

Einführung in die Logik

Skript zur Vorlesung
von
Thomas Piecha

Sommersemester 2026
Universität Tübingen
Fachbereich Informatik

Inhaltsverzeichnis

1	Einleitung	5
2	Die formale Sprache der Aussagenlogik	12
3	Semantik der Aussagenlogik	18
3.1	Wahrheitstafelverfahren	22
3.2	Logische Folgerung	24
3.3	Logische Äquivalenz	27
4	Normalformen und funktionale Vollständigkeit	30
4.1	Normalformen	31
4.2	Funktionale Vollständigkeit	36
5	Resolutionsverfahren	41
5.1	Resolutionskalkül	42
5.2	Resolutionswiderlegung	43
6	Kalkül des natürlichen Schließens	47
6.1	Strukturregeln	58
6.2	Widerspruchsregel und reductio ad absurdum	60
6.3	Korrektheit und Vollständigkeit	61
7	Die formale Sprache der Quantorenlogik	63
7.1	Einleitende Bemerkungen	63
7.1.1	Prädikatsymbole und Individuenkonstanten	64
7.1.2	Individuenvariablen	66
7.1.3	Quantoren	67
7.2	Syntax der Quantorenlogik	69
7.3	Strukturelle Induktion	71
7.4	Mathematische Induktion	73
8	Semantik der Quantorenlogik	76
8.1	Strukturen	76
8.2	Variablenbelegungen	78
8.3	Gültigkeit	79
8.4	Erfüllbarkeit und Allgemeingültigkeit	83
8.5	Logische Folgerung	85
9	Gesetze der Quantorenlogik, pränexe Normalform und Skolemisierung	87
9.1	Gesetze der Quantorenlogik	87
9.2	Pränexe Normalform	94
9.3	Skolemisierung	96
10	Kalkül des natürlichen Schließens	99
10.1	Der quantorenlogische Kalkül NK	99
10.2	Zur Notwendigkeit der Eigenvariablenbedingung bei $(\forall E)$ und $(\exists B)$. .	103
10.3	Definierbarkeit von $\forall$ bzw. $\exists$	105
10.4	Korrektheit und Vollständigkeit	106

11 Quantorenlogische Resolution	108
11.1 Unifikation und quantorenlogische Resolution	109
11.2 Allgemeinste Unifikatoren, Unifikationsalgorithmus	114
11.3 SLD-Resolution	116
Literaturverzeichnis	125

1 Einleitung

Unter *Logik* versteht man im Wesentlichen die Lehre vom gültigen Schließen. Die Beantwortung der Frage, was ein gültiger Schluss ist, wird ein Hauptgegenstand der “Einführung in die Logik” sein. *Logik*

Im Folgenden werden anhand von Beispielen Begriffe wie (*gültiger*) *Schluss*, *Argument*, *Aussage* und *logische Konstante* vorgestellt, mit denen wir uns dann später genauer auseinandersetzen werden.

Beispiel. Folgendes ist ein gültiger Schluss:

Alle Menschen sind sterblich. Sokrates ist ein Mensch. Also ist Sokrates sterblich.

Er enthält die beiden *Prämissen* “Alle Menschen sind sterblich” und “Sokrates ist ein Mensch”, sowie die *Konklusion* “Sokrates ist sterblich”. Das “Also” zeigt an, dass es sich um die Konklusion handelt (man sagt stattdessen z. B. auch “Folglich . . .”, “Deshalb . . .” oder “Daher . . .”). *Prämissen*
Konklusion

Statt von Schlüssen spricht man auch von Argumenten.

Definition 1.1 Ein *Schluss*, bzw. *Argument*, besteht aus einer Menge von Aussagen (den Prämissen) und einer als Konklusion gekennzeichneten Aussage. *Schluss*
Argument

Schematisch werden Schlüsse auch unter Verwendung eines Schluss-Strichs dargestellt:

$$\begin{array}{c} \text{Alle Menschen sind sterblich.} \\ \text{Sokrates ist ein Mensch.} \\ \hline \text{Sokrates ist sterblich.} \end{array}$$

Die Prämissen stehen über dem Schluss-Strich, die Konklusion darunter.
Im Allgemeinen hat ein Schluss die Form

$$\begin{array}{c} \langle \text{Prämisse}_1 \rangle \\ \vdots \\ \langle \text{Prämisse}_n \rangle \\ \hline \langle \text{Konklusion} \rangle \end{array}$$

wobei $n \geq 0$. Ein Schluss hat genau eine Konklusion. Die Anzahl der Prämissen ist hier endlich; es könnten aber auch unendlich viele Prämissen zugelassen werden.

Ein Schluss kann auch aus nur einer Prämisse und einer Konklusion, oder aus einer Konklusion allein bestehen. Beispiele gültiger Schlüsse dieser Art sind

$$\begin{array}{c} \text{Sokrates ist sterblich.} \\ \hline \text{Sokrates ist sterblich.} \end{array}$$

und

$$\begin{array}{c} \hline \text{Sokrates ist sterblich oder Sokrates ist nicht sterblich.} \end{array}$$

Oft werden Prämissen, die für die Gültigkeit eines Schlusses notwendig sind, stillschweigend vorausgesetzt. Der Schluss

Turing ist ein Mensch.

Turing ist ein Eukaryot.

mag Biologen gültig erscheinen, da sie wissen, dass alle Menschen Eukaryoten sind. Ohne dieses Wissen wird man den Schluss für ungültig halten. Einen davon unabhängig gültigen Schluss erhält man, indem man die vorher implizit verwendete Prämisse, dass alle Menschen Eukaryoten sind, explizit macht:

Alle Menschen sind Eukaryoten.

Turing ist ein Mensch.

Turing ist ein Eukaryot.

Die Gültigkeit des vervollständigten Schlusses erkennt man ohne biologisches Wissen (auch wenn man nicht weiß, was Menschen und Eukaryoten sind, bzw. was oder wer Turing ist, wird man den Schluss für gültig halten). Dies ist möglich, da es für die Gültigkeit eines Schlusses lediglich auf seine Form ankommt.

Für die Gültigkeit eines Schlusses ist es nicht maßgeblich, ob alle in ihm vorkommenden Aussagen wahr sind. Im folgenden Schluss kommen nur wahre Aussagen vor (dass Turing nicht mehr lebt, sei unerheblich):

Einige Menschen sind Logiker.

Turing ist ein Mensch.

Turing ist ein Logiker.

Es handelt sich jedoch nicht um einen gültigen Schluss. Halten wir an der Form dieses Schlusses fest, ersetzen aber “Logiker” durch “Ägypter”, so erhalten wir folgenden ungültigen Schluss:

Einige Menschen sind Ägypter.

Turing ist ein Mensch.

Turing ist ein Ägypter.

Die Ersetzung hat dazu geführt, dass zwar die beiden Prämissen wieder wahre Aussagen sind, die Konklusion nun aber eine falsche Aussage ist. Die Ersetzung hat den Effekt einer Uminterpretation des Wortes “Logiker”: Statt der Standardinterpretation, in der das Wort “Logiker” Menschen bezeichnet, die zur Logik arbeiten, bezeichne es nun Menschen, die in Ägypten leben. Diese neue Interpretation haben wir durch die Ersetzung explizit vorgenommen.

Bemerkung. Die Ersetzung, bzw. Uminterpretation, wurde hierbei einheitlich auf den ganzen Schluss angewandt. Das heißt, Gleiches wurde durch Gleiches ersetzt, und nicht etwa nur eine Ersetzung in der Prämisse oder nur in der Konklusion vorgenommen. Uneinheitliche Ersetzungen lassen wir nicht zu, da diese die Form des Schlusses verändern würden.

Eine Präzisierung dessen, was wir einen gültigen Schluss nennen, kann damit wie folgt lauten:

Präzisierung 1.2 Ein Schluss ist *gültig*, falls es keine Interpretation gibt, für die zwar die Prämissen wahr sind, die Konklusion aber falsch ist. *gültiger Schluss*

Mit anderen Worten: Ein Schluss ist *gültig*, falls für jede Interpretation, für welche die Prämissen wahr sind, auch die Konklusion wahr ist.

Beispiel. Der Schluss

Alle Menschen sind sterblich.	
Sokrates ist ein Mensch.	
Sokrates ist sterblich.	

ist gemäß dieser Präzisierung gültig.

Man spricht auch von *Wahrheitskonservierung*, da sich die Wahrheit von den Prämissen auf die Konklusion “überträgt”, und in diesem Sinne erhalten bleibt.

*Wahrheits-
konservierung*

- Bemerkungen.** (i) Wenn ein Schluss gültig ist, und alle Prämissen wahr sind, dann ist auch die Konklusion wahr.
- (ii) Daraus, dass alle Prämissen und die Konklusion eines Schlusses wahr sind, folgt jedoch nicht, dass der Schluss auch gültig ist.
- (iii) Ein gültiger Schluss kann falsche Prämissen enthalten. In diesem Fall kann auch die Konklusion eine falsche Aussage sein.

