigen:

Satz 2.1-2:

Die Klasse der von Turingmaschinen akzeptierten Mengen ist gegenüber Vereinigungsbildung und Schnittbildung abgeschlossen.

Die Klasse der von Turingmaschinen akzeptierten Mengen ist gegenüber Komplementbildung nicht abgeschlossen: Wenn $L \subseteq \Sigma^*$ von einer Turingmaschine TM akzeptiert wird, d.h. $L = L(TM)$, muß das Komplement $\Sigma^* \setminus L$ von L nicht auch notwendigerweise von einer Turingmaschine akzeptiert werden.

Für eine k -DTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ und eine Eingabe $w \in L(TM)$ gelte

$$(q_0, (w, 1), (e, 1), \dots, (e, 1)) \Rightarrow^m K_{accept}$$

mit einer Endkonfiguration $K_{accept} = (q_{accept}, (a_1, i_1), \dots, (a_k, i_k))$. Dann wird durch $t_{TM}(w) = m$ eine partielle Funktion $t_{TM} : \Sigma^* \rightarrow \mathbb{N}$ definiert, die angibt, **wieviele Überführungen TM macht, um w zu akzeptieren**. Es handelt sich hierbei um eine partielle Funktion, da t_{TM} für Wörter $w \notin L(TM)$ nicht definiert ist.

Die **Zeitkomplexität** von TM (**im schlechtesten Fall, worst case**) wird definiert durch die partielle Funktion $T_{TM} : \mathbb{N} \rightarrow \mathbb{N}$ mit

$$T_{TM}(n) = \max \{ t_{TM}(w) \mid w \in \Sigma^* \text{ und } |w| \leq n \}.$$

Bemerkung: In der Literatur findet man auch die Definition

$$T_{TM}(n) = \max \{ t_{TM}(w) \mid w \in \Sigma^* \text{ und } |w| = n \}$$

Entsprechend kann man die **Platzkomplexität** von TM (**im schlechtesten Fall, worst case**) $S_{TM}(n)$ als den maximalen Abstand vom linken Ende eines Bandes definieren, den ein Schreib/Lesekopf bei der Akzeptanz eines Worts $w \in L(TM)$ mit $|w| \leq n$ erreicht.

Beispiele:

Die oben angegebene Turingmaschine zur Akzeptanz der Palindrome über $\{0,1\}$ hat eine Zeitkomplexität der Größe $T_{TM}(n) = 4n + 3$ und eine Platzkomplexität $S_{TM}(n) = n + 2$.

Die oben angegebene Turingmaschine zur Berechnung der Funktion $f: \begin{cases} \mathbf{N} & \rightarrow & \mathbf{N} \\ n & \rightarrow & n+1 \end{cases}$ hat folgende Zeitkomplexität: Der Wert n wird in Binärdarstellung auf das Eingabeband gegeben. Die Eingabe $w = \text{bin}(n)$ belegt daher für $n \geq 1$ $|w| = |\text{bin}(n)| = \lfloor \log_2(n) \rfloor + 1$ viele Zeichen bzw. ein Zeichen für $n = 0$. Dann führt die Turingmaschine $O(|w|) = O(\log(n))$ viele Schritte aus. Die Zeitkomplexität ist linear in der Länge der Eingabe.

Bei der Betrachtung der Zeitkomplexität werden in diesem Beispiel also die **Bitoperationen** gezählt, die benötigt werden, um $f(n)$ zu berechnen, wenn n in Binärdarstellung gegeben ist. Die Turingmaschine dieses Beispiels läßt sich leicht zu einer Turingmaschine modifizieren,

die die Funktion $SUM: \begin{cases} \mathbf{N} \times \mathbf{N} & \rightarrow & \mathbf{N} \\ (n, m) & \rightarrow & n + m \end{cases}$ berechnet. Hierbei werden die Zahlen n und m in

Binärdarstellung, getrennt durch das Zeichen #, auf das Eingabeband geschrieben, d.h. die Eingabe hat die Form $w = \text{bin}(n)\#\text{bin}(m)$. Dann werden die Zahlen n und m stellengerecht

addiert. Auch für die arithmetischen Funktionen $DIF: \begin{cases} \mathbf{N} \times \mathbf{N} & \rightarrow & \mathbf{N} \\ (n, m) & \rightarrow & \begin{cases} n - m & \text{für } n \geq m, \\ 0 & \text{für } n < m \end{cases} \end{cases}$,

$MULT: \begin{cases} \mathbf{N} \times \mathbf{N} & \rightarrow & \mathbf{N} \\ (n, m) & \rightarrow & n \cdot m \end{cases}$ und $DIV: \begin{cases} \mathbf{N} \times \mathbf{N}_{>0} & \rightarrow & \mathbf{N} \\ (n, m) & \rightarrow & \lfloor n/m \rfloor \end{cases}$ lassen sich entsprechende Tu-

ringmaschinen angeben. Dabei wird beispielsweise die Berechnung von $MULT$ auf sukzessive Anwendung der Berechnung für SUM zurückgeführt. In jedem Fall erfolgt die Eingabe in der Form $w = \text{bin}(n)\#\text{bin}(m)$, und die Zeitkomplexitäten geben an, wieviele Bitoperationen zur Berechnung der jeweiligen Funktionen erforderlich sind. Die Ergebnisse sind in folgendem Satz zusammengefaßt.

Satz 2.1-3:

Es seien n und m natürliche Zahlen. Mit $\text{size}(n)$ werde die Länge der Binärdarstellung der Zahl n (ohne führende binäre Nullen) bezeichnet. Dabei sei $\text{size}(0) = 1$. Es gelte $|n| \geq |m|$, $k = \text{size}(n) \geq \text{size}(m)$, $k \in O(\log(n))$. Dann gibt es Turingmaschinen, die die Funktionen $\text{SUM}(n, m)$, $\text{DIF}(n, m)$, $\text{MULT}(n, m)$ und $\text{DIV}(n, m)$ berechnen und folgende Zeitkomplexitäten besitzen:

- (i) Die Addition und Differenzenbildung der Zahlen n und m (Berechnung von $\text{SUM}(n, m)$ und $\text{DIF}(n, m)$) erfordert jeweils $O(k)$, also $O(\log(n))$ viele Bitoperationen.
- (ii) Die Multiplikation der Zahlen n und m (Berechnung von $\text{MULT}(n, m)$) kann in $O(k^2)$, also in $O((\log(n))^2)$ vielen Bitoperationen durchgeführt werden.
- (iii) Ist $\text{size}(n) \geq 2 \cdot \text{size}(m)$, dann kann die Berechnung des ganzzahligen Quotienten $\lfloor n/m \rfloor$ (Berechnung von $\text{DIV}(n, m)$) in $O(k^2)$, also in $O((\log(n))^2)$ vielen Bitoperationen durchgeführt werden.

Meist ist man nicht am exakten Wert der Anzahl der Konfigurationsänderungen interessiert, sondern nur an der **Größenordnung der Zeitkomplexität** (in Abhängigkeit von der Größe der Eingabe). Bei der Analyse wird man meist eine obere Schranke $g(n)$ für $T_{TM}(n)$ herleiten, d.h. man wird eine Aussage der Form $T_{TM}(n) \leq g(n)$ begründen. In diesem Fall ist $T_{TM}(n) \in O(g(n))$. Eine zusätzliche Aussage $T_{TM}(n) \leq h(n) \leq g(n)$ mit einer Funktion $h(n)$ führt auf die verbesserte Abschätzung $T_{TM}(n) \in O(h(n))$. Man ist also bei der worst case-Analyse an einer möglichst kleinen oberen Schranke für $T_{TM}(n)$ interessiert.

Es gibt eine Reihe von **Varianten von Turingmaschinen**:

- Die Turingmaschine besitzt Bänder, die nach links und rechts unendlich lang sind.
- Die Turingmaschine besitzt ein einziges nach links und rechts oder nur zu einer Seite unendlich langes Band.
- Die Turingmaschine besitzt mehrdimensionale Bänder bzw. ein mehrdimensionales Band.
- Das Eingabealphabet der Turingmaschine besteht aus zwei Zeichen, etwa $I = \{0, 1\}$.
- Das Arbeitsalphabet der Turingmaschine besteht aus zwei Zeichen, etwa $\Sigma = \{0, 1\}$.
- Die Turingmaschine hat endlich viele Endzustände.

Es zeigt sich, daß alle Definitionen **algorithmisch äquivalent** sind, d.h. eine Turingmaschinenvariante kann eine andere Turingmaschinenvariante simulieren. Allerdings ändert sich da-

bei u.U. die Zeitkomplexität, in der von der jeweiligen Turingmaschinenvariante Sprachen erkannt werden. Beispielsweise kann eine k -DTM TM mit Zeitkomplexität $T_{TM}(n)$ durch eine 1-DTM mit Zeitkomplexität $O(T_{TM}^2(n))$ simuliert werden.

2.2 Random Access Maschinen

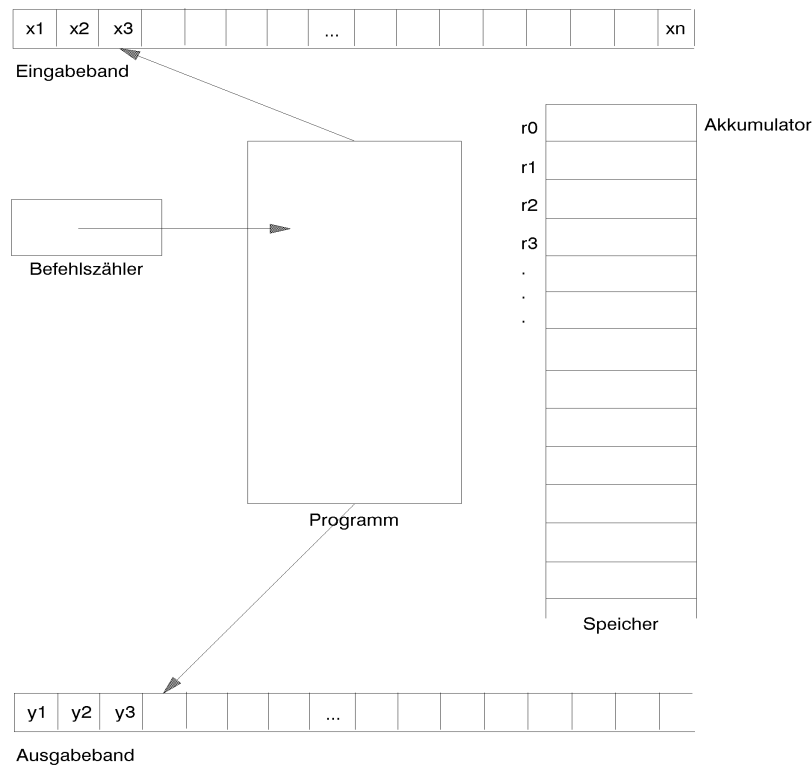
Eine **Random Access Maschine (RAM)** modelliert einen Algorithmus in Form eines sehr einfachen Computers, der ein Programm ausführt, das sich nicht selbst modifizieren kann und fest in einem Programmspeicher geladen ist. Das Konzept der RAM ist zunächst als Modell für die Berechnung partieller Funktionen $f: \mathbb{Z}^n \rightarrow \mathbb{Z}^m$ gedacht, die n -stellige Zahlenfolgen (ganzer Zahlen) auf m -stellige Zahlenfolgen abbilden. Eine Turingmaschine dagegen ist zur Akzeptanz von Sprachen $L \subseteq \Sigma^*$ bzw. zur Berechnung von partieller Funktionen $f: \Sigma^* \rightarrow \Sigma^*$ konzipiert, die Zeichenketten (Wörtern) über einem endlichen Alphabet auf Zeichenketten über eventuell einem anderen Alphabet abbilden. Es wird sich jedoch zeigen, daß beide Modelle äquivalent sind. Beide Modelle sind in der Lage, numerische Funktionen zu berechnen und als Sprachakzeptoren eingesetzt zu werden. Während die „Programmierung“ einer Turingmaschine in der Definition der Überföhrungsfunktion besteht, eine Vorgehensweise, die verglichen mit der gängigen Programmierung eines Computers zumindest gewöhnungsbedürftig ist, entspricht die Programmierung einer RAM eher der Programmierung eines Computers auf Maschinensprachebene.

Eine RAM hat folgende Bestandteile:

- **Eingabeband:** eine Folge von Zellen, die als Eingabe n ganze (positive oder negative) Zahlen $x_1, \dots, x_n$ enthalten und von einem Lesekopf nur gelesen werden können. Nachdem eine Zahl x_i gelesen wurde, rückt der Lesekopf auf die benachbarte Zelle, die die ganze Zahl x_{i+1} enthält.
- **Ausgabeband:** eine Folge von Zellen, über die ein Schreibkopf von links nach rechts wandert, der nacheinander ganze Zahlen y_j schreibt. Ist eine Zahl geschrieben, kann sie nicht mehr verändert werden; der Schreibkopf rückt auf das rechts benachbarte Feld.
- **Speicher:** eine unendliche Folge $r_0, r_1, r_2, \dots$ von **Registern**, die jeweils in der Lage sind, eine beliebig große ganze Zahl aufzunehmen. Das Register r_0 wird als **Akkumulator** bezeichnet. In ihm finden alle arithmetische Operationen statt.
- **Programm:** eine Folge aufsteigend numerierter Anweisungen, die fest in der RAM „verdrahtet“ sind. Zur besseren Lesbarkeit eines RAM-Programms können einzelne Anweisungen auch mit **symbolischen Marken** versehen werden. Jede RAM stellt ein eigenes Programm dar, das natürlich mit unterschiedlichen Eingaben konfrontiert werden kann.

Das Programm kann sich nicht modifizieren. Ein Programm ist aus einfachen Befehlen aufgebaut (siehe unten). Die einzelnen Anweisungen werden nacheinander ausgeführt.

- **Befehlszähler:** enthält die Nummer der nächsten auszuführenden Anweisung.



Um die Bedeutung des **Befehlsvorrats** einer RAM festzulegen, wird die **Speicherabbildungsfunktion** $c: \mathbf{N} \rightarrow \mathbf{Z}$ verwendet. $c(i)$ gibt zu jedem Zeitpunkt des Programmlaufs den Inhalt des Registers r_i an. Zu Beginn des Programmlaufs wird jedes Register mit dem Wert 0 initialisiert, d.h. es ist $c(i) = 0$ für jedes $i \in \mathbf{N}$.

Jeder **Anweisung (Befehl)** ist aus einem **Operationscode** und einem **Operanden** aufgebaut. Ein Operand kann eine der folgenden Formen aufweisen:

1. i
2. $c(i)$, für $i \geq 0$; bei $i < 0$ stoppt die RAM
3. $c(c(i))$, für $i \geq 0$ und $c(i) \geq 0$; bei $i < 0$ oder $c(i) < 0$ stoppt die RAM.

Die RAM-Anweisungen und ihre Bedeutung sind folgender Tabelle zu entnehmen. Dabei steht i für eine natürliche Zahl und m für eine Anweisungsnummer bzw. für eine Anweisungsmarke im Programm der RAM.

RAM-Anweisung	Bedeutung
LOAD i LOAD $c(i)$ für $i \geq 0$ LOAD $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := i$ $c(0) := c(i)$ $c(0) := c(c(i))$
STORE $c(i)$ für $i \geq 0$ STORE $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(i) := c(0)$ $c(c(i)) := c(0)$
ADD i ADD $c(i)$ für $i \geq 0$ ADD $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) + i$ $c(0) := c(0) + c(i)$ $c(0) := c(0) + c(c(i))$
SUB i SUB $c(i)$ für $i \geq 0$ SUB $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) - i$ $c(0) := c(0) - c(i)$ $c(0) := c(0) - c(c(i))$
MULT i MULT $c(i)$ für $i \geq 0$ MULT $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) * i$ $c(0) := c(0) * c(i)$ $c(0) := c(0) * c(c(i))$
DIV i für $i \neq 0$ DIV $c(i)$ für $i \geq 0$ DIV $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := \lfloor c(0) / i \rfloor$ $c(0) := \lfloor c(0) / c(i) \rfloor$, falls $c(i) \neq 0$ ist, sonst stoppt die RAM $c(0) := \lfloor c(0) / c(c(i)) \rfloor$, falls $c(c(i)) \neq 0$ ist, sonst stoppt die RAM
READ $c(i)$ für $i \geq 0$ READ $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(i) :=$ momentaner Eingabewert, und der Lesekopf rückt eine Zelle nach rechts $c(c(i)) :=$ momentaner Eingabewert, und der Lesekopf rückt eine Zelle nach rechts

../..

WRITE i	i wird auf das Ausgabeband gedruckt, und der Schreibkopf rückt um eine Zelle nach rechts
WRITE $c(i)$ für $i \geq 0$	$c(i)$ wird auf das Ausgabeband gedruckt, und der Schreibkopf rückt um eine Zelle nach rechts
WRITE $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(c(i))$ wird auf das Ausgabeband gedruckt, und der Schreibkopf rückt um eine Zelle nach rechts
JUMP m	Der Befehlszähler wird mit der Nummer der Anweisung geladen, die die Marke m trägt (nächste auszuführende Anweisung)
JGTZ m	Der Befehlszähler wird mit der Nummer der Anweisung geladen, die die Marke m trägt (nächste auszuführende Anweisung), falls $c(0) > 0$ ist, sonst wird der Inhalt des Befehlszählers um 1 erhöht
JZERO m	Der Befehlszähler wird mit der Nummer der Anweisung geladen, die die Marke m trägt (nächste auszuführende Anweisung), falls $c(0) = 0$ ist, sonst wird der Inhalt des Befehlszählers um 1 erhöht
HALT	Die RAM stoppt

Falls eine Anweisung nicht definiert ist, z.B. STORE 3 oder STORE $c(c(3))$ mit $c(3) = -10$, dann stoppt die RAM, desgleichen bei einer Division durch 0.

Die RAM RAM definiert bei Beschriftung der ersten n Zellen des Eingabebandes mit $x_1, \dots, x_n$ eine partielle Funktion $f_{RAM} : \mathbf{Z}^n \rightarrow \mathbf{Z}^m$ (die Funktion ist partiell, da RAM nicht bei jeder Eingabe stoppen muß): falls RAM bei Eingabe von $x_1, \dots, x_n$ stoppt, nachdem $y_1, \dots, y_m$ in die m ersten Zellen des Ausgabebands geschrieben wurde, dann ist $f_{RAM}(x_1, \dots, x_n) = (y_1, \dots, y_m)$.

$$\text{Eine RAM zur Berechnung der Funktion } f : \begin{cases} \mathbf{N} & \rightarrow \\ n & \rightarrow \end{cases} \begin{cases} \mathbf{N} \\ 1 & \text{für } n = 0 \\ n^n & \text{für } n \geq 1 \end{cases}$$

Ein entsprechendes RAM-Programm lautet:

RAM-Anweisungsfolge		Bemerkung zur Bedeutung
Zeilennummer	RAM-Befehl	
1	READ $c(1)$	$r_1 := n$, Eingabe
2	LOAD $c(1)$	IF $r_1 \leq 0$ THEN <i>Write</i> (0)
3	JGTZ pos	
4	WRITE 0	
5	JUMP endif	
pos	LOAD $c(1)$	
7	STORE $c(2)$	$r_2 := r_1$
8	LOAD $c(1)$	$r_3 := r_1 - 1$
9	SUB 1	
10	STORE $c(3)$	
while	LOAD $c(3)$	WHILE $r_3 > 0$ DO
12	JGTZ continue	
13	JUMP endwhile	
continue	LOAD $c(2)$	$r_2 := r_2 * r_1$
15	MULT $c(1)$	
16	STORE $c(2)$	
17	LOAD $c(3)$	$r_3 := r_3 - 1$
18	SUB 1	
19	STORE $c(3)$	
20	JUMP while	
endwhile	WRITE $c(2)$	<i>Write</i> (r_2)
endif	HALT	

Eine RAM kann auch als **Akzeptor einer Sprache** interpretiert werden. Die Elemente des endlichen Alphabets Σ werden dazu mit den Zahlen $1, 2, \dots, |\Sigma|$ identifiziert. Das RAM-Programm *RAM* akzeptiert $L \subseteq \Sigma^*$ auf folgende Weise. In die ersten n Zellen des Eingabebandes werden Zahlen geschrieben, die den Buchstaben $x_1, \dots, x_n$ eines Wortes $w \in \Sigma^*$, $w = x_1 \dots x_n$, entsprechen. In die $(n+1)$ -te Zelle kommt als Endmarkierung der Wert 0. *RAM* akzeptiert w , falls alle den Buchstaben von w entsprechenden Zahlen und die Endmarkierung 0 gelesen wurden, von *RAM* eine 1 auf das Ausgabeband geschrieben wurde und *RAM* stoppt, dann ist $w \in L$. Falls $w \notin L$ ist, dann stoppt *RAM* mit einer Ausgabe ungleich 1, oder *RAM* stoppt nicht. Die auf diese Weise akzeptierten Wörter aus Σ^* bilden die Menge $L(\text{RAM})$.

Ein RAM-Programm P zur Akzeptanz von

$$L(P) = \{ w \mid w \in \{1, 2\}^* \text{ und die Anzahl der 1'en ist gleich der Anzahl der 2'en} \}$$

P liest jedes Eingabesymbol nach Register r_1 (hier wird angenommen, daß nur die Zahlen 0, 1 und 2 auf dem Eingabeband stehen) und berechnet in Register r_2 die Differenz d der Anzahlen der bisher gelesenen 1'en und 2'en. Wenn beim Lesen der Endmarkierung 0 diese Differenz gleich 0 ist, wird eine 1 auf das Ausgabeband geschrieben, und P stoppt.

RAM-Anweisungsfolge		Bemerkung zur Bedeutung
Zeilennummer	RAM-Befehl	
	LOAD 0 STORE $c(2)$	$d := 0$
	READ $c(1)$	<i>Read</i> (x)
while	LOAD $c(1)$ JZERO endwhile	WHILE $x \neq 0$ DO
	LOAD $c(1)$ SUB 1 JZERO one	IF $x \neq 1$
	LOAD $c(2)$ SUB 1 STORE $c(2)$	THEN $d := d - 1$
	JUMP endif	
one	LOAD $c(2)$ ADD 1 STORE $c(2)$	ELSE $d := d + 1$
endif	READ $c(1)$ JUMP while	<i>Read</i> (x)
endwhile	LOAD $c(2)$ JZERO output HALT	IF $d = 0$ THEN <i>Write</i> (1)
output	WRITE 1 HALT	

Auch beim RAM-Modell interessieren **Zeit- und Raumkomplexität einer Berechnung**.

Mit $\mathbf{Z}^n$ werde das n -fache kartesische Produkt von $\mathbf{Z}$ bezeichnet, d.h. $\mathbf{Z}^n = \underbrace{\mathbf{Z} \times \dots \times \mathbf{Z}}_{n\text{-mal}}$. Die

Menge $\mathbf{Z}^*$ aller endlichen Folgen ganzer Zahlen sei definiert durch $\mathbf{Z}^* = \bigcup_{n \geq 0} \mathbf{Z}^n$.

Ein RAM-Programm RAM lese die Eingabe $x_1, \dots, x_n$ und durchlaufe bis zum Erreichen der HALT-Anweisung m viele Anweisungen (einschließlich der HALT-Anweisung). Dann wird durch $t_{RAM}(x_1, \dots, x_n) = m$ eine partielle Funktion $t_{RAM} : \mathbf{Z}^* \rightarrow \mathbf{N}$ definiert. Falls RAM bei Eingabe von $x_1, \dots, x_n$ nicht anhält, dann ist $t_{RAM}(x_1, \dots, x_n)$ nicht definiert.

In diesem Modell werden zwei Kostenkriterien unterschieden:

Beim **uniformen Kostenkriterium** benötigt die Ausführung jeder Anweisung eine Zeiteinheit, und jedes Register belegt eine Platzeinheit. Die **Zeitkomplexität** von RAM (**im schlechtesten Fall, worst case**) wird **unter dem uniformen Kostenkriterium** definiert durch die partielle Funktion $T_{RAM} : \mathbf{N} \rightarrow \mathbf{N}$ mit

$$T_{RAM}(n) = \max\{t_{RAM}(w) \mid w \text{ besteht aus bis zu } n \text{ vielen natürlichen Zahlen}\}.$$

Entsprechend wird die **Platzkomplexität** von RAM (**im schlechtesten Fall, worst case**) $S_{RAM}(n)$ **unter dem uniformen Kostenkriterium** als die während der Rechnung benötigte maximale Anzahl an Registern definiert, wenn bis zu n viele Zahlen eingelesen werden.

Ein Register bzw. eine Zelle des Eingabe- und Ausgabebandes kann im RAM-Modell eine beliebig große Zahl aufnehmen. Eine Berechnung mit n „großen Zahlen“ ist unter dem uniformen Kostenkriterium genauso komplex wie eine Berechnung mit n „kleinen Zahlen“. Diese Sichtweise entspricht häufig nicht den praktischen Gegebenheiten. Dies gilt in der Praxis besonders dann, wenn Speicherzellen eine beschränkte Größe aufweisen, so daß für die Handhabung großer Zahlen pro Zahl mehrere Speicherzellen erforderlich sind. Es erscheint daher sinnvoll, die Größenordnung der bei einer Berechnung beteiligten Zahlen bzw. Operanden mit zu berücksichtigen. Dieses führt auf das **logarithmische Kostenkriterium**, das die Größen der bei einer Berechnung beteiligten Zahlen (gemessen in der Anzahl der Stellen zur Darstellung einer Zahl in einem geeigneten Stellenwertsystem) berücksichtigt.

Mit $l(i)$ werde die **Länge** einer ganzen Zahl i bezeichnet:

$$l(i) = \begin{cases} \lfloor \log|i| \rfloor + 1 & \text{für } i \neq 0 \\ 1 & \text{für } i = 0 \end{cases}$$

Die **Kosten eines Operanden** einer RAM-Anweisung faßt folgende Tabelle zusammen.

Operand	Kosten
i	$l(i)$
$c(i)$	$l(i) + l(c(i))$
$c(c(i))$	$l(i) + l(c(i)) + l(c(c(i)))$

Beispielsweise erfordert die Ausführung einer Anweisung ADD $c(c(i))$ folgende Einzelschritte:

1. „Decodieren“ des Operanden i : Kosten $l(i)$
2. „Lesen“ von $c(i)$: Kosten $l(c(i))$
3. „Lesen“ von $c(c(i))$: Kosten $l(c(c(i)))$
4. Ausführung von ADD $c(c(i))$: Kosten $l(c(0)) + l(i) + l(c(i)) + l(c(c(i)))$.

Die folgende Tabelle zeigt die Kosten unter dem logarithmischen Kostenkriterium der Ausführung für jede RAM-Anweisung.

RAM-Anweisung	Bedeutung	Kosten der Ausführung
LOAD i	$c(0) := i$	$l(i)$
LOAD $c(i)$ für $i \geq 0$	$c(0) := c(i)$	$l(i) + l(c(i))$
LOAD $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(c(i))$	$l(i) + l(c(i)) + l(c(c(i)))$
STORE $c(i)$ für $i \geq 0$	$c(i) := c(0)$	$l(c(0)) + l(i)$
STORE $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(c(i)) := c(0)$	$l(c(0)) + l(i) + l(c(i))$
ADD i	$c(0) := c(0) + i$	$l(c(0)) + l(i)$
ADD $c(i)$ für $i \geq 0$	$c(0) := c(0) + c(i)$	$l(c(0)) + l(i) + l(c(i))$
ADD $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) + c(c(i))$	$l(c(0)) + l(i) + l(c(i)) + l(c(c(i)))$

../..

SUB i	$c(0) := c(0) - i$	$l(c(0)) + l(i)$
SUB $c(i)$ für $i \geq 0$	$c(0) := c(0) - c(i)$	$l(c(0)) + l(i) + l(c(i))$
SUB $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) - c(c(i))$	$l(c(0)) + l(i) + l(c(i)) + l(c(c(i)))$
MULT i	$c(0) := c(0) * i$	$l(c(0)) + l(i)$
MULT $c(i)$ für $i \geq 0$	$c(0) := c(0) * c(i)$	$l(c(0)) + l(i) + l(c(i))$
MULT $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := c(0) * c(c(i))$	$l(c(0)) + l(i) + l(c(i)) + l(c(c(i)))$
DIV i für $i \neq 0$	$c(0) := \lfloor c(0) / i \rfloor$	$l(c(0)) + l(i)$
DIV $c(i)$ für $i \geq 0$	$c(0) := \lfloor c(0) / c(i) \rfloor$	$l(c(0)) + l(i) + l(c(i))$
DIV $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(0) := \lfloor c(0) / c(c(i)) \rfloor$	$l(c(0)) + l(i) + l(c(i)) + l(c(c(i)))$
READ $c(i)$ für $i \geq 0$	$c(i) := \text{Eingabewert}$	$l(\text{Eingabewert}) + l(i)$
READ $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	$c(c(i)) := \text{Eingabewert}$	$l(\text{Eingabewert}) + l(c(i))$
WRITE i	Ausgabe von i	$l(i)$
WRITE $c(i)$ für $i \geq 0$	Ausgabe von $c(i)$	$l(c(i))$
WRITE $c(c(i))$ für $i \geq 0$ und $c(i) \geq 0$	Ausgabe von $c(c(i))$	$l(c(c(i)))$
JUMP m	Befehlszähler := m	1
JGTZ m	Befehlszähler := m , falls $c(0) > 0$ ist	$l(c(0))$
JZERO m	Befehlszähler := m , falls $c(0) = 0$ ist	$l(c(0))$
HALT	Die RAM stoppt	1

Die **Zeitkosten eines RAM-Programms** bei gegebenen Eingabewerten **unter dem logarithmischen Kostenkriterium** ist gleich der Summe der Kosten der abzuarbeitenden Anweisungen. Entsprechend sind die **Platzkosten eines RAM-Programms** bei gegebenen Eingabewerten **unter dem logarithmischen Kostenkriterium** gleich dem Produkt der während der Rechnung benötigte maximalen Anzahl an Registern und der Länge der längsten während der Abarbeitung abgespeicherten Zahl.

Kosten des RAM-Programms zur Berechnung der Funktion

$$f : \begin{cases} \mathbb{N} & \rightarrow & \mathbb{N} \\ n & \rightarrow & \begin{cases} 1 & \text{für } n = 0 \\ n^n & \text{für } n \geq 1 \end{cases} \end{cases}$$

	uniformes Kostenkriterium	Logarithmisches Kostenkriterium
Zeitkomplexität	$O(n)$	$O(n^2 \cdot \log(n))$
Platzkomplexität	$O(1)$	$O(\log(n))$

Kosten des RAM-Programms P zur Akzeptanz von

$$L(P) = \left\{ w \mid w \in \{1, 2\}^* \text{ und die Anzahl der 1'en ist gleich der Anzahl der 2'en} \right\}$$

	uniformes Kostenkriterium	Logarithmisches Kostenkriterium
Zeitkomplexität	$O(n)$	$O(n \cdot \log(n))$
Platzkomplexität	$O(1)$	$O(\log(n))$

Eine RAM kann eine deterministische Turingmaschine mit k Bändern (k -DTM) TM simulieren. Dazu wird jede Zelle von TM durch ein Register von RAM nachgebildet. Genauer: der Inhalt der i -ten Zelle des j -ten Bandes von TM kann in Register r_{ik+j+c} gespeichert werden. Hierbei wird angenommen, daß die Symbole des Alphabets Σ von TM für RAM mit den Zahlen $1, \dots, |\Sigma|$ identifiziert werden, so daß man davon sprechen kann, daß „ein Symbol aus Σ in einem Register oder einer Zelle des Eingabe- oder Ausgabebandes von RAM gespeichert werden kann“. Der Wert $c \geq k$ ist eine Konstante, durch die die Register $r_1, \dots, r_c$ als zusätzliche Arbeitsregister für RAM reserviert werden. Unter diesen Arbeitsregistern werden k Register, etwa $r_1, \dots, r_k$, dazu verwendet, die jeweilige Kopfposition von TM während der Simulation festzuhalten (r_j enthält die Kopfposition des j -ten Bandes). Um den Inhalt einer Zelle von TM in der Simulation zu lesen, wird indirekte Adressierung eingesetzt. Beispielsweise kann der Inhalt der simulierten Zelle, über der sich gerade der Schreib/Lesekopf des j -ten Bandes befindet, in den Akkumulator durch die RAM-Anweisung $LOAD\ c(c(j))$ geladen werden.

Satz 2.2-1:

Hat die k -DTM TM die Zeitkomplexität $T_{TM}(n) \geq n$, dann gibt es eine RAM RAM , die die Eingabe von TM in die Register liest, die das erste Band nachbilden, und anschließend die $T_{TM}(n)$ Schritte von TM simuliert. Dieser Vorgang verläuft unter uniformem Kostenkriterium in $O(T_{TM}(n))$ vielen Schritten, unter logarithmischem Kostenkriterium in $O(T_{TM}(n) \cdot \log(T_{TM}(n)))$ vielen Schritten.

Umgekehrt gilt:

Es sei L eine Sprache, die von einem RAM-Programm der Zeitkomplexität $T_{RAM}(n)$ unter dem logarithmischen Kostenkriterium akzeptiert wird. Falls das RAM-Programm keine Multiplikationen oder Divisionen verwendet, dann wird L von einer k -DTM (für ein geeignetes k) mit Zeitkomplexität $O(T_{RAM}^2(n))$ akzeptiert.

Falls das RAM-Programm Multiplikationen oder Divisionen einsetzt, werden diese Operationen jeweils durch Unterprogramme simuliert, in denen als arithmetische Operationen nur Additionen und Subtraktionen verwendet werden. Es läßt sich zeigen, daß diese Unterprogramme so entworfen werden können, daß die logarithmischen Kosten zur Ausführung des jeweiligen Unterprogramms höchstens das Quadrat der logarithmischen Kosten der Anweisung betragen, die es nachbildet.

Die folgende Übersicht zeigt (in Ergänzung mit den Aussagen aus Kapitel 2.1) die Anzahl der Schritte (Zeitkomplexität), die Maschinenvariante A benötigt, um die Ausführung einer Rechnung zu simulieren, der auf Maschinenvariante B bei einem akzeptierten Eingabewort w der Länge n eine Schrittzahl (Anzahl der Überführungen bis zum Akzeptieren) von $T(n)$ benötigt. Dabei wird für das RAM-Modell das logarithmische Kostenkriterium angenommen.

Simulierte Maschine B	simulierende Maschine A		
	1-DTM	k -DTM	RAM
1-DTM	-	$O(T(n))$	$O(T(n) \log(T(n)))$
k -DTM	$O(T^2(n))$	-	$O(T(n) \log(T(n)))$
RAM	$O(T^3(n))$	$O(T^2(n))$	-

2.3 Programmiersprachen

In den meisten Anwendungen werden Algorithmen mit Hilfe der gängigen Programmiersprachen formuliert (siehe Kapitel 1.3). Es zeigt sich (siehe angegebene Literatur), daß man sehr einfache Programmiersprachen definieren kann, so daß man mit in diesen Sprachen formulierten Algorithmen Turingmaschinen simulieren kann. Umgekehrt kann man mit Turingmaschinen das Verhalten dieser Algorithmen nachbilden. Da eine Turingmaschine einen unendlich großen Speicher besitzt, muß man für Algorithmen in diesen Programmiersprachen unendlich viele Variablen zulassen bzw. voraussetzen, daß die Algorithmen auf Rechnern ablaufen, die einen unendlich großen Speicher besitzen. Dieses Prinzip wurde ja auch im RAM-Modell verwirklicht.

Eine Programmiersprache, deren Programme die Turingberechenbarkeit nachzubilden in der Lage sind, benötigt die Definition von:

- abzählbar vielen Variablen, die Werte aus $\mathbf{N}$ annehmen können (ganzzahlige und rationale Werte werden durch natürlichzahlige Werte nachgebildet, vgl. die Festpunktdarstellung von Zahlen, reelle Werte durch natürlichzahlige Werte approximiert, vgl. die Gleitpunktdarstellung von Zahlen)
- elementaren Anweisungen wie Wertzuweisungen an Variablen, einfachen arithmetischen Operationen (Addition, Subtraktion, Multiplikation, ganzzahlige Division) zwischen Variablen und Konstanten
- zusammengesetzten Anweisungen (Sequenz *anweisung₁; anweisung₂*), Blockbildung (**BEGIN** *anweisung₁*; ...; *anweisung_n* **END**), bedingte Anweisungen (**IF** *bedingung* **THEN** *anweisung₁* **ELSE** *anweisung₂*; hierbei ist *bedingung* ein Boolescher Ausdruck, der ein logisches Prädikat mit Variablen darstellt), Wiederholungsanweisungen (**WHILE** *bedingung* **DO** *anweisung*;))
- einfachen Ein/Ausgabeansweisungen (**READ** (*x*), **WRITE** (*x*)).

Auf eine formale Beschreibung dieses Ansatzes soll hier verzichtet werden.

Auch bei der Formulierung eines Algorithmus mit Hilfe einer Programmiersprache kann man nach der Zeit- und Raumkomplexität fragen. Dabei muß man wieder zwischen uniformen und logarithmischen Kostenkriterium unterscheiden.

Ein Programm *PROG* habe die Eingabe $x = [x_1, \dots, x_n]$, die sich aus n Parametern (Formalparameter bei der Spezifikation des Programms, Aktualparameter bei Aufruf des Programms) zusammensetzt. Beispielsweise kann x ein Graph sein, der n Knoten besitzt. *PROG* berechne eine Ausgabe $y = [y_1, \dots, y_m]$, die sich aus m Teilen zusammensetzt. Beispielsweise kann y

das Ergebnis der Berechnung einer Funktion f sein, d.h. $[y_1, \dots, y_m] = f(x_1, \dots, x_n)$. Oder $PROG$ soll entscheiden, ob die Eingabe x eine spezifizierte Eigenschaft besitzt; dann ist $y \in \{\text{TRUE}, \text{FALSE}\}$. In der Regel wird die Zeit- und Raumkomplexität von $PROG$ in Abhängigkeit der Eingabe $x = [x_1, \dots, x_n]$ gemessen.

Bei der Berechnung der **uniformen Zeitkomplexität** von $PROG$ bei Eingabe von x wird die Anzahl der durchlaufenen Anweisungen gezählt. Bei der Berechnung der **uniformen Raumkomplexität** von $PROG$ bei Eingabe von x wird die Anzahl der benötigten Variablen gezählt.

Der Ansatz bietet sich immer dann an, wenn die so ermittelten Werte der Zeit- und Raumkomplexität nur von der Anzahl n der Komponenten der Eingabe $x = [x_1, \dots, x_n]$ abhängen. Diese Situation liegt beispielsweise vor, wenn $PROG$ die Aufgabe hat, die Komponenten der Eingabe (nach einem „einfachen“ Kriterium) zu sortieren.

Das folgende Beispiel zeigt jedoch, daß der Ansatz nicht immer adäquat ist. Die Pascal-Prozedur `zweier_potenz` berechnet bei Eingabe einer Zahl $n > 0$ den Wert $c = 2^{2^n} - 1$.

```

PROCEDURE zweier_potenz (    n : INTEGER;
                           VAR c : INTEGER);

VAR idx : INTEGER;
    p   : INTEGER;

BEGIN {zweier-potenz }
    idx := n;           { Anweisung 1 }
    p   := 2;           { Anweisung 2 }
    WHILE idx > 0 DO     { Anweisung 3 }
        BEGIN
            p := p*p;     { Anweisung 4 }
            idx := idx - 1; { Anweisung 5 }
        END;
    c := p - 1;          { Anweisung 6 }
END   { zweier-potenz };

```

Werden nur die Anzahl der ausgeführten Anweisungen gezählt, so ergibt sich bei Eingabe der Zahl $n > 0$ eine Zeitkomplexität der Ordnung $O(n)$ und eine Raumkomplexität der Ordnung $O(1)$. Das Ergebnis $c = 2^{2^n} - 1$ belegt jedoch 2^n viele binäre Einsen und benötigt zu seiner Erzeugung mindestens 2^n viele (elementare) Schritte. Vergrößert man die Eingabe n um eine Konstante k , d.h. betrachtet man die Eingabe $n + k$, so bleibt die Laufzeit in der Ordnung $O(n)$, während das Ergebnis um den Faktor 2^k größer wird.

Es bietet sich daher eine etwas sorgfältigere Definition der Zeit- und Raumkomplexität an. Diese wird hier aufgezeigt, wenn die Ein- und Ausgabe eines Programms aus numerischen Daten bzw. aus Daten besteht, die als numerische Daten dargestellt werden können (beispielsweise die Datentypen **BOOLEAN** oder **SET OF ...** in Pascal).

Wenn die Zeit- und Raumkomplexität nicht nur von der Anzahl n der Komponenten der Eingabe $x = [x_1, \dots, x_n]$ abhängt, sondern wie im Beispiel der Prozedur `zweier_potenz` von den Zahlenwerten $x_1, \dots, x_n$ selbst, wird folgender Ansatz gewählt.

Für eine ganze Zahl z sei die **Größe** $size(z) = |bin(z)|$ die Anzahl signifikanter Bits, die benötigt werden, um z im Binärsystem darzustellen. Für negative Werte von z wird die Komplementdarstellung gewählt. Dabei sei $size(0) = 1$. Es gilt also $size(z) = \begin{cases} \lfloor \log_2(|z|) \rfloor + 1 & \text{für } z \neq 0 \\ 1 & \text{für } z = 0 \end{cases}$ und $size(z) \in O(\log(|z|))$. Für $x = [x_1, \dots, x_n]$ sei $size(x) = n \cdot \max\{size(x_1), \dots, size(x_n)\}$.

Der Wert $size(x)$ repräsentiert damit die Anzahl der Bits, um die Zahlenwerte darzustellen, die in x vorkommen.

Die **logarithmische Zeitkomplexität** eines Programms *PROG* bei Eingabe von x ist die Anzahl der durchlaufenen Anweisungen, gemessen in Abhängigkeit von der so definierten Größe $size(x)$. Bei der Berechnung der **logarithmischen Raumkomplexität** von *PROG* bei Eingabe von x wird die Anzahl der Bits gezählt, um alle von *PROG* benötigten Variablen abzuspeichern; dieser Wert wird in Abhängigkeit von $size(x)$ genommen.

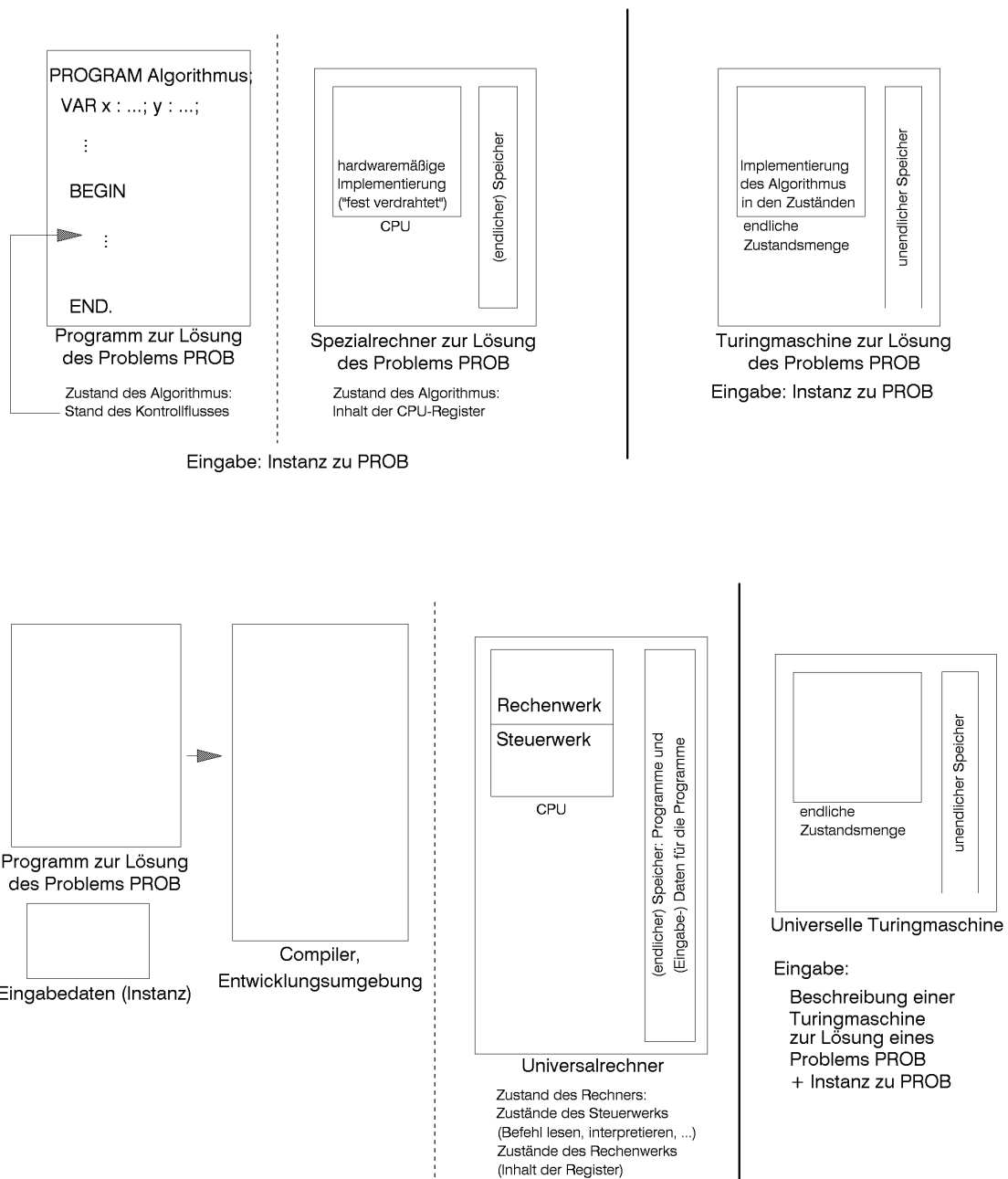
In obiger Pascal-Prozedur `zweier_potenz` gilt für die Eingabe n die Beziehung $k = size(n) = c \cdot \log_2(n)$; die Prozedur hat eine Laufzeit der Ordnung

$O(n) = O(2^{\log_2(n)}) = O(2^{1/c} \cdot 2^{c \cdot \log_2(n)}) = O(2^k)$, also exponentielle Laufzeit, gemessen in der Größe der Eingabe. Seine Raumkomplexität wird durch die Anzahl der Bits bestimmt, um die Variable p zu speichern. Diese ist von der Ordnung $O(2^n) = O(2^{2^k})$. Insgesamt gibt das logarithmische Kostenkriterium die Realität exakt wieder.

2.4 Universelle Turingmaschinen

Jede Turingmaschine und jeder in einer Programmiersprache formulierte Algorithmus ist zur Lösung eines spezifischen Problems entworfen. Dabei ist natürlich die Eingabe wechselnder Werte der Problemparameter möglich. Entsprechend könnte man in dieser Situation den Algorithmus hardwaremäßig implementieren und hätte dadurch einen „Spezialrechner“, der in der Lage ist, das spezifische Problem, für das er entworfen ist, zu lösen. Ein Computer wird

jedoch üblicherweise anders eingesetzt: ein Algorithmus wird in Form eines in einer Programmiersprache formulierten Quelltextes einem Compiler (etwa innerhalb einer Entwicklungsumgebung) vorgelegt. Das Programm wird im Computer in Maschinensprache übersetzt, so daß er in der Lage ist, das Programm „zu interpretieren“. Anschließend wird es von diesem Computer mit entsprechenden eingegebenen Problemparametern ausgeführt. Der Computer ist also in der Lage, nicht nur einen speziellen Algorithmus zur Lösung eines spezifischen Problems, sondern jede Art von Algorithmen auszuführen, solange sie syntaktisch korrekt formuliert sind.



Diese Art der Universalität ist auch im Turingmaschinen-Modell möglich. Im folgenden wird dazu eine **universelle Turingmaschine** UTM definiert. Diese erhält als Eingabe die Beschreibung einer Turingmaschine TM (ein Problemlösungsalgorithmus) und einen Eingabedatensatz w für TM . Die universelle Turingmaschine verhält sich dann genauso, wie sich TM bei Eingabe von w verhalten würde.

Im folgenden wird zunächst gezeigt, wie man die Beschreibung einer Turingmaschine aus ihrer formalen Definition erhält.

Eine deterministische k -Band-Turingmaschine $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ kann als endliches Wort über $\{0, 1\}$ kodiert werden. Dazu wird angegeben, wie eine mögliche Kodierung von TM aussehen kann (die hier vorgestellte Kodierung ist natürlich nur eine von vielen möglichen):

Es sei $\#$ ein neues Zeichen, das in $Q \cup \Sigma$ nicht vorkommt. Das Zeichen $\#$ dient als syntaktischen Begrenzungssymbol innerhalb der Kodierung.

Die Zustandsmenge von TM sei $Q = \{q_0, \dots, q_{|Q|-1}\}$, ihr Arbeitsalphabet $\Sigma = \{a_0, \dots, a_{|\Sigma|-1}\}$. Man kann annehmen, daß der Startzustand q_0 und der akzeptierende Zustand $q_{accept} = q_{|Q|-1}$ ist. Das Blankzeichen b kann mit a_0 identifiziert werden. Die Anzahl der Zustände, die Anzahl der Elemente in Σ und die Anzahl der Bänder werden in einem Wort $w_0 \in \{0, 1, \#\}^*$ festgehalten:
 $w_0 = \# \text{bin}(|Q|) \# \text{bin}(|\Sigma|) \# \text{bin}(k) \# \#$.

Hierbei bezeichnet wieder $\text{bin}(n)$ die Binärdarstellung von n .

Die Überföhrungsfunktion d habe d viele Zeilen. Die t -te Zeile laute:

$$d(q_i, a_{i_1}, \dots, a_{i_k}) = (q_j, (a_{j_1}, d_1), \dots, (a_{j_k}, d_k)).$$

Es sind $a_{i_l} \in \Sigma$, $a_{j_m} \in \Sigma$ und $d_i \in \{L, R, S\}$. Diese Zeile wird zunächst als Zeichenkette

$$w_t = (\text{bin}(i) \# \text{bin}(i_1) \# \dots \# \text{bin}(i_k)) (\text{bin}(j) \# (\text{bin}(j_1) \# d_1) \# \dots \# (\text{bin}(j_k) \# d_k)) \#$$

kodiert. Diese Zeichenkette ist ein Wort über dem Alphabet $\{0, 1, (,), \#, L, R, S\}$. Zu beachten ist, daß w_t sich öfhnende und schließende Klammern enthält; der Ausdruck $\text{bin}(i)$ enthält keine Klammern, sondern ist eine Zeichenkette über $\{0, 1\}$.

Die gesamte Turingmaschine TM läßt sich durch Konkatination w_{TM} der Wörter $w_0, w_1, \dots, w_d$ kodieren: $w_{TM} = w_0 w_1 \dots w_d$. Es ist $w_{TM} \in \{0, 1, (,), \#, L, R, S\}^*$. Die 8 Buchstaben dieses Alphabets werden buchstabenweise gemäß der Tabelle

Zeichen	Umsetzung	Zeichen	Umsetzung
0	000	#	100
1	001	L	101
(	010	R	110
)	011	S	111

umgesetzt. Das Resultat der Umsetzung von w_{TM} in eine Zeichenkette über $\{0, 1\}$ wird mit $code(TM)$ bezeichnet.

Beispiel: Kodierung einer Turingmaschine

Gegeben sei die 1-DTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ mit $Q = \{q_0, q_1, q_2\}$, $\Sigma = \{0, 1, b\}$, $I = \{0, 1\}$, $q_{accept} = q_2$ und d gemäß der folgenden Tabelle:

Überföhrungsfunktion $d(q_i, a_i, \dots, a_k)$ $= (q_j, (a_j, d_1), \dots, (a_k, d_k))$	Zwischenkodierung $w_t = (bin(i) \# bin(i_1) \# \dots \# bin(i_k))$ $(bin(j) \# (bin(j_1) \# d_1) \# \dots \# (bin(j_k) \# d_k))$; entsprechend dem Wort über $\{0, 1\}$
$d(q_0, 1) = (q_1, (0, R))$	$w_1 = (0 \# 1)(1 \# (0 \# R)) \#$ 0100001000010111010001100010000100110011011100
$d(q_1, 0) = (q_0, (1, R))$	$w_2 = (1 \# 0)(0 \# (1 \# R)) \#$ 010001100000011010000100010001100110011011100
$d(q_1, 1) = (q_2, (0, R))$	$w_3 = (1 \# 1)(10 \# (0 \# R)) \#$ 0100011000010111010001000100010000100110011011100
$d(q_1, b) = (q_1, (1, L))$	$w_4 = (0 \# 11)(1 \# (1 \# L)) \#$ 010001100001001011010001100010001100101011011100

Das Wort w_0 lautet: $w_0 = \#11\#11\#\#$. Dieses Wort wird übersetzt in eine 0-1-Kodierung und ergibt 1000010011000010011001100100.

Die Turingmaschine TM ist also kodiert durch $code(TM) =$

100001001100001001100110010001000010000101101000110001000010011001101110
001000110000001101000010001000110011001101110001000110000101101000100010
0010000100110011011100010001100001001011010001100010001100101011011100.

Die Turingmaschine TM_{min} mit der kürzesten Kodierung ist die k-DTM $TM_{min} = (Q, \Sigma, I, d, b, q_0, q_{accept})$ mit $k = 0$, $Q = \{q_0\}$, $\Sigma = \{b\}$ und $d = \emptyset$. Das dieser Turingmaschine entsprechende Wort w_0 lautet: $w_0 = \#1\#1\#0\#\#$, übersetzt in eine 0-1-Kodierung

100001100001100000100100. Da $d = \emptyset$ ist, stimmt $code(TM_{min})$ mit dieser 0-1-Folge bereits überein.

Zu jeder Turingmaschine TM gibt es also ein Wort $w \in \{0,1\}^*$, $w = code(TM)$, das TM kodiert. Offensichtlich ist $\{code(TM) \mid TM \text{ ist eine Turingmaschine}\} \subseteq \{0,1\}^*$.

Für unterschiedliche Turingmaschinen TM_1 und TM_2 gilt dabei $code(TM_1) \neq code(TM_2)$, d.h. die so definierte Abbildung $code$ ist injektiv.

Andererseits kann man einem Wort $w \in \{0,1\}^*$ „ansehen“, ob es eine Turingmaschine kodiert. Nicht jedes Wort $w \in \{0,1\}^*$ kodiert eine Turingmaschine, und natürlich kann es vorkommen, daß die durch ein Wort $w \in \{0,1\}^*$ kodierte Turingmaschine wenig Sinnvolles leistet. Die oben informell beschriebenen Regeln der Kodierung einer Turingmaschine TM durch ein Wort $code(TM)$ erlauben ein (deterministisches) algorithmisches Verfahren der syntaktischen Analyse, das feststellt, ob ein Wort $w \in \{0,1\}^*$ die Kodierung einer Turingmaschine darstellt. Dieses Verfahren, auf dessen detaillierte Darstellung hier verzichtet werden soll, werde mit $VERIFIZIERE_TM$ bezeichnet.

Für ein Wort $w \in \{0,1\}^*$ ist

$$VERIFIZIERE_TM(w) = \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Turingmaschine darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Turingmaschine darstellt} \end{cases}$$

Falls ein Wort $w \in \{0,1\}^*$ eine Turingmaschine TM kodiert, d.h. $w = code(TM)$, dann bezeichnet man mit $TM = \overset{-1}{code}(w)$ **die durch w kodierte Turingmaschine** (diese Notation ist zulässig, da $code$ injektiv ist). Wir schreiben dafür kürzer $TM = K_w$.

Mit Hilfe der Kodierungen von Turingmaschinen kann man eine **lineare Ordnung auf der Menge der Turingmaschinen** definieren. Es sei $w \in \{0,1\}^*$. Für $i \in \mathbb{N}_{>0}$ heißt **w Kodierung der i -ten Turingmaschine**, falls gilt:

- (i) $VERIFIZIERE_TM(w) = \text{TRUE}$
- (ii) die Menge $\left\{ x \mid \begin{array}{l} x \in \{0,1\}^* \text{ und } x \text{ liegt in der lexikographischen Ordnung vor } w \\ \text{und } VERIFIZIERE_TM(x) = \text{TRUE} \end{array} \right\}$ enthält genau $i-1$ viele Elemente.

Falls $w = \text{code}(TM)$ die Kodierung der i -ten Turingmaschine ist, dann heit $TM = K_w$ die **i -te Turingmaschine**.

Der folgende Algorithmus ermittelt bei Eingabe einer Zahl $i \in \mathbf{N}_{>0}$ die Kodierung der i -ten Turingmaschine.

Eingabe: Eine natrliche Zahl $i \in \mathbf{N}_{>0}$

Verfahren: Aufruf der Funktion `Generiere_TM (i)`

Ausgabe: $w \in \{0,1\}^*$, w ist die Kodierung der i -ten Turingmaschine.

Die Funktion `Generiere_TM` wird in Pseudocode beschrieben.

```

FUNCTION Generiere_TM (i : INTEGER) : STRING;

VAR x : STRING;
    y : STRING;
    k : INTEGER;

BEGIN { Generiere_TM }
  x := '0'; {  $x \in \{0,1\}^*$  }
  k := 0;

  WHILE k < i DO
    BEGIN
      IF VERIFIZIERE_TM(x)
      THEN BEGIN
        k := k + 1;
        y := x;
      END;
      x := Nachfolger von x in der lexikographischen Reihenfolge von  $\{0,1\}^*$ ;
    END;

  Generiere_TM := y;
END { Generiere_TM }

```

Der folgende Algorithmus ermittelt bei Eingabe von $w \in \{0,1\}^*$ die Nummer i in der oben definierten Ordnung auf der Menge der Turingmaschinen, falls w berhaupt eine Turingmaschine kodiert; ansonsten gibt er den Wert 0 aus.

Eingabe: $w \in \{0,1\}^*$

Verfahren: Aufruf der Funktion `Ordnung_TM (w)`

Ausgabe: $i > 0$, falls w die Kodierung der i -ten Turingmaschine ist,
 $i = 0$, sonst.

Die Funktion `Ordnung_TM` wird in Pseudocode beschrieben.

```

FUNCTION Ordnung_TM (w : STRING) : INTEGER;

VAR x : STRING;
    k : INTEGER;

BEGIN { Ordnung_TM }
  IF NOT VERIFIZIERE_TM(w)
  THEN BEGIN
    Ordnung_TM := 0;
    Exit;
  END;

  x := '0'; {  $x \in \{0,1\}^*$  }
  k := 1;

  WHILE NOT (x = w) DO
    BEGIN
      IF VERIFIZIERE_TM(x)
      THEN k := k + 1;
      x := Nachfolger von x in der lexikographischen Reihenfolge von  $\{0,1\}^*$ ;
    END;

  Ordnung_TM := k;
END { Ordnung_TM }

```

Ein Wort $w \in \{0,1\}^*$ kann damit sowohl als die Kodierung der i -ten Turingmaschine als auch als ein Eingabewort für eine gegebene Turingmaschine oder auch als Binärzahl interpretiert werden. Die Nummer i kann aus w algorithmisch deterministisch ermittelt werden (falls w überhaupt eine Turingmaschine kodiert); ebenso kann die Kodierung der i -ten Turingmaschine erzeugt werden.

Satz 2.4-1:

Es gibt eine **universelle Turingmaschine** UTM mit Eingaben der Form $z = u\#w$ mit $u \in \{0,1\}^*$ und $w \in \{0,1\}^*$, so daß $z \in L(UTM)$ genau dann gilt, wenn $w = code(TM)$ für eine Turingmaschine TM ist, d.h. $TM = K_w$, und $u \in L(K_w)$ ist.

Die universelle Turingmaschine UTM simuliert das Verhalten von K_w auf u : Sie betrachtet w als die Turingmaschine K_w und verhält sich dann auf u so, wie sich K_w auf u verhält.

Beweis:

Zunächst prüft UTM , ob die Eingabe $z \in \{0,1,\#\}^*$ genau ein Zeichen $\#$ enthält. Ist dieses nicht der Fall, so wird z nicht akzeptiert. Hat z die Form $z = u\#w$ mit $u \in \{0,1\}^*$ und $w \in \{0,1\}^*$, so überprüft UTM , ob w eine Turingmaschine kodiert (siehe Funktion $VERIFIZIERE_TM(w)$). Falls dieses nicht zutrifft, wird z nicht akzeptiert.

Andernfalls hat w die Form $w = 100v_1100v_2100v_3100100v$, wobei sich v_1 , v_2 und v_3 ausschließlich aus der Konkatination von Zeichenfolgen 000 und 001 zusammensetzen und $v \in \{0,1\}^*$ ist. Nimmt man von v_1 jedes dritte Zeichen, so erhält man eine Binärzahl, die die Anzahl der Zustände von K_w angibt. Der Zustand mit der höchsten Nummer ist der akzeptierende Zustand, der Zustand mit der Nummer 0 der Anfangszustand. Entsprechend erhält man aus v_2 die Anzahl der Zeichen des Arbeitsalphabets und damit implizit das Arbeitsalphabet Σ und aus v_3 die Anzahl k der Bänder von K_w . In v ist die Überföhrungsfunktion von K_w kodiert.

UTM erzeugt jetzt auf einem weiteren Band, das als Konfigurations-Simulationsband bezeichnet werden soll, die Kodierung der Anfangskonfiguration von K_w mit Eingabewort u ,

d.h. die Kodierung von $K_0 = \left(q_0, \underbrace{(u,1), (e,1), \dots, (e,1)}_k \right)$. Hierbei kann ein ähnlicher Umsetzungsmechanismus wie zur Erzeugung der Kodierung einer Turingmaschine verwendet werden.

Eine Konfiguration von K_w der Form $K = (q_t, (a_1, i_1), \dots, (a_k, i_k))$, wobei q_t der t -te Zustand von K_w und $a_j \in \Sigma^*$, $i_j \geq 1$ für $j = 1, \dots, k$ ist, wird zunächst (gedanklich) umgesetzt in $(bin(t)\#(b_1\#bin(i_1))\# \dots \#(b_k\#bin(i_k)))$. Hierbei erhält man b_j aus a_j (für $j = 1, \dots, k$), indem man jeden Buchstaben durch den Binärwert seiner Nummer in Σ , gefolgt vom Zeichen $\#$ ersetzt. Ist beispielsweise $a_j = aacca$ und sind a bzw. c die Zeichen in Σ mit den Nummern 1 bzw. 3, so ist $b_j = 1\#1\#1\#1\#1\#$. Die Zeichenfolge $(bin(t)\#(b_1\#bin(i_1))\# \dots \#(b_k\#bin(i_k)))$ wird mittels anfangs angegebener Tabelle in eine Folge über $\{0,1\}$ umgesetzt.

Ist K_w etwa eine 3-DTM mit $\Sigma = \{b, 0, 1, a\}$, wobei b das Blankzeichen (mit Nummer 0) bezeichnet und die Zeichen 0, 1 und a die Nummern 1 bis 3 tragen, so lautet die Anfangskonfiguration K_0 von K_w mit dem Eingabewort $u = 11001$:

$K_0 = (q_0, (11001, 1), (e, 1), (e, 1), (e, 1))$ bzw.

$(0\#(10\#10\#1\#1\#10\#\#1)\#(0\#\#1)\#(0\#\#1))$ und in eine 0-1-Folge kodiert:

010000100010001000100001000100001100001100001000100100001011100010000100100
001011100010000100100001011011.

Im Endstück v des Wortes $w = 100v_1100v_2100v_3100100v$ ist die Überföhrungsfunktion von K_w kodiert. UTM kann mit Hilfe der dort enthaltenen Informationen das Verhalten von K_w simulieren, indem UTM die Einträge auf dem Konfigurations-Simulationsband entsprechend der aus v gelesenen Überföhrungsfunktion schrittweise ändert. Auf dem Anfangsstück $010m100$ des Konfigurations-Simulationsbands (hierbei besteht m ausschließlich aus Teilzeichenketten der Form 000 und 001) steht dabei jeweils die Kodierung des aktuellen Zustands von K_w . Endet die Simulation mit einem Wert m , der dem akzeptierenden Zustand von K_w entspricht, akzeptiert UTM die Eingabe z . ///

Die universelle Turingmaschine UTM kann so konstruiert werden, daß sie mit einem Band auskommt, da sich jede Turingmaschinen mit mehreren Bändern durch eine einbändige Turingmaschine simulieren läßt.

Das Konzept einer universellen Turingmaschine wird u.a. dazu verwendet, die Grenzen der Berechenbarkeit aufzuzeigen (Kapitel 3.2).

2.5 Nichtdeterminismus

Der Begriff „deterministisch“ in der Definition einer k -DTM drückt aus, daß die Nachfolgekonfiguration einer Konfiguration K in einer Berechnung der k -DTM eindeutig durch den in K vorkommenden Zustand q , die von den Schreib/Leseköpfen gerade gelesenen Zeichen $a_1, \dots, a_k$, den (eindeutigen) Wert $d(q, a_1, \dots, a_k)$ der Überföhrungsfunktion und die gegenwärtigen Positionen der Schreib/Leseköpfe bestimmt ist. In diesem Kapitel wird das Modell der nichtdeterministischen Turingmaschine eingeföhrt. Eine nichtdeterministische Turingmaschine unterscheidet sich von einer deterministischen Turingmaschine in der Definition der Überföhrungsfunktion d .

Eine **nichtdeterministische k -Band-Turingmaschine (k -NDTM)** TM ist definiert durch

$$TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$$

mit:

1. Q ist eine endliche nichtleere Menge: die **Zustandsmenge**
2. Σ ist eine endliche nichtleere Menge: das **Arbeitsalphabet**
3. $I \subseteq \Sigma$ ist eine endliche nichtleere Menge: das **Eingabealphabet**
4. $b \in \Sigma \setminus I$ ist das **Leerzeichen**
5. $q_0 \in Q$ ist der **Anfangszustand** (oder **Startzustand**)
6. $q_{accept} \in Q$ ist der **akzeptierende Zustand** (**Endzustand**)
7. $d : (Q \setminus \{q_{accept}\}) \times \Sigma^k \rightarrow \mathcal{P}(Q \times (\Sigma \times \{L, R, S\})^k)$ ist eine partielle Funktion, die **Überföhrungsfunktion**; insbesondere ist $d(q, a_1, \dots, a_k)$ für $q = q_{accept}$ nicht definiert; zu beachten ist, daß d für einige weitere Argumente eventuell nicht definiert ist.

Die Überföhrungsfunktion ordnet also jedem Argument $(q, a_1, \dots, a_k)$ nicht einen einzigen (eindeutigen) Wert $(q', (b_1, d_1), \dots, (b_k, d_k))$ zu, sondern eine *endliche Menge von Werten* der Form $(q', (b_1, d_1), \dots, (b_k, d_k))$, von denen in einer Berechnung ein Wert genommen werden kann (natürlich kann diese endliche Menge auch aus einem einzigen Element bestehen).

Damit wird auch das Aussehen einer Berechnung einer nichtdeterministischen Turingmaschine TM neu definiert. Auch hierbei wird wieder der Begriff der **Konfiguration** verwendet, um den gegenwärtigen Gesamtzustand von TM zu beschreiben, d.h. den gegenwärtigen Zustand zusammen mit den Bandinhalten und den Positionen der Schreib/Leseköpfe:

$$K = (q, (a_1, i_1), \dots, (a_k, i_k)) \text{ mit } q \in Q, a_j \in \Sigma^*, i_j \geq 1 \text{ für } j = 1, \dots, k.$$

TM startet wie im deterministischen Fall im Anfangszustand mit einer **Anfangskonfiguration** $K_0 = (q_0, (w, 1), (e, 1), \dots, (e, 1))$. Ist TM in eine Konfiguration $K = (q, (a_1, i_1), \dots, (a_k, i_k))$ gekommen und ist $d(q, a_1, \dots, a_k)$ definiert, dann besteht $d(q, a_1, \dots, a_k)$ aus einer endlichen Menge von Werten der Form $(q', (b_1, d_1), \dots, (b_k, d_k))$, d.h.

$$d(q, a_1, \dots, a_k) = \{ (q_1, (b_{11}, d_{11}), \dots, (b_{1k}, d_{1k})), \dots, (q_t, (b_{t1}, d_{t1}), \dots, (b_{tk}, d_{tk})) \}.$$

Als **Folgekonfiguration** von K wird eine der t möglichen Konfigurationen

$$K_1 = (q_1, (b_{11}, i'_{11}), \dots, (b_{1k}, i'_{1k})), \dots, K_t = (q_t, (b_{t1}, i'_{t1}), \dots, (b_{tk}, i'_{tk}))$$

genommen. K_i für $i = 1, \dots, t$ entsteht aus K dadurch, daß man wie im deterministischen Fall q durch q_i und $a_1, \dots, a_k$ durch $b_{i1}, \dots, b_{ik}$ ersetzt und die Köpfe so bewegt, wie es $(q_i, (b_{i1}, d_{i1}), \dots, (b_{ik}, d_{ik}))$ in $d(q, a_1, \dots, a_k)$ angibt. Welche der t möglichen Folgekonfigurationen auf die Konfiguration K folgt, wird nicht gesagt (wird nichtdeterministisch festgelegt); es wird „eine geeignete Folgekonfiguration“ genommen. Man schreibt dann

$$K \Rightarrow K_i.$$

Wie im deterministischen Fall heißt eine Konfiguration K_{accept} , die den akzeptierenden Zustand q_{accept} enthält, d.h. eine Konfiguration der Form

$$K_{accept} = (q_{accept}, (a_1, i_1), \dots, (a_k, i_k)),$$

akzeptierende Konfiguration (Endkonfiguration).

Wie im deterministischen Fall schreibt man $K \Rightarrow^m K'$ mit $m \in \mathbf{N}$, wenn K' aus K durch m Konfigurationsänderungen hervorgegangen ist, d.h. wenn es m Konfigurationen $K_1, \dots, K_m$ gibt mit

$$K \Rightarrow K_1 \Rightarrow \dots \Rightarrow K_m = K'.$$

Für $m = 0$ ist dabei $K = K'$.

Die von einer k -NDTM TM **akzeptierte Sprache** ist die Menge

$$\begin{aligned} L(TM) &= \{ w \mid w \in I^* \text{ und } w \text{ wird von } TM \text{ akzeptiert} \} \\ &= \{ w \mid w \in I^* \text{ und } (q_0, (w, 1), (e, 1), \dots, (e, 1)) \Rightarrow^* (q_{accept}, (a_1, i_1), \dots, (a_k, i_k)) \}. \end{aligned}$$

Wie im deterministischen Fall kann man die Überföhrungsfunktion d modifizieren, indem man die Zustandsmenge Q um einen neuen Zustand q_{reject} erweitert und d um entsprechende Zeilen ergnzt (siehe Kapitel 2.1), so da alle Berechnungen, die in einem Zustand $q \neq q_{accept}$, f r die $d(q, \dots)$ bisher nicht definiert ist, um eine  berf hrung in den Zustand q_{reject} fortgesetzt werden k nnen. F r q_{reject} ist d nicht definiert.