Bei einer Interpretation von “Mensch(en)” als Logiker, “sterblich” als Eukaryot(en) und “Sokrates” als Turing ergibt sich der (ebenfalls gültige) Schluss

Alle Logiker sind Eukaryoten.	
Turing ist ein Logiker.	
Turing ist ein Eukaryot.	

Dieser Schluss enthält ausschließlich wahre Aussagen (soweit wir wissen).
Für die Interpretation von “Mensch(en)” als Eukaryot(en), “sterblich” als Logiker und “Sokrates” als Turing, ergibt sich der Schluss

Alle Eukaryoten sind Logiker.	
Turing ist ein Eukaryot.	
Turing ist ein Logiker.	

In diesem Schluss ist die erste Prämisse falsch (denn nicht alle Eukaryoten sind Logiker). Dies schließt jedoch nicht aus, dass der Schluss gemäß unserer Präzisierung 1.2 gültig ist. Denn eine Interpretation, für die nicht alle Prämissen wahr sind, ist für die Frage nach Interpretationen, für welche die Prämissen wahr sind, die Konklusion jedoch falsch ist, irrelevant.

Bemerkung. Schlüsse mit widersprüchlichen Prämissen sind stets gültig, da es in diesem Fall keine Interpretation geben kann, für die alle Prämissen wahr sind.

- Definition 1.3** (i) Eine Aussage heißt *kontradiktorisch*, falls sie für keine Interpretation wahr ist. *kontradiktorisch*
- (ii) Eine Menge von Aussagen heißt *kontradiktorisch*, falls es keine Interpretation gibt, für die alle Aussagen der Menge wahr sind.

- (iii) Ist eine Aussage, bzw. eine Menge von Aussagen, nicht kontradiktorisch, so ist sie *konsistent*. (Dies ist genau dann der Fall, wenn es mindestens eine Interpretation gibt, für welche die Aussage, bzw. alle Aussagen in der Menge, wahr sind.) *konsistent*

Bemerkung. Ein ungültiger Schluss kann einfach dadurch in einen gültigen Schluss verwandelt werden, dass man die Menge der Prämissen zu einer kontradiktorischen Menge erweitert. Zum Beispiel erhalten wir aus dem ungültigen Schluss

$$\frac{\text{Turing ist ein Mensch.}}{\text{Turing ist ein Eukaryot.}}$$

durch Hinzufügen der Prämisse “Turing ist kein Mensch” den gültigen Schluss

$$\frac{\begin{array}{l} \text{Turing ist kein Mensch.} \\ \text{Turing ist ein Mensch.} \end{array}}{\text{Turing ist ein Eukaryot.}}$$

Allerdings ist auch folgender Schluss gültig:

$$\frac{\begin{array}{l} \text{Turing ist kein Mensch.} \\ \text{Turing ist ein Mensch.} \end{array}}{\text{Turing ist kein Eukaryot.}}$$

Er verwendet dieselben Prämissen, aber die Konklusion widerspricht der des vorherigen Schlusses.

Die beiden gültigen Schlüsse sind Beispiele für das Prinzip *ex contradictione (sequitur) quodlibet* (auch: *ex falso (sequitur) quodlibet*).

Um nachzuweisen, dass ein Schluss gültig ist, müssen wir für jede Interpretation, für welche die Prämissen wahr sind, zeigen, dass für sie die Konklusion ebenfalls wahr ist. Eine Interpretation, bei der dies nicht der Fall ist (für die also die Prämissen wahr sind, und die Konklusion falsch ist), ist ein Gegenbeispiel zur Gültigkeit des Schlusses.

Da für die Gültigkeit eines Schlusses beliebige Interpretationen zu betrachten sind, können wir von konkreten Prämissen und Konklusionen abstrahieren, und stattdessen zu Variablen P, Q, C übergehen. Für den eingangs behandelten Schluss erhalten wir dann:

$$\frac{\begin{array}{l} \text{Alle } P \text{ sind } Q. \\ C \text{ ist } P. \end{array}}{C \text{ ist } Q.}$$

Die Gültigkeit eines Schlusses beruht somit allein auf der Form des Schlusses. Egal welche konkreten Einsetzungen wir für P, Q und C vornehmen, es darf nie der Fall einer Einsetzung eintreten, bei der die Prämissen wahr sind und die Konklusion falsch ist.

Zu beachten ist, dass wir in den in einem Schluss vorkommenden Aussagen nur bestimmte Komponenten uminterpretieren, bzw. Einsetzungen nur an bestimmten Stellen (hier P, Q und C) zulassen. Hingegen haben wir z. B. “Alle” konstant beibehalten, und nicht etwa durch “Einige” oder etwas anderes ersetzt. “Alle” und “Einige” sind Beispiele für *logische Konstanten*; sie haben eine festgelegte Bedeutung. Hingegen stehen die *Variablen* P, Q und C jeweils für Ausdrücke, die unterschiedliche Bedeutungen haben können.

logische Konstanten

Die *Quantifikatoren* “Alle” und “Einige”, kurz: *Quantoren*, werden in der *Quantorenlogik*

Quantoren

(auch: *Prädikatenlogik* oder *Logik erster Stufe*, engl. *first-order logic*) behandelt. Damit beschäftigen wir uns im zweiten Teil der “Einführung in die Logik”.

Im ersten Teil wird es um die *Aussagenlogik* (engl. *propositional logic*; auch: *Junktorenlogik*) gehen. In ihr werden Aussagen behandelt, die mit *Konnektiven* (*Junktoren*) zu komplexeren Aussagen zusammengesetzt sein können. Beispiele für Konnektive sind die *Konjunktion* “und”, die *Disjunktion* “oder” und die *Implikation* “Wenn . . . , dann . . . ”; auch die *Negation* “nicht” wird als Konnektiv behandelt. Die Konnektive sind die logischen Konstanten der Aussagenlogik.

Konnektive

Beispiele. Die Aussagen

Die Sonne scheint.

und

$3 \times 7 = 21$.

können z. B. zu den Aussagen

Die Sonne scheint, *und* $3 \times 7 = 21$.

oder

Wenn $3 \times 7 = 21$, *dann* scheint die Sonne.

zusammengesetzt werden.

Unter Verwendung der Negation kann z. B. die Aussage

Die Sonne scheint *nicht*.

bzw. gleichbedeutend die Aussage

Es ist nicht der Fall, dass die Sonne scheint.

gebildet werden.

In der klassischen Aussagenlogik setzt man voraus, dass eine Aussage entweder wahr oder falsch ist. Man sagt auch, dass eine Aussage immer genau einen der beiden *Wahrheitswerte* *wahr* (bzw. “das Wahre”) oder *falsch* (bzw. “das Falsche”) hat. Diese Voraussetzung bezeichnet man als *Bivalenzprinzip*.

Wahrheitswerte

Bivalenzprinzip

Bemerkung. Das Bivalenzprinzip impliziert weder, dass wir Kenntnis des Wahrheitswerts einer Aussage haben, noch, dass wir in jedem Fall entscheiden können, welchen Wahrheitswert eine Aussage hat. Falls jedoch etwas nicht genau einen der beiden Wahrheitswerte hat, kann es keine Aussage im Sinne des Bivalenzprinzips sein.

Der Wahrheitswert von mit Konnektiven aus Aussagen zusammengesetzten Aussagen hängt sowohl von den Wahrheitswerten der Teilaussagen als auch von der Bedeutung der verwendeten Konnektive ab. (Zumindest wird dies bei den in der Aussagenlogik behandelten Konnektiven immer der Fall sein.)

Beispiel. Die Aussage

Sokrates ist Ägypter.

ist falsch. Die Aussage

Menschen sind sterblich.

ist wahr. Die zusammengesetzte Aussage

Sokrates ist Ägypter, *und* Menschen sind sterblich.

hat dann aufgrund der Bedeutung von “und” den Wahrheitswert *falsch*. Die Aussage

Sokrates ist Ägypter, *oder* Menschen sind sterblich.

hat hingegen den Wahrheitswert *wahr*.

Bemerkung. Der auch als “Lügner” bezeichnete Satz

Dieser Satz ist falsch.

ist keine Aussage gemäß dem Bivalenzprinzip, da er weder den Wahrheitswert *wahr* noch den Wahrheitswert *falsch* haben kann:

Angenommen der Satz ist eine Aussage. Dann muss der Satz entweder wahr sein oder er muss falsch sein.

1. Fall: Angenommen der Satz ist wahr. Dann ist es wahr, dass der Satz falsch ist. Also ist der Satz falsch. Widerspruch zur Annahme. Folglich ist der Satz nicht wahr.
2. Fall: Angenommen der Satz ist falsch. Dann ist es falsch, dass der Satz falsch ist. Also ist der Satz nicht falsch. Widerspruch zur Annahme. Folglich ist der Satz nicht falsch.

Somit ist der Satz weder wahr noch falsch, im Widerspruch zur Annahme, der Satz sei eine Aussage. Also ist der Satz keine Aussage.

Neben der Untersuchung semantischer Eigenschaften, wie der Gültigkeit von Schlüssen oder den Wahrheitswerten von Aussagen, sind Kalküle von großer Bedeutung. In der Vorlesung werden wir den *Kalkül des natürlichen Schließens* behandeln. Dieser besteht aus Regeln, mit denen wir von Aussagen schrittweise zu neuen Aussagen übergehen können. Ob eine bestimmte Regel auf gegebene Aussagen angewendet werden kann, hängt dabei nicht von semantischen Eigenschaften (z. B. den Wahrheitswerten) der Aussagen ab, sondern lediglich von deren logischer Form. Im Kalkül kann so eher die Struktur tatsächlicher Argumentationen abgebildet werden. Für rechnergestützte Beweisverfahren sind Kalküle wesentlich. In diesem Zusammenhang ist insbesondere der *Resolutionskalkül* von Bedeutung, den wir ebenfalls behandeln.

*Kalkül des
natürlichen
Schließens*

Resolutionskalkül

Literatur

Es gibt eine Vielzahl von einführenden Lehrbüchern zur Logik, von denen hier nur die folgenden genannt seien:

- M. Ben-Ari, *Mathematical Logic for Computer Science. Third Edition*, Springer, 2012.
- D. van Dalen, *Logic and Structure. Fifth Edition*, Springer, 2013.
- U. Schöning, *Logik für Informatiker*, 5. Auflage, Spektrum Akademischer Verlag, 2000.
- R. M. Smullyan, *A Beginner's Guide to Mathematical Logic*, Dover Publications, 2014.

Das Skript bedient sich verschiedener Quellen (siehe [Literaturverzeichnis](#)). Es wird darauf verzichtet, die Verwendung dieser Quellen immer kenntlich zu machen.

Bei der Behandlung der Quantorenlogik werden wir verstärkt mengentheoretische Notation verwenden. Siehe dazu z. B. [Makinson \(2020, § 1 und § 2\)](#).

2 Die formale Sprache der Aussagenlogik

Um Sätze auszuschließen, die keine Aussagen im Sinne des Bivalenzprinzips darstellen, werden wir für die Aussagenlogik eine künstliche, formale Sprache einführen. Deren Ausdrucksstärke ist zwar gegenüber natürlichen Sprachen stark eingeschränkt; für die formale Sprache kann aber sichergestellt werden, dass jede in ihr formulierte Aussage genau einen Wahrheitswert hat.

Die Übersetzung natürlichsprachlicher Sätze oder Argumente in die formale Sprache der Aussagenlogik (*Formalisierung*) kann eine Präzisierung bedeuten, und die Anwendung von Methoden und Resultaten der Aussagenlogik ermöglichen. Falls die Übersetzung adäquat ist, kann dann z. B. von der Gültigkeit des formalisierten Arguments auf die Gültigkeit des natürlichsprachlichen Ausgangsarguments zurückgeschlossen werden. Unser Hauptaugenmerk wird sich jedoch auf die Untersuchung der formalen Sprache selbst richten.

Da die formale Sprache der Aussagenlogik den Gegenstand unserer Untersuchung darstellen wird, spielt sie die Rolle der *Objektsprache*. Die Untersuchung selbst erfolgt in der sogenannten *Metasprache*. In unserem Fall ist dies die (mathematisch angereicherte) natürliche Sprache. Es wird also eine Trennung in Objekt- und Metasprache vorgenommen. “Objektsprache” und “Metasprache” sind keine absoluten Begriffe. Eine Objektsprache kann selbst als Metasprache relativ zu einer anderen Objektsprache fungieren, und eine Metasprache kann relativ zu einer anderen Metasprache die Objektsprache sein (*Sprachstufenhierarchie*).

Objektsprache
Metasprache

In der Metasprache *erwähnt* man Ausdrücke der Objektsprache durch *Verwendung* von Ausdrücken der Metasprache.

Verwendung und Erwähnung von Ausdrücken

Um über etwas zu reden, muss man Ausdrücke (d. h. sprachliche Objekte) *verwenden*. Die Gegenstände, über die man redet, werden dabei *erwähnt*.

Beispiel. In dem Satz

Tübingen liegt am Neckar.

wird der Ausdruck “Tübingen” verwendet, um etwas über die Stadt Tübingen zu sagen. Die Stadt Tübingen wird hierbei erwähnt.

Um über Ausdrücke zu reden (d. h., um sie zu erwähnen), verwendet man *Anführungsnamen*. Charakteristisch für Anführungsnamen ist, dass sie das, was sie erwähnen, als Bestandteil enthalten.

Anführungsnamen

Beispiele.

(i) In dem Satz

$$\overbrace{\text{“Tübingen”}}^{\text{verwendet}} \text{ hat 8 Buchstaben.}$$

$$\underbrace{\text{“Tübingen”}}_{\text{erwähnt}}$$

wird ein Anführungsname verwendet. Erwähnt wird nicht die Stadt Tübingen, sondern der Ausdruck “Tübingen”, d. h. eine Zeichenkette mit 8 Buchstaben, die mit “T” beginnt und mit “n” endet.

(ii) Es gilt:

- Die Stadt Tübingen ist nicht Bestandteil des Ausdrucks “Tübingen”.
- Der Ausdruck “Tübingen” ist Bestandteil des Ausdrucks ““Tübingen””.

Ausdrücke können auch ohne Verwendung eines Anführungsnamens erwähnt werden:

Beispiele.

(i) Es gilt:

Der aus dem 21., 8. und 21. Buchstaben des Alphabets (in dieser Reihenfolge) gebildete Ausdruck ist ein Palindrom.

(ii) Sei “ABCD” ein Name für “Uhu”. Dann gilt:

ABCD ist ein Palindrom.

Die Unterscheidung von Verwendung und Erwähnung von Ausdrücken ist insbesondere bei der Einführung einer Sprache wichtig. Um z. B. festzulegen, wie aus zwei Sätzen mit dem Konnektiv “oder” ein neuer Satz gebildet wird, kann man sagen

Wenn A und B Sätze sind, dann ist auch “ A oder B ” ein Satz.

Ist A der Satz “Die Sonne scheint” und B der Satz “ $3 \times 7 = 21$ ”, dann ist nach dieser Festlegung auch “Die Sonne scheint oder $3 \times 7 = 21$ ” ein Satz. Nun sei A der Satz “Die Sonne scheint oder $3 \times 7 = 21$ ” und B der Satz “Tübingen liegt am Neckar”; dann ist auch “Die Sonne scheint oder $3 \times 7 = 21$ oder Tübingen liegt am Neckar” ein Satz. Es können also beliebig lange Sätze gebildet werden.

Zu beachten ist, dass der Ausdruck “ A oder B ” selbst kein Satz ist, da die Ausdrücke “ A ” und “ B ” als solche keine Sätze sind. Nur wenn “ A ” und “ B ” durch Sätze ersetzt werden, ist auch “ A oder B ” (nach dieser Ersetzung) ein Satz. Die Ausdrücke “ A ” und “ B ” werden hier als *metasprachliche Variablen* (d. h. als Variablen in der Metasprache), verwendet, die für Ausdrücke der Objektsprache stehen (gemäß der “Wenn”-Bedingung in unserer Festlegung sind diese Ausdrücke hier Sätze).

*metasprachliche
Variablen*

Ohne Verwendung von metasprachlichen Variablen kann obige Festlegung wie folgt formuliert werden:

Ein Satz gefolgt von “oder”, gefolgt von einem weiteren Satz, ist ebenfalls ein Satz.

Bei der Untersuchung formaler Sprachen ist die Verwendung von metasprachlichen Variablen jedoch vorzuziehen, da sich dadurch metasprachliche Aussagen einfacher formulieren lassen.

Eine weitere Vereinfachung ergibt sich durch die Konvention, Ausdrücke der Objektsprache als *autonym* zu behandeln. Namen von Ausdrücken der Objektsprache sind hierbei die Ausdrücke selbst. Um einen Ausdruck zu erwähnen, kann dieser dann einfach verwendet werden; auf Anführungsnamen kann also verzichtet werden.

autonym

Beispiel. Statt

“($p \rightarrow q$)” ist ein Ausdruck der Objektsprache.

kann man dann sagen:

($p \rightarrow q$) ist ein Ausdruck der Objektsprache.

Durch die Verschiedenartigkeit von objekt- und metasprachlichen Ausdrücken sollten Missverständnisse hierbei ausgeschlossen sein.

Nach diesen Vorbemerkungen führen wir nun die formale Sprache der Aussagenlogik ein. Dazu legen wir zunächst ein Alphabet fest, und definieren dann, welche Ausdrücke über diesem Alphabet wir als Elemente der Sprache zulassen wollen. Eine solche Definition bezeichnet man auch als *Syntax*. Die so eingeführte Sprache ist die Objektsprache; die zu ihrer Einführung verwendete Metasprache ist die natürliche Sprache.