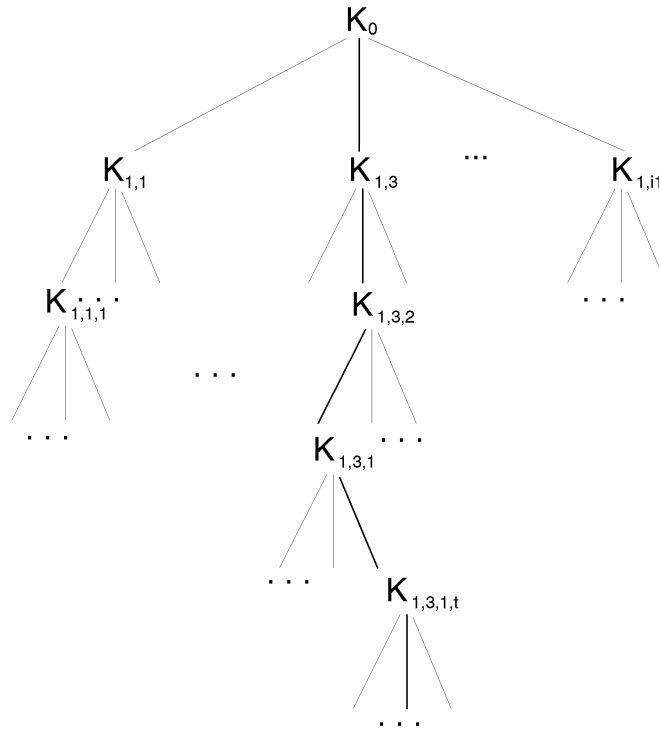
Im *deterministischen* Fall kann man sich den Ablauf einer Berechnung als eine Folge

$$K_0 \Rightarrow \dots \Rightarrow K \Rightarrow K' \Rightarrow \dots$$

vorstellen, wobei jeder Schritt $K \Rightarrow K'$ eindeutig durch die  berf hrungsfunktion d feststeht. Im *nichtdeterministischen* Fall hat man in einem Schritt $K \Rightarrow K'$ eventuell mehrere Alternativen, so da sich eine Berechnung hier als ein *Pfad* durch einen „Berechnungsbaum“ darstellt. Jeder Knoten dieses Berechnungsbaums ist mit einer Konfiguration markiert. Die Wurzel ist mit einer Anfangskonfiguration $K_0 = (q_0, (w, 1), (e, 1), \dots, (e, 1))$ mit einem Wort $w \in I^*$ markiert. Ist die Konfiguration $K = (q, (a_1, i_1), \dots, (a_k, i_k))$ die Markierung eines Knotens und kann hier der Eintrag $d(q, a_1, \dots, a_k)$ der  berf hrungsfunktion angewendet werden, etwa

$d(q, a_1, \dots, a_k) = \{ (q_1, (b_{11}, d_{11}), \dots, (b_{1k}, d_{1k})), \dots, (q_t, (b_{t1}, d_{t1}), \dots, (b_{tk}, d_{tk})) \}$, dann enth lt dieser Knoten des Berechnungsbaums t direkte Nachfolger, die mit den sich ergebenden Nachfolge-

konfigurationen $K_1 = (q_1, (b_{11}, i'_{11}), \dots, (b_{1k}, i'_{1k}))$, ..., $K_t = (q_t, (b_{t1}, i'_{t1}), \dots, (b_{tk}, i'_{tk}))$ markiert sind. Führt einer der Pfade von einer Anfangskonfiguration K_0 zu einer Endkonfiguration K_{accept} , so wird das in der Anfangskonfiguration stehende Wort w akzeptiert.



Das Partitionenproblem mit ganzzahligen Eingabewerten

Instanz: $I = \{a_1, \dots, a_n\}$

I ist eine Menge von n natürlichen Zahlen.

Lösung: Entscheidung „ja“, falls es eine Teilmenge $J \subseteq \{1, \dots, n\}$ mit $\sum_{i \in J} a_i = \sum_{i \notin J} a_i$ gibt (d.h.

wenn sich die Eingabeinstanz in zwei Teile zerlegen lässt, deren jeweilige Summen gleich groß sind).

Entscheidung „nein“ sonst.

Offensichtlich lautet bei einer Instanz $I = \{a_1, \dots, a_n\}$ mit ungeradem $B = \sum_{i=1}^n a_i$ die Entscheidung „nein“. Daher kann B als gerade vorausgesetzt werden.

Zunächst wird eine 3-NDTM TM angegeben, die dieses Problem löst, wenn die Eingaben in „unärer“ Kodierung eingegeben werden: Eine Eingabeinstanz $I = \{a_1, a_2, \dots, a_n\}$ wird dabei als Wort $w = 10^{a_1} 10^{a_2} \dots 10^{a_n}$ kodiert.

Die 3-NDTM $TM = (\{q_0, \dots, q_5\}, \{0, 1, b, \$\}, \{0, 1\}, d, b, q_0, q_5)$ arbeitet folgendermaßen:

1. Das Eingabewort wird von links nach rechts gelesen. Jedesmal, wenn eine Folge 10^{a_i} erreicht wird, wird die Folge der Nullen entweder auf das 2. oder das 3. Band kopiert
2. Wenn das Ende des Eingabeworts erreicht ist, wird geprüft, ob auf dem 2. und 3. Band die gleiche Anzahl von Nullen steht; dieses geschieht durch simultanes Vorrücken der Köpfe nach links.

Die Überföhrungsfunktion d wird durch folgende Tabelle gegeben.

momentaner Zustand	Symbol unter dem Schreib/Lesekopf auf			neuer Zustand	neues Symbol, Kopfbewegung auf		
	Band 1	Band 2	Band 3		Band 1	Band 2	Band 3
q_0	1	b	b	q_1	1, S	\$, R	\$, R
q_1	1	b	b	q_2	1, R	b , S	b , S
				q_3	1, R	b , S	b , S
q_2	0	b	b	q_2	0, R	0, R	b , S
	1	b	b	q_1	1, S	b , S	b , S
	b	b	b	q_4	b , S	b , L	b , L
q_3	0	b	b	q_3	0, R	b , S	0, R
	1	b	b	q_1	1, S	b , S	b , S
	b	b	b	q_4	b , S	b , L	b , L
q_4	b	0	0	q_4	b , S	0, L	0, L
	b	\$	\$	q_5	b , S	\$, S	\$, S

Nichtdeterministische Strategien können den Ablauf von Berechnungen vereinfachen. Als Beispiel werde die Sprache

$$L = \{x\#y \mid x \in \{0,1\}^*, y \in \{0,1\}^*, x \neq e, x \neq y\}$$

betrachtet. Es sei $w \in \{0,1\}^*$. Eine deterministische Turingmaschine würde bei Eingabe von w zunächst an der Position des Zeichens $\#$ feststellen, wo y beginnt, und dann x und y buchstabenweise auf einen Unterschied hin vergleichen (falls w kein oder mehr als ein Zeichen $\#$ enthält, würde w nicht akzeptiert werden). Eine nichtdeterministische Turingmaschine NTM würde ebenfalls (deterministisch) zunächst prüfen, ob w genau ein Zeichen $\#$ enthält. Ist dieses nicht der Fall, wird w nicht akzeptiert. Ist $w = x\#y$ und besteht x aus n Buchstaben und y aus m Buchstaben, etwa $x = x_1 \dots x_n$ und $y = y_1 \dots y_m$, so überprüft NTM (deterministisch), ob $n \neq m$ gilt (in diesem Fall wird w akzeptiert). Andernfalls erzeugt NTM auf einem seiner Arbeitsbänder *nichtdeterministisch* den Wert $\text{bin}(i)$ mit $i \leq n = m$. Anschließend überprüft NTM

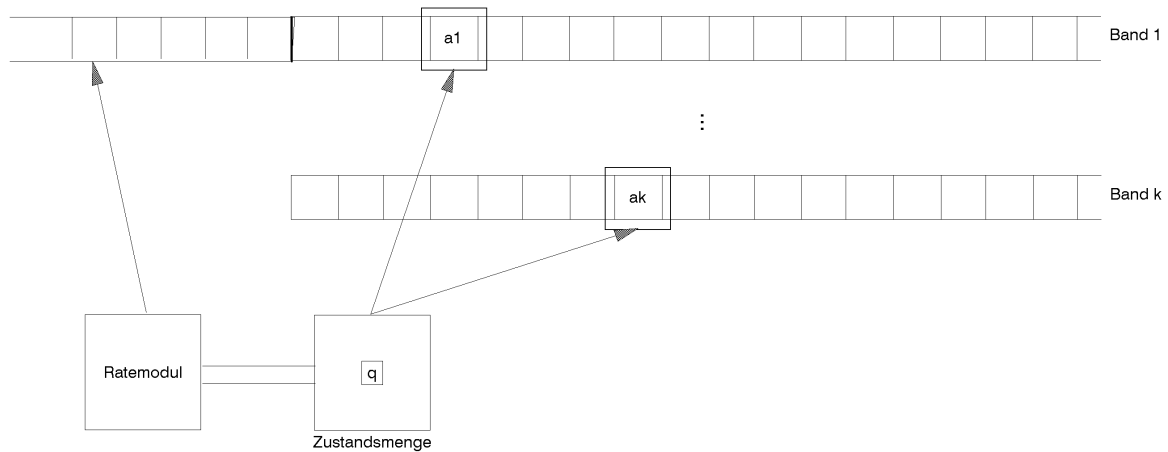
die Buchstaben x_i und y_i . Das Wort w wird genau dann akzeptiert, wenn $x_i \neq y_i$ ist. Für die nichtdeterministische Erzeugung von $bin(i)$ muß in die entsprechende Zelle des Arbeitsbands entweder 0 oder 1 geschrieben werden. *NTM* „rät (nichtdeterministisch)“ die Position, an der sich x und y unterscheiden und verifiziert anschließend die Richtigkeit des Ratevorgangs. Der geratene Wert $bin(i)$ kann auch als Beweis (Zertifikat) dafür angesehen werden, daß $x \neq y$ bzw. $w \in L$ gilt.

Die in diesem Beispiel beschriebene Arbeitsweise ist für eine nichtdeterministische Turingmaschine typisch:

Die nichtdeterministische Turingmaschine *NTM* sei konzipiert, um die Menge $L \subseteq \Sigma^*$ zu akzeptieren. Bei einer Eingabe $x \in \Sigma^*$ **rät** *NTM* **auf nichtdeterministische Weise** einen **Beweis** B (ein **Zertifikat**) dafür, daß $x \in L$ gilt. Der Beweis ist eine Zeichenkette über einem endlichen Alphabet Σ_0 , d.h. $B \in \Sigma_0^*$. Anschließend **verifiziert** *NTM* **den Beweis**.

Eine nichtdeterministische Berechnung kann auch so organisiert werden, daß zunächst alle „nichtdeterministischen Schritte“ ausgeführt werden und anschließend nur noch deterministische Schritte erfolgen. Dieser Ansatz führt auf ein **Nichtstandardmodell** der NDTM:

Eine nichtdeterministische k -Band-Turingmaschine (k -NDTM) TM ist definiert durch $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$. Die Überföhrungsfunktion ist eine partielle Abbildung $d : (Q \setminus \{q_{accept}\}) \times \Sigma^k \rightarrow \mathbf{P}(Q \times (\Sigma \times \{L, R, S\})^k)$, die sich in zwei Teile d_1 und d_2 ($d = d_1 \cup d_2$) zerlegen läßt. Alle nichtdeterministischen Teile von d sind in d_1 zusammengefaßt, d_2 besteht ausschließlich aus deterministischen Teilen. Das Eingabeband (1. Band) wird nach links abzählbar unendlich erweitert wird; die Zellen werden mit 0, -1, -2, ... numeriert. Zusätzlich verfügt TM über ein „Ratemodul“.



Die Arbeitsweise von TM verläuft in zwei Phasen. Ein Eingabewort $w = a_1 \dots a_n$, $w \in I^*$, wird auf das 1. Band in die Zellen mit Nummern 1, ..., n eingegeben (allen anderen Zellen enthalten das Leerzeichen, die Schreib/Leseköpfe stehen auf den jeweiligen Bändern über der ersten Zelle, der Schreibkopf des Ratemoduls steht über der Zelle mit Nummer -1). Nun beginnt Phase 1. Das Ratemodul steuert getaktet den Schreibkopf: in jedem Takt wird ein Symbol aus Σ auf das 1. Band geschrieben und der Schreibkopf um eine Zelle nach links bewegt, oder der Ratemodul stoppt und wird „inaktiv“. Dieser Vorgang wird durch den nichtdeterministischen Teil d_1 von d gesteuert. Die so vom Ratemodul auf dem 1. Band erzeugte Zeichenkette wird auch als **Beweis (Zertifikat)** bezeichnet. TM befindet sich im Anfangszustand q_0 . Eventuell stoppt das Ratemodul nicht, so daß die Turingmaschine bei der entsprechenden Eingabe nicht anhält. Sobald das Ratemodul stoppt und TM in den Zustand q_0 geht, beginnt Phase 2. TM verhält sich deterministisch, gesteuert durch den deterministischen Teil d_2 von d , wie eine „normale“ k -DTM, wobei die Zeichenkette, die das Ratemodul in Phase 1 erzeugt hat, in die Berechnung einbezogen wird.

Phase 2 kann auch als **Verifikationsphase** von TM bezeichnet werden. Es wird nämlich die „Brauchbarkeit“ der in Phase 1 erzeugten („geratenen“) Zeichenkette, der in Phase 1 erzeugte Beweis, für die Berechnung verifiziert.

Die Zeitkomplexität und die Platzkomplexität (im schlechtesten Fall, worst case) einer nicht-deterministischen Turingmaschine wird ähnlich wie im deterministischen Fall definiert:

Für eine k -NDTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ und eine Eingabe $w \in L(TM)$ gelte

$$(q_0, (w, 1), (e, 1), \dots, (e, 1)) \Rightarrow^m K_{accept}$$

mit einer Endkonfiguration K_{accept} . Hierbei sei m der *kleinste* Wert, so daß eine Endkonfiguration erreicht wird. Dann wird durch $t_{TM}(w) = m$ eine partielle Funktion $t_{TM} : \Sigma^* \rightarrow \mathbb{N}$ definiert, die angibt, wieviele Überführungen TM in der kürzesten Berechnung macht, um w zu akzeptieren (eventuell ist $t_{TM}(w)$ nicht definiert, nämlich dann, wenn die Eingabe von w nicht auf eine Endkonfiguration führt).

Die **Zeitkomplexität** von TM (**im schlechtesten Fall, worst case**) wird definiert durch die partielle Funktion $T_{TM} : \mathbb{N} \rightarrow \mathbb{N}$ mit

$$T_{TM}(n) = \max \{ t_{TM}(w) \mid w \in \Sigma^* \text{ und } |w| \leq n \}.$$

Entsprechend kann man die **Platzkomplexität** von TM (**im schlechtesten Fall, worst case**) $S_{TM}(n)$ als den maximalen Abstand vom linken Ende eines Bandes definieren, den ein Schreib/Lesekopf bei der Akzeptanz eines Worts $w \in L(TM)$ mit $|w| \leq n$ *mindestens* erreicht.

Beispielsweise ist die oben angegebene Turingmaschine zur Lösung des Partitionenproblems $(2n + 2)$ -zeitbeschränkt und $(n + 1)$ -raumbeschränkt.

Nichtdeterministisches Verhalten läßt sich deterministisch simulieren:

Satz 2.5-1:

Falls L von einer k -NDTM TM akzeptiert wird, dann gibt es eine k' -DTM TM' mit $L(TM') = L(TM)$. Man kann TM' so konstruieren, daß gilt:

Falls TM $T(n)$ -zeitbeschränkt ist mit $T(n) \geq n$, dann ist $TM' \in O(T(n) \cdot d^{T(n)})$ -zeitbeschränkt mit einer Konstanten $d > 0$.

Beweis:

Es sei TM eine k -NDTM, die $T(n)$ -zeitbeschränkt mit $T(n) \geq n$ ist. Für jedes Wort $w \in L(TM)$ mit $|w| = n$ gibt es also eine Folge von Konfigurationen, die von einer Anfangskonfiguration K_0 mit w zu einer akzeptierenden Konfiguration K_{accept} führt und deren Länge $\leq T(n)$ ist.

Jede Zeile der Tabelle, mit der die Überföhrungsfunktion d von TM gegeben wird (Überföhrungstabelle), hat die Form

$$d(q, a_1, \dots, a_k) = \{ (q_1, (b_{11}, d_{11}), \dots, (b_{1k}, d_{1k})), \dots, (q_t, (b_{t1}, d_{t1}), \dots, (b_{tk}, d_{tk})) \}$$

(hier stehen t „Teilzeilen“). Wenn in einer Konfigurationsfolge $d(q, a_1, \dots, a_k)$ angewendet wird, gibt es t mögliche Folgekonfigurationen. Der maximale Wert t an Teilzeilen in einer Zeile der Überführungstabelle sei d . In einer Berechnung von TM gibt es für eine Konfiguration maximal d mögliche Folgekonfigurationen. Es sei $D = \{0, 1, \dots, d-1\}$. Dann kann man eine von TM ausgeführte Berechnung bzw. die dabei durchlaufene Konfigurationsfolge der Länge l durch ein Wort $x = d_1 \dots d_l$ über D beschreiben: d_i gibt an, daß in der i -ten Überführung die $(d_i + 1)$ -te Alternative in der entsprechenden Zeile der Überführungstabelle verwendet wird. Eventuell beschreibt nicht jedes Wort über D gültige Konfigurationsfolgen von TM .

Eine Simulation des Verhaltens von TM durch eine deterministische Turingmaschine TM' kann folgendermaßen durchgeführt werden. Eine Eingabe $w \in I^*$ wird auf einem Band von TM' zwischengespeichert (es wird eine Kopie angelegt). TM' erzeugt ein Wort $x \in D^*$ mit $|x| = l \leq T(n)$, $x = d_1 \dots d_l$, falls es noch ein nicht bereits erzeugtes Wort dieser Art gibt. Insgesamt können die Wörter in lexikographischer Reihenfolge erzeugt werden. Dann kopiert TM' das Eingabewort w aus dem Zwischenspeicher auf das Eingabeband von TM und simuliert auf deterministische Weise das Verhalten von TM für l Schritte, wobei in der i -ten Überführung die $(d_i + 1)$ -te Alternative in der entsprechenden Zeile der Überführungstabelle verwendet wird. Falls die Simulation auf den akzeptierenden Zustand führt, wird das Wort w akzeptiert, ansonsten wird das nächste Wort $x \in D^*$ erzeugt und die Simulation erneut ausgeführt.

Um ein Wort $x \in D^*$ mit $|x| = l \leq T(n)$ zu erzeugen, benötigt man $O(l)$ viele Schritte (entsprechend der Addition einer 1 auf eine Zahl mit l Stellen). Anschließend wird w in $O(n)$ vielen Schritten aus dem Zwischenspeicher auf das Eingabeband von TM kopiert. Das Verhalten von TM wird dann für l viele Schritte simuliert. Es gibt d^l viele Möglichkeiten für $x \in D^*$ mit $|x| = l$. Insgesamt ergibt sich ein zeitlicher Aufwand der Größe

$$T'(n) = \sum_{l=0}^{T(n)} (c_1 \cdot l + c_2 \cdot n + c_3 \cdot l) \cdot d^l$$

(der erste Term unter dem Summenzeichen beschreibt den zeitlichen Aufwand, um x zu erzeugen, der zweite Term, um w aus dem Zwischenspeicher auf das Eingabeband von TM zu kopieren, der dritte Term, um l Schritte von TM zu simulieren). Es gilt

$$T'(n) \leq c \left(\sum_{l=0}^{T(n)} \left(l \cdot d^l + n \cdot \sum_{l=0}^{T(n)} d^l \right) \right) = c \left(\frac{d - (T(n)+1)d^{T(n)+1} + T(n)d^{T(n)+2}}{(d-1)^2} + n \frac{d^{T(n)+1} - 1}{d-1} \right) \text{ mit ei-}$$

ner geeigneten Konstanten c . Der Ausdruck rechts ist von der Ordnung $O(T(n) \cdot d^{T(n)})$; hierbei wurde $n \leq T(n)$ verwendet. ///

Die deterministische Simulation der k -NDTM TM erfolgt also zum Preis einer exponentiellen Steigerung des Laufzeitverhaltens, denn $O(T(n) \cdot d^{T(n)}) \subseteq O(c^{T(n)})$ mit einer Konstanten $c > 0$.

Bisher ist keine nicht-triviale untere Schranke für die Simulation einer NDTM durch eine DTM bekannt. Insbesondere ist nicht bekannt, ob eine NDTM, deren Laufzeit durch ein Polynom begrenzt ist, durch eine DTM simuliert werden kann, deren Laufzeit ebenfalls ein Polynom ist (eventuell von einem sehr viel höheren Grad).

Anders verhält es sich bei Betrachtung der Platzkomplexität:

Eine Funktion $S : \mathbb{N} \rightarrow \mathbb{N}$ heißt **platz-konstruierbar**, wenn es eine k -DTM TM gibt, die bei Eingabe eines Wortes der Länge n ein spezielles Symbol in die $S(n)$ -te Zelle eines ihrer Bänder schreibt, ohne jeweils mehr als $S(n)$ viele Zellen auf allen Bändern zu verwenden.

Satz 2.5-2:

Ist TM eine k -NDTM mit einer platz-konstruierbaren Speicherplatzkomplexität $S(n)$, dann gibt es eine k' -DTM TM' mit einer Speicherplatzkomplexität der Ordnung $O(S^2(n))$ und $L(TM') = L(TM)$.

Beweis:

Es sei $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ eine k -NDTM mit einer platz-konstruierbaren Speicherplatzkomplexität $S(n)$. Es wird ein deterministischer Algorithmus, d.h. eine deterministische k' -DTM TM' , angegeben, der das Verhalten von TM bei Eingabe eines Wortes w mit $|w| = n$ simuliert und dessen Speicherplatzbedarf durch einen Wert der Größenordnung $O(S^2(n))$ beschränkt ist. Es sei $K = (q, (a_1, i_1), \dots, (a_k, i_k))$ eine Konfiguration in einer Berechnung von TM bei Eingabe von w . Wegen $|a_j| \leq S(n)$ und $1 \leq i_j \leq S(n)$ für $j = 1, \dots, k$ ist die Anzahl verschiedener Konfigurationen durch $|Q| \cdot |\Sigma|^{k \cdot S(n)} \cdot (S(n))^k \leq c^{S(n)}$ mit einer Konstanten $c > 0$ begrenzt. Falls in einer Berechnung von TM also $K_1 \Rightarrow^* K_2$ gilt, dann kann man annehmen, daß dabei höchstens $c^{S(n)}$ viele Überführungen vorkommen. Es gilt: $K_1 \Rightarrow^* K_2$ in höchstens i Schritten genau dann, wenn es eine Konfiguration K_3 gibt, so daß $K_1 \Rightarrow^* K_3$ in höchstens $\lceil i/2 \rceil$ vielen Schritten und $K_3 \Rightarrow^* K_2$ in höchstens $\lfloor i/2 \rfloor$ vielen Schritten abläuft. Für K_3 kommen nur $c^{S(n)}$ viele Möglichkeiten in Frage. Diese Überlegung führt auf folgenden Algorithmus:

Eingabe: Eine k -NDTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ mit einer platz-konstruierbaren Speicherplatzkomplexität $S(n)$ und ein Wort $w \in I^*$ mit $|w| = n$

Verfahren: Aufruf der Funktion `space_Simulation (TM, w)`

Ausgabe: TRUE, falls $w \in L(TM)$, FALSE sonst.

Die Funktion `space_Simulation` wird in Pseudocode beschrieben; sie besitzt zwei Formalparameter TM und w , über die die Beschreibung einer k -NDTM und ein Eingabewort für diese Turingmaschine eingegeben werden. Innerhalb von `space_Simulation` wird eine Funktion `test` verwendet, die als Parameter zwei Konfigurationen $K1$ und $K2$ und eine natürliche Zahl i hat. Auf die genaue syntaktische Spezifikation soll hier verzichtet werden.

```

FUNCTION space_Simulation (TM : ...;
                           w  : ...) : BOOLEAN;

VAR Kf : ...;
    K0 : ...;
    OK : BOOLEAN;

    FUNCTION test (K1 : ...;
                  K2 : ...;
                  i  : INTEGER) : BOOLEAN;
    VAR resultat : BOOLEAN;
        K3       : ...;
    BEGIN { test }
        resultat := FALSE;
        IF i = 1 THEN BEGIN
            IF (K1  $\Rightarrow$  K2) OR (K1 = K2)
            THEN resultat := TRUE;
        END
        ELSE BEGIN
            Für jede Konfiguration  $K3 = (q, (a_1, i_1), \dots, (a_k, i_k))$ 
            mit  $|a_j| \leq S(n)$  und  $1 \leq i_j \leq S(n)$  DO
            BEGIN
                resultat := test (K1, K3, (i+1) DIV 2)
                AND
                test (K3, K2, i DIV 2);
                IF resultat = TRUE THEN Break;
            END;
        END;
        test := resultat;
    END { test };

```

```

BEGIN { space_Simulation }
  K0 := (q0, (w,1), (e,1), ..., (e,1))
  OK := FALSE;
  FOR_EACH Endkonfiguration Kf = (qf, (a1, |a1|+1), ..., (ak, |ak|+1)) mit |aj| ≤ S(n)
    DO BEGIN
      OK := test (K0, Kf, cS(n));
      IF OK THEN Break;
    END;
  space_Simulation := OK;
END { space_Simulation };

```

Eine genauere Untersuchung der Aufrufe der rekursiven Funktion `test` zeigt, daß bei jedem Aufruf innerhalb von `test` der dritte Parameter im wesentlichen halbiert wird. Der maximale Wert des dritten Parameters ist $i = c^{S(n)}$. Daher werden zu keinem Zeitpunkt der (rekursiven) Abarbeitung von `test` mehr als $1 + \log\lceil c^{S(n)} \rceil \in O(S(n))$ viele Aktivierungsrecords auf dem Stack benötigt. Jeder Aktivierungsrecord hat eine Größe der Ordnung $O(S(n))$, so daß der Speicherplatzbedarf von der Ordnung $O(S^2(n))$ ist. ///

Der Begriff des Nichtdeterminismus läßt sich auf Algorithmen (formuliert als Programm in einer Programmiersprache, vgl. Kapitel 2.3) übertragen und führt auf die Definition **eines nichtdeterministischen Algorithmus unter Einsatz eines Verifizierers**. In Kapitel 1.3 und Kapitel 2.1 wird ein Algorithmus für ein Entscheidungsproblem wie folgt definiert (da der Algorithmus als Programm in einer Programmiersprache beschrieben wird, handelt es sich um einen *deterministischen* Algorithmus).

Ein **deterministischer Algorithmus** A_{L_Π} für ein Entscheidungsproblem Π mit der Menge

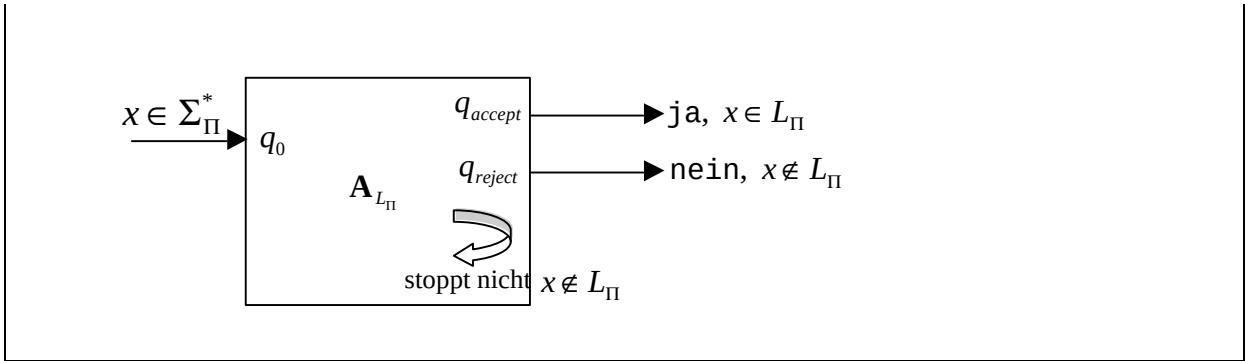
$L_\Pi \subseteq \Sigma_\Pi^*$ hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_\Pi^*$

Ausgabe: ja (accept), falls $x \in L_\Pi$ gilt

nein (reject), falls $x \notin L_\Pi$ gilt

falls A_{L_Π} bei einer Eingabe $x \in \Sigma_\Pi^*$ nicht hält, gilt ebenfalls $x \notin L_\Pi$.



Ein **Verifizierer für das Entscheidungsproblem** Π über einem Alphabet Σ_Π ist ein deterministischer Algorithmus V_{L_Π} , der eine Eingabe $x \in \Sigma_\Pi^*$ und eine **Zusatzinformation** (einen **Beweis**) $B \in \Sigma_0^*$ lesen kann und die Frage „ $x \in L_\Pi$?“ mit Hilfe des Beweises B entscheidet. Die Bereitstellung des Beweises B ist nicht Aufgabe des Verifizierers; B wird „von außen“ vorgegeben. Eventuell stoppt der Verifizierer bei Eingabe von $x \in \Sigma_\Pi^*$ nicht. **Der Verifizierer repräsentiert also die Verifikationsphase** im Nichtstandardmodell einer nichtdeterministischen Turingmaschine; der erforderliche Beweis wird zuvor vom Ratemodul auf nichtdeterministische Weise erzeugt.

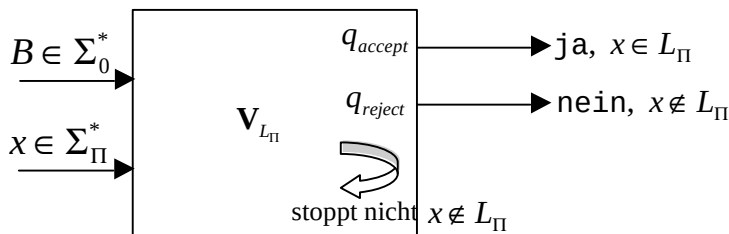
Ein **Verifizierer V_{L_Π} für das Entscheidungsproblem Π** mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_\Pi^*, B \in \Sigma_0^*$

Ausgabe: ja (accept), falls $x \in L_\Pi$ gilt

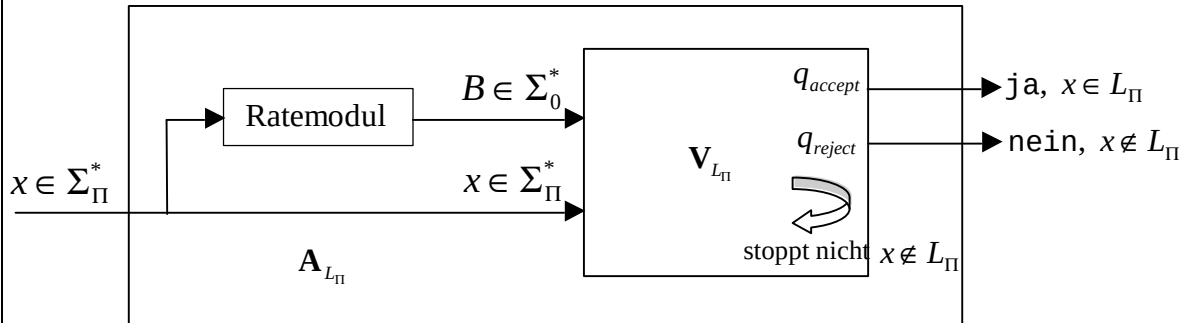
nein (reject), falls $x \notin L_\Pi$ gilt.

falls V_{L_Π} bei einer Eingabe $x \in \Sigma_\Pi^*$ nicht hält, gilt ebenfalls $x \notin L_\Pi$.



Falls V_{L_Π} bei Eingabe von $x \in \Sigma_\Pi^*$ stoppt, wird mit $V_{L_\Pi}(x, B)$, $V_{L_\Pi}(x, B) \in \{\text{ja}, \text{nein}\}$, die Entscheidung von V_{L_Π} bezeichnet.

Ein **nichtdeterministischer Algorithmus** A_{L_Π} für ein Entscheidungsproblem Π hat die Form



Es gilt also für $x \in \Sigma_\Pi^*$:

$x \in L_\Pi$, falls es einen Beweis $B_x \in \Sigma_0^*$ gibt, so daß V_{L_Π} bei Eingabe von x und B_x mit $V_{L_\Pi}(x, B_x) = \text{ja}$ stoppt;

$x \notin L_\Pi$, falls für jeden Beweis $B \in \Sigma_0^*$ gilt: entweder stoppt V_{L_Π} bei Eingabe von x und B nicht, oder V_{L_Π} stoppt mit $V_{L_\Pi}(x, B) = \text{nein}$.

Beispiel:

Erfüllbarkeit Boolescher Ausdrücke

Nicht jeder Boolesche Ausdruck, selbst wenn er syntaktisch korrekt ist, ist erfüllbar. Mit Hilfe eines deterministischen Algorithmus kann man beispielsweise die Erfüllbarkeit eines Booleschen Ausdrucks F mit n Variablen dadurch testen, daß man systematisch nacheinander alle 2^n möglichen Belegungen der Variablen in F mit Wahrheitswerten TRUE bzw. FALSE erzeugt. Immer, wenn eine neue Belegung generiert worden ist, wird diese in die Variablen eingesetzt und der Boolesche Ausdruck ausgewertet. Die Entscheidung, ob F erfüllbar ist oder nicht, kann u.U. erst nach Überprüfung aller 2^n Belegungen erfolgen. Das Verfahren ist daher von der Ordnung $O(2^n)$.

Ein nichtdeterministischer Algorithmus würde bei Eingabe eines Booleschen Ausdrucks F mit n Variablen (in geeigneter Kodierung) in Phase 1 zunächst nichtdeterministisch eine 0-1-Folge der Länge n erzeugen, die einer Belegung der n Variablen entspricht. In Phase 2 wird dann diese Belegung in F eingesetzt, F ausgewertet und die Entscheidung „ F ist erfüllbar“ genau dann getroffen, wenn bei dieser Auswertung der Wahrheitswert TRUE entsteht. In Phase 1 kann, falls F erfüllbar ist, in der Tat eine erfüllende Belegung der Variablen von F er-

zeugt werden. Das geht nicht, wenn F nicht erfüllbar ist; in diesem Fall kann in Phase 1 eine beliebige 0-1-Folge erzeugt werden, die, eingesetzt in die Variablen von F , stets den Wahrheitswert FALSE ergibt.

Ein Verifizierer zur Überprüfung der Erfüllbarkeit eines Booleschen Ausdrucks F mit n Variablen erhält als Eingabe den Ausdruck F in geeigneter Kodierung und als Beweis B_F eine 0-1-Folge der Länge n , die als Belegung der Variablen in F interpretiert wird. Der Verifizierer setzt die Belegung in die Variablen von F ein. Falls F mit B_F erfüllbar ist, trifft der Verifizierer die Entscheidung ja, sonst nein.

Der Verifizierer für das Erfüllbarkeitsproblem können so entworfen werden, daß er ein Laufzeitverhalten der Ordnung $O(n)$ hat. Insgesamt ist die Laufzeit dieses nichtdeterministischen Algorithmus also von der Ordnung $O(n)$. Es ist nicht bekannt, ob es einen deterministischen Algorithmus gibt, der die Erfüllbarkeit Boolescher Ausdrücke mit einem Laufzeitverhalten der Ordnung $O(p(n))$ mit einem Polynom p testet.

3 Grenzen der Berechenbarkeit

In Kapitel 2 werden zwei Ansätze vorgestellt, um Berechenbarkeit zu modellieren. Beide Modelle haben sich als äquivalent erwiesen in dem Sinn, daß die Fähigkeit, etwas zu berechnen, beiden Modellen in gleicher Weise zukommt. Selbst das Konzept des Nichtdeterminismus erweitert nicht die Berechenbarkeit, da nichtdeterministisches Verhalten deterministisch simuliert werden kann, auch wenn dabei die Dauer einer Berechnung exponentiell wächst. Von den Modellen, die in Kapitel 2 behandelt werden, zeichnet sich das Modell der Turingmaschine aufgrund seiner Einfachheit und Universalität und nicht zuletzt aus historischen Gründen gegenüber den anderen Modellen aus.

Im folgenden wird unter dem Begriff Turingmaschine eine k -DTM verstanden (siehe Kapitel 2.1) mit einer dem jeweiligen Problem angemessenen geeigneten Anzahl k an Bändern.

Es sei Σ ein endliches Alphabet und $L \subseteq \Sigma^*$. Die Menge L heißt **rekursiv aufzählbar**, wenn es eine Turingmaschine TM gibt mit $L = L(TM)$.

Die Bezeichnung *rekursiv aufzählbar* leitet sich aus der Tatsache her, daß eine von einer Turingmaschine akzeptierte Menge $L \subseteq \Sigma^*$ durch eine totale (d.h. überall definierte) berechenbare Funktion „aufgezählt“ werden kann. Genauer:

Satz 3-1:

$L \subseteq \Sigma^*$ mit $L \neq \emptyset$ ist genau dann rekursiv aufzählbar, wenn es eine totale berechenbare Funktion $h : (\Sigma \cup \{0, 1, \#\})^* \rightarrow \Sigma^*$ gibt mit $L = h((\Sigma \cup \{0, 1, \#\})^*)$.

Die Funktion h „zählt L rekursiv auf“ (erzeugt L).

Beweis:

Diese Aussage enthält zwei „Richtungen“, deren Beweise beide eine genauere Betrachtung verdienen:

Die eine „Richtung“ dieser Aussage, nämlich der Nachweis der Existenz einer totalen und berechenbaren Funktion $h : (\Sigma \cup \{0, 1, \#\})^* \rightarrow \Sigma^*$ mit $L = h((\Sigma \cup \{0, 1, \#\})^*)$, d.h. einer Funktion h , die L aufzählt, wird folgendermaßen bewiesen. Wegen $L \neq \emptyset$ gibt es ein Wort $w_0 \in L$. Da L als rekursiv aufzählbar angenommen wird, gibt es eine Turingmaschine TM mit $L = L(TM)$. Es wird eine Turingmaschine TM' angegeben, die zwei Bänder mehr als TM

besitzt (nämlich ein zusätzliches Eingabeband und ein Ausgabeband), bei jeder Eingabe stoppt und die gesuchte Funktion $h : (\Sigma \cup \{0, 1, \#\})^* \rightarrow \Sigma^*$ berechnet:

Bei Eingabe von $w \in (\Sigma \cup \{0, 1, \#\})^*$ prüft TM' zunächst, ob w die Form $w = u\#bin(i)$ mit $u \in \Sigma^*$ und der 0-1-Folge $bin(i)$ hat. Falls nicht, schreibt TM' das Wort w_0 auf sein Ausgabeband und stoppt. Falls w diese Form hat, schreibt TM' den Teil u auf das 1. Band von TM und simuliert das Verhalten von TM für höchstens i viele Schritte. Falls TM das Wort u innerhalb dieser Schrittzahl akzeptiert, wird u auf das Ausgabeband geschrieben, und TM' stoppt. Falls TM das Wort u innerhalb dieser Schrittzahl nicht akzeptiert, schreibt TM' das Wort w_0 auf sein Ausgabeband und stoppt. Es sei h die von TM' berechnete Funktion. Dann gilt $L = h((\Sigma \cup \{0, 1, \#\})^*)$:

Ist nämlich $u \in L$, dann akzeptiert TM das Wort in einer endlichen Anzahl i von Schritten. Mit $w = u\#bin(i)$ gilt $h(w) = u$. Das bedeutet $L \subseteq h((\Sigma \cup \{0, 1, \#\})^*)$.

Ist umgekehrt $u \in h((\Sigma \cup \{0, 1, \#\})^*)$, etwa $u = h(w)$ für ein Wort $w \in (\Sigma \cup \{0, 1, \#\})^*$, dann ist entweder $u = w_0$ oder $w = u\#bin(i)$ und TM' hat höchstens i viele Schritte von TM simuliert und dabei $u \in L$ festgestellt. In beiden Fällen ist also $u \in L$, d.h. $h((\Sigma \cup \{0, 1, \#\})^*) \subseteq L$.

Zum Beweis der anderen „Richtung“ der Aussage wird angenommen, daß $L = h((\Sigma \cup \{0, 1, \#\})^*)$ mit einer totalen und berechenbaren Funktion h gilt, und es ist zu zeigen, daß es eine Turingmaschine TM gibt, die genau L akzeptiert:

Bei Eingabe von $w \in \Sigma^*$ verhält sich TM wie folgt. TM erzeugt alle Wörter $u \in (\Sigma \cup \{0, 1, \#\})^*$ in lexikographischer Reihenfolge. Sobald ein Wort u erzeugt ist, berechnet TM den Wert $h(u)$. Gilt $h(u) = w$, so stoppt TM und akzeptiert w . Andernfalls wird das nächste Wort u erzeugt. Es läßt sich $L = L(TM)$ zeigen:

Ist $w \in L$, dann ist nach Voraussetzung $w = h(u)$ für ein Wort $u \in (\Sigma \cup \{0, 1, \#\})^*$. Da $|u| < \infty$ ist, findet TM dieses Wort u (bei der Erzeugung aller Wörter in lexikographischer Reihenfolge). Da h total und berechenbar ist, kann TM den Wert $h(u)$ ermitteln und feststellen, daß $h(u) = w$ gilt. Daher wird w von TM akzeptiert, d.h. $L \subseteq L(TM)$.

Ist $w \notin L$, dann gilt für jedes $u \in (\Sigma \cup \{0, 1, \#\})^*$: $h(u) \neq w$. Dann stoppt TM bei Eingabe von w nicht, d.h. $w \notin L(TM)$. Daher gilt $L(TM) \subseteq L$. ///

In Kapitel 1.1 wurde der Begriff der Abzählbarkeit definiert: Eine Menge M ist abzählbar unendlich, wenn es eine bijektive Abbildung $f : \mathbb{N} \rightarrow M$ gibt, d.h. $M = \{f(i) \mid i \in \mathbb{N}\}$, und jedes Element $m \in M$ trägt eine eindeutige Nummer i : $m = f(i)$. Die Elemente von M können mit natürlichen Zahlen durchnumeriert bzw. indiziert werden. Ersetzt man in dieser Definition den Begriff „bijektiv“ durch „total und berechenbar“ und $\mathbb{N}$ durch $(\Sigma \cup \{0, 1, \#\})^*$, so kommt man auf den Begriff der rekursiven Aufzählbarkeit. Für eine rekursiv aufzählbare Menge L

gilt $L = h\left((\Sigma \cup \{0, 1, \#\})^*\right)$ mit einer totalen und berechenbaren Funktion $h : (\Sigma \cup \{0, 1, \#\})^* \rightarrow \Sigma^*$. Die Elemente $u \in L$ können aus Elementen gewonnen werden, die als Zusatzinformation eine obere Schranke für die Schrittzahl enthalten, die eine Turingmaschine benötigt, um das jeweilige Wort zu akzeptieren.

Die Buchstaben des endlichen Alphabets $\Sigma = \{a_0, \dots, a_{|\Sigma|-1}\}$ einer Turingmaschine lassen sich als Zeichenkette über dem Alphabet $\{0, 1\}$ kodieren: der Buchstabe a_i hat dabei die Kodierung $\text{bin}(a_i)$. Ein Wort $w = a_1 \dots a_n$ kann man dann zunächst in $\text{bin}(a_1)\# \dots \# \text{bin}(a_n)\#$ umsetzen und dann die darin vorkommenden einzelnen Zeichen aus $\{0, 1, \#\}$ in eine 0-1-Folge, wie es etwa in Kapitel 2.4 beschrieben wurde, umkodieren. Entsprechend kann man Paare (w, v) von Worten $w = a_1 \dots a_n$ und $v = b_1 \dots b_m$ zunächst in $(\text{bin}(a_1)\# \dots \# \text{bin}(a_n)\# \text{bin}(b_1)\# \dots \# \text{bin}(b_m)\#)$ umsetzen und die darin vorkommenden Zeichen aus $\{0, 1, \#, (,)\}$ ähnlich wie in Kapitel 2.4 in eine 0-1-Folge umkodieren. Je nach Anwendung soll es daher erlaubt sein, daß eine Turingmaschine Eingaben der Form $w \in \Sigma^*$, aber auch Eingaben der Form $(w, v) \in \Sigma^* \times \Sigma^*$ verarbeitet. Letztlich lassen sich diese aus Paaren bestehenden Eingaben mit „normalen“ 0-1-Folgen identifizieren.

Der obige Satz kann dann auch so formuliert werden:

Satz 3-2:

$L \subseteq \Sigma^*$ mit $L \neq \emptyset$ ist genau dann rekursiv aufzählbar, wenn es eine totale berechenbare Funktion $h : \{0, 1\}^* \rightarrow \Sigma^*$ gibt mit $L = h(\{0, 1\}^*)$.

Die Funktion h „zählt L rekursiv auf“ (erzeugt L).

Im folgenden wird die Klasse der Sprachen, die von Turingmaschinen erkannt werden, d.h. die Klasse der rekursiv aufzählbaren Sprachen, genauer beschrieben. Es ist klar, daß jede rekursiv aufzählbare Menge über einem Alphabet Σ endlich oder abzählbar unendlich ist:

Satz 3-3:

Die Klasse der rekursiv aufzählbaren Mengen über einem Alphabet Σ ist in der Klasse abzählbaren Teilmengen von Σ^* enthalten.

Es fragt sich, ob umgekehrt jede abzählbare Teilmenge von Σ^* bereits rekursiv aufzählbar ist. Folgende Überlegungen zeigen, daß diese Frage verneint werden muß.

Zu jeder rekursiv aufzählbaren Teilmenge von Σ^* gibt es nach Definition eine Turingmaschine, die die Menge akzeptiert. Natürlich kann eine rekursiv aufzählbare Menge von verschiedenen Turingmaschinen akzeptiert werden. Andererseits ist jede von einer Turingmaschine akzeptierte Menge rekursiv aufzählbar. Daher gibt es höchstens so viele rekursiv aufzählbare Teilmengen von Σ^* wie Turingmaschinen mit dem Alphabet Σ . Die Menge der Turingmaschinen mit dem Alphabet Σ ist abzählbar; denn jede Turingmaschine TM läßt sich mit Hilfe der injektiven Funktion *code* auf ein endlich langes Wort $code(TM)$ über $\{0,1\}$ abbilden, (siehe Kapitel 2.4). Satz 1.1-4 sagt aus, daß es überabzählbar viele Teilmengen von Σ^* gibt. Daher existieren auch nicht-rekursiv aufzählbare Teilmengen von Σ^* .

Eine derartige nicht rekursiv aufzählbare Teilmenge von Σ^* kann auch konstruktiv durch eine Technik angegeben werden, die man als **Diagonalisierung** bezeichnet. Der Einfachheit halber soll hier $\Sigma = \{0,1\}$ angenommen werden. Die Elemente von $\Sigma^* = \{0,1\}^*$ werden in lexikographischer Reihenfolge aufgezählt:

Nummer i	Wort $w_i \in \{0,1\}^*$	Nummer i	Wort $w_i \in \{0,1\}^*$
0	ϵ	11	100
1	0	12	101
2	1	13	110
3	00	14	111
4	01	15	0000
5	10	16	0001
6	11	17	0010
7	000	18	0011
8	001	19	0100
9	010	20	0101
10	011	...	...

Das i -te Wort w_i kann als Eingabe für eine Turingmaschine mit Eingabealphabet $\{0,1\}$ und auch, falls $VERIFIZIERE_TM(w_i) = \text{TRUE}$ gilt, als die von w_i kodierte Turingmaschine K_{w_i} interpretiert werden (siehe Kapitel 2.4). Es läßt sich daher eine Matrix M definieren, deren Zeilen und Spalten mit den Wörtern w_i markiert sind. Die Zeilenmarkierung w_i stellt ein derartiges Eingabewort für eine Turingmaschine dar, die Spaltenmarkierung w_j bezeichnet eventuell die Turingmaschine K_{w_j} . Im Schnittpunkt der i -ten Zeile von M mit der j -ten Spalte wird entweder der Wert 0 oder der Wert 1 eingetragen, und zwar so, daß dort

$$\begin{cases} 0 & \text{für } VERIFIZIERE_TM(w_j) = \text{FALSE} \text{ oder } w_i \notin L(K_{w_j}) \\ 1 & \text{für } VERIFIZIERE_TM(w_j) = \text{TRUE} \text{ und } w_i \in L(K_{w_j}) \end{cases}$$

steht. Da die ersten Spaltenmarkierungen $w_0, w_1, w_2, \dots$ nicht sinnvolle Turingmaschinen kodieren, enthält die Matrix M im linken Teil nur die Einträge 0.

Satz 3-4:

Die Menge Menge $L_d \subseteq \{0,1\}^*$ sei definiert durch:

Es ist $w \in L_d$ genau dann, wenn $w = w_i$ ist (das i -te Wort in der lexikographischen Reihenfolge) und M in der i -ten Zeile und i -ten Spalte den Eintrag 0 hat.

Dann gilt:

Die so definierte Menge L_d ist nicht rekursiv aufzählbar, d.h. es gibt keine Turingmaschine TM mit $L(TM) = L_d$.

Beweis:

Offensichtlich ist $L_d \neq \emptyset$.

Angenommen, L_d ist rekursiv aufzählbar. Dann gibt es eine Turingmaschine TM mit $L(TM) = L_d$. Die Kodierung dieser Turingmaschine sei w_i , d.h. $TM = K_{w_i}$ und $L_d = L(K_{w_i})$. Gilt nun $w_i \in L_d$? Falls $w_i \in L_d$ gilt, dann enthält nach Definition von L_d die Matrix M in der i -ten Zeile und i -ten Spalte den Eintrag 0. Nach Definition von M und wegen $VERIFIZIERE_TM(w_i) = \mathbf{TRUE}$ bedeutet das aber: $w_i \notin L(K_{w_i})$, also $w_i \notin L_d$. Dieser Widerspruch erlaubt nur die Alternative $w_i \notin L_d$. Nach Definition von L_d enthält die Matrix M in der i -ten Zeile und i -ten Spalte den Eintrag 1. Das bedeutet nach Definition von M : $w_i \in L(K_{w_i})$ und damit $w_i \in L_d$. Dieser Widerspruch zeigt, daß die ursprüngliche Annahme, daß L_d rekursiv aufzählbar sei, falsch ist. ///

3.1 Eigenschaften rekursiv aufzählbarer und rekursiver Mengen

In Kapitel 2.1 wurden bereits einige Eigenschaften rekursiv aufzählbarer Mengen angegeben:

Satz 3.1-1:

1. Sind $L_1 \subseteq \Sigma^*$ und $L_2 \subseteq \Sigma^*$ rekursiv aufzählbar, dann sind auch $L_1 \cup L_2$ und $L_1 \cap L_2$ rekursiv aufzählbar. Die Klasse der rekursiv aufzählbaren Mengen über einem endlichen Alphabet Σ ist abgeschlossen gegenüber Vereinigung- und Schnittbildung.
2. Ist $L \subseteq \Sigma^*$ rekursiv aufzählbar, dann ist nicht notwendigerweise auch $\Sigma^* \setminus L$ rekursiv aufzählbar. Die Klasse der rekursiv aufzählbaren Mengen über einem endlichen Alphabet Σ ist nicht abgeschlossen gegenüber Komplementbildung.

Die Gültigkeit der 2. Aussage ergibt sich aus den folgenden Sätzen 3.1-4 und 3.1-6.

Die Menge $L \subseteq \Sigma^*$ werde von der k -DTM TM akzeptiert. Wird $w \in L$ auf das Eingabeband von TM gegeben, dann hält TM nach endlich vielen Schritten im Zustand q_{accept} an. Wird ein Wort w mit $w \in \Sigma^* \setminus L$ auf das Eingabeband gegeben, dann hält TM eventuell nicht an. Es wäre natürlich wünschenswert, daß TM auch in diesem Fall anhält, z.B. im Zustand q_{reject} , der ausdrückt, daß $w \notin L$ gilt. Gibt es also zu jeder von einer Turingmaschine TM akzeptierten Sprache $L \subseteq \Sigma^*$ eine Turingmaschine TM' mit $L(TM') = L$ und der Eigenschaft, daß TM' bei jedem Wort $w \in \Sigma^*$ anhält: im Zustand q_{accept} , falls $w \in L$ ist, und im Zustand q_{reject} , falls $w \notin L$ ist (nichtstoppende Berechnungen kommen nicht vor)? Diese Frage führt auf die folgende Definition:

Die Menge $L \subseteq \Sigma^*$ heißt **entscheidbar** (oder **rekursiv**), wenn es eine Turingmaschine TM gibt mit der Eigenschaft:

TM stoppt bei jeder Eingabe $w \in \Sigma^*$, und TM akzeptiert w genau dann, wenn $w \in L$ ist.

In praktischen Anwendungen sind die entscheidbaren Mengen gerade diejenigen Sprachen, mit denen man „umgehen“ kann. Denn man möchte ja bei Eingabe von $w \in \Sigma^*$ in einen Entscheidungsalgorithmus nach endlicher Zeit die Frage geklärt haben, ob w eine spezifizierte Eigenschaft hat, d.h. ob $w \in L$ ist, oder nicht. Ist L entscheidbar, läßt sich diese Entscheidung definitiv herbeiführen. Bei einer rekursiv aufzählbaren Menge L und einer passenden Turingmaschine (Entscheidungsverfahren), der man eine Eingabe $w \in \Sigma^*$ vorgelegt hat und die bereits einige Rechenschritte vollzogen hat, ohne bisher eine Entscheidung zu treffen, kann man nicht sicher sein, ob überhaupt jemals eine Entscheidung getroffen wird (vgl. dazu Satz 3.1-4).

Satz 3.1-2:

1. Jede entscheidbare Menge ist rekursiv aufzählbar.
Die Klasse der entscheidbaren Mengen über einem endlichen Alphabet Σ ist in der Klasse der rekursiv aufzählbaren Mengen über diesem Alphabet enthalten.
2. Jede endliche Menge ist entscheidbar.
3. Ist $L \subseteq \Sigma^*$ entscheidbar, dann ist auch $\Sigma^* \setminus L$ entscheidbar.
Die Klasse der entscheidbaren Mengen über einem endlichen Alphabet Σ ist abgeschlossen gegenüber Komplementbildung.
4. Sind $L_1 \subseteq \Sigma^*$ und $L_2 \subseteq \Sigma^*$ entscheidbar, dann sind auch $L_1 \cup L_2$ und $L_1 \cap L_2$ entscheidbar.
Die Klasse der entscheidbaren Mengen über einem endlichen Alphabet Σ ist abgeschlossen gegenüber Vereinigung- und Schnittbildung.
5. Es sei $f : \Sigma^* \rightarrow \Sigma'^*$ eine totale und von einer Turingmaschine berechenbare Funktion. Ist $L' \subseteq \Sigma'^*$ entscheidbar, dann ist auch $f^{-1}(L')$ entscheidbar.
6. Es sei $h : \Sigma^* \rightarrow \Sigma'^*$ eine totale und von einer Turingmaschine berechenbare Funktion. Für die Sprachen $L \subseteq \Sigma^*$ und $L' \subseteq \Sigma'^*$ gelte die Eigenschaft

$$w \in L \Leftrightarrow h(w) \in L'$$

Ist L' entscheidbar, dann ist auch L entscheidbar.
Ist L nicht entscheidbar, dann ist auch L' nicht entscheidbar.

Ist L' rekursiv aufzählbar, dann ist auch L rekursiv aufzählbar.

Bemerkung: Man sagt in diesem Fall, daß L auf L' **reduzierbar** sei und schreibt dafür $L \leq L'$.

Beweis:

Zu 1. Die Aussage ergibt sich direkt aus der Definition.

Zu 2. Ist $L = \{w_1, \dots, w_n\}$ eine endliche Menge, $L \subseteq \Sigma^*$, so ist eine auf allen Eingaben $w \in \Sigma^*$ stoppen Turingmaschine TM anzugeben, die genau dann im akzeptierenden Zustand hält, wenn w mit einem der endlich vielen Werte in L übereinstimmt. Es ist also eine Turingmaschine zu definieren, die das (Pseudocode-) Statement

```

CASE w OF
  w1 : accept;
  ...
  wn : accept;
ELSE reject;
END;
realisiert.

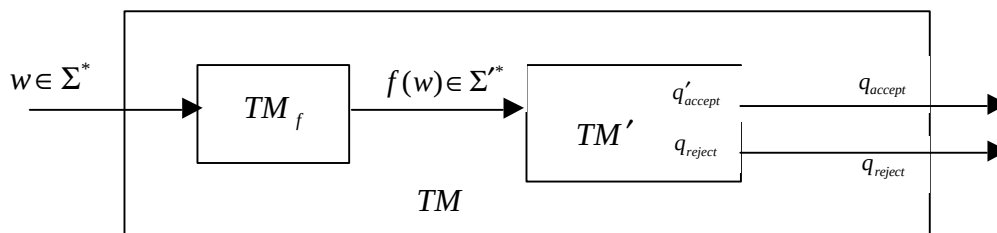
```

Zu 3. Es sei TM eine auf allen Eingaben $w \in \Sigma^*$ stoppende Turingmaschine mit $L = L(TM)$. Man kann annehmen, daß TM den akzeptierenden Zustand q_{accept} hat, der erreicht wird, wenn $w \in L$ gilt, und den verwerfenden Zustand q_{reject} hat, wenn $w \notin L$ ist. Eine auf allen Eingaben $w \in \Sigma^*$ stoppende Turingmaschine TM' mit $L(TM') = \Sigma^* \setminus L$ erhält man aus TM , indem man die Rollen von q_{accept} und q_{reject} vertauscht.

Zu 4. Diese Aussage wird genauso gezeigt wie bei rekursiv aufzählbaren Mengen.

Zu 5. Es sei TM' eine auf allen Eingaben $u \in \Sigma'^*$ stoppende Turingmaschine mit $L(TM') = L'$. Die Funktion $f : \Sigma^* \rightarrow \Sigma'^*$ werde durch die Turingmaschine TM_f berechnet. Durch Hintereinanderschalten von TM_f und TM' erhält man eine Turingmaschine TM , die auf allen Eingaben $w \in \Sigma^*$ stoppt (weil f total ist und TM' auf allen Eingaben stoppt) und für die

$$L(TM) = f^{-1}(L') = \{w \mid f(w) \in L'\} \text{ gilt:}$$



Zu 6. Aus der Eigenschaft $w \in L \Leftrightarrow h(w) \in L'$ folgt $h^{-1}(L') = L$. Mit 5. ergibt sich die Aussage über die Entscheidbarkeit von L . Auf ähnliche Weise wird die Aussage über die rekursive Aufzählbarkeit gezeigt (im vorherigen Bild wird dazu nur der mit q_{accept} markierte Ausgang betrachtet). ///

Satz 3.1-3:

Die Klasse der entscheidbaren Mengen ist eine echte Untermenge der Klasse der rekursiv aufzählbaren Mengen.

Dazu ist mindestens ein Beispiel einer rekursiv aufzählbaren Menge, die nicht entscheidbar sind, anzugeben. Dieses kann konstruktiv geschehen:

Zur Erinnerung (siehe Kapitel 2.4): Für ein Wort $w \in \{0,1\}^*$ ist $VERIFIZIERE_TM(w) = \text{TRUE}$ genau dann, wenn w die Kodierung einer Turingmaschine ist. Die durch w kodierte Turingmaschine wird mit K_w bezeichnet.

Die zum **Halteproblem gehörende Sprache** $L_H \subseteq \{0,1,\#\}^*$ wird definiert als die Menge

$$L_H = \left\{ u\#w \mid \begin{array}{l} u \in \{0,1\}^*, w \in \{0,1\}^*, VERIFIZIERE_TM(w) = \text{TRUE} \text{ und} \\ K_w \text{ stoppt bei Eingabe von } u \end{array} \right\}.$$

Es gilt:

Satz 3.1-4:

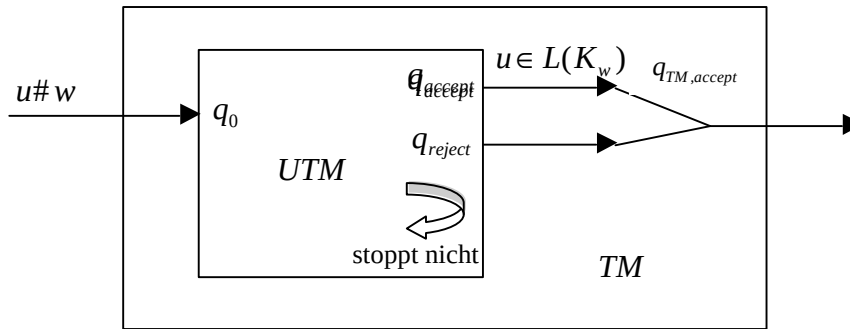
L_H ist rekursiv aufzählbar, aber nicht entscheidbar.

Insbesondere heißt das: es gibt keine immer anhaltende Turingmaschine TM mit akzeptierendem Zustand q_{accept} und nicht-akzeptierendem Zustand q_{reject} mit folgender Eigenschaft: Bei Eingabe von $u\#w$ stoppt TM im akzeptierenden Zustand q_{accept} , falls K_w auf u stoppt, und TM stoppt im nicht-akzeptierenden Zustand q_{reject} , falls K_w auf u nicht stoppt.

Beweis:

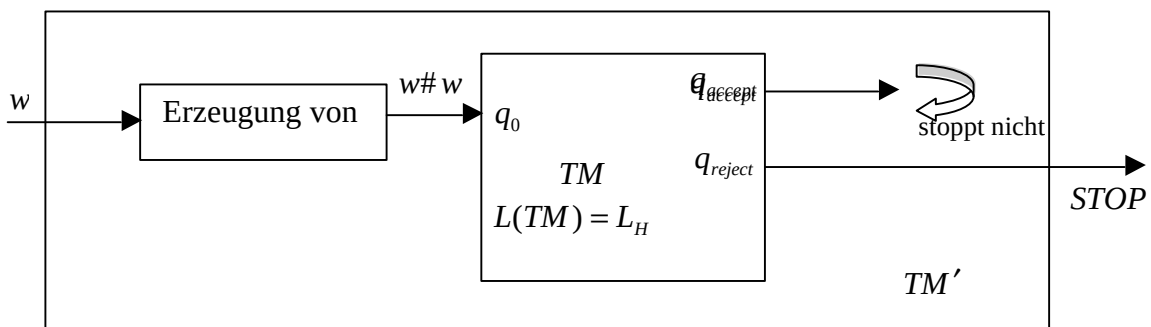
Um die rekursive Aufzählbarkeit von L_H zu zeigen, ist eine Turingmaschine TM mit $L(TM) = L_H$ anzugeben. TM ist eine einfache Modifikation der in Kapitel 2.4 beschriebenen universellen Turingmaschine UTM : Man kann annehmen, daß UTM nur ein Band hat und einen akzeptierenden Zustand q_{accept} und einen nicht-akzeptierenden Zustand q_{reject} besitzt. Kommt UTM in einen dieser beiden Zustände, stoppt UTM . Eventuell stoppt UTM jedoch bei einer Eingabe nicht. Es wird ein neuer Zustand $q_{TM,accept}$ in TM eingeführt und die Überföhrungsfunktion d_{UTM} von UTM so geändert, daß noch eine Überföhrung in den Zustand $q_{TM,accept}$ ausgeföhrt wird, wenn UTM in q_{accept} oder q_{reject} kommt. Dazu wird d_{UTM} um die Zeilen $d_{UTM}(q_{accept}, a) = (q_{TM,accept}, (a, S))$ und $d_{UTM}(q_{reject}, a) = (q_{TM,accept}, (a, S))$ für jedes

Symbol a aus dem Arbeitsalphabet von UTM erweitert. Eine Eingabe $u\#w \in \{0,1,\#\}^*$ für TM wird in UTM eingegeben. Ist $u\#w \in L(UTM)$, d.h. $VERIFIZIERE_TM(w) = \mathbf{TRUE}$ und $u \in L(K_w)$, dann stoppt UTM im Zustand q_{accept} , und TM geht in den Zustand $q_{TM,accept}$ und stoppt; ist $u\#w \notin L(UTM)$, weil UTM in den Zustand q_{reject} kommt, dann geht TM ebenfalls in den Zustand $q_{TM,accept}$ und stoppt. Offensichtlich gilt $L(TM) = L_H$.



Zum Beweis der Nichtentscheidbarkeit von L_H wird wieder die Methode der Diagonalisierung eingesetzt:

Angenommen, es gibt eine auf allen Eingaben $z \in \{0,1,\#\}^*$ stoppende Turingmaschine TM mit $L(TM) = L_H$. Dann wird TM leicht abgewandelt zu einer Turingmaschine TM' , die wie folgt arbeitet: TM' liest eine Eingabe $w \in \{0,1\}^*$ und erzeugt mit einer Kopie von w das Wort $w\#w$. Dann verhält sich TM' wie TM bei Eingabe von $w\#w$. Falls TM im akzeptierenden Zustand q_{accept} stoppt, geht TM' in einen neuen Zustand q' über und stoppt nicht³. Falls TM im nicht-akzeptierenden Zustand q_{reject} stoppt, dann stoppt TM' ; dabei ist es irrelevant, ob TM' in einem akzeptierenden oder einem nicht-akzeptierenden Zustand stoppt.



³ Ist TM' eine k -DTM mit der Überföhrungsfunktion d' , dann ist $d'(q', a_1, \dots, a_k) = (q', (a_1, S), \dots, (a_k, S))$ für $a_1 \in \Sigma, \dots, a_k \in \Sigma$.

Falls TM existiert, dann auch TM' , und zwar mit einer Kodierung $u \in \{0,1\}^*$, d.h. $TM' = K_u$. Nun gibt es zwei Alternativen: TM' stoppt auf der eigenen Kodierung u , oder TM' stoppt auf der eigenen Kodierung u nicht. Im ersten Fall ist $u\#u \notin L_H$; nach Definition von L_H stoppt $K_u = TM'$ bei Eingabe von u nicht. Dieser Widerspruch läßt nur die zweite Alternative zu; das bedeutet aber $u\#u \in L_H$, also stoppt $K_u = TM'$ bei Eingabe von u . Beide Alternativen führen auf einen Widerspruch. Daher existiert TM nicht. ///

Die zum **Halteproblem mit leerem Eingabeband gehörende Sprache** $L_{H_0} \subseteq \{0,1\}^*$ wird definiert als die Menge

$$L_{H_0} = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } K_w \text{ stoppt bei leerem Eingabeband}\}.$$

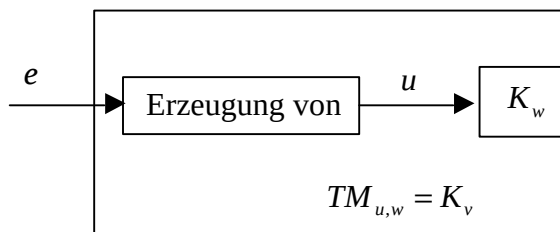
Satz 3.1-5:

L_{H_0} ist rekursiv aufzählbar, aber nicht entscheidbar.

Beweis:

Man zeigt, daß L_H nach L_{H_0} reduzierbar ist, d.h. daß $L_H \leq L_{H_0}$ gilt. Dazu ist eine totale und berechenbare Funktion $h: \{0,1,\#\}^* \rightarrow \{0,1\}^*$ anzugeben, für die $z \in L_H$ genau dann gilt, wenn $h(z) \in L_{H_0}$ ist. Aus Satz 3.1-2, Teil 6 folgt dann die Aussage des Satzes.

Zu $u \in \{0,1\}^*$ und $w \in \{0,1\}^*$ mit $\text{VERIFIZIERE_TM}(w) = \text{TRUE}$ sei $TM_{u,w}$ eine Turingmaschine, die folgendes Verhalten aufweist: $TM_{u,w}$ startet mit leerem Eingabeband und erzeugt das Wort $u \in \{0,1\}^*$. Dann simuliert $TM_{u,w}$ das Verhalten von K_w . Die Kodierung von $TM_{u,w}$ sei v . Es gilt $v \in \{0,1\}^*$ und $\text{VERIFIZIERE_TM}(v) = \text{TRUE}$.



Mit Hilfe eines entsprechenden Algorithmus, eines „Turingmaschinengenerators“, kann man bei Vorgabe von u und w die Kodierung v von $TM_{u,w}$ auf deterministische Weise erzeugen. Die Funktion h wird definiert durch

$$h: \begin{cases} \{0, 1, \#\}^* & \rightarrow \{0, 1\}^* \\ z & \rightarrow \begin{cases} code(TM_{u,w}) & \text{falls } z = u\#w \text{ mit } u \in \{0, 1\}^* \text{ und } w \in \{0, 1\}^* \\ & \text{und } VERIFIZIERE_TM(w) = \text{TRUE} \\ 1 & \text{sonst} \end{cases} \end{cases}$$

Diese ist total und berechenbar. Nach Definition von L_H gilt:

$$\begin{aligned} u\#w \in L_H &\Leftrightarrow K_w \text{ stoppt bei Eingabe von } u \\ &\Leftrightarrow TM_{u,w} = K_v \text{ stoppt bei leerem Eingabeband} \\ &\Leftrightarrow h(u\#w) = v \text{ und } v \in L_{H_0} . /// \end{aligned}$$

Eine Charakterisierung der Entscheidbarkeit einer Menge liefert der folgende Satz:

Satz 3.1-6:

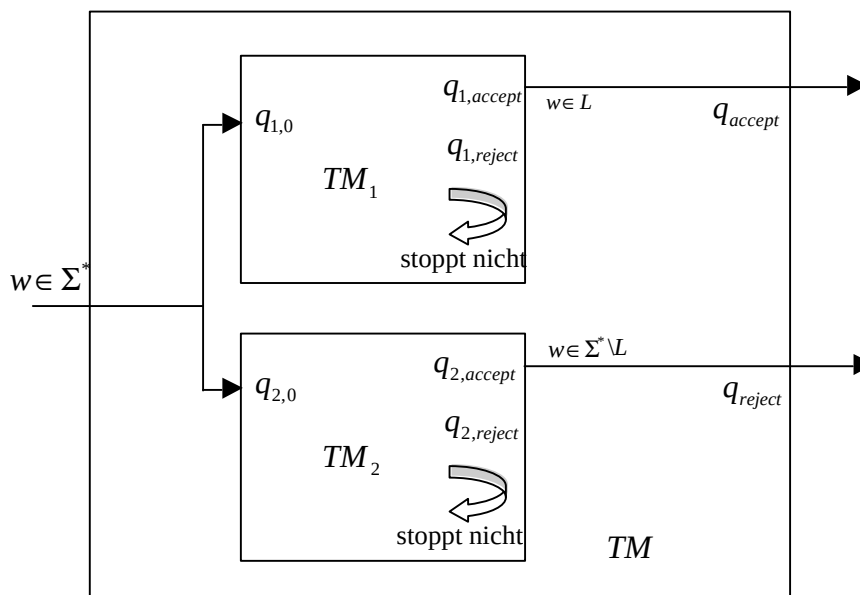
Eine Menge $L \subseteq \Sigma^*$ ist genau dann entscheidbar, wenn sowohl L als auch das Komplement $\Sigma^* \setminus L$ von L rekursiv aufzählbar sind.

Beweis:

Auch dieser Satz beinhaltet wieder zwei „Richtungen“:

Ist die Menge $L \subseteq \Sigma^*$ entscheidbar, dann ist sie rekursiv aufzählbar. Satz 3.1-2 Teil 3. besagt, daß mit L auch $\Sigma^* \setminus L$ entscheidbar und damit rekursiv aufzählbar ist.

Gilt umgekehrt für eine Menge $L \subseteq \Sigma^*$, daß mit ihr auch das Komplement $\Sigma^* \setminus L$ rekursiv aufzählbar ist, dann gibt es zwei Turingmaschinen TM_1 und TM_2 mit $L = L(TM_1)$ und $\Sigma^* \setminus L = L(TM_2)$. Aus diesen beiden Turingmaschinen wird durch Parallelschaltung eine auf allen Eingaben $w \in \Sigma^*$ stoppende Turingmaschine TM konstruiert, die w genau dann akzeptiert, wenn $w \in L$ gilt. Dabei wird eine ähnliche Konstruktion gewählt, wie sie zum Nachweis der Abgeschlossenheit der rekursiv aufzählbaren Mengen bezüglich Vereinigungsbildung angegeben wurde (Kapitel 2.1):



Es gilt für $w \in \Sigma^*$: Ist $w \in L$, dann stoppt TM_1 im Zustand $q_{1,accept}$ und damit stoppt TM im Zustand q_{accept} . Ist $w \notin L$, dann stoppt TM_2 im Zustand $q_{2,accept}$ und damit stoppt TM im Zustand q_{reject} . Da entweder $w \in L$ oder $w \notin L$ gilt, stoppt TM bei allen Eingaben $w \in \Sigma^*$ und akzeptiert genau die Menge L . ///

Die oben aufgeführte zum Halteproblem gehörige Menge L_H ist rekursiv aufzählbar, aber nicht entscheidbar. Das Komplement von L_H , die Menge $\{0,1,\#\}^* \setminus L_H$, kann daher nicht rekursiv aufzählbar sein. Daher ist die Klasse der rekursiv aufzählbaren Mengen gegenüber Komplementbildung nicht abgeschlossen (das ist Satz 3.1-1 Teil 2.).