Syntax

Definition 2.1 Das *Alphabet der Sprache der Aussagenlogik* besteht aus folgenden Zeichen:

Alphabet

- (i) *Aussagesymbole*: $p, q, r, \dots$, auch mit Indizes: $p_0, p_1, p_2, \dots$

Aussagesymbole

Es sei $AS = \{p, q, r, \dots\}$ die Menge der *Aussagesymbole*.

- (ii) *Konnektive (logische Konstanten, Junktoren)*: $\top$ (Verum), $\perp$ (Falsum), $\neg$ (Negation), $\wedge$ (Konjunktion), $\vee$ (Disjunktion), $\rightarrow$ (Implikation) und $\leftrightarrow$ (Bimplikation).

Konnektive

- (iii) Klammern: $(,)$.

Definition 2.2 Die *(aussagenlogischen) Formeln* über $AS = \{p, q, r, \dots\}$ sind wie folgt definiert:

Formeln

- (i) Jedes Aussagesymbol in AS ist eine Formel. $\top$ und $\perp$ sind ebenfalls Formeln.

- (ii) Wenn A eine Formel ist, dann auch $\neg A$.

- (iii) Wenn A und B Formeln sind, dann auch $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$ und $(A \leftrightarrow B)$.

Die *Sprache der Aussagenlogik* ist die Menge aller (aussagenlogischen) Formeln. Diese Menge bezeichnen wir auch mit $FORMEL$.

Sprache der Aussagenlogik

Bemerkung. Wir fassen die Klauseln (i)-(iii) als Regeln zur Erzeugung von Formeln auf, so dass Formeln genau die mit diesen Regeln erzeugbaren Ausdrücke sind. Wir verzichten daher auf die weitere Bedingung, dass nichts sonst eine Formel sei, oder dass die Menge der Formeln gemäß (i)-(iii) die kleinste derartige Menge sei. (Wir werden das so auch bei anderen Definitionen entsprechender Form handhaben.)

Bemerkung. Die logischen Konstanten haben noch keine Bedeutung erhalten; dies geschieht erst später, in der Semantik der Aussagenlogik. Um Formeln besser mitteilen oder lesen zu können, verwenden wir aber jetzt schon natürlichsprachliche Ausdrücke:

<i>Formel(schema)</i>	<i>lies/sprich</i>
$\top$	Verum; das Wahre
$\perp$	Falsum; das Falsche
$\neg A$	nicht A ; Negation A
$(A \wedge B)$	A und B ; A Konjunktion B
$(A \vee B)$	A oder B ; A Disjunktion B
$(A \rightarrow B)$	wenn A , dann B ; A impliziert B ; A Pfeil B
$(A \leftrightarrow B)$	A genau dann, wenn B ; A biimpliziert B ; A Doppelpfeil B

Dabei soll jedoch nicht vorausgesetzt werden, dass die logischen Konstanten dieselbe Bedeutung wie die verwendeten natürlichsprachlichen Ausdrücke hätten.

Beispiele. Die folgenden Ausdrücke sind aussagenlogische Formeln:

- (i) p_{176}
- (ii) $(\perp \rightarrow \top)$
- (iii) $\neg\neg\neg q$
- (iv) $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$

Folgende Ausdrücke sind hingegen *keine* Formeln:

- (i) $(\rightarrow \perp)$ (das Vorderglied der Implikation fehlt);
- (ii) $p \leftrightarrow q$ (die Außenklammern fehlen);
- (iii) $(\neg p)$ (Klammern zu viel);
- (iv) $(p_0 \vee p_1 \vee p_2)$ (es fehlen Klammern).

Bemerkungen. (i) Die Sprache der Aussagenlogik ist unsere *Objektsprache*. In ihr kommen ausschließlich Zeichen des Alphabets gemäß Definition 2.1 vor. Wir verwenden $A, B, C, \dots$ (ggf. mit Indizes) als *metasprachliche Variablen* (auch: *Metavariablen*) für in der Objektsprache ausgedrückte Formeln.

metasprachliche Variablen

(ii) Die Aussagesymbole werden auch als *atomare Formeln*, kurz: *Atome*, bezeichnet. Ebenso $\top$ und $\perp$.

atomare Formel

(iii) Nicht-atomare Formeln heißen *komplex*.

komplexe Formel

(iv) Aussagenlogische Formeln werden im Folgenden auch einfach als *Aussagen* bezeichnet.

Aussagen

Bemerkung. Zur *Klammerersparnis* gelten folgende Regeln:

Klammerersparnis

(i) Außenklammern können weggelassen werden.

(ii) *Bindungsstärke*: $\neg$ bindet am stärksten, $\wedge$ und $\vee$ binden stärker als $\rightarrow$ und $\leftrightarrow$.

Bindungsstärke

(iii) Wir verwenden bei Iterationen von $\wedge$, bzw. $\vee$, *Linksklammerung*. Das heißt

Linksklammerung

- $(A_1 \wedge A_2 \wedge A_3 \wedge \dots \wedge A_n)$ steht für $(\dots((A_1 \wedge A_2) \wedge A_3) \dots) \wedge A_n$,
- $(A_1 \vee A_2 \vee A_3 \vee \dots \vee A_n)$ steht für $(\dots((A_1 \vee A_2) \vee A_3) \dots) \vee A_n$.

Beispiele. (i) Die Formel mit Klammerersparnis $p \rightarrow q \wedge r$ steht für $(p \rightarrow (q \wedge r))$.

Die (Außen-)klammern in $(q \wedge r)$ können weggelassen werden, da $\wedge$ stärker bindet als $\rightarrow$ (dies schließt aus, dass $p \rightarrow q \wedge r$ für $((p \rightarrow q) \wedge r)$ stehen kann). Durch Weglassen der Außenklammern in $(p \rightarrow q \wedge r)$ erhält man $p \rightarrow q \wedge r$.

(ii) $p \vee q \vee r$ steht für $((p \vee q) \vee r)$.

Wegen Linksklammerung steht $p \vee q \vee r$ für $(p \vee q) \vee r$, was wegen Weglassung der Außenklammern für $((p \vee q) \vee r)$ steht.

(iii) Die Formel mit Klammerersparnis

$$q \wedge p_1 \wedge p_7 \wedge p_2 \rightarrow q \vee p_1 \vee p_7 \vee p_2$$

steht für die Formel

$$(((q \wedge p_1) \wedge p_7) \wedge p_2) \rightarrow (((q \vee p_1) \vee p_7) \vee p_2))$$

Bemerkung. Bei einer mit den Regeln zur Klammerersparnis erzeugten Formel muss immer *eindeutig* ersichtlich sein, welche Formel gemäß Definition 2.2 der so erzeugten Formel zugrunde liegt.

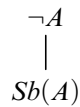
Die Formel $((p \vee q) \vee (p \vee r))$ kann nach den Regeln zur Klammerersparnis zu $p \vee q \vee (p \vee r)$ abgekürzt werden, aber nicht zu $p \vee q \vee p \vee r$, da letztere Formel wegen Linksklammerung für $((p \vee q) \vee p) \vee r$ steht.

Definition 2.3 Die Erzeugung von Formeln kann als *Strukturbaum* dargestellt werden. Dieser ist rekursiv definiert wie folgt:

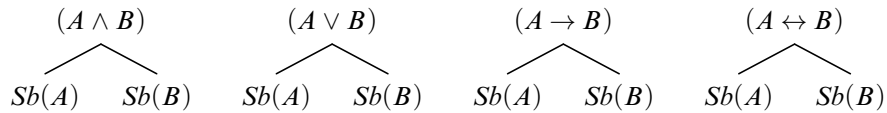
- (i) Der Strukturbaum $Sb(A)$ einer atomaren Formel A ist der Knoten

A

- (ii) Der Strukturbaum $Sb(\neg A)$ einer negierten Formel $\neg A$ ist



- (iii) Die Strukturbäume $Sb((A \wedge B))$, $Sb((A \vee B))$, $Sb((A \rightarrow B))$ bzw. $Sb((A \leftrightarrow B))$ von Formeln $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$ bzw. $(A \leftrightarrow B)$ sind



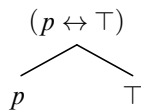
Bemerkung. Sb ist also eine Funktion, die Formeln auf (nach unten verzweigende) binäre Bäume abbildet. Der Wurzelknoten ist jeweils die Formel selbst; die Blätter sind die in der Formel vorkommenden atomaren Formeln.

- Beispiele.** (i) Der Strukturbaum der Formel $\perp$ ist:

$\perp$

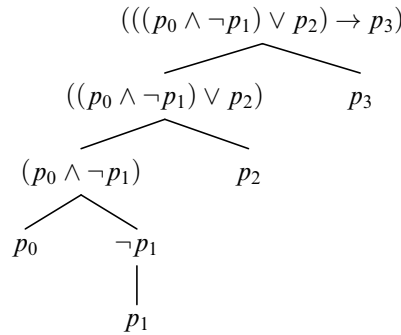
(Denn $\perp$ ist eine atomare Formel; der Strukturbaum $Sb(\perp)$ ist somit der Knoten $\perp$.)