Da die Rollen von L und $\Sigma^* \setminus L$ im letzten Satz „symmetrisch“ sind, ergibt sich als Konsequenz:

Es sei $L \subseteq \Sigma^*$. Dann gilt

Satz 3.1-7:

Entweder

1. sowohl L als auch $\Sigma^* \setminus L$ ist entscheidbar

oder

2. weder L als noch $\Sigma^* \setminus L$ ist entscheidbar

oder

3. entweder L oder $\Sigma^* \setminus L$ ist rekursiv aufzählbar, aber nicht entscheidbar, und die jeweils andere Menge ist nicht rekursiv aufzählbar.

Die Klasse der rekursiv aufzählbaren Mengen über einem Alphabet Σ ist eine echte Teilklasse von $\mathbf{P}(\Sigma^*)$. Die Klasse der entscheidbaren Mengen über Σ ist echt enthalten in der Klasse der rekursiv aufzählbaren Mengen. Es stellt sich die Frage, ob man mit Hilfe eines durch eine Turingmaschine definierten algorithmischen Verfahrens die Klasse der entscheidbaren bzw. die Klasse der nicht-entscheidbaren Mengen erkennen kann oder wenigstens rekursiv aufzählen kann. Genauer: Es wird danach gefragt, ob es eine Turingmaschine gibt, die die Kodierungen aller derjenigen Turingmaschinen erzeugt (rekursiv aufzählt), deren akzeptierte Sprachen gerade die entscheidbaren bzw. nichtentscheidbaren Mengen sind. Leider ist das nicht möglich, wie folgender Satz zeigt:

Satz 3.1-8:

Es seien

$$L_e = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) \text{ ist entscheidbar}\} \text{ und}$$

$$L_{ne} = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) \text{ ist nicht entscheidbar}\}.$$

Dann ist weder L_e noch L_{ne} rekursiv aufzählbar (und damit auch nicht entscheidbar).

Beweis:

Es wird die von der universellen Turingmaschine UTM akzeptierte Sprache

$$L_{uni} = L(UTM) = \left\{ u \# w \mid \begin{array}{l} u \in \{0,1\}^*, w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE und } u \in L(K_w) \end{array} \right\} \subseteq \{0,1,\#\}^* \text{ herangezogen.}$$

L_{uni} ist rekursiv aufzählbar, aber nicht entscheidbar. Daher ist das Komplement

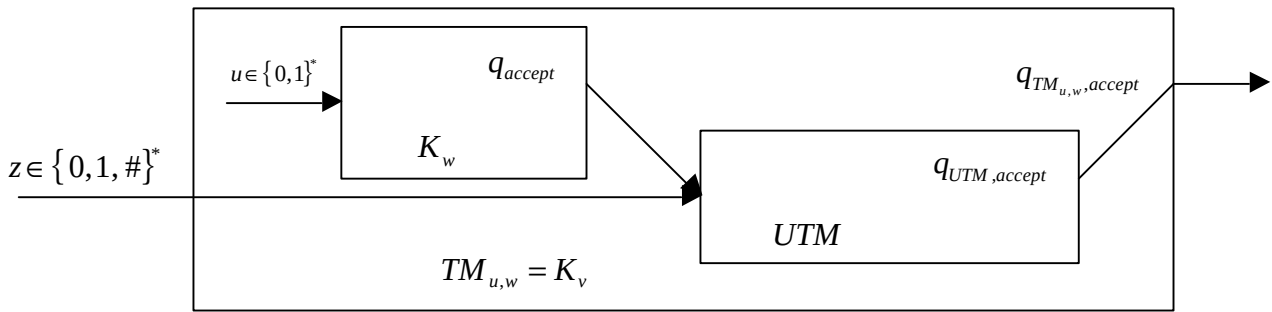
$\bar{L}_{uni} = \{0,1,\#\}^* \setminus L_{uni}$ nicht rekursiv aufzählbar. Es läßt sich zeigen, daß $\bar{L}_{uni} \leq L_e$ gilt; das

Aussehen der dabei beteiligten totalen und berechenbaren Funktion $h: \{0,1,\#\}^* \rightarrow \{0,1\}^*$ wird im folgenden skizziert. Aus Satz 3.1-2 Teil 6. folgt, daß L_e nicht rekursiv aufzählbar ist.

Der Beweis für L_{ne} verläuft analog.

Ähnlich wie oben wird zu $u \in \{0,1\}^*$ und $w \in \{0,1\}^*$ mit $\text{VERIFIZIERE_TM}(w) = \text{TRUE}$ mit Hilfe eines deterministisch arbeitenden Turingmaschinengenerators eine Turingmaschine $TM_{u,w}$ konstruiert, die wie folgt arbeitet:

Bei Eingabe von $z \in \{0,1,\#\}^*$ wird die Eingabe zunächst nicht berücksichtigt, sondern $TM_{u,w}$ verhält sich wie K_w auf der Eingabe u . Falls $u \in L(K_w)$ festgestellt wird, simuliert $TM_{u,w}$ das Verhalten von UTM auf der Eingabe z . Die Kodierung von $TM_{u,w}$ sei v . Es gilt $v \in \{0,1\}^*$ und $\text{VERIFIZIERE_TM}(v) = \text{TRUE}$. Das Wort v ist deterministisch berechenbar, da nur Beschreibungen von Turingmaschinen zusammengestellt wurden.



$$\text{Es gilt } L(K_v) = L(TM_{u,w}) = \begin{cases} L(UTM) = L_{uni} & \text{falls } u \in L(K_w) \text{ ist} \\ \emptyset & \text{falls } u \notin L(K_w) \text{ ist} \end{cases}$$

Bemerkung: Offensichtlich ist $L(K_v)$ genau dann entscheidbar, wenn $u \notin L(K_w)$ ist; denn dann ist $L(K_v) = \emptyset$. Für $u \in L(K_w)$ ist $L(K_v) = L_{uni}$, und diese Menge ist nicht entscheidbar.

Es sei $w_0 \in \{0,1\}^*$ die Kodierung einer Turingmaschine mit $L(K_{w_0}) = \{0,1\}^*$. Es ist $w_0 \in L_e$.

Die gesuchte Funktion h wird definiert durch

$$h: \begin{cases} \{0,1,\#\}^* & \rightarrow \{0,1\}^* \\ z & \rightarrow \begin{cases} code(TM_{u,w}) & \text{falls } z = u\#w \text{ mit } u \in \{0,1\}^* \text{ und } w \in \{0,1\}^* \\ & \text{und } VERIFIZIERE_TM(w) = TRUE \\ w_0 & \text{sonst} \end{cases} \end{cases}$$

Wie obige Ausführungen zeigen, ist h total und berechenbar. Außerdem gilt

$$z \in \bar{L}_{uni} \Leftrightarrow z \text{ hat nicht die Form } z = u\#w \text{ mit } u \in \{0,1\}^* \text{ und } w \in \{0,1\}^*$$

$$\text{oder } VERIFIZIERE_TM(w) = FALSE$$

$$\text{oder } z \notin L_{uni}, \text{ d.h. } u \notin L(K_w).$$

In den ersten beiden Fällen ist $h(z) = w_0$, also $h(z) \in L_e$. Im dritten Fall ist $h(z) = code(TM_{u,w}) = v$ und nach obiger Bemerkung und nach Definition von L_e : $h(z) \in L_e$, insgesamt also $z \in \bar{L}_{uni} \Leftrightarrow h(z) \in L_e$, und damit ist $\bar{L}_{uni} \leq L_e$ gezeigt. ///

Ähnlich verhält es sich mit den totalen berechenbaren Funktionen. Für ein Wort $w \in \{0,1\}^*$ sei $f_{K_w}: \{0,1\}^* \rightarrow \{0,1\}^*$ die von der Turingmaschine K_w berechnete Funktion. Dann gilt:

Satz 3.1-9:

Die Menge $L_t = \{w \mid f_{K_w} \text{ ist eine totale Funktion}\}$ ist nicht rekursiv aufzählbar.

Im folgenden wird nach einem allgemeinen **algorithmischen Verfahren** gefragt, das einer Turingmaschine irgendeine nichttriviale Eigenschaft ansieht. Unter einem algorithmischen Verfahren ist hierbei eine auf allen Eingaben (entweder akzeptierend oder nicht akzeptierend) stoppende Turingmaschine zu verstehen. Der Begriff „nichttriviale Eigenschaft“ wird wie folgt definiert:

Eine Menge $L \subseteq \{0,1\}^*$ von Kodierungen von Turingmaschinen heißt **nichttriviales Entscheidungsproblem über Turingmaschinen**, wenn gilt:

- (i) $L \neq \emptyset$
- (ii) L enthält nicht die Kodierungen aller Turingmaschinen
- (iii) für zwei Turingmaschinen TM_1 und TM_2 , die durch w_1 bzw. w_2 kodiert werden, d.h. $TM_1 = K_{w_1}$ und $TM_2 = K_{w_2}$, impliziert $L(TM_1) = L(TM_2)$: Es gilt $w_1 \in L$ genau dann, wenn $w_2 \in L$ gilt.

Die **Frage nach der Entscheidbarkeit eines nichttrivialen Entscheidungsproblems L über Turingmaschinen** lautet dann in unterschiedlichen äquivalenten Formulierungen:

- Ist L entscheidbar?
- Gibt es ein algorithmisches Verfahren, das bei Eingabe eines Wortes $w \in \{0,1\}^*$ nach endlich vielen Schritten entscheidet, ob $w \in L$ gilt oder nicht?
- Gibt es ein algorithmisches Verfahren, das bei Eingabe der Beschreibung einer Turingmaschine nach endlich vielen Schritten entscheidet, ob die Kodierung der Turingmaschine zu L gehört oder nicht?
- Die Turingmaschinen, deren Kodierungen zu L gehören, haben alle eine durch L beschriebene Eigenschaft E_L . **Gibt es ein algorithmisches Verfahren, das bei Eingabe einer Turingmaschine nach endlich vielen Schritten entscheidet, ob die Turingmaschine die Eigenschaft E_L aufweist oder nicht?**
- Kann man mit Hilfe eines algorithmischen Verfahrens entscheiden, ob eine Turingmaschine eine definierte Eigenschaft E_L aufweist oder nicht?

Beispiele nichttrivialer Entscheidungsprobleme:

1. $L = \{ w \mid \text{die durch } w \text{ kodierte Turingmaschine } K_w \text{ berechnet eine konstante Funktion} \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob eine Turingmaschine eine konstante Funktion berechnet?
2. Es sei g eine Turingberechenbare Funktion.
 $L = \{ w \mid \text{die durch } w \text{ kodierte Turingmaschine } K_w \text{ berechnet eine mit } g \text{ identische Funktion} \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob die von einer Turingmaschine berechnete Funktion mit g übereinstimmt?
3. $L_{=\emptyset} = \{ w \mid \text{für die durch } w \text{ kodierte Turingmaschine } K_w \text{ gilt } L(K_w) = \emptyset \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob die von einer Turingmaschine akzeptierte Menge leer ist?
4. $L_{\neq\emptyset} = \{ w \mid \text{für die durch } w \text{ kodierte Turingmaschine } K_w \text{ gilt } L(K_w) \neq \emptyset \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob die von einer Turingmaschine akzeptierte Menge mindestens ein Wort enthält?
5. $L = \{ w \mid \text{für die durch } w \text{ kodierte Turingmaschine } K_w \text{ gilt } : L(K_w) \text{ ist endlich} \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob die von einer Turingmaschine akzeptierte Menge endlich ist?
6. Es sei $L_0 \subseteq \{0,1\}^*$ eine entscheidbare Menge und
 $L = \{ w \mid \text{für die durch } w \text{ kodierte Turingmaschine } K_w \text{ gilt } : L(K_w) = L_0 \}$; das Entscheidungsproblem lautet: Ist es entscheidbar, ob eine Turingmaschine die Menge L_0 akzeptiert?

Satz 3.1-10:

Jedes nichttriviale Entscheidungsproblem $L \subseteq \{0,1\}^*$ über Turingmaschinen ist nicht entscheidbar.

Beweis:

Es wird gezeigt, daß entweder $L_{H_0} \leq L$ oder $L_{H_0} \leq \{0,1\}^* \setminus L$ gilt. Da L_{H_0} nicht entscheidbar ist, ist im ersten Fall L nicht entscheidbar (Satz 3.1-2 Teil 6.), im zweiten Fall ist $\{0,1\}^* \setminus L$ nicht entscheidbar und damit auch L nicht (Satz 3.1-2 Teil 3.).

Der Beweis folgt dem gleichen Schema wie die Beweise der Sätze 3.1-5 und 3.1-8:

Es sei TM_0 eine Turingmaschine mit $L(TM_0) = \emptyset$ und $w_0 = \text{code}(TM_0)$, d.h. $TM_0 = K_{w_0}$. Es werden die beiden Fälle

1. Fall: $w_0 \in L$ und

2. Fall: $w_0 \notin L$

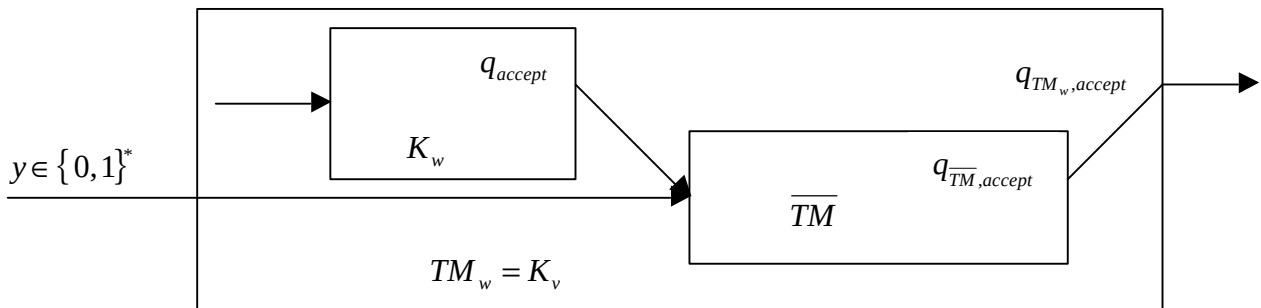
unterschieden.

Zum 1. Fall ($w_0 \in L$):

Nach obiger Bedingung (ii) gibt es eine Turingmaschine $\overline{TM}$ mit Kodierung $\overline{w} \in \{0,1\}^*$ und $\overline{w} \notin L$. Es wird die Gültigkeit von $L_{H_0} \leq \{0,1\}^* \setminus L$ gezeigt, indem eine totale und berechenbare Funktion $h: \{0,1\}^* \rightarrow \{0,1\}^*$ angegeben wird, für die $h \in L_{H_0} \Leftrightarrow h \in \{0,1\}^* \setminus L \Leftrightarrow h \notin L$ gilt.

Es sei $w \in \{0,1\}^*$ mit $VERIFIZIERE_TM(w) = \text{TRUE}$. Mit Hilfe eines deterministisch arbeitenden Turingmaschinengenerators wird aus w eine Turingmaschine TM_w konstruiert, die wie folgt arbeitet:

Bei Eingabe von $y \in \{0,1\}^*$ wird die Eingabe zunächst nicht berücksichtigt, sondern TM_w simuliert das Verhalten von K_w bei leerem Eingabeband. Falls K_w stoppt, wird y in $\overline{TM}$ eingegeben, und TM_w simuliert das Verhalten von $\overline{TM}$ auf der Eingabe y . Die Kodierung von TM_w sei v . Es gilt $v \in \{0,1\}^*$ und $VERIFIZIERE_TM(v) = \text{TRUE}$. Das Wort v ist deterministisch berechenbar, da nur Beschreibungen von Turingmaschinen zusammengestellt wurden.



Man sieht:

$$L(K_v) = L(TM_w) = \begin{cases} L(\overline{TM}) & \text{falls } w \in L_{H_0} \text{ ist} \\ \emptyset & \text{sonst} \end{cases}.$$

Die Funktion h wird definiert durch:

$$h: \begin{cases} \{0,1\}^* & \rightarrow \{0,1\}^* \\ w & \rightarrow \begin{cases} code(TM_w) & \text{falls } VERIFIZIERE_TM(w) = \text{TRUE} \\ w_0 & \text{sonst} \end{cases} \end{cases}$$

Die Funktion h ist total und berechenbar. Außerdem ist diese Funktion für den Nachweis der Relation $L_{H_0} \leq \{0,1\}^* \setminus L$ geeignet:

Ist $w \in L_{H_0}$, dann ist (nach Definition von L_{H_0}) $VERIFIZIERE_TM(w) = \text{TRUE}$ und $h(w) = \text{code}(TM_w) = v$. Außerdem gilt (siehe oben) $L(K_v) = L(\overline{TM}) = L(K_{\overline{w}})$. Nach Bedingung (iii) ist wegen $\overline{w} \notin L$ auch $v \notin L$, d.h. $h(w) \notin L$ bzw. $h(w) \in \{0,1\}^* \setminus L$.

Ist $w \notin L_{H_0}$ und ist $VERIFIZIERE_TM(w) = \text{FALSE}$, dann ist $h(w) = w_0$ und $h(w) \in L$ bzw. $h(w) \notin \{0,1\}^* \setminus L$.

Ist $w \notin L_{H_0}$ und ist $VERIFIZIERE_TM(w) = \text{TRUE}$, dann ist $h(w) = \text{code}(TM_w) = v$ und $L(K_v) = \emptyset = L(K_{w_0})$. Bedingung (iii), jetzt jedoch zusammen mit der Annahme $w_0 \in L$, impliziert $h(w) \in L$ bzw. $h(w) \notin \{0,1\}^* \setminus L$.

Zum 2. Fall ($w_0 \notin L$):

Nach obiger Bedingung (ii) gibt es eine Turingmaschine $\overline{\overline{TM}}$ mit Kodierung $\overline{\overline{w}} \in \{0,1\}^*$ und $\overline{\overline{w}} \in L$. Es wird die Gültigkeit von $L_{H_0} \leq L$ gezeigt, indem eine totale und berechenbare Funktion $h: \{0,1\}^* \rightarrow \{0,1\}^*$ angegeben wird, für die $h \in L_{H_0} \Leftrightarrow h \in L$ gilt. Der Beweis erfolgt wie im 1. Fall; dabei wird die dortige Rolle von $\overline{TM}$ von $\overline{\overline{TM}}$ übernommen. ///

Satz 3.1-10 läßt sich auch so formulieren:

Satz 3.1-11:

Ist $L \subseteq \{0,1\}^*$ eine Menge von Kodierungen von Turingmaschinen. Dann ist L nur in den Fällen entscheidbar, daß $L = \emptyset$ ist oder daß L aus der Menge der Kodierungen aller Turingmaschinen besteht.

Sämtliche oben angeführten Entscheidungsprobleme sind nicht entscheidbar, d.h. es ist beispielsweise nicht entscheidbar, ob

- eine Turingmaschine eine konstante Funktion berechnet
- eine Turingmaschine eine vorgegebene Funktion berechnet
- die von einer Turingmaschine akzeptierte Sprache leer ist
- die von einer Turingmaschine akzeptierte Sprache endlich ist
- die von einer Turingmaschine akzeptierte Sprache mit einer entscheidbaren Menge übereinstimmt.

Der letzte Punkt kann in der Praxis folgendermaßen verwendet werden. Es zeigt sich nämlich, daß Programmverifikation bzw. die Verifikation einer Spezifikation algorithmisch unmöglich ist.

Satz 3.1-12:

Für kein algorithmisch lösbares Problem (für keine entscheidbare Menge) läßt sich durch einen einzigen Algorithmus testen, ob ein entworfener Algorithmus eine korrekte Lösung des Problems liefert.

Die Bedingungen, unter denen Mengen aus Kodierungen von Turingmaschinen rekursiv aufzählbar ist, sind wesentlich komplexer:

Satz 3.1-13:

Es sei $L \subseteq \{0,1\}^*$ eine Menge von Kodierungen von Turingmaschinen. Dann ist L genau dann rekursiv aufzählbar, wenn folgende Eigenschaften (i) – (iii) gelten:

- (i) Ist $w_1 \in L$, $L_1 = L(K_{w_1})$ und ist L_2 eine rekursiv aufzählbare Menge, etwa $L_2 = L(K_{w_2})$ für ein Wort $w_2 \in \{0,1\}^*$, mit $L_1 \subseteq L_2$, dann ist auch $w_2 \in L$.
- (ii) Ist $w_1 \in L$ und $L_1 = L(K_{w_1})$ eine unendliche Menge, dann gibt es eine endliche Teilmenge $L_2 \subseteq L_1$ mit $L_2 = L(K_{w_2})$ und $w_2 \in L$.
- (iii) Es sei $E \subseteq L$ die Menge der Kodierungen, die Turingmaschinen kodieren, deren akzeptierte Sprachen endlich sind; dann ist E entscheidbar.

Beispielsweise folgt aus Satz 3.1-10, daß die Mengen

$$L_{\neq \emptyset} = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) \neq \emptyset\} \text{ und}$$

$$L_{=\emptyset} = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) = \emptyset\}$$

nicht entscheidbar sind. Es wird jedoch nichts darüber gesagt, ob sie rekursiv aufzählbar sind oder nicht. Aus Satz 3.1-13 folgt, daß $L_{=\emptyset}$ nicht rekursiv aufzählbar ist; denn Bedingung (i) ist verletzt.

In diesem Fall kann man auch anders argumentieren: Man zeigt direkt, daß die Sprache $L_{\neq \emptyset}$ rekursiv aufzählbar ist und folglich $L_{=\emptyset}$ nicht rekursiv aufzählbar ist (denn sonst wäre nach Satz 3.1-6 $L_{=\emptyset}$ entscheidbar). Dazu ist eine Turingmaschine TM anzugeben mit $L(TM) = L_{\neq \emptyset}$.

TM arbeitet wie folgt: Als Eingabe erhält TM ein Wort $w \in \{0,1\}^*$. Falls

$\text{VERIFIZIERE_TM}(w) = \text{FALSE}$ ist, wird w nicht akzeptiert. Andernfalls beginnt TM , alle Worte $z \in \{0,1,\#\}^*$ der Form $z = v\#bin(i)$ mit $v \in \{0,1\}^*$ in lexikographischer Reihenfolge zu erzeugen. Jedesmal, wenn ein derartiges Wort generiert worden ist, simuliert TM das Ver-

halten von K_w bei Eingabe von v für höchstens i Schritte. Wird v dabei akzeptiert, akzeptiert TM die Eingabe w . Andernfalls wird das nächste Wort $v'bin(i')$ erzeugt und getestet. Es gilt:

Ist $w \in L(TM)$, dann gibt es ein Wort $v \in \{0,1\}^*$, so daß K_w dieses Wort in höchstens i Schritten für ein $i \in \mathbb{N}$ akzeptiert. Daher ist $L(K_w) \neq \emptyset$ und $w \in L_{\neq \emptyset}$.

Ist umgekehrt $w \in L_{\neq \emptyset}$, d.h. $L(K_w) \neq \emptyset$, dann gibt es ein Wort $v \in L(K_w)$. Dieses Wort wird in endlich vielen, etwa j Schritten akzeptiert. Wenn TM also bei Eingabe von w das Wort $v\#bin(j)$ generiert hat, stoppt K_w nach der Simulation von j Schritten im akzeptierenden Zustand, und TM akzeptiert w , d.h. $w \in L(TM)$.

Für die in Satz 3.1-8 untersuchten Sprachen

$$L_e = \left\{ w \mid w \in \{0,1\}^*, VERIFIZIERE_TM(w) = \text{TRUE und } L(K_w) \text{ ist entscheidbar} \right\} \text{ und}$$

$$L_{ne} = \left\{ w \mid w \in \{0,1\}^*, VERIFIZIERE_TM(w) = \text{TRUE und } L(K_w) \text{ ist nicht entscheidbar} \right\}$$

folgt ebenfalls aus Satz 3.1-13, daß sie nicht rekursiv aufzählbar sind; in beiden Fällen ist wieder Bedingung (i) verletzt.

Zusammenfassung wichtiger Beispiele:

nicht entscheidbar, aber rekursiv aufzählbar	nicht rekursiv aufzählbar
$L_H = \left\{ u\#w \mid \begin{array}{l} u \in \{0,1\}^*, w \in \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{TRUE}, \\ K_w \text{ stoppt bei Eingabe von } u \end{array} \right\}$	
$L_{H_0} = \left\{ w \mid \begin{array}{l} w \in \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{TRUE}, \\ K_w \text{ stoppt bei leerem Eingabeband} \end{array} \right\}$	
$\bar{L}_d = \left\{ w \mid \begin{array}{l} w \in \{0,1\}^*, w = w_i \text{ ist das } i\text{-te Wort} \\ \text{in der lexikographischen Ordnung} \\ \text{aller Worte aus } \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{TRUE}, \\ w_i \in L(K_{w_i}) \end{array} \right\}$	$L_d = \left\{ w \mid \begin{array}{l} w \in \{0,1\}^*, w = w_i \text{ ist das } i\text{-te Wort} \\ \text{in der lexikographischen Ordnung} \\ \text{aller Worte aus } \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{FALSE} \\ \text{oder } w_i \notin L(K_{w_i}) \end{array} \right\}$
$L_{uni} = \left\{ u\#w \mid \begin{array}{l} u \in \{0,1\}^*, w \in \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{TRUE}, \\ u \in L(K_w) \end{array} \right\}$	$\bar{L}_{uni} = \left\{ u\#w \mid \begin{array}{l} u \in \{0,1\}^*, w \in \{0,1\}^*, \\ VERIFIZIERE_TM(w) = \text{TRUE}, \\ u \notin L(K_w) \end{array} \right\}$

../..

$L_{\neq \emptyset} = \left\{ w \left \begin{array}{l} w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE}, \\ L(K_w) \neq \emptyset \end{array} \right. \right\}$	$L_{=\emptyset} = \left\{ w \left \begin{array}{l} w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE}, \\ L(K_w) = \emptyset \end{array} \right. \right\}$
	$L_e = \left\{ w \left \begin{array}{l} w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE}, \\ L(K_w) \text{ ist entscheidbar} \end{array} \right. \right\}$
	$L_{ne} = \left\{ w \left \begin{array}{l} w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE}, \\ L(K_w) \text{ ist nicht entscheidbar} \end{array} \right. \right\}$
	$L_t = \left\{ w \left \begin{array}{l} w \in \{0,1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE}, \\ f_{K_w} \text{ ist eine totale Funktion} \end{array} \right. \right\}$

4 Elemente der Theorie Formaler Sprachen und der Automatentheorie

Bisher wurden Sprachen über die Akzeptanz durch Turingmaschinen definiert. Das Ergebnis ist die Klasse der rekursiv aufzählbaren Sprachen. Diese soll nun weiter strukturiert werden, indem das Modell der Turingmaschine eingeschränkt wird. Eine Einschränkung wurde bereits in Kapitel 3.1 behandelt: es werden dort Turingmaschinen betrachtet, die bei jeder Eingabe stoppen. Diese Spezialisierung führte auf die Klasse der entscheidbaren Sprachen, die eine echte Teilklasse der Klasse der rekursiv aufzählbaren Sprachen ist.

Insgesamt werden in den folgenden Unterkapiteln vier Sprachklassen beschrieben. Diese Einteilung ist nach Noam Chomsky benannt, der diese Klassen als mögliche Modelle für natürliche Sprachen charakterisiert hat. Allerdings definiert die **Chomskyhierarchie** Sprachen nicht über die Akzeptanz durch geeignete Maschinenmodelle, sondern definiert jede Klasse durch die Form von syntaktischen Regeln, nach denen Wörter der jeweiligen Sprache „erzeugt“ werden können. Für jede Sprache wird eine entsprechende (formale) Grammatik festgelegt. Je nach Art der erzeugenden Regeln ist die erzeugte Sprache eine Typ-0-Sprache, Typ-1-Sprache, Typ-2-Sprache bzw. Typ-3-Sprache. Die zugrundeliegende Theorie heißt **Theorie der Formalen Sprachen**. Zu jedem Sprachtypen der Chomskyhierarchie gibt es ein entsprechendes Berechnungsmodell (Automatentyp), das sich aus dem Modell der Turingmaschine (durch die erwähnten Einschränkungen) ableitet. Somit besteht ein enger Zusammenhang zwischen der Theorie der Formalen Sprachen und der **Automatentheorie**, die sich wesentlich mit der Definition von Modellen der Berechenbarkeit und Übersetzbarkeit beschäftigt. Beide Ansätze werden in den folgenden Unterkapiteln gegenübergestellt.

4.1 Grammatiken und formale Sprachen

Die Akzeptanz (das Erkennen) einer Sprache $L \subseteq \Sigma^*$ über einem endlichen Alphabet Σ mit Hilfe eines Berechnungsmodells wie der Turingmaschine kann als „analytischer“ Ansatz bezeichnet werden. Dem gegenüber steht ein „synthetischer“ Ansatz, der beschreibt, wie die Wörter einer Sprache mit Hilfe von Regeln erzeugt werden können. Dieser Ansatz wird in der **Theorie der formalen Sprachen** verfolgt.

Eine **Grammatik** $G = (\Sigma, N, S, R)$ wird definiert durch

1. das endliche Alphabet Σ der **Terminalsymbole**
2. das endliche Alphabet N der **Nichtterminalsymbole (Variablen)**
3. das **Startsymbol** $S \in N$
4. eine endliche Menge R von **Erzeugungsregeln (Ableitungsregeln, Produktionen)** mit

$$R \subseteq \{v \rightarrow w \mid v \in (N \cup \Sigma)^+, w \in (N \cup \Sigma)^*, |v| \geq 1\}.$$

Für $x \in (N \cup \Sigma)^*$, $y \in (N \cup \Sigma)^*$, $v \in (N \cup \Sigma)^*$, $w \in (N \cup \Sigma)^*$ heißt das Wort xwy **von G in einem Schritt aus xvy erzeugt**, wenn es eine Erzeugungsregel $v \rightarrow w$ in R gibt. Man schreibt dann: $xvy \xRightarrow{v \rightarrow w} xwy$ oder $xvy \xRightarrow[G]{} xwy$.

Wenn der Zusammenhang klar ist, werden die Super- bzw. Subskripte auch weggelassen.

Ein Wort $w \in (N \cup \Sigma)^*$ heißt **von G aus $x \in (N \cup \Sigma)^*$ erzeugt**, wenn es eine Folge von Wörtern $x_i \in (N \cup \Sigma)^*$, $i = 0, \dots, t$, gibt mit $x = x_0$, $x_i \xRightarrow[G]{} x_{i+1}$ für $i = 0, \dots, t-1$, $w = x_t$. Man schreibt dann auch $x \xRightarrow[G]^* w$. Die Menge $L(G) = \left\{ w \mid w \in \Sigma^* \text{ und } S \xRightarrow[G]^* w \right\}$ heißt **die von G erzeugte Sprache**.

Beispiele:

Grammatik für einige Sprachen über $\{a, b, c\}$

1. $G_1 = (\{a, b\}, \{S, A, B\}, S, \{S \rightarrow AB, A \rightarrow aA, A \rightarrow e, B \rightarrow bB, B \rightarrow e\})$ erzeugt die Sprache $L_1 = \{a^i b^j \mid i \in \mathbf{N}, j \in \mathbf{N}\}$. Die Sprache L_1 wird auch von der Grammatik $G'_1 = (\{a, b\}, \{S, B\}, S, \{S \rightarrow e, S \rightarrow aS, S \rightarrow bB, B \rightarrow bB, B \rightarrow e\})$ erzeugt.
2. $G_2 = (\{a, b\}, \{S\}, S, \{S \rightarrow e, S \rightarrow aSb\})$ erzeugt die Sprache $L_2 = \{a^n b^n \mid n \in \mathbf{N}\}$
3. $G_3 = (\{a, b\}, \{S, A, B\}, S, \{S \rightarrow AB, A \rightarrow e, A \rightarrow aAb, B \rightarrow e, B \rightarrow cB\})$ erzeugt die Sprache $L_3 = \{a^n b^n c^i \mid n \in \mathbf{N}, i \in \mathbf{N}\}$
4. $G_4 = (\{a, b, c\}, \{S, B\}, S, \{S \rightarrow aSbC, S \rightarrow abc, cB \rightarrow Bc, bB \rightarrow bb\})$ erzeugt die Sprache $L_4 = \{a^n b^n c^n \mid n \in \mathbf{N}, n \geq 1\}$. Die Sprache L_4 wird auch von der Grammatik $G'_4 = (\{a, b, c\}, \{S, B, C\}, S, \{S \rightarrow aSbC, S \rightarrow aBC, CB \rightarrow BC, aB \rightarrow ab, bB \rightarrow bb, bC \rightarrow bc, cC \rightarrow cc\})$ erzeugt.
5. $G_5 = (\{0, 1\}, \{S, A, B\}, S, \{S \rightarrow 1A, S \rightarrow 0B, A \rightarrow 0, A \rightarrow 0S, A \rightarrow 1AA, B \rightarrow 1, B \rightarrow 1S, B \rightarrow 0BB\})$ erzeugt die Sprache

$$L_5 = \{w \mid w \in \{0,1\}^+ \text{ und die Anzahl der Zeichen 0 in } w \text{ ist gleich der Anzahl der Zeichen 1}\}$$

Eine Grammatik $G = (\Sigma, N, S, R)$ kann ähnlich wie eine Turingmaschine (siehe Kapitel 2.4) durch eine Zeichenkette aus $\{0,1\}^*$ (algorithmisch auf deterministische Weise) kodiert werden. Die Kodierung ist dabei eine injektive Abbildung, die jeder Grammatik G ein Wort $code(G) \in \{0,1\}^*$ zuordnet. Zusätzlich kann die Codierung so entworfen werden, daß man entscheiden kann, ob ein Wort $w \in \{0,1\}^*$ die Kodierung einer Grammatik darstellt, und daß man aus der Kodierung einer Grammatik diese rekonstruieren kann. Auf Details der Darstellung und der Verfahren soll hier verzichtet werden.

Für ein Wort $w \in \{0,1\}^*$ ist

$$VERIFIZIERE_G(w) = \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Grammatik darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Grammatik darstellt} \end{cases}$$

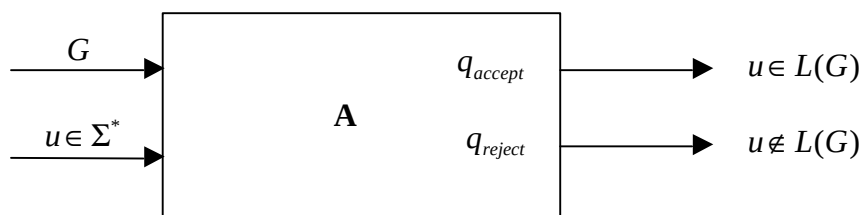
Die durch ein Wort $w \in \{0,1\}^*$ mit $VERIFIZIERE_G(w) = \text{TRUE}$ kodierte Grammatik sei KG_w .

Die **zum Wortproblem gehörende Menge** L_{Wort} wird definiert durch

$$L_{\text{Wort}} = \{u\#w \mid u \in \Sigma^*, w \in \{0,1\}^*, VERIFIZIERE_G(w) = \text{TRUE} \text{ und } u \in L(KG_w)\}.$$

Das **Wortproblem ist entscheidbar**, wenn L_{Wort} entscheidbar ist. In diesem Fall gibt es einen auf allen Eingaben der Form $u\#w$ mit $u \in \Sigma^*$ und $w \in \{0,1\}^*$ stoppenden Algorithmus, der die Eingabe genau dann akzeptiert, wenn $u\#w \in L_{\text{Wort}}$ ist.

Vereinfacht ausgedrückt ist **das Wortproblem entscheidbar**, wenn es einen auf jeder Eingabe stoppenden Algorithmus **A** gibt, der eine aus zwei Teilen bestehende Eingabe, nämlich bestehend aus (der Kodierung) einer Grammatik $G = (\Sigma, N, S, R)$ und einer Zeichenkette $u \in \Sigma^*$, erhält und die Eingabe genau dann akzeptiert, wenn $u \in L(G)$ gilt, d.h. wenn sich u aus dem Startsymbol von G durch Anwendung der in G definierten Ableitungsregeln herleiten läßt.

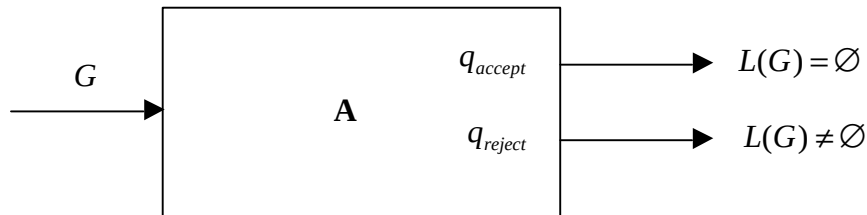


Entsprechend kann man das Leerheitsproblem definieren:

Die **zum Leerheitsproblem gehörende Menge** L_{leer} wird definiert durch

$$L_{\text{leer}} = \left\{ w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_}G(w) = \text{TRUE und } L(KG_w) = \emptyset \right\}.$$

Das **Leerheitsproblem ist entscheidbar** (hier in der vereinfachten Darstellung), wenn es einen auf jeder Eingabe stoppenden Algorithmus **A** gibt, der als Eingabe eine Grammatik $G = (\Sigma, N, S, R)$ erhält und die Eingabe genau dann akzeptiert, wenn $L(G) = \emptyset$ gilt.



4.2 Typ-0-Sprachen

Eine Sprache, die von einer Grammatik $G = (\Sigma, N, S, R)$ erzeugt wird, für die

$$R \subseteq \left\{ v \rightarrow w \mid v \in (N \cup \Sigma)^+ \setminus \Sigma^*, w \in (N \cup \Sigma)^* \right\}$$

gilt, heißt **Typ-0-Sprache**. Die Regeln $v \rightarrow w$ einer Typ-0-Sprache erfüllen also die Bedingung $|v| \geq 1$, und die linke Seite v einer Regel enthält mindestens ein nichtterminales Symbol; außerdem kann in einer Ableitung durch Anwendung einer Regel keine einmal entstandene rein terminale Zeichenkette mehr ersetzt werden.

In einem Ableitungsschritt $xvy \xRightarrow{v \rightarrow w} xwy$ kann es aber durchaus vorkommen, daß das Ergebnis xwy des Ableitungsschritts kürzer als die Ausgangszeichenkette ist.

Man kann zeigen, daß jede Typ-0-Sprache mit einem Alphabet Σ von einer Turingmaschine mit Eingabealphabet Σ akzeptiert werden kann. Dazu wird der Erzeugungsvorgang einer Grammatik mit Hilfe einer nichtdeterministischen Turingmaschine simuliert. Umgekehrt läßt sich jede rekursiv aufzählbare Menge durch eine Typ-0-Grammatik erzeugen. Daher gilt:

Satz 4.2-1:

Die Klasse der rekursiv aufzählbaren Mengen (von Turingmaschinen akzeptierte Mengen) über einem endlichen Alphabet Σ ist mit der Klasse der Typ-0-Sprachen mit Alphabet Σ identisch.

In Kapitel 3.1 wurde gezeigt, daß die Menge

$$L_{uni} = \{u\#w \mid u \in \{0,1\}^*, w \in \{0,1\}^*, VERIFIZIERE_TM(w) = \text{TRUE} \text{ und } u \in L(K_w)\}$$

nicht entscheidbar ist. Satz 4.2-1 legt nahe, in der Definition von L_{uni} das Prädikat $VERIFIZIERE_TM$ durch $VERIFIZIERE_G$ zu ersetzen und dieses dann durch das leicht modifizierte Prädikat

$$VERIFIZIERE_G_TYP_0(w)$$

$$= \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Typ-0-Grammatik darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Typ-0-Grammatik darstellt} \end{cases}$$

Man erhält dann das **Wortproblem für Typ-0-Grammatiken**, das danach fragt, ob die Menge

$$L_{Wort_Typ-0} = \{u\#w \mid u \in \Sigma^*, w \in \{0,1\}^*, VERIFIZIERE_G_TYP_0(w) = \text{TRUE} \text{ und } u \in L(KG_w)\}$$

entscheidbar ist. Diese Frage muß verneint werden.

Ähnlich verhält es sich mit dem **Leerheitsproblem für Typ-0-Grammatiken**. Die Menge

$$L_{leer_Typ-0} = \{w \mid w \in \{0,1\}^*, VERIFIZIERE_G_TYP_0(w) = \text{TRUE} \text{ und } L(KG_w) = \emptyset\}$$

ist nicht entscheidbar, da (vgl. Kapitel 3.1) die Menge

$$L_{=\emptyset} = \{w \mid w \in \{0,1\}^*, VERIFIZIERE_TM(w) = \text{TRUE} \text{ und } L(K_w) = \emptyset\}$$

nicht rekursiv aufzählbar und damit auch nicht entscheidbar ist.

Insgesamt ergibt sich

Satz 4.2-2:

Das Wortproblem und das Leerheitsproblem für Typ-0-Grammatiken sind nicht entscheidbar:

Es gibt keinen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-0-Grammatik $G = (\Sigma, N, S, R)$ und eines Wortes $u \in \Sigma^*$ entscheidet, ob $u \in L(G)$ ist.

Es gibt keinen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-0-Grammatik $G = (\Sigma, N, S, R)$ entscheidet, ob $L(G) = \emptyset$ gilt.

4.3 Typ-1-Sprachen

Es sei $G = (\Sigma, N, S, R)$ eine Grammatik, in der die Erzeugungsregeln

$R \subseteq \{v \rightarrow w \mid v \in (N \cup \Sigma)^+ \setminus \Sigma^*, w \in (N \cup \Sigma)^*\}$ folgender zusätzlicher Einschränkung unterliegen:

Für alle Regeln $v \rightarrow w$ gilt $|v| \leq |w|$. Als einzige Ausnahme ist die Regel $S \rightarrow e$ zugelassen; wenn diese Regel vorkommt, darf S auf keiner rechten Seite einer Regel vorkommen.

Wie bei einer Typ-0-Grammatik dürfen auf der linken Seite einer Regel in einer Typ-1-Grammatik also sowohl terminale als auch nichtterminale Zeichen vorkommen, wobei links mindestens ein nichtterminales Zeichen steht. Die Länge der rechten Seite einer Regel ist bis auf die Ausnahme $S \rightarrow e$ mindestens so groß wie die Länge der linken Seite. Daher wird in

einem Ableitungsschritt $xvy \xRightarrow{v \rightarrow w} xwy$ die Länge des Ableitungsergebnisses nicht kürzer. In einer Ableitung $S \xRightarrow{*} w$ eines Wortes $w \in \Sigma^+$, etwa $S \Rightarrow x_1 \Rightarrow \dots \Rightarrow x_i \Rightarrow x_{i+1} \Rightarrow \dots \Rightarrow w$, gilt für alle Zwischenschritte $1 = |S| \leq |x_1| \leq \dots \leq |x_i| \leq |x_{i+1}| \leq \dots \leq |w|$. Diese Beobachtung wird wichtig, wenn man den Turingmaschinentyp charakterisieren möchte, der eine von einer Typ-1-Grammatik erzeugten Sprache erkennt.

Enthält R die Regel $S \rightarrow e$, dann ist $e \in L(G)$, und die Anwendung der Regel $S \rightarrow e$ stellt die einzige Möglichkeit dar, um e aus S abzuleiten.

Die von einer derartigen Grammatik erzeugte Sprache heißt **Typ-1-Sprache** oder **kontext-sensitive Sprache**. Die Grammatik heißt **kontextsensitive Grammatik**. Beispielsweise ist die Sprache $L_4 = \{a^n b^n c^n \mid n \in \mathbf{N}, n \geq 1\}$ aus Kapitel 4.1 eine kontextsensitive Sprache (Typ-1-Sprache).

Jede Typ-1-Sprache ist natürlich auch eine Typ-0-Sprache.

Es stellt sich die Frage, ob es Typ-0-Sprachen gibt, die nicht kontextsensitiv sind. Zur Beantwortung dieser Frage wird folgender Ansatz gewählt: Es wird ein Berechnungsmodell vorgestellt, das sich aus dem Modell der Turingmaschine ableitet und genau kontextsensitive Sprachen akzeptiert, und dann untersucht, wie sich dieses Berechnungsmodell zum Modell der Turingmaschine verhält:

Ein **linear beschränkter Automat LBA** ist eine nichtdeterministische Turingmaschine, deren Schreib/Leseköpfe auf allen Bändern bei Eingabe eines Wortes w mit $|w| = n$ jeweils nicht mehr als n Zellen auf den Bändern verwenden. Insbesondere bewegt sich kein Kopf weiter nach rechts als bis zur Position $n+1$ (an dem Blank auf dem Eingabeband bei Position $n+1$ wird das Ende des Eingabeworts erkannt). Die von einem LBA akzeptierte Menge $L(LBA)$ wird wie bei Turingmaschinen definiert.

Man kann zeigen, daß es zu jeder von einem linear beschränkten Automaten LBA akzeptierten Sprache $L = L(LBA)$ eine kontextsensitive Grammatik G_{LBA} gibt mit $L(G_{LBA}) = L(LBA)$. Umgekehrt gibt es zu jeder von einer kontextsensitiven Grammatik G erzeugten Sprache $L' = L(G)$ einen linear beschränkten Automaten LBA_G mit $L(LBA_G) = L(G)$. In diese Überlegungen geht wesentlich ein, daß die Produktionen $v \rightarrow w$ der kontextsensitiven Grammatik der Bedingung Regeln $|v| \leq |w|$ genügen, so daß zur Akzeptanz auf allen Bändern jeweils ein durch die Länge des Eingabeworts beschränkter Speicherplatz ausreicht.

Satz 4.3-1:

Die Klasse der von linear beschränkten Automaten akzeptierten Sprachen über einem endlichen Alphabet Σ ist mit der Klasse der kontextsensitiven Sprachen (Typ-1-Sprachen) über Σ identisch.

Möchte man daher Aussagen über kontextsensitive Sprachen beweisen, so kann man dazu mit kontextsensitiven Grammatiken oder linear beschränkten Automaten argumentieren. Satz 4.3-1 läßt vermuten, daß die Klasse der kontextsensitiven Sprachen über Σ eine echte Teilklasse der rekursiv aufzählbaren Sprachen über Σ ist. Der folgende Satz bestätigt diese Vermutung.

Satz 4.3-2:

Die von einem LBA akzeptierte Menge $L(LBA)$ über einem Alphabet Σ ist entscheidbar.

Jede kontextsensitive Sprache über einem Alphabet Σ ist entscheidbar.

Die Entscheidung kann bei einem Wort $w \in \Sigma^*$ mit Länge n in $O(2^{O(n)})$ vielen Schritten getroffen werden; das Entscheidungsverfahren hat also exponentielle Laufzeit.

Beweis:

Es sei $G = (\Sigma, N, S, R)$ eine kontextsensitive Grammatik. Zu zeigen ist: $L(G)$ ist entscheidbar. Dazu wird eine auf allen Eingaben $w \in \Sigma^*$ stoppende Turingmaschine TM_G angegeben, die w genau dann akzeptiert, wenn $w \in L(G)$ ist. Die Arbeitsweise von TM_G wird hier informell in Form eines Algorithmus beschrieben:

Bei Eingabe von $w \in \Sigma^*$ mit $|w| = n$ erzeugt TM_G einen Graphen, dessen Knoten mit den Zeichenketten in $(N \cup \Sigma)^*$ markiert sind, die eine Länge besitzen, die kleiner oder gleich n ist

(andere Zeichenketten kommen in einer Ableitung $S \Rightarrow^* w$ nicht vor). Ist dabei der Knoten K_i mit der Zeichenkette $a \in (N \cup \Sigma)^*$ und der Knoten K_j mit $b \in (N \cup \Sigma)^*$ markiert und kann man in G durch Anwendung einer Regel b aus a in einem Ableitungsschritt herleiten, d.h. $a \Rightarrow_G b$, dann wird eine Kante von K_i nach K_j eingefügt. Ein Knoten ist mit S markiert und einer mit w . Es ist $w \in L(G)$ genau dann, wenn es einen Pfad von dem mit S markierten Knoten zu dem mit w markierten Knoten gibt. Um diese Tatsache festzustellen, kann man einen der bekannten Algorithmen zum Auffinden von Pfaden in Graphen zwischen definierten Knoten anwenden. Da der beschriebene Graph $O(c^n)$ viele Knoten (mit einer Konstanten $c = c(G)$) besitzt und sich die Laufzeit des Pfadsuchalgorithmus durch eine Funktion der Ordnung $O(m^3)$ beschränken läßt, wobei $m \in O(c^n)$ die Anzahl der Knoten im Graphen angibt, ist das gesamte Entscheidungsverfahren von der Ordnung $O(2^{O(n)})$. ///

Die Klasse der kontextsensitiven Sprachen ist eine echte Teilklasse der Klasse der entscheidbaren Sprachen, wie folgende Sätze zeigen.

Mit der Diagonalisierungstechnik läßt sich zunächst folgender (technischer) Hilfssatz zeigen:

Satz 4.3-3:

Es sei

$$L_0 = \left\{ w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } K_w \text{ stoppt auf jeder Eingabe} \right\}$$

eine rekursiv aufzählbare Menge. Dann gibt es eine entscheidbare Sprache L , die von einer auf allen Eingaben stoppenden Turingmaschine TM erkannt wird, deren Kodierung in L_0 nicht vorkommt.

Bemerkung: Da die Menge

$$L_e = \left\{ w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) \text{ ist entscheidbar} \right\}$$

aus Kapitel 3.1 nicht rekursiv aufzählbar ist, gilt $L_0 \neq L_e$.

Beweis:

Da L_0 als rekursiv aufzählbar angenommen wird, gibt es nach Satz 3-2 eine totale berechenbare Funktion $h: \{0,1\}^* \rightarrow \Sigma^*$ mit $L_0 = h(\{0,1\}^*)$.

Es wird $L = \left\{ w \mid w \in \{0,1\}^* \text{ und } w \notin L(K_{h(w)}) \right\}$ gesetzt.

Dann ist L entscheidbar: Eine auf allen Eingaben $w \in \{0,1\}^*$ stoppende Turingmaschine TM mit $L(TM) = L$, d.h. die entscheidet, ob $w \in L$ gilt oder nicht, arbeitet wie folgt: TM berech-

net zunächst $h(w)$. Dabei ist $h(w) \in L_0$, insbesondere $VERIFIZIERE_TM(h(w)) = \text{TRUE}$. Jetzt simuliert TM das Verhalten von $K_{h(w)}$ bei Eingabe von w . Da $K_{h(w)}$ nach Definition von L_0 auf allen Eingaben stoppt, kann TM feststellen, ob $w \in L(K_{h(w)})$ gilt oder nicht. Das Wort w wird von TM genau dann akzeptiert, wenn bei dieser Simulation $w \in L(K_{h(w)})$ festgestellt wird.

TM habe die Kodierung w_L . Falls $w_L \in L_0$ gilt, dann sei $w \in \{0,1\}^*$ so gewählt, daß $h(w) = w_L$ ist. Dann folgt (nach Definition von L) der Widerspruch

$$\begin{aligned} w \in L &\Leftrightarrow w \notin L(K_{h(w)}) \quad (\text{nach Definition von } L) \\ &\Leftrightarrow w \notin L(K_{w_L}) \quad (\text{wegen } h(w) = w_L) \\ &\Leftrightarrow w \notin L(TM) \quad (\text{da } w_L \text{ die Kodierung von } TM \text{ ist}) \\ &\Leftrightarrow w \notin L \quad (\text{wegen } L(TM) = L). \end{aligned}$$

Daher kommt die Kodierung w_L von TM in L_0 nicht vor. ///

Satz 4.3-3 kann genutzt werden, um zu zeigen, daß es eine entscheidbare Menge gibt, die nicht kontextsensitiv ist:

Jede kontextsensitive Grammatik kann ähnlich wie eine Turingmaschine (siehe Kapitel 2.4 und Kapitel 4) durch eine Zeichenkette aus $\{0,1\}^*$ (algorithmisch auf deterministische Weise) kodiert werden. Wie in Kapitel 2.4 auf der Menge der Turingmaschinen kann man dann auf der Menge der kontextsensitiven Grammatiken auf Basis ihrer Kodierungen eine lineare Ordnung definieren. Man kann also von der i -ten kontextsensitiven Grammatik G_i sprechen. Nach Satz 4.3-2 ist $L(G_i)$ entscheidbar mit einer (dort angegebenen) Turingmaschine TM_{G_i} ; die Turingmaschine TM_{G_i} stoppt (nach Konstruktion) auf jeder Eingabe. Deren Kodierung sei (gemäß dem Vorgehen aus Kapitel 2.4) $code(TM_{G_i})$. Die Berechnung des Werts $code(TM_{G_i})$ bei Vorgabe von i bzw. von $bin(i)$ erfolgt insgesamt auf deterministische Weise und definiert eine Funktion $h: \{0,1\}^* \rightarrow \{0,1\}^*$:

$$(i \leftrightarrow) bin(i) \rightarrow i\text{-te kontextsensitive Grammatik } G_i \rightarrow TM_{G_i} \rightarrow code(TM_{G_i}).$$

Diese Abbildung ist injektiv, wie man leicht nachprüfen kann. Durch h wird jeder Zeichenkette $u = bin(i)$ die Kodierung einer Turingmaschine zugeordnet, die auf jeder Eingabe stoppt. Setzt man $L_0 = h(\{0,1\}^*)$, so sind die Voraussetzungen in Satz 4.3-3 erfüllt. Daher gilt der folgende Satz.

Satz 4.3-4:

Die Klasse der kontextsensitiven Sprachen (Typ-1-Sprachen) über einem endlichen Alphabet Σ ist eine echte Untermenge der entscheidbaren Sprachen über Σ und damit auch der Typ-0-Sprachen über Σ .

Wie im allgemeinen Fall der Turingmaschine unterscheidet man auch bei linear beschränkten Automaten nichtdeterministisches und deterministisches Verhalten. Ein **nichtdeterministischer linear beschränkter Automat** liegt vor, wenn die Überföhrungsfunktion die Form $d : Q \times \Sigma^k \rightarrow \mathbf{P}(Q \times (\Sigma \times \{L, R, S\})^k)$ hat, ein **deterministischer linear beschränkter Automat** liegt vor, wenn die Überföhrungsfunktion die Form $d : Q \times \Sigma^k \rightarrow Q \times (\Sigma \times \{L, R, S\})^k$ hat. Im deterministischen Fall ist die Folgekonfiguration (falls sie überhaupt existiert) einer Konfiguration eindeutig bestimmt; zum Nichtdeterminismus vgl. Kapitel 2.5. Es ist nicht bekannt, ob jede von einem *nichtdeterministischen* linear beschränkten Automaten auch von einem *deterministischen* linear beschränkten Automaten akzeptiert wird. Dieses offene Problem heißt **LBA-Problem**. Zu beachten ist dabei folgendes: Ist L eine kontextsensitive Sprache, dann gibt es einen nichtdeterministischen linear beschränkten Automaten LBA mit $L = L(LBA)$. Die Akzeptanz eines Wortes w mit $|w| = n$ benötigt eine Anzahl von Zellen, die durch die lineare Funktion $S(n) = n + 1$ gegeben ist. Das nichtdeterministische Verhalten einer Turingmaschine, die durch eine Funktion der Ordnung $O(f(n))$ platzbeschränkt ist, kann deterministisch simuliert werden, wobei dabei der Speicherplatzbedarf die Ordnung $O(f^2(n))$ annimmt. Da ein linear beschränkter Automat auch eine nichtdeterministische Turingmaschine ist, könnte man versuchen, diese deterministische Simulation hier zu verwenden. Diese führt jedoch aus der Klasse der linear beschränkten Automaten heraus, da sie quadratischen Speicherplatz der Ordnung $O(S^2(n)) = O(n^2)$ benötigt. Daher trägt die deterministische Simulation einer nichtdeterministischen Turingmaschine in dieser Allgemeinheit zur Lösung des LBA-Problems nichts bei.

Die Klasse der von deterministischen linear beschränkten Automaten akzeptierten Sprachen ist abgeschlossen gegen Komplementbildung (das zeigt man wie in Satz 3.1-2 Teil 3.). Man hat lange Zeit vermutet, daß diese Abschlußeigenschaft für die Klasse der von nichtdeterministischen linear beschränkten Automaten akzeptierten Sprachen nicht zutrifft. Die Richtigkeit dieser Vermutung hätte die (negative) Lösung des LBA-Problems nach sich gezogen. Inzwischen weiß man jedoch, daß auch die Klasse der von nichtdeterministischen linear beschränkten Automaten akzeptierten Sprachen abgeschlossen gegen Komplementbildung ist. Das LBA-Problem ist weiterhin ungelöst.

Die Definition des Wortproblems für Typ-0-Grammatiken kann man auf Typ-1-Grammatiken übertragen: Dazu wird das Prädikat

$VERIFIZIERE_G_TYP_1(w)$

$$= \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Typ-1-Grammatik darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Typ-1-Grammatik darstellt} \end{cases}$$

definiert. Dieses Prädikat ist berechenbar. Dazu ist unter anderem zu prüfen, ob die Erzeugungsregeln in KG_w den Bedingungen einer kontextsensitiven Grammatik genügen.

Das **Wortproblem für Typ-1-Grammatiken** fragt danach, ob die Menge

$$L_{Wort_Typ-1} = \left\{ u\#w \mid u \in \Sigma^*, w \in \{0,1\}^*, VERIFIZIERE_G_TYP_1(w) = \text{TRUE} \text{ und } u \in L(KG_w) \right\}$$

entscheidbar ist.

Aus dem Beweis von Satz 4.3-2 folgt, daß das Wortproblem für kontextsensitive Grammatiken entscheidbar ist. Das **Leerheitsproblem für Typ-1-Grammatiken**, nämlich die Frage, ob die Menge

$$L_{leer_Typ-1} = \left\{ w \mid w \in \{0,1\}^*, VERIFIZIERE_G_TYP_1(w) = \text{TRUE} \text{ und } L(KG_w) = \emptyset \right\}$$

entscheidbar ist, muß wie bei Typ-0-Grammatiken verneint werden (der Beweis findet sich in der angegebenen Literatur).

Zusammenfassend ergibt sich

Satz 4.3-5:

Das Wortproblem für Typ-1-Grammatiken ist entscheidbar; das Leerheitsproblem für Typ-1-Grammatiken ist nicht entscheidbar:

Es gibt einen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-1-Grammatik $G = (\Sigma, N, S, R)$ und eines Wortes $u \in \Sigma^*$ entscheidet, ob $u \in L(G)$ ist.

Es gibt keinen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-1-Grammatik $G = (\Sigma, N, S, R)$ entscheidet, ob $L(G) = \emptyset$ gilt.

4.4 Typ-2-Sprachen

Es sei $G = (\Sigma, N, S, R)$ eine Grammatik, in der für die Erzeugungsregeln $R \subseteq \{A \rightarrow w \mid A \in N \text{ und } w \in (N \cup \Sigma)^*\}$ gilt. Auf der linken Seite einer Regel steht immer genau ein nichtterminales Symbol, rechts kann auch die leere Zeichenkette vorkommen.

Die von einer derartigen Grammatik erzeugte Sprache heißt **Typ-2-Sprache** oder **kontextfreie Sprache**. Die zugehörige Grammatik heißt **kontextfreie Grammatik**. Beispielsweise sind die Sprache $L_2 = \{a^n b^n \mid n \in \mathbf{N}\}$, $L_3 = \{a^n b^n c^i \mid n \in \mathbf{N}, i \in \mathbf{N}\}$ und

$L_5 = \{w \mid w \in \{0,1\}^+ \text{ und die Anzahl der Zeichen 0 in } w \text{ ist gleich der Anzahl der Zeichen 1}\}$ aus Kapitel 4.1 kontextfreie Sprachen (Typ-2-Sprachen).

Die kontextfreien Sprachen zählen zu den am intensivsten erforschten Sprachen. Einen über-
ragenden Erfolg haben sie in der Anwendung und der Theorie des Compilerbaus. Program-
miersprachen wie Pascal und ihre Nachfolger sind erst nach intensiver Erforschung der kon-
textfreien Sprachen und unter Anwendung dieser Theorie entwickelt worden. Die Syntax der
meisten heute üblichen Programmiersprachen wird in Form einer Grammatik definiert, die
weitestgehend kontextfrei ist. Man kann jedoch formal zeigen, daß eine Programmiersprache,
in der Variablen mit Datentypen deklariert werden, so daß Typverträglichkeit verlangt wird,
in der die Anzahl von Formal- und Aktualparametern in Prozeduren übereinstimmen müssen,
in der ausschließlich die Verwendung vorher deklarerter Objekte zulässig ist usw., nicht
komplett durch eine kontextfreie Grammatik beschrieben werden kann.

Beispiel:

Ausschnitt aus der Sprachdefinition der Programmiersprache Object Pascal

Die Syntax der Sprache Object Pascal wird wie bei vielen anderen Programmiersprachen
(weitgehend) durch eine kontextfreie Grammatik definiert. Nichtterminale Symbole werden
dabei durch Bezeichner angegeben, die mit einem Großbuchstaben beginnen und weitere
Kleinbuchstaben enthalten. Bezeichner, die nur Großbuchstaben enthalten, stehen für jeweils
ein einziges terminales Symbol. So bezeichnet im folgenden Ausschnitt der Bezeichner `ziel`
das Startsymbol der Sprache, während der Bezeichner `UNIT` für ein einziges terminales Sym-
bol steht.

Eine Regelangabe der Form $A \rightarrow a_1 | a_2 | \dots | a_n$ steht für die n Regeln $A \rightarrow a_1$, $A \rightarrow a_2$, ...,
 $A \rightarrow a_n$. Ein Angabe in eckigen Klammern innerhalb einer Regel bezeichnet einen optionalen
Teil, d.h. $A \rightarrow a[b]g$ steht für die Regeln $A \rightarrow abg$ und $A \rightarrow ag$.

Ziel -> (Programm | Package | Bibliothek | Unit)

Programm -> [PROGRAM Bezeichner ['('Bezeichnerliste')'] ';']

Programmblock '.'

Unit -> UNIT Bezeichner ';'

interface-Abschnitt

implementation-Abschnitt

initialization-Abschnitt '.'

Package -> PACKAGE Bezeichner ';'

[requires-Klausel]

[contains-Klausel]

END '.'

Bibliothek -> LIBRARY Bezeichner ';'

Programmblock '.'

Programmblock -> [uses-Klausel]

Block

uses-Klausel -> USES Bezeichnerliste ';'

interface-Abschnitt -> INTERFACE

[uses-Klausel]

[interface-Deklaration]...

interface-Deklaration -> const-Abschnitt

-> type-Abschnitt

-> var-Abschnitt

-> exported-Kopf

exported-Kopf -> Prozedurkopf ';' [Direktive]

-> Funktionskopf ';' [Direktive]

implementation-Abschnitt -> IMPLEMENTATION

[uses-Klausel]

[Deklarationsabschnitt]...

Block -> [Deklarationsabschnitt]

Verbundanweisung

Deklarationsabschnitt -> Label-Deklarationsabschnitt

-> const-Abschnitt

-> type-Abschnitt

-> var-Abschnitt

-> Prozedurdeklarationsabschnitt

...

Einfache Anweisung -> Designator ['(' Ausdrucksliste ')']

-> Designator ':'=' Ausdruck

-> INHERITED

-> GOTO Label-Bezeichner

Strukturierte Anweisung -> Verbundanweisung

-> Bedingte Anweisung

-> Schleifenanweisung

-> with-Anweisung

Verbundanweisung -> BEGIN Anweisungsliste END

Bedingte Anweisung -> if-Anweisung

-> case-Anweisung

```
if-Anweisung -> IF Ausdruck THEN Ausdruck [ELSE Ausdruck]
case-Anweisung -> CASE Ausdruck OF case-Selektor/';'... [ELSE Ausdruck] [';'] END
case-Selektor -> case-Label/','... ':' Anweisung
case-Label -> Konstanter Ausdruck ['..' Konstanter Ausdruck]
Schleifenanweisung -> repeat-Anweisung
```

-> while-Anweisung
-> for-Anweisung

```
repeat-Anweisung -> REPEAT Anweisung UNTIL Ausdruck
while-Anweisung -> WHILE Ausdruck DO Anweisung
for-Anweisung -> FOR Qualifizierter Bezeichner ':'= 'Ausdruck (TO | DOWNT0)
    Ausdruck DO Anweisung
with-Anweisung -> WITH Bezeichnerliste DO Anweisung
Prozedurdeklarationsabschnitt -> Prozedurdeklaration
```

-> Funktionsdeklaration

...

Klassentyp -> CLASS [Klassenvererbung]

```
[Klassenfelderliste]
[Klassenmethodenliste]
[Klasseneigenschaftenliste]
END
```

```
Klassenvererbung -> '(' Bezeichnerliste ')'
Klassensichtbarkeit -> [PUBLIC | PROTECTED | PRIVATE | PUBLISHED]
Klassenfelderliste -> (Klassensichtbarkeit Objektfelderliste)/';'...
Klassenmethodenliste -> (Klassensichtbarkeit Methodenliste)/';'...
Klasseneigenschaftenliste -> (Klassensichtbarkeit Eigenschaftenliste ';')...
Eigenschaftenliste -> PROPERTY Bezeichner [Eigenschaftsschnittstelle]
    Eigenschaftsbezeichner
```

```
Eigenschaftsschnittstelle -> [Eigenschaftsparameterliste] ':' Bezeichner
Eigenschaftsparameterliste -> '[' (Bezeichnerliste ':' Typbezeichner)/';'... ']'
Eigenschaftsbezeichner -> [INDEX Konstanter Ausdruck]
```

```
[READ Ident]
[WRITE Bezeichner]
[STORED Bezeichner | Konstante]]
[(DEFAULT Konstanter Ausdruck) | NODEFAULT]
[IMPLEMENTS Typbezeichner]
```

Schnittstellentyp -> INTERFACE [Schnittstellenvererbung]

```
[Klassenmethodenliste]
[Klasseneigenschaftenliste]
END
```

```
Schnittstellenvererbung -> '(' Bezeichnerliste ')'
requires-Klausel -> REQUIRES Bezeichnerliste... ';'
contains-Klausel -> CONTAINS Bezeichnerliste... ';'
Bezeichnerliste -> Bezeichner/','...
```

```

Qualifizierter Bezeichner -> [Unit-Bezeichner '.' ] Bezeichner
Typbezeichner -> [Unit-Bezeichner '.' ] <Typbezeichner>
Ident -> <Bezeichner>
ConstExpr -> <Konstanter Ausdruck>
UnitId -> <Unit-Bezeichner>
LabelId -> <Label-Bezeichner>

Number -> <Nummer>
String -> <String>

```

Die Erzeugungsregeln einer kontextfreien Grammatik erfüllen (auf den ersten Blick) auch die Bedingung, die man an eine kontextsensitive Grammatik stellt. In einer kontextfreien Grammatik können jedoch auch Produktionen der Form $A \rightarrow e$ vorkommen, wobei $A \neq S$ oder zusätzlich A auf der rechten Seite einer Produktion vorkommt. In diesem Fall kann man in einem endlichen Verfahren die Regeln und nichtterminalen Symbole der Grammatik so abändern, daß keine Produktionen der Form $A \rightarrow e$ mehr vorkommen außer wenn eventuell A das Startsymbol ist (A steht dann auf keiner rechten Seite einer Produktion). Die Änderung kann so erfolgen, daß weiterhin dieselbe Sprache erzeugt wird. Es gilt daher:

Satz 4.4-1:

Zu jeder kontextfreien Grammatik G gibt es eine kontextsensitive Grammatik G' , die dieselbe Sprache erzeugt:
 Die Klasse der kontextfreien Sprachen (Typ-1-Sprachen) über einem endlichen Alphabet Σ ist eine echte Teilmenge der kontextsensitiven Sprachen über Σ .

Die Tatsache, daß die Klasse der kontextfreien Sprachen über einem endlichen Alphabet Σ echt in der Klasse der kontextsensitiven Sprachen über Σ enthalten ist, wird weiter unten bewiesen.

Der folgende Satz, der auch hier ohne Beweis aufgeführt wird, besagt, daß man jede kontextfreie Grammatik in eine Grammatik in **Normalform** überführen kann, deren Regeln eine einfache Struktur aufweisen.

Satz 4.4-2:

Jede kontextfreie Grammatik kann so umgeformt werden, daß alle Regeln die Form $S \rightarrow e$ oder $A \rightarrow BC$ mit $A \in N$, $B \in N$, $C \in N$ oder $A \rightarrow a$ mit $A \in N$, $a \in \Sigma$ haben (Chomsky-Normalform) und dabei dieselbe Sprache erzeugt wird.

Der folgende Satz (**uvwxy-Theorem, pumping lemma**) liefert ein Beweismittel, mit dessen Hilfe man zeigen kann, daß eine Sprache nicht kontextfrei ist.

Satz 4.4-3:

Zu jeder kontextfreien Sprache L gibt es eine Zahl $n_0 > 0$ mit der Eigenschaft:

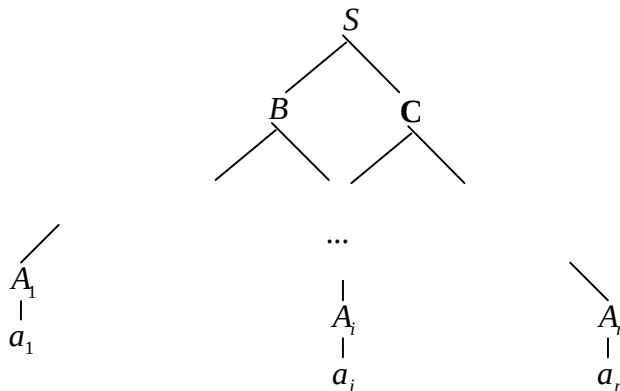
jedes $z \in L$ mit $|z| \geq n_0$ lässt sich zerlegen in $z = uvwxy$ mit

- (i) $|vwx| \leq n_0$
- (ii) $|vx| > 0$
- (iii) $uv^kwx^ky \in L$ für jedes $k \in \mathbb{N}$.

Beweis:

Es sei $G = (\Sigma, N, S, R)$ eine kontextfreie Grammatik in Chomsky-Normalform mit $L = L(G)$ und $n_0 = 2^{|N|+1}$.

Für ein Wort $z \in L$, etwa $z = a_1 \dots a_n$ mit $a_1 \in \Sigma, \dots, a_n \in \Sigma$ und $|z| = n \geq n_0$, hat eine Ableitung aus S die Form $S \Rightarrow BC \Rightarrow \dots \Rightarrow a_1 \dots a_n$. Diese Ableitung kann als Ableitungsbaum geschrieben werden, in dem jeder Knoten durch ein in der Ableitung vorkommendes Symbol markiert ist: Die Wurzel des Ableitungsbaums ist mit S markiert. Wird in der Ableitung eine Regel der Form $A \rightarrow BC$ mit $B \in N$ und $C \in N$ angewendet, so gibt es im Ableitungsbaum einen dieser Regelanwendung entsprechenden mit A markierten Knoten, dessen direkte Nachfolger mit B bzw. C markiert sind. Wird eine Regel der Form $A \rightarrow a$ mit $a \in \Sigma$ angewendet, so gibt es im Ableitungsbaum einen dieser Regelanwendung entsprechenden mit A markierten Knoten, dessen direkter Nachfolger mit a markiert ist. Da G Chomsky-Normalform hat, ist dieser Ableitungsbaum ein Binärbaum, in dem jeder innere Knoten genau zwei Nachfolger hat und der Anwendung einer Regel der Form $A \rightarrow BC$ entspricht. Die Blätter des Baums sind mit $a_1, \dots, a_n$ markiert. Ein mit a_i markiertes Blatt mit seinem mit A_i markierten Vorgänger entspricht der Anwendung der Regel $A_i \rightarrow a_i$. Nur an den Blättern werden Regeln dieser Form angewendet.



Für die Anzahl n der Blätter dieses Ableitungsbaums gilt $n = |z| \geq n_0 = 2^{|N|+1}$. Dann hat der Baum nach Satz 1.1-9 Teil 4. eine Mindesthöhe von $\lceil \log_2(n) + 1 \rceil \geq \lceil \log_2(2^{|N|+1}) + 1 \rceil = |N| + 2$.

Es gibt also einen Pfad von der Wurzel zu einem Blatt, das mit einem terminalen Zeichen a_i markiert ist, auf dem mindestens $|N|+1$ viele innere Knoten liegen, die mit nichtterminalen Symbol markiert sind. Da die Grammatik nur $|N|$ viele nichtterminale Symbole enthält, kommt auf diesem Pfad ein $A \in N$ mindestens zweimal vor. Es werde hier die Situation betrachtet, bei der dieses zum ersten Mal geschieht:

$$S \Rightarrow^* w_1 A w_2 \Rightarrow^+ w_1 w_3 A w_4 w_2 \Rightarrow^* uvwxy \quad \text{mit} \quad w_1 \Rightarrow^* u, \quad w_3 \Rightarrow^* v, \quad A \Rightarrow^* w, \quad w_4 \Rightarrow^* x \quad \text{und} \quad w_2 \Rightarrow^* y.$$

Da G kontextfrei ist, hängen diese Ableitungen nicht zusammen, man kann sie unabhängig voneinander in einer beliebigen Reihenfolge ausführen. Daher sind in G auch folgende Ableitungen möglich:

$$S \Rightarrow^* uAy \Rightarrow^+ uw_3Aw_4y \Rightarrow^* uvAxy \Rightarrow^* uvwxy.$$

Der erste Schritt in der Teildableitung $uAy \Rightarrow^+ uw_3Aw_4y$ erfolgt durch Anwendung einer Regel der Form $A \rightarrow BC$, daher ergibt sich

$$uAy \Rightarrow uBCy \Rightarrow^* uw_3Aw_4y \quad \text{und} \quad BC \Rightarrow^* w_3Aw_4 \Rightarrow^* vAx. \quad \text{Wären beide Teilworte } v \text{ und } x \text{ leer, so könnte man in } G \text{ eine Ableitung } BC \Rightarrow^* A \text{ durchführen. Dieses ist nicht möglich, da } G \text{ in Chomsky-Normalform vorliegt. Daher gilt Eigenschaft (ii).}$$

Die Teildableitung $A \Rightarrow^* vwx$ hat höchstens $|N|$ viele Ableitungsschritte, da in G wegen der Chomsky-Normalform keine Ableitungen der Form $A \Rightarrow^+ B$ möglich sind. Daher gilt $|vwx| \leq 2^{|N|} < n_0$ (Eigenschaft (i)).

Die in G mögliche Ableitung $A \Rightarrow^* vAx$ kann man beliebig oft in die Ableitung $S \Rightarrow^* uvwxy$ einbauen und mit der Ableitung $A \Rightarrow^* w$ kombinieren:

$$S \Rightarrow^* uAy \Rightarrow^* uvAxy \Rightarrow^* uvvAxxxy \Rightarrow^* uv^k Ax^k y \Rightarrow^* uv^k wx^k y \quad \text{für } k \geq 1; \quad \text{außerdem ist die Ableitung } S \Rightarrow^* uAy \Rightarrow^* uwy \text{ möglich. Insgesamt ist dieses die Eigenschaft (iii).} //$$

Mit Hilfe dieses Satzes läßt sich zeigen, daß die Sprache $L_4 = \{a^n b^n c^n \mid n \in \mathbb{N}, n \geq 1\}$ (aus Kapitel 4.1) nicht kontextfrei (aber kontextsensitiv ist). Es sei $n = n_0$. Für das Wort $z = a^n b^n c^n$ ist $|z| = 3n_0 \geq n_0$. Dann läßt sich z zerlegen in $z = uvwxy$ mit den obigen Eigenschaften (i), (ii) und (iii). Der Teil vwx enthält höchstens n_0 viele Zeichen und kann daher nicht gleichzeitig aus a 's, b 's und c 's bestehen. Eigenschaft (iii) besagt, daß auch $uv^0wx^0y \in L_4$ ist, d.h. $uwy \in L_4$. Es werden drei Fälle unterschieden:

1. Fall: vwx enthält kein Zeichen c . Dann ist $uvw = a^n b^m$, $y = b^{n-m} c^n$. Da $|vx| > 0$ ist, enthält vx $k \geq 1$ Zeichen a oder b . Das Wort uwy enthält $n + m - k + n - m = 2n - k < 2n$ Zeichen a oder b und n Zeichen c . Daher ist $uwy \notin L_4$.

2. Fall: vwx enthält mindestens ein Zeichen c . Wegen $|vwx| \leq n_0$ enthält es kein Zeichen a . Es ist $u = a^n b^m$, $vwx = b^{n-m} c^l$, $y = c^{n-l}$ mit $l \geq 1$. Da $|vx| > 0$ ist, enthält vx $k \geq 1$ Zeichen b oder c . Das Wort uwy enthält n Zeichen a , $m + n - m + l - k = n + l - k$ Zeichen b oder c und $n - l$ Zeichen c . Daher ist $uwy \notin L_4$.

In beiden Fällen ergibt sich ein Widerspruch. Daher ist Sprache $L_4 = \{a^n b^n c^n \mid n \in \mathbb{N}, n \geq 1\}$ nicht kontextfrei.

Ein weiteres Beispiel ist die kontextsensitive Sprache $L = \{a^{2^i} \mid i \geq 1\}$.

Es gilt daher die in Satz 4.4-1 formulierte Aussage, daß die Klasse der kontextfreien Sprachen über einem endlichen Alphabet Σ eine echte Teilmenge der Klasse der kontextsensitiven Sprachen über Σ ist.

Exemplarisch für viele interessante Entscheidungsprobleme im Zusammenhang mit kontextfreien Grammatiken und kontextfreien Sprachen sollen wieder das Wortproblem und das Leerheitsproblem, jetzt bezogen auf kontextfreie Sprachen, betrachtet werden.

Wieder wird (analog zu den Typ-0- und Typ-1-Grammatiken) ein Prädikat $VERIFIZIERE_G_TYP_2(w)$

$$= \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Typ-2-Grammatik darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Typ-2-Grammatik darstellt} \end{cases}$$

definiert. Dieses Prädikat ist berechenbar. Dazu ist unter anderem zu prüfen, ob die Erzeugungsregeln in KG_w den Bedingungen einer kontextfreien Grammatik genügen.

Das **Wortproblem für Typ-2-Grammatiken** fragt danach, ob die Menge

$$L_{Wort_Typ-2} = \{u \# w \mid u \in \Sigma^*, w \in \{0,1\}^*, VERIFIZIERE_G_TYP_2(w) = \text{TRUE} \text{ und } u \in L(KG_w)\}$$

entscheidbar ist. In vereinfachter Darstellung wird also danach gefragt, ob es einen auf allen Eingaben stoppenden Algorithmus gibt, der bei Eingabe (der Kodierung) einer kontextfreien Grammatik G über dem Alphabet Σ und eines Worts $u \in \Sigma^*$ entscheidet, ob $u \in L(G)$ gilt oder nicht. Da dieses Entscheidungsproblem für kontextsensitive Sprachen entscheidbar ist, ist es auch im kontextfreien Fall entscheidbar. Das Entscheidungsverfahren aus Kapitel 4.3 ist im Falle einer kontextfreien Grammatik jedoch zu aufwendig (exponentielles Laufzeitverhalten). In der Theorie des Compilerbaus, die sich intensiv mit der Frage beschäftigt, ob ein Wort zur Sprache einer gegebenen Grammatik gehört, wird gezeigt, daß die Frage sogar mit einem Zeitaufwand der Ordnung $O(|u|^3)$ bei gegebener Grammatik G entschieden werden kann. Auf Details soll hier verzichtet werden.

Das **Leerheitsproblem**, das für Typ-0- und Typ-1-Grammatiken nicht entscheidbar ist, fragt für **Typ-2-Grammatiken**, ob die Menge

$$L_{\text{leer_Typ-2}} = \left\{ w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_G_TYP_2}(w) = \text{TRUE und } L(KG_w) = \emptyset \right\}$$

entscheidbar ist, bzw. in vereinfachter Darstellung, ob es einen auf allen Eingaben stoppenden Algorithmus gibt, der bei Eingabe (der Kodierung) einer Grammatik G entscheidet, ob $L(G) = \emptyset$ gilt oder nicht.

Der folgende (Pseudocode-) Algorithmus entscheidet bei Eingabe einer kontextfreien Grammatik die $G = (\Sigma, N, S, R)$ die Frage „ $L(G) \neq \emptyset$?“. Aus diesem Algorithmus lässt sich leicht ein Entscheidungsalgorithmus für das Leerheitsproblem für Typ-2-Grammatiken gewinnen.

Eingabe: Eine kontextfreie Grammatik $G = (\Sigma, N, S, R)$

Verfahren: Aufruf der Funktion `ist_nichtleer (G)`

Ausgabe: TRUE, falls $L(G) \neq \emptyset$, FALSE sonst.

FUNCTION `ist_nichtleer (G): BOOLEAN;`

`{ G = (Σ, N, S, R) }`

VAR `V : ...;`

`W : ...;`

BEGIN `{ ist_nichtleer }`

`V := ∅;`

`W := { A | A ∈ N und A → w ist eine Erzeugungsregel mit w ∈ Σ* };`

WHILE NOT (V = W) DO

BEGIN

`V := W;`

`W := { A | A ∈ N und A → w ist eine Erzeugungsregel mit w ∈ (V ∪ Σ)* } ∪ V;`

END;

IF `S ∈ W` THEN `ist_nichtleer := TRUE`

ELSE `ist_nichtleer := FALSE;`

END `{ ist_nichtleer };`

Es läßt sich zeigen, daß die WHILE-Schleife höchstens n -mal durchlaufen wird, wenn G n Erzeugungsregeln enthält. Daher bricht das Verfahren ab. Die Korrektheit des Verfahrens ergibt sich aus der Behauptung

Ein nichtterminales Symbol $A \in N$ wird im i -ten Durchlauf der WHILE-Schleife genau dann in W aufgenommen, wenn es ein $u \in \Sigma^*$ gibt mit $A \Rightarrow^* u$.

Zusammenfassend ergibt sich

Satz 4.4-4:

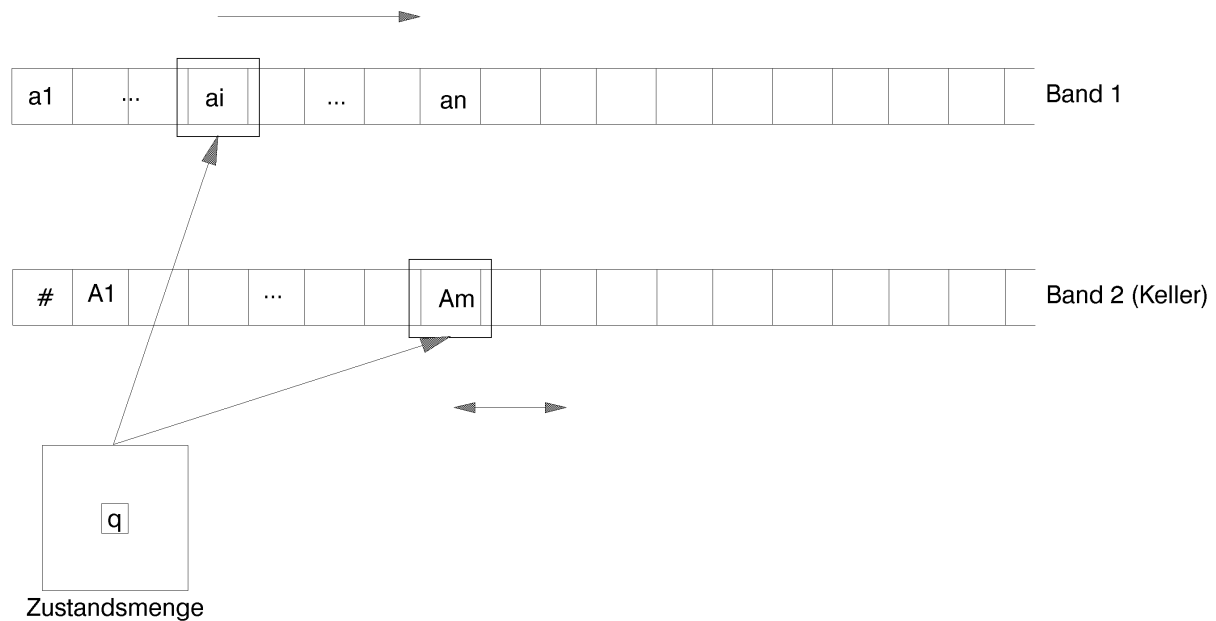
Das Wortproblem und das Leerheitsproblem für Typ-2-Grammatiken sind entscheidbar:

Es gibt einen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-2-Grammatik $G = (\Sigma, N, S, R)$ und eines Wortes $u \in \Sigma^*$ entscheidet, ob $u \in L(G)$ ist.

Es gibt einen auf allen Eingaben stoppenden Algorithmus, der bei Eingabe einer Typ-2-Grammatik $G = (\Sigma, N, S, R)$ entscheidet, ob $L(G) = \emptyset$ gilt.

Auch zu den kontextfreien Sprachen gibt es ein Berechnungsmodell, das genau diese Sprachen erkennt und als eine Einschränkung einer Turingmaschine gesehen werden.

Ein **nichtdeterministischer Kellerautomat (NKA)** ist eine 2-NDTM, die ihre beiden Bänder in einer speziellen Weise nutzt. Vor der Angabe einer formalen Definition eines nichtdeterministischen Kellerautomaten wird das Modell informell beschrieben. Das 1. Band (Eingabeband) eines nichtdeterministischen Kellerautomaten NKA kann nur gelesen werden, wobei der Lesekopf nach dem Lesen eines Zeichens um eine Zelle nach rechts rückt. Das 1. Band enthält bei Start von NKA ein Wort über einem Eingabealphabet, der Lesekopf steht über dem ersten Zeichen des Eingabeworts. NKA kann auch tätig werden, wenn das Eingabeband eine leere Zeichenkette enthält; dann findet ein „spontaner (Konfigurations-) Übergang“ statt. Er kann auch Zustandsänderungen durchführen, ohne ein Symbol des Eingabebands zu lesen; man nennt diese Konfigurationsübergänge e -Überführungen. Das 2. Band heißt **Keller**; auf dem Keller bewegt sich ein Schreib/Lesekopf. Der Keller enthält zu Beginn der Berechnungsfolge von NKA ein spezielles Symbol $\#$, und der Schreib/Lesekopf steht über diesem Symbol. Bei jeder Konfigurationsänderung wird das Symbol unter dem Schreib/Lesekopf durch ein endlich langes Wort ersetzt, das mit Hilfe eines Kelleralphabets gebildet wird, und der Schreib/Lesekopf rückt über das letzte Zeichen der nun im Keller befindlichen Zeichenkette. Es kann auch das leere Wort in den Keller geschrieben werden; dann wird der Inhalt des Kellers verkürzt. Auf diese Weise kann immer nur auf das zuletzt in den Keller gebrachte Zeichen zugegriffen werden. Der Schreib/Lesekopf im Keller rückt nicht über das Ende des Kellers auf Zeichen, die weiter links stehen (last-in-first-out-Prinzip).



Eine formale Definition eines nichtdeterministischen Kellerautomaten wandelt die Definition einer nichtdeterministischen Turingmaschine ab. Die Definition einer Konfiguration und des Konfigurationsübergangs wird der besonderen Arbeitsweise eines Kellerautomaten angepaßt. In der Literatur finden sich häufig Varianten der hier gegebenen Definitionen, die aber auf dieselbe Sprachklasse führen.

Ein **nichtdeterministischer Kellerautomat** *NKA* ist definiert durch

$NKA = (Q, \Sigma, \Gamma, d, q_0, \#, F)$ mit:

1. Q ist eine endliche nichtleere Menge: die **Zustandsmenge**
2. Σ ist eine endliche nichtleere Menge: das **Eingabealphabet**
3. Γ ist eine endliche nichtleere Menge: das **Kelleralphabet**
4. $d : Q \times (\Sigma \cup \{e\}) \times \Gamma \rightarrow \mathbf{P}_e(Q \times \Gamma^*)$ ist eine partielle Funktion, die **Überföhrungsfunktion**;
hierbei bezeichnet $\mathbf{P}_e(Q \times \Gamma^*)$ die Menge aller endlichen Teilmengen von $Q \times \Gamma^*$
5. $q_0 \in Q$ ist der **Anfangszustand** (oder **Startzustand**)
6. $\# \in \Gamma$ ist das **Startsymbol im Keller**
7. $F \subseteq Q$ ist die **Menge der Endzustände**.

In jedem Schritt eines nichtdeterministischen Kellerautomaten ist klar, wo sich der Kopf des jeweiligen Bandes befindet: Wenn das Eingabewort w die Form $w = w_1 w_2$ hat und der Kellerautomat bereits die Zeichen in w_1 gelesen hat, dann steht der Lesekopf des 1. Bandes über

dem ersten Zeichen von w_2 , und auf die Zeichen von w_1 kann der Kopf nicht mehr zugreifen. Anstelle daher in einer Konfiguration für das 1. Band den gesamten Bandinhalt und die Position des Kopfes in der Form $(w_1 w_2, |w_1| + 1)$ festzuhalten, wird der Inhalt des 1. Bandes nur durch die Zeichenkette w_2 beschrieben; implizit wird angenommen, daß der Kopf über dem ersten Zeichen von w_2 steht. Entsprechend braucht man in einer Konfiguration für das 2. Band nicht den gesamten Bandinhalt und die Position des Kopfes in der Form $(a, |a|)$ festzuhalten, sondern es genügt die Zeichenkette a zusammen mit der impliziten Annahme, daß der Kopf über dem zuletzt in den Keller geschriebenen Zeichen steht. Daher wird **eine Konfiguration für einen nichtdeterministischen Kellerautomaten** in der Form (q, w, a) mit $q \in Q$, $w \in \Sigma^*$ und $a \in \Gamma^*$ notiert. Die Konfiguration drückt aus, daß sich der Kellerautomat gegenwärtig im Zustand q befindet, daß das Eingabeband eine Zeichenkette $w_1 w$ enthält, von der bisher die Zeichen in $w = a_1 \dots a_n$ noch nicht gelesen wurden, daß der Lesekopf des 1. Bandes über dem ersten Zeichen a_1 steht, daß sich im Keller die Zeichenkette $a = A_1 \dots A_m$ befindet, und daß der Schreib/Lesekopf des 2. Bandes über A_m steht. Falls $d(q, a_1, A_m)$ definiert ist (das ist eine endliche Teilmenge von $Q \times \Gamma^*$) und ein Element $(q', B_1 \dots B_k)$ enthält, dann kann der Kellerautomat in eine **Nachfolgekonfiguration** übergehen, die den Zustand q' enthält, den Lesekopf auf dem 1. Band um eine Position nach rechts verrückt, im Keller das Zeichen A_m durch $B_1 \dots B_k$ ersetzt und den Schreiblesekopf im Keller auf B_k positioniert hat. Formal wird für Konfigurationen die Übergangsrelation $\Rightarrow$ definiert durch

$$(q, a_1 \dots a_n, A_1 \dots A_m) \Rightarrow \begin{cases} (q', a_2 \dots a_n, A_1 \dots A_{m-1} B_1 \dots B_k) & \text{falls } (q', B_1 \dots B_k) \in d(q, a_1, A_m) \\ (q', a_1 a_2 \dots a_n, A_1 \dots A_{m-1} B_1 \dots B_k) & \text{falls } (q', B_1 \dots B_k) \in d(q, e, A_m) \end{cases}$$

Entsprechend bedeutet für Konfigurationen K und K' die Beziehung $K \Rightarrow^* K'$:

Es gibt Konfigurationen $K_0, K_1, \dots, K_l$ mit $l \geq 0$ und $K = K_0, K' = K_l$ und $K_i \Rightarrow K_{i+1}$ für $i = 0, \dots, l-1$.

Eine **Anfangskonfiguration für einen nichtdeterministischen Kellerautomaten** NKA hat die Form $(q_0, w, \#)$, eine **Endkonfiguration** hat die Form (q, e, e) mit $q \in F$.

Ein Wort $w \in \Sigma^*$ wird von NKA **akzeptiert**, wenn $(q_0, w, \#) \Rightarrow^* (q, e, e)$ mit $q \in F$ gilt. Die von NKA **akzeptierte Sprache** ist $L(NKA) = \{w \mid w \in \Sigma^*, \text{ und } w \text{ wird von } NKA \text{ akzeptiert}\}$.

Beispiel:

Ein nichtdeterministischer Kellerautomat zum Erkennen der Sprache

$$L = \{ w \mid w = a_1 \dots a_m a_m \dots a_1, a_i \in \{a, b\} \}$$

Die Sprache $L = \{ w \mid w = a_1 \dots a_m a_m \dots a_1, a_i \in \{a, b\} \}$ läßt sich beispielsweise durch die kontextfreie Grammatik $G = (\{a, b\}, \{S\}, S, \{S \rightarrow aSa, S \rightarrow bSb, S \rightarrow aa, S \rightarrow bb\})$ erzeugen.

Der nichtdeterministische Kellerautomaten NKA wird gegeben durch

$NKA = (\{q_0, q_1, q_2\}, \{a, b\}, \{A, B, \#\}, d, q_0, \#, \{q_2\})$ mit der durch folgende Tabelle festgelegten Überföhrungsfunktion d :

momentaner Zustand	Symbol unter dem Kopf auf		neuer Zustand	neues Wort im Keller
	Band 1	Keller		
q_0	a	X mit $X \in \{A, B, \#\}$	q_0	XA
	a	X mit $X \in \{A, B, \#\}$	q_1	XA
	b	X mit $X \in \{A, B, \#\}$	q_0	XB
	b	X mit $X \in \{A, B, \#\}$	q_1	XB
q_1	a	A	q_1	e
	b	B	q_1	e
	e	$\#$	q_2	e

Der Nichtdeterminismus wird in diesem Beispiel eingesetzt, um die Mitte eines Eingabewortes „zu raten“.

Auch im Modell des Kellerautomaten kann man Nichtdeterminismus und Determinismus unterscheiden:

Ein **deterministischer Kellerautomat** DKA ist wie ein nichtdeterministischer Kellerautomat definiert mit dem Unterschied, daß die Überföhrungsfunktion die Form

$$d : Q \times (\Sigma \cup \{e\}) \times \Gamma \rightarrow Q \times \Gamma^*$$

aufweist.

Beispiel:

Ein deterministischer Kellerautomat zum Erkennen der Sprache

$$L = \{a^n b^n \mid n \in \mathbf{N}\}$$

Die Sprache $L = \{a^n b^n \mid n \in \mathbf{N}\}$ läßt sich beispielsweise durch die kontextfreie Grammatik $G = (\{a, b\}, \{S\}, S, \{S \rightarrow aSb, S \rightarrow e\})$ erzeugen.

Der deterministische Kellerautomaten DKA wird gegeben durch

$DKA = (\{q_0, q_1, q_2\}, \{a, b\}, \{a, b, \#\}, d, q_0, \#, \{q_0\})$ mit der durch folgende Tabelle festgelegten Überföhrungsfunktion d :

momentaner Zustand	Symbol unter dem Kopf auf		neuer Zu- stand	neues Wort im Keller
	Band 1	Keller		
q_0	a	$\#$	q_1	$\#a$
	e	$\#$	q_0	e
q_1	a	a	q_1	aa
	b	a	q_2	e
q_2	b	A	q_2	e
	e	$\#$	q_0	e

Satz 4.4-5:

Die Klasse der von nichtdeterministischen Kellerautomaten akzeptierten Sprachen über einem endlichen Alphabet Σ ist mit der Klasse der kontextfreien Sprachen (Typ-2-Sprachen) über Σ identisch.

Beweis:

Es ist zu zeigen, daß es zu jeder von einem nichtdeterministischen Kellerautomaten NKA akzeptierten Sprache $L = L(NKA)$ eine kontextfreie Grammatik G_{NKA} gibt mit $L = L(G_{NKA})$. Die umgekehrte Richtung, nämlich der Nachweis, daß es zu jeder von einer kontextfreien Grammatik G erzeugten Sprache $L' = L(G)$ einen nichtdeterministischen Kellerautomaten NKA_G mit $L(NKA_G) = L'$ gibt, ist von größerer praktischer Relevanz, da hierdurch implizit gezeigt wird, wie man zu einer vorgegebenen Grammatik einen Algorithmus (hier in Form eines

Kellerautomaten) entwerfen kann, der die Wörter der von der Grammatik erzeugten Sprache erkennt. Diese Beweisrichtung soll daher skizziert werden.

Es sei $G = (\Sigma, N, S, R)$ eine kontextfreie Grammatik. Ein nichtdeterministischer Kellerautomat $NKA_G = (Q, \Sigma, \Gamma, d, q_0, \#, F)$ mit $L(NKA_G) = L(G)$ wird gegeben durch die Zustandsmenge $Q = \{z\}$, das Kelleralphabet $\Gamma = N \cup \Sigma$, den Anfangszustand $q_0 = z$, das Startsymbol im Keller $\# = S$, die Menge der Endzustände $F = \{z\}$ und die Überföhrungsfunktion d , die folgendermaßen definiert ist:

- (i) Für jede Erzeugungsregel $A \rightarrow a$ aus R wird (z, a^R) in $d(z, e, A)$ aufgenommen; hierbei ist a^R die Konkation der Buchstaben von a in umgekehrter Reihenfolge (Spiegelung von a)
- (ii) Für jedes $a \in \Sigma$ wird (z, e) in $d(z, a, a)$ aufgenommen.

Wegen (i) kann immer dann, wenn das im Keller gelesene Symbol ein nichtterminales Zeichen ist, dieses durch die rechte Seite einer Regel (in gespiegelter Reihenfolge) ersetzt werden. Ein Terminalzeichen, das gerade im Keller gelesen wird, wird entfernt, wenn es mit dem nächsten Eingabesymbol übereinstimmt. Man kann $L(NKA_G) = L(G)$ zeigen. ///

Beispiel:

Ein nichtdeterministischer Kellerautomat für eine kontextfreie Grammatik

Gegeben sei die Grammatik $G = (\Sigma, N, S, R)$ mit $\Sigma = \{a, b, c\}$, $N = \{S, A\}$ und der Produktionsmenge $R = \{S \rightarrow A, A \rightarrow aAb, A \rightarrow c\}$. Die in der Konstruktion beschriebene Überföhrungsfunktion wird durch folgende Tabelle gegeben:

momentaner Zustand	Symbol unter dem Kopf auf		neuer Zu- stand	neues Wort im Keller
	Band 1	Keller		
z	e	S	z	A
	e	A	z	bAa
	e	A	z	c
	A	a	z	e
	B	b	z	e
	C	c	z	e

Das Wort $aacbb \in L(G)$ wird durch folgenden Konfigurationsübergänge akzeptiert:

$$\begin{aligned}
 (z, aacbb, S) &\Rightarrow (z, aacbb, A) \Rightarrow (z, aacbb, bAa) \Rightarrow (z, acbb, bA) \Rightarrow (z, acbb, bbAa) \\
 &\Rightarrow (z, cbb, bbA) \Rightarrow (z, cbb, bbc) \Rightarrow (z, bb, bb) \Rightarrow (z, b, b) \\
 &\Rightarrow (z, e, e).
 \end{aligned}$$

Jede von einem deterministischen Kellerautomaten akzeptierte Sprache wird auch von einem nichtdeterministischen Kellerautomaten akzeptiert. Umgekehrt gibt es Sprachen, etwa $L = \{ w \mid w = a_1 \dots a_m a_m \dots a_1, a_i \in \{a, b\} \}$, die nicht von einem deterministischen Kellerautomaten akzeptiert werden können. Insbesondere weisen die Klassen der deterministisch kontextfreien und die Klasse der nichtdeterministisch kontextfreien Sprachen unterschiedliche Abschlußeigenschaften auf (siehe Zusammenstellung in Kapitel 4.6).

Satz 4.4-5:

Die Klasse der deterministisch kontextfreien Sprachen über einem endlichen Alphabet Σ ist eine echte Teilmenge der Klasse der nichtdeterministisch kontextfreien Sprachen Σ .

Die Klasse der deterministisch kontextfreien Sprachen spielt eine besondere Rolle im Compilerbau. Eine derartige Sprache erlaubt die Syntaxanalyse eines Wortes (in der Anwendung ist dieses ein Programm in einer Programmiersprache), indem nur wenige Zeichen des Wortes während des Analysevorgangs im Vorgriff gelesen werden, um festzustellen, welche Produktion in einer Ableitung als nächstes anzuwenden ist bzw. daß das Wort nicht zu der von der Grammatik erzeugten Sprache gehört. Für deterministisch kontextfreie Sprachen ist das Wortproblem für ein Wort u mit einem Algorithmus entscheidbar, der eine Zeitkomplexität der Ordnung $O(|u|)$ aufweist. Die meisten Programmiersprachen sind weitgehend deterministisch kontextfrei, so daß in der Praxis die Syntaxanalyse bei Programmiersprachen sehr effizient abläuft.

4.5 Typ-3-Sprachen

Es sei $G = (\Sigma, N, S, R)$ eine Grammatik, in der die Erzeugungsregeln

$R \subseteq \{ A \rightarrow w \mid A \in N \text{ und } w \in (N \cup \Sigma)^* \}$ folgende zusätzliche Bedingung erfüllen:

Alle Regeln haben die Form $A \rightarrow aB$ mit $A \in N$, $B \in N$, $a \in \Sigma$ oder $A \rightarrow e$.

Die von einer derartigen Grammatik erzeugte Sprache heißt **Typ-3-Sprache** oder **rechtslineare Sprache** oder **reguläre Sprache**. Die zugehörige Grammatik heißt **Typ-3-Grammatik** oder **rechtslineare Grammatik**.

Beispielsweise ist die Sprache $L_1 = \{ a^i b^j \mid i \in \mathbf{N}, j \in \mathbf{N} \}$ aus Kapitel 4.1 regulär (rechtslinear, Typ-3-Sprache), da sie durch die Grammatik

$G'_1 = (\{a, b\}, \{S, B\}, S, \{S \rightarrow e, S \rightarrow aS, S \rightarrow bB, B \rightarrow bB, B \rightarrow e\})$ erzeugt wird.

Der Name „rechtslineare Sprache“ erklärt sich durch die Form der Erzeugungsregeln $A \rightarrow aB$ der zugehörigen Grammatik: Betrachtet man die Nichtterminalsymbole als „Variablen“ (im Sinne eines Gleichungssystems) und die Terminalsymbole als „Konstanten“, so liegt in der Form $A \rightarrow aB$ eine lineare Beziehung zwischen den Variablen vor, in der die Variablen rechts der Konstanten stehen.

Der Name „reguläre Sprache“ weist auf eine alternative Möglichkeit hin, um eine Typ-3-Sprache zu definieren: eine derartige Sprache lässt sich durch einen „regulären Ausdruck“ repräsentieren. Der Vollständigkeit soll hier die Definition eines regulären Ausdrucks und der durch ihn repräsentierten Sprache angeführt werden, wenn auch im folgenden auf dieses Konzept nicht weiter eingegangen wird.

Ein **regulärer Ausdruck über Σ** und **die durch ihn repräsentierte Sprache** wird rekursiv definiert durch:

- (i) $\emptyset$ ist ein regulärer Ausdruck, der die reguläre Sprache $\emptyset$ repräsentiert
- (ii) e ist ein regulärer Ausdruck, der die reguläre Sprache $\{e\}$ repräsentiert
- (iii) a aus Σ ist ein regulärer Ausdruck, der die reguläre Sprache $\{a\}$ repräsentiert
- (iv) Sind p bzw. q reguläre Ausdrücke, die die regulären Sprachen P bzw. Q repräsentieren, dann ist
 - $(p + q)$ ein regulärer Ausdruck, der die reguläre Sprache $P \cup Q$ repräsentiert,
 - (pq) ein regulärer Ausdruck, der die reguläre Sprache $P \cdot Q$ repräsentiert,
 - $(p)^*$ ein regulärer Ausdruck, der die reguläre Sprache P^* repräsentiert
- (v) Ein regulärer Ausdruck und die von ihm repräsentierte Sprache wird genau durch die Punkte (i) bis (iv) definiert.

Folgender Satz lässt sich beweisen:

Satz 4.5-1:

Eine Sprache L über einem endlichen Alphabet Σ wird genau dann durch einen regulären Ausdruck repräsentiert, wenn es eine Typ-3-Grammatik (rechtslineare Grammatik) $G = (\Sigma, N, S, R)$ mit $L = L(G)$ gibt.

Im folgenden werden daher die Begriffe „Typ-3-Sprache“, „rechtslineare Sprache“ und „reguläre Sprache“ synonym verwendet.

Auch für die regulären Sprachen lässt sich ein Berechnungsmodell zur Akzeptanz gerade dieses Typs angeben. Man erhält es durch Einschränkung eines Kellerautomaten. Der Keller ei-

nes Kellerautomaten erlaubt die Zwischenablage von (evtl. transformierten) Eingabezeichen. Entfernt man den Keller, so kommt man auf das folgende Modell:

Ein **nichtdeterministischer endlicher Automat** *NEA* ist definiert durch

$NEA = (Q, \Sigma, d, q_0, F)$ mit:

1. Q ist eine endliche nichtleere Menge: die **Zustandsmenge**
2. Σ ist eine endliche nichtleere Menge: das **Eingabealphabet**
3. $d : Q \times \Sigma \rightarrow \mathbf{P}(Q)$ ist eine partielle Funktion, die **Überföhrungsfunktion**; hierbei bezeichnet $\mathbf{P}(Q)$ die Menge aller Teilmengen von Q .
4. $q_0 \in Q$ ist der **Anfangszustand** (oder **Startzustand**)
5. $F \subseteq Q$ ist die **Menge der Endzustände**.

Entsprechend liegt ein **deterministischer endlicher Automat** *DEA* vor, wenn die Überföhrungsfunktion d die Form $d : Q \times \Sigma \rightarrow Q$ aufweist.

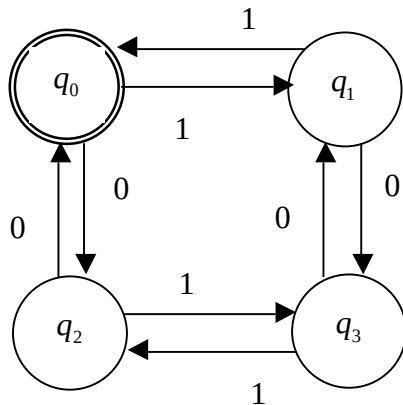
Ein deterministischer bzw. nichtdeterministischer endlicher Automat läßt sich auch in Form eines endlichen gerichteten Graphen G mit markierten Kanten als **Transitionsdiagramm** darstellen. Die Knotenmenge von G ist Q . Ist $z' = d(z, a)$ bzw. im nichtdeterministischen Fall $z' \in d(z, a)$, so wird eine mit a markierte Kante in G aufgenommen. Die Menge der den Endzuständen entsprechenden Knoten werden explizit markiert.

Beispiel:

Es sei $DEA = (Q, \Sigma, d, q_0, F)$ der deterministische endliche Automat mit $Q = \{q_0, q_1, q_2, q_3\}$, $\Sigma = \{0, 1\}$, $F = \{q_0\}$, d gemäß folgender Tabelle:

momentaner Zustand	gelesenes Symbol	neuer Zustand
q_0	0	q_2
	1	q_1
q_1	0	q_3
	1	q_0
q_2	0	q_0
	1	q_3
q_3	0	q_1
	1	q_2

Das entsprechende Transitionsdiagramm ist



Es ist

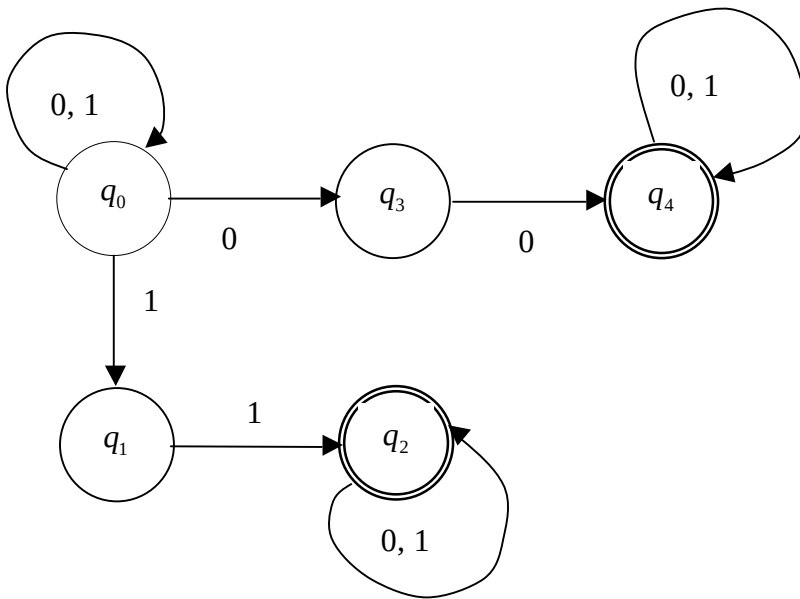
$$L(DEA) = \{ w \mid w \in \{0,1\}^* \text{ und } w \text{ enthält eine gerade Anzahl von Zeichen 0 und Zeichen 1} \}.$$

Ein Beispiel eines nichtdeterministischen endlichen Automaten mit dem zugehörigen Transitionsdiagramm ist das folgende:

$NEA = (Q, \Sigma, d, q_0, F)$ mit $Q = \{q_0, q_1, q_2, q_3, q_4\}$, $\Sigma = \{0, 1\}$, $F = \{q_2, q_4\}$, d gemäß folgender Tabelle:

momentaner Zustand	gelesenes Symbol	neuer Zustand
q_0	0	q_0
	1	q_3
q_1	1	q_2
q_2	0	q_2
	1	q_2
q_3	0	q_4
q_4	0	q_4
	1	q_4

Das zugehörige Transitionsdiagramm lautet



Es ist

$$L(NEA) = \left\{ w \mid \begin{array}{l} w \in \{0,1\}^+ \\ \text{und } w \text{ enthält mindestens zwei aufeinanderfolgende Zeichen 0 oder Zeichen 1} \end{array} \right\}.$$

Wie bei einem Kellerautomaten durch Weglassen des Kellers kann man auch hier Konfigurationen und Konfigurationsübergänge definieren. Ein Wort $w \in \Sigma^*$ wird von *NEA* (*DEA*) **akzeptiert**, wenn $(q_0, w) \Rightarrow^* (q, e)$ mit $q \in F$ gilt. Die von *NEA* (*DEA*) **akzeptierte Sprache** ist $L(NEA) = \{ w \mid w \in \Sigma^*, \text{ und } w \text{ wird von } NEA \text{ akzeptiert} \}$ (entsprechend für den deterministischen Fall).

Bei Turingmaschinen fallen Determinismus und Nichtdeterminismus zusammen, jedoch zum Preis einer exponentiell wachsenden Berechnungs- bzw. Laufzeitkomplexität. Der folgende Satz zeigt, daß auch bei den regulären Sprachen Nichtdeterminismus keine zusätzliche Berechnungsfähigkeit gegenüber dem Determinismus bringt. In den dazwischenliegenden Sprachklassen der kontextfreien und kontextsensitiven Sprachen ist der Determinismus vom Nichtdeterminismus zu unterscheiden bzw. der Unterschied zwischen Determinismus und Nichtdeterminismus ist nicht bekannt.

Satz 4.5-2:

Für jede von einem nichtdeterministischen endlichen Automaten akzeptierte Sprache über einem endlichen Alphabet Σ gibt es einen deterministischen endlichen Automaten, der die Sprache akzeptiert.

Beweis:

Es ist bei gegebenem nichtdeterministischen endlichen Automaten $NEA = (Q, \Sigma, d, q_0, F)$ ein deterministischer endlicher Automat $DEA = (Q', \Sigma, d', q'_0, F')$ anzugeben, der dieselbe Sprache akzeptiert. DEA wird definiert durch:

$$Q' = \mathbf{P}(Q), \quad q'_0 = \{q_0\}, \quad d'(Z', a) = \bigcup_{z \in Z'} d(z, a) \text{ für } Z' \subseteq Q \text{ und } a \in \Sigma,$$

$$F' = \{Z' \mid Z' \subseteq Q \text{ und } Z' \cap F \neq \emptyset\}. \quad ///$$

Der folgende Satz besagt, daß das Konzept der endlichen erkennenden Automaten das adäquate Modell für die regulären Mengen ist.

Satz 4.5-3:

Die Klasse der von (nichtdeterministischen bzw. deterministischen) endlichen Automaten akzeptierten Sprachen über einem endlichen Alphabet Σ ist mit der Klasse der regulären Sprachen (Typ-3-Sprachen) über Σ identisch.

Beweis:

Es sei $NEA = (Q', \Sigma, d', q'_0, F')$ ein nichtdeterministischer endlicher Automat. Dann gibt es nach Satz 4.5-2 einen deterministischen endlichen Automaten $DEA = (Q, \Sigma, d, q_0, F)$ mit $L(DEA) = L(NEA)$. Es wird eine rechtslineare Grammatik $G = (\Sigma, N, S, R)$ angegeben, für die $L(G) = L(DEA)$ gilt:

$N = Q$, $S = q_0$, die Menge R der Regeln wird definiert durch:

R enthält genau dann eine Erzeugungsregel $q \rightarrow aq'$ mit $q \in Q$, $q' \in Q$ und $a \in \Sigma$, wenn $d(q, a) = q'$ ist; für $q \in F$ enthält R außerdem die Regel $q \rightarrow e$.

Es sei umgekehrt $G = (\Sigma, N, S, R)$ eine rechtslineare Grammatik. Dann erkennt der folgende Automat $NEA = (Q, \Sigma, d, q_0, F)$ die Sprache $L(G)$:

Es sei A ein nichtterminales Symbol mit $A \notin N \cup \Sigma$. Es ist

$Q = N \cup \{A\}, \quad q_0 = S, \quad F = \begin{cases} \{A\} & \text{für } e \notin L(G) \\ \{S, A\} & \text{für } e \in L(G) \end{cases},$ die Überföhrungsfunktion

$d : Q \times (\Sigma \cup \{e\}) \rightarrow \mathbf{P}(Q)$ wird definiert durch:

$d(B, a)$ enthalt genau dann C , wenn R die Regel $B \rightarrow aC$ enthalt; f r jede Regel der Form $B \rightarrow e$ wird $d(B, e) = \{A\}$ gesetzt. Der Automat $NEA = (Q, \Sigma, d, q_0, F)$ erf llt noch nicht die Definition eines endlichen Automaten; denn diese lsst keine e - bergnge zu, d.h. ein endlicher Automat macht keine  berf hrungen, in denen kein Eingabesymbol gelesen wird. Der hier definierte Automat lsst jedoch eventuell derartige  berf hrungen zu. Es lsst sich jedoch zeigen (siehe angegebene Literatur), da  man NEA so zu einem nichtdeterministischen endlichen Automaten $NEA' = (Q', \Sigma, d', q'_0, F')$ modifizieren kann, so da  dessen  berf hrungsfunktion $d' : Q' \times \Sigma \rightarrow \mathbf{P}(Q')$ die Definition eines nichtdeterministischen endlichen Automaten erf llt und $L(NEA') = L(G)$ gilt. ///

Auch f r regulre Sprachen gibt es einen dem $uvwxy$ -Theorem entsprechenden Satz, der wieder ein Beweismittel liefert, mit dessen Hilfe man zeigen kann, da  eine Sprache nicht regulr ist:

Satz 4.5-4:

Zu jeder regulren (Typ-3-, rechtslinearen) Sprache L gibt es eine Zahl $n_0 > 0$ mit der Eigenschaft:

jedes $z \in L$ mit $|z| \geq n_0$ lsst sich zerlegen in $z = uvw$ mit

(i) $|uv| \leq n_0$

(ii) $|v| > 0$

(iii) $uv^k w \in L$ f r jedes $k \in \mathbf{N}$.

Beweis:

Die Argumentation kann hnlich  ber die Ableitung eines Wortes mit Hilfe einer regulren Grammatik verlaufen wie im Beweis des $uvwxy$ -Theorems. Einfacher geht es hier  ber die Erkennung eines Wortes mit Hilfe eines deterministischen endlichen Automaten:

Es sei regulren Sprache L und DEA ein deterministischer endlicher Automat mit $L(DEA) = L$. DEA enthalte n_0 viele Zustnde. Bei der Akzeptanz eines Wortes $z \in L$ mit $|z| \geq n_0$ durchluft DEA einschlie lich des Anfangszustands $|z| + 1$ viele Zustnde. Hierbei geht wesentlich ein, da  DEA bei jeder  berf hrung ein Eingabesymbol liest. In der Folge der Zustands berf hrungen vom Anfangszustand q_0 bis zu einem akzeptierenden Zustand $q_{accept} \in F$ mu  sich also ein Zustand q_i wiederholen. Es sei u das Anfangsteilwort von z , das

in den Überführungen gelesen wurde, die *DEA* von q_0 aus bis zum ersten Erreichen von q_i gemacht hat. Das Teilwort v von z besteht aus den Symbolen, die *DEA* dann von q_i aus bis zum nächsten Erreichen von q_i gelesen hat. Schließlich ist w das Teilwort von z , das *DEA* liest, bis der akzeptierende Zustand q_{accept} erreicht ist. Mit dieser Zerlegung von z gelten die angegebenen drei Eigenschaften. ///

Mit Hilfe dieses Satzes läßt sich zeigen, daß die Menge der Palindrome

$L = \{w \mid w \in \{0,1\}^* \text{ und } w = a_1 \dots a_{n-1} a_n = a_n a_{n-1} \dots a_1\} \cup \{e\}$ nicht regulär ist; L ist jedoch kontextfrei. Es sei $n_0 > 0$ der Wert aus obigem Satz und $n \geq n_0$ die kleinste gerade Zahl $\geq n_0$. Das Wort $z = \underbrace{1010 \dots 100101 \dots 01}_{n \text{ Zeichen}} \underbrace{}_{n \text{ Zeichen}}$ liegt in L und läßt sich in der Form $z = uvw$ mit den Eigenschaften (i), (ii) und (iii) zerlegen. Zu beachten ist, daß die einzige Stelle, an der zwei gleiche Zeichen (00) aufeinanderfolgen, in der Mitte von z liegt. Das Teilwort uv ist Teil von $\underbrace{1010 \dots 10}_{n \text{ Zeichen}}$.

1. Fall: $v = 1$: Das Wort uv^2w ist in L und enthält zwei aufeinanderfolgende Zeichen 1. Damit uv^2w ein Palindrom ist müßten diese beiden Zeichen 1 jedoch in der Mitte des Worts liegen, da weiter rechts und links keine aufeinanderfolgenden Zeichen 1 stehen. Rechts der beiden Zeichen 1 gibt es noch die beiden Zeichen 0 (die ursprüngliche Mitte des Worts). Zu ihnen korrespondieren links der beiden Zeichen 1 keine entsprechende Zeichenfolge 00. Daher ist $uv^2w \notin L$.
2. Fall: $v = 0$: Das Wort uv^3w ist in L und enthält drei aufeinanderfolgende Zeichen 0. Damit uv^3w ein Palindrom ist müßten diese drei Zeichen 0 jedoch in der Mitte des Worts liegen, da weiter rechts und links keine aufeinanderfolgenden drei Zeichen 0 stehen. Zur ursprünglichen Mitte des Worts korrespondieren links der drei Zeichen 0 keine entsprechende Zeichenfolge 00. Daher ist $uv^3w \notin L$.
3. Fall: $|v| > 1$ und $v = 10 \dots 1$: Dann enthält Wort uv^2w zwei aufeinanderfolgende Zeichen 1 und man kann wie im 1. Fall argumentieren.
4. Fall: $|v| > 1$ und $v = 10 \dots 10$: Dann enthält Wort uv^2w zwei aufeinanderfolgende Zeichen 0 (die ursprüngliche Mitte) und links davon nur Teilzeichenketten 01 als rechts, und zwar mehr als rechts. Daher ist $uv^2w \notin L$.

In allen Fällen ergibt sich ein Widerspruch.

Die Regeln einer rechtslinearen Grammatik erfüllen die Bedingungen, die an die Regeln einer kontextfreien Grammatik gestellt werden. Jede von einer rechtslinearen Grammatik erzeugte Sprache ist daher auch eine kontextfreie Sprache. Die Sätze 4.5-2 und 4.5-3 implizieren, daß es zu jeder regulären Sprache über einem endlichen Alphabet Σ einen deterministischen

Kellerautomaten gibt, der die Sprache akzeptiert. Andererseits läßt sich mit Satz 4.5-4 zeigen, daß die deterministisch kontextfreie Sprache $L = \{a^n b^n \mid n \in \mathbf{N}\}$ nicht regulär ist. Es gilt daher:

Satz 4.5-5:

Die Klasse der regulären (Typ-3-, rechtslinearen) Sprachen über einem endlichen Alphabet Σ ist eine echte Teilmenge der deterministisch kontextfreien Sprachen über Σ .

Auch für Typ-3-Sprachen gibt es wieder viele interessante Entscheidungsprobleme. Hier sollen jedoch wieder nur das Wortproblem und das Leerheitsproblem, jetzt bezogen auf Typ-3-Sprachen, betrachtet werden. Dazu wird (analog zu den Typ-0-, Typ-1- und Typ-2-Grammatiken) ein Prädikat

$VERIFIZIERE_G_TYP_3(w)$

$$= \begin{cases} \text{TRUE} & \text{falls } w \text{ die Kodierung einer Typ-3-Grammatik darstellt} \\ \text{FALSE} & \text{falls } w \text{ nicht die Kodierung einer Typ-3-Grammatik darstellt} \end{cases}$$

definiert. Dieses Prädikat ist berechenbar. Dazu ist unter anderem zu prüfen, ob die Erzeugungsregeln in KG_w den Bedingungen einer Typ-3-Grammatik genügen.

Das **Wortproblem für Typ-3-Grammatiken** fragt danach, ob die Menge

$$L_{Wort_Typ-3} = \left\{ u \# w \mid u \in \Sigma^*, w \in \{0,1\}^*, VERIFIZIERE_G_TYP_3(w) = \text{TRUE} \text{ und } u \in L(KG_w) \right\}$$

entscheidbar ist. In vereinfachter Darstellung wird also danach gefragt, ob es einen auf allen Eingaben stoppenden Algorithmus gibt, der bei Eingabe (der Kodierung) einer rechtslinearen Grammatik G über dem Alphabet Σ und eines Worts $u \in \Sigma^*$ entscheidet, ob $u \in L(G)$ gilt oder nicht. Da dieses Problem für Typ-2-Grammatiken entscheidbar ist, gilt dieses auch für Typ-3-Grammatiken. In diesem Fall bietet sich jedoch ein einfacher Algorithmus an: Anstelle der Überprüfung $u \in L(G)$ wird untersucht, ob $u \in L(DEA)$ gilt, wobei $DEA = (Q, \Sigma, d, q_0, F)$ ein deterministischer endlicher Automat mit $L(DEA) = L(G)$ ist:

Eingabe: Ein deterministischer endlicher Automat $DEA = (Q, \Sigma, d, q_0, F)$ und ein Wort $u \in \Sigma^*$

Verfahren: Aufruf der Funktion `accept (DEA, u)`

Ausgabe: TRUE, falls $u \in L(G)$, FALSE sonst.

```
FUNCTION accept (DEA : ...;
                 u : STRING) : BOOLEAN;
{ DEA = (Q, Σ, d, q0, F),
```

$$u = a_1 \dots a_n \quad \}$$

```

VAR q : ...;
    v : STRING;
    a : CHAR;
    i : INTEGER;

BEGIN { accept }
    q := q0;
    v := u;

    FOR i := 1 TO Length(u) DO
        BEGIN
            a := Copy (v, 1, 1);
            Delete (v, 1, 1);
            q := d(q,a);
        END;

    IF q ∈ F THEN accept := TRUE
        ELSE accept := FALSE;

END { accept };

```

Diese Algorithmus hat eine Laufzeit der Ordnung $O(|u|)$, d.h. lineare Laufzeit.

Das **Leerheitsproblem** fragt für **Typ-3-Grammatiken**, ob die Menge

$$L_{\text{leer_Typ-3}} = \left\{ w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_G_TYP_3}(w) = \text{TRUE und } L(KG_w) = \emptyset \right\}$$

entscheidbar ist. Diese Frage ist natürlich positiv zu beantworten, da das Leerheitsproblem für Typ-2-Grammatiken entscheidbar ist. Im Fall einer regulären L Menge kann man mit Satz 4.5-4 argumentieren. Es sei n_0 die in Satz 4.5-4 angegebene natürliche Zahl. Dann gilt:

$$L \neq \emptyset \text{ genau dann, wenn es ein Wort } u \in L \text{ gibt mit } |u| < n_0$$

Eine Überprüfung der Frage „ $L = \emptyset$?“ kann also so ablaufen, daß man alle Wörter $u \in \Sigma^*$ mit $|u| < n_0$ darauf hin überprüft, ob $u \in L$ gilt (Wortproblem).

4.6 Eigenschaften im tabellarischen Überblick

Die folgenden Tabellen stellen wichtige Eigenschaften der behandelten Sprachklassen zusammen. Nicht alle aufgeführten Eigenschaften wurden in den vorherigen Kapiteln bewiesen; es wird vielmehr auf die angegebene Literatur verwiesen.

Beschreibungsmittel	
Typ 0	Typ-0-Grammatik Turingmaschine
Typ 1	kontextsensitive Grammatik linear beschränkter Automat
Typ 2	kontextfreie Grammatik Kellerautomat
deterministisch kontextfrei	$LR(k)$ -Grammatik deterministischer Kellerautomat
Typ 3	regulärer Ausdruck, rechtslineare Grammatik endlicher Automat

Sind das nichtdeterministische und das deterministische Modell äquivalent?	
Turingmaschine	ja
Linear beschränkter Automat	?
Kellerautomat	nein
Endlicher Automat	ja

Abschlußeigenschaften					
Operation	Schnitt $L_1 \cap L_2$	Vereinigung $L_1 \cup L_2$	Komplement $\Sigma^* \setminus L$	Produkt $L_1 \cdot L_2$	Stern L^*
Typ 0	ja	ja	nein	ja	ja
Typ 1	ja	ja	ja	ja	ja
Typ 2	nein	ja	nein	ja	ja
det. kontextfrei	nein	nein	ja	nein	nein
Typ 3	ja	ja	ja	ja	ja

Wortproblem	
Typ 0	unlösbar
Typ 1	lösbar mit Komplexität der Ordnung $O(2^{O(n)})$
Typ 2 (bei gegebener kfr. Grammatik)	lösbar mit Komplexität der Ordnung $O(n^3)$
deterministisch kontextfrei	lösbar mit Komplexität der Ordnung $O(n)$
Typ 3	lösbar mit Komplexität der Ordnung $O(n)$

Leerheitsproblem	
Typ 0	unlösbar
Typ 1	unlösbar
Typ 2 (bei gegebener kfr. Grammatik)	lösbar
deterministisch kontextfrei	lösbar
Typ 3	lösbar

5 Praktische Berechenbarkeit

Im vorliegenden Kapitel geht es um die Lösung von Problemen „aus der Praxis“, insbesondere um die Untersuchung des Aufwandes, den diese Lösungen erfordern. **Alle hier behandelten (Entscheidungs-) Probleme sind entscheidbar, d.h. es gibt jeweils einen Algorithmus, der bei jeder Eingabe mit einer akzeptierenden oder verwerfenden Entscheidung stoppt.** Wie ein derartiger Algorithmus formuliert wird, ob als deterministische oder nichtdeterministische Turingmaschine, RAM oder Programm in einer Programmiersprache, ergibt sich jeweils aus dem Zusammenhang; das für die Darstellung geeignetste Modell wird verwendet.

Der Begriff der Entscheidbarkeit bezieht sich auf Entscheidungsprobleme, in der Praxis hat man jedoch häufig mit Optimierungsproblemen zu tun. Beispielsweise wurde in Kapitel 1.2 das Problem des Handlungsreisenden als Optimierungsproblem, Entscheidungsproblem und Berechnungsproblem beschrieben. Der Zusammenhang zwischen diesen Problemtypen wird in Kapitel 5.1 genauer betrachtet.

Im folgenden werden zwei Beispiele für Optimierungsprobleme angeführt, die sich auf Graphen beziehen. Für beide Probleme werden Lösungsalgorithmen angegeben, die sich in ihrer Laufzeit wesentlich unterscheiden.

Das Problem der Wege mit minimalem Gewicht in gerichteten Graphen

Instanz: $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter gerichteter Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

Lösung: Die minimalen Gewichte d_i der Wege vom Knoten v_1 zu allen anderen Knoten v_i des Graphs für $i = 1, \dots, n$.

Zur algorithmischen Behandlung wird der Graph $G = (V, E, w)$ in Form seiner **Adjazenzmatrix** $A(G)$ gespeichert. Diese ist definiert durch

$$A(G) = A_{(n,n)} = [a_{i,j}]_{(n,n)} \text{ mit } a_{i,j} = \begin{cases} w((v_i, v_j)) & \text{falls } (v_i, v_j) \in E \\ \infty & \text{sonst} \end{cases}.$$

Es werden folgende Datentypen verwendet:

```

CONST n          = ...;      { Problemgröße }
      unendlich = ...;      { hier kann der maximal mögliche
                              REAL-Wert verwendet werden      }

TYPE knotenbereich = 1..n;
   matrix          = ARRAY [knotenbereich, knotenbereich] OF REAL;
   bit_feld        = ARRAY [knotenbereich] OF BOOLEAN;
   feld            = ARRAY [knotenbereich] OF REAL;

```

Die folgende Funktion wird zur Ermittlung des Minimums zweier REAL-Zahlen eingesetzt:

```

FUNCTION min (a, b: REAL):REAL;

BEGIN { min }
  IF a <= b THEN min := a
    ELSE min := b;
END   { min };

```

Der folgende Algorithmus zur Lösung des Problems der Wege mit minimalem Gewicht in gerichteten Graphen ist ein Beispiel für ein allgemeines Lösungsprinzip: die **Greedy-Methode**. Eine (global) optimale Lösung wird schrittweise gefunden, indem Entscheidungen gefällt werden, die auf „lokalen“ Informationen beruhen. Für eine Entscheidung wird hierbei nur ein kleiner (lokaler) Teil der Informationen genutzt, die über das gesamte Problem zur Verfügung stehen. Es wird die in der jeweiligen Situation optimal erscheinende Entscheidung getroffen, unabhängig von vorhergehenden und nachfolgenden Entscheidungen. Die einmal getroffenen Entscheidungen werden im Laufe des Rechenprozesses nicht mehr revidiert. Diese Methode ist jedoch nicht auf alle Optimierungsprobleme anwendbar.

Zunächst gelten alle Knoten v_i bis auf den Knoten v_1 als „nicht behandelt“. Aus den nicht behandelten Knoten wird ein Knoten ausgewählt, der zu einer bisherigen Teillösung hinzugenommen wird, und zwar so, daß dadurch eine bzgl. der Teillösung, d.h. bzgl. der behandelten Knoten, optimale Lösung entsteht.

Eingabe: $G = (V, E, w)$ ein gewichteter gerichteter Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht

```

VAR A      : matrix;      { Adjazenzmatrix }
    i      : knotenbereich;
    j      : knotenbereich;
    distanz : feld;

```

```

FOR i := 1 TO n DO
  FOR j := 1 TO n DO
    IF  $(v_i, v_j) \in E$  THEN  $A[i, j] := w((v_i, v_j))$ 
    ELSE  $A[i, j] := \text{unendlich}$ ;

```

Verfahren: Aufruf der Prozedur `minimale_wege (A, distanz)`;

Ausgabe: $d_i = \text{distanz}[i]$ für $i = 1, \dots, n$.

```

PROCEDURE minimale_wege ( A          : matrix;
                          VAR distanz : feld);

```

```

VAR rest_knoten : bit_feld;      { Kennzeichnung der noch nicht
                                behandelten Knoten: ist  $v_i$  ein
                                nicht behandelter Knoten, so
                                ist  $\text{rest\_knoten}[i] = \text{TRUE}$       }

idx          : knoten_bereich;
u            : knoten_bereich;
exist       : BOOLEAN;

```

```

PROCEDURE auswahl (VAR knoten  : knoten_bereich;
                  VAR gefunden : BOOLEAN);

```

```

{ die Prozedur wählt unter den noch nicht behandelten
  Knoten einen Knoten mit minimalem distanz-Wert aus
  (in knoten); gibt es einen derartigen Knoten, dann ist
  gefunden = TRUE, sonst ist gefunden = FALSE      }

```

```

VAR i    : knoten_bereich;
    min  : REAL;

```

```

BEGIN { auswahl }
  gefunden := FALSE;
  min      := unendlich;

```

```

FOR i := 1 TO n DO
  IF rest_knoten [i]
  THEN BEGIN
    IF distanz [i] < min
    THEN BEGIN
      gefunden := TRUE;
      knoten   := i;
      min      := distanz [i]
    END
  END

```

```

END
END { auswahl };

```

```

BEGIN { minimale_wege }
  { Gewichte zum Startknoten initialisieren und alle Knoten,
    bis auf den Startknoten, als nicht behandelt kennzeichnen: }
  FOR idx := 1 TO n DO
    BEGIN
      distanz[idx]      := A[1, idx];
      rest_knoten[idx] := TRUE
    END;
  distanz[1]      := 0;
  rest_knoten[1] := FALSE;

  { einen Knoten mit minimalem Distanzwert aus den Restknoten
    auswählen: }
  auswahl (u, exist);

  WHILE exist DO
    BEGIN
      { den gefundenen Knoten aus den noch nicht behandelten
        Knoten herausnehmen: }
      rest_knoten[u] := FALSE;

      { für jeden noch nicht behandelten Knoten idx, der mit
        dem Knoten u verbunden ist, den distanz-Wert auf den
        neuesten Stand bringen: }
      FOR idx := 1 TO n DO
        IF rest_knoten[idx]
          AND
          (A[u, idx] <> unendlich)
          THEN distanz [idx] := min (distanz[idx],
                                   distanz[u] + A[u, idx]);

      { einen neuen Knoten mit minimalem Distanzwert aus den
        Restknoten auswählen: }
      auswahl (u, exist);

    END {WHILE exist }

  END { minimale_wege };

```

Es gilt bei Eintritt in die WHILE exist DO-Schleife folgende Schleifeninvariante:

- (1) Ist der Knoten v_i ein bereits behandelter Knoten, d.h. $\text{rest_knoten}[i]=\text{FALSE}$, dann ist $\text{distanz}[i]$ das minimale Gewicht eines Weges von v_1 nach v_i , der nur behandelte Knoten v_j enthält, für die also $\text{rest_knoten}[j]=\text{FALSE}$ ist.

u wird jeweils beim Aufruf der Prozedur `auswahl (u, exist)` ausgewählt:

u	nichtbehandelte Knoten	distanz [1]	distanz [2]	distanz [3]	distanz [4]	distanz [5]	distanz [6]	distanz [7]	distanz [8]
	2 3 4 5 6 7 8	---	28	2	∞	1	∞	∞	∞
5	2 3 4 6 7 8		9	2	∞	---	27	∞	∞
3	2 4 6 7 8		9	---	∞		26	∞	29
2	4 6 7 8		---		18		19	∞	∞
4	6 7 8				---		19	26	25
6	7 8						---	26	20
8	7							26	---
7	leer							---	

Das zweite Beispiel löst das Handlungsreisenden-Optimierungsproblem nach der sogenannten **Branch-and-Bound-Methode**. Potentielle, aber nicht unbedingt optimale Lösungen werden systematisch erzeugt. Diese werden in Teilmengen aufgeteilt, die sich auf Knoten eines Entscheidungsbaums abbilden lassen. Es wird eine Abschätzung für das Optimum mitgeführt und während des Verfahrensverlaufs laufend aktualisiert. Potentielle Lösungen, deren Zielfunktion einen „weit von der Abschätzung entfernten Wert“ aufweisen, werden nicht weiter betrachtet, d.h. der Teilbaum, der bei einer derartigen Lösung beginnt, muß nicht weiter durchlaufen werden.

Problem des Handlungsreisenden auf Graphen als Optimierungsproblem

Instanz: $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

Lösung: Eine **Tour** durch G , d.h. eine geschlossene Kantenfolge

$$\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle \text{ mit } (v_{i_j}, v_{i_{j+1}}) \in E \text{ für } j = 1, \dots, n-1,$$

in der jeder Knoten des Graphen (als Anfangsknoten einer Kante) genau einmal vorkommt, die minimale Kosten (unter allen möglichen Touren durch G) verursacht. Man kann o.B.d.A. $v_{i_1} = v_1$ setzen.

Die Kosten einer Tour $\langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$ ist dabei die Summe der Gewichte der Kanten auf der Tour, d.h. der Ausdruck

$$\sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1}).$$

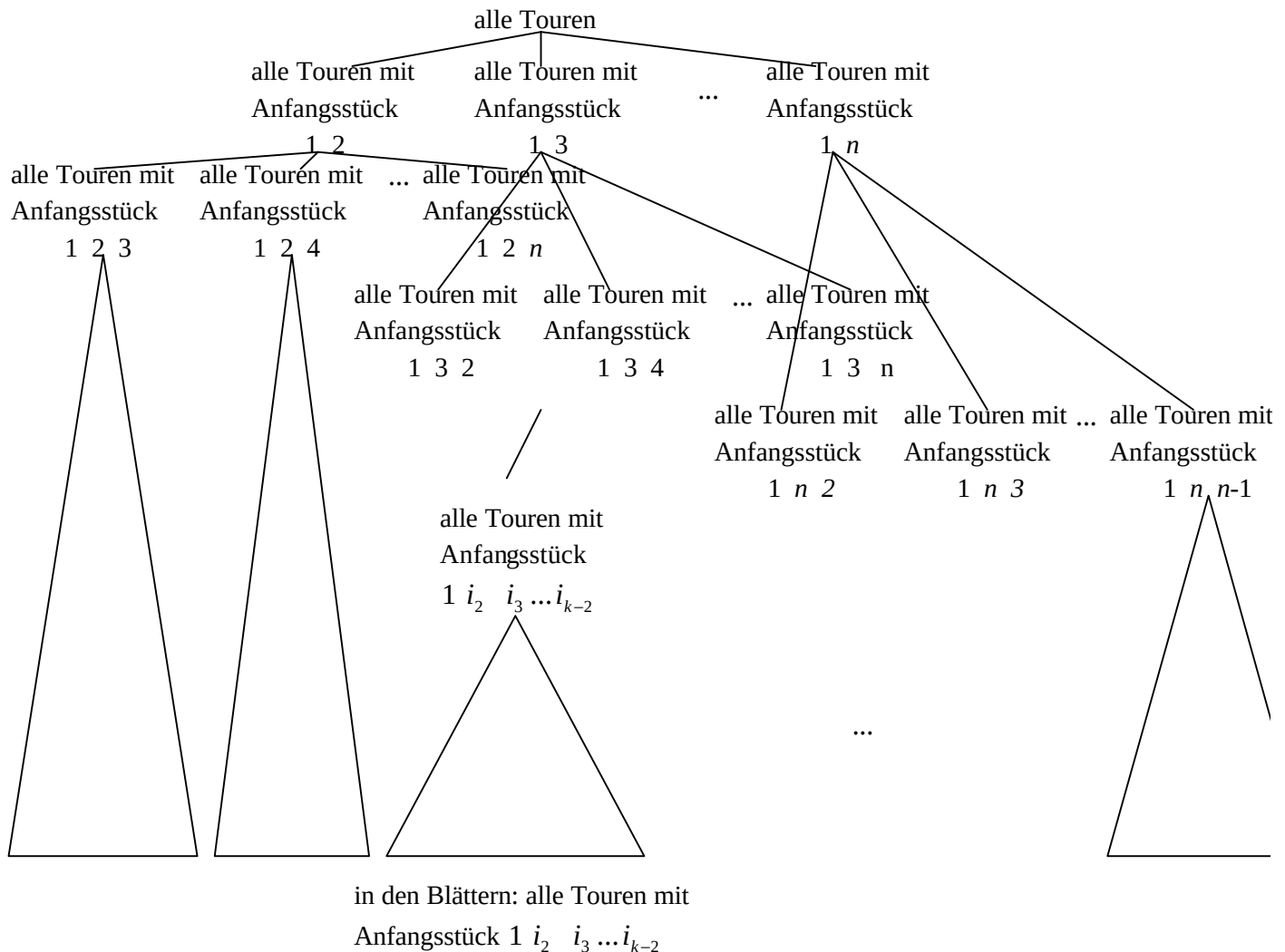
Es wird angenommen, daß eine Tour bei v_1 beginnt und endet. Außerdem wird angenommen, daß der Graph **vollständig** ist, d.h. daß $w((v_i, v_j))$ für jedes $v_i \in V$ und $v_j \in V$ definiert ist (eventuell ist $w((v_i, v_j)) = \infty$).

Im folgenden wird (zur Vereinfachung der Darstellung) die Knotenmenge $V = \{v_1, \dots, v_n\}$ mit der Menge der Zahlen $\{1, \dots, n\}$ gleichgesetzt (dem Knoten v_i entspricht die Zahl i). Eine Tour durch G läßt sich dann als Zahlenfolge

$$\langle 1 \ i_1 \ i_2 \ \dots \ i_{n-1} \ 1 \rangle$$

darstellen, wobei alle i_j paarweise verschieden (und verschieden von 1) sind. Alle Touren erhält man, wenn man für i_j in $\langle 1 \ i_1 \ i_2 \ \dots \ i_{n-1} \ 1 \rangle$ alle $(n-1)!$ Permutationen der Zahlen $2, \dots, n$ einsetzt. Manche dieser Touren haben eventuell das Gewicht ∞ .

Alle Touren (Zahlenfolgen der beschriebenen Art) lassen sich als Blätter eines Baums darstellen:



Alle Touren (Zahlenfolgen der beschriebenen Art) werden systematisch erzeugt (siehe unten). Dabei wird eine obere Abschätzung für das Gewicht einer optimalen Tour (Tour mit minimalem Gewicht) in der globalen Variablen bound mitgeführt (Anfangswert $\text{bound} = \infty$). Wird ein Anfangsstück einer Tour erzeugt, deren Gewicht größer als der Wert in der Variablen bound ist, dann brauchen alle Touren, die mit diesem Anfangsstück beginnen, nicht weiter betrachtet zu werden. Es wird also ein ganzer Teilbaum aus dem Baum aller Touren abgeschnitten. Ist schließlich eine Tour gefunden, deren Gewicht kleiner oder gleich dem Wert in der Variablen bound ist, so erhält die Variable bound als neuen Wert dieses Gewicht, und die Tour wird als temporär optimale Tour vermerkt. Eine Variable `opttour` nimmt dabei die Knotennummern einer temporär optimalen Tour auf; die Variable `teilstück` dient der Aufnahme des Anfangsstücks einer Tour.

Ist bereits ein Anfangsstück $\langle 1 \ i_1 \ i_2 \dots i_{k-1} \rangle$ einer Tour erzeugt, so sind die Zahlen $1, i_1, i_2, \dots, i_{k-1}$ alle paarweise verschieden. Eine weitere Zahl i kann in die Folge aufgenommen werden, wenn

- (1) i unter den Zahlen $1, i_1, i_2, \dots, i_{k-1}$ nicht vorkommt und
- (2) $w((v_{i_{k-1}}, v_i)) < \infty$ ist und
- (3) das Gewicht der neu entstehenden Teiltour $\langle 1, i_1, i_2, \dots, i_{k-1} \ i \rangle$ kleiner oder gleich dem Wert in der Variablen bound ist.