- (ii) Der Strukturbaum der Formel $(p \leftrightarrow \top)$ ist:



(Denn es ist $Sb((p \leftrightarrow \top)) = \begin{array}{c} (p \leftrightarrow \top) \\ \swarrow \quad \searrow \\ Sb(p) \quad Sb(\top) \end{array}$, wobei $Sb(p) = p$ und $Sb(\top) = \top$.)

- (iii) Der Strukturbaum der Formel $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$ ist:



- Definition 2.4** (i) A ist *unmittelbare Teilformel* von $\neg A$. *unmittelbare Teilformel*
- (ii) A und B sind *unmittelbare Teilformeln* von $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$ und $(A \leftrightarrow B)$. *Teilformel*
- (iii) A ist *Teilformel* von B , falls A und B identisch sind, oder A Teilformel einer unmittelbaren Teilformel von B ist. *Teilformel*

Bemerkungen. (i) Im Strukturbaum von A sind die unmittelbaren Teilformeln von A die direkt auf den Wurzelknoten folgenden Knoten. Die Teilformeln von A sind die Knoten, einschließlich des Wurzelknotens.

- (ii) Teilformeln A von B , die nicht identisch mit B sind, bezeichnet man auch als *echte Teilformeln* von B . *echte Teilformel*

- Beispiele.** (i) Die unmittelbaren Teilformeln von $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$ sind $((p_0 \wedge \neg p_1) \vee p_2)$ und p_3 . Die Teilformeln sind
- $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$,
 - $((p_0 \wedge \neg p_1) \vee p_2)$,
 - $(p_0 \wedge \neg p_1)$,
 - $\neg p_1$,
 - p_0, p_1, p_2, p_3 .

Die echten Teilformeln von $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$ sind $((p_0 \wedge \neg p_1) \vee p_2)$, $(p_0 \wedge \neg p_1)$, $\neg p_1$, p_0 , p_1 , p_2 , p_3 .

- (ii) Die Menge der Teilformeln von $(p \vee \neg p)$ ist $\{(p \vee \neg p), \neg p, p\}$. Die Formel $(p \vee \neg p)$ enthält zwei Vorkommen der Teilformel p .

Definition 2.5 Das *Hauptkonnektiv* einer Formel A ist das bei der Erzeugung von A zuletzt eingeführte Konnektiv. (Im Strukturbaum von A ist dies das nur im Wurzelknoten vorkommende Konnektiv.) *Hauptkonnektiv*

- Beispiele.** (i) Das Hauptkonnektiv der Formel $((p_0 \wedge \neg p_1) \vee p_2) \rightarrow p_3$ ist $\rightarrow$.
- (ii) Das Hauptkonnektiv von $(p \wedge (q \wedge r))$ ist das linke Vorkommen von $\wedge$.
- (iii) Das Hauptkonnektiv von $p \wedge q \wedge r$ ist das rechte Vorkommen von $\wedge$. Denn wegen Linksklammerung steht die Formel mit Klammerersparnis $p \wedge q \wedge r$ für die Formel $((p \wedge q) \wedge r)$.

3 Semantik der Aussagenlogik

In der *Semantik* einer Sprache wird jedem Ausdruck der Sprache eine Bedeutung zugewiesen. Die Semantik stellt also eine Interpretation der Sprache dar.

Semantik

In der Semantik der Aussagenlogik werden die Aussagesymbole durch Wahrheitswerte interpretiert. Dies geschieht durch Bewertungen, die jedem Aussagesymbol einen der beiden Wahrheitswerte 0 (“falsch”; als Gegenstand: *das Falsche*) oder 1 (“wahr”; als Gegenstand: *das Wahre*) zuordnen.

Die Bedeutung der logischen Konstanten $\top$, $\perp$, $\neg$, $\wedge$, $\vee$, $\rightarrow$ und $\leftrightarrow$ wird durch Funktionen von Wahrheitswerten festgelegt (*Wahrheitsfunktionalität*). Für $\top$ und $\perp$ sind dies konstante Funktionen. Im Fall der Konnektive $\neg$, $\wedge$, $\vee$, $\rightarrow$ und $\leftrightarrow$ ist der Wahrheitswert der mit dem jeweiligen Konnektiv als Hauptkonnektiv gebildeten komplexen Formel eine Funktion der Wahrheitswerte der unmittelbaren Teilformeln. Als Funktionswerte sind ebenfalls nur die beiden Wahrheitswerte 0 oder 1 zugelassen (*Bivalenzprinzip*).

Wahrheitsfunktionalität

Bivalenzprinzip

Die Voraussetzung von Bivalenz und Wahrheitsfunktionalität ist charakteristisch für die *klassische* (Aussagen-)Logik.

Bemerkung. Die Forderung von Wahrheitsfunktionalität schließt die Behandlung gewisser Aussagen aus. Sei z. B. p die Aussage “Jeder Gegenstand ist mit sich selbst identisch” und q die Aussage “Tübingen liegt am Neckar”. Sowohl p als auch q sind wahr. Jedoch ist die Aussage “Es ist notwendig, dass p ” wahr, während die Aussage “Es ist notwendig, dass q ” falsch ist. Der Operator “Es ist notwendig, dass A ” bildet also den Wahrheitswert der Teilaussage A nicht eindeutig auf einen der beiden Wahrheitswerte ab, und kann daher keine Funktion sein. “Es ist notwendig, dass” ist somit kein wahrheitsfunktionaler Operator, der im Rahmen der aussagenlogischen Semantik behandelt werden könnte.

Definition 3.1 Eine *Bewertung* $\mathcal{I}$ ist eine Funktion, die jedem Aussagesymbol aus $AS = \{p, q, r, \dots\}$ einen der Wahrheitswerte 0 oder 1 zuordnet, d. h. $\mathcal{I} : AS \rightarrow \{0, 1\}$. Für $A \in AS$ heißt A *wahr unter* $\mathcal{I}$, falls $\mathcal{I}(A) = 1$, und *falsch unter* $\mathcal{I}$, falls $\mathcal{I}(A) = 0$.

Bewertung

Bemerkungen. (i) Wir verwenden die für Funktionen übliche Notation. Das heißt

$$\mathcal{I} : AS \rightarrow \{0, 1\}$$

drückt aus, dass $\mathcal{I}$ jedem Element der Definitionsmenge AS genau ein Element der Zielmenge $\{0, 1\}$ zuordnet. Der in diesem Kontext die Zuordnung kennzeichnende Pfeil “ $\rightarrow$ ” hat also nichts mit dem Konnektiv der Implikation zu tun.

(ii) Um eine konkrete Bewertung anzugeben, muss jedem Aussagesymbol in AS einer der beiden Wahrheitswerte 0 oder 1 zugeordnet werden.

(iii) Statt von Bewertungen kann auch von *Interpretationen* $\mathcal{I}$ gesprochen werden.

Interpretation

(In der Literatur ist auch die Bezeichnung “Belegung” gebräuchlich.)

Beispiele. (i) $\mathcal{I}(A) = 1$ für alle $A \in AS$.

Diese Bewertung ordnet allen Aussagesymbolen den Wahrheitswert 1 zu.

(ii) Durch die alleinige Angabe von z. B. $\mathcal{I}(p) = 0$ und $\mathcal{I}(q) = 1$ ist noch *keine* Bewertung $\mathcal{I}$ festgelegt, da für alle Aussagesymbole außer p und q eine Wahrheitswertzuordnung fehlt.

- (iii) $\mathcal{I}(p) = 1, \mathcal{I}(q) = 1, \mathcal{I}(p_3) = 1, \mathcal{I}(A) = 0$ für alle $A \in \text{AS} \setminus \{p, q, p_3\}$. (Wobei $\text{AS} \setminus \{p, q, p_3\}$ für das Komplement der Menge $\{p, q, p_3\}$ in der Menge AS steht, d. h. für die Menge aller Aussagesymbole AS mit Ausnahme von p, q und p_3 .)

Dies ist eine Bewertung; sie ordnet den Aussagesymbolen p, q und p_3 den Wahrheitswert 1 zu, und allen übrigen Aussagesymbolen den Wahrheitswert 0.

Bemerkung. Wir verwenden im Folgenden das Zeichen “:=”, um auszudrücken, dass ein Wert auf der linken Seite des Zeichens per Definition gleich dem Wert auf der rechten Seite ist.

Entsprechend werden wir das Zeichen “: $\Longleftrightarrow$ ” verwenden, um auszudrücken, dass die linke Seite per Definition der Fall ist genau dann, wenn die rechte Seite der Fall ist.

Beide Zeichen gehören zur Metasprache; ebenso das, was links und rechts von ihnen steht. Die linke Seite ist das *Definiendum* (d. h. das, was definiert wird), die rechte das *Definiens* (d. h. das, wodurch definiert wird). Ausdrücke der Form

$$\langle \text{Definiendum} \rangle := \langle \text{Definiens} \rangle$$

und

$$\langle \text{Definiendum} \rangle : \Longleftrightarrow \langle \text{Definiens} \rangle$$

sind immer metasprachliche Aussagen.