Treffen (1) oder (2) nicht zu, kann i nicht als Fortsetzung der Teiltour $1, i_1, i_2, \dots, i_{k-1}$ genommen werden, denn es entsteht keine Tour. Trifft (3) nicht zu (aber (1) und (2)), dann brauchen alle Touren, die mit dem Anfangsstück $\langle 1, i_1, i_2, \dots, i_{k-1} \ i \rangle$ beginnen, nicht weiter berücksichtigt zu werden.

Algorithmus zur Lösung des Problems des Handlungsreisenden nach der Branch-and-bound-Methode:

Es werden wieder die bereits bekannten Deklarationen (ergänzt um neue Deklarationen) verwendet:

```

CONST n          = ...;      { Problemgröße }
    unendlich = ...;      { hier kann der maximal mögliche
                           REAL-Wert verwendet werden      }

TYPE knotenbereich = 1..n;
    matrix          = ARRAY [knotenbereich, knotenbereich] OF REAL;
    tourfeld        = ARRAY [1..(n+1)] OF INTEGER;

```

Die Knotennummern einer optimalen Tour werden im Feld `opttour` abgelegt.

Eingabe: $G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht.

```

VAR A          : matrix;      { Adjazenzmatrix }
    i          : knotenbereich;
    j          : knotenbereich;
    opttour    : tourfeld;
    bound      : REAL;

```

```

FOR i := 1 TO n DO
  FOR j := 1 TO n DO
    IF  $(v_i, v_j) \in E$  THEN  $A[i, j] := w((v_i, v_j))$ 
    ELSE  $A[i, j] := \text{unendlich}$ ;

```

Verfahren: Aufruf der Prozedur salesman_bab (A, opttour, bound).

Ausgabe: bound gibt das minimale Gewicht einer Tour durch G an, die beim Knoten v_1 beginnt und endet; opttour enthält die Knotennummern einer optimalen Tour.

PROCEDURE salesman_bab

```

  (A          : matrix;
   VAR opttour : tourfeld; { optimale Tour      }
   VAR bound   : REAL      { Länge der Tour     }
  );

```

VAR i : INTEGER;

PROCEDURE bab_g

```

  (teilstück : tourfeld;
                                     { zu ergänzendes Anfangsstück
                                     einer Tour                          }
   gewicht   : REAL;   { Gewicht des Anfangsstücks }
   position  : INTEGER { Position, an der zu
                                     ergänzen ist          }
  );

```

```

VAR i : INTEGER;
    j : INTEGER;
    ok : BOOLEAN;
    w  : REAL;

```

BEGIN {bab_g }

IF position = n + 1

THEN BEGIN

 w := A[teilstück[n], 1];

 IF (w < unendlich)

 AND

 (gewicht + w < bound)

 THEN BEGIN

 teilstück[position] := 1;

 bound := gewicht + w;

 opttour := teilstück;

 END;

END

```

ELSE BEGIN
  FOR i := 2 TO n DO
    BEGIN
      w := A[teilstad[position - 1], i];
      ok := TRUE;
      { Bedingung (2): }
      FOR j := 2 TO position - 1 DO
        IF teilstad[j] = i
          THEN BEGIN
            ok := FALSE;
            Break;
          END;
      IF ok
        AND
        (w < unendlich)
        AND
        (gewicht + w < bound)
      THEN BEGIN
        teilstad[position] := i;
        bab_g (teilstad, gewicht + w,
              position + 1);
      END;
    END;
  END;

END {bab_g };

BEGIN { salesman_bab }

  FOR i := 2 TO n + 1 DO
    opttour[i] := 0;
    opttour[1] := 1;
    bound := unendlich;

    bab_g(opttour, 0, 2);

  END {salesman_bab };

```

Das Verfahren erzeugt u.U. alle $(n-1)!$ Folgen $\langle 1 i_1 i_2 \dots i_{n-1} 1 \rangle$, bis es eine optimale Tour gefunden hat. In der Praxis werden jedoch schnell Folgen mit Anfangsstücken, die ein zu großes Gewicht aufweisen, ausgeschlossen. Die (worst-case-) Zeitkomplexität des Verfahrens bleibt jedoch mindestens exponentiell in der Anzahl der Knoten des Graphen in der Eingabeinstanz.

In der Praxis gilt ein Algorithmus mit exponentiellem Laufzeitverhalten als **schwer durchführbar (intractable)**, ein Algorithmus mit polynomielltem Laufzeitverhalten als **leicht**

durchführbar (tractable). Natürlich kann dabei ein Algorithmus mit exponentiellem Laufzeitverhalten für einige Eingabeinstanzen schnell ablaufen. Trotzdem bleibt die Ursache zu untersuchen, die dazu führt, daß manche Probleme „leicht lösbar“ sind (d.h. einen Algorithmus mit polynomielltem Laufzeitverhalten besitzen), während für andere Probleme bisher nur Lösungsalgorithmen mit exponentiellem Laufzeitverhalten bekannt sind.

Der Grund dafür, daß ein Verfahren mit exponentielle Laufzeit als nicht durchführbar gilt, wird aus den folgenden Tabelle sichtbar. Es werden dort jeweils fünf Funktionen $h_i : \mathbf{N}_{>0} \rightarrow \mathbf{R}, i = 1, \dots, 5$ und einige ausgewählte (gerundete) Funktionswerte gezeigt. Jede Funktion soll als Laufzeit für Algorithmen interpretiert werden: Der Funktionswert $h_i(n)$ gibt die Anzahl der Rechenschritte an, die bei einer Eingabeinstanz der Größe n durchlaufen werden. Selbst kleine Problemgrößen führen bereits auf eine immense Laufzeit.

Spalte 1	Spalte 2	Spalte 3	Spalte 4	Spalte 5
i	$h_i(n)$	$h_i(10)$	$h_i(100)$	$h_i(1000)$
1	$\log_2(n)$	3,3219	6,6439	9,9658
2	$\sqrt{n}$	3,1623	10	31,6228
3	n	10	100	1000
4	n^2	100	10.000	1.000.000
5	2^n	1024	$1,2676506 \cdot 10^{30}$	$> 10^{693}$

Die folgende Tabelle zeigt noch einmal die fünf Funktionen $h_i : \mathbf{N}_{>0} \rightarrow \mathbf{R}, i = 1, \dots, 5$. Es sei $y_0 > 0$ ein fester Wert. Die dritte Spalte zeigt für jede der fünf Funktionen Werte n_i mit $h_i(n_i) = y_0$. In der vierten Spalte sind diejenigen Werte $\bar{n}_i$ aufgeführt, für die $h_i(\bar{n}_i) = 10 \cdot y_0$ gilt, d.h. dort ist angegeben, auf welchen Wert man die Problemgröße n_i vergrößern kann, wenn man die 10-fachen Laufzeit in Kauf nimmt. Wie man sieht, kann man bei der Logarithmusfunktion wegen ihres langsamen Wachstums den Wert stark vergrößern, während bei der schnell anwachsenden Exponentialfunktion nur eine additive konstante Steigerung um ca. 3,3 möglich ist. Das bedeutet, daß selbst bei einer Verzehnfachung der Rechenleistung nur eine geringfügig vergrößerte Problemgröße zu bewältigen ist.

Spalte 1	Spalte 2	Spalte 3	Spalte 4
i	$h_i(n)$	n_i mit $h_i(n_i) = y_0$	$\bar{n}_i$ mit $h_i(\bar{n}_i) = 10 \cdot y_0$
1	$\log_2(n)$	n_1	$(n_1)^{10}$
2	$\sqrt{n}$	n_2	$100 \cdot n_2$
3	n	n_3	$10 \cdot n_3$
4	n^2	n_4	$\approx 3,162 \cdot n_4$
5	2^n	n_5	$n_5 + 3,322$

5.1 Der Zusammenhang zwischen Optimierungs- und Entscheidungsproblemen

Optimierungsprobleme spielen in der Anwendung eine bedeutende Rolle; in der Theorie der Berechenbarkeit sind es eher die Entscheidungsprobleme. Zwischen beiden Typen besteht jedoch ein enger Zusammenhang.

Es sei Π ein Optimierungsproblem mit einer reellwertigen Zielfunktion:

Instanz: 1. $x \in \Sigma_{\Pi}^*$

2. Spezifikation einer Funktion SOL_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ eine Menge zulässiger Lösungen zuordnet
3. Spezifikation einer Zielfunktion m_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ und $y \in \text{SOL}_{\Pi}(x)$ einen Wert $m_{\Pi}(x, y)$, den Wert einer zulässigen Lösung, zuordnet
4. $\text{goal}_{\Pi} \in \{\min, \max\}$, je nachdem, ob es sich um ein Minimierungs- oder ein Maximierungsproblem handelt.

Lösung: $y^* \in \text{SOL}_{\Pi}(x)$ mit $m_{\Pi}(x, y^*) = \min\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Minimierungsproblem (d.h. $\text{goal}_{\Pi} = \min$)

bzw.

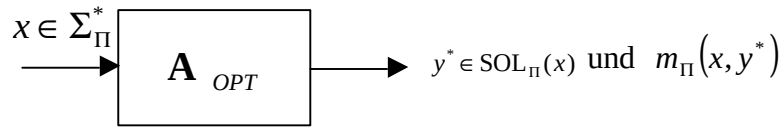
$m_{\Pi}(x, y^*) = \max\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Maximierungsproblem (d.h. $\text{goal}_{\Pi} = \max$).

Der Wert $m_{\Pi}(x, y^*)$ einer optimalen Lösung wird auch mit $m_{\Pi}^*(x)$ bezeichnet.

Eine Instanz x ist eine Zeichenkette $x \in \Sigma_{\Pi}^*$. Hier wird das Alphabet Σ_{Π} für das Problem „geeignet“ gewählt. Für ein Graphenproblem ist sicherlich dabei ein anderes Alphabet geeignet

als zur Darstellung Boolescher Ausdrücke. Letztlich kann man sich natürlich auf ein allgemein gültiges „Grundalphabet“ verständigen, etwa $\Sigma_0 = \{0, 1\}$.

Ein Lösungsalgorithmus $\mathbf{A}_{OPT}$ für Π hat die Form



Die von $\mathbf{A}_{OPT}$ bei Eingabe von $x \in \Sigma_{\Pi}^*$ ermittelte Lösung ist $\mathbf{A}_{OPT}(x) = (y^*, m_{\Pi}(x, y^*))$, d.h. sie besteht aus einer optimalen Lösung y^* und dem Wert der Zielfunktion $m_{\Pi}(x, y^*) = m_{\Pi}^*(x)$ bei einer optimalen Lösung. Natürlich kann es für ein Optimierungsproblem mehrere optimale Lösungen geben; der Wert $m_{\Pi}^*(x)$ ist jedoch eindeutig.

Das zu Π **zugehörige Entscheidungsproblem** Π_{ENT} wird definiert durch

Instanz: $[x, K]$

mit $x \in \Sigma_{\Pi}^*$ und $K \in \mathbf{R}$

Lösung: Entscheidung „ja“, falls für den Wert $m_{\Pi}^*(x)$ einer optimalen Lösung

$m_{\Pi}^*(x) \geq K$ bei einem Maximierungsproblem

(d.h. $goal_{\Pi} = \max$)

bzw.

$m_{\Pi}^*(x) \leq K$ bei einem Minimierungsproblem

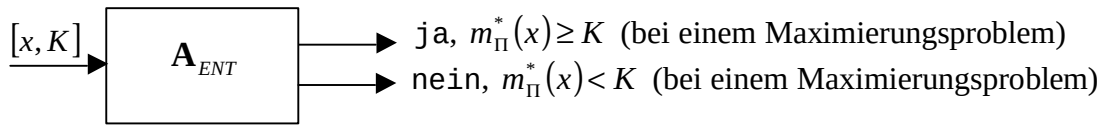
(d.h. $goal_{\Pi} = \min$)

gilt,

Entscheidung „nein“, sonst.

Hierbei ist zu beachten, daß beim Entscheidungsproblem weder nach einer optimalen Lösung noch nach dem Wert $m_{\Pi}^*(x)$ der Zielfunktion bei einer optimalen Lösung gefragt wird.

Ein Lösungsalgorithmus für Π_{ENT} hat zwei mögliche Ausgänge, nämlich einen akzeptierenden und einen verwerfenden Ausgang. Der erste Ausgang wird erreicht (aktiviert), wenn bei Eingabe einer Instanz für Π_{ENT} die Entscheidung „ja“ getroffen wird, der zweite Ausgang entspricht der Entscheidung „nein“. Ein Lösungsalgorithmus $\mathbf{A}_{ENT}$ für Π_{ENT} , hier für ein Maximierungsproblem dargestellt, hat daher die Form



Aus der Kenntnis eines Algorithmus $\mathbf{A}_{OPT}$ für ein Optimierungsproblem läßt sich leicht ein Algorithmus $\mathbf{A}_{ENT}$ für das zugehörige Entscheidungsproblem konstruieren, der im wesentlichen dieselbe Komplexität besitzt:

Eingabe: $[x, K]$

mit $x \in \Sigma_{\Pi}^*$ und $K \in \mathbf{R}$

Verfahren: Man berechne $\mathbf{A}_{OPT}(x) = (y^*, m_{\Pi}(x, y^*))$ und vergleiche das Resultat $m_{\Pi}(x, y^*) = m_{\Pi}^*(x)$ mit K :

Ausgabe: $\mathbf{A}_{ENT}([x, K]) = \text{ja}$, falls $m_{\Pi}^*(x) \geq K$ bei einem Maximierungsproblem ist bzw.

$\mathbf{A}_{ENT}([x, K]) = \text{ja}$, falls $m_{\Pi}^*(x) \leq K$ bei einem Minimierungsproblem ist,

$\mathbf{A}_{ENT}([x, K]) = \text{nein}$ sonst.

Daher gilt:

Satz 5.1-1:

Das zu einem Optimierungsproblem Π zugehörige Entscheidungsproblem Π_{ENT} ist im wesentlichen *nicht schwieriger zu lösen* als das Optimierungsproblem.

Falls man andererseits bereits weiß, daß das Entscheidungsproblem Π_{ENT} immer nur „schwer lösbar“ ist, z.B. beweisbar nur Lösungsverfahren mit exponentiellem Laufzeitverhalten besitzt, dann ist das Optimierungsproblem ebenfalls nur „schwer lösbar“.

In den folgenden Kapiteln wird für einige *Entscheidungsprobleme* Π_{ENT} gezeigt, daß sie (vermutlich) keine schnell laufenden Lösungsverfahren besitzen. Daher können die zugehörigen Optimierungsprobleme, an deren Lösung man in der Praxis interessiert ist, ebenfalls nur mit sehr lang laufenden Verfahren angegangen werden. Die algorithmische Untersuchung von Entscheidungsproblemen ist in diesen Fällen daher für die Praxis von großem Interesse. Man hat zudem abzuwägen, ob man nicht mit einer Lösung zufrieden sein kann, die „schnell“ zu

finden ist und sich der optimalen Lösung annähert. Diese Fragestellung führt auf das Gebiet der **Approximationsalgorithmen** in Kapitel 6.

In einigen Fällen läßt sich auch die „umgekehrte“ Argumentationsrichtung zeigen: Aus einem Algorithmus $\mathbf{A}_{ENT}$ für das zu einem Optimierungsproblem zugehörige Entscheidungsproblem läßt sich ein Algorithmus $\mathbf{A}_{OPT}$ für das Optimierungsproblem konstruieren, dessen Laufzeitverhalten im wesentlichen dieselbe Komplexität aufweist. Insbesondere ist man dabei an Optimierungsproblemen interessiert, für die gilt: Hat $\mathbf{A}_{ENT}$ zur Lösung des zu einem Optimierungsproblem zugehörigen Entscheidungsproblem polynomielles Laufzeitverhalten (in der Größe der Eingabe), so hat auch der aus $\mathbf{A}_{ENT}$ konstruierte Algorithmus $\mathbf{A}_{OPT}$ zur Lösung des Optimierungsproblems polynomielles Laufzeitverhalten.

Ein Beispiel ist das Problem des Handlungsreisenden auf Graphen mit natürlichzahligen Gewichten:

Problem des Handlungsreisenden auf Graphen mit ganzzahligen Gewichten als Optimierungsproblem

Instanz: 1. $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{N}$ gibt jeder Kante $e \in E$ ein **natürlichzahliges Gewicht**; es ist $w((i, j)) = \infty$ für $(i, j) \notin E$

Als Problemgröße wird hier der Wert $k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$ genommen, da in das Verfahren wesentlich die numerischen Größen der beteiligten Kantengewichte eingeht

2. $\text{SOL}(G) = \{T \mid T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle \text{ ist eine Tour durch } G\}$

3. für $T \in \text{SOL}(G)$, $T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$, ist die Zielfunktion

$$\text{definiert durch } m(G, T) = \sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1})$$

4. $\text{goal} = \min$

Lösung: Eine Tour $T^* \in \text{SOL}(G)$, die minimale Kosten (unter allen möglichen Touren durch G) verursacht, und $m(G, T^*)$.

Ein Algorithmus, der bei Eingabe einer Instanz $G = (V, E, w)$ eine optimale Lösung erzeugt, werde mit $\mathbf{A}_{TSP_{OPT}}$ bezeichnet. Die von $\mathbf{A}_{TSP_{OPT}}$ bei Eingabe von G ermittelte Tour mit minimalen Kosten sei T^* , die ermittelten minimalen Kosten $m^*(G)$:

$$\mathbf{A}_{TSP_{OPT}}(G) = (T^*, m^*(G)).$$

Das zugehörige Entscheidungsproblem lautet:

Problem des Handlungsreisenden auf Graphen mit ganzzahligen Gewichten als Entscheidungsproblem

Instanz: $[G, K]$

$G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbb{N}$ gibt jeder Kante $e \in E$ ein natürlichzahliges Gewicht; es ist $w((i, j)) = \infty$ für $(i, j) \notin E$;

$K \in \mathbb{N}$.

Die Problemgröße ist $k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$.

Lösung: Entscheidung „ja“, falls es eine Tour durch G gibt mit minimalen Kosten $m^*(G) \leq K$ gibt,

Entscheidung „nein“, sonst.

Ein Algorithmus, der bei Eingabe einer Instanz $[G, K]$ eine entsprechende Entscheidung fällt, werde mit $\mathbf{A}_{TSP_{ENT}}$ bezeichnet. Die von $\mathbf{A}_{TSP_{ENT}}$ bei Eingabe von $[G, K]$ getroffene ja/nein-Entscheidung sei $\mathbf{A}_{TSP_{ENT}}([G, K])$. Mit Hilfe von $\mathbf{A}_{TSP_{ENT}}$ wird ein Algorithmus $\mathbf{A}_{TSP_{OPT}}$ zur Lösung des Optimierungsproblem konstruiert. Dabei wird wiederholt Binärsuche eingesetzt, um zunächst die Kosten einer optimalen Tour zu ermitteln.

Der Algorithmus $\mathbf{A}_{TSP_{OPT}}$ zur Lösung des Optimierungsproblem verwendet eine als Pseudocode formulierte Prozedur $\mathbf{A}_{TSP_{OPT}}$. Der Eingabegraph kann wieder in Form seiner Adjazenzmatrix verarbeitet werden. Da Implementierungsdetails hier nicht weiter betrachtet werden sollen, können wir annehmen, daß es zur Darstellung eines gewichteten Graphen einen geeigneten Datentyp

TYPE gewichteter_Graph = ...;

und zur Speicherung einer Tour (Knoten und Kanten) in dem Graphen einen Datentyp

TYPE tour = ...;

gibt. Der Algorithmus $\mathbf{A}_{TSP_{OPT}}$ hat folgendes Aussehen:

Eingabe: $G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{N}$ gibt jeder Kante $e \in E$ ein natürlichzahliges Gewicht; es ist $w((i, j)) = \infty$ für $(i, j) \notin E$

```
VAR G      : gewichter_Graph;
    opttour : tour;
    opt     : INTEGER;
```

```
G := G = (V, E, w)
```

Verfahren: Aufruf der Prozedur A_TSP0PT (G, opt, opttour).

Ausgabe: opttour = T^* , d.h. eine Tour mit minimalen Kosten, opt = die ermittelten minimalen Kosten $m^*(G)$.

```
PROCEDURE A_TSP0PT (G      : gewichter_Graph;
                    { Graph mit natürlichzahligen Kantengewichten }
                    VAR opt : INTEGER;
                    { Gewicht einer optimalen Tour }
                    VAR opttour : graph;
                    { ermittelte Tour mit minimalem Gewicht }
                    )
```

```
VAR s : INTEGER;
```

```
PROCEDURE TSP0PT (G      : gewichter_graph;
                  VAR opt : INTEGER;)
{ ermittelt im Parameter opt das Gewicht einer
  optimalen Tour in G : }
```

```
VAR min : INTEGER;
    max : INTEGER;
    t   : INTEGER;
```

```
BEGIN { TSP0PT }
```

```
min := 0;
```

```
max :=  $\sum_{e \in E} w(e)$ ;
```

```
{ Binärsuche auf dem Intervall [0..max]: }
```

```
WHILE max - min >= 1 DO
```

```
  BEGIN
```

```

    t :=  $\left\lceil \frac{\min + \max}{2} \right\rceil$ ;
    IF  $A_{TSPENT}([G, t]) = \text{„ja“}$ 
    THEN max := t
    ELSE min := t + 1;
    END;

    { t enthält nun das Gewicht einer optimalen Tour in G: }
    opt := t;
    END { TSPOPT };

BEGIN { A_TSPOPT }
    TSPOPT (G, opt);

    FOR alle  $e \in E$  DO
        BEGIN
            ersetze in G das Gewicht  $w(e)$  der Kante  $e$  durch  $w(e) + 1$ ;
            TSPOPT (G, s);
            IF  $s > \text{opt}$ 
            THEN { es gibt keine optimale Tour in G, die  $e$  nicht enthält }
                ersetze in G das Gewicht  $w(e)$  der Kante  $e$  durch  $w(e) - 1$ ;
            END;

            { alle Kanten mit nicht erhöhten Gewichten bestimmen eine Tour
              mit minimalen Kosten }
            tour := Menge der Kanten mit nicht erhöhten Gewichten und zugehörige Knoten;
        END { A_TSPOPT };

```

Der Graph G habe die Problemgröße $k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$. Es sei $m = \sum_{e \in E} w(e)$. Dann sind in obigem Verfahren höchstens $\lceil \log_2(m) \rceil$ viele Aufrufe des Ent-

scheidungsverfahrens $A_{TSPENT}([G, t])$ erforderlich (Binärsuche). Es gilt

$$\lceil \log_2(m) \rceil \leq \log_2(m) + 1 \leq \sum_{e \in E} \log_2(w(e)) + 1 \leq \sum_{e \in E} \lceil \log_2(w(e)) \rceil + 1 \leq n^2 \cdot B + 1 \in O(k^2),$$

d.h. $\lceil \log_2(m) \rceil \in O(k^2)$.

Falls man also weiß, daß A_{TSPENT} ein in der Größe der Eingabe polynomiell zeitbeschränkter Algorithmus ist, ist das gesamte Verfahren zur Bestimmung einer optimalen Lösung polynomiell zeitbeschränkt.

Falls man umgekehrt weiß, daß es beweisbar keinen schnellen Algorithmus zur Lösung des Problems des Handlungsreisenden auf Graphen mit ganzzahligen Gewichten als Optimierungsproblem gibt, gibt es einen solchen auch nicht für das zugehörige Entscheidungspro-

blem. In diesem Fall ist das Entscheidungsproblem genauso schwer zu lösen wie das Optimierungsproblem.

5.2 Komplexitätsklassen

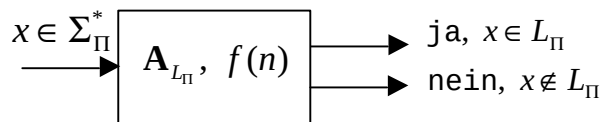
Eine Eingabeinstanz $G = (V, E, w)$ des Handlungsreisenden-Minimierungsproblem besteht aus natürlichen Zahlen für die Knotennummern, Paare natürlicher Zahlen für die Kanten mit ihren Gewichten und trennenden Klammern und dem Kommazeichen. Die Zahlen lassen sich im Binärdarstellung angeben, so daß man als Alphabet zur Formulierung der Eingabeinstanzen die Menge $\Sigma_{\Pi} = \{0, 1, (,)\} \cup \text{Kommazeichen}$ wählen kann. Als Problemgröße der Eingabeinstanz wird in Kapitel 5.1 der Wert $k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$ genommen, d.h. ein Wert, der proportional zu der Anzahl an Zeichen ist, um eine Eingabeinstanz zu kodieren. Dieser Ansatz berücksichtigt, daß für die Berechnung der Zeitkomplexität eines Algorithmus in der Regel das uniforme Kostenkriterium zugrunde gelegt wird, die Laufzeit aber in einigen Fällen von der numerischen Größe der Eingabe abhängt, so daß eigentlich das logarithmische Kostenkriterium angewendet werden müßte. Dadurch daß eine numerische Eingabe in der Anzahl der Zeichen, um sie darzustellen, gemessen wird, ist implizit die Laufzeitberechnung nach dem korrekt anzuwendenden Kostenkriterium gegeben.

Im folgenden wird (bei allgemeinen Betrachtungen) als **Größe einer Eingabeinstanz** $x \in \Sigma_{\Pi}^*$ für ein Entscheidungsproblem Π mit der Menge $L_{\Pi} \subseteq \Sigma_{\Pi}^*$ die Anzahl der Zeichen in x genommen. Die Laufzeit eines Entscheidungsalgorithmus wird dann in Abhängigkeit von $n = |x|$ beschrieben.

Für ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ gilt $L_\Pi \in TIME(f(n))$, wenn es einen (deterministischen) Algorithmus A_{L_Π} folgender Form gibt:

Eingabe: $x \in \Sigma_\Pi^*$ mit $|x| = n$

Ausgabe: ja, falls $x \in L_\Pi$ gilt
nein, falls $x \notin L_\Pi$ gilt.



Hierbei wird die Entscheidung $A_{L_\Pi}(x)$ nach einer Anzahl von Schritten getroffen, die in $O(f(n))$ liegt.

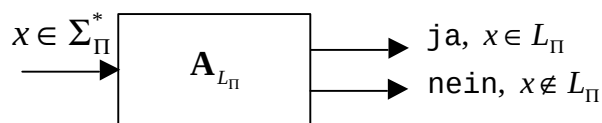
Die Aussage „das Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in $TIME(f(n))$ “ steht synonym für $L_\Pi \in TIME(f(n))$; man sagt auch abkürzend „das Entscheidungsproblem Π liegt in $TIME(f(n))$ “. In diesem Sinn bezeichnet $TIME(f(n))$ die **Klasse der Entscheidungsprobleme, die in der Zeitkomplexität $O(f(n))$ gelöst werden können**.

Eine entsprechende Definition gilt für die Speicherplatzkomplexität:

Für ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ gilt $L_\Pi \in SPACE(f(n))$, wenn es einen Algorithmus A_{L_Π} folgender Form gibt:

Eingabe: $x \in \Sigma_\Pi^*$ mit $|x| = n$

Ausgabe: ja, falls $x \in L_\Pi$ gilt
nein, falls $x \notin L_\Pi$ gilt.



Hierbei werden zur Findung der Entscheidung $A_{L_\Pi}(x)$ eine Anzahl von Speicherzellen verwendet, die in $O(f(n))$ liegt.

Die Aussage „das Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in $SPACE(f(n))$ “ steht synonym für $L_\Pi \in SPACE(f(n))$; man sagt auch „das Entscheidungsproblem Π liegt in $SPACE(f(n))$ “. In diesem Sinn bezeichnet $SPACE(f(n))$ die **Klasse der Entscheidungsprobleme, die mit Speicherplatzkomplexität $O(f(n))$ gelöst werden können**.

Wichtige Komplexitätsklassen sind

- die Klasse der Entscheidungsprobleme, die in einer Zeit gelöst werden können, die proportional zu einem Polynom in der Größe der Eingabe ist:

$$\mathbf{P} = \bigcup_{k=0}^{\infty} TIME(n^k)$$

Beispielsweise gibt es für jede von einer kontextfreien Grammatik G erzeugten Sprache $L(G) \subseteq \Sigma^*$ ein Entscheidungsverfahren, das für ein Wort $u \in \Sigma^*$ mit $|u| = n$ in $O(n^3)$ vielen Schritten entscheidet, ob $u \in L(G)$ gilt oder nicht (vgl. Kapitel 4.4). Daher ist die Klasse der kontextfreien Sprachen in $\mathbf{P}$ enthalten. Diese Inklusion ist echt, wie das Beispiel der Sprache $L_4 = \{a^n b^n c^n \mid n \in \mathbf{N}, n \geq 1\}$ aus Kapitel 4.1 zeigt, die in $\mathbf{P}$ liegt, aber nicht kontextfrei ist.

- die Klasse der Entscheidungsprobleme, die in mit einem Speicherplatzbedarf gelöst werden können, der proportional zu einem Polynom in der Größe der Eingabe ist:

$$\mathbf{PSPACE} = \bigcup_{k=0}^{\infty} SPACE(n^k)$$

- die Klasse der Entscheidungsprobleme, die in einer Zeit gelöst werden können, die proportional zu einer Exponentialfunktion in der Größe der Eingabe ist:

$$\mathbf{EXP} = \bigcup_{k=0}^{\infty} TIME(2^{n^k}).$$

Es gilt $\mathbf{P} \subseteq \mathbf{PSPACE} \subseteq \mathbf{EXP}$. Die Frage, ob eine dieser Inklusionen echt ist, d.h. ob $\mathbf{P} \subset \mathbf{PSPACE}$ oder $\mathbf{PSPACE} \subset \mathbf{EXP}$ gilt, ist ein bisher ungelöstes Problem der Komplexitätstheorie. Man weiß jedoch, daß $\mathbf{P} \subset \mathbf{EXP}$ gilt.

In Kapitel 2.5 wird eine nichtdeterministische Turingmaschinen $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ in zunächst unterschiedlichen Modell-Sichtweisen definiert. Entweder wird die Überföhrungsfunktion als partielle Funktion der Form $d : Q \times \Sigma^k \rightarrow \mathbf{P}(Q \times (\Sigma \times \{L, R, S\})^k)$ angegeben; während der Berechnung kann dann die Turingmaschine nichtdeterministische Schritte

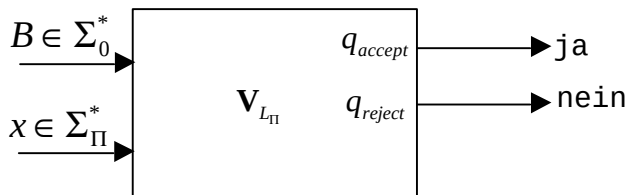
ausführen, indem sie bei Erreichen einer Konfiguration aus mehreren möglichen Folgekonfigurationen „die richtige“ auswählt. Oder die nichtdeterministische Turingmaschine wird als deterministische Turingmaschine gesehen, die mit einem „Ratemodul“ ausgestattet ist, das in nichtdeterministischer Weise zunächst eine Zusatzinformation, deren Brauchbarkeit anschließend mit Hilfe einer Überföhrungsfunktion der Form $d : Q \times \Sigma^k \rightarrow Q \times (\Sigma \times \{L, R, S\})^k$ deterministisch verifiziert wird. Diese Überlegung führte schließlich auf die Definition eines Verifizierers. Die Definition wird hier noch einmal wiederholt. Da hier nur entscheidbare Probleme behandelt werden, stoppt hier der Verifizierer bei allen Eingaben.

Ein **Verifizierer für das Entscheidungsproblem** Π über einem Alphabet Σ_Π ist ein (deterministischer) Algorithmus V_{L_Π} , der eine Eingabe $x \in \Sigma_\Pi^*$ und eine **Zusatzinformation** (einen **Beweis**) $B \in \Sigma_0^*$ lesen kann und die Frage „ $x \in L_\Pi$?“ mit Hilfe des Beweises B entscheidet. Die Bereitstellung des Beweises B ist nicht Aufgabe des Verifizierers; B wird von außen vorgegeben.

Ein **Verifizierer** V_{L_Π} für das Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ hat die Schnittstellen und die Form

Eingabe: $x \in \Sigma_\Pi^*, B \in \Sigma_0^*$

Ausgabe: ja (accept)
nein (reject)



Mit $V_{L_\Pi}(x, B), V_{L_\Pi}(x, B) \in \{\text{ja}, \text{nein}\}$ wird die Entscheidung von V_{L_Π} bezeichnet.

Für $x \in \Sigma_\Pi^*$ gilt die Entscheidung:

$x \in L_\Pi$, falls es einen Beweis $B_x \in \Sigma_0^*$ gibt, so daß V_{L_Π} bei Eingabe von x und B_x mit $V_{L_\Pi}(x, B_x) = \text{ja}$ stoppt;

$x \notin L_\Pi$, falls für alle Beweise $B \in \Sigma_0^*$ gilt: entweder stoppt V_{L_Π} bei Eingabe von x und B nicht, oder V_{L_Π} stoppt mit $V_{L_\Pi}(x, B) = \text{nein}$.

Der Beweis $B \in \Sigma_0^*$, den der Algorithmus V_{L_Π} zusammen mit der Eingabe $x \in \Sigma_\Pi^*$ liest, wird **auf nichtdeterministische Weise** vorgegeben. Wie er zu konstruieren ist, liegt außerhalb des Algorithmus V_{L_Π} . Der Verifizierer V_{L_Π} realisiert daher einen **nichtdeterministischen Algorithmus zur Akzeptanz (Erkennung, Entscheidung)** der Menge $L_\Pi \subseteq \Sigma_\Pi^*$.

Es sei $f: \mathbb{N} \rightarrow \mathbb{N}$ eine Funktion. Falls der nichtdeterministische Algorithmus mit Hilfe des Verifizierers V_{L_Π} die ja/nein-Entscheidung bei Eingabe von $x \in \Sigma_\Pi^*$ mit $|x| = n$ nach einer Anzahl von Schritten trifft, die in $O(f(n))$ liegt, so ist der Algorithmus (der Verifizierer V_{L_Π}) **$f(n)$ -zeitbeschränkt**.

Zu beachten ist, daß die Laufzeit des nichtdeterministischen Algorithmus bzw. die Laufzeit des Verifizierers V_{L_Π} in Abhängigkeit von der Größe von $x \in \Sigma_\Pi^*$ gemessen wird. Natürlich ist im Zeitaufwand, den ein Verifizierer für seine Entscheidung benötigt, auch die Zeit enthalten, um die Zeichen des Beweises B zu lesen. Ist dieser $f(n)$ -zeitbeschränkt, so kann er höchstens $C \cdot f(n)$ viele Zeichen des Beweises lesen (hierbei ist C eine Konstante). Wenn es also überhaupt einen Beweis B_x mit $V_{L_\Pi}(x, B_x) = \text{ja}$ gibt, dann gibt es auch einen Beweis mit Länge $\leq C \cdot f(n)$, so daß man für den Beweis gleich eine durch $C \cdot f(n)$ beschränkte Länge annehmen kann.

Für ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ gilt $L_\Pi \in \text{NTIME}(f(n))$, wenn es einen nichtdeterministischen $f(n)$ -zeitbeschränkten Algorithmus (einen Verifizierer) zur Akzeptanz von L_Π gibt.

Die Aussage „das Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in $\text{NTIME}(f(n))$ “ steht synonym für $L_\Pi \in \text{NTIME}(f(n))$. In diesem Sinn bezeichnet $\text{NTIME}(f(n))$ die **Klasse der Entscheidungsprobleme, die auf nichtdeterministische Weise in der Zeitkomplexität $O(f(n))$ gelöst werden können**.

Eine der wichtigsten Klassen, die Klasse **NP** der Entscheidungsprobleme, die nichtdeterministisch mit polynomieller Zeitkomplexität gelöst werden können, ergibt sich, wenn f ein Polynom ist:

$$\mathbf{NP} = \bigcup_{k=0}^{\infty} \text{NTIME}(n^k)$$

In Kapitel 2.5 wird gezeigt, daß eine $T(n)$ -zeitbeschränkte nichtdeterministische Turingmaschine durch eine $O(c^{T(n)})$ -zeitbeschränkte deterministische Turingmaschine simuliert werden

kann. Setzt man $T(n)=n^k$, so sieht man unmittelbar, daß $NTIME(n^k) \subseteq TIME(2^{n^k})$ gilt. Daraus folgt $\mathbf{NP} = \bigcup_{k=0}^{\infty} NTIME(n^k) \subseteq \bigcup_{k=0}^{\infty} TIME(2^{n^k}) = \mathbf{EXP}$.

Da ein polynomiell zeitbeschränkter Algorithmus $\mathbf{A}_{L_{\Pi}}$, wie er in der Definition von $\mathbf{P}$ vorkommt, ein spezieller polynomiell zeitbeschränkter Verifizierer ist, der für seine Entscheidung ohne die Zuhilfenahme eines Beweises auskommt, gilt

$$\mathbf{P} \subseteq \mathbf{NP}.$$

Insgesamt ergibt sich $\mathbf{P} \subseteq \mathbf{NP} \subseteq \mathbf{EXP}$. Zusammen mit $\mathbf{P} \subset \mathbf{EXP}$ fragt sich, welche der Inklusionen echt ist. Die dabei wichtigste Frage $\mathbf{P} = \mathbf{NP}$ bzw. $\mathbf{P} \neq \mathbf{NP}$ ist bisher ungelöst (**P-NP-Problem**). Vieles spricht dafür, daß $\mathbf{P} \neq \mathbf{NP}$ gilt.

5.3 Die Klassen P und NP

Das vorliegende Kapitel behandelt Beispiele aus den Klassen $\mathbf{P}$ und $\mathbf{NP}$.

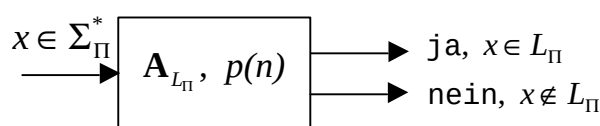
Zur Verdeutlichung des Unterschieds zwischen $\mathbf{P}$ und $\mathbf{NP}$ werden die entsprechenden Definitionen der beiden Klassen noch einmal gegenübergestellt:

Ein Entscheidungsproblem Π mit der Menge $L_{\Pi} \subseteq \Sigma_{\Pi}^*$ liegt in $\mathbf{P}$, wenn es einen Algorithmus $\mathbf{A}_{L_{\Pi}}$ und ein Polynom $p(n)$ gibt, so daß ein Entscheidungsverfahren für Π folgende Form hat:

Eingabe: $x \in \Sigma_{\Pi}^*$ mit $|x| = n$

Ausgabe: Entscheidung „ $x \in L_{\Pi}$ “, falls $\mathbf{A}_{L_{\Pi}}(x) = \text{ja}$ gilt,
 Entscheidung „ $x \notin L_{\Pi}$ “, falls $\mathbf{A}_{L_{\Pi}}(x) = \text{nein}$ gilt.

Hierbei wird die Entscheidung $\mathbf{A}_{L_{\Pi}}(x)$ nach einer Anzahl von Schritten getroffen, die in $O(p(n))$ liegt.



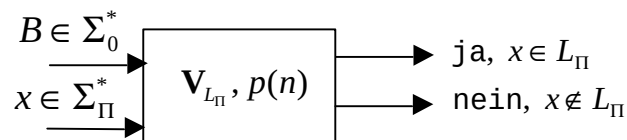
Ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in **NP**, wenn es einen Verifizierer V_{L_Π} und ein Polynom $p(n)$ gibt, so daß ein Entscheidungsverfahren für Π folgende Form hat:

Eingabe: $x \in \Sigma_\Pi^*$ mit $|x| = n$

Ausgabe: Entscheidung „ $x \in L_\Pi$ “, falls es einen Beweis $B_x \in \Sigma_0^*$ gibt mit $|B_x| \leq C \cdot p(n)$ und $V_{L_\Pi}(x, B_x) = \text{ja}$;

Entscheidung „ $x \notin L_\Pi$ “ falls für alle Beweise $B \in \Sigma_0^*$ mit $|B| \leq C \cdot p(n)$ gilt: $V_{L_\Pi}(x, B) = \text{nein}$.

Hierbei wird die Entscheidung $V_{L_\Pi}(x, B)$ nach einer Anzahl von Schritten getroffen, die in $O(p(n))$ liegt.



Umgangssprachlich kann man die Klassen **P** und **NP** etwa folgendermaßen beschreiben:

Die Klasse **P** ist die Klasse der Entscheidungsprobleme, für die bei Vorgabe einer Instanz **in polynomieller Zeit** (in Abhängigkeit von der Größe der Instanz) **entschieden wird, ob die Instanz eine das Problem definierende Eigenschaft besitzt oder nicht**.

Die Klasse **NP** ist die Klasse der Entscheidungsprobleme, für die bei Vorgabe einer Instanz und eines Beweises, der zeigen soll, daß die Instanz eine das Problem definierende Eigenschaft besitzt, **der Beweis in polynomieller Zeit** (in Abhängigkeit von der Größe der Instanz) **verifiziert werden**.

Man kann die Klassen **P** bzw. **NP** auch durch deterministische bzw. nichtdeterministische Turingmaschinen charakterisieren.

Ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in **P**, wenn es eine deterministische Turingmaschine TM_{L_Π} gibt, die bei Eingabe eines Wortes $x \in \Sigma_\Pi^*$ entscheidet, ob $x \in L_\Pi$ gilt oder nicht. Diese Entscheidung wird in $p_{L_\Pi}(|x|)$ vielen Schritten getroffen, wobei p_{L_Π} ein

Polynom mit $p_{L_\Pi}(n) \geq n$ ist. Ist $x \in L_\Pi$, dann stoppt TM_{L_Π} im Zustand q_{accept} („ja-Zustand“, „ja-Entscheidung“). Ist $x \notin L_\Pi$, dann befindet sich TM_{L_Π} nach höchstens $p_{L_\Pi}(|x|)$ in einem Zustand $q \neq q_{accept}$, für den es keinen Nachfolgezustand gibt. Dann stoppt TM_{L_Π} in einem nichtakzeptierenden Zustand („nein-Zustand“, „nein-Entscheidung“). Man kann darüber hinaus annehmen, daß TM_{L_Π} nur ein Band besitzt, da jede k -DTM durch eine 1-DTM simuliert werden kann; diese Simulation ist von der Ordnung $O(p_{L_\Pi}^2(n))$, also wieder polynomiell zeitbeschränkt, so daß man für die weiteren Betrachtungen dann das Polynom $C \cdot p_{L_\Pi}^2(n)$ anstelle von $p_{L_\Pi}(n)$ nimmt. Da TM_{L_Π} also nicht mehr als $p_{L_\Pi}(|x|)$ viele Schritte ausführt, werden von TM_{L_Π} auch höchstens nur die Zellen mit den Nummern $1, \dots, p_{L_\Pi}(|x|)+1$ besucht.

Ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ liegt in **NP**, wenn die soeben beschriebene Turingmaschine mit einer nichtdeterministischen Überföhrungsfunktion arbeitet. Legt man dabei das in Kapitel 2.5 beschriebene Nichtstandardmodell einer Turingmaschine TM_{L_Π} mit Ratemodul zugrunde, dann kann man annehmen, daß das Band von TM_{L_Π} beidseitig unbegrenzt ist; das Eingabewort x belegt anfangs die Zellen mit den Nummern $1, \dots, |x|$, in Zelle 0 wird eine linke Markierung $\#$ geschrieben. Die Zellen mit negativen Nummern werden vom Ratemodul zur anfänglichen Erzeugung eines Beweises verwendet, der dann anschließend deterministisch verifiziert wird. In diesem Fall besucht TM_{L_Π} höchstens die Zellen $-p_{L_\Pi}(|x|), \dots, 0, 1, \dots, p_{L_\Pi}(|x|)+1$.

Die Arbeitsweise eines nichtdeterministischen polynomiellen Entscheidungsalgorithmus kann man durch einen „normalen“ deterministischen Algorithmus simulieren. Es ist jedoch bisher keine Simulation bekannt, die dabei in polynomieller Zeit arbeitet. Genauer gilt:

Satz 5.3-1:

Für jedes Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ aus **NP** mit Verifizierer V_{L_Π} und polynomieller Zeitkomplexität $p(n)$ gibt es einen deterministischen Algorithmus A_{L_Π} , der für jedes $x \in \Sigma_\Pi^*$ mit $|x| = n$ entscheidet, ob $x \in L_\Pi$ gilt oder nicht und die Zeitkomplexität $O(p(n) \cdot 2^{O(p(n))})$ besitzt.

Die Aussage „**P** = **NP**“ würde bedeuten, daß es für jedes Problem Π aus **NP** mit polynomiell zeitbeschränktem Verifizierer V_{L_Π} sogar einen *polynomiell zeitbeschränkten* deterministischen Algorithmus A_{L_Π} gibt, der dasselbe Ergebnis wie V_{L_Π} liefert. Das bzw. die Unmöglichkeit einer derartigen Simulation sind jedoch nicht bekannt.

Beweis:

Der folgende Pseudocode für $A_{L_{\Pi}}$ beschreibt die wesentlichen Aspekte der Simulation:

Bemerkung: Ein Beweis B , den der Verifizierer liest, ist ein Wort über einem Alphabet Σ_0 .

FUNCTION $A_{L_{\Pi}}(x)$;

```
VAR n      : INTEGER;
    lng     : INTEGER;
    B       :  $\Sigma_0^*$ ;
    antwort : BOOLEAN;
    weiter  : BOOLEAN;
```

BEGIN

```
  n      := size(x);      { Größe der Eingabe          }
  lng     := C * p(n);    { maximale Länge eines Beweises }
  antwort := FALSE;
  weiter  := TRUE;
```

WHILE weiter DO

```
  IF (es sind bereits alle Wörter aus  $\Sigma_0^*$  mit Länge  $\leq$  lng erzeugt worden)
  THEN weiter := FALSE
  ELSE
```

 BEGIN

```
      B := nächstes Wort aus  $\Sigma_0^*$  mit Länge  $\leq$  lng;
```

```
      IF  $V_{\Pi}(x, B) = \text{ja}$ 
```

```
      THEN BEGIN
```

```
          antwort := TRUE;
```

```
          weiter  := FALSE;
```

```
      END;
```

```
    END;
```

CASE antwort OF

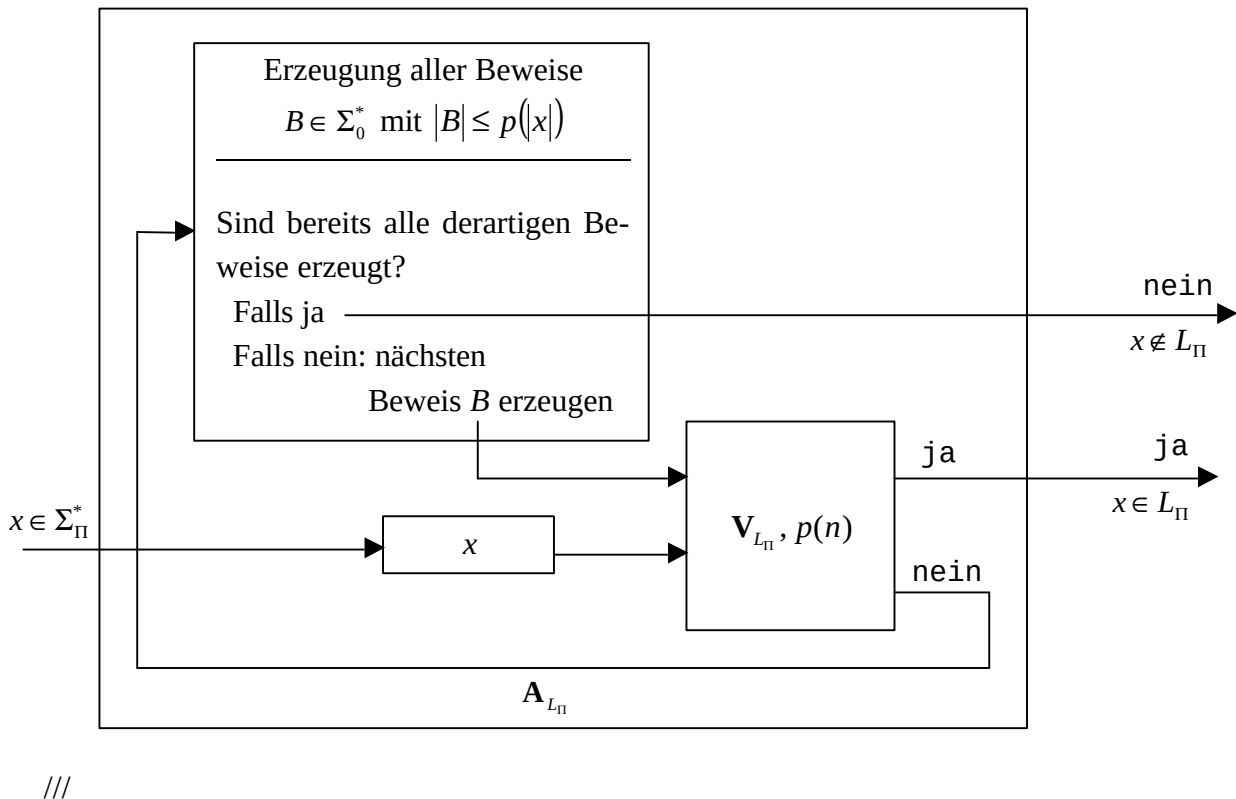
```
  TRUE  :  $A_{\Pi}(x) := \text{ja}$ ;
```

```
  FALSE :  $A_{\Pi}(x) := \text{nein}$ ;
```

```
END;
```

END;

Graphisch läßt sich die Simulation folgendermaßen darstellen:



///

Einige Beispiele für Probleme aus der Klasse $\mathbf{P} = \bigcup_{k=0}^{\infty} \text{TIME}(n^k)$ wurden bereits in den vorherigen Kapiteln behandelt. Dazu gehört das Problem, festzustellen, ob zwischen zwei Knoten eines gewichteten Graphen ein Pfad mit minimalem Gewicht existiert, das eine vorgegebene Schranke nicht überschreitet.

Entscheidungsprobleme leiten sich häufig von Optimierungsproblemen ab. Leider stellt sich heraus, daß für eine Vielzahl dieser Entscheidungsprobleme keine Lösungsalgorithmen bekannt sind, die polynomielles Laufzeitverhalten aufweisen. Das gilt insbesondere für viele Entscheidungsprobleme, die zu Optimierungsproblemen gehören, die für die Praxis relevant sind (Problem des Handlungsreisenden, Partitionenproblem usw.). Für diese (Optimierungs-) Probleme verfügt man meist über deterministische Lösungsalgorithmen, deren Laufzeitverhalten exponentiell in der Größe der Eingabe ist, bzw. man kann den Nachweis erbringen, daß die zugehörigen Entscheidungsprobleme in **NP** liegen. Derartige Verfahren werden als praktisch nicht durchführbar (intractable) angesehen, obwohl durch den Einsatz immer schnellerer Rechner immer größere Probleme behandelt werden können. Es stellt sich daher die Frage, wieso trotz intensiver Suche nach polynomiellen Lösungsverfahren derartige schnelle Algorithmen (bisher) nicht gefunden wurden. Seit Anfang der 1970'er Jahre erklärt eine inzwischen gut etablierte Theorie, die **Theorie der NP-Vollständigkeit**, Ursachen dieses Phänomens. Eine Einführung in diese Theorie gibt Kapitel 5.4.

Ein wichtiges Beispiel für Probleme auf der Grenze zwischen **P** und **NP** ist das Problem der Erfüllbarkeit Boolescher Ausdrücke (siehe Kapitel 1.1).

Das folgende Entscheidungsproblem liegt in **P**:

Erfüllende Wahrheitsbelegung (erfSAT)

Instanz: $[F, f]$

F ist ein Boolescher Ausdruck in konjunktiver Normalform mit der Variablenmenge $V = \{x_1, \dots, x_n\}$, $f : V \rightarrow \{\text{TRUE}, \text{FALSCH}\}$ ist eine Belegung der Variablen mit Wahrheitswerten.

Lösung: Entscheidung „ja“, falls die durch f gegebene Belegung dem Ausdruck F den Wahrheitswert TRUE gibt,
Entscheidung „nein“, sonst.

Ein Lösungsalgorithmus setzt einfach die in der Eingabe-Instanz $[F, f]$ gelieferte Belegung f in die Formel F ein und ermittelt den Wahrheitswert der Formel gemäß den Auswertungsregeln für die beteiligten Junktoren.

Das folgende Entscheidungsproblem liegt in **PSPACE**:

Erfüllbarkeitsproblem für Boolesche Ausdrücke in konjunktiver Normalform (CSAT)

Instanz: F

F ist ein Boolescher Ausdruck in konjunktiver Normalform mit der Variablenmenge $V = \{x_1, \dots, x_n\}$.

Lösung: Entscheidung „ja“, falls es eine Belegung der Variablen von F gibt, so daß sich bei Auswertung der Formel F der Wahrheitswert TRUE ergibt,
Entscheidung „nein“, sonst.

CSAT liegt in **PSPACE** (und damit in **EXP**). Für den Beweis genügt es zu zeigen, daß nacheinander alle möglichen 2^n Belegungen der n Variablen in einer Eingabe-Instanz F mit einem Speicherplatzverbrauch von polynomiell vielen Speicherzellen erzeugt, in die Formel F eingesetzt und ausgewertet werden können.

Es ist nicht bekannt, ob CSAT in **P** liegt (viele sprechen dagegen).

Die Probleme erfSAT mit einer Eingabeinstanz $[F, f]$ und CSAT mit einer Eingabeinstanz F unterscheiden sich grundsätzlich dadurch, daß bei ersterem in der Eingabeinstanz $[F, f]$ eine wesentliche Zusatzinformation, nämlich eine potentiell erfüllende Belegung f der Variablen vorgegeben ist, die nur noch daraufhin überprüft werden muß, ob sie wirklich die in der Eingabeinstanz enthaltene Formel F erfüllt. Zur Entscheidung, ob eine Eingabeinstanz F von CSAT erfüllbar ist, muß also entweder eine erfüllende Belegung f konstruiert bzw. aufgrund geeigneter Argumente die Erfüllbarkeit gezeigt werden, oder es muß der Nachweis erbracht werden, daß keine Belegung der Variablen von F die Formel erfüllt. Wenn dieser Nachweis nur dadurch gelingt, daß *alle* 2^n möglichen Belegungen überprüft werden, ist mit einem polynomiellen Entscheidungsalgorithmus nicht zu rechnen. Intuitiv ist diese Entscheidungsaufgabe also schwieriger zu bewältigen, weil weniger Anfangsinformationen vorliegen, als lediglich die Verifikation einer potentiellen Lösung.

Eine einfache Überlegung zeigt, daß CSAT in **NP** liegt. Ein Verifizierer für CSAT arbeitet wie folgt: Er erhält mit einer Eingabeinstanz f (mit n Variablen) für CSAT als Zusatzinformation eine Belegung B_f der Variablen in f . Wenn f erfüllbar ist, nimmt man für B_f gerade eine erfüllende Belegung. Wenn f nicht erfüllbar ist, liefert keine Belegung B_f der Variablen den Wahrheitswert TRUE. Der Verifizierer überprüft lediglich (auf deterministische Weise) in $O(n)$ vielen Schritten, ob sich der Wert TRUE ergibt, wenn man die in B_f vorkommenden Wahrheitsheitswerte in f einsetzt; in diesem Fall wird f akzeptiert, ansonsten nicht. Insgesamt liegt polynomielles Laufzeitverhalten vor.

Da bei diesem Verfahren nicht wesentlich eingeht, ob f in konjunktiver Normalform vorliegt, sondern lediglich, ob f ein korrekter Boolescher Ausdruck und erfüllbar ist, ergibt sich, daß auch das folgende allgemeinere Problem SAT in **NP** liegt.

Erfüllbarkeitsproblem für Boolesche Ausdrücke in allgemeiner Form (SAT)

Instanz: F

F ist ein Boolescher Ausdruck mit der Variablenmenge $V = \{x_1, \dots, x_n\}$.

Lösung: Entscheidung „ja“, falls es eine Belegung der Variablen von F gibt, so daß sich bei Auswertung der Formel F der Wahrheitswert TRUE ergibt, Entscheidung „nein“, sonst.

SAT liegt in **NP**.

Schränkt man das Problem CSAT auf diejenigen Booleschen Ausdrücke in konjunktiver Normalform ein, in denen jede Klausel genau 2 Literale enthält, so erhält man das Problem 2-

SAT. Entsprechend ist 3-SAT dadurch definiert, daß hier jede Klausel genau 3 Literale enthält.

Erfüllbarkeitsproblem für Boolesche Ausdrücke in konjunktiver Normalform mit genau 2 Literalen pro Klausel (2-CSAT)

Instanz: F

$F = F_1 \wedge \dots \wedge F_m$ ist ein Boolescher Ausdruck in konjunktiver Normalform mit der Variablenmenge $V = \{x_1, \dots, x_n\}$ mit der zusätzlichen Eigenschaft, daß jedes F_i genau zwei Literale enthält.

Lösung: Entscheidung „ja“, falls es eine Belegung der Variablen von F gibt, so daß sich bei Auswertung der Formel F der Wahrheitswert TRUE ergibt, Entscheidung „nein“, sonst.

Man kann zeigen, daß 2-SAT in **P** liegt:

Satz 5.3-2:

Ist F eine Instanz für 2-CSAT mit n Variablen und m Klauseln, dann liefert das beschriebene Verfahren mit der Prozedur `Erfuellbarkeit_2CSAT` eine korrekte ja/nein-Entscheidung. Die (worst-case-) Zeitkomplexität des Verfahrens ist von der Ordnung $O(n \cdot m)$. Dieser Aufwand ist polynomiell in der Größe der Eingabe.

Beweis:

Ein Algorithmus zur Ermittlung einer erfüllenden Belegung kann etwa nach folgender Strategie vorgehen:

Man nehme eine bisher noch nicht betrachtete Klausel $F_i = (y_1 \vee y_2)$ von F . Falls eines der Literale während des bisherigen Ablaufs bereits den Wahrheitswert TRUE erhalten hat, dann wird F_i als erfüllt erklärt. Falls eines der Literale, etwa y_1 , den Wert FALSE hat, dann erhält das Literal y_2 den Wahrheitswert TRUE und $\neg y_2$ den Wahrheitswert FALSE. F gilt dann als erfüllt. Dabei kann jedoch ein Konflikt auftreten, nämlich daß ein Literal durch die Zuweisung einen Wahrheitswert bekommen soll, jedoch den komplementären Wahrheitswert bereits vorher erhalten hat. In diesem Fall wird die vorherige Zuweisung rückgängig gemacht. Falls es dann wieder zu einem Konflikt mit diesem Literal kommt, ist die Formel nicht erfüllbar.

Der folgende Pseudocode-Algorithmus implementiert diese Strategie; aus Gründen der Lesbarkeit wird auf die exakte Deklaration der lokalen Variablen und der verwendeten Datentypen und auf deren Implementierung verzichtet.

Algorithmus zur Lösung des Erfüllbarkeitsproblems für Boolesche Ausdrücke in konjunktiver Normalform mit genau 2 Literalen pro Klausel (2-CSAT)

Eingabe: $F = F_1 \wedge \dots \wedge F_m$
 F ist ein Boolescher Ausdruck in konjunktiver Normalform mit der Variablenmenge $V = \{x_1, \dots, x_n\}$ mit der zusätzlichen Eigenschaft, daß jedes F_i genau zwei Literale enthält, d.h. jedes F_i hat die Form $F_i = (y_{i_1} \vee y_{i_2})$
 $size(F) = n$.

Der Datentyp

TYPE KNF_typ = ...;

beschreibe den Typ einer Formel in konjunktiver Normalform, der Datentyp

TYPE Literal_typ = ...;

den Typ eines Literals bzw. einer Variablen in einem Booleschen Ausdruck.

VAR F : KNF_typ;

Entscheidung : Entscheidungs_typ;

F := $F_1 \wedge \dots \wedge F_m$

Verfahren: Aufruf der Prozedur
 Erfuellbarkeit_2CSAT (F, Entscheidung);

Im Ablauf der Prozedur werden Variablen Wahrheitswerte TRUE bzw. FALSE zugewiesen. Zu Beginn des Ablaufs hat jede Variable noch einen undefinierten Wert; in diesem Fall wird die Variable als „noch nicht zugewiesen“ bezeichnet. Jede Klausel F_i von F gilt zu Beginn des Prozedurablaufs als „unerfüllt“ (da ihren Literalen ja noch kein Wahrheitswert zugewiesen wurde). Sobald einem Literal in F_i der Wahrheitswert TRUE zugewiesen wurde, gilt die Klausel als „erfüllt“.

Ausgabe: Entscheidung = ja, falls F erfüllbar ist,
 Entscheidung = nein sonst.

PROCEDURE Erfuellbarkeit_2CSAT (F : KNF_typ;
 VAR Entscheidung:Entscheidungs_typ);
 VAR C : SET OF KNF_typ;

```

V          : SET OF Literal_typ;
x          : Literal_typ;
firstguess : BOOLEAN;

BEGIN { Erfuellbarkeit_2CSAT }
  {  $F = F_1 \wedge \dots \wedge F_m$  }
  C := { $F_1, \dots, F_m$ };
  erkläre alle Klauseln in C als unerfüllt;
  V := Menge der in F vorkommenden Variablen;
  erkläre alle Variablen in V als nicht zugewiesen;

  WHILE (V enthält eine Variable x) DO
    BEGIN
      x          := TRUE;
      firstguess := TRUE;
      WHILE (C enthält eine unerfüllte Klausel  $F_i = (y_{i_1} \vee y_{i_2})$ , wobei mindestens einem
        Literal ein Wahrheitswert zugewiesen ist) DO
        BEGIN
          IF ( $y_{i_1} = \text{TRUE}$ ) OR ( $y_{i_2} = \text{TRUE}$ )
          THEN erkläre  $F_i$  als erfüllt
          ELSE IF ( $y_{i_1} = \text{FALSE}$ ) AND ( $y_{i_2} = \text{FALSE}$ )
          THEN BEGIN
            IF NOT firstguess
            THEN BEGIN
              Entscheidung := nein;
              Exit
            END
            ELSE BEGIN
              erkläre alle Klauseln in C als unerfüllt;
              erkläre alle Variablen in V als nicht zugewiesen;
              x          := FALSE;
              firstguess := FALSE;
            END;
          END
          ELSE BEGIN
            IF  $y_{i_1} = \text{FALSE}$ 
            THEN  $y_{i_2} := \text{TRUE}$ 
            ELSE  $y_{i_1} := \text{TRUE}$ ;
            erkläre  $F_i$  als erfüllt;
          END;
        END
      { WHILE C enthält eine unerfüllte Klausel };
      entferne aus C die erfüllten Klauseln;
    END
  END

```

```

    entferne aus V die Variablen, denen ein Wahrheitswert zugewiesen wurde;
  END { WHILE V enthält eine Variable x };
  Entscheidung := ja;

END { Erfuellbarkeit_2CSAT };
///

```

Für das zu 2-CSAT analoge Problem 3-CSAT der Erfüllbarkeit, bei dem jede Klausel genau 3 Literale enthält, ist kein polynomielles Entscheidungsverfahren bekannt. Es wird vermutet, daß es auch kein derartiges Verfahren gibt. Natürlich liegt 3-CSAT in **NP**.

Im folgenden werden einige wenige **Beispiele für Probleme Π aus NP** aufgeführt, von denen nicht bekannt ist, ob sie in **P** liegen. Dabei wird für Instanzen $x \in \Sigma_{\Pi}^*$ jeweils der Beweis $B_x \in \Sigma_0^*$ angegeben, den der Verifizierer zur ja/nein-Entscheidung heranzieht. Die Angabe von B_x erfolgt informell, sie muß entsprechend (beispielsweise in eine binäre Zeichenkette) übersetzt werden. Mit $size(x)$ wird die Größe einer Eingabeinstanz angegeben.

- Erfüllbarkeitsproblem für Boolesche Ausdrücke in konjunktiver Normalform (CSAT) bzw. Boolescher Ausdrücke in allgemeiner Form (SAT):

Instanz: F ist ein Boolescher Ausdruck in konjunktiver Normalform bzw. in allgemeiner Form mit der Variablenmenge $V = \{x_1, \dots, x_n\}$; $size(F) = n$

Beweis: B_F ist eine 0-1-Folge der Länge n

Arbeitsweise des Verifizierers: Der i -te Wert in B_F wird als Belegung von x_i interpretiert, und zwar wird der Wert 0 in FALSE und der Wert 1 in TRUE übersetzt. Die so entstehende Belegung der Variablen wird in F eingesetzt und die Formel ausgewertet. Falls sich bei dieser Auswertung von F der Wert TRUE ergibt, wird ja ausgegeben, ansonsten nein.
- Problem des Handlungsreisenden:

Instanz: $[G, K]$, $G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein Gewicht; $K \in \mathbf{R}_{\geq 0}$; $size([G, K]) = k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$

Beweis: $B_{[G, K]}$ ist eine Permutation der Zahlen $1, \dots, n$ in Binärkodierung (mit Länge in $O(n \cdot \log(n))$)

Arbeitsweise des Verifizierers: Überprüfung, ob die Permutation $B_{[G, K]}$ eine Rundreise in G beschreibt, deren Gewicht $\leq K$ ist. In diesem Fall wird ja ausgegeben, an-

sonsten nein. Mit Binärsuche lässt sich sogar in polynomieller Zeit überprüfen, ob es eine Rundreise mit minimalen Gewicht $\leq K$ gibt.

- Partitionenproblem:

Instanz: $I = \{a_1, \dots, a_n\}$

I ist eine Menge von n natürlichen Zahlen; $size(I) = n \cdot \log_2(B)$ mit $B = \max\{\lceil \log(a_i) \rceil \mid i = 1, \dots, n\}$

Beweis: B_I ist eine Teilmenge von $\{1, \dots, n\}$ in Binärkodierung (mit Länge in $O(n \cdot \log(n))$)

Arbeitsweise des Verifizierers: Überprüfung, ob $\sum_{j \in B_I} a_j = \sum_{j \notin B_I} a_j$ gilt. In diesem Fall wird ja ausgegeben, ansonsten nein.

- 0/1-Rucksackproblem:

Instanz: $I = (A, M)$

$A = \{a_1, \dots, a_n\}$ ist eine Menge von n natürlichen Zahlen; $M \in \mathbf{N}$; $size(I) = n \cdot \log_2(B)$ mit $B = \max(\{\lceil \log(a_i) \rceil \mid i = 1, \dots, n\} \cup \{\log(M)\})$

Beweis: B_I ist eine 0-1-Folge $x_1, \dots, x_n$ (mit Länge n)

Arbeitsweise des Verifizierers: Überprüfung, ob $\sum_{i=1}^n x_i \cdot a_i = M$ gilt. In diesem Fall wird ja ausgegeben, ansonsten nein.

- Problem der $\{0,1\}$ -Linearen Programmierung

Instanz: $[A, \vec{b}, \vec{c}, K]$, $A \in \mathbf{Z}^{m \cdot n}$ ist eine ganzzahlige Matrix mit m Zeilen und n Spalten, $\vec{b} \in \mathbf{Z}^m$ ist ein ganzzahliger Vektor, $\vec{c} \in \mathbf{N}^n$ ist ein nichtnegativer ganzzahliger Vektor, $K \in \mathbf{N}$, $size([A, \vec{b}, \vec{c}]) = k = m \cdot n \cdot B$ mit

$$B = \max\{\lceil \log(e) \rceil \mid e \in A \vee e \in \vec{b} \vee e \in \vec{c}\}$$

Beweis: $B_{[A, \vec{b}, \vec{c}]}$ ist eine 0-1-Folge der Länge n (für den Vektor $\vec{x}$)

Arbeitsweise des Verifizierers: Überprüfung, ob für $\vec{x} = B_{[A, \vec{b}, \vec{c}]}$ die Bedingungen $A \cdot \vec{x} \geq \vec{b}$

und $\sum_{i=1}^n c_i \cdot x_i \leq K$ gelten. In diesem Fall wird ja ausgegeben, ansonsten nein.

- Problem des längsten Wegs in einem gewichteten Graphen:

Instanz: $[G, v_i, v_j, K]$, $G = (V, E, w)$ ist ein gewichteter (gerichteter oder ungerichteter) Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; $v_i \in V$; $v_j \in V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein Gewicht; $K \in \mathbf{R}_{\geq 0}$; $\text{size}([G, K]) = k = n \cdot B$ mit $B = \max\{\lceil \log(w(e)) \rceil \mid e \in E\}$

Gesucht: ein einfacher Weg (d.h. ohne Knotenwiederholungen) von v_i nach v_j mit Gewicht $\geq K$.

Beweis: $B_{[G, v_i, v_j, K]}$ ist eine Folge $i_1, i_2, \dots, i_k$ von natürlichen Zahlen aus $\{1, \dots, n\}$ (mit Länge in $O(n \cdot \log(n))$)

Arbeitsweise des Verifizierers: Überprüfung, ob $v_{i_1}, v_{i_2}, \dots, v_{i_k}$ einen einfachen Weg (d.h. ohne Knotenwiederholungen) von v_{i_1} nach v_{i_k} mit Gewicht $\geq K$ beschreibt. In diesem Fall wird ja ausgegeben, ansonsten nein.

Man kennt heute mehrere tausend Probleme aus **NP**. In der angegebenen Literatur werden geordnet nach Anwendungsgebieten umfangreiche Listen von **NP**-Problemen aufgeführt.

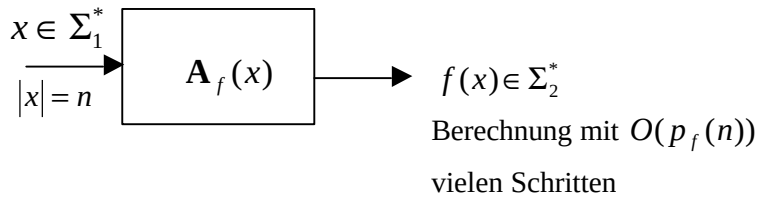
5.4 NP-Vollständigkeit

Die bisher beschriebenen Probleme entstammen verschiedenen Anwendungsgebieten. Dementsprechend unterscheiden sich die Alphabete, mit denen man die jeweiligen Instanzen bildet. Boolesche Ausdrücke werden mit einem anderen Alphabet kodiert als gewichtete Graphen oder Instanzen für das Partitionenproblem. Selbstverständlich lassen sich letztlich alle Probleme mit Hilfe des Alphabets $\{0, 1\}$ kodieren, so daß man mit einem einzigen Alphabet auskommen könnte. Aber selbst, wenn die Eingabeinstanzen unterschiedlicher Probleme mit dem Alphabet $\{0, 1\}$ kodiert werden, weisen die dann so kodierten Bestandteile der betrachteten Objekte wie Graphen, Boolesche Variablen oder Zahlen immer noch grundlegend verschiedene Eigenschaften auf, so daß man weiterhin davon ausgehen kann, daß man unterschiedliche Anwendungen mit Hilfe unterschiedlicher Alphabete kodiert. Im folgenden wird eine Verbindung zwischen den unterschiedlichen Problemen und ihren zugrundeliegenden Alphabeten hergestellt.

Eine Funktion $f: \Sigma_1^* \rightarrow \Sigma_2^*$, die Wörter über dem endlichen Alphabet Σ_1 auf Wörter über dem endlichen Alphabet Σ_2 abbildet, heißt **durch einen deterministischen Algorithmus in**

polynomieller Zeit berechenbar, wenn gilt: Es gibt einen deterministischen Algorithmus A_f mit Eingabemenge Σ_1^* und Ausgabemenge Σ_2^* und ein Polynom p_f mit folgenden Eigenschaften:

bei Eingabe von $x \in \Sigma_1^*$ mit der Größe $|x|=n$ erzeugt der Algorithmus die Ausgabe $f(x) \in \Sigma_2^*$ und benötigt dazu höchstens $O(p_f(n))$ viele Schritte.