Die folgende Definition legt die *Semantik der Aussagenlogik* fest. Die Semantik ist offensichtlich wahrheitsfunktional und erfüllt das Bivalenzprinzip.

Semantik der Aussagenlogik

Definition 3.2 Durch eine Funktion $\llbracket \cdot \rrbracket^{\mathcal{I}} : \text{FORMEL} \rightarrow \{0, 1\}$ wird jeder Formel über $\text{AS} = \{p, q, r, \dots\}$ ein Wahrheitswert 0 oder 1 zugeordnet. Der *Wahrheitswert* $\llbracket A \rrbracket^{\mathcal{I}}$ einer Formel A unter der Bewertung $\mathcal{I}$ ist wie folgt über dem Aufbau von Formeln definiert:
Für Aussagesymbole $A \in \text{AS}$:

Wahrheitswert

$$\llbracket A \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \mathcal{I}(A) = 1, \\ 0 & \text{sonst.} \end{cases}$$

Für die übrigen aussagenlogischen Formeln:

$$\llbracket \top \rrbracket^{\mathcal{I}} := 1$$

$$\llbracket \perp \rrbracket^{\mathcal{I}} := 0$$

$$\llbracket \neg A \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \llbracket A \rrbracket^{\mathcal{I}} = 0, \\ 0 & \text{sonst.} \end{cases}$$

$$\llbracket A \wedge B \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \llbracket A \rrbracket^{\mathcal{I}} = 1 \text{ und } \llbracket B \rrbracket^{\mathcal{I}} = 1, \\ 0 & \text{sonst.} \end{cases}$$

$$\llbracket A \vee B \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \llbracket A \rrbracket^{\mathcal{I}} = 1 \text{ oder } \llbracket B \rrbracket^{\mathcal{I}} = 1, \\ 0 & \text{sonst.} \end{cases}$$

$$\llbracket A \rightarrow B \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \llbracket A \rrbracket^{\mathcal{I}} = 0 \text{ oder } \llbracket B \rrbracket^{\mathcal{I}} = 1, \\ 0 & \text{sonst.} \end{cases}$$

$$\llbracket A \leftrightarrow B \rrbracket^{\mathcal{I}} := \begin{cases} 1 & \text{falls } \llbracket A \rrbracket^{\mathcal{I}} = \llbracket B \rrbracket^{\mathcal{I}} \\ 0 & \text{sonst.} \end{cases}$$

Beispiel. Sei $\mathcal{I}$ eine Bewertung, so dass $\mathcal{I}(p) = 0$, und $\mathcal{I}(A) = 1$ für alle $A \in \text{AS} \setminus \{p\}$. Gilt $\llbracket p \vee \neg p \rrbracket^{\mathcal{I}} = 1$? Es ist $\llbracket p \vee \neg p \rrbracket^{\mathcal{I}} = 1$, falls $\llbracket p \rrbracket^{\mathcal{I}} = 1$ oder $\llbracket \neg p \rrbracket^{\mathcal{I}} = 1$. Es ist $\llbracket p \rrbracket^{\mathcal{I}} = 1$, falls $\mathcal{I}(p) = 1$; dies ist jedoch nicht der Fall. Es ist aber $\llbracket \neg p \rrbracket^{\mathcal{I}} = 1$. Denn $\llbracket \neg p \rrbracket^{\mathcal{I}} = 1$, falls $\llbracket p \rrbracket^{\mathcal{I}} = 0$. Letzteres gilt, da $\mathcal{I}(p) = 0$. Es gilt also $\llbracket p \vee \neg p \rrbracket^{\mathcal{I}} = 1$.

Die Wahrheitswertabhängigkeit bei mit logischen Konstanten gebildeten Formeln kann auch durch *Wahrheitstafeln* ausgedrückt werden:

Wahrheitstafeln

<i>Verum</i>		<i>Falsum</i>		<i>Negation</i>		<i>Konjunktion</i>		
$\top$		$\perp$		A	$\neg A$	A	B	$A \wedge B$
1		0		0	1	0	0	0
				1	0	0	1	0
						1	0	0
						1	1	1

<i>Disjunktion</i>			<i>Implikation</i>			<i>Biimplikation</i>		
A	B	$A \vee B$	A	B	$A \rightarrow B$	A	B	$A \leftrightarrow B$
0	0	0	0	0	1	0	0	1
0	1	1	0	1	1	0	1	0
1	0	1	1	0	0	1	0	0
1	1	1	1	1	1	1	1	1

Durch die linken Spalten werden jeweils alle Kombinationen von Wahrheitswerten der unmittelbaren Teilformeln A, B einer Formel (z. B. $A \wedge B$) mit angegebenem Hauptkonnektiv erfasst. In der rechten Spalte sind jeweils die Wahrheitswerte der Formel (z. B. $A \wedge B$) für die jeweilige Kombination (eine pro Zeile) von Wahrheitswerten der unmittelbaren Teilformeln A, B angegeben. Im Fall der Negation $\neg A$ gibt es nur eine linke Spalte, da es nur eine unmittelbare Teilformel A gibt; bei Verum und Falsum gibt es keine linke Spalte, da $\top$ und $\perp$ keine unmittelbaren Teilformeln haben.

Beispiel. Sei $\mathcal{I}$ die Bewertung mit $\mathcal{I}(p) = 1$, $\mathcal{I}(q) = 0$ und $\mathcal{I}(A) = 1$ für alle übrigen $A \in \text{AS}$. Wir bestimmen den Wahrheitswert der Formel $p \wedge \neg q$ unter $\mathcal{I}$ (d. h. $\llbracket p \wedge \neg q \rrbracket^{\mathcal{I}}$) schrittweise wie folgt:

- (1) Das Hauptkonnektiv von $p \wedge \neg q$ ist $\wedge$. Um den Wahrheitswert der Formel zu bestimmen, müssen wir zunächst die Wahrheitswerte der unmittelbaren Teilformeln p und $\neg q$ bestimmen. Der Wahrheitswert von p ist durch die Bewertung $\mathcal{I}$ gegeben. Der Wahrheitswert von $\neg q$ hängt vom Wahrheitswert von q ab. Dieser ist ebenfalls durch $\mathcal{I}$ gegeben.

Im ersten Schritt notieren wir die Wahrheitswerte der in der Formel vorkommenden Aussagesymbole gemäß $\mathcal{I}$ (die Wahrheitswerte der übrigen Aussagesymbole müssen nicht berücksichtigt werden):

p	q	$p \wedge \neg q$
1	0	1 0

- (2) Aus der Wahrheitstafel für die Negation entnimmt man den Wahrheitswert von $\neg q$ für $\llbracket q \rrbracket^{\mathcal{I}} = 0$. Es ist $\llbracket \neg q \rrbracket^{\mathcal{I}} = 1$, was wir unter $\neg$ notieren:

p	q	$p \wedge \neg q$
1	0	1 1 0

- (3) Aus der Wahrheitstafel für die Konjunktion entnimmt man den Wahrheitswert von $\llbracket p \wedge \neg q \rrbracket^{\mathcal{I}}$ für $\llbracket p \rrbracket^{\mathcal{I}} = 1$ und $\llbracket \neg q \rrbracket^{\mathcal{I}} = 1$, und notiert ihn unter $\wedge$:

p	q	$p \wedge \neg q$
1	0	1
1	1	0

Der unter dem Hauptkonnektiv notierte Wahrheitswert ist der gesuchte Wahrheitswert der Formel. Es ist also $\llbracket p \wedge \neg q \rrbracket^{\mathcal{I}} = 1$.

In der Semantik der Aussagenlogik haben wir zunächst Bewertungen

$$\mathcal{I} : \text{AS} \rightarrow \{0, 1\}$$

eingeführt, d. h. Funktionen, die jedem Aussagesymbol einen der Wahrheitswerte 0 oder 1 zuordnen. Auf der Grundlage von Bewertungen haben wir dann die Bedeutung der logischen Konstanten festgelegt, indem wir die Funktionen $\mathcal{I}$ zu einer Funktion

$$\llbracket \cdot \rrbracket^{\mathcal{I}} : \text{FORMEL} \rightarrow \{0, 1\}$$

erweitert haben, die jeder Formel A einen Wahrheitswert $\llbracket A \rrbracket^{\mathcal{I}}$ zuordnet. Die Bedeutung der logischen Konstanten ist dabei jeweils als Funktion der Wahrheitswerte der unmittelbaren Teilformeln gegeben. Eine Ausnahme bilden Verum und Falsum, deren Bedeutungen durch konstante Funktionen gegeben sind. Jede dieser Funktionen haben wir durch eine Wahrheitstafel dargestellt.

Bemerkung. Die Bezeichnung “logische Konstante” für Funktionen von Wahrheitswerten mag verwirrend erscheinen. Die logischen Konstanten sind jedoch konstant in dem Sinn, dass ihre jeweiligen Bedeutungen durch diese Funktionen *festgelegt* sind. Im Gegensatz dazu sind die Bedeutungen der Aussagesymbole *nicht* konstant, da sie je nach Bewertung variiert werden. Verschiedene Bewertungen stellen also verschiedene Interpretationen der Aussagesymbole dar.