Es seien $L_1 \subseteq \Sigma_1^*$ und $L_2 \subseteq \Sigma_2^*$ zwei Mengen aus Zeichenketten (Wörtern) über jeweils zwei endlichen Alphabeten. L_1 heißt **polynomiell (many-one) reduzierbar** auf L_2 , geschrieben

$$L_1 \leq_m^p L_2,$$

wenn gilt: Es gibt eine Funktion $f: \Sigma_1^* \rightarrow \Sigma_2^*$, die durch einen deterministischen Algorithmus in polynomieller Zeit berechenbar ist und für die gilt:

$$x \in L_1 \Leftrightarrow f(x) \in L_2.$$

Diese Eigenschaft kann auch so formuliert werden:

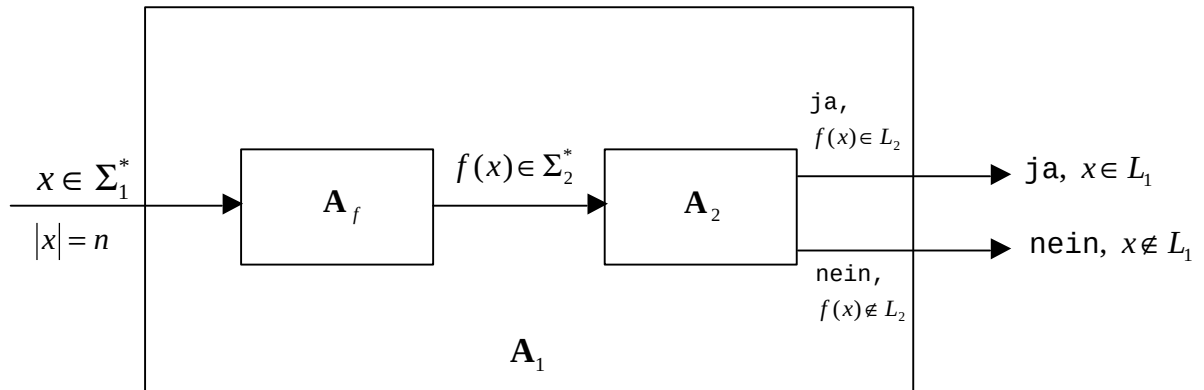
$$x \in L_1 \Rightarrow f(x) \in L_2 \text{ und } x \notin L_1 \Rightarrow f(x) \notin L_2.$$

Bemerkung: Es gibt noch andere Formen der Reduzierbarkeit zwischen Mengen, z.B. die allgemeinere Form der Turing-Reduzierbarkeit mit Hilfe von Orakel-Turingmaschinen.

Die **Bedeutung der Reduzierbarkeit** $\leq_m^p$ zeigt folgende Überlegung.

Für die Mengen $L_1 \subseteq \Sigma_1^*$ und $L_2 \subseteq \Sigma_2^*$ gelte $L_1 \leq_m^p L_2$ mittels der Funktion $f: \Sigma_1^* \rightarrow \Sigma_2^*$ bzw. des Algorithmus A_f . Seine Zeitkomplexität sei das Polynom p_f . Für die Menge L_2 gebe es einen polynomiell zeitbeschränkten deterministischen Algorithmus A_2 , der L_2 erkennt; seine Zeitkomplexität sei das Polynom p_2 :

Mit Hilfe von A_f und A_2 läßt sich durch Hintereinanderschaltung beider Algorithmen ein polynomiell zeitbeschränkter deterministischer Algorithmus A_1 konstruieren, der L_1 erkennt:



Erkennung von L_1 mit Zeitaufwand der Größenordnung $O(p_f(n) + p_2(p_f(n)))$, d.h. polynomielltem Zeitaufwand

Ein $x_2 \in \Sigma_2^*$ kann dazu dienen, für mehrere (*many*) $x \in \Sigma_1^*$ die Frage „ $x \in L_1$?“ zu entscheiden (nämlich für alle diejenigen $x \in \Sigma_1^*$, für die $f(x) = x_2$ gilt). Allerdings darf man für jede Eingabe $f(x)$ den Algorithmus A_2 nur einmal (*one*) verwenden, nämlich bei der Entscheidung von $f(x)$.

Gilt für ein Problem Π_0 zur Entscheidung der Menge $L_{\Pi_0} \subseteq \Sigma_{\Pi_0}^*$ und für ein Problem Π zur Entscheidung der Menge $L_{\Pi} \subseteq \Sigma_{\Pi}^*$ die Relation $L_{\Pi} \leq_m^p L_{\Pi_0}$, so heißt das Problem Π auf das Problem Π_0 **polynomiell (many-one) reduzierbar**, geschrieben $\Pi \leq_m^p \Pi_0$.

Beispiel:

$\{0,1\}$ -Lineare Programmierung:

Instanz: $[A, \vec{b}, \vec{z}]$

$A \in \mathbb{Z}^{m \times n}$ ist eine ganzzahlige Matrix mit m Zeilen und n Spalten, $\vec{b} \in \mathbb{Z}^m$, $\vec{z}$ ist ein Vektor von n Variablen, die die Werte 0 oder 1 annehmen können.

Lösung: Entscheidung „ja“, falls gilt:

Es gibt eine Zuweisung der Werte 0 oder 1 an die Variablen in $\vec{z}$, so daß das lineare Ungleichungssystem $A \cdot \vec{z} \leq / \geq \vec{b}$ erfüllt ist. Die Bezeichnung $\leq / \geq$ steht hier für entweder $\leq$ oder $\geq$ in einer Ungleichung.

Es gilt $\text{CSAT} \leq_m^p \{0,1\}$ -Lineare Programmierung.

Einige leicht nachzuweisende Eigenschaften der polynomiellen many-one-Reduzierbarkeit faßt der nächste Satz zusammen.

Satz 5.4-1:

Es seien $L \subseteq \Sigma^*$, $L_1 \subseteq \Sigma_1^*$, $L_2 \subseteq \Sigma_2^*$ und $L_3 \subseteq \Sigma_3^*$ Sprachen über endlichen Alphabeten. Dann gilt:

1. $L \leq_m^p L$, d.h. die Relation $\leq_m^p$ ist reflexiv.
2. Aus $L_1 \leq_m^p L_2$ und $L_2 \leq_m^p L_3$ folgt $L_1 \leq_m^p L_3$, d.h. die Relation $\leq_m^p$ ist transitiv.
3. Gilt $L_1 \leq_m^p L_2$ und $L_2 \in \mathbf{P}$, dann ist $L_1 \in \mathbf{P}$.
4. Gilt $L_1 \leq_m^p L_2$ und $L_2 \in \mathbf{NP}$, dann ist $L_1 \in \mathbf{NP}$.

Eine der zentralen Definitionen in der Angewandten Komplexitätstheorie ist die **NP**-Vollständigkeit:

Ein Entscheidungsproblem Π_0 zur Entscheidung (Akzeptanz, Erkennung) der Menge $L_{\Pi_0} \subseteq \Sigma_{\Pi_0}^*$ heißt **NP-vollständig**, wenn gilt:

$L_{\Pi_0} \in \mathbf{NP}$, und für jedes Probleme Π zur Entscheidung (Akzeptanz, Erkennung) der Menge $L_{\Pi} \subseteq \Sigma_{\Pi}^*$ mit $L_{\Pi} \in \mathbf{NP}$ gilt $L_{\Pi} \leq_m^p L_{\Pi_0}$.

Mit obiger Definition der Reduzierbarkeit zwischen Problemen kann man auch sagen:

Das Entscheidungsproblem Π_0 ist **NP**-vollständig, wenn Π_0 in **NP** liegt und für jedes Entscheidungsprobleme Π in **NP** die Relation $\Pi \leq_m^p \Pi_0$ gilt.

Das erste Beispiel eines **NP**-vollständigen Problems wurde 1971 von Stephen Cook gefunden. Es gilt:

Satz 5.4-2:

Das Erfüllbarkeitsproblem für Boolesche Ausdrücke in allgemeiner Form (SAT) ist **NP**-vollständig.

Beweis:

Da es sich bei dieser Aussage um den zentralen Satz der angewandten Komplexitätstheorie handelt, soll die Beweisidee skizziert werden:

1. Die zu SAT gehörige Sprache ist $L_{\text{SAT}} = \{F \mid F \text{ ist ein erfüllbarer Boolescher Ausdruck}\}$. Bezeichnet Σ_{BOOLE}^* die Menge $\{F \mid F \text{ ist ein Boolescher Ausdruck}\}$, dann ist $L_{\text{SAT}} \subseteq \Sigma_{\text{BOOLE}}^*$. In Kapitel 5.3 wird gezeigt, daß SAT in **NP** liegt.
2. Für jedes Problem Π mit der Menge $L_{\Pi} \subseteq \Sigma_{\Pi}^*$ in **NP** ist die Relation $L_{\Pi} \leq_m^p L_{\text{SAT}}$ zu zeigen. Die einzige Eigenschaft, die bezüglich L_{Π} in einem Beweis genutzt werden kann, ist die Tatsache, daß es eine 1-NDTM $TM_{L_{\Pi}}$ gibt, die bei Eingabe eines Wortes $x \in \Sigma_{\Pi}^*$ entscheidet, ob $x \in L_{\Pi}$ gilt oder nicht. Diese Entscheidung wird in $p_{L_{\Pi}}(|x|)$ vielen Schritten getroffen, wobei $p_{L_{\Pi}}$ ein Polynom mit $p_{L_{\Pi}}(n) \geq n$ ist. Ist $x \in L_{\Pi}$, dann stoppt $TM_{L_{\Pi}}$ im akzeptierenden Zustand q_{accept} . Ist $x \notin L_{\Pi}$, dann stoppt $TM_{L_{\Pi}}$ in einem Zustand $q \neq q_{\text{accept}}$. $TM_{L_{\Pi}}$ besucht dabei höchstens die Zellen mit den Nummern $-p_{L_{\Pi}}(|x|), \dots, 0, 1, \dots, p_{L_{\Pi}}(|x|)+1$; hierbei wird das Nichtstandardmodell einer nichtdeterministischen Turingmaschine genommen (vgl. Kapitel 5.3).

Zum Nachweis von $L_{\Pi} \leq_m^p L_{\text{SAT}}$ ist eine Funktion $f_{\Pi} : \Sigma_{\Pi}^* \rightarrow \Sigma_{\text{BOOLE}}^*$ anzugeben, die die Eigenschaft „ $x \in L_{\Pi} \Leftrightarrow f_{\Pi}(x)$ ist erfüllbar“ besitzt ($f_{\Pi}(x)$ ist ein Boolescher Ausdruck); außerdem muß $f_{\Pi}(x)$ in $p_{f_{\Pi}}(|x|)$ vielen Schritten deterministisch berechenbar (konstruierbar) sein mit einem Polynom $p_{f_{\Pi}}$. Im folgenden wird beschrieben, wie $f_{\Pi}(x)$ aus x erzeugt werden kann. Der Boolesche Ausdruck $f_{\Pi}(x)$ beschreibt im wesentlichen, wie eine Berechnung von $TM_{L_{\Pi}}$ bei Eingabe von $x \in \Sigma_{\Pi}^*$ abläuft.

Die Zustandsmenge von $TM_{L_{\Pi}}$ sei $Q = \{q_0, \dots, q_k\}$, das Arbeitsalphabet von $TM_{L_{\Pi}}$ sei $\Sigma_{\Pi} = \{a_0, \dots, a_l\}$. Auf dem Eingabeband stehe das Wort $x \in \Sigma_{\Pi}^*$, $x = x_1 \dots x_n$ mit $x_i \in \Sigma_{\Pi}$ für $i = 1, \dots, n$. Es wird eine Reihe Boolescher Variablen erzeugt, deren Interpretation folgender Tabelle zu entnehmen ist:

Variable	Indizes	Interpretation
$zust_{t,q}$	$t = 0, \dots, p_{L_{\Pi}}(n)$ $q \in Q$	$zust_{t,q} = \text{TRUE}$ genau dann, wenn sich $TM_{L_{\Pi}}$ im Schritt t im Zustand q befindet Anzahl an Variablen $zust_{t,q} : (k+1) \cdot (p_{L_{\Pi}}(n)+1)$
$pos_{t,i}$	$t = 0, \dots, p_{L_{\Pi}}(n)$ $i = -p_{L_{\Pi}}(n), \dots, p_{L_{\Pi}}(n)+1$	$pos_{t,i} = \text{TRUE}$ genau dann, wenn sich der Schreib/Lesekopf von $TM_{L_{\Pi}}$ im Schritt t über der Zelle mit Nummer i befindet Anzahl an Variablen $pos_{t,i} : 2 \cdot (p_{L_{\Pi}}(n)+1)^2$
$band_{t,i,a}$	$t = 0, \dots, p_{L_{\Pi}}(n)$ $i = -p_{L_{\Pi}}(n), \dots, p_{L_{\Pi}}(n)+1,$ $a \in \Sigma_{\Pi}$	$band_{t,i,a} = \text{TRUE}$ genau dann, wenn sich im Schritt t in der Zelle mit Nummer i das Zeichen a befindet Anzahl an Variablen $band_{t,i,a} : 2(l+1)(p_{L_{\Pi}}(n)+1)^2$

Der zu konstruierende Boolesche Ausdruck $f_{\Pi}(x)$ besteht aus mehreren Teilen und enthält insbesondere mehrmals eine Teilformeln G , die genau dann den Wahrheitswert TRUE erhält, wenn genau eine der in G vorkommenden Variablen den Wahrheitswert TRUE trägt. Sind $y_1, \dots, y_m$ die in G vorkommenden Variablen, so lautet G :

$$\begin{aligned}
 G = G(y_1, \dots, y_m) &= \left(\bigvee_{i=1}^m y_i \right) \wedge \left(\bigwedge_{j=1}^{m-1} \bigwedge_{i=j+1}^m \neg(y_j \wedge y_i) \right) \\
 &= \left(\bigvee_{i=1}^m y_i \right) \wedge \left(\bigwedge_{j=1}^{m-1} \bigwedge_{i=j+1}^m (\neg y_j \vee \neg y_i) \right).
 \end{aligned}$$

Die zweite Zeile zeigt, daß $G = G(y_1, \dots, y_m)$ in konjunktiver Normalform formulierbar ist, ohne die Anzahl an Literalen zu ändern, und m^2 viele Literale enthält.

In $f_{\Pi}(x)$ kommen mehrere Teilformeln, die gemäß G aufgebaut sind, mit unterschiedlichen Variablen aus obiger Tabelle vor.

$f_{\Pi}(x)$ hat die Bauart $f_{\Pi}(x) = R \wedge A \wedge \ddot{U}_1 \wedge \ddot{U}_2 \wedge E$. Die Teilformel R beschreibt Randbedingungen, A Anfangsbedingungen, $\ddot{U}_1$ und $\ddot{U}_2$ beschreiben Übergangsbedingungen, und E beschreibt Endbedingungen.

In R wird ausgedrückt, daß $TM_{L_{\Pi}}$ zu jedem Zeitpunkt t in genau einem Zustand q ist, daß der Schreib/Lesekopf über genau einer Zelle mit einer Nummer i aus dem Intervall von $-p_{L_{\Pi}}(n)$ bis $p_{L_{\Pi}}(n)+1$ steht, und daß jede Zelle genau ein Zeichen $a \in \Sigma_{\Pi}$ enthält:

$$R = \bigwedge_{t=0}^{p_{L_{\Pi}}(n)} \left[G(zust_{t,q_0}, \dots, zust_{t,q_k}) \wedge G(pos_{t,-p_{L_{\Pi}}(n)}, \dots, pos_{t,p_{L_{\Pi}}(n)+1}) \wedge \left(\bigwedge_{i=-p_{L_{\Pi}}(n)}^{p_{L_{\Pi}}(n)+1} G(band_{t,i,a_0}, \dots, band_{t,i,a_l}) \right) \right]$$

Die Anzahl an Literalen in R beträgt

$$(p_{L_{\Pi}}(n)+1) \cdot ((k+1)^2 + 4(p_{L_{\Pi}}(n)+1)^2 + 2(p_{L_{\Pi}}(n)+1)(l+1)^2) \in O((p_{L_{\Pi}}(n))^3).$$

A beschreibt die Situation zum Zeitpunkt $t = 0$ (hierbei ist $b \in \Sigma_{\Pi}$ das Blankzeichen):

$$A = zust_{0,q_0} \wedge pos_{0,1} \wedge \left(\bigwedge_{i=1}^n band_{0,i,x_i} \right) \wedge \left(\bigwedge_{i=-p_{L_{\Pi}}(n)}^0 band_{0,i,b} \right) \wedge \left(\bigwedge_{i=n+1}^{p_{L_{\Pi}}(n)+1} band_{0,i,b} \right).$$

A enthält $2(p_{L_{\Pi}}(n)+2) \in O(p_{L_{\Pi}}(n))$ viele Literale.

$\ddot{U}_1$ beschreibt den Übergang von der zum Zeitpunkt t bestehenden Konfiguration zur Konfiguration zum Zeitpunkt $t+1$ für $t = 0, \dots, p_{L_{\Pi}}(n)$. Anstelle der Kopfbewegung L, S bzw. R werden hier die Werte $y = -1, y = 0$ bzw. $y = 1$ verwendet:

$$\ddot{U}_1 = \bigwedge_{t,q,i,a} \left[zust_{t,q} \wedge pos_{t,i} \wedge band_{t,i,a} \Rightarrow \bigvee \left\{ zust_{t+1,q'} \wedge pos_{t+1,i+y} \wedge band_{t+1,i,a'} \mid (q',a',y) \in d(q,a) \right\} \right]$$

Die Indizes nehmen die Werte $t = 0, \dots, p_{L_{\Pi}}(n)$, $q \in Q$, $i = -p_{L_{\Pi}}(n), \dots, p_{L_{\Pi}}(n)+1$ und $a \in \Sigma_{\Pi}$ an.

Die geschweifte Klammer in $\ddot{U}_1$ enthält für feste Werte t, q, a und i höchstens $3(k+1)(l+1)$ viele Literale, in der eckigen Klammer sind es daher höchstens $3(k+1)(l+1)+3$. Insgesamt enthält $\ddot{U}_1$ höchstens

$$2(p_{L_{\Pi}}(n)+1)^2 \cdot (k+1) \cdot (l+1)(3(k+1)(l+1)+3) \in O((p_{L_{\Pi}}(n))^2) \text{ viele Literale.}$$

Um zu zeigen, wie sich auch $\ddot{U}_1$ in eine äquivalente Formel in konjunktiver Normalform umformen läßt, soll exemplarisch ein Ausdruck

$zust_{t,q} \wedge pos_{t,i} \wedge band_{t,i,a} \Rightarrow \bigvee \left\{ zust_{t+1,q'} \wedge pos_{t+1,i+y} \wedge band_{t+1,i,a'} \mid (q',a',y) \in d(q,a) \right\}$ innerhalb der eckigen Klammer, der im rechten Teil zwei Alternativen

$zust_{t+1,q'} \wedge pos_{t+1,i+y'} \wedge band_{t+1,i,a'}$ und $zust_{t+1,q''} \wedge pos_{t+1,i+y''} \wedge band_{t+1,i,a''}$ enthält, umgeformt

werden. Um den Vorgang übersichtlich zu halten, wird dieser Ausdruck in der Form

$X_1 \wedge X_2 \wedge X_3 \Rightarrow (Y_1 \wedge Y_2 \wedge Y_3) \vee (Z_1 \wedge Z_2 \wedge Z_3)$ geschrieben. Hier stehen X_i, Y_i und Z_i

für jeweils drei Variablen. Es gilt

$$\begin{aligned} X_1 \wedge X_2 \wedge X_3 &\Rightarrow (Y_1 \wedge Y_2 \wedge Y_3) \vee (Z_1 \wedge Z_2 \wedge Z_3) \\ &= (\neg X_1 \vee \neg X_2 \vee \neg X_3) \vee (Y_1 \wedge Y_2 \wedge Y_3) \vee (Z_1 \wedge Z_2 \wedge Z_3) \\ &= (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Y_1) \wedge (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Y_2) \wedge (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Y_3) \\ &\quad \wedge (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Z_1) \wedge (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Z_2) \wedge (\neg X_1 \vee \neg X_2 \vee \neg X_3 \vee Z_3). \end{aligned}$$

Wie man sieht, werden für jedes Y_i und Z_i vier Literale notiert, so daß sich für die Anzahl der Literale in der äquivalenten Formel innerhalb eines Ausdrucks in der eckigen Klammer im allgemeinen Fall eine Obergrenze von $12(k+1)(l+1)$ angeben läßt. Daher

bleibt die Anzahl der Literale in $\ddot{U}_1$, selbst in der konjunktiven Normalform, von der Ordnung $O((p_{L_\Pi}(n))^2)$.

$\ddot{U}_2$ besagt, daß sich der Inhalt von Zellen, über denen der Schreib/Lesekopf nicht steht, nicht ändert:

$$\ddot{U}_2 = \bigwedge_{t,i,a} [(\neg pos_{t,i} \wedge band_{t,i,a}) \Rightarrow band_{t+1,i,a}].$$

Hierbei nehmen die Indizes die Werte $t = 0, \dots, p_{L_\Pi}(n)$, $i = -p_{L_\Pi}(n), \dots, p_{L_\Pi}(n)+1$ und $a \in \Sigma_\Pi$ an.

Die Anzahl an Literalen in $\ddot{U}_2$ beträgt $6(p_{L_\Pi}(n)+1)^2 \cdot (l+1) \in O((p_{L_\Pi}(n))^2)$.

$\ddot{U}_2$ läßt sich in eine äquivalente Formel in konfunktiver Normalform umformen, ohne die Anzahl an Literalen zu ändern:

$$\begin{aligned} \ddot{U}_2 &= \bigwedge_{t,i,a} [(\neg pos_{t,i} \wedge band_{t,i,a}) \Rightarrow band_{t+1,i,a}] \\ &= \bigwedge_{t,i,a} [(pos_{t,i} \vee \neg band_{t,i,a} \vee band_{t+1,i,a})]. \end{aligned}$$

E prüft nach, ob der Endzustand q_f zum Zeitpunkt $t = p_{L_\Pi}(n)$ erreicht ist:

$$E = zust_{p_{L_\Pi}(n), q_f}.$$

Insgesamt enthält $f_\Pi(x)$ eine Anzahl von Literalen der Ordnung $O((p_{L_\Pi}(n))^3)$, d.h. einer in der Länge des Eingabewortes polynomiellen Ordnung. Werden diese Literale durchnumeriert, so daß man $O((p_{L_\Pi}(n))^3)$ viele Literalpositionen bekommt, und dann binär kodiert (jede Zahl hat dann eine Länge der Ordnung $O(\log((p_{L_\Pi}(n))^3)) = O(\log(p_{L_\Pi}(n)))$), so hat $f_\Pi(x)$ eine Länge der Ordnung $O((p_{L_\Pi}(n))^4)$. Daher ist $f_\Pi(x)$ bei Vorgabe von $x \in \Sigma_\Pi^*$ (für festes Π) deterministisch in polynomieller Zeit berechenbar.

Es läßt sich leicht nachweisen, daß die Eigenschaft „ $x \in L_\Pi \Leftrightarrow f_\Pi(x)$ ist erfüllbar“ gilt. ///

Die Beweisskizze zeigt sogar:

Satz 5.4-3:

Das Erfüllbarkeitsproblem für Boolesche Ausdrücke in konjunktiver Normalform (CSAT) ist **NP**-vollständig.

NP-vollständige Probleme kann man innerhalb der Klasse **NP** als die am schwersten zu lösenden Probleme betrachten, denn sie entscheiden die **P-NP**-Frage:

Satz 5.4-4:

Gibt es mindestens ein **NP**-vollständiges Problem, das in **P** liegt, so ist **P** = **NP**.

Beweis:

Wegen **P** $\subseteq$ **NP** ist die umgekehrte Inklusion **NP** $\subseteq$ **P** zu zeigen. Es sei L_0 ein **NP**-vollständiges Problem, das in **P** liegt. Es sei $L \in \mathbf{NP}$. Zu zeigen ist $L \in \mathbf{P}$. Da L_0 **NP**-vollständig ist, gilt $L \leq_m^p L_0$. Mit Satz 5.4-1 folgt (wegen der Annahme $L_0 \in \mathbf{P}$) $L \in \mathbf{P}$. ///

Aus **NP**-vollständigen Problemen lassen sich aufgrund der Transitivität der $\leq_m$ -Relation (Satz 5.4-1) weitere **NP**-vollständige Probleme ableiten:

Satz 5.4-5:

Ist das Problem zur Entscheidung einer Menge $L_0 \subseteq \Sigma_0^*$ **NP**-vollständig und ist $L_0 \leq_m^p L_1$ für eine Menge $L_1 \subseteq \Sigma_1^*$ so gilt:
Ist L_1 in **NP**, so ist L_1 ebenfalls **NP**-vollständig.

Satz 5.4-6:

Bei **P** $\neq$ **NP** gilt:

Ist ein Entscheidungsproblem Π zur Entscheidung der Menge $L_\Pi \subseteq \Sigma_\Pi^*$ **NP**-vollständig, so ist das Problem Π schwer lösbar (intractable), d.h. es gibt keinen polynomiell zeitbeschränkten deterministischen Lösungsalgorithmus zur Entscheidung von L_Π . Insbesondere ist das zugehörige Optimierungsproblem, falls es ein solches gibt, erst recht schwer (d.h. nur mit mindestens exponentiellem Aufwand) lösbar.

NP-vollständige Probleme mit praktischer Relevanz sind heute aus vielen Gebieten bekannt, z.B. aus der Graphentheorie, dem Netzwerk-Design, der Theorie von Mengen und Partitionen, der Datenspeicherung, dem Scheduling, der Maschinenbelegung, der Personaleinsatzplanung, der mathematischen Programmierung und Optimierung, der Analysis und Zahlentheorie, der Kryptologie, der Logik, der Automatentheorie und der Theorie Formaler Sprachen, der Programm- und Codeoptimierung usw. (siehe Literatur). Die folgende Zusammenstellung listet einige wenige Beispiele auf, die z.T. bereits behandelt wurden.

Beispiele für NP-vollständige Entscheidungsprobleme

- **Erfüllbarkeitsproblem der Aussagenlogik (SAT):**

Instanz: F ,

F ist ein Boolescher Ausdruck (Formel der Aussagenlogik)

Lösung: Entscheidung „ja“, falls gilt:

F ist erfüllbar, d.h. es gibt eine Belegung der Variablen in F mit Werten TRUE bzw. FALSE, wobei gleiche Variablen mit gleichen Werten belegt werden, so daß sich bei Auswertung der Formel F der Wert TRUE ergibt.

Kodiert man Formeln der Aussagenlogik über dem Alphabet $A = \{\wedge, \vee, \neg, (,), x, 0, 1\}$, so ist eine Formel der Aussagenlogik ein Wort über dem Alphabet A . Nicht jedes Wort über dem Alphabet A ist eine Formel der Aussagenlogik (weil es eventuell syntaktisch nicht korrekt ist) und nicht jede Formel der Aussagenlogik als Wort über dem Alphabet A ist erfüllbar.

$L_{\text{SAT}} = \{F \mid F \text{ ist eine } \mathbf{erfüllbare} \text{ Formel der Aussagenlogik}\}.$

- **Erfüllbarkeitsproblem der Aussagenlogik mit Formeln in konjunktiver Normalform mit höchstens 3 Literalen pro Klausel (3-CSAT):**

Instanz: F ,

F ist eine Formel der Aussagenlogik in konjunktiver Normalform mit höchstens 3 Literalen pro Klausel, d.h. sind $x_1, \dots, x_n$ die verschiedenen in F vorkommenden Booleschen Variablen, so hat F die Form

$$F = F_1 \wedge F_2 \wedge \dots \wedge F_m$$

Hierbei hat jedes F_i die Form $F_i = (y_{i_1} \vee y_{i_2} \vee y_{i_3})$ oder $F_i = (y_{i_1} \vee y_{i_2})$ oder $F_i = (y_{i_1})$, und y_{i_j} steht für eine Boolesche Variable (d.h. $y_{i_j} = x_l$) oder für eine negierte Boolesche Variable (d.h. $y_{i_j} = \neg x_l$) oder für eine Konstante 0 (d.h. $y_{i_j} = 0$) bzw. 1 (d.h. $y_{i_j} = 1$).

Lösung: Entscheidung „ja“, falls gilt:

F ist erfüllbar.

Kodiert man die Formeln wie bei SAT, so gilt $L_{3\text{-CSAT}} \subseteq L_{\text{SAT}}$.

- **Kliquenproblem (KLIQUE):**

Instanz: $[G, k]$,
 $G = (V, E)$ ist ein ungerichteter Graph und k eine natürliche Zahl.

Lösung: Entscheidung „ja“, falls gilt:
 G besitzt eine „Klique“ der Größe k . Dieses ist eine Teilmenge $V' \subseteq V$ der Knotenmenge mit $|V'| = k$, und für alle $u \in V'$ und alle $v \in V'$ mit $u \neq v$ gilt $(u, v) \in E$.

- **0/1-Rucksack-Entscheidungsproblem (RUCKSACK):**

Instanz: $[a_1, \dots, a_n, b]$,
 $a_1, \dots, a_n$ und b sind natürliche Zahlen.

Lösung: Entscheidung „ja“, falls gilt:
 Es gibt eine Teilmenge $I \subseteq \{1, \dots, n\}$ mit $\sum_{i \in I} a_i = b$.

- **Partitionenproblem (PARTITION):**

Instanz: $[a_1, \dots, a_n]$,
 $a_1, \dots, a_n$ sind natürliche Zahlen.

Lösung: Entscheidung „ja“, falls gilt:
 Es gibt eine Teilmenge $I \subseteq \{1, \dots, n\}$ mit $\sum_{i \in I} a_i = \sum_{i \notin I} a_i$.

- **Packungsproblem (BINPACKING):**

Instanz: $[a_1, \dots, a_n, b, k]$,
 $a_1, \dots, a_n$ („Objekte“) sind natürliche Zahlen mit $a_i \leq b$ für $i = 1, \dots, n$, $b \in \mathbb{N}$ („Behältergröße“), $k \in \mathbb{N}$.

Lösung: Entscheidung „ja“, falls gilt:
 Die Objekte können so auf k Behälter der Füllhöhe b verteilt werden, so daß kein Behälter überläuft, d.h. es gibt eine Abbildung $f : \{1, \dots, n\} \rightarrow \{1, \dots, k\}$,

so daß für alle $j \in \{1, \dots, k\}$ gilt: $\sum_{f(i)=j} a_i \leq b$

- **Problem des Hamiltonschen Kreises in einem gerichteten (bzw. ungerichteten) Graphen (GERICHTETER bzw. UNGERICHTETER HAMILTONKREIS):**

Instanz: G ,

$G = (V, E)$ ist ein gerichteter bzw. ungerichteter Graph mit $V = \{v_1, \dots, v_n\}$.

Lösung: Entscheidung „ja“, falls gilt:

G besitzt einen Hamiltonschen Kreis. Dieses ist eine Anordnung $(v_{p(1)}, v_{p(2)}, \dots, v_{p(n)})$ der Knoten mittels einer Permutation π der Knotenindizes, so daß für $i = 1, \dots, n-1$ gilt: $(v_{p(i)}, v_{p(i+1)}) \in E$ und $(v_{p(n)}, v_{p(1)}) \in E$.

- **Problem des Handlungsreisenden (HANDLUNGSREISENDER):**

Instanz: $[M, k]$,

$M = (M_{i,j})$ ist eine $(n \times n)$ -Matrix von „Entfernungen“ zwischen n „Städten“ und eine Zahl k .

Lösung: Entscheidung „ja“, falls gilt:

Es gibt eine Permutation π (eine Tour, „Rundreise“), so daß

$$\sum_{i=1}^{n-1} M_{p(i), p(i+1)} + M_{p(n), p(1)} \leq k \text{ gilt?}$$

Zum Nachweis der **NP**-Vollständigkeit für eines dieser Probleme Π ist neben der Zugehörigkeit von Π zu **NP** jeweils die Relation $\Pi_0 \leq_m^p \Pi$ zu zeigen, wobei hier Π_0 ein Problem ist, für das bereits bekannt ist, daß es **NP**-vollständig ist, und das eine „ähnliche“ Struktur aufweist. Häufig beweist man

$$\text{SAT} \leq_m^p \text{3-CSAT} \leq_m^p \text{RUCKSACK} \leq_m^p \text{PARTITION} \leq_m^p \text{BINPACKING},$$

$$\text{3-CSAT} \leq_m^p \text{KLIQUE und}$$

$$\text{3-CSAT} \leq_m^p \text{GERICHTETER HAMILTONKREIS}$$

$$\leq_m^p \text{UNGERICHTETER HAMILTONKREIS} \leq_m^p \text{HANDLUNGSREISENDER.}$$

5.5 Bemerkungen zur Struktur von NP

Es sei **NPC** die Menge der **NP**-vollständigen Sprachen über einem endlichen Alphabet Σ . $\mathbf{NPC} \subseteq \mathbf{NP}$. Unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ gilt $\mathbf{NPC} \cap \mathbf{P} = \emptyset$.

Es gilt folgender Satz:

Satz 5.5-1:

Es sei B eine entscheidbare Menge mit $B \notin \mathbf{P}$. Dann gibt es eine Sprache $D \in \mathbf{P}$, so daß $A = D \cap B$ nicht zu $\mathbf{P}$ gehört, $A \leq_m B$, aber nicht $B \leq_m^p A$ gelten.

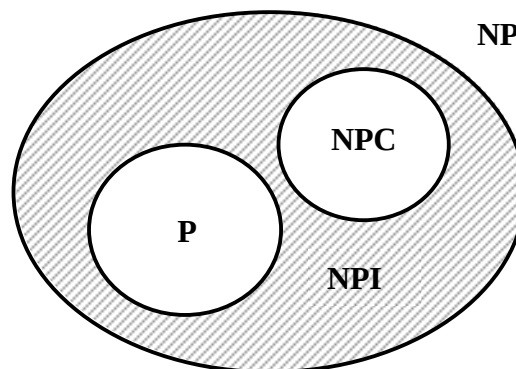
Satz 5.5-1 läßt sich folgendermaßen anwenden:

Es sei B eine **NP**-vollständige Sprache. Unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ gilt $B \notin \mathbf{P}$. Die Sprache $A = D \cap B$ gehört zu **NP**, da $D \in \mathbf{P}$ und $B \in \mathbf{NP}$ sind. Nach dem obigen Satz gilt nicht $B \leq_m^p A$, also ist A nicht **NP**-vollständig. Folglich gilt:

Satz 5.5-2:

Unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ ist $\mathbf{NP} \setminus (\mathbf{P} \cup \mathbf{NPC}) \neq \emptyset$.

Bezeichnet man $\mathbf{NP} \setminus (\mathbf{P} \cup \mathbf{NPC})$ als die Menge **NPI** der **NP-unvollständigen Sprachen**, so zeigt **NP** folgende Struktur:



Die Sprachen in **NPI** liegen bezüglich ihrer Komplexität also zwischen den „leichten“ Sprachen in **P** und den „schweren“ Sprachen in **NPC**. Obwohl es (bei $\mathbf{P} \neq \mathbf{NP}$) unendlich viele Sprachen in **NPI** geben muß, ist die Angabe konkreter Beispiele schwierig. Bisher kennt man kein Problem aus **NPI**. Von dem folgenden Graphenisomorphieproblem ISO wird vermutet, daß es in **NPI** liegt, da bisher (trotz großer Anstrengung) weder der Nachweis für $\text{ISO} \in \mathbf{P}$ noch der Nachweis $\text{ISO} \in \mathbf{NPC}$ gelungen ist.

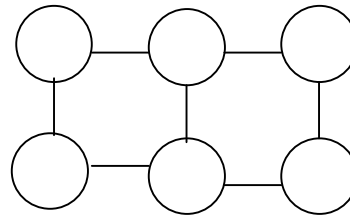
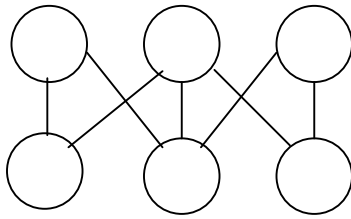
Graphenisomorphieproblem (ISO):

Instanz: Graphen $G_1 = (V_1, E_1)$ und $G_2 = (V_2, E_2)$

Lösung: Entscheidung „ja“, falls gilt:

G_1 und G_2 sind isomorph, d.h. es gibt eine Abbildung $f : V_1 \rightarrow V_2$ mit der Eigenschaft $(v, w) \in E_1 \Leftrightarrow (f(v), f(w)) \in E_2$.

Beispielsweise sind die folgenden beiden Graphen isomorph:



Die folgende Sprachklasse besteht aus Sprachen, deren Komplemente in **NP** liegen:

$$\mathbf{co-NP} = \{ \Sigma^* \setminus L \mid L \text{ ist eine Sprache über } \Sigma \text{ und } L \in \mathbf{NP} \}.$$

Für das Primzahlproblem PRIMES kann man zeigen, daß es in **NP** und in **co-NP** liegt (es wird vermutet, daß PRIMES in **P** liegt):

Primzahlproblem (PRIMES):

Instanz: $n \in \mathbf{N}$ in binärer Darstellung, $\text{size}(n) = \log(n)$.

Lösung: Entscheidung „ja“, falls n eine Primzahl ist.

Lösung: Entscheidung „ja“, falls n eine Primzahl ist.

Hierzu wird ein zahlentheoretischer Satz benötigt, der an dieser Stelle angeführt wird und dessen Beweis in der angegebenen Literatur nachgelesen werden kann:

Satz 5.5-3:

Die Zahl $n \in \mathbf{N}$ mit $n > 2$ ist genau dann eine Primzahl, wenn es eine Zahl $a \in \mathbf{N}$ mit $1 < a \leq n-1$ gibt, die folgende Eigenschaften besitzt:

- (i) $a^{n-1} \equiv 1 \pmod{n}$
- (ii) $a^{(n-1)/p} \not\equiv 1 \pmod{n}$ für jeden Primteiler p von $n-1$.

Satz 5.5-4:

$\text{PRIMES} \in \mathbf{NP} \cap \mathbf{co-NP}$.

Beweis:

Zu zeigen ist, daß es sowohl für $L_1 = \{n \mid n \in \mathbf{N} \text{ und } n \text{ ist eine Primzahl}\}$ als auch für $L_2 = \{n \mid n \in \mathbf{N} \text{ und } n \text{ ist keine Primzahl}\}$ jeweils einen polynomiell zeitbeschränkten Verifizierer gibt. Der Entwurf des Verifizierers für L_1 basiert auf Satz 5.5-3, ein Verifizierer für L_2 ergibt sich direkt aus den Eigenschaften von zusammengesetzten Zahlen. Beide Verifizierer werden hier in Form von Pseudocode beschrieben.

Ein Verifizierer V_{L_2} für $L_2 = \{n \mid n \in \mathbf{N} \text{ und } n \text{ ist keine Primzahl}\}$ erhält als Eingabeinstanz eine Zahl $n \in \mathbf{N}$ (in Binärkodierung) und einen Beweis B , der in diesem Fall ebenfalls eine Zahl $B \in \mathbf{N}$ (in Binärkodierung) mit $\text{size}(B) \leq \text{size}(n)$ bzw. $1 < B \leq n-1$ ist. Die Arbeitsweise von V_{L_2} lautet:

```

FUNCTION  $V_{L_2}$  ( $n$  : INTEGER;
                 $B$  : INTEGER) : ...;

BEGIN {  $V_{L_2}$  }
  IF  $1 < B \leq n-1$                                 { Zeile 1 }
  THEN IF  $(n \bmod B) = 0$                             { Zeile 2 }
    THEN  $V_{L_2} := \text{ja}$ 
    ELSE  $V_{L_2} := \text{nein}$ 
  ELSE  $V_{L_2} := \text{nein};$ 
END {  $V_{L_2}$  };

```

Die Eingabe n für V_{L_2} belegt k Bits mit $k \in O(\log(n))$, $k = \text{size}(n)$. Die Überprüfung in Zeile 1 benötigt $O((\log(n))) = O(k)$ Operationen. Die Überprüfung in Zeile 2 (die Berechnung des

Wertes $(n \bmod B)$ kann mit $O((\log(n))^2) = O(k^2)$ vielen Bitoperationen durchgeführt werden; denn $(n \bmod B) = n - (n \text{ DIV } B) \cdot B$, und alle arithmetische Operationen benötigen nach Satz 2.1-3 höchstens quadratisch viele Bitoperationen. Daher arbeitet V_{L_2} in polynomieller Zeit, gemessen in der Größe der Eingabe n .

Ist n keine Primzahl, dann besitzt n einen Teiler $a \in \mathbf{N}$ mit $1 < a \leq n-1$. Gibt man diesen Wert (Beweis) in V_{L_2} für den Parameter B ein, so antwortet V_{L_2} mit $V_{L_2}(n, a) = \text{ja}$.

Ist n eine Primzahl, dann teilt kein $a \in \mathbf{N}$ mit $1 < a \leq n-1$ die Zahl n , d.h. für alle $a \in \mathbf{N}$ mit $1 < a \leq n-1$ ist $(n \bmod a) \neq 0$. Daher wird V_{L_2} bei Eingabe von n und eines beliebigen Beweises B immer die Antwort $V_{L_2}(n, B) = \text{nein}$ geben.

Insgesamt ist damit $L_2 \in \mathbf{NP}$ gezeigt, also $\text{PRIMES} \in \mathbf{co-NP}$.

Ein Verifizierer V_{L_1} für $L_1 = \{n \mid n \in \mathbf{N} \text{ und } n \text{ ist eine Primzahl}\}$ ist komplizierter zu konstruieren und setzt die in Satz 5.5-3 beschriebene Charakterisierung von Primzahlen um.

```

FUNCTION  $V_{L_1}$  ( $n$  : INTEGER;
                 $B$  : INTEGER) : ...;

BEGIN {  $V_{L_1}$  }
  IF  $n = 2$ 
  THEN  $V_{L_1} := \text{ja}$ 
  ELSE IF (NOT ( $1 < B \leq n-1$ )) OR ( $B^{n-1} \neq 1 \pmod{n}$ )           { Zeile 1 }
  THEN  $V_{L_1} := \text{nein}$ 
  ELSE BEGIN
    erzeuge  $k \leq \log(n-1)$  Zahlen  $p_1, \dots, p_k$  mit  $3 \leq p_i \leq (n-1)/2$ 
    und  $k$  Zahlen  $B_1, \dots, B_k$  mit  $1 < B_i \leq p_i - 1$ 
    für  $i = 1, \dots, k$ ;                                           { Zeile 2 }
    IF NOT( $(n-1 \bmod p_i) \neq 0$  für  $i = 1, \dots, k$ )                 { Zeile 3 }
    THEN  $V_{L_1} := \text{nein}$ 
    ELSE BEGIN
      IF (( $V_{L_1}(p_i, B_i) = \text{ja}$ ) für  $i = 1, \dots, k$ )           { Zeile 4 }
      THEN IF ( $B^{(n-1)/p_i} \neq 1 \pmod{n}$ ) für  $i = 1, \dots, k$ ) { Zeile 5 }
      THEN  $V_{L_1} := \text{ja}$ 
      ELSE  $V_{L_1} := \text{nein}$ 
      ELSE  $V_{L_1} := \text{nein}$ ;
    END;
  END;

```

```

                END;
END    {  $V_{L_1}$  };

```

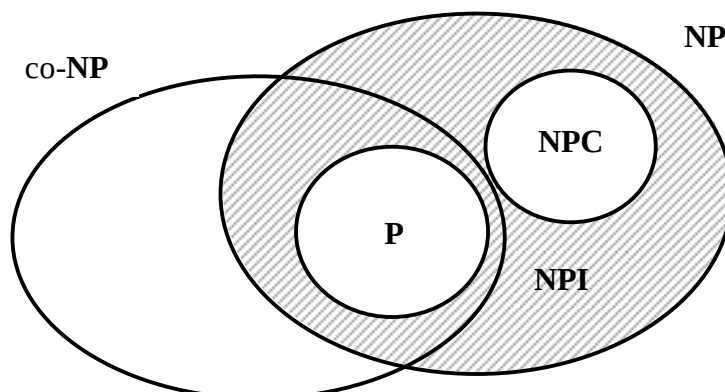
Für die Korrektheit ist zu beachten, daß in Zeile 2 auf nichtdeterministische Weise die Menge der Primfaktoren von $n-1$ erzeugt wird. Die Werte B_i sind die zugehörigen Beweise. Die Anzahl der Primfaktoren von $n-1$ ist durch $\log(n-1)$ beschränkt. Außerdem gilt in Zeile 2 wegen $n \geq 3$, daß jeder Primfaktor von $n-1$ die Abschätzung $3 \leq p_i \leq (n-2)/2$ erfüllt. In Zeile 4 wird also V_{L_1} mit kleineren Werten rekursiv aufgerufen, deren Stellenzahl echt kleiner als die Stellenzahl von n ist. In Zeile 3 wird geprüft, ob die erzeugten Zahlen p_i Teiler von $n-1$ sind, und in Zeile 4 wird geprüft, ob die Zahlen p_i Primzahlen sind. Zeile 5 prüft Bedingung (ii) aus Satz 5.5-3. Es läßt sich zeigen (siehe angegebene Literatur), daß die Berechnung in den Zeilen 1 und 5 durch ein Laufzeitverhalten der Ordnung $O((\log(n))^3)$ abgeschätzt werden kann. Insgesamt läßt sich damit zeigen, daß die Laufzeit von V_{L_1} polynomiell in $size(n)$ ist. ///

Man hat für viele Probleme in **co-NP** jedoch nicht nachweisen können, daß sie in **NP** liegen (PRIMES gehört nicht dazu). Daher wird angenommen (obwohl es noch keinen Beweis dafür gibt), daß $\mathbf{NP} \neq \mathbf{co-NP}$ gilt. Aus der Gültigkeit von $\mathbf{NP} \neq \mathbf{co-NP}$ würde übrigens folgen, daß $\mathbf{P} \neq \mathbf{NP}$ ist, da wegen der Abgeschlossenheit der Klasse **P** gegenüber Komplementbildung $\mathbf{P} = \mathbf{co-P}$ gilt.

Satz 5.5-5:

Gibt es eine Sprache $L \in \mathbf{NPC}$ mit $\Sigma^* \setminus L \in \mathbf{NP}$, dann ist $\mathbf{NP} = \mathbf{co-NP}$.

Unter den Annahmen $\mathbf{P} \neq \mathbf{NP}$ und $\mathbf{NP} \neq \mathbf{co-NP}$ ergibt sich folgendes Gesamtbild der beschriebenen Klassen:



6 Approximation von Optimierungsaufgaben

Gegeben sei das Optimierungsproblem Π :

Instanz: 1. $x \in \Sigma_{\Pi}^*$

2. Spezifikation einer Funktion SOL_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ eine Menge zulässiger Lösungen zuordnet
3. Spezifikation einer Zielfunktion m_{Π} , die jedem $x \in \Sigma_{\Pi}^*$ und $y \in \text{SOL}_{\Pi}(x)$ den Wert $m_{\Pi}(x, y)$ einer zulässigen Lösung zuordnet
4. $\text{goal}_{\Pi} \in \{\min, \max\}$, je nachdem, ob es sich um ein Minimierungs- oder ein Maximierungsproblem handelt.

Lösung: $y^* \in \text{SOL}_{\Pi}(x)$ mit $m_{\Pi}(x, y^*) = \min\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Minimierungsproblem (d.h. $\text{goal}_{\Pi} = \min$) bzw. $m_{\Pi}(x, y^*) = \max\{m_{\Pi}(x, y) \mid y \in \text{SOL}_{\Pi}(x)\}$ bei einem Maximierungsproblem (d.h. $\text{goal}_{\Pi} = \max$).

Der Wert $m_{\Pi}(x, y^*)$ einer optimalen Lösung wird auch mit $m_{\Pi}^*(x)$ bezeichnet.

Im folgenden (und in den vorhergehenden Beispielen) werden Optimierungsproblemen untersucht, die auf der Grenze zwischen praktischer Lösbarkeit (tractability) und praktischer Unlösbarkeit (intractability) stehen. In Analogie zu Entscheidungsproblemen in **NP** bilden diese die Klasse **NPO**:

Das Optimierungsproblem Π gehört zur **Klasse NPO**, wenn gilt:

1. Die Menge der Instanzen $x \in \Sigma_{\Pi}^*$ ist in polynomieller Zeit entscheidbar
2. Es gibt ein Polynom q mit der Eigenschaft: für jedes $x \in \Sigma_{\Pi}^*$ und jede zulässige Lösung $y \in \text{SOL}_{\Pi}(x)$ gilt $|y| \leq q(|x|)$, und für jedes y mit $|y| \leq q(|x|)$ ist in polynomieller Zeit entscheidbar, ob $y \in \text{SOL}_{\Pi}(x)$ ist
3. Die Zielfunktion m_{Π} ist in polynomieller Zeit berechenbar.

Alle bisher behandelten Beispiele für Optimierungsprobleme liegen in der Klasse **NPO**. Dazu gehören insbesondere auch solche, deren zugehöriges Entscheidungsproblem **NP**-vollständig ist.

Den Zusammenhang zwischen den Klassen **NPO** und **NP** beschreibt der Satz

Satz 6-1:

Für jedes Optimierungsproblem in **NPO** ist das zugehörige Entscheidungsproblem in **NP**.

Beweis:

Der Beweis wird hier nur für ein Maximierungsproblem erbracht, für ein Minimierungsproblem kann man in ähnlicher Weise argumentieren.

Es sei Π ein Maximierungsproblem in **NPO** (die im Beweis verwendeten Funktionen werden in den obigen Definitionen benannt). Das zugehörige Entscheidungsproblem sei Π_{Ent} . Hierbei soll bei Vorlage einer Instanz $[x, K]$ mit $x \in \Sigma_{\Pi}^*$ und $K \in \mathbf{N}$ genau dann die Entscheidung $x \in L_{\Pi_{Ent}}$ getroffen werden, wenn $m_{\Pi}^*(x) \geq K$ ist. Um zu zeigen, daß Π_{Ent} in **NP** liegt, ist ein Verifizierer **V** für Π_{Ent} anzugeben, der bei Eingabe einer Instanz $[x, K]$ und eines Beweises B diesen in polynomieller Zeit verifiziert. Ein Beweis ist hier eine Zeichenkette, die über dem Alphabet Σ_0 gebildet wird, mit dem man auch die zulässigen Lösungen von x formuliert.

Der Verifizierer **V** für Π_{Ent} wird durch den folgenden Pseudocode gegeben:

```

FUNCTION V ( [x, K] : ...;
              B   : ... ) : ...;
{ x ∈ ΣΠ*, K ∈ N, B ∈ Σ0* }

BEGIN { V }
  IF ( |B| ≤ q(|x|) ) AND ( B ∈ SOLΠ(x) )           { Zeile 1 }
  THEN IF mΠ(x, B) ≥ K THEN V := „ja“              { Zeile 2 }
        ELSE V := „nein“
  ELSE V := „nein“;
END { V };

```

Zu zeigen ist

1. **V** arbeitet in polynomieller Zeit, gemessen in $|x|$
2. $[x, K] \in L_{\Pi_{Ent}} \Leftrightarrow$ es gibt $B_x \in \Sigma_0^*$ mit $V(x, B_x) = \text{ja}$.

Zu 1.: Da Π in **NPO** ist, lassen sich die Berechnungen in den Zeilen 1 und 2 in polynomieller Zeit ausführen; dieses wird durch die Punkte 2. und 3. in der Definition der Klasse **NPO** gesichert.

Zu 2.: Es sei $[x, K] \in L_{\Pi_{Ent}}$. Dann gibt es eine optimale Lösung $y^* \in \text{SOL}_{\Pi}(x)$ mit $m_{\Pi}^*(x) = m_{\Pi}(x, y^*) \geq K$. Da Π in **NPO** ist, gilt $|y^*| \leq q(|x|)$. Bei Eingabe von $[x, K]$ und y^* in **V** ergibt sich $\mathbf{V}([x, K], y^*) = \text{ja}$.

Es gelte $[x, K] \notin L_{\Pi_{Ent}}$. Dann gilt für jedes $y \in \text{SOL}_{\Pi}(x)$: $m_{\Pi}(x, y) \leq m_{\Pi}^*(x) < K$. Es sei $B \in \Sigma_0^*$ mit $|B| \leq q(|x|)$. Ist $B \in \text{SOL}_{\Pi}(x)$, dann antwortet **V** wegen $m_{\Pi}(x, B) < K$ in Zeile 2 mit $\mathbf{V}([x, K], B) = \text{nein}$. Ist $B \notin \text{SOL}_{\Pi}(x)$, dann antwortet **V** wegen Zeile 1 mit $\mathbf{V}([x, K], B) = \text{nein}$. ///

Analog zur Definition der Klasse **P** innerhalb **NP** lässt sich innerhalb **NPO** eine Klasse **PO** definieren:

Ein Optimierungsproblem Π aus **NPO** gehört zur **Klasse PO**, wenn es einen polynomiell zeitbeschränkten Algorithmus gibt, der für jede Instanz $x \in \Sigma_{\Pi}^*$ eine optimale Lösung $y^* \in \text{SOL}_{\Pi}(x)$ zusammen mit dem optimalen Wert $m_{\Pi}^*(x)$ der Zielfunktion ermittelt.

Offensichtlich ist $\mathbf{PO} \subseteq \mathbf{NPO}$.

Es gilt:

Satz 6-2:

Ist $\mathbf{P} \neq \mathbf{NP}$, dann ist $\mathbf{PO} \neq \mathbf{NPO}$.

Im Laufe dieses Kapitels wird die Struktur der Klasse **NPO** genauer untersucht. Im folgenden liegen daher alle behandelten Optimierungsprobleme in **NPO**.

Bei Optimierungsaufgaben gibt man sich häufig mit Näherungen an die optimale Lösung zufrieden, insbesondere dann, wenn diese „leicht“ zu berechnen sind und vom Wert der optimalen Lösung nicht zu sehr abweichen. Unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ gibt es für ein Optimierungsproblem, dessen zugehöriges Entscheidungsproblem **NP**-vollständig ist, kein Verfahren, das eine optimale Lösung in polynomieller Laufzeit ermittelt. Gerade diese Probleme sind in der Praxis jedoch häufig von großem Interesse.

Für eine Instanz $x \in \Sigma_{\Pi}^*$ und für eine zulässige Lösung $y \in \text{SOL}_{\Pi}(x)$ bezeichnet

$$D(x, y) = |m_{\Pi}^*(x) - m_{\Pi}(x, y)|$$

den **absoluten Fehler von y bezüglich x** .

Ein Algorithmus **A** ist ein **Approximationsalgorithmus (Näherungsalgorithmus)** für Π , wenn er bei Eingabe von $x \in \Sigma_{\Pi}^*$ eine zulässige Lösung liefert, d.h. wenn $A(x) \in \text{SOL}_{\Pi}(x)$ gilt. **A** heißt **absoluter Approximationsalgorithmus (absoluter Näherungsalgorithmus)**, wenn es eine Konstante k gibt mit $D(x, A(x)) = |m_{\Pi}^*(x) - m_{\Pi}(x, A(x))| \leq k$. Das Optimierungsproblem Π heißt in diesem Fall **absolut approximierbar**.

Ist Π ein Optimierungsproblem, dessen zugehöriges Entscheidungsproblem **NP**-vollständig ist, so sucht man natürlich nach absoluten Approximationsalgorithmen für Π , die polynomielle Laufzeit aufweisen und für die der Wert $D(x, A(x))$ möglichst klein ist. Das folgende Beispiel zeigt jedoch, daß unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ nicht jedes derartige Problem absolut approximierbar ist. Dazu werde der folgende Spezialfall des 0/1-Rucksack-Maximierungsproblems betrachtet:

Das ganzzahlige 0/1-Rucksack-Maximierungsproblem

Instanz: 1. $I = (A, M)$

$A = \{a_1, \dots, a_n\}$ ist eine Menge von n Objekten und $M \in \mathbf{N}$ die „Rucksackkapazität“. Jedes Objekt a_i , $i = 1, \dots, n$, hat die Form $a_i = (w_i, p_i)$; hierbei bezeichnet $w_i \in \mathbf{N}$ das Gewicht und $p_i \in \mathbf{N}$ den Wert (Profit) des Objekts a_i .

$$\text{size}(I) = n$$

$$2. \text{ SOL}(I) = \left\{ (x_1, \dots, x_n) \mid x_i = 0 \text{ oder } x_i = 1 \text{ für } i = 1, \dots, n \text{ und } \sum_{i=1}^n x_i \cdot w_i \leq M \right\}$$

$$3. m(I, (x_1, \dots, x_n)) = \sum_{i=1}^n x_i \cdot p_i \text{ für } (x_1, \dots, x_n) \in \text{SOL}(I)$$

$$4. \text{ goal} = \max$$

Lösung: Eine Folge $x_1^*, \dots, x_n^*$ von Zahlen mit

$$(1) x_i^* = 0 \text{ oder } x_i^* = 1 \text{ für } i = 1, \dots, n$$

$$(2) \sum_{i=1}^n x_i^* \cdot w_i \leq M$$

$$(3) m(I, (x_1^*, \dots, x_n^*)) = \sum_{i=1}^n x_i^* \cdot p_i \text{ ist maximal unter allen möglichen Auswahlen } x_1, \dots, x_n, \text{ die (1) und (2) erfüllen.}$$

Das zugehörige Entscheidungsproblem ist **NP**-vollständig, so daß dieses Optimierungsproblem unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ keinen polynomiell zeitbeschränkten Lösungsalgorithmus (der eine *optimale* Lösung ermittelt) besitzt. Es gilt sogar:

Satz 6-3:

Es sei k eine vorgegebene Konstante. Unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ gibt es keinen polynomiell zeitbeschränkten Approximationsalgorithmus $\mathbf{A}$ für das ganzzahlige 0/1-Rucksack-Maximierungsproblem, der bei Eingabe einer Instanz $I = (A, M)$ eine zulässige Lösung $\mathbf{A}(I) \in \text{SOL}(I)$ berechnet, für deren absoluter Fehler

$$D(I, \mathbf{A}(I)) = |m^*(I) - m(I, \mathbf{A}(I))| \leq k \quad \text{gilt.}$$

Beweis:

Die verwendete Beweistechnik ist auch auf andere Probleme übertragbar.

Es wird angenommen, daß es (bei Vorgabe der Konstanten k) einen derartigen polynomiell zeitbeschränkten Approximationsalgorithmus $\mathbf{A}$ gibt und zeigt dann, daß dieser (mit einer einfachen Modifikation) dazu verwendet werden kann, in polynomieller Zeit eine optimale Lösung für das ganzzahlige 0/1-Rucksack-Maximierungsproblem zu ermitteln. Damit kann man dann das zu diesem Problem gehörende Entscheidungsproblem in polynomieller Zeit lösen. Da dieses **NP**-vollständig ist, folgt $\mathbf{P} = \mathbf{NP}$ (vgl. Satz 5.1-1).

Bei Annahme der Existenz von $\mathbf{A}$ kann ein Algorithmus $\tilde{\mathbf{A}}$ definiert werden, der folgendermaßen arbeitet:

Aus einer Eingabe einer Instanz $I = (A, M)$ mit $A = \{(w_1, p_1), \dots, (w_n, p_n)\}$ für das ganzzahlige 0/1-Rucksack-Maximierungsproblem erzeugt $\tilde{\mathbf{A}}$ eine Instanz $\tilde{I} = (\tilde{A}, M)$ mit $\tilde{A} = \{(w_1, (k+1) \cdot p_1), \dots, (w_n, (k+1) \cdot p_n)\}$ (es werden dabei also lediglich alle Profite p_i durch $(k+1) \cdot p_i$ ersetzt). Diese neue Instanz $\tilde{I} = (\tilde{A}, M)$ wird in $\mathbf{A}$ eingegeben. $\mathbf{A}$ ermittelt eine approximative Lösung $\mathbf{A}(\tilde{I}) = (x_1, \dots, x_n)$ mit $\sum_{i=1}^n x_i \cdot w_i \leq M$ und

$$|m^*(\tilde{I}) - m(\tilde{I}, \mathbf{A}(\tilde{I}))| = \left| m^*(\tilde{I}) - \sum_{i=1}^n x_i \cdot (k+1) \cdot p_i \right| \leq k \quad (\text{Ungleichung } (*)).$$

$\tilde{\mathbf{A}}$ gibt die von $\mathbf{A}$ bei Eingabe von $\tilde{I}$ ermittelte Lösung $\tilde{\mathbf{A}}(I) = \mathbf{A}(\tilde{I}) = (x_1, \dots, x_n)$ mit (dem Wert der Zielfunktion) $m(I, \tilde{\mathbf{A}}(I)) = \frac{m(\tilde{I}, \mathbf{A}(\tilde{I}))}{k+1}$ aus. Es ist $\tilde{\mathbf{A}}(I) \in \text{SOL}(I)$. Es gilt sogar $m(I, \tilde{\mathbf{A}}(I)) = m^*(I)$, d.h. $\tilde{\mathbf{A}}$ liefert in polynomieller Zeit (da $\mathbf{A}$ in polynomieller Zeit arbeitet) eine optimale Lösung für die Eingabeinstanz I :

Da $m^*(\tilde{I})$ ein Vielfaches von $k+1$ ist, denn jeder Profit in $\tilde{I}$ ist ein Vielfaches von $k+1$, kann man den Faktor $k+1$ auf der linken Seite des $\leq$ -Zeichens in der Ungleichung (*) ausklammern und erhält

$$\left| m^*(\tilde{I}) - m(\tilde{I}, \mathbf{A}(\tilde{I})) \right| = \left| m^*(\tilde{I}) - \sum_{i=1}^n x_i \cdot (k+1) \cdot p_i \right| = (k+1) \cdot R \leq k$$

mit einer natürlichen Zahl R . Diese Ungleichung ist nur mit $R=0$ möglich, so daß

$$\left| m^*(\tilde{I}) - m(\tilde{I}, \mathbf{A}(\tilde{I})) \right| = 0 \text{ folgt, d.h. } \mathbf{A}(\tilde{I}) \text{ ist eine optimale Lösung für } \tilde{I}.$$

Jede zulässige Lösung für $\tilde{I}$ ist auch eine zulässige Lösung für I , jedoch mit dem $(k+1)$ -fachen Profit. Umgekehrt ist jede zulässige Lösung für I eine zulässige Lösung für $\tilde{I}$. Damit folgt die Optimalität von $\tilde{\mathbf{A}}(I)$ (dazu ist $m(I, \tilde{\mathbf{A}}(I)) \geq m(I, y)$ für jedes $y \in \text{SOL}(I)$ zu zeigen):

$$m(I, \tilde{\mathbf{A}}(I)) = \frac{m(\tilde{I}, \mathbf{A}(\tilde{I}))}{k+1} = \frac{m^*(\tilde{I})}{k+1} \geq \frac{m(\tilde{I}, y)}{k+1} = m(I, y) \text{ für jedes } y \in \text{SOL}(I). \quad ///$$

Die Forderung nach der Garantie der Einhaltung eines absoluten Fehlers ist also häufig zu stark. Es bietet sich daher an, einen Approximationsalgorithmus nach seiner relativen Approximationsgüte zu beurteilen.

Es sei ein Π wieder ein Optimierungsproblem und $\mathbf{A}$ ein Approximationsalgorithmus für Π (siehe oben), der eine zulässige Lösung $\mathbf{A}(x)$ ermittelt. Ist $x \in \Sigma_\Pi^*$ eine Instanz von Π , so gilt trivialerweise $m(x, \mathbf{A}(x)) \leq m_\Pi^*(x)$ bei einem Maximierungsproblem bzw. $m_\Pi^*(x) \leq m(x, \mathbf{A}(x))$ bei einem Minimierungsproblem.

Die **relative Approximationsgüte** $R_A(x)$ von $\mathbf{A}$ bei Eingabe einer Instanz $x \in \Sigma_\Pi^*$ wird definiert durch

$$R_A(x) = \frac{m_\Pi^*(x)}{m(x, \mathbf{A}(x))} \text{ bei einem Maximierungsproblem}$$

bzw.

$$R_A(x) = \frac{m(x, \mathbf{A}(x))}{m_\Pi^*(x)} \text{ bei einem Minimierungsproblem.}$$

Bemerkung: Um nicht zwischen Maximierungs- und Minimierungsproblem in der Definition unterscheiden zu müssen, kann man die relative Approximationsgüte von $\mathbf{A}$ bei Eingabe einer Instanz $x \in \Sigma_\Pi^*$ auch durch

$$R_A(x) = \max \left\{ \frac{m_\Pi^*(x)}{m(x, \mathbf{A}(x))}, \frac{m(x, \mathbf{A}(x))}{m_\Pi^*(x)} \right\}$$

definieren.

Offensichtlich gilt immer $1 \leq R_A(x)$. Je dichter $R_A(x)$ bei 1 liegt, um so besser ist die Approximation.

Die Aussage „ $R_A(x) \leq c$ “ mit einer Konstanten $c \geq 1$ für alle Instanzen $x \in \Sigma_\Pi^*$ bedeutet bei einem Maximierungsproblem $1 \leq \frac{m_\Pi^*(x)}{m(x, \mathbf{A}(x))} \leq c$, also $m(x, \mathbf{A}(x)) \geq \frac{1}{c} \cdot m_\Pi^*(x)$, d.h. der Approximationsalgorithmus liefert eine Lösung, die sich mindestens um $\frac{1}{c}$ dem Optimum nähert. Gewünscht ist also ein möglichst großer Wert von $\frac{1}{c}$ bzw. ein Wert von c , der möglichst klein (d.h. dicht bei 1) ist.

Bei einem Minimierungsproblem impliziert die Aussage „ $R_A(x) \leq c$ “ die Beziehung $m(x, \mathbf{A}(x)) \leq c \cdot m_\Pi^*(x)$, d.h. der Wert der Approximation überschreitet das Optimum um höchstens das c -fache. Auch hier ist also ein Wert von c gewünscht, der möglichst klein (d.h. dicht bei 1) ist.

6.1 Relativ approximierbare Probleme

Es sei $r \geq 1$. Der Approximationsalgorithmus $\mathbf{A}$ für das Optimierungsproblem Π aus **NPO** heißt **r -approximativer Algorithmus**, wenn $R_A(x) \leq r$ für jede Instanz $x \in \Sigma_\Pi^*$ gilt.

Bemerkung: Es sei $\mathbf{A}$ ein Approximationsalgorithmus für das Minimierungsproblem Π , und es gelte $m(x, \mathbf{A}(x)) \leq r \cdot m_\Pi^*(x) + k$ für alle Instanzen $x \in \Sigma_\Pi^*$ (mit Konstanten r und k). Dann ist $\mathbf{A}$ lediglich $(r+k)$ -approximativ und nicht etwa r -approximativ, jedoch **asymptotisch** r -approximativ (siehe Kapitel 6.2).

Die **Klasse APX** besteht aus denjenigen Optimierungsproblemen aus **NPO**, für die es einen r -approximativen Algorithmus für ein $r \geq 1$ gibt.

Offensichtlich ist **APX** $\subseteq$ **NPO**.

Das folgende Optimierungsproblem liegt in **APX**:

Binpacking-Minimierungsproblem:

- Instanz:
1. $I = [a_1, \dots, a_n]$
 $a_1, \dots, a_n$ („Objekte“) sind rationale Zahlen mit $0 < a_i \leq 1$ für $i = 1, \dots, n$
 $\text{size}(I) = n$
 2. $\text{SOL}(I) = \left\{ [B_1, \dots, B_k] \mid \begin{array}{l} [B_1, \dots, B_k] \text{ ist eine Partition (disjunkte Zerlegung)} \\ \text{von } I \text{ mit } \sum_{a_i \in B_j} a_i \leq 1 \text{ für } j = 1, \dots, k \end{array} \right\}$, d.h.
 in einer zulässigen Lösung werden die Objekte so auf k „Behälter“ der Höhe 1 verteilt, daß kein Behälter „überläuft“.
 3. Für $[B_1, \dots, B_k] \in \text{SOL}(I)$ ist $m(I, [B_1, \dots, B_k]) = k$, d.h. als Zielfunktion wird die Anzahl der benötigten Behälter definiert
 4. $\text{goal} = \min$

Lösung: Eine Partition der Objekte in möglichst wenige Teile $B_1, \dots, B_{k^*}$ und (implizit) die Anzahl k^* der benötigten Teile.

Bemerkung: Da das Laufzeitverhalten der folgenden Approximationsalgorithmen nicht von den Größen der in den Instanzen vorkommenden Objekte abhängt, sondern nur von deren Anzahl, kann man für eine Eingabeinstanz $I = [a_1, \dots, a_n]$ als Größe den Wert $\text{size}(I) = n$ wählen.

Das Binpacking-Minimierungsproblem ist eines der am besten untersuchten Minimierungsprobleme einschließlich der Verallgemeinerungen auf mehrdimensionale Objekte. Das zugehörige Entscheidungsproblem ist **NP**-vollständig, so daß unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ kein polynomiell zeitbeschränkter Optimierungsalgorithmus erwartet werden kann.

Der folgende in Pseudocode formulierte Algorithmus approximiert eine optimale Lösung:

Nextfit-Algorithmus zur Approximation von Binpacking:

Eingabe: $I = [a_1, \dots, a_n]$,
 $a_1, \dots, a_n$ („Objekte“) sind rationale Zahlen mit $0 < a_i \leq 1$ für $i = 1, \dots, n$

Verfahren: $B_1 := a_1$;

Sind $a_1, \dots, a_i$ bereits in die Behälter $B_1, \dots, B_j$ gelegt worden, ohne daß einer dieser Behälter überläuft, so lege a_{i+1} nach B_j (in den Behälter mit dem höchsten Index), falls dadurch B_j nicht überläuft, d.h. $\sum_{a_l \in B_j} a_l \leq 1$ gilt; andernfalls lege a_{i+1} in einen neuen (bisher leeren) Behälter B_{j+1} .

Ausgabe: $\mathbf{NF}(I) =$ die benötigten nichtleeren Behälter $[B_1, \dots, B_k]$.

Satz 6.1-1:

Ist $I = [a_1, \dots, a_n]$ eine Instanz des Binpacking-Minimierungsproblems, so gilt $R_{\mathbf{NF}}(I) \leq 2$, d.h. der Nextfit-Algorithmus ist 2-approximativ ($m(I, \mathbf{NF}(I)) \leq 2 \cdot m^*(I)$). Das Binpacking -Minimierungsproblem liegt in **APX**.

Beweis:

Interessanterweise kann man die Aussage $m(I, \mathbf{NF}(I)) \leq 2 \cdot m^*(I)$ machen, ohne $m^*(I)$ oder $m(I, \mathbf{NF}(I))$ zu kennen. Dazu wird $m^*(I)$ abgeschätzt:

Für eine Instanz I habe der Nextfit-Algorithmus k Behälter ermittelt, d.h. $m(I, \mathbf{NF}(I)) = k$. Die Füllhöhe im j -ten Behälter sei u_j für $j = 1, \dots, k$. Das erste Element, das der Nextfit-Algorithmus in den Behälter B_j gelegt hat, sei b_j .

Man betrachte die beiden Behälter B_j und B_{j+1} (für $1 \leq j < k$): Da b_{j+1} nicht mehr in den Behälter B_j paßte, gilt $1 - u_j < b_j \leq u_{j+1}$ und damit $u_j + u_{j+1} > 1$. Summiert man diese $k-1$ Ungleichungen auf, so erhält man

$$(u_1 + u_2) + (u_2 + u_3) + \dots + (u_{k-1} + u_k) = u_1 + 2u_2 + \dots + 2u_{k-1} + u_k > k - 1.$$

Auf beide Seiten werden die Füllhöhen des ersten und letzten Behälters addiert, und man erhält $2 \cdot \sum_{j=1}^k u_j > k - 1 + u_1 + u_k$. Die Summe der Füllhöhen in den Behältern ist gleich der Summe der Objekte, die gepackt wurden. Damit folgt

$$k < 2 \cdot \sum_{j=1}^k u_j + 1 - (u_1 + u_k) = 2 \cdot \sum_{i=1}^n a_i + 1 - (u_1 + u_k) \leq 2 \cdot \sum_{i=1}^n a_i + 1 \text{ und } k \leq 2 \cdot \sum_{i=1}^n a_i.$$

Trivialerweise gilt $m^*(I) \geq \sum_{i=1}^n a_i$ (hier gilt „ $\geq$ “, wenn in einer optimalen Packung alle Behälter bis zur maximalen Füllhöhe 1 aufgefüllt werden). Damit ergibt sich schließlich $R_{\mathbf{NF}}(I) = k / m^*(I) \leq 2$. ///

Die Grenze 2 im Nextfit-Algorithmus ist asymptotisch optimal: es gibt Instanzen I , für die $R_{\text{NF}}(I)$ beliebig dicht an 2 herankommt. Dazu betrachte man etwa eine Eingabeinstanz I mit $n = 2m$ vielen Objekten, wobei m eine gerade Zahl ist: $I = [a_1, \dots, a_{2m}]$. Die Werte a_i seien definiert durch $a_i = \begin{cases} 1/2 - 1/3m & \text{für ungerades } i \\ 1/m & \text{für gerades } i \end{cases}$. Mit dieser Instanz gilt $m(I, \text{NF}(I)) = m$ und $m^*(I) = m/2 + 1$, so daß $R_{\text{NF}}(I) = 2 \cdot \frac{m}{m+2}$ folgt, und dieser Wert kann für große m beliebig dicht an 2 herangehen.