Alternativ kann die Semantik der Aussagenlogik durch Definition einer Relation $\mathcal{I} \models A$ zwischen Bewertungen $\mathcal{I}$ und Formeln A angegeben werden:

Definition 3.3 Die Relation $\mathcal{I} \models A$ (“ A gilt in $\mathcal{I}$ ”, “ A ist wahr unter $\mathcal{I}$ ”), bzw. $\mathcal{I} \not\models A$ (“in $\mathcal{I}$ gilt A nicht”, “ A ist falsch unter $\mathcal{I}$ ”) ist wie folgt definiert:

Für Aussagesymbole $A \in \text{AS}$:

$$\mathcal{I} \models A :\iff \mathcal{I}(A) = 1$$

(Entsprechend sei $\mathcal{I} \not\models A :\iff \mathcal{I}(A) = 0$.)

Für die übrigen aussagenlogischen Formeln:

$$\mathcal{I} \models \top$$

$$\mathcal{I} \not\models \perp$$

$$\mathcal{I} \models \neg A :\iff \mathcal{I} \not\models A \quad (\text{nicht } \mathcal{I} \models A)$$

$$\mathcal{I} \models A \wedge B :\iff \mathcal{I} \models A \text{ und } \mathcal{I} \models B$$

$$\mathcal{I} \models A \vee B :\iff \mathcal{I} \models A \text{ oder } \mathcal{I} \models B$$

$$\mathcal{I} \models A \rightarrow B :\iff \mathcal{I} \not\models A \text{ oder } \mathcal{I} \models B$$

$$\iff \text{Wenn } \mathcal{I} \models A, \text{ dann } \mathcal{I} \models B$$

$$\mathcal{I} \models A \leftrightarrow B :\iff \mathcal{I} \models A \text{ genau dann, wenn } \mathcal{I} \models B$$

Bemerkung. Es gilt $\mathcal{I} \models A$ genau dann, wenn $\llbracket A \rrbracket^{\mathcal{I}} = 1$, und es gilt $\mathcal{I} \not\models A$ genau dann, wenn $\llbracket A \rrbracket^{\mathcal{I}} = 0$.

In der Logik interessieren wir uns für die Frage, ob der Wahrheitswert einer Formel für verschiedene Bewertungen variiert oder für alle Bewertungen gleich ist. Wir möchten Formeln also auf die folgenden Eigenschaften untersuchen:

- | | |
|--|--|
| <p>Definition 3.4 (i) A heißt <i>allgemeingültig</i> (oder <i>tautologisch</i> oder <i>logisch wahr</i>), falls A unter allen Bewertungen wahr ist (d. h. falls $\mathcal{I} \models A$ für alle $\mathcal{I}$ gilt). Notation: $\models A$.</p> <p>(ii) A heißt <i>erfüllbar</i> (oder <i>konsistent</i>), falls A unter mindestens einer Bewertung wahr ist (d. h. falls es eine Bewertung $\mathcal{I}$ gibt, so dass $\mathcal{I} \models A$).</p> <p>(iii) A heißt <i>unerfüllbar</i> (oder <i>inkonsistent</i> oder <i>kontradiktorisch</i> oder <i>logisch falsch</i>), falls A unter keiner Bewertung wahr ist (d. h. falls $\mathcal{I} \not\models A$ für alle $\mathcal{I}$ gilt).</p> <p>(iv) A heißt <i>kontingent</i>, falls A weder allgemeingültig noch unerfüllbar ist.</p> | <p><i>allgemeingültig</i></p> <p><i>erfüllbar</i></p> <p><i>unerfüllbar</i></p> <p><i>kontingent</i></p> |
|--|--|

Beispiel. Wir haben schon gesehen, dass $\mathcal{I} \models p \vee \neg p$ (d. h. $\llbracket p \vee \neg p \rrbracket^{\mathcal{I}} = 1$) für beliebige Bewertungen mit $\mathcal{I}(p) = 0$ gilt. Für beliebige Bewertungen $\mathcal{I}'$ mit $\mathcal{I}'(p) = 1$ gilt $\mathcal{I}' \models p \vee \neg p$, da $\mathcal{I}' \models p$. Somit ist $p \vee \neg p$ unter allen Bewertungen wahr, d. h. $\models p \vee \neg p$.

3.1 Wahrheitstafelverfahren

Im *Wahrheitstafelverfahren* wird der Wahrheitswert einer Formel systematisch für alle Bewertungen der in dieser Formel vorkommenden Aussagesymbole bestimmt. Für den Wahrheitswert einer Formel A unter einer Bewertung $\mathcal{I}$ sind dabei nur jene Werte von $\mathcal{I}$ relevant, die den in A vorkommenden Aussagesymbolen zugeordnet sind. Man überlegt sich leicht, dass Folgendes gilt:

Wahrheitstafelverfahren

Theorem 3.5 (Koinzidenz) Sei A eine beliebige Formel, und seien $\mathcal{I}$ und $\mathcal{I}'$ zwei Bewertungen, so dass für alle in A vorkommenden Aussagesymbole B gilt: $\mathcal{I}(B) = \mathcal{I}'(B)$. Dann gilt $\llbracket A \rrbracket^{\mathcal{I}} = \llbracket A \rrbracket^{\mathcal{I}'}$.

Enthält eine Formel A z. B. nur die Aussagesymbole p und q , so genügen zur Bestimmung des Wahrheitswerts von A unter einer Bewertung $\mathcal{I}$ die Werte $\mathcal{I}(p)$ und $\mathcal{I}(q)$. Die Werte $\mathcal{I}(B)$ für $B \in \text{AS} \setminus \{p, q\}$ sind irrelevant.

Um festzustellen, ob eine Formel allgemeingültig, erfüllbar, unerfüllbar oder kontingent ist, müssen i. A. alle Bewertungen betrachtet werden. Das heißt, es muss für jede Bewertung der Wahrheitswert der Formel unter der jeweiligen Bewertung bestimmt werden. Aufgrund von Bivalenz gibt es für jedes Aussagesymbol 2 Bewertungen. Enthält eine Formel $n > 0$ verschiedene Aussagesymbole, so gibt es 2^n Bewertungen, die für den jeweiligen Wahrheitswert der Formel relevant sind. Mit jedem zusätzlichen in einer Formel vorkommenden Aussagesymbol verdoppelt sich also die Anzahl der zu betrachtenden relevanten Bewertungen.

Im Wahrheitstafelverfahren listet man zunächst alle diese Bewertungen zeilenweise auf, wobei für jedes Aussagesymbol eine Spalte anzulegen ist. Dann bestimmt man den Wahrheitswert der Formel schrittweise für jede Bewertung, d. h. zeilenweise. Unter dem Hauptkonjunktiv der Formel stehen dann die jeweiligen Wahrheitswerte der Formel in einer Spalte, der *Hauptspalte*.

Anhand der Hauptspalte kann nun eine Aussage über Allgemeingültigkeit, Erfüllbarkeit, Unerfüllbarkeit oder Kontingenz der Formel gemacht werden.

Beispiele. (Wir verzichten der Übersichtlichkeit halber auf die Wiederholung der Wahrheitswerte der Aussagesymbole auf der rechten Seite, und schreiben die Wahrheitswerte in der Hauptspalte fett.)

- (i) Die Formel $\neg\neg p \rightarrow p$ enthält als einziges Aussagesymbol p , d. h. es sind $2^1 = 2$ Bewertungen zu betrachten. Man legt also eine Wahrheitstafel für $\neg\neg p \rightarrow p$ an wie folgt:

p	$\neg\neg p \rightarrow p$
0	
1	

Nun bestimmt man den Wahrheitswert der Formel zeilenweise Schritt für Schritt (in derselben Wahrheitstafel):

(1)

p	$\neg\neg p \rightarrow p$
0	1
1	

(2)

p	$\neg\neg p \rightarrow p$
0	0 1
1	

(3)

p	$\neg\neg p \rightarrow p$
0	0 1 1
1	

An dieser Stelle können wir schon eine Aussage über die Erfüllbarkeit der Formel machen: Es gibt eine Bewertung (die erste) unter der die Formel wahr ist; die Formel ist also erfüllbar. Um noch herauszufinden, ob die Formel auch allgemeingültig oder kontingent ist, müssen wir die verbleibende Bewertung untersuchen:

(4)

p	$\neg\neg p \rightarrow p$
0	0 1 1
1	0

(5)

p	$\neg\neg p \rightarrow p$
0	0 1 1
1	1 0

(6)

p	$\neg\neg p \rightarrow p$
0	0 1 1
1	1 0 1

Die Hauptspalte enthält für jede Bewertung den Wahrheitswert 1; die Formel ist also auch allgemeingültig.

- (ii) Die Formel $p \vee q \rightarrow p$ enthält 2 Aussagesymbole p, q , d. h. es sind $2^2 = 4$ Bewertungen zu betrachten. Wahrheitstafel für $p \vee q \rightarrow p$:

p	q	$p \vee q \rightarrow p$
0	0	0 1
0	1	1 0
1	0	1 1
1	1	1 1

Die Formel ist erfüllbar und kontingent.