Es gilt sogar eine genauere Abschätzung der relativen Approximationsgüte, falls die Größe aller Objekte beschränkt ist: Mit $a_{\max} = \max\{a_i \mid i = 1, \dots, n\}$ ist

$$m(I, \text{NF}(I)) \leq \begin{cases} (1 + a_{\max} / (1 - a_{\max})) \cdot m^*(I) & \text{falls } 0 < a_{\max} \leq 1/2 \\ 2 \cdot m^*(I) & \text{falls } 1/2 < a_{\max} \leq 1 \end{cases}$$

Zu beachten ist, daß der Nextfit-Algorithmus ein online-Algorithmus ist, d.h. die Objekte der Reihenfolge nach inspiziert, und sofort eine Entscheidung trifft, in welchen Behälter ein Objekt zu legen ist, ohne alle Objekte gesehen zu haben. Die Laufzeit des Nextfit-Algorithmus bei einer Eingabeinstanz der Größe n liegt in $O(n)$.

Eine asymptotische Verbesserung der Approximation liefert der Firstfit-Algorithmus, der ebenfalls ein online-Algorithmus ist und bei geeigneter Implementierung ein Laufzeitverhalten der Ordnung $O(n \cdot \log(n))$ hat:

Firstfit-Algorithmus zur Approximation von Binpacking:

Eingabe: $I = [a_1, \dots, a_n]$,
 $a_1, \dots, a_n$ („Objekte“) sind rationale Zahlen mit $0 < a_i \leq 1$ für $i = 1, \dots, n$

Verfahren: $B_1 := a_1$;
 Sind $a_1, \dots, a_i$ bereits in die Behälter $B_1, \dots, B_j$ gelegt worden, ohne daß einer dieser Behälter überläuft, so lege a_{i+1} in den Behälter unter $B_1, \dots, B_j$ mit dem kleinsten Index, in den das Objekt a_{i+1} noch paßt, ohne daß er überläuft. Falls

es einen derartigen Behälter unter $B_1, \dots, B_j$ nicht gibt, lege a_{i+1} in einen neuen (bisher leeren) Behälter B_{j+1} .

Ausgabe: $\mathbf{FF}(I)$ = die benötigten nichtleeren Behälter $[B_1, \dots, B_k]$.

Satz 6.1-2:

Ist $I = [a_1, \dots, a_n]$ eine Instanz des Binpacking-Minimierungsproblems, so gilt $m(I, \mathbf{FF}(I)) \leq 1,7 \cdot m^*(I) + 2$.

Der Beweis verwendet eine „Gewichtung“ der Objekte der Eingabeinstanz. Wegen der Länge des Beweises muß auf die Literatur verwiesen werden.

Die Grenze 1,7 im Firstfit-Algorithmus ist ebenfalls asymptotisch optimal: es gibt Instanzen I mit beliebig großem Wert $m^*(I)$, für die $m(I, \mathbf{FF}(I)) \geq 1,7 \cdot (m^*(I) - 1)$ gilt. Daher kommt $R_{\mathbf{FF}}(I)$ asymptotisch beliebig dicht an 1,7 heran.

Zu beachten ist weiterhin, daß der Firstfit-Algorithmus kein 1,7-approximativer Algorithmus ist, sondern nur asymptotisch 1,7-approximativ (siehe Kapitel 6.2) ist.

Eine weitere asymptotische Verbesserung erhält man, indem man die Objekte vor der Aufteilung auf Behälter nach absteigender Größe sortiert. Die entstehenden Approximationsalgorithmen sind dann jedoch offline-Algorithmen, da zunächst alle Objekte vor der Aufteilung bekannt sein müssen.

FirstfitDecreasing-Algorithmus zur Approximation von Binpacking:

Eingabe: $I = [a_1, \dots, a_n]$,
 $a_1, \dots, a_n$ („Objekte“) sind rationale Zahlen mit $0 < a_i \leq 1$ für $i = 1, \dots, n$

Verfahren: Sortiere die Objekte $a_1, \dots, a_n$ nach absteigender Größe. Die Objekte erhalten im folgenden wieder die Benennung $a_1, \dots, a_n$, d.h. es gilt $a_1 \geq \dots \geq a_n$.

$B_1 := a_1$;

Sind $a_1, \dots, a_i$ bereits in die Behälter $B_1, \dots, B_j$ gelegt worden, ohne daß einer dieser Behälter überläuft, so lege a_{i+1} in den Behälter unter $B_1, \dots, B_j$ mit dem kleinsten Index, in den das Objekt a_{i+1} noch paßt, ohne daß er überläuft. Falls es einen derartigen Behälter unter $B_1, \dots, B_j$ nicht gibt, lege a_{i+1} in einen neuen (bisher leeren) Behälter B_{j+1} .

Ausgabe: **FFD**(I) = die benötigten nichtleeren Behälter $[B_1, \dots, B_k]$.

Satz 6.1-3:

Ist $I = [a_1, \dots, a_n]$ eine Instanz des Binpacking-Minimierungsproblems, so gilt $m(I, \mathbf{FFD}(I)) \leq 1,5 \cdot m^*(I) + 1$.

Beweis:

Die Objekte der Eingabeinstanz $I = [a_1, \dots, a_n]$ seien nach absteigender Größe sortiert, d.h. $a_1 \geq \dots \geq a_n$. Sie werden in vier Größenordnungen eingeteilt. Dazu wird

$$A = \{a_i \mid a_i \in I \text{ und } a_i > 2/3\},$$

$$B = \{a_i \mid a_i \in I \text{ und } 2/3 \geq a_i > 1/2\},$$

$$C = \{a_i \mid a_i \in I \text{ und } 1/2 \geq a_i > 1/3\} \text{ und}$$

$$D = \{a_i \mid a_i \in I \text{ und } 1/3 \geq a_i > 0\} \text{ gesetzt.}$$

Der FirstfitDecreasing-Algorithmus habe eine Packung mit k Behältern ermittelt, d.h. $m(I, \mathbf{FFD}(I)) = k$. Die Füllhöhe des j -ten Behälters B_j sei u_j .

1. Fall: Es gibt mindestens einen Behälter, der nur Objekte aus D enthält.

Dann sind alle anderen Behälter, bis eventuell auf den Behälter B_k zu mindestens $2/3$ gefüllt. Es gilt dann

$$\sum_{i=1}^n a_i = \sum_{j=1}^{k-1} u_j + u_k \geq \sum_{j=1}^{k-1} 2/3 = 2/3 \cdot (k-1) \text{ und folglich}$$

$$m(I, \mathbf{FFD}(I)) = k \leq 3/2 \cdot \sum_{i=1}^n a_i + 1 \leq 3/2 \cdot m^*(I) + 1.$$

2. Fall: Kein Behälter enthält nur Objekte aus D .

Es sei $\tilde{I} = I \setminus D$. Alle Objekte in $\tilde{I}$ haben eine Größe von mehr als $1/3$. Dann gilt $\mathbf{FFD}(\tilde{I}) = \mathbf{FFD}(I)$, da die Objekte in D noch auf die übrigen Behälter verteilt werden können und erst dann verteilt werden, wenn alle Objekte in A , B und C verteilt sind.

Es gilt sogar $m(\tilde{I}, \mathbf{FFD}(\tilde{I})) = m^*(\tilde{I})$, d.h. der FirstfitDecreasing-Algorithmus liefert bei Eingabe von $\tilde{I}$ eine optimale Packung:

Dazu wird das Aussehen einer optimalen Packung von $\tilde{I}$ betrachtet:

- Es gibt t_A Behälter, die nur ein einziges Objekt aus A enthalten. Objekte aus B oder C passen dort nicht mehr hinein.
- Kein Behälter enthält 3 oder mehr Objekte.
- Es gibt t Behälter mit 2 Objekten, davon jeweils höchstens eines aus B , eventuell beide aus C . Die Anzahl der Behälter mit 2 Objekten, von denen eines aus B und eines aus C kommt, sei t_{BC} ; die Anzahl der Behälter mit 2 Objekten, von denen beide aus C kommt, sei t_{CC} .
- Es gibt t_B Behälter, die nur ein Objekt aus B enthalten.
- Es gibt t_C Behälter, die nur ein Objekt aus C enthalten; $t_C \leq 1$.

$$m^*(\tilde{I}) = t_A + t_{BC} + t_{CC} + t_B + t_C \text{ und } |C| = t_{BC} + 2 \cdot t_{CC} + t_C.$$

Bei der Abarbeitung von $\tilde{I}$ durch den FirstfitDecreasing-Algorithmus werden zuerst die Objekte aus $A \cup B$ gepackt; dazu werden $t_A + t_{BC} + t_B$ Behälter benötigt. Das erste Objekt $c \in C$ kommt in den Behälter, der ein Objekt aus B enthält, und zwar das größte Objekt $b \in B$ mit $b + c \leq 1$. Ein Objekt aus C kommt erst dann in einen neuen Behälter oder in einen Behälter, der bereits ein Element aus C enthält, wenn es keinen Behälter gibt, der genau ein Element aus B enthält und zu dem man es legen könnte. Daher kommen auch t_{BC} Objekte aus C in Behälter mit einem Element aus B . Für die übrigen Objekte aus C benötigt der FirstfitDecreasing-Algorithmus noch $2 \cdot t_{CC} + t_C$ Behälter. Insgesamt ergibt sich

$$m(\tilde{I}, \mathbf{FFD}(\tilde{I})) = t_A + t_{BC} + t_{CC} + t_B + t_C = m^*(\tilde{I}).$$

Damit ergibt sich $m^*(I) \leq m(I, \mathbf{FFD}(I)) = m(\tilde{I}, \mathbf{FFD}(\tilde{I})) = m^*(\tilde{I}) \leq m^*(I)$; die letzte Ungleichung folgt aus der Inklusion $\tilde{I} \subseteq I$. Das bedeutet $m(I, \mathbf{FFD}(I)) = m^*(I)$. ///

Auch der FirstfitDecreasing-Algorithmus kein 1,5-approximativer Algorithmus, sondern nur asymptotisch 1,5-approximativ (siehe Kapitel 6.2).

Eine genauere Analyse des FirstfitDecreasing-Algorithmus zeigt, daß die in Satz 6.1-3 angegebene Schranke 1,5 verbessert werden kann. Es läßt sich zeigen, daß für jede Instanz $I = [a_1, \dots, a_n]$ des Binpacking-Minimierungsproblems die Abschätzung

$m(I, \mathbf{FFD}(I)) \leq 11/9 \cdot m^*(I) + 4$ gilt. Das folgende Beispiel zeigt, daß die 11/9-Grenze nicht verbessert werden kann: Man betrachte die Instanz I , die aus $n = 5m$ vielen Objekten besteht, und zwar m Objekte der Größe $1/2 + d$ mit $0 < d < 1/40$, m Objekte der Größe $1/4 + 2d$, m Objekte der Größe $1/4 + d$ und $2m$ Objekte der Größe $1/4 - 2d$. Die Objekte dieser Instanz

sind nach absteigender Größe sortiert. Es ist $m^*(I) = 3m/2$ und $m(I, \text{FFD}(I)) = 11m/6$, also $m(I, \text{FFD}(I)) = 11/9 \cdot m^*(I)$.

BestfitDecreasing-Algorithmus zur Approximation von Binpacking:

Eingabe: $I = [a_1, \dots, a_n]$,

$a_1, \dots, a_n$ („Objekte“) sind rationale Zahlen mit $0 < a_i \leq 1$ für $i = 1, \dots, n$

Verfahren: Sortiere die Objekte $a_1, \dots, a_n$ nach absteigender Größe. Die Objekte erhalten im folgenden wieder die Benennung $a_1, \dots, a_n$, d.h. es gilt $a_1 \geq \dots \geq a_n$.

$B_1 := a_1$;

Sind $a_1, \dots, a_i$ bereits in die Behälter $B_1, \dots, B_j$ gelegt worden, ohne daß einer dieser Behälter überläuft, so lege a_{i+1} in den Behälter unter $B_1, \dots, B_j$, der den kleinsten freien Platz aufweist. Falls a_{i+1} in keinen der Behälter $B_1, \dots, B_j$ paßt, lege a_{i+1} in einen neuen (bisher leeren) Behälter B_{j+1} .

Ausgabe: $\text{BFD}(I) =$ die benötigten nichtleeren Behälter $[B_1, \dots, B_k]$.

Satz 6.1-4:

Ist $I = [a_1, \dots, a_n]$ eine Instanz des Binpacking-Minimierungsproblems, so gilt $m(I, \text{BFD}(I)) \leq \frac{11}{9} \cdot m^*(I) + 4$.

Auch der BestfitDecreasing-Algorithmus nur asymptotisch $11/9$ -approximativ (siehe Kapitel 6.2). Das obige Beispiel zeigt, daß auch für diesen Algorithmus die $11/9$ -Grenze nicht verbessert werden kann.

Die folgende Zusammenstellung zeigt noch einmal die mit den verschiedenen Approximationsalgorithmen für das Binpacking-Problem zu erzielenden Approximationsgüten. In der letzten Zeile ist dabei ein Algorithmus erwähnt, der in einem gewissen Sinne (siehe Kapitel 6.2) unter allen approximativen Algorithmen mit polynomieller Laufzeit eine optimale relative Approximationsgüte erzielt. Zu beachten ist ferner, daß im Sinne der Definition die Algorithmen Firstfit, FirstfitDecreasing und BestfitDecreasing nicht r -approximativ (mit $r = 1,7$ bzw. $r = 1,5$ bzw. $r = 1,22$), sondern nur asymptotisch r -approximativ sind.

Approximationsalgorithmus	Approximationsgüte
Nextfit	$m(I, \mathbf{NF}(I)) \leq \begin{cases} (1 + a_{\max}/(1 - a_{\max})) \cdot m^*(I) & \text{falls } 0 < a_{\max} \leq 1/2 \\ 2 \cdot m^*(I) & \text{falls } 1/2 < a_{\max} \leq 1 \end{cases}$
Firstfit	$m(I, \mathbf{FF}(I)) \leq 1,7 \cdot m^*(I) + 2$
FirstfitDecreasing	$m(I, \mathbf{FFD}(I)) \leq 1,5 \cdot m^*(I) + 1$ (Satz 6.1-3) $m(I, \mathbf{FFD}(I)) \leq 11/9 \cdot m^*(I) + 4$; $11/9 = 1,222$
BestFitDecreasing	$m(I, \mathbf{BFD}(I)) \leq 11/9 \cdot m^*(I) + 4$
Simchi-Levi, 1994	$m(I, \mathbf{SL}(I)) \leq 1,5 \cdot m^*(I)$

In Kapitel 6 wurde gezeigt, daß das 0/1-Rucksack-Maximierungsproblem unter der Annahme $\mathbf{P} \neq \mathbf{NP}$ nicht absolut approximierbar ist. Es gibt jedoch für dieses Problem einen 2-approximativen Algorithmus:

Das 0/1-Rucksackproblem als Maximierungsproblem (maximum 0/1 knapsack problem)

Instanz: 1. $I = (A, M)$

$A = \{a_1, \dots, a_n\}$ ist eine Menge von n Objekten und $M \in \mathbf{R}_{>0}$ die „Rucksackkapazität“. Jedes Objekt a_i , $i = 1, \dots, n$, hat die Form $a_i = (w_i, p_i)$; hierbei bezeichnet $w_i \in \mathbf{R}_{>0}$ das Gewicht und $p_i \in \mathbf{R}_{>0}$ den Wert (Profit) des Objekts a_i .
 $size(I) = n$

2. $SOL(I) = \left\{ (x_1, \dots, x_n) \mid x_i = 0 \text{ oder } x_i = 1 \text{ für } i = 1, \dots, n \text{ und } \sum_{i=1}^n x_i \cdot w_i \leq M \right\}$; man

beachte, daß hier nur Werte $x_i = 0$ oder $x_i = 1$ zulässig sind

3. $m(I, (x_1, \dots, x_n)) = \sum_{i=1}^n x_i \cdot p_i$ für $(x_1, \dots, x_n) \in SOL(I)$

4. $goal = \max$

Lösung: Eine Folge $x_1^*, \dots, x_n^*$ von Zahlen mit

(1) $x_i^* = 0$ oder $x_i^* = 1$ für $i = 1, \dots, n$

(2) $\sum_{i=1}^n x_i^* \cdot w_i \leq M$

- (3) $m(I, (x_1^*, \dots, x_n^*)) = \sum_{i=1}^n x_i^* \cdot p_i$ ist maximal unter allen möglichen Auswahlen $x_1, \dots, x_n$, die (1) und (2) erfüllen.

Der Approximationsalgorithmus **RUCK_APP** wird hier informell beschrieben:

1. Bei Eingabe einer Instanz $I = (A, M)$ sortiere man die Objekte nach absteigenden Werten p_i/w_i . Die Objekte werden in der durch diese Sortierung bestimmten Reihenfolge bearbeitet. Ein Objekt a_i wird dabei in den Rucksack gelegt, wenn es noch paßt; in diesem Fall wird $x_i = 1$ gesetzt. Paßt das Objekt a_i nicht, dann wird $x_i = 0$ gesetzt und zum nächsten Objekt übergegangen.
2. Es sei $p_{\max} = \max\{p_i \mid i = 1, \dots, n\}$. Man vergleiche das Ergebnis im Schritt 1 mit der Rucksackfüllung, die man erhält, wenn man nur das Objekt mit dem Gewinn $p_{\max}$ allein in den Rucksack legt. Man nehme diejenige Rucksackfüllung, die den größeren Wert der Zielfunktion liefert.

Das Ergebnis des Verfahrens ist eine 0-1-Folge $\mathbf{RUCK_APP}(I) = (x_1, \dots, x_n)$ mit dem Wert $m(I, \mathbf{RUCK_APP}(I))$ der Zielfunktion.

Satz 6.1-5:

Für jede Instanz $I = (A, M)$ des 0/1-Rucksack-Maximierungsproblems gilt:
 $1 \leq m^*(I)/m(I, \mathbf{RUCK_APP}(I)) \leq 2$, d.h. das 0/1-Rucksack-Maximierungsproblem liegt in **APX**.

Beweis:

Es sei a_j das erste Objekt, das im 1. Teil des Algorithmus **RUCK_APP** nicht mehr in den Rucksack paßt. Zu diesem Zeitpunkt hat der Rucksack eine Füllung $\bar{w} = \sum_{i=1}^{j-1} w_i \leq M$. Es gilt

also $w_j > M - \bar{w}$. Der bisher entstandene Profit ist $\bar{p} = \sum_{i=1}^{j-1} p_i$.

Es werde eine modifizierte Aufgabenstellung des 0/1-Rucksack-Maximierungsproblems betrachtet, in dem die Forderung „ $x_i = 0$ oder $x_i = 1$ “ durch „ $0 \leq x_i \leq 1$ “ ersetzt ist. Für die so modifizierte Aufgabenstellung kann im 1. Teil des Algorithmus **RUCK_APP** das Objekt a_j noch in den Rucksack gelegt werden, jedoch nur mit einem Anteil $x_j = (M - \bar{w})/w_j$. Der Rucksack ist dann gefüllt, d.h. im 1. Teil von **RUCK_APP** wird für die modifizierte Aufga-

benstellung $x_{j+1} = \dots = x_n = 0$ gesetzt. Es läßt sich zeigen, daß jetzt bereits eine optimale Lösung des modifizierten 0/1-Rucksack-Maximierungsproblems gefunden ist, und zwar mit einem Profit $m_{MOD} = \bar{p} + ((M - \bar{w})/w_j) \cdot p_j$. Da jede zulässige Lösung des (ursprünglichen) 0/1-Rucksack-Maximierungsproblems eine zulässige Lösung des modifizierten 0/1-Rucksack-Maximierungsproblems ist, folgt $m^*(I) \leq \bar{p} + ((M - \bar{w})/w_j) \cdot p_j < \bar{p} + p_j$ (die letzte Ungleichung ergibt sich aus $w_j > M - \bar{w}$).

Außerdem gilt $\bar{p} \leq \max\{\bar{p}, p_{\max}\} \leq m(I, \mathbf{RUCK_APP}(I)) \leq m^*(I)$.

Es werden zwei Fälle unterschieden:

1. Fall: $p_j \leq \bar{p}$

Dann gilt $m^*(I) < 2 \cdot \bar{p} \leq 2 \cdot m(I, \mathbf{RUCK_APP}(I))$.

2. Fall: $p_j > \bar{p}$

Dann gilt $p_{\max} \geq p_j > \bar{p}$ und $m^*(I) < \bar{p} + p_j < 2 \cdot p_{\max} \leq 2 \cdot m(I, \mathbf{RUCK_APP}(I))$.

In beiden Fällen folgt die Abschätzung $1 \leq m^*(I)/m(I, \mathbf{RUCK_APP}(I)) \leq 2$. ///

Gibt man die Instanz $I = \left(\left(\frac{M}{2} + 1, \frac{M}{2} + 2 \right), \left(\frac{M}{2}, \frac{M}{2} \right), \left(\frac{M}{2}, \frac{M}{2} \right), M \right)$ mit $M > 4$ in den Algorithmus **RUCK_APP** ein, so liefert er das Ergebnis $x_1 = 1$, $x_2 = 0$ und $x_3 = 0$ und den Profit $m(I, \mathbf{RUCK_APP}(I)) = M/2 + 2$. Die optimale Lösung ist jedoch Ergebnis $x_1^* = 0$, $x_2^* = 1$ und $x_3^* = 1$ mit dem Profit $m^*(I) = M > M/2 + 2$. Damit ist

$$m^*(I) / m(I, \mathbf{RUCK_APP}(I)) = \frac{M}{M/2 + 2} = \frac{2M}{M + 4} = \frac{2 \cdot (M + 4) - 8}{M + 4} = 2 - \frac{8}{M + 4}.$$

Bei genügend großem M kommt dieser Wert beliebig dicht an 2 heran, so daß die Abschätzung in Satz 6.1-5 nicht verbessert werden kann.

Der folgende Satz zeigt, daß es unter der Voraussetzung $\mathbf{P} \neq \mathbf{NP}$ nicht für jedes Optimierungsproblem aus **NPO** einen approximativen Algorithmus gibt. Dazu sei noch einmal das Handlungsreisenden-Minimierungsproblem mit ungerichteten Graphen angegeben:

Das Handlungsreisenden-Minimierungsproblem

Instanz: 1. $G = (V, E, w)$

$G = (V, E, w)$ ist ein gewichteter ungerichteter Graph mit der Knotenmenge $V = \{v_1, \dots, v_n\}$, bestehend aus n Knoten, und der Kantenmenge $E \subseteq V \times V$; die Funktion $w: E \rightarrow \mathbf{R}_{\geq 0}$ gibt jeder Kante $e \in E$ ein nichtnegatives Gewicht

2. $\text{SOL}(G) = \{T \mid T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle \text{ ist eine Tour durch } G\}$

3. für $T \in \text{SOL}(G)$, $T = \langle (v_{i_1}, v_{i_2}), (v_{i_2}, v_{i_3}), \dots, (v_{i_{n-1}}, v_{i_n}), (v_{i_n}, v_{i_1}) \rangle$, ist die Zielfunktion definiert durch $m(G, T) = \sum_{j=1}^{n-1} w(v_{i_j}, v_{i_{j+1}}) + w(v_{i_n}, v_{i_1})$

4. $\text{goal} = \min$

Lösung: Eine Tour $T^* \in \text{SOL}(G)$, die minimale Kosten (unter allen möglichen Touren durch G) verursacht, und $m(G, T^*)$.

Satz 6.1-6:

Ist $\mathbf{P} \neq \mathbf{NP}$, so gibt es keinen r -approximativen Algorithmus für das Handlungsreisenden-Minimierungsproblem (mit $r \in \mathbf{R}_{\geq 1}$).

Ist $\mathbf{P} \neq \mathbf{NP}$, so ist $\mathbf{APX} \subset \mathbf{NPO}$.

Beweis:

Die Argumentation folgt der Idee aus dem Beweis von Satz 6-3.

Es wird angenommen, daß es bei Vorgabe des Wertes $r \in \mathbf{R}_{\geq 1}$ einen r -approximativen Algorithmus **A** für das Handlungsreisenden-Minimierungsproblem gibt und zeigt dann, daß dieser (mit einer einfachen Modifikation) dazu verwendet werden kann, das Problem des Hamiltonschen Kreises in einem ungerichteten Graphen (UNGERICHTETER HAMILTONKREIS, vgl. Kapitel 5.4) in polynomialer Zeit zu entscheiden. Da dieses **NP**-vollständig ist, folgt $\mathbf{P} = \mathbf{NP}$ (vgl. Satz 5.1-1).

Das Problem des Hamiltonschen Kreises in einem ungerichteten Graphen (UNGERICHTETER HAMILTONKREIS) lautet wie folgt:

Instanz: G ,
 $G = (V, E)$ ist ein ungerichteter Graph mit $V = \{v_1, \dots, v_n\}$.

Lösung: Entscheidung „ja“, falls gilt:

G besitzt einen Hamiltonschen Kreis. Dieses ist eine Anordnung $(v_{p(1)}, v_{p(2)}, \dots, v_{p(n)})$ der Knoten mittels einer Permutation π der Knotenindizes, so daß für $i = 1, \dots, n-1$ gilt: $(v_{p(i)}, v_{p(i+1)}) \in E$ und $(v_{p(n)}, v_{p(1)}) \in E$.

Es sei $\mathbf{A}$ ein r -approximativen Algorithmus $\mathbf{A}$ für das Handlungsreisenden-Minimierungsproblem, d.h. für alle Instanzen I dieses Problems gilt $m(I, \mathbf{A}(I)) \leq r \cdot m^*(I)$. Man kann $r > 1$ annehmen, denn für $r = 1$ ermittelt $\mathbf{A}$ bereits eine optimale Lösung. Es wird ein Algorithmus $\tilde{\mathbf{A}}$ für UNGERICHTETER HAMILTONKREIS konstruiert, der folgendermaßen arbeitet:

Bei Eingabe eines ungerichteten Graphen $G = (V, E)$ mit $V = \{v_1, \dots, v_n\}$ konstruiert $\tilde{\mathbf{A}}$ einen vollständigen bewerteten ungerichteten Graphen $\tilde{G} = (V, \tilde{E})$ (die Knotenmenge wird beibehalten) mit $\tilde{E} = \{(v_i, v_j) \mid v_i \in V \text{ und } v_j \in V\}$ und der Kantenbewertung $w: \tilde{E} \rightarrow \mathbf{R}_{\geq 0}$, die durch

$w(v_i, v_j) = \begin{cases} 1 & \text{falls } (v_i, v_j) \in E \\ nr & \text{sonst} \end{cases}$ definiert ist. Diese Konstruktion erfolgt in polynomieller

Zeit, da höchstens $O(n^2)$ viele Kanten hinzuzufügen und $O(n^2)$ viele Kanten zu bewerten sind. Alle Bewertungen haben eine Länge der Ordnung $O(\log(n))$. Der Graph $\tilde{G} = (V, \tilde{E})$ wird in den Algorithmus $\mathbf{A}$ eingegeben und das Ergebnis $m(\tilde{G}, \mathbf{A}(\tilde{G}))$ mit dem Wert $r \cdot n$ verglichen. $\tilde{\mathbf{A}}$ trifft die Entscheidung

$$\tilde{\mathbf{A}}(G) = \begin{cases} \text{ja} & \text{für } m(\tilde{G}, \mathbf{A}(\tilde{G})) \leq r \cdot n \\ \text{nein} & \text{für } m(\tilde{G}, \mathbf{A}(\tilde{G})) > r \cdot n \end{cases}.$$

Falls $\mathbf{A}$ ein polynomiell zeitbeschränkter Algorithmus ist, dann auch $\tilde{\mathbf{A}}$.

Besitzt G einen Hamiltonkreis (mit Länge n), dann ist $m^*(\tilde{G}) = n$, und $\tilde{\mathbf{A}}(G) = \text{ja}$, da $\mathbf{A}$ ein r -approximativen Algorithmus $\mathbf{A}$ für das Handlungsreisenden-Minimierungsproblem ist und damit $m(\tilde{G}, \mathbf{A}(\tilde{G})) \leq r \cdot m^*(\tilde{G}) = r \cdot n$ ist.

Besitzt G keinen Hamiltonkreis, dann enthält jeder Hamiltonkreis in $\tilde{G}$ mindestens eine Kante, die mit $r \cdot n$ bewertet ist (da $\tilde{G}$ ein vollständiger Graph ist, enthält $\tilde{G}$ einen Hamiltonkreis). Damit ergibt sich $m(\tilde{G}, \mathbf{A}(\tilde{G})) \geq m^*(\tilde{G}) \geq n - 1 + r \cdot n > r \cdot n$, d.h. $\tilde{\mathbf{A}}(G) = \text{nein}$.

In beiden Fällen trifft $\tilde{\mathbf{A}}$ die korrekte Entscheidung. ///

Das **metrische Handlungsreisenden-Minimierungsproblem** liegt jedoch in **APX**. Dieses stellt eine zusätzliche Bedingung an die Gewichtsfunktion einer Eingabeinstanz, nämlich die Gültigkeit der **Dreiecksungleichung**:

Für $v_i \in V$, $v_j \in V$ und $v_k \in V$ gilt $w(v_i, v_j) \leq w(v_i, v_k) + w(v_k, v_j)$.

Auch hierbei ist das zugehörige Entscheidungsproblem **NP**-vollständig. Es gibt (unabhängig von der Annahme $\mathbf{P} \neq \mathbf{NP}$) im Gegensatz zum (allgemeinen) Handlungsreisenden-Minimierungsproblem für dieses Problem einen 1,5-approximativen Algorithmus (siehe Literatur). Es ist nicht bekannt, ob es einen Approximationsalgorithmus mit einer kleineren relativen Approximationsgüte gibt oder ob aus der Existenz eines derartigen Algorithmus bereits $\mathbf{P} = \mathbf{NP}$ folgt.

Hat man für ein Optimierungsproblem aus **APX** einen r -approximativen Algorithmus gefunden, so stellt sich die Frage, ob dieser noch verbessert werden kann, d.h. ob es einen t -approximativen Algorithmus mit $1 \leq t < r$ gibt. Der folgende Satz (gap theorem) besagt, daß man unter Umständen an Grenzen stößt, daß es nämlich Optimierungsprobleme in **APX** gibt, die in polynomieller Zeit nicht beliebig dicht approximiert werden können, außer es gilt $\mathbf{P} = \mathbf{NP}$.

Satz 6.1-7:

Es sei Π' ein **NP**-vollständiges Entscheidungsproblem über Σ'^* und Π aus **NPO** ein Minimierungsproblem über Σ^* . Es gebe zwei in polynomieller Zeit berechenbare Funktionen $f: \Sigma'^* \rightarrow \Sigma^*$ und $c: \Sigma'^* \rightarrow \mathbf{N}$ und eine Konstante $e > 0$ mit folgender Eigenschaft:

$$\text{Für jedes } x \in \Sigma'^* \text{ ist } m^*(f(x)) = \begin{cases} c(x) & \text{für } x \in L_{\Pi'} \\ c(x) \cdot (1+e) & \text{für } x \notin L_{\Pi'} \end{cases}.$$

Dann gibt es keinen polynomiell zeitbeschränkten r -approximativen Algorithmus für Π mit $r < 1+e$, außer $\mathbf{P} = \mathbf{NP}$.

Beweis:

Angenommen, es gibt einen polynomiell zeitbeschränkten r -approximativen Algorithmus $\mathbf{A}$ für Π mit $r < 1+e$. Dann kann man daraus einen polynomiell zeitbeschränkten Algorithmus $\mathbf{A}'$ für Π' konstruieren, der genau $L_{\Pi'}$ erkennt. Das bedeutet $\mathbf{P} = \mathbf{NP}$.

Die Arbeitsweise von A' wird informell beschrieben:

Bei Eingabe von $x \in \Sigma'^*$ in A' werden in polynomieller Zeit die Werte $f(x)$ und $c(x)$ berechnet. Der Wert $f(x)$ wird in den polynomiell zeitbeschränkten r -approximativen Algorithmus A für Π eingegeben und der Näherungswert $m(f(x), A(f(x)))$ mit $c(x) \cdot (1+e)$ verglichen. A' trifft die Entscheidung

$$A'(x) = \begin{cases} \text{ja} & \text{für } m(f(x), A(f(x))) < c(x) \cdot (1+e) \\ \text{nein} & \text{für } m(f(x), A(f(x))) \geq c(x) \cdot (1+e) \end{cases}.$$

A' ist ein polynomiell zeitbeschränkter Entscheidungsalgorithmus. Zu zeigen ist, daß die von A' getroffene Entscheidung korrekt ist, d.h. daß $A'(x) = \text{ja}$ genau dann gilt, wenn $x \in L_{\Pi'}$ ist.

Es sei $x \in L_{\Pi'}$. Da A r -approximativ ist, gilt $\frac{m(f(x), A(f(x)))}{m^*(f(x))} \leq r < 1+e$. Wegen $x \in L_{\Pi'}$ ist $m^*(f(x)) = c(x)$. Also ist $m(f(x), A(f(x))) < m^*(f(x)) \cdot (1+e) = c(x) \cdot (1+e)$, und $A'(x) = \text{ja}$.

Es sei $x \notin L_{\Pi'}$. Dann folgt $m^*(f(x)) = c(x) \cdot (1+e)$,
 $m(f(x), A(f(x))) \geq m^*(f(x)) = c(x) \cdot (1+e)$ und $A'(x) = \text{nein}$.

In beiden Fällen trifft A' die korrekte Entscheidung. ///

Bemerkung: Der Beweis von Satz 6.1-7 zeigt, daß Satz 6.1-7 gültig bleibt, wenn die dort formulierte Voraussetzung über $m^*(f(x))$ ersetzt durch die Voraussetzung

$$\text{Für jedes } x \in \Sigma'^* \text{ ist } m^*(f(x)) \begin{cases} = c(x) & \text{für } x \in L_{\Pi'} \\ \geq c(x) \cdot (1+e) & \text{für } x \notin L_{\Pi'} \end{cases}.$$

Ein entsprechender Satz gilt für Maximierungsprobleme:

Satz 6.1-8:

Es sei Π' ein **NP**-vollständiges Entscheidungsproblem über Σ'^* und Π aus **NPO** ein Maximierungsproblem über Σ^* . Es gebe zwei in polynomieller Zeit berechenbare Funktionen $f: \Sigma'^* \rightarrow \Sigma^*$ und $c: \Sigma'^* \rightarrow \mathbb{N}$ und eine Konstante $e > 0$ mit folgender Eigenschaft:

$$\text{Für jedes } x \in \Sigma'^* \text{ ist } m^*(f(x)) = \begin{cases} c(x) & \text{für } x \in L_{\Pi'} \\ c(x) \cdot (1-e) & \text{für } x \notin L_{\Pi'} \end{cases}.$$

Dann gibt es keinen polynomiell zeitbeschränkten r -approximativen Algorithmus für Π mit $r < 1/(1-e)$, außer $\mathbf{P} = \mathbf{NP}$.

Mit Hilfe dieser Sätze läßt sich zeigen beispielsweise, daß die Grenze $r = 1,5$ für einen r -approximativen (polynomiell zeitbeschränkten) Algorithmus unter der Annahme $\mathbf{P} \neq \mathbf{NP}$ für das Binpacking-Minimierungsproblem optimal ist:

Satz 6.1-9:

Ist $\mathbf{P} \neq \mathbf{NP}$, dann gibt es keinen r -approximativen Algorithmus für das Binpacking-Minimierungsproblem mit $r \leq 3/2 - e$ für $e > 0$.

Beweis:

In Satz 6.1-7 übernimmt das Binpacking-Minimierungsproblem die Rolle von Π ; für Π' wird das **NP**-vollständige Partitionenproblem (siehe Kapitel 5.4) genommen. Es wird definiert durch

Instanz: $I = [a_1, \dots, a_n]$,
 $a_1, \dots, a_n$ sind natürliche Zahlen.

Lösung: Entscheidung „ja“ genau dann, wenn gilt:

Es gibt eine Aufteilung der Menge $\{a_1, \dots, a_n\}$ in zwei disjunkte Teile I_1 und I_2 mit $\sum_{a_i \in I_1} a_i = \sum_{a_i \in I_2} a_i$.

Es werden zwei Abbildungen f und c definiert, die den Bedingungen in Satz 6.1-7 genügen. Die Abbildung f ordnet jeder Instanz $I = [a_1, \dots, a_n]$ des Partitionenproblems eine Instanz $f(I)$ des Binpacking-Minimierungsproblems zu.

Die Funktion c mit $c(I) = 2$ für alle Instanzen I des Partitionenproblems ist trivialerweise in polynomieller Zeit berechenbar.

Für eine Instanz $I = [a_1, \dots, a_n]$ des Partitionenproblems sei $B = \sum_{a_i \in I} a_i$. Es wird $f(I) = [(2a_1)/B, \dots, (2a_n)/B]$ definiert, falls sich dadurch eine Instanz des Binpacking-Minimierungsproblems ergibt, d.h. falls $(2a_i)/B \leq 1$ für $i = 1, \dots, n$ gilt. Ansonsten wird $f(I) = [1, 1, 1]$ gesetzt. In jedem Fall ist $f(I)$ eine Instanz des Binpacking-Minimierungsproblems. Da dieses in **NPO** liegt und wegen Satz 2.1-3 ist $f(I)$ in polynomieller Zeit berechenbar.

Ist $I \in L_{\Pi'}$, d.h. I ist eine „ja“-Instanz des Partitionenproblems bzw. es gibt eine Aufteilung der Menge $\{a_1, \dots, a_n\}$ in zwei disjunkte Teile I_1 und I_2 mit $\sum_{a_i \in I_1} a_i = \sum_{a_i \in I_2} a_i$, dann gilt $a_i \leq B/2$ für $i = 1, \dots, n$, denn sonst könnte man I nicht in zwei gleichgroße Teile aufteilen. Daher ist $(2a_i)/B \leq 1$, $\sum_{a_i \in I_1} a_i = \sum_{a_i \in I_2} a_i = B/2$ und $\sum_{a_i \in I_1} (2a_i)/B = \sum_{a_i \in I_2} (2a_i)/B = 1$. Zur Packung der Instanz $f(I) = [(2a_1)/B, \dots, (2a_n)/B]$ des Binpacking-Minimierungsproblems sind genau 2 Behälter erforderlich, d.h. $m^*(f(I)) = 2 = c(I)$.

Ist $I \notin L_{\Pi'}$ und $[(2a_1)/B, \dots, (2a_n)/B]$ keine Instanz des Binpacking-Minimierungsproblems, d.h. es gibt mindestens ein a_i mit $(2a_i)/B > 1$, dann ist $f(I) = [1, 1, 1]$ und $m^*(f(I)) = 3$.

Ist $I \notin L_{\Pi'}$ und $[(2a_1)/B, \dots, (2a_n)/B]$ eine Instanz des Binpacking-Minimierungsproblems, dann ist $f(I) = [(2a_1)/B, \dots, (2a_n)/B]$ und $m^*(f(I)) \geq 3$.

In beiden Fällen gilt $m^*(f(I)) \geq 3 = 2 \cdot \frac{3}{2} = c(I) \cdot \left(1 + \frac{1}{2}\right)$.

Aus Satz 6.1-7 zusammen mit der sich dort anschließenden Bemerkung folgt, daß es keinen polynomiell zeitbeschränkten r -approximativen Algorithmus für das Binpacking-Minimierungsproblem mit $r < 1 + 1/2 = 1.5$ gibt, außer **P** = **NP**. ///

6.2 Polynomiell zeitbeschränkte und asymptotische Approximationsschemata

In vielen praktischen Anwendungen möchte man die relative Approximationsgüte verbessern. Dabei ist man sogar bereit, längere Laufzeiten der Approximationsalgorithmen in Kauf zu nehmen, solange sie noch polynomielles Laufzeitverhalten bezüglich der Größe der Eingabeinstanzen haben. Bezüglich der relativen Approximationsgüte r akzeptiert man eventuell

sogar ein Laufzeitverhalten, das exponentiell von $1/(r-1)$ abhängt: Je besser die Approximation ist, um so größer ist die Laufzeit. In vielen Fällen kann man so eine optimale Lösung beliebig dicht approximieren, allerdings zum Preis eines dramatischen Anstiegs der Rechenzeit.

Es sei Π ein Optimierungsproblem aus **NPO**. Ein Algorithmus **A** heißt **polynomiell zeitbeschränktes Approximationsschema** (polynomial-time approximation scheme) für Π , wenn er für jede Eingabeinstanz x von Π und für jede rationale Zahl $r > 1$ bei Eingabe von (x, r) eine r -approximative Lösung für x in einer Laufzeit liefert, die polynomiell von $size(x)$ abhängt.

Mit **PTAS** werde die Klasse der Optimierungsprobleme in **NPO** bezeichnet, für die es ein polynomiell zeitbeschränktes Approximationsschema gibt. Daß diese Klasse nicht leer ist, zeigt das folgende Beispiel.

Partitionen-Minimierungsproblem

- Instanz: 1. $I = [a_1, \dots, a_n]$
 $a_1, \dots, a_n$ („Objekte“) sind natürliche Zahlen mit $a_i > 0$ für $i = 1, \dots, n$
 $size(I) = n$
 2. $SOL(I) = \{[Y_1, Y_2] \mid [Y_1, Y_2] \text{ ist eine Partition (disjunkte Zerlegung) von } I\}$
 3. Für $[Y_1, Y_2] \in SOL(I)$ ist $m(I, [Y_1, Y_2]) = \max\left\{\sum_{a_i \in Y_1} a_i, \sum_{a_j \in Y_2} a_j\right\}$
 4. $goal = \min$

Lösung: Eine Partition der Objekte in zwei Teile $[Y_1, Y_2]$, so daß sich die Summen der Objekte in beiden Teilen möglichst wenig unterscheiden.

Der folgende in Pseudocode formulierte Algorithmus ist ein polynomiell zeitbeschränktes Approximationsschema für das Partitionen-Minimierungsproblem.

Polynomiell zeitbeschränktes Approximationsschema für das Partitionen-Minimierungsproblem

Eingabe: $I = [a_1, \dots, a_n]$, rationale Zahl $r > 1$,
 $a_1, \dots, a_n$ („Objekte“) sind natürliche Zahlen mit $a_i > 0$ für $i = 1, \dots, n$

Verfahren: VAR Y_1 : SET OF INTEGER;
 Y_2 : SET OF INTEGER;
 k : REAL;

```

j : INTEGER;

BEGIN
  IF r >= 2
  THEN BEGIN
    Y1 := {a1, ..., an};
    Y2 := {};
  END
  ELSE BEGIN
    Sortiere die Objekte nach absteigender Größe; die dabei
    entstehende Folge sei (x1, ..., xn);
    k := ⌈(2-r)/(r-1)⌉;
    { Phase 1: }
    finde eine optimale Partition [Y1, Y2] für [x1, ..., xk];
    { Phase 2: }
    FOR j := k+1 TO n DO
      IF  $\sum_{x_i \in Y_1} x_i \leq \sum_{x_i \in Y_2} x_i$ 
      THEN Y1 := Y1 ∪ {xj}
      ELSE Y2 := Y2 ∪ {xj}
    END;
  END;

```

Ausgabe: [Y₁, Y₂].

Satz 6.2-1:

Das Partitionen-Minimierungsproblem liegt in **PTAS**.

Beweis:

Es ist zu zeigen, daß der angegebene Algorithmus, der mit **A** bezeichnet werde, bei Eingabe einer Instanz $I = [a_1, \dots, a_n]$ für das Partitionen-Minimierungsproblem und einer rationalen Zahl $r > 1$ eine r -approximative Lösung für I in einer Laufzeit liefert, die polynomiell von n abhängt.

Es werden $I = [a_1, \dots, a_n]$ und r in den angegebenen Algorithmus eingegeben. Die Ausgabe sei $[Y_1, Y_2]$. O.B.d.A. sei $\sum_{a_i \in Y_1} a_i \geq \sum_{a_j \in Y_2} a_j$. Dann ist $m(I, \mathbf{A}(I)) = \sum_{a_i \in Y_1} a_i$. Es sei $w(I)$ die

Summe aller Objekte der Eingabeinstanz I : $w(I) = \sum_{a_i \in I} a_i$, $w(Y_2)$ die Summe aller Objekte in Y_2 : $w(Y_2) = \sum_{a_i \in Y_2} a_i$.

1. Fall: $r \geq 2$

Dann ist $m^*(I) \geq w(I)/2 \geq m(I, \mathbf{A}(I))/2$, also $m(I, \mathbf{A}(I)) \leq 2 \cdot m^*(I) \leq r \cdot m^*(I)$.

2. Fall: $1 \leq r < 2$

Es sei a_h das letzte zu Y_1 hinzugefügte Objekt.

Wurde a_h in Phase 1 des Algorithmus in Y_1 aufgenommen, so läßt sich $m(I, \mathbf{A}(I)) = m^*(I)$ zeigen. Es wird daher angenommen, daß a_h in Phase 2 des Algorithmus in Y_1 aufgenommen wurde. Es folgt nacheinander:

$$m(I, \mathbf{A}(I)) - a_h \leq w(Y_2),$$

$$2 \cdot m(I, \mathbf{A}(I)) - a_h \leq w(Y_2) + m(I, \mathbf{A}(I)) = w(I) \text{ und}$$

$$m(I, \mathbf{A}(I)) - w(I)/2 \leq a_h/2.$$

Außerdem gilt wegen der Sortierung der Objekte $a_h \leq a_j$ für $j = 1, \dots, k$. Diese Ungleichungen werden auf a_h aufsummiert zu dem Ergebnis:

$$(k+1) \cdot a_h = a_h + k \cdot a_h \leq \sum_{j=1}^k a_j + a_h \leq w(I), \text{ also } a_h \leq w(I)/(k+1).$$

Es gilt $m(I, \mathbf{A}(I)) \geq w(I)/2 \geq w(Y_2)/2$ und $m^*(I) \geq w(I)/2$.

Insgesamt ergibt sich:

$$\frac{m(I, \mathbf{A}(I))}{m^*(I)} \leq \frac{m(I, \mathbf{A}(I))}{w(I)/2} \leq \frac{a_h/2 + w(I)/2}{w(I)/2} = 1 + \frac{a_h}{w(I)} \leq 1 + \frac{1}{k+1} \leq r;$$

die letzte Ungleichung folgt aus der Wahl von k , nämlich $k = \lceil (2-r)/(r-1) \rceil$:

$$k \geq (2-r)/(r-1), \text{ d.h. } k+1 \geq (2-r+r-1)/(r-1) = 1/(r-1) > 1/r.$$

Mit $k(r) = \lceil (2-r)/(r-1) \rceil$ ist das Laufzeitverhalten von der Ordnung $O(n \cdot \log(n) + n^{k(r)})$. Der erste Term gibt die Laufzeit zum Sortieren der Objekte der Eingabeinstanz an, der zweite Term die Laufzeit zur Ermittlung einer optimalen Lösung für die ersten k Elemente in Phase 1. Die Laufzeit für Phase 2 ist von der Ordnung $O(n)$. Da $k(r) \in O(1/(r-1))$ ist das Laufzeitverhalten bei festem r polynomiell in der Größe n der Eingabeinstanz, jedoch exponentiell in n und $1/(r-1)$. ///

Offensichtlich ist $\mathbf{PTAS} \subseteq \mathbf{APX}$. In Kapitel 6.1 wurde erwähnt, daß es unter der Annahme $\mathbf{P} \neq \mathbf{NP}$ keinen r -approximativen Algorithmus für das Binpacking-Minimierungsproblem mit $r \leq 3/2 - e$ für $e > 0$ gibt. Daraus folgt:

Satz 6.2-2:

Ist $\mathbf{P} \neq \mathbf{NP}$, so gilt $\mathbf{PTAS} \subset \mathbf{APX} \subset \mathbf{NPO}$.

Es gibt in $\mathbf{PTAS}$ Optimierungsprobleme, die ein polynomiell zeitbeschränktes Approximationsschema $\mathbf{A}$ zulassen, das bei Eingabe einer Instanz x von Π und einer rationalen Zahl $r > 1$ eine r -approximative Lösung für x in einer Laufzeit liefert, die polynomiell von $|x|$ und $1/(r-1)$ abhängt.

Einen derartigen Algorithmus nennt man **voll polynomiell zeitbeschränktes Approximationsschema** (fully polynomial-time approximation scheme). Die Optimierungsprobleme, die einen derartigen Algorithmus zulassen, bilden die Klasse $\mathbf{FPTAS}$.

In der Literatur werden Beispiele für Optimierungsprobleme genannt, die in $\mathbf{FPTAS}$ liegen.

Es läßt sich zeigen:

Satz 6.2-3:

Ist $\mathbf{P} \neq \mathbf{NP}$, so gilt $\mathbf{FPTAS} \subset \mathbf{PTAS} \subset \mathbf{APX} \subset \mathbf{NPO}$.

In Kapitel 6.1 wurden mehrere Approximationsalgorithmen für das Binpacking-Minimierungsproblem angegeben. Die Approximationsgüte $m(I, \mathbf{BFD}(I)) \leq \frac{11}{9} \cdot m^*(I) + 4$ von Best-fitDecreasing zeigt, daß man (unabhängig von der Annahme $\mathbf{P} \neq \mathbf{NP}$) durchaus einen Approximationsalgorithmus entwerfen kann, dessen relative Approximationsgüte unterhalb der überhaupt für einen r -approximativen Algorithmus möglichen Untergrenze liegt. Eventuell gibt es sogar in einem erweiterten Sinne ein polynomiell zeitbeschränktes Approximationsschema für das Binpacking-Minimierungsproblem und andere Optimierungsprobleme. Diese Überlegung führt auf folgende Definition:

Es sei Π ein Optimierungsproblem aus $\mathbf{NPO}$. Ein Algorithmus $\mathbf{A}$ heißt **asymptotisches Approximationsschema** für Π , wenn es eine Konstante k gibt, so daß gilt:

für jede Eingabeinstanz x von Π und für jede rationale Zahl $r > 1$ liefert $\mathbf{A}$ bei Eingabe von (x, r) eine (zulässige) Lösung, deren relative Approximationsgüte $R_{\mathbf{A}}(x)$ die Bedingung $R_{\mathbf{A}}(x) \leq r + k/m_{\Pi}^*(x)$ erfüllt. Außerdem ist die Laufzeit von $\mathbf{A}$ für jedes feste r polynomiell in der Größe $\text{size}(x)$ der Eingabeinstanz.

Zur Erinnerung: die relative Approximationsgüte wurde definiert durch

$$R_{\mathbf{A}}(x) = \frac{m_{\Pi}^*(x)}{m(x, \mathbf{A}(x))} \text{ bei einem Maximierungsproblem}$$

bzw.

$R_A(x) = \frac{m(x, \mathbf{A}(x))}{m_{\Pi}^*(x)}$ bei einem Minimierungsproblem.

Die Bedingung $R_A(x) \leq r + k/m_{\Pi}^*(x)$ besagt also

bei einem Maximierungsproblem:

$$m(x, \mathbf{A}(x)) \geq \frac{1}{r} \cdot m_{\Pi}^*(x) - k' \quad \text{mit} \quad k' = k / (r + k/m_{\Pi}^*(x)) \leq k/r,$$

$$\text{daher } m(x, \mathbf{A}(x)) \geq \frac{1}{r} \cdot m_{\Pi}^*(x) - k/r$$

bzw.

bei einem Minimierungsproblem: $m(x, \mathbf{A}(x)) \leq r \cdot m_{\Pi}^*(x) + k$.

Die Bezeichnung „asymptotisches Approximationsschema“ ist aus der Tatsache zu erklären, daß für „große“ Eingabeinstanzen x der Wert $m_{\Pi}^*(x)$ der Zielfunktion einer optimalen Lösung ebenfalls groß ist. Daher gilt in diesem Fall $\lim_{\text{size}(x) \rightarrow \infty} R_A(x) \leq r$.

Die Klasse aller Optimierungsprobleme, die ein asymptotisches Approximationsschema zulassen, wird mit **PTAS**[∞] bezeichnet.

Satz 6.2-4:

Das Binpacking-Minimierungsproblem liegt in **PTAS**[∞], d.h. es kann asymptotisch beliebig dicht in polynomieller Zeit approximiert werden (auch wenn **P** ≠ **NP** gilt).

Es gilt sogar: Es gibt ein asymptotisches Approximationsschema, das polynomiell in der Problemgröße und in $1/(r-1)$ ist.

Die Klasse **PTAS**[∞] ordnet sich in die übrigen Approximationsklassen ein:

Satz 6.2-5:

Ist **P** ≠ **NP**, so gilt **FPTAS** ⊂ **PTAS** ⊂ **PTAS**[∞] ⊂ **APX** ⊂ **NPO**.

7 Weiterführende Konzepte

Die Analyse des Laufzeitverhaltens eines Algorithmus **A** im schlechtesten Fall $T_A(n) = \max\{t_A(x) \mid x \in \Sigma^* \text{ und } |x| \leq n\}$ liefert eine *Garantie* (obere Schranke) für die Zeit, die er bei einer Eingabe zur Lösung benötigt⁴. Für jede Eingabe $x \in \Sigma^*$ mit $|x| = n$ gilt $t_A(x) \leq T_A(n)$. Dieses Verhalten ist oft jedoch nicht charakteristisch für das Verhalten bei „den meisten“ Eingaben. Ist eine Wahrscheinlichkeitsverteilung P der Eingabewerte bekannt oder werden alle Eingaben als gleichwahrscheinlich aufgefaßt, so kann man das Laufzeitverhalten von **A** untersuchen, wie es sich im Mittel darstellt. Die **Zeitkomplexität** von **A** im **Mittel** wird definiert als

$$T_A^{avg}(n) = \mathbf{E}[t_A(x) \mid |x| = n] = \int_{|x|=n} t_A(x) \cdot dP(x).$$

Die Untersuchung des mittleren Laufzeitverhaltens von Algorithmen ist in den meisten Fällen sehr viel schwieriger als die worst-case-Analyse. Trotzdem gibt es sehr ermutigende Ergebnisse, insbesondere bei der Lösung einiger „klassischer“ Probleme. Der Einbau von Zufallsexperimenten in Algorithmen hat auf dem Gebiet der Approximationsalgorithmen für Optimierungsprobleme aus **NPO** gute Ergebnisse geliefert. In neuerer Zeit hat der Einsatz von probabilistischen Modellen zu neuen Erkenntnissen auf dem Weg zur Lösung der **P-NP**-Frage geführt.

7.1 Randomisierte Algorithmen

Im ansonsten deterministischen Algorithmus werden als zulässige Elementaroperationen Zufallsexperimente zugelassen. Ein derartiges Zufallsexperiment kann beispielsweise mit Hilfe eines Zufallszahlengenerators ausgeführt werden. So wird beispielsweise auf diese Weise entschieden, in welchem Teil einer Programmverzweigung der Algorithmus während seines Ablaufs fortgesetzt wird. Andere Möglichkeiten zum Einsatz eines Zufallsexperiments bestehen bei Entscheidungen zur Auswahl möglicher Elemente, die im weiteren Ablauf des Algorithmus als nächstes untersucht werden sollen.

Man nennt derartige Algorithmen **randomisierte Algorithmen**.

Grundsätzlich gibt es zwei Klassen randomisierter Algorithmen: **Las-Vegas-Verfahren**, die stets – wie von deterministischen Algorithmen gewohnt – ein korrektes Ergebnis berechnen.

⁴ $t_A(x)$ = Anzahl der Anweisungen, die von **A** zur Berechnung von **A**(x) durchlaufen werden.

Daneben gibt es **Monte-Carlo-Verfahren**, die ein korrektes Ergebnis nur mit einer gewissen Fehlerwahrscheinlichkeit, aber in jedem Fall effizient, bestimmen. Häufig findet man bei randomisierten Algorithmen einen Trade-off zwischen Korrektheit und Effizienz.

Beispiele für randomisierte Algorithmen vom Las-Vegas-Typ sind:

- Einfügen von Primärschlüsselwerten in einen binären Suchbaum: Statt die Schlüssel $S_1, \dots, S_n$ in dieser Reihenfolge sequentiell in den binären Suchbaum einzufügen, wird jeweils der nächste einzufügende Schlüssel aus den restlichen, d.h. noch nicht eingefügten Schlüsseln zufällig ausgewählt und in den binären Suchbaum eingefügt. Das Ergebnis ist eine mittlere Baumhöhe der Ordnung $O(\log(n))$
- Quicksort zum Sortieren Elementen, auf denen eine lineare Ordnung definiert ist.

Sortieren von Elementen, auf denen eine Ordnungsrelation besteht

Instanz: $x = \langle a_1, \dots, a_n \rangle$

x ist eine Folge von Elementen der Form $a_i = (key_i, info_i)$; die Elemente sind bezüglich ihrer *key*-Komponente vergleichbar, d.h. es gilt für $a_i = (key_i, info_i)$ und $a_j = (key_j, info_j)$ die Beziehung $a_i \leq a_j$ genau dann, wenn $key_i \leq key_j$ ist.

Lösung: Eine Umordnung $p : [1:n] \rightarrow [1:n]$ der Elemente in x , so daß $\langle a_{p(1)}, \dots, a_{p(n)} \rangle$ nach aufsteigenden Werten der *key*-Komponente sortiert ist, d.h. es gilt: $a_{p(i)} \leq a_{p(i+1)}$ für $i = 1, \dots, n-1$.

Als Randbedingung gilt, daß die Sortierung innerhalb der Folge x erfolgen soll, also insbesondere nicht zusätzlicher Speicherplatz der Größenordnung $O(n)$ erforderlich wird (**internes Sortierv Verfahren**).

Zur algorithmischen Behandlung wird eine Instanz x der Größe n in einem Feld a gespeichert, das definiert wird durch

```
CONST n          = ...;          { Problemgröße }

TYPE idx_bereich = 1..n;
   key_typ       = ...;          { Typ der key-Komponente           }
   info_typ      = ...;          { Typ der info-Komponente          }
   entry_typ     = RECORD
       key   : key_typ;
       info  : info_typ;
   END;
```

```
feld_typ    = ARRAY [idx_bereich] OF entry_typ;
```

Die folgende Prozedur interchange vertauscht zwei Einträge miteinander:

```
PROCEDURE interchange (VAR x, y : entry_typ);
{ die Werte von x und y werden miteinander vertauscht }

VAR z : entry_typ;

BEGIN { interchange }
  z := x;
  x := y;
  y := z;
END   { interchange };
```

Der Algorithmus quicksort beruht auf der Idee der Divide-and-Conquer-Methode und erzeugt die umsortierte Eingabe $\langle a_{p(1)}, \dots, a_{p(n)} \rangle$:

Besteht x aus keinem oder nur aus einem einzigen Element, dann ist nichts zu tun. Besteht x aus mehr Elementen, dann wird aus x ein Element e (**Pivot-Element**) ausgewählt und das Problem der Sortierung von x in zwei kleinere Probleme aufgeteilt, nämlich der Sortierung aller Elemente, die kleiner als e sind (das sei das Problem der Sortierung der Folge $x(\text{lower})$), und die Sortierung der Elemente die größer oder gleich e sind (das sei das Problem der Sortierung der Folge $x(\text{upper})$); dabei wird als „Raum zur Sortierung“ die Instanz x verwendet. Zuvor wird e an diejenige Position innerhalb von x geschoben, an der es in der endgültigen Sortierung stehen wird. Das Problem der Sortierung der erzeugten kleineren Probleme $x(\text{lower})$ und $x(\text{upper})$ wird (rekursiv) nach dem gleichen Prinzip gelöst. Die Position, an der e innerhalb von x in der endgültigen Sortierung stehen wird und damit die kleineren Probleme $x(\text{lower})$ und $x(\text{upper})$ bestimmt, sind dadurch gekennzeichnet, daß gilt:
 $k < e$ für alle $k \in x(\text{lower})$ und $k \geq e$ für alle $k \in x(\text{upper})$.

Sortieralgorithmus mit quicksort:

Eingabe: $x = \langle a_1, \dots, a_n \rangle$ ist eine Folge von Elementen der Form $a_i = (key_i, info_i)$; die Elemente sind bezüglich ihrer key -Komponente vergleichbar, d.h. es gilt für $a_i = (key_i, info_i)$ und $a_j = (key_j, info_j)$ die Beziehung $a_i \leq a_j$ genau dann, wenn $key_i \leq key_j$ ist

```
VAR a : feld_typ;          { Eingabefeld }
    i : idx_bereich;
```

```

FOR i := 1 TO n DO
  BEGIN
    a[i].key := keyi;
    a[i].info := infoi
  END;

```

Verfahren: Aufruf der Prozedur quicksort (a, 1, n);

Ausgabe: a mit $a[i].key \leq a[i + 1].key$ für $i = 1, \dots, n-1$

```

PROCEDURE quicksort (VAR a      : feld_typ;
                     lower : idx_bereich;
                     upper : idx_bereich);

  VAR t      : entry_typ;
      m      : idx_bereich;
      idx     : idx_bereich;
      zufalls_idx : idx_bereich;

  BEGIN { quicksort }
    IF lower < upper
    THEN BEGIN { ein Feldelement zufällig aussuchen und
               mit a[lower] austauschen }
      zufalls_idx := lower
                  + Trunc(Random * (upper - lower + 1));
      interchange (a[lower], a[zufalls_idx]);

      t := a[lower];
      m := lower;

      FOR idx := lower + 1 TO upper DO
        { es gilt: a[lower+1] , ..., a[m] < t und
                  a[m+1], ... a[idx-1] >= t }
        IF a[idx].key < t.key
        THEN BEGIN
          m := m+1;
          { vertauschen von a[m] und a[idx] }
          interchange (a[m], a[idx])
        END;

        { a[lower] und a[m] vertauschen }
        interchange (a[lower], a[m]);

        { es gilt: a[lower], ..., a[m-1] < a[m] und
                  a[m] <= a[m+1], ..., a[upper] }

```

```
        quicksort (a, lower, m-1);
        quicksort (a, m+1, upper);

    END { IF}

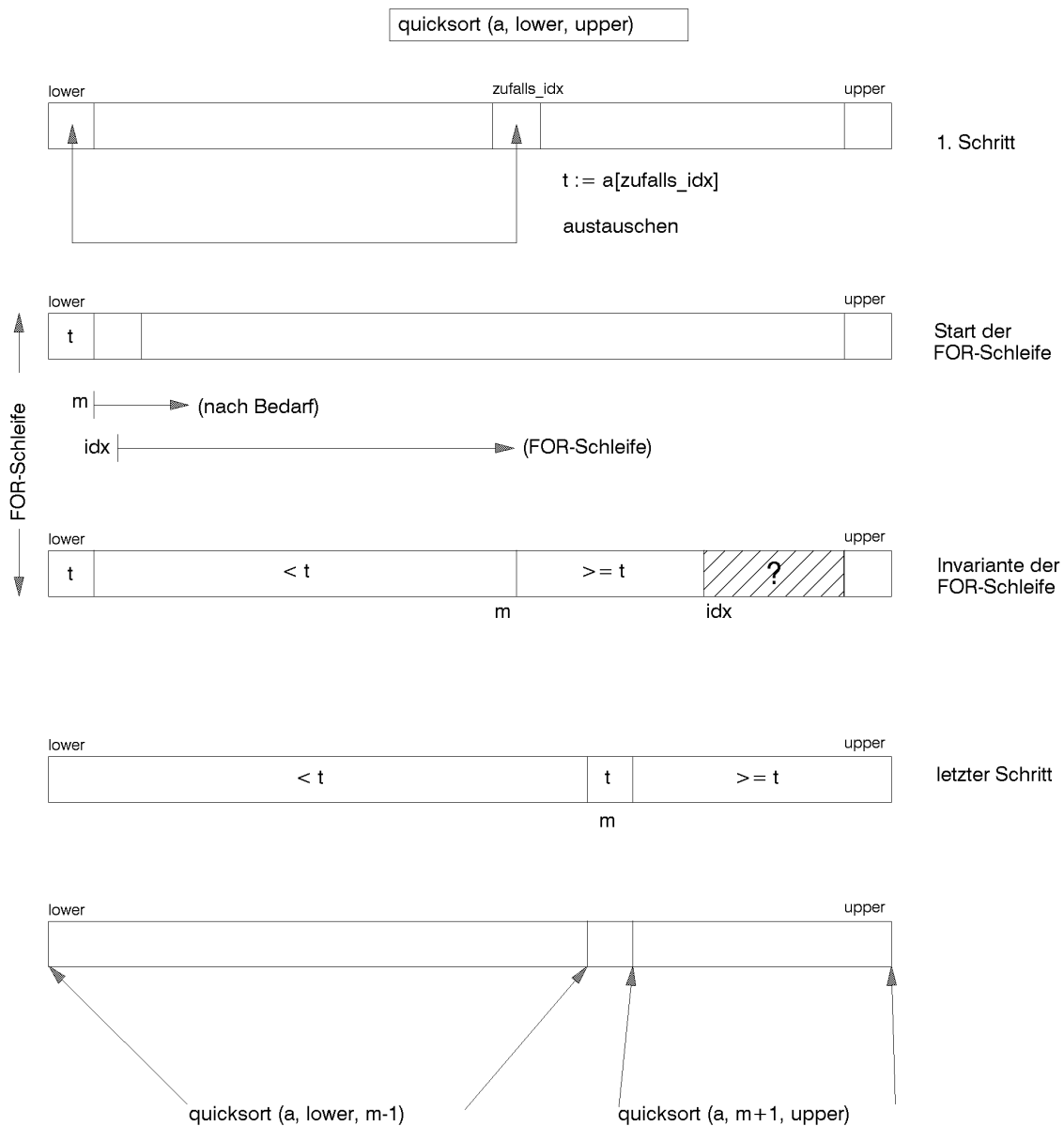
END { quicksort };
```

Der Algorithmus stoppt, da in jedem quicksort-Aufruf ein Element an die endgültige Position geschoben wird und die Anzahl der Elemente in den quicksort-Aufrufen in den beiden restlichen Teilfolgen zusammen um 1 verringert ist. Die Korrektheit des Algorithmus folgt aus der Invariante der FOR-Schleife.

Satz 7.1-1:

Das beschriebene Verfahren mit der Prozedur quicksort sortiert eine Folge $x = \langle a_1, \dots, a_n \rangle$ aus n Elementen mit einer (worst-case-) Zeitkomplexität der Ordnung $O(n^2)$. Diese Schranke wird erreicht, wenn zufällig in jedem quicksort-Aufruf das Pivot-Element so gewählt wird, daß die Teilfolgen $x(\text{lower})$ kein Element und $x(\text{upper})$ alle restlichen Elemente enthalten (bzw. umgekehrt). Im Mittel ist das Laufzeitverhalten jedoch wesentlich besser, nämlich von der Ordnung $O(n \cdot \log(n))$.

Es läßt sich zeigen, daß jedes Sortierverfahren, das auf dem Vergleich von Elementgrößen beruht, eine untere Zeitkomplexität mindestens der Ordnung $O(n \cdot \log(n))$ und eine mittlere Zeitkomplexität ebenfalls mindestens der Ordnung $O(n \cdot \log(n))$ hat, so daß das Verfahren mit der Prozedur quicksort optimales Laufzeitverhalten im Mittel aufweist. Daß das Verfahren mit der Prozedur quicksort in der Praxis ein Laufzeitverhalten zeigt, das fast alle anderen Sortierverfahren schlägt, liegt daran, daß die „ungünstigen“ Eingaben für quicksort mit gegen 0 gehender Wahrscheinlichkeit vorkommen und die Auswahl des Pivot-Elements zufällig erfolgt.



Beispiele für randomisierte Algorithmen vom Monte-Carlo-Typ sind die heute üblichen Primzahltests, die hier informell beschrieben werden sollen (Details findet man in der angegebenen Literatur).

Um eine natürliche Zahl $n > 2$ auf Primzahleigenschaft zu testen, könnte man alle Primzahlen von 2 bis $\lfloor \sqrt{n} \rfloor$ daraufhin untersuchen, ob es eine von ihnen gibt, die n teilt. Wenn die Zahl n

nämlich zusammengesetzt ist, d.h. keine Primzahl ist, hat sie einen Primteiler p mit $p \leq \sqrt{n}$. Umgekehrt, falls alle Primzahlen p mit $p \leq \sqrt{n}$ keine Teiler von n sind, dann ist n selbst eine Primzahl. Die Primzahlen könnte man etwa systematisch erzeugen (siehe Literatur unter dem Stichwort „Sieb des Eratosthenes“) oder man könnte sie einer Primzahltafel (falls vorhanden) entnehmen. Allerdings ist dieser Ansatz für sehr große Werte von n nicht praktikabel und benötigt exponentiellen Rechenaufwand (in der Anzahl der Stellen von n): Die Anzahl der Stellen $b(n)$ einer Zahl $n \geq 1$ im Zahlensystem zur Basis B ist $b(n) = \lfloor \log_B(n) \rfloor + 1 = \lceil \log_B(n+1) \rceil$, d.h. $b(n) \in O(\log(n))$. Die Anzahl der Primzahlen unterhalb $\lfloor \sqrt{n} \rfloor$ beträgt für große n nach dem Primzahlsatz der Zahlentheorie $\pi(\sqrt{n}) \sim \frac{n^{1/2}}{\ln(n^{1/2})}$, also ein Wert der Ordnung $O\left(\frac{2^{b(n)/2}}{b(n)}\right)$. Jede in Frage kommende Primzahl muß daraufhin untersucht werden, ob sie n teilt. Dazu sind mindestens $O((b(n))^2)$ viele Bitoperationen erforderlich. Daher ist der Gesamtaufwand mindestens von der Ordnung $O(b(n) \cdot 2^{b(n)/2})$.

Durch Anwendung zahlentheoretischer Erkenntnisse hat man versucht, effiziente Primzahltests zu entwickeln. Der bisher bekannte schnellste Algorithmus zur Überprüfung einer Zahl n auf Primzahleigenschaft, der APRCL-Test, hat eine Laufzeit der Ordnung $O((\log(n))^{c(\log(\log(\log(n))))})$ mit einer Konstanten $c > 0$. Auch das ist keine polynomielle Laufzeit.

Es sind daher andere Ansätze für effiziente Primzahltests erforderlich. Als erfolgreich haben sich hierbei **probabilistische Algorithmen** erwiesen.

Um zu testen, ob eine Zahl n eine Primzahl ist oder nicht, versucht man, einen **Zeugen (witness) für die Primzahleigenschaft von n** zu finden. Ein Zeuge ist dabei eine Zahl a mit $1 \leq a \leq n-1$, der eine bestimmte Eigenschaft zukommt, aus der man *vermuten* kann, daß n Primzahl ist. Dabei muß diese Eigenschaft bei Vorgabe von a einfach zu überprüfen sein, und nach Auffinden einiger weniger Zeugen für die Primzahleigenschaft von n muß der Schluß gültig sein, daß n mit hoher Wahrscheinlichkeit eine Primzahl ist. Die Eigenschaft „die Primzahl p mit $2 \leq p \leq \lfloor \sqrt{n} \rfloor$ ist kein Teiler von n “ ist dabei nicht geeignet, da man im allgemeinen zu viele derartige Zeugen zwischen 2 und $\lfloor \sqrt{n} \rfloor$ bemühen müßte, um sicher auf die Primzahleigenschaft schließen zu können.

Es sei $E(a)$ eine (noch genauer zu definierende) Eigenschaft, die einer Zahl a mit $1 \leq a \leq n-1$ zukommen kann und für die gilt:

- (i) $E(a)$ ist algorithmisch mit geringem Aufwand zu überprüfen
- (ii) falls n Primzahl ist, dann trifft $E(a)$ für alle Zahlen a mit $1 \leq a \leq n-1$ zu
- (iii) falls n keine Primzahl ist, dann trifft $E(a)$ für weniger als die Hälfte aller Zahlen a mit $1 \leq a \leq n-1$ zu.

Falls $E(a)$ gilt, dann heißt a **Zeuge (witness) für die Primzahleigenschaft von n** . Dann läßt sich folgender randomisierte Primzahltest definieren:

Eingabe: $n \in \mathbb{N}$, n ist ungerade, $m \in \mathbb{N}$

Verfahren: Aufruf der Funktion `random_is_prime` (n : INTEGER;
 m : INTEGER): BOOLEAN;

Ausgabe: n wird als Primzahl angesehen, wenn `random_is_prime` (n , m) = TRUE ist, ansonsten wird n nicht als Primzahl angesehen.

```

FUNCTION random_is_prime ( $n$  : INTEGER;
                            $m$  : INTEGER): BOOLEAN;

  VAR idx      : INTEGER;
      a        : INTEGER;
      is_prime : BOOLEAN;

  BEGIN { random_is_prime }
    is_prime := TRUE;

    FOR idx := 1 TO  $m$  DO
      BEGIN
        wähle eine Zufallszahl  $a$  zwischen 1 und  $n-1$ ;
        IF (  $E(a)$  trifft nicht zu )
          THEN BEGIN
            is_prime := FALSE;
            Break;
          END;
      END;

    random_is_prim := is_prime;

  END { random_is_prime };

```

Der Algorithmus versucht also, m Zeugen für die Primzahleigenschaft von n zu finden. Wird dabei zufällig eine Zahl a mit $1 \leq a \leq n-1$ erzeugt, für die $E(a)$ nicht zutrifft, dann wird wegen (ii) die korrekte Antwort gegeben. Ist n Primzahl, dann gibt der Algorithmus ebenfalls wegen (ii) die korrekte Antwort. Wurden m Zeugen für die Primzahleigenschaft von n festge-

stellt, kann es trotzdem sein, daß n keine Primzahl ist, obwohl der Algorithmus angibt, n sei Primzahl. Die Wahrscheinlichkeit, bei einer Zahl n , die nicht Primzahl ist, m Zeugen zu finden, ist wegen (iii) kleiner als $(1/2)^m$, d.h. die Wahrscheinlichkeit einer fehlerhaften Entscheidung ist in diesem Fall kleiner als $(1/2)^m$. Insgesamt ist die Wahrscheinlichkeit einer korrekten Antwort des Algorithmus größer als $1 - (1/2)^m$. Da wegen (i) die Eigenschaft $E(a)$ leicht zu überprüfen ist, ist hiermit ein effizientes Verfahren beschrieben, das mit beliebig hoher Wahrscheinlichkeit eine korrekte Antwort liefert. Diese ist beispielsweise für $m = 20$ größer als 0,999999.

Die Frage stellt sich, ob es geeignete Zeugeneigenschaften $E(a)$ gibt. Einen ersten Versuch legt der Satz von Fermat nahe:

Satz 7.1-2:

Ist $n \in \mathbf{N}$ eine Primzahl und $a \in \mathbf{N}$ eine Zahl mit $\text{ggT}(a, n) = 1$, dann ist $a^{n-1} \equiv 1 \pmod{n}$.

Definiert man $E(a) = „\text{ggT}(a, n) = 1 \text{ und } a^{n-1} \equiv 1 \pmod{n}“$, dann sieht man, daß (i) und (ii) gelten. Leider gilt (iii) für diese Eigenschaft $E(a)$ nicht. Es gibt nämlich unendlich viele Zahlen n , die **Carmichael-Zahlen**, die folgende Eigenschaft besitzen: n ist *keine* Primzahl, und es ist $a^{n-1} \equiv 1 \pmod{n}$ für alle a mit $\text{ggT}(a, n) = 1$. Ist n also eine Carmichael-Zahl und a eine Zahl mit $1 \leq a \leq n-1$ und $\text{ggT}(a, n) = 1$, dann ist $a^{n-1} \equiv 1 \pmod{n}$. In diesem Fall gilt also für alle Zahlen a mit $1 \leq a \leq n-1$ und $\text{ggT}(a, n) = 1$ die Eigenschaft $E(a)$, und das sind mehr als die Hälfte aller Zahlen a mit $1 \leq a \leq n-1$.

Ein zweiter Versuch zur geeigneten Definition einer Zeugeneigenschaften $E(a)$ erweitert den ersten Versuch und wird durch die folgenden beiden Sätze begründet:

Satz 7.1-3:

Es sei n eine Primzahl. Dann gilt $x^2 \equiv 1 \pmod{n}$ genau dann, wenn $x \equiv 1 \pmod{n}$ oder $x \equiv -1 \equiv n-1 \pmod{n}$ ist.

Es sei n eine Primzahl mit $n > 2$. Dann ist n ungerade, d.h. $n = 1 + 2^j \cdot r$ mit ungeradem r und $j > 0$ bzw. $n-1 = 2^j \cdot r$. Ist a eine Zahl mit $1 \leq a \leq n-1$, dann ist $\text{ggT}(a, n) = 1$ und folglich $a^{n-1} \equiv 1 \pmod{n}$. Wegen $a^{n-1} = a^{(2^{j-1} \cdot r) \cdot 2} = (a^{2^{j-1} \cdot r})^2 \equiv 1 \pmod{n}$ folgt mit Satz 7.1-3 (dort wird $x = a^{2^{j-1} \cdot r}$ gesetzt): $a^{2^{j-1} \cdot r} \equiv 1 \pmod{n}$ oder $a^{2^{j-1} \cdot r} \equiv -1 \pmod{n}$. Ist hierbei $j-1 > 0$ und

$a^{2^{j-1} \cdot r} \equiv 1 \pmod{n}$, dann kann man den Vorgang des Wurzelziehens wiederholen: In Satz 7.1-3 wird $x = a^{2^{j-2} \cdot r}$ gesetzt usw. Der Vorgang ist spätestens dann beendet, wenn $a^{2^{j-j} \cdot r} = a^r$ erreicht ist. Es gilt daher der folgende Satz:

Satz 7.1-4:

Es sei n eine ungerade Primzahl, $n = 1 + 2^j \cdot r$ mit ungeradem r und $j > 0$. Dann gilt für jedes a mit $1 \leq a \leq n-1$:

Die Folge $(a^r, a^{2r}, a^{4r}, \dots, a^{2^{j-1} \cdot r}, a^{2^j \cdot r})$ der Länge $j+1$, wobei alle Werte modulo n reduziert werden, hat eine der Formen

$(1, 1, 1, \dots, 1, 1)$ oder

$(*, *, \dots, *, -1, 1, \dots, 1, 1)$.

Hierbei steht das Zeichen „*“ für eine Zahl, die verschieden von 1 oder -1 ist.

Wenn die in Satz 7.1-4 beschriebene Folge eine der drei Formen

$(*, *, \dots, *, 1, 1, \dots, 1, 1)$,

$(*, *, \dots, *, -1)$ oder

$(*, *, \dots, *, *, \dots, *)$

aufweist, dann ist n mit Sicherheit keine Primzahl. Andererseits ist es nicht ausgeschlossen, daß für eine ungerade zusammengesetzte Zahl n und eine Zahl a mit $1 \leq a \leq n-1$ die Folge $(a^r, a^{2r}, a^{4r}, \dots, a^{2^{j-1} \cdot r}, a^{2^j \cdot r}) \pmod{n}$ eine der beiden Formen $(1, 1, 1, \dots, 1, 1)$ oder $(*, *, \dots, *, -1, 1, \dots, 1, 1)$ hat. In diesem Fall heißt n **streng pseudoprim zu Basis a** . Es gilt folgender Satz (Beweis siehe Literatur):

Satz 7.1-5:

Es sei n eine ungerade zusammengesetzte Zahl. Dann ist n streng pseudoprim für höchstens ein Viertel aller Basen a mit $1 \leq a \leq n-1$.

Eine geeignete Zeigeneigenschaften $E(a)$ für die Funktion

```
random_is_prime (n : INTEGER;
                 m : INTEGER): BOOLEAN;
```

ist daher die folgende Bedingung:

$E(a) = „ggT(a, n) = 1 \text{ und}$

$a^{n-1} \equiv 1 \pmod{n} \text{ und}$

die Folge $(a^r, a^{2r}, a^{4r}, \dots, a^{2^{j-1} \cdot r}, a^{2^j \cdot r}) \pmod{n}$ hat eine der Formen

$(1, 1, 1, \dots, 1, 1)$ oder $(*, *, \dots, *, -1, 1, \dots, 1, 1)“$.

Die Wahrscheinlichkeit einer korrekten Antwort des Algorithmus mit dieser Zeugeneigenschaft ist wegen Satz 7.1-5 größer als $1 - (1/4)^m$.

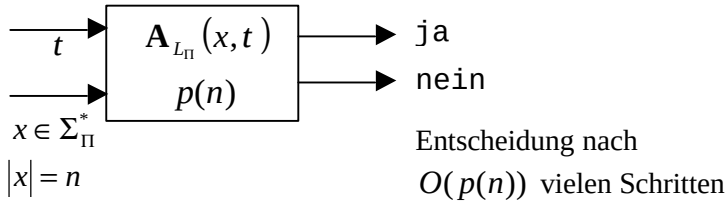
Es ist übrigens nicht notwendig, für eine große Anzahl von Zahlen a mit $1 \leq a \leq n-1$ die Zeugeneigenschaft zu überprüfen um sicher zu gehen, daß n eine Primzahl ist: Nur eine zusammengesetzte Zahl $n < 2,5 \cdot 10^{10}$, nämlich $n = 3.215.031.751$, ist streng pseudoprim zu den vier Basen $a = 2, 3, 5$ und 7 . Für praktische Belange ist daher das Verfahren ein effizienter Primzahltest. Untersucht man große Zahlen, die spezielle Formen aufweisen, etwa Mersenne-Zahlen, auf Primzahleigenschaft bieten sich speziell angepaßte Testverfahren an. Schließlich gibt es eine Reihe von Testverfahren, die andere zahlentheoretische Eigenschaften nutzen.

7.2 Modelle randomisierter Algorithmen

In Zusammenführung der obigen Ansätze mit den Modellen aus der Theorie der Berechenbarkeit (Turingmaschinen, deterministische und nichtdeterministische Algorithmen) wurde eine Reihe weiterer Berechnungsmodelle entwickelt. Im folgenden werden wieder Entscheidungsprobleme betrachtet.

Ausgehend von der Klasse **P** der deterministisch polynomiell entscheidbaren Probleme erweitert man deren Algorithmen nicht um die Möglichkeit der Verwendung nichtdeterministisch erzeugter Zusatzinformationen (Beweise) wie beim Übergang zu den polynomiellen Verifizierern, sondern läßt zu, daß bei Verzweigungen während des Ablaufs der Algorithmen (Verzweigungen) ein Zufallsexperiment darüber entscheidet, welche Alternative für den weiteren Ablauf gewählt wird. Man gelangt so zu der Klasse **RP** (randomized polynomial time).

Es sei Π ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$. L_Π liegt in **RP** (bzw. „ Π liegt in **RP**“), wenn es einen Algorithmus (**RP-Akzeptor**) A_{L_Π} gibt, der neben der Eingabe $x \in \Sigma_\Pi^*$ eine Folge $t \in \{0,1\}^*$ zufälliger Bits liest, wobei jedes Bit unabhängig von den vorher gelesenen Bits ist und $P(0) = P(1) = 1/2$ gilt. Nach polynomiell vielen Schritten (in Abhängigkeit von der Größe $|x|$ der Eingabe) kommt der Akzeptor auf die ja-Entscheidung bzw. auf die nein-Entscheidung. Eingaben für A_{L_Π} sind also die Wörter $x \in \Sigma_\Pi^*$ und eine Folge $t \in \{0,1\}^*$ zufälliger Bits:



Für jedes $x \in L_\Pi$ ist $P(A_{L_\Pi}(x, t) = \text{ja}) \geq 1/2$,

für jedes $x \notin L_\Pi$ ist $P(A_{L_\Pi}(x, t) = \text{nein}) = 1$

Man läßt also zur Akzeptanz von $x \in L_\Pi$ einen einseitigen Fehler zu. Jedoch muß bei $x \in L_\Pi$ der Algorithmus A_Π bei mindestens der Hälfte aller möglichen Zufallsfolgen τ auf die ja-Entscheidung kommen. Für die übrigen darf er auch auf die nein-Entscheidung führen. Für $x \notin L_\Pi$ muß er aber immer die nein-Entscheidung treffen. Der einseitige Fehler kann durch Einsatz geeigneter Replikations-Techniken beliebig klein gehalten werden.

Da dem Akzeptor nur polynomielle Zeit zur Verfügung steht, kann er auch nur polynomiell viele Bits der Zufallsfolge t lesen, so daß man annehmen kann, daß t aus polynomiell vielen Bits besteht.

Ein Beispiel für eine Sprache aus **RP** ist die Menge $L_2 = \{n \mid n \in \mathbf{N} \text{ und } n \text{ ist keine Primzahl}\}$ aus Kapitel 5.5. Dort wird $L_2 \in \mathbf{NP}$ gezeigt. Einen **RP**-Akzeptor A_{L_2} für L_2 erhält man aus der Funktion `random_is_prim` aus Kapitel 7.1. A_{L_2} liest zwei Eingaben, nämlich eine Zahl $n \in \mathbf{N}$ und eine Folge $t \in \{0, 1\}^*$ zufälliger Bits der Länge $k = \text{size}(n)$. Diese Zufallsfolge wird als Binärcode einer Zufallszahl a mit $1 \leq a \leq n-1$ interpretiert (die Fälle $a = 0$ und $a = n$ sollen zur vereinfachten Darstellung hier ausgeschlossen sein). Dazu wird angenommen, daß es eine Funktion `make_INTEGER(t, size(n))` gibt, die die Zufallsfolge t in eine natürliche Zahl a mit $1 \leq a \leq n-1$ umwandelt. Dann wird die Bedingung $E(a)$ (siehe Kapitel 7.1) überprüft und eine Entscheidung getroffen:

```

FUNCTION  $A_{L_2}$  ( $n$  : INTEGER;
                $t$  : ... ) : ...;

VAR  $a$  : INTEGER;

BEGIN {  $A_{L_2}$  }
   $a$  := make_INTEGER( $t$ , size( $n$ ));
  IF ( $E(a)$  trifft nicht zu)
  THEN  $A_{L_2}$  := ja
  ELSE  $A_{L_2}$  := nein;
END {  $A_{L_2}$  };

```

Ist $n \in L_2$, d.h. n ist keine Primzahl, dann trifft nach Definition $E(a)$ für weniger als die Hälfte aller Zahlen a mit $1 \leq a \leq n-1$ zu. Daher trifft $E(a)$ für mehr als die Hälfte aller Zahlen a mit $1 \leq a \leq n-1$ nicht zu, und es gilt $P(A_{L_2}(n, t) = \text{ja}) > 1/2$.

Ist $n \notin L_2$, d.h. n ist eine Primzahl, dann trifft nach Definition $E(a)$ für alle Zahlen a mit $1 \leq a \leq n-1$ zu. Daher antwortet A_{L_2} mit $A_{L_2} := \text{nein}$.

Satz 7.2-1:

Es gilt $\mathbf{P} \subseteq \mathbf{RP}$ und $\mathbf{RP} \subseteq \mathbf{NP}$.

Ob diese Inklusionen echt sind, ist nicht bekannt, vieles spricht jedoch dafür.

Beweis:

Es sei $L_\Pi \in \mathbf{P}$, $L_\Pi \subseteq \Sigma_\Pi^*$. Dann gibt es einen deterministischen polynomiell zeitbeschränkten Entscheidungsalgorithmus für L_Π . Dieser kann als **RP**-Akzeptor betrachtet werden, der ohne Zufallsfolge auskommt. Daher gilt $L_\Pi \in \mathbf{RP}$.

Es sei $L_\Pi \in \mathbf{RP}$, $L_\Pi \subseteq \Sigma_\Pi^*$. Dann gibt es einen **RP**-Akzeptor A_{L_Π} für L_Π . Zu zeigen ist, daß es auch einen Verifizierer, wie er für die Klasse **NP** erforderlich ist, für L_Π gibt („**NP**-Verifizierer“). Der **RP**-Akzeptor A_{L_Π} kann als Verifizierer angesehen werden. Bei Eingabe von $x \in \Sigma_\Pi^*$ bekommt dieser als Beweis B_x eine Zufallsfolge t . Ist $x \in L_\Pi$, dann gibt es eine Zufallsfolge t mit $A_{L_\Pi}(x, t) = \text{ja}$ (da $x \in L_\Pi$ und in diesem Fall $P(A_{L_\Pi}(x, t) = \text{ja}) \geq 1/2$ gelten, führen mindestens die Hälfte aller Zufallsfolgen auf die ja-Entscheidung). Ist $x \notin L_\Pi$, dann führen wegen $P(A_{L_\Pi}(x, t) = \text{nein}) = 1$ alle Zufallsfolgen auf die nein-Entscheidung.

///

Die Zusammenführung der Konzepte des Zufalls und des Nichtdeterminismus bei polynomiell zeitbeschränktem Laufzeitverhalten führt auf die Klasse **PCP** (probabilistically checkable proofs). Hierbei werden Verifizierer, wie sie bei der Definition der Klasse **NP** eingeführt wurden, um die Möglichkeit erweitert, Zufallsexperimente auszuführen:

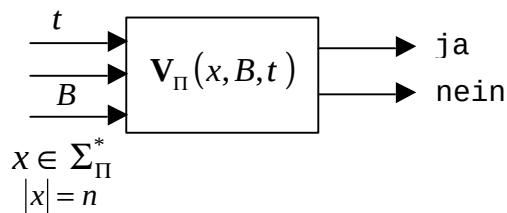
Es seien $r : \mathbf{N} \rightarrow \mathbf{N}$ und $q : \mathbf{N} \rightarrow \mathbf{N}$ Funktionen. Ein Verifizierer (nichtdeterministischer Algorithmus) V_Π heißt $(r(n), q(n))$ -**beschränkter Verifizierer für das Entscheidungsproblem** Π über einem Alphabet Σ_Π , wenn er Zugriff auf eine Eingabe $x \in \Sigma_\Pi^*$ mit $|x| = n$, einen Beweis $B \in \Sigma_0^*$ und eine Zufallsfolge $t \in \{0, 1\}^*$ hat und sich dabei folgendermaßen verhält:

1. V_Π liest zunächst die Eingabe $x \in \Sigma_\Pi^*$ (mit $|x| = n$) und $O(r(n))$ viele Bits aus der Zufallsfolge $t \in \{0, 1\}^*$.
2. Aus diesen Informationen berechnet V_Π $O(q(n))$ viele Positionen der Zeichen von $B \in \Sigma_0^*$, die überhaupt gelesen (erfragt) werden sollen.
3. In Abhängigkeit von den gelesenen Zeichen in B (und der Eingabe x) kommt V_Π auf die ja- bzw. nein-Entscheidung.

Die Entscheidung, die V_Π bei Eingabe von $x \in \Sigma_\Pi^*$, $B \in \Sigma_0^*$ und $t \in \{0, 1\}^*$ liefert, wird mit $V_\Pi(x, B, t)$ bezeichnet.

Es sei Π ein Entscheidungsproblem Π mit der Menge $L_\Pi \subseteq \Sigma_\Pi^*$. $L_\Pi \in \mathbf{PCP}(r(n), q(n))$ („ Π liegt in der Klasse $\mathbf{PCP}(r(n), q(n))$ “), wenn es einen $(r(n), q(n))$ -beschränkten Verifizierer V_Π gibt, der L_Π in folgender Weise akzeptiert:

Für jedes $x \in L_\Pi$ gibt es einen Beweis B_x mit $P(V_\Pi(x, B_x, t) = \text{ja}) = 1$,
für jedes $x \notin L_\Pi$ und alle Beweise B gilt $P(V_\Pi(x, B, t) = \text{nein}) \geq 1/2$.



Für Eingaben $x \in L_\Pi$ gibt es also einen Beweis, der immer akzeptiert wird. Der Versuch, eine Eingabe $x \notin L_\Pi$ zu akzeptieren, scheitert mindestens mit einer Wahrscheinlichkeit $\geq 1/2$.

Im folgenden bezeichnet $\text{poly}(n)$ die Klasse der Polynome. Dann gilt beispielsweise

$$\mathbf{P} = \bigcup_{k=0}^{\infty} \text{TIME}(n^k) = \text{TIME}(\text{poly}(n)) \text{ und } \mathbf{NP} = \bigcup_{k=0}^{\infty} \text{NTIME}(n^k) = \text{NTIME}(\text{poly}(n)).$$

Satz 7.2-2:

$$\mathbf{PCP}(r(n), q(n)) \subseteq \text{NTIME}(q(n) \cdot 2^{O(r(n))})$$

Beweis:

Es sei $L_\Pi \in \mathbf{PCP}(r(n), q(n))$. Dann gibt es einen $(r(n), q(n))$ -beschränkten Verifizierer $\mathbf{V}_\Pi$ für L_Π . Aus diesem kann man auf folgende Weise einen Verifizierer $\tilde{\mathbf{V}}_\Pi$ (im Sinne der Definition einer nichtdeterministischen Turingmaschine) für L_Π konstruieren:

Es sei $x \in \Sigma_\Pi^*$ mit $|x| = n$ und B_x ein Beweis. Für jede der $2^{O(r(n))}$ vielen Zufallsfolgen der Länge $O(r(n))$ simuliert $\tilde{\mathbf{V}}_\Pi$ in jeweils polynomiell vielen Schritten die Berechnung des Verifizierers $\mathbf{V}_\Pi$ und akzeptiert x genau dann, wenn $\mathbf{V}_\Pi$ die Eingabe x für jede Zufallsfolge akzeptiert. $\tilde{\mathbf{V}}_\Pi$ ist daher eine nichtdeterministische Turingmaschine, die x genau dann akzeptiert, wenn $x \in L_\Pi$ gilt. ///

Setzt man im speziellen für $r(n)$ und $q(n)$ Polynome ein, so gilt sogar:

Satz 7.2-3:

$$\mathbf{PCP}(\text{poly}(n), \text{poly}(n)) = \text{NTIME}(2^{\text{poly}(n)})$$

Weiter gelten folgende Aussagen:

Satz 7.2-4:

$$\mathbf{PCP}(0, 0) = \mathbf{P},$$

$$\mathbf{PCP}(0, \text{poly}(n)) = \mathbf{NP},$$

$$\mathbf{PCP}(\log(n), \text{poly}(n)) = \mathbf{NP}$$

Aus $\mathbf{PCP}(r(n), q(n)) \subseteq \mathbf{NTIME}(q(n) \cdot 2^{O(r(n))})$ folgt unmittelbar:

Satz 7.2-5:

$$\mathbf{PCP}(\log(n), 1) \subseteq \mathbf{NP}$$

Die Umkehrung dieser Aussage wird als das wichtigste Resultat der Theoretischen Informatik der letzten 10 Jahre angesehen (Arora, Lund, Motwani, Sudan, Szegedy, 1992). Es hat zu zahlreichen Konsequenzen für die Nicht-Approximierbarkeit von Optimierungsaufgaben aus **NPO** geführt.

Satz 7.2-6:

$$\mathbf{PCP}(\log(n), 1) = \mathbf{NP}$$

„Wie man Beweise verifiziert, ohne sie zu lesen“

Das Ergebnis besagt, daß es zu jeder Menge L_Π für ein Entscheidungsproblem Π aus **NP** einen Verifizierer gibt, der bei jeder Eingabe nur konstant viele Stellen des Beweises liest, die er unter Zuhilfenahme von $O(\log(n))$ vielen zufälligen Bits auswählt, um mit hoher Wahrscheinlichkeit richtig zu entscheiden.

Ein Beispiel für die Anwendung von Satz 7.2-6 ist der Beweis des folgenden Satzes. In ihm wird gezeigt, daß das 3-SAT-Maximierungsproblem nicht in **PTAS** liegt (siehe Kapitel 6.2), d.h. kein polynomiell zeitbeschränktes Approximationsschema besitzt, außer **P** = **NP**.

3-SAT-Maximierungsproblem

Instanz: 1. $I = (K, V)$

$K = \{F_1, F_2, \dots, F_m\}$ ist eine Menge von Klauseln, die aus Booleschen Variablen aus der Menge $V = \{x_1, \dots, x_n\}$ gebildet werden und jeweils die Form $F_i = (y_{i_1} \vee y_{i_2} \vee y_{i_3})$ für $i = 1, \dots, m$ besitzen. Dabei steht y_{i_j} für eine Boolesche Variable (d.h. $y_{i_j} = x_l$) oder für eine negierte Boolesche Variable (d.h.

$$y_{i_j} = \neg x_l)$$

$size(I) = \text{Anzahl der Zeichen in } I$

2. $SOL(I)$ = Belegung der Variablen in V mit Wahrheitswerten TRUE oder FALSE, d.h. $SOL(I)$ ist eine Abbildung $f : V \rightarrow \{TRUE, FALSE\}$
3. Für $f \in SOL(I)$ ist $m(I, f)$ = Anzahl der Klauseln in K , die durch f erfüllt werden
4. $goal = \max$

Es läßt sich ein 2-approximativer Algorithmus für das 3-SAT-Maximierungsproblem angeben (siehe Literatur), d.h. diese Problem liegt in **APX**.

Satz 7.2-7:

Das 3-SAT-Maximierungsproblem liegt nicht in **PTAS**, d.h. es besitzt kein polynomiell zeitbeschränktes Approximationsschema, außer **P = NP**.

Beweis:

Es sei $L_{SAT} = \{F \mid F \text{ ist ein erfüllbarer Boolescher Ausdruck}\}$.

$L_{SAT} \subseteq \Sigma_{BOOLE}^* = \{F \mid F \text{ ist ein Boolescher Ausdruck}\}$, L_{SAT} ist **NP**-vollständig.

Es wird eine Abbildung f definiert, die jedem $F \in \Sigma_{BOOLE}^*$ eine Instanz $f(F) = (K, V)$ für das 3-SAT-Maximierungsproblem zuordnet, wobei $f(F)$ in polynomieller Zeit aus F berechnet werden kann, und die folgende Eigenschaft besitzt:

Ist $F \in L_{SAT}$, dann sind alle Klauseln in K erfüllbar.

Ist $F \notin L_{SAT}$, dann gibt es eine Konstante $e > 0$, so daß mindestens ein Anteil der Größe e aller Klauseln in K nicht erfüllbar ist. Das bedeutet mit $c(F) = |K|$ für $f(F) = (K, V)$:

$$m^*(f(F)) = \begin{cases} = c(F) & \text{für } F \in L_{SAT} \\ \leq c(F) \cdot (1-e) & \text{für } F \notin L_{SAT} \end{cases} . \text{ Mit Satz 6.1-8 folgt, daß es keinen polynomiell}$$

zeitbeschränkten r -approximativen Algorithmus für das 3-SAT-Maximierungsproblem mit $r < 1/(1-e)$ gibt, außer **P = NP**. Daher besitzt das 3-SAT-Maximierungsproblem kein polynomiell zeitbeschränktes Approximationsschema.

Da L_{SAT} **NP**-vollständig ist, gibt es für L_{SAT} nach Satz 7.2-6 einen $(\log(n), 1)$ -beschränkten Verifizierer V_{SAT} . In V_{SAT} werden eine Formel F mit $size(F) = n$, ein Beweis B und eine Zufallsfolge t eingegeben. Man kann annehmen, daß das Alphabet, mit dem Beweise formuliert werden, das Alphabet $\Sigma_0 = \{0, 1\}$ ist, d.h. jeder Beweis B ist eine 0-1-Folge. Man kann weiterhin annehmen, daß V_{SAT} genau $q > 2$ Zeichen von B erfragt (auch wenn nicht alle Zeichen des Beweises zur Verifikation benötigt werden). Da höchstens $c \cdot \log(n)$ Bits (mit einer Kon-

stanten c) aus t und dann q Positionen in B gelesen werden, brauchen nur Beweise betrachtet zu werden, die nicht länger als $q \cdot 2^{c \cdot \log(n)} = q \cdot n^c$ sind.

Es wird eine Menge V von Booleschen Variablen definiert: Für jede Position eines Beweises wird eine Boolesche Variable in V aufgenommen, d.h. $V = \{x_1, x_2, x_3, \dots, x_k\}$ mit $k = q \cdot n^c$. V enthält polynomiell viele Variablen und ist in polynomieller Zeit aus F konstruierbar. Zwischen den Belegungen der Variablenmenge V und den Werten eines Beweises B besteht eine bijektive Abbildung g : Der Wert der j -ten Variablen x_j ist genau dann $x_j = \text{TRUE}$, wenn B an der j -ten Position den Wert $b_j = 1$ hat. Mit $g(x_j)$ werde der Wert an der j -ten Position in B bezeichnet; entsprechend sei $g^{-1}(b_j)$ der Wert der j -ten Variablen in V . Die Abbildung g werde auf Literale $\neg x_j$ durch $g(\neg x_j) = g(x_j)$ erweitert.

Für die Zufallsfolge t seien die q Positionen, die in einem Beweis erfragt werden, die Positionen $t_1, \dots, t_q$. Mit einigen der möglichen Wert-Kombinationen an diesen Positionen kommt $\mathbf{V}_{\text{SAT}}$ bei Auswertung zur ja-Entscheidung, mit anderen kommt $\mathbf{V}_{\text{SAT}}$ zur nein-Entscheidung. Mit A_t wird die Menge derjenigen 0-1-Kombinationen bezeichnet, für die $\mathbf{V}_{\text{SAT}}$ zur nein-Entscheidung kommt. Offensichtlich ist $|A_t| \leq |\Sigma_0|^q = 2^q$. Für jedes $(b_1, \dots, b_q) \in A_t$ wird eine Formel $(y_{t_1} \vee \dots \vee y_{t_q})$ gebildet mit $y_{t_i} = \begin{cases} x_{t_i} & \text{für } b_i = 0 \\ \neg x_{t_i} & \text{für } b_i = 1 \end{cases}$.

Es sei eine Belegung der Variablen in V gegeben. Dann gilt:

(*) Unter dieser Belegung hat die Formel $(y_{t_1} \vee \dots \vee y_{t_q})$, die aus $(b_1, \dots, b_q) \in A_t$ gebildet wurde, genau dann den Wert TRUE, wenn $(g(y_{t_1}), \dots, g(y_{t_q})) \neq (b_1, \dots, b_q)$ ist.

Denn $(y_{t_1} \vee \dots \vee y_{t_q})$ hat genau dann den Wert TRUE, wenn mindestens eines der Literale den Wert TRUE besitzt, etwa $y_{t_j} = \text{TRUE}$. Im Fall $y_{t_j} = x_{t_j}$ bedeutet dieses (nach Definition von y_{t_j}) $b_j = 0$ und $g(y_{t_j}) = 1$, im Fall $y_{t_j} = \neg x_{t_j}$ ist $x_{t_j} = \text{FALSE}$, d.h. $g(y_{t_j}) = 0$, und $b_j = 1$.

Alle so entstehenden Formeln (für alle Zufallsfolgen t) bilden die Formelmenge K' . Diese enthält höchstens $2^{c \cdot \log(n)} \cdot 2^q = 2^q n^c$ viele Formeln. Alle Formeln der Menge K' lassen sich so in Klauseln umformen, daß jede neue Klausel genau 3 Literale enthält. Gilt $q = 3$, haben alle Formeln $(y_{t_1} \vee \dots \vee y_{t_q})$ bereits die gewünschte Form. Für $q > 3$ werden für jede Formel $G = (y_{t_1} \vee \dots \vee y_{t_q})$ $q - 3$ neue Variablen $z_{G,1}, \dots, z_{G,q-3}$ eingeführt und G durch $(y_{t_1} \vee y_{t_2} \vee z_{G,1})$, $(\neg z_{G,1} \vee y_{t_3} \vee z_{G,2})$, $\dots$, $(\neg z_{G,q-4} \vee y_{t_{q-2}} \vee z_{G,q-3})$, $(\neg z_{G,q-3} \vee y_{t_{q-1}} \vee y_{t_q})$ er-

setzt. G ist genau dann erfüllbar, wenn alle so entstandenen Formeln erfüllbar sind. Die Variablenmenge V wird um die neu hinzugenommen Variablen erweitert; entsprechend wird unter einer Belegung $g(z_{G,j})=1$ genau dann gesetzt, wenn $z_{G,j} = \text{TRUE}$ ist.

Diese eventuell erweiterte Klauselmenge bildet die Menge K . Es ist $|K| \leq (q-2) \cdot |K'|$, und auch V behält eine polynomielle Größe.

Die Berechnung von $f(F) = (K, V)$ erfolgt in polynomieller Zeit.

Ist $F \in L_{\text{SAT}}$, dann gibt es einen Beweis B_F , so daß der Verifizierer V_{SAT} für jede Zufallsfolgen t die Entscheidung $V_{\text{SAT}}(F, B_F, t) = \text{ja}$ trifft, denn $P(V_{\text{SAT}}(F, B_F, t) = \text{ja}) = 1$. Ist für eine Zufallsfolgen t die definierte Menge $A_t = \emptyset$, so wurde keine Formel in K' bzw. K aufgenommen. Ist $A_t \neq \emptyset$ und sind $t_1, \dots, t_q$ die Positionen in B_F , die von V_{SAT} aufgrund der Zufallsfolge t erfragt werden, etwa mit den Werten $a_1, \dots, a_q$, dann wird

$$x_{t_i} = \begin{cases} \text{TRUE} & \text{für } a_i = 1 \\ \text{FALSE} & \text{für } a_i = 0 \end{cases} \text{ gesetzt. Es sei } (b_1, \dots, b_q) \in A_t \text{ mit der dazu gebildeten Formel } (y_{t_1} \vee \dots \vee y_{t_q}) \in K'.$$

Dann ist $(a_1, \dots, a_q) \neq (b_1, \dots, b_q)$, denn $(a_1, \dots, a_q)$ führt auf eine ja-Entscheidung, und alle $(b_1, \dots, b_q) \in A_t$ führen auf eine nein-Entscheidung. An mindestens einer Position, etwa an der Position j ist $a_j \neq b_j$. Ist $b_j = 0$, dann ist $a_j = 1$. Das Literal $y_{t_j} = x_{t_j}$ wurde auf TRUE gesetzt. Ist $b_j = 1$, dann ist $a_j = 0$. Das Literal $y_{t_j} = \neg x_{t_j}$ wurde auf TRUE gesetzt. In beiden Fällen ist $(y_{t_1} \vee \dots \vee y_{t_q})$ erfüllt. Das zeigt, daß alle Klauseln in K' und damit alle Klauseln in K erfüllbar sind und damit $m^*(f(F)) = |K| = c(F)$ gilt.

Ist $F \notin L_{\text{SAT}}$, dann gilt für jeden Beweis B , daß der Verifizierer V_{SAT} für mindestens die Hälfte aller $2^{c \cdot \log(n)} = n^c$ Zufallsfolgen die Entscheidung $V_{\text{SAT}}(F, B, t) = \text{nein}$ trifft, denn für jeden Beweis B ist in diesem Fall $P(V_{\text{SAT}}(F, B, t) = \text{nein}) \geq 1/2$. Es sei eine Belegung von V gegeben, der mittels g zugehörige Beweis sei B . t sei eine der $2^{c \cdot \log(n)}/2$ Zufallsfolgen, die auf die nein-Entscheidung führen. Die Positionen in B , die von V_{SAT} aufgrund der Zufallsfolge t erfragt werden, seien $t_1, \dots, t_q$, die Werte an diesen Positionen in B seien $a_1, \dots, a_q$. Dann ist nach Definition $(a_1, \dots, a_q) \in A_t$. Die aus $(a_1, \dots, a_q)$ gebildete Formel $(y_{t_1} \vee \dots \vee y_{t_q})$ hat wegen (*) den Wert FALSE. Für $q > 3$ wurde durch obige Konstruktion die Formel $(y_{t_1} \vee \dots \vee y_{t_q})$ durch $q-3$ Formeln ersetzt. Es läßt sich zeigen, daß im vorliegenden Fall mindestens eine dieser Formeln den Wert FALSE trägt, unabhängig davon, wie die dort enthaltenen neuen Variablen $z_{G,j}$ belegt sind. Zu jeder der $2^{c \cdot \log(n)}/2$ Zufallsfolgen, die auf die

nein-Entscheidung führen, gibt es also mindestens eine Formel in K , die nicht erfüllt ist. Daher beträgt der Anteil nicht erfüllbarer Klauseln mindestens

$$2^{c \cdot \log(n)} / (2 \cdot |K|) \geq 2^{c \cdot \log(n)} / (2 \cdot (q-2) \cdot 2^{q+c \cdot \log(n)}) = 1 / ((q-2) \cdot 2^{q+1}).$$

Mit $e = 1 / ((q-2) \cdot 2^{q+1})$ folgt die Behauptung. ///

8 Übungsaufgaben

Im folgenden werden Übungsaufgaben zu den einzelnen Kapiteln bereitgestellt. Einige der Übungsaufgaben werden in der Veranstaltung behandelt. Aufgaben, die eventuell einiges Nachdenken und Nachschlagen in der Fachliteratur erfordern, sind mit dem Zeichen $\square$ gekennzeichnet.

Aufgabe 1.1.1: Es sei $\Sigma = \{a, b, c\}$. Bestimmen Sie Σ^3 und Σ^* .

Aufgabe 1.1.2: Es sei a ein Buchstabe eines endlichen Alphabets. Bestimmen Sie $\{a^2\}^+$ und $\{a^7\}^* \cdot \{a^3\}$.

Aufgabe 1.1.3: Es seien A und B Mengen. Zeigen Sie: $(A^*)^* = A^*$, $(A \cup B)^* = (A^* \cdot B^*)^*$ und $\emptyset^* = \{e\}$.

Aufgabe 1.1.4: Können L^* bzw. L^+ jemals die leere Menge sein? Unter welchen Umständen sind L^* bzw. L^+ endliche Mengen?

Aufgabe 1.1.5: Es seien A und B Mengen. Beweisen oder widerlegen Sie durch ein entsprechendes Gegenbeispiel die Behauptungen

$$(A \cup B)^* = A^* \cup B^*, \quad (A \cdot B \cup A)^* \cdot A = A \cdot (B \cdot A \cup A)^*$$

Aufgabe 1.1.6: $\square$ Für ein endliches Alphabet $\Sigma = \{a_1, \dots, a_n\}$ kann man die Wörter aus Σ^* in lexikographischer Reihenfolge durchnummerieren (siehe Kapitel 1.1). Welche Nummer enthält in dieser Numerierung das Wort $a_{i_1} \dots a_{i_k}$ der Länge k ? Überlegen Sie sich die Lösung der Aufgabe zunächst für eine zweielementige Menge $\Sigma = \{0, 1\}$.

Aufgabe 1.1.7: $\square$ Zeigen Sie, daß die Menge der rationalen Zahlen gleichmächtig mit der Menge der natürlichen Zahlen ist. Zur Erinnerung: Man kann die rationalen Zahlen durch $\mathbb{Q} = \left\{ \frac{r}{t} \mid r \in \mathbb{Z} \text{ und } t \in \mathbb{N}_{>0} \right\} = \{(r, t) \mid r \in \mathbb{Z} \text{ und } t \in \mathbb{N}_{>0}\}$ definieren.

Aufgabe 1.1.8: Zeigen Sie: Die Mächtigkeit einer unendlich abzählbaren Menge ändert sich nicht, wenn man endlich viele Elemente entfernt.

Aufgabe 1.1.9: ☐ Zeigen Sie: Die Vereinigung abzählbar vieler abzählbarer Mengen ist wieder abzählbar.

Aufgabe 1.1.10: Es seien $f : \mathbf{N} \rightarrow \mathbf{R}$ und $g : \mathbf{N} \rightarrow \mathbf{R}$ zwei Funktionen. Zeigen Sie:

$$O(f(n)) \cdot O(g(n)) \subseteq O(f(n) \cdot g(n)),$$

$$O(f(n) \cdot g(n)) \subseteq |f(n)| \cdot O(g(n)) \text{ und}$$

$$O(f(n)) + O(g(n)) \subseteq O(|f(n)| + |g(n)|).$$

Aufgabe 2.1.1: Geben Sie eine Turingmaschine mit 2 Bändern und dem Eingabealphabet $I = \{0, 1\}$ durch ihre Überföhrungsfunktion an, die zunächst das (neue) Zeichen # auf das 2. Band schreibt und dann das Eingabewort vom 1. Band auf das 2. Band kopiert.

Aufgabe 2.1.2: Geben Sie eine Turingmaschine mit dem Eingabealphabet $I = \{0, 1\}$ durch ihre Überföhrungsfunktion an, die das Eingabewort auf dem 1. Band dupliziert, d.h. steht das Wort $w = a_1 \dots a_n$ im Startzustand auf dem Eingabeband, so steht dort abschließend das Wort $ww = a_1 \dots a_n a_1 \dots a_n$.

Aufgabe 2.1.3: ☐ Geben Sie eine Turingmaschine TM mit $L(TM) = \{a^n b^n c^n \mid n \in \mathbf{N}\}$ an. Welche Speicherplatzkomplexität hat Ihre Turingmaschine?

Aufgabe 2.1.4: Beschreiben Sie die Arbeitsweise einer Turingmaschine, die die Funktion

$$f : \begin{cases} \mathbf{N} \times \mathbf{N} & \rightarrow \mathbf{N} \\ (n, m) & \rightarrow n + m \end{cases} \text{ berechnet. Die Eingabe dieser Turingmaschine auf dem 1. Band}$$

habe die Form $\text{bin}(n)\#\text{bin}(m)$, wobei # ein von 0 und 1 verschiedenes Zeichen ist. Am Ende der Berechnung soll auf dem Ausgabeband $\text{bin}(n+m)\#$ stehen.

Aufgabe 2.2.1: Zeigen Sie, daß das in der Vorlesung behandelte RAM-Programm **P** zur Akzeptanz von

$L(P) = \{ w \mid w \in \{1, 2\}^* \text{ und die Anzahl der 1'en ist gleich der Anzahl der 2'en} \}$ folgende Zeit- und Platzkomplexität hat:

	uniformes Kostenkriterium	logarithmisches Kostenkriterium
Zeitkomplexität	$O(n)$	$O(n \cdot \log(n))$
Platzkomplexität	$O(1)$	$O(\log(n))$

Aufgabe 2.2.2: □ Die folgende Pascal-Prozedur berechnet bei Eingabe einer Zahl $n > 0$ den Wert $c = 2^{2^n} - 1$.

```

PROCEDURE zweier_potenz (    n : INTEGER;
                           VAR c : INTEGER);

VAR idx : INTEGER;

BEGIN {zweier-potenz }
  idx := n;           { Anweisung 1 }
  c   := 2;           { Anweisung 2 }
  WHILE idx > 0 DO    { Anweisung 3 }
    BEGIN
      c := c*c;        { Anweisung 4 }
      idx := idx - 1;   { Anweisung 5 }
    END;
  c := c - 1;          { Anweisung 6 }
END   { zweier-potenz };

```

Man könnte sich die Prozedur als in ein RAM-Programm umgesetzt denken. Dieses hätte dann nicht 6, sondern $c \cdot 6$ viele Anweisungen mit einer Konstanten $c > 0$. Anstatt die Kosten des entsprechenden RAM-Programms zu ermitteln, kann man dieses auch gleich für die Pascal-Prozedur machen, wenn man ohnehin nur an der Größenordnung der Komplexität interessiert ist. Beim uniformen Kostenkriterium zählt man dabei die Anzahl ausgeführter Pascal-Anweisungen; beim logarithmischen Kostenkriterium berücksichtigt man die Stellenzahl der in den beteiligten Variablen gespeicherten Werte. Enthält die Variable idx beispielsweise den Wert m , so verursacht die einmalige Ausführung der Anweisung $idx := idx - 1$ Kosten der Ordnung $O(\log(m))$. Enthält die Variable c den Wert k , so verursacht die einmalige Ausführung der Anweisung $c := c * c$ Kosten der Ordnung $O(\log(k))$.

- (a) Bestimmen Sie die Kosten der Ausführung der Prozedur bei Eingabe von n nach dem uniformen Kostenkriterium. Was fällt Ihnen auf?

- (b) Bestimmen Sie die Kosten der Ausführung der Prozedur bei Eingabe von n nach dem logarithmischen Kostenkriterium.
- (c) Falls die Laufzeit eines Verfahrens nicht nur von der Anzahl der Eingabeparameter, sondern auch von deren numerischen Werten abhängt, dann nimmt man als Meßgröße für die Laufzeitkomplexität die Anzahl der Binärstellen, die benötigt werden, um die Eingabeparameter darzustellen (oder einen Wert, der proportional hierzu ist). Man wendet dann das uniforme Kostenkriterium an. Im vorliegenden Fall kann man als Größe der Eingabe den Wert $k = \lceil \log_2(n+1) \rceil$. Welche Laufzeitkomplexität hat die Pascal-Prozedur unter diesen Voraussetzungen? Wie verbindet sich das Ergebnis mit den Resultaten aus (a) und (b)?

Aufgabe 2.2.3: ☐ Entwerfen Sie ein RAM-Programm mit einer Zeitkomplexität der Ordnung $O(\log(n))$ unter dem uniformen Kostenkriterium zur Berechnung von n^n .

Aufgabe 2.6.1: In der Vorlesung wird eine 3-NDTM zur Lösung des Partitionenproblems mit ganzzahligen Eingabewerten angegeben. Das Partitionenproblem ist folgendes Entscheidungsproblem:

Instanz: $I = \{a_1, \dots, a_n\}$

I ist eine Menge von n natürlichen Zahlen.

Lösung: Entscheidung „ja“, falls es eine Teilmenge $J \subseteq \{1, \dots, n\}$ mit $\sum_{i \in J} a_i = \sum_{i \notin J} a_i$ gibt

(d.h. wenn sich die Eingabeinstanz in zwei Teile zerlegen läßt, deren jeweilige Summen gleich groß sind).

Entscheidung „nein“ sonst.

Für eine Instanz $I = \{a_1, \dots, a_n\}$ sei $B = \sum_{i=1}^n a_i$. Falls B ungerade ist, lautet die Entscheidung für I „nein“, also kann B im folgenden als gerade angenommen werden. Es wird ein Boolescher Ausdruck $T[i, j]$ definiert durch

$T[i, j] = \text{TRUE}$ genau dann, wenn es eine Teilmenge von $\{a_1, \dots, a_i\}$ gibt, deren Elemente sich auf genau j aufsummieren.

- (a) Welche Wahrheitswerte haben $T[1, 0]$, $T[1, a_1]$ und $T[1, j]$ für $j \neq 0$ und $j \neq a_1$?
- (b) Welcher Zusammenhang besteht zwischen der ja/nein-Entscheidung für eine Instanz $I = \{a_1, \dots, a_n\}$ und dem Wert $T[n, B/2]$?
- (c) Zeigen Sie: $T[i, j] = \text{TRUE}$ genau dann, wenn entweder $T[i-1, j] = \text{TRUE}$ oder wenn $a_i \leq j$ und $T[i-1, j-a_i] = \text{TRUE}$ ist.

- (d) Es sei $I = \{1, 9, 5, 3, 8\}$. Berechnen Sie $T[i, j]$ für $i = 1, \dots, n$ und $j = 0, \dots, B/2$.
- (e) Entwickeln Sie einen deterministischen Algorithmus, der bei Eingabe einer Instanz $I = \{a_1, \dots, a_n\}$ für das Partitionenproblem eine ja/nein-Entscheidung trifft und dabei eine Anzahl von Anweisungen der Größenordnung $O(n \cdot B)$ ausführt. Warum handelt es sich hierbei *nicht* um einen Algorithmus, dessen Laufzeit polynomiell in der Größe der Eingabe ist?

Aufgabe 2.6.2: Zeigen Sie, daß die folgenden Funktionen $S_i : \mathbf{N} \rightarrow \mathbf{N}$ platzkonstruierbar sind, indem Sie deterministische Turingmaschinen angeben, die bei Eingabe eines Wortes der Länge n ein spezielles Symbol in die $S_i(n)$ -te Zelle eines ihrer Bänder schreibt, ohne jeweils mehr als $S_i(n)$ viele Zellen auf allen Bändern zu verwenden.

- (a) $S_1(n) = n^2$
 (b) $S_2(n) = n^3 - n^2 + 1$
 (c) $S_3(n) = 2^n$
 (d) $S_4(n) = n!$

Aufgabe 2.6.3: □ Zeigen Sie:

Wird L von einer k -NDTM $TM = (Q, \Sigma, I, d, b, q_0, q_{accept})$ mit Zeitkomplexität $T(n)$ akzeptiert, dann wird L von einer 1-NDTM TM' mit Zeitkomplexität der Ordnung $O(T^2(n))$ akzeptiert. Wie läßt sich das Ergebnis auf deterministische Turingmaschinen übertragen?

Aufgabe 3.1.2: Zeigen Sie, daß die folgenden Mengen entscheidbar sind. Dabei werden natürliche Zahlen in die entsprechenden Turingmaschinen in Binärokodierung eingegeben.

- (a) $\{n^2 \mid n \in \mathbf{N}\}$
 (b) $\{2^n \mid n \in \mathbf{N}\}$
 (c) $\{n \mid n \in \mathbf{N}, n \geq 2 \text{ und } 2n = p + q \text{ mit Primzahlen } p \text{ und } q\}$; ist die Menge $\{n \mid n \in \mathbf{N}, n \geq 2 \text{ und } 2n = p - q \text{ mit Primzahlen } p \text{ und } q\}$ entscheidbar?

Aufgabe 3.1.3: Zeigen Sie:

- (a) Die Menge

$$L_{uni} = L(UTM) = \left\{ u \# w \mid \begin{array}{l} u \in \{0, 1\}^*, w \in \{0, 1\}^*, \\ \text{VERIFIZIERE_TM}(w) = \text{TRUE und } u \in L(K_w) \end{array} \right\} \subseteq \{0, 1, \#\}^* \text{ ist}$$

rekursiv aufzählbar.

- (b) Die Menge $\{0,1\}^* \setminus L_d$ ist nicht entscheidbar; hierbei ist L_d die in der Vorlesung definierte Menge $L_d = \{w \mid w = w_i \text{ und } w_i \notin L(K_{w_i})\}$.
- (c) Die Menge L_{uni} aus Aufgabe (a) ist nicht entscheidbar.

Aufgabe 3.1.4: $\square$ Zeigen Sie:

Die Menge

$$L_{ne} = \{w \mid w \in \{0,1\}^*, \text{VERIFIZIERE_TM}(w) = \text{TRUE und } L(K_w) \text{ ist nicht entscheidbar}\}$$

ist nicht rekursiv aufzählbar. Gehen Sie dabei vor wie im Beweis zu $L_e = \{w \mid L(K_w) \text{ ist entscheidbar}\}$. Zu gegebenen $u \in \{0,1\}^*$ und $w \in \{0,1\}^*$ mit $\text{VERIFIZIERE_TM}(w) = \text{TRUE}$ entwerfen Sie eine Turingmaschine $TM_{u,w}$ mit

$$L(TM_{u,w}) = \begin{cases} \{0,1\}^* \# \{0,1\}^* & \text{falls } u \in L(K_w) \\ L_{uni} & \text{falls } u \notin L(K_w) \end{cases}.$$

Aufgabe 3.1.5: Es sei Σ ein endliches Alphabet und $L \subseteq \Sigma^*$ rekursiv aufzählbar. Zeigen Sie, daß dann $L \leq L_H$ gilt. Wie kann man diese Aussage interpretieren?

Aufgabe 3.1.6: Es sei Σ ein endliches Alphabet und $L \subseteq \Sigma^*$ entscheidbar. Für die Menge $M \subseteq \Sigma^*$ gelte $M \neq \emptyset$ und $M \neq \Sigma^*$. Zeigen Sie, daß dann $L \leq M$ gilt. Wie kann man diese Aussage interpretieren?

Aufgabe 3.1.7: Formalisieren Sie die folgenden umgangssprachlichen Feststellungen und beweisen Sie die entsprechenden Aussagen:

- (a) Es ist nicht entscheidbar, ob ein beliebiges Programm (in einer realen Programmiersprache) für eine Eingaben in eine unendliche Schleife läuft.
- (b) Es ist nicht entscheidbar, ob ein beliebiges Programm (in einer realen Programmiersprache) überhaupt eine Ausgabe erzeugt.

Aufgabe 4.1.1: Geben Sie eine Grammatik zur Erzeugung der Sprache $(0^*1^*)^*$ an.

Aufgabe 4.3.1: Zeigen Sie, daß die Sprache $L = \{0^{2^i} \mid i \geq 0\}$ kontextsensitiv ist.

Aufgabe 4.4.1: Zeigen Sie, daß die Sprache $L = \{0^{2^i} \mid i \geq 0\}$ nicht kontextfrei ist.

Aufgabe 4.4.2: Es seien a , b und c paarweise verschiedene Symbole. Welche der folgenden Sprachen sind kontextfrei? Begründen Sie Ihre Antwort.

- (a) $\{a^k b^k c^i \mid k \in \mathbf{N}, i \in \mathbf{N}\}$
- (b) $\{a^i b^j c^k \mid i \in \mathbf{N}, j \in \mathbf{N}, k \in \mathbf{N}, i > j \text{ oder } j > k\}$
- (c) $\{a^i b^j a^i b^j \mid i \in \mathbf{N}, j \in \mathbf{N}\}$
- (d) $\{ww \mid w \in \{a\}^*\}$
- (e) $\{ww \mid w \in \{a, b\}^*\}$

Aufgabe 4.4.3: Beschreiben Sie, wie ein nichtdeterministischer Kellerautomat mit zwei Kellern das Verhalten einer Turingmaschine simulieren kann.

Aufgabe 4.5.1: □ Beweisen Sie Satz 4.5-1.

Aufgabe 4.5.2: □ Zeigen Sie: Wird die Sprache L von einem nichtdeterministischen endlichen Automaten erkannt, dann gibt es einen deterministischen Automaten, der L erkennt.

Aufgabe 4.5.3: Es seien a und b verschiedene Buchstaben. Zeigen Sie, daß die folgenden Sprachen regulär sind:

- (a) $\{w \mid w \in \{a, b\}^*, \text{ die Anzahl von } a\text{'s ist gerade, die Anzahl von } b\text{'s ist ungerade}\}$
- (b) $\{w \mid w \in \{a, b\}^*, w \text{ enthält höchstens ein Paar aufeinanderfolgende } a\text{'s oder } b\text{'s}\}$
- (c) $\{w \mid w \in \{a, b\}^*, w \text{ enthält nicht die Zeichenkette } aba\}$.

Aufgabe 4.5.4: Es seien a und b verschiedene Buchstaben. Zeigen Sie, daß die folgenden Sprachen nicht regulär sind:

- (a) $\{a^n b^n \mid n \in \mathbf{N}\}$
- (b) $\{a^n \mid n \in \mathbf{N}, n \text{ ist eine Quadratzahl}\}$
- (c) $\{a^n \mid n \in \mathbf{N}, n \text{ ist eine Primzahl}\}$.

Aufgabe 4.5.5: ☐ Es sei L eine kontextfreie Sprache und R eine reguläre Sprache. Zeigen Sie:

- (a) $L \cap R$ ist kontextfrei.
- (b) Es seien a und b verschiedene Buchstaben. Mit Hilfe des uvwyx-Theorems für kontextfreie Sprachen läßt sich zeigen, daß die Sprache $L_1 = \{a^n b^m a^n b^m \mid n \in \mathbf{N}, m \in \mathbf{N}\}$ nicht kontextfrei ist. Warum ist auch $L_2 = \{ww \mid w \in \{a, b\}^+\}$ nicht kontextfrei?
- (c) Es seien $c_1, \dots, c_n$ neue paarweise verschiedene Buchstaben. Warum ist $L_3 = \{x_1 w x_2 w x_3 \mid w \in \{a, b\}^+, x_i \in \{c_1, \dots, c_n\}^* \text{ für } i = 1, 2, 3\}$ nicht kontextfrei?
- (d) Es sei L_p die Menge aller korrekten Pascal-Programme (dasselbe kann man mit Java-, C++-, Cobolprogrammen usw. machen). Es wird eine Teilmenge L_4 korrekter Pascal-Programme definiert durch

$$L_4 = \{\text{PROGRAM A VAR } w : \text{INTEGER}; \text{BEGIN } w := 1; \text{END.} \mid w \in \{a, b\}^+\}.$$
 Zusätzlich wird die reguläre Menge R definiert durch

$$R = \{\text{PROGRAM A VAR } w_1 : \text{INTEGER}; \text{BEGIN } w_2 := 1; \text{END.} \mid w_i \in \{a, b\}^+, i = 1, 2\}$$
 Zeigen Sie mit Hilfe von L_4 und R , daß L_p nicht kontextfrei ist.

Aufgabe 4.6.1: Begründen Sie die Gültigkeit der in Kapitel 4.6 aufgeführten Eigenschaften der verschiedenen Sprachklassen.

Aufgabe 5.3.1: Zeigen Sie:

- (a) Die Klasse **NP** ist abgeschlossen bezüglich Schnitt und Vereinigung.
- (b) Ist L **NP**-vollständig und ist $L_0 \in \mathbf{NP}$ mit $L \leq_m L_0$, dann ist auch L_0 **NP**-vollständig.
- (c) Mit $L \in \mathbf{NP}$ gilt auch $L^* \in \mathbf{NP}$.

Aufgabe 5.3.2: ☐ Zeigen Sie, daß die Sprache

$L = \{(u, w, 0^t) \mid \text{die nichtdeterministische Turingmaschine } K_w \text{ akzeptiert } u \text{ in höchstens } t \text{ Schritten}\}$ **NP**-vollständig ist.

Aufgabe 5.3.3: Zeigen Sie: $\mathbf{P} = \mathbf{NP}$ impliziert, daß $L = \{0, 1\}$ **NP**-vollständig ist.

Aufgabe 5.4.1: $\square$ Zeigen Sie (indem Sie beispielsweise die einschlägige Fachliteratur konsultieren) die Gültigkeit der nicht in der Vorlesung behandelten Reduktionen:

$\text{SAT} \leq_m^p \text{3-CSAT} \leq_m^p \text{RUCKSACK} \leq_m^p \text{PARTITION} \leq_m^p \text{BIN PACKING}$ und

$\text{3-CSAT} \leq_m^p \text{KLIQUE}$ und

$\text{3-CSAT} \leq_m^p \text{GERICHTETER HAMILTONKREIS}$

$\leq_m^p \text{UNGERICHTETER HAMILTONKREIS} \leq_m^p \text{HANDUNGSREISENDER}.$

Literaturauswahl

Die folgende Literaturlauswahl wurde im vorliegenden Text verwendet. Im Text sind die entsprechenden Stellen jedoch nicht explizit gekennzeichnet.

Aho, A.V.; Hopcroft, J.E.; Ullman, J.D.: **The Theory of Parsing, Translation, and Compiling, Vol. I: Parsing**, Addison-Wesley, 1972.

Aho, A.V.; Hopcroft, J.E.; Ullman, J.D.: **The Design and Analysis of Computer Algorithms**, Addison-Wesley, 1974.

Ausiello, G.; Crescenzi, P.; Gambosi, G.; Kann, V.; Marchetti-Spaccamela, A.; Protasi, M.: **Complexity and Approximation**, Springer, 1999.

Bartholomé, A.; Rung, J.; Kern, H.: **Zahlentheorie für Einsteiger**, Vieweg, 1995.

Blum, N.: **Theoretische Informatik**, Oldenbourg, 1998.

Bovet, D.P.; Crescenzi, P.: **Introduction to the Theory of Complexity**, Prentice Hall, 1994.

Dewdney, A.K.: **Der Turing Omnibus**, Springer, 1995.

Erk, K.; Priesel, L.: **Theoretische Informatik**, 2. Aufl., Springer 2002.

Garey, M.R.; Johnson, D.: **Computers and Intractability, A Guide to the Theory of NP-Completeness**, Freeman, 1979.

Harel, D.: **Algorithmics**, 2nd Ed., Addison Wesley, 1992.

Hopcroft, J.E.; Ullman, J.D.: **Einführung in die Automatentheorie, formale Sprachen und Komplexitätstheorie**, Addison-Wesley, 1989.

Hromkovic, J.: **Algorithmische Konzepte der Informatik**, Teubner, 2001.

Paul, W.J.: **Einführung in die Komplexitätstheorie**, Teubner, 1990.

Salomaa, A.: **Public-Key Cryptography**, 2. Aufl., Springer, 1996.

Schöning, U.: **Theoretische Informatik kurzgefaßt**, 3. Aufl., Spektrum Akademischer Verlag, 1997.

Vossen, G.; Witt, K.-U.: **Grundlagen der Theoretischen Informatik mit Anwendungen**, Vieweg, 2000.

Yan, S.Y.: **Number Theory for Computing**, Springer, 2000.

In der Reihe FINAL sind bisher erschienen:

1. Jahrgang 1991:

1. Hinrich E. G. Bonin; Softwaretechnik, Heft 1, 1991 (ersetzt durch Heft 2, 1992).
2. Hinrich E. G. Bonin (Herausgeber); Konturen der Verwaltungsinformatik, Heft 2, 1991 (überarbeitet und erschienen im Wissenschaftsverlag, Bibliographisches Institut & F. A. Brockhaus AG, Mannheim 1992, ISBN 3-411-15671-6).

2. Jahrgang 1992:

1. Hinrich E. G. Bonin; Produktionshilfen zur Softwaretechnik --- Computer-Aided Software Engineering --- CASE, Materialien zum Seminar 1992, Heft 1, 1992.
2. Hinrich E. G. Bonin; Arbeitstechniken für die Softwareentwicklung, Heft 2, 1992 (3. überarbeitete Auflage Februar 1994), PDF-Format (Passwort: arbeiten).
3. Hinrich E. G. Bonin; Object-Orientedness --- a New Boxologie, Heft 3, 1992.
4. Hinrich E. G. Bonin; Objekt-orientierte Analyse, Entwurf und Programmierung, Materialien zum Seminar 1992, Heft 4, 1992.
5. Hinrich E. G. Bonin; Kooperative Produktion von Dokumenten, Materialien zum Seminar 1992, Heft 5, 1992.

3. Jahrgang 1993:

1. Hinrich E. G. Bonin; Systems Engineering in Public Administration, Proceedings IFIP TC8/ WG8.5: Governmental and Municipal Information Systems, March 3--5, 1993, Lüneburg, Heft 1, 1993 (überarbeitet und erschienen bei North-Holland, IFIP Transactions A-36, ISSN 0926-5473).
2. Antje Binder, Ralf Linhart, Jürgen Schultz, Frank Sperschneider, Thomas True, Bernd Willenbockel; COTEXT --- ein Prototyp für die kooperative Produktion von Dokumenten, 19. März 1993, Heft 2, 1993.
3. Gareth Harries; An Introduction to Artificial Intelligence, April 1993, Heft 3, 1993.
4. Jens Benecke, Jürgen Grothmann, Mark Hilmer, Manfred Hölzen, Heiko Köster, Peter Mattfeld, Andre Peters, Harald Weiss; ConFusion --- Das Produkt des AWÖ-Projektes 1992/93, 1. August 1993, Heft 4, 1993.
5. Hinrich E. G. Bonin; The Joy of Computer Science --- Skript zur Vorlesung EDV ---, September 1993, Heft 5, 1993 (4. ergänzte Auflage März 1995).
6. Hans-Joachim Blanke; UNIX to UNIX Copy --- Interactive application for installation and configuration of UUCP ---, Oktober 1993, Heft 6, 1993.

4. Jahrgang 1994:

1. Andre Peters, Harald Weiss; COMO 1.0 --- Programmierumgebung für die Sprache COBOL --- Benutzerhandbuch, Februar 1994, Heft 1, 1994.

2. Manfred Hölzen; UNIX-Mail --- Schnelleinstieg und Handbuch ---, März 1994, Heft 2, 1994.
3. Norbert Kröger, Roland Seen; EBrain --- Documentation of the 1994 AWÖ-Project Prototype ---, June 11, 1994, Heft 3, 1994.
4. Dirk Mayer, Rainer Saalfeld; ADLATUS --- Documentation of the 1994 AWÖ-Project Prototype -- -, July 26, 1994, Heft 4, 1994.
5. Ulrich Hoffmann; Datenverarbeitungssystem 1, September 1994, Heft 5, 1994. (2. überarbeitete Auflage Dezember 1994)
6. Karl Goede; EDV-gestützte Kommunikation und Hochschulorganisation, Oktober 1994, Heft 6 (Teil 1), 1994.
7. Ulrich Hoffmann; Zur Situation der Informatik, Oktober 1994, Heft 6 (Teil 2), 1994.

5. Jahrgang 1995:

1. Horst Meyer-Wachsmuth; Systemprogrammierung 1, Januar 1995, Heft 1, 1995.
2. Ulrich Hoffmann; Datenverarbeitungssystem 2, Februar 1995, Heft 2, 1995.
3. Michael Guder / Kersten Kalischefski / Jörg Meier / Ralf Stöver / Cheikh Zeine; OFFICE-LINK --- Das Produkt des AWÖ-Projektes 1994/95, März 1995, Heft 3, 1995.
4. Dieter Riebesehl; Lineare Optimierung und Operations Research, März 1995, Heft 4, 1995.
5. Jürgen Mattern / Mark Hilmer; Sicherheitsrahmen einer UTM-Anwendung, April 1995, Heft 5, 1995.
6. Hinrich E. G. Bonin; Publizieren im World-Wide Web --- HyperText Markup Language und die Kunst der Programmierung ---, Mai 1995, Heft 6, 1995
7. Dieter Riebesehl; Einführung in Grundlagen der theoretischen Informatik, Juli 1995, Heft 7, 1995
8. Jürgen Jacobs; Anwendungsprogrammierung mit Embedded-SQL, August 1995, Heft 8, 1995
9. Ulrich Hoffmann; Systemnahe Programmierung, September 1995, Heft 9, 1995 (ersetzt durch Heft 1, 1999).
10. Klaus Lindner; Neuere statistische Ergebnisse, Dezember 1995, Heft 10, 1995

6. Jahrgang 1996:

1. Jürgen Jacobs / Dieter Riebesehl; Computergestütztes Repetitorium der Elementarmathematik, Februar 1996, Heft 1, 1996
2. Hinrich E. G. Bonin; "Schlanker Staat" & Informatik, März 1996, Heft 2, 1996
3. Jürgen Jacobs; Datenmodellierung mit dem Entity-Relationship-Ansatz, Mai 1996, Heft 3, 1996
4. Ulrich Hoffmann; Systemnahe Programmierung, (2. überarbeitete Auflage von Heft 9, 1995), September 1996, Heft 4, 1996 (ersetzt durch Heft 1, 1999).
5. Dieter Riebesehl; Prolog und relationale Datenbanken als Grundlagen zur Implementierung einer NF2-Datenbank (Sommer 1995), November 1996, Heft 5, 1996

7. Jahrgang 1997:

1. Jan Binge, Hinrich E. G. Bonin, Volker Neumann, Ingo Stadtsholte, Jürgen Utz; Intranet-/Internet- Technologie für die Öffentliche Verwaltung --- Das AÖW-Projekt im WS96/97 --- (Anwendungen in der Öffentlichen Verwaltung), Februar 1997, Heft 1, 1997
2. Hinrich E. G. Bonin; Auswirkungen des Java-Konzeptes für Verwaltungen, FTVI'97, Oktober 1997, Heft 2, 1997

8. Jahrgang 1998:

1. Hinrich E. G. Bonin; Der Java-Coach, Oktober 1998, Heft 1, 1998 (CD-ROM, PDF-Format; aktuelle Fassung)
2. Hinrich E. G. Bonin (Hrsg.); Anwendungsentwicklung WS 1997/98 --- Programmierbeispiele in COBOL & Java mit Oracle, Dokumentation in HTML und tcl/tk, September 1998, Heft 2, 1998 (CD-ROM)
3. Hinrich E. G. Bonin (Hrsg.); Anwendungsentwicklung SS 1998 --- Innovator, SNIFF+, Java, Tools, Oktober 1998, Heft 3, 1998 (CD-ROM)
4. Hinrich E. G. Bonin (Hrsg.); Anwendungsentwicklung WS 1998 --- Innovator, SNIFF+, Java, Mail und andere Tools, November 1998, Heft 4, 1998 (CD-ROM)
5. Hinrich E. G. Bonin; Persistente Objekte --- Der Elchtest für ein Java-Programm, Dezember 1998, Heft 5, 1998 (CD-ROM)

9. Jahrgang 1999:

1. Ulrich Hoffmann; Systemnahe Programmierung (3. überarbeitete Auflage von Heft 9, 1995), Juli 1999, Heft 1, 1999 (CD-ROM und Papierform), Postscript-Format, zip-Postscript-Format, PDF-Format und zip-PDF-Format.

10. Jahrgang 2000:

1. Hinrich E. G. Bonin; Citizen Relationship Management, September 2000, Heft 1, 2000 (CD-ROM und Papierform), PDF-Format --- Password: arbeiten
2. Hinrich E. G. Bonin; WI>DATA --- Eine Einführung in die Wirtschaftsinformatik auf der Basis der Web_Technologie, September 2000, Heft 2, 2000 (CD-ROM und Papierform), PDF-Format --- Password: arbeiten
3. Ulrich Hoffmann; Angewandte Komplexitätstheorie, November 2000, Heft 3, 2000 (CD-ROM und Papierform), PDF-Format
4. Hinrich E. G. Bonin; Der kleine XMLer, Dezember 2000, Heft 4, 2000 (CD-ROM und Papierform), PDF-Format, aktuelle Fassung --- Password: arbeiten

11. Jahrgang 2001:

1. Hinrich E. G. Bonin (Hrsg.); 4. SAP-Anwenderforum der FHNON, März 2001, (CD-ROM und Papierform), Downloads & Videos.
2. J. Jacobs / G. Weinrich; Bonitätsklassifikation kleiner Unternehmen mit multivariater linear Diskriminanzanalyse und Neuronalen Netzen; Mai 2001, Heft 2, 2001, (CD-ROM und Papierform), PDF-Format und MS Word DOC-Format --- Password: arbeiten
3. K. Lindner; Simultanttestprozedur für globale Nullhypothesen bei beliebiger Abhängigkeitsstruktur der Einzeltests, September 2001, Heft 3, 2001 (CD-ROM und Papierform).

12. Jahrgang 2002:

1. Hinrich E. G. Bonin: Aspect-Oriented Software Development. März 2002, Heft 1, 2002 (CD-ROM und Papierform), PDF-Format --- Password: arbeiten.
2. Hinrich E. G. Bonin: WAP & WML --- Das Projekt Jagdzeit ---. April 2002, Heft 2, 2002 (CD-ROM und Papierform), PDF-Format --- Password: arbeiten.

Herausgeber:

Prof. Dr. Dipl.-Ing. Dipl.-Wirtsch.-Ing. Hinrich E. G. Bonin Fachhochschule
Nordostniedersachsen (FH NON), Volgershall 1, D-21339 Lüneburg, email:
bonin@fhnon.de

Verlag:

Eigenverlag (Fotographische Vervielfältigung), FH NON

Erscheinungsweise:

ca. 4 Hefte pro Jahr Für unverlangt eingesendete Manuskripte wird nicht gehaftet.
Sie sind aber willkommen.

Copyright:

All rights, including translation into other languages reserved by the authors. No part of this report may be reproduced or used in any form or by any means --- graphic, electronic, or mechanical, including photocopying, recording, taping, or information and retrieval systems --- without written permission from the authors, except for noncommercial, educational use, including classroom teaching purposes.

Copyright Bonin Apr-1995,..., May-2002 all rights reserved