- (iii) Die Formel $\top \rightarrow \perp$ enthält kein Aussagesymbol. Es sind daher keine Bewertungen zu betrachten; der Wahrheitswert der Formel hängt allein von der Bedeutung der logischen Konstanten $\top$, $\perp$ und $\rightarrow$ ab. Wahrheitstafel für $\top \rightarrow \perp$:

$\top \rightarrow \perp$
1 0 0

Es gibt keine Bewertung, unter der die Formel wahr ist. Die Formel ist also unerfüllbar.

(iv) Wahrheitstafel für $(p \rightarrow q) \leftrightarrow (\neg q \rightarrow \neg p)$:

p	q	$(p \rightarrow q) \leftrightarrow (\neg q \rightarrow \neg p)$			
0	0	1	1	1	1
0	1	1	1	0	1
1	0	0	1	1	0
1	1	1	1	0	1

Die Formel ist erfüllbar und allgemeingültig.

3.2 Logische Folgerung

Der Begriff der logischen Folgerung ist einer der zentralen Begriffe der Logik. Er entspricht dem für natürlichsprachliche Aussagen angegebenen Begriff der Gültigkeit eines Schlusses, wobei an die Stelle natürlichsprachlicher Aussagen nun Formeln treten.

Definition 3.6 Ist eine Formel A unter einer Bewertung $\mathcal{I}$ wahr (d. h. $\llbracket A \rrbracket^{\mathcal{I}} = 1$, bzw. $\mathcal{I} \models A$), so heißt $\mathcal{I}$ *Modell von A* . (Man sagt dann auch, dass A die Bewertung $\mathcal{I}$ als Modell hat.)

Modell von A

Sei Γ eine Menge von Formeln und $\mathcal{I}$ eine Bewertung. Dann heißt $\mathcal{I}$ *Modell von Γ* , falls $\mathcal{I}$ ein Modell aller Formeln $A \in \Gamma$ ist (Notation: $\mathcal{I} \models \Gamma$).

Definition 3.7 Aus Γ folgt (aussagen-)logisch A (Notation: $\Gamma \models A$), falls jedes Modell von Γ ein Modell von A ist. Das heißt

logische Folgerung

$$\Gamma \models A :\iff \text{Für alle Bewertungen } \mathcal{I}: \text{ Wenn } \mathcal{I} \models \Gamma, \text{ dann } \mathcal{I} \models A.$$

Man sagt auch: Γ *impliziert (logisch) A* . Die Formeln in Γ heißen in diesem Zusammenhang *Prämissen*, und die Formel A heißt *Konklusion*.

Bemerkungen. (i) Wir verwenden das Zeichen “ $\models$ ” sowohl für die Modellbeziehung $\mathcal{I} \models A$, bzw. $\mathcal{I} \models \Gamma$, als auch für die logische Folgerung ($\Gamma \models A$). Die jeweilige Bedeutung ist aber durch den Bezug auf entweder ein Modell $\mathcal{I}$ oder auf eine Formelmenge Γ eindeutig.

(ii) Mengen von Formeln schreiben wir auch als kommaseparierte Listen der in der jeweiligen Menge enthaltenen Formeln. Die Liste $A_1, \dots, A_n$ steht für die Menge $\{A_1, \dots, A_n\}$; die Liste Γ, A steht für die Menge $\Gamma \cup \{A\}$ (d. h. für die Vereinigung von Γ und $\{A\}$). Ist Γ die leere Menge $\emptyset$, schreiben wir $\models A$ statt $\emptyset \models A$.

(iii) Die Definition der logischen Folgerung beruht auf der Idee der Wahrheitskonservierung: Eine Folgerungsbehauptung $\Gamma \models A$ gilt genau dann, wenn stets für wahre Prämissen Γ auch die Konklusion A wahr ist.

Beispiele. Überprüfung der logischen Folgerung unter Verwendung von Wahrheitstafeln:

(i) Gilt die Folgerungsbehauptung $p \vee q, \neg p \models q$?

	p	q	$p \vee q$	$\neg p$	$\models$	q
$\rightarrow$	0	0	0	1		0
	0	1	1	1	✓	1
	1	0	1	0		0
	1	1	1	0		1

Es gibt nur eine Bewertung, unter der alle Prämissen wahr sind (zweite Zeile). Unter dieser Bewertung ist auch die Konklusion q wahr. Die Konklusion q ist also unter allen Bewertungen wahr, unter denen auch die Prämissen $p \vee q, \neg p$ wahr sind. Somit gilt $p \vee q, \neg p \models q$.

(ii) Gilt die Folgerungsbehauptung $p \vee q, p \models q$?

	p	q	$p \vee q$	p	$\models$	q
	0	0	0	0		0
	0	1	1	0		1
$\rightarrow$	1	0	1	1	$\times$	0
$\rightarrow$	1	1	1	1	$\checkmark$	1

Es gibt zwei Bewertungen, unter denen alle Prämissen wahr sind (dritte und vierte Zeile). Allerdings ist die Konklusion q unter der dritten Bewertung falsch. Es ist also $p \vee q, p \not\models q$, d. h. $p \vee q, p \models q$ gilt nicht.

(iii) Gilt die Folgerungsbehauptung $p \rightarrow (q \rightarrow r), p \wedge q \models r$?

Es kommen 3 Aussagesymbole vor, die jeweils einen der 2 Wahrheitswerte 0 und 1 annehmen können. Das heißt, es müssen $2^3 = 8$ Bewertungen betrachtet werden; wir benötigen also eine Wahrheitstafel mit 8 Zeilen.

	p	q	r	$p \rightarrow (q \rightarrow r)$	$p \wedge q$	$\models$	r
	0	0	0	1	0		0
	0	0	1	1	0		1
	0	1	0	1	0		0
	0	1	1	1	0		1
	1	0	0	1	0		0
	1	0	1	1	0		1
	1	1	0	0	1		0
$\rightarrow$	1	1	1	1	1	$\checkmark$	1

Die Folgerungsbehauptung gilt. Es gibt nur eine Bewertung, unter der alle Prämissen wahr sind (letzte Zeile), und unter dieser Bewertung ist auch die Konklusion wahr.

Für das Verhältnis zwischen logischer Folgerung und Implikation gilt Folgendes:

Theorem 3.8 (Import-Export) $A_1, \dots, A_n \models B$ genau dann, wenn $\models A_1 \wedge \dots \wedge A_n \rightarrow B$.

Beweis. Das Theorem enthält die beiden Aussagen

(i) Wenn $A_1, \dots, A_n \models B$, dann $\models A_1 \wedge \dots \wedge A_n \rightarrow B$. (*Import*)

(ii) Wenn $\models A_1 \wedge \dots \wedge A_n \rightarrow B$, dann $A_1, \dots, A_n \models B$. (*Export*)

Zu (i): Angenommen $A_1, \dots, A_n \models B$.

Es sei $\mathcal{I}$ eine beliebige Bewertung. Dann sind zwei Fälle zu unterscheiden:

1. Fall: $\mathcal{I}$ ist ein Modell von $A_1 \wedge \dots \wedge A_n$, d. h. $\mathcal{I} \models A_1 \wedge \dots \wedge A_n$.

Nach Definition der Konjunktion ist dann $\mathcal{I} \models A_1, \dots, \mathcal{I} \models A_n$. Aufgrund der Annahme ist $\mathcal{I}$ dann auch ein Modell von B , d. h. $\mathcal{I} \models B$.

Mit der Definition der Implikation folgt dann $\mathcal{I} \models A_1 \wedge \dots \wedge A_n \rightarrow B$.

2. Fall: $\mathcal{I}$ ist kein Modell von $A_1 \wedge \dots \wedge A_n$, d. h. $\mathcal{I} \not\models A_1 \wedge \dots \wedge A_n$.

Nach Definition der Implikation gilt dann $\mathcal{I} \models A_1 \wedge \dots \wedge A_n \rightarrow B$.

In beiden Fällen gilt $\mathcal{I} \models A_1 \wedge \dots \wedge A_n \rightarrow B$.

Da $\mathcal{I}$ beliebig ist, ist $A_1 \wedge \dots \wedge A_n \rightarrow B$ allgemeingültig, d. h. $\models A_1 \wedge \dots \wedge A_n \rightarrow B$.

Zu (ii): Angenommen $\models A_1 \wedge \dots \wedge A_n \rightarrow B$.

Es sei $\mathcal{I}$ ein Modell von $A_1, \dots, A_n$, d. h. $\mathcal{I} \models A_1, \dots, \mathcal{I} \models A_n$. Nach Definition der Konjunktion gilt dann $\mathcal{I} \models A_1 \wedge \dots \wedge A_n$.

Nach Annahme ist $A_1 \wedge \dots \wedge A_n \rightarrow B$ allgemeingültig. Dies schließt wegen $\mathcal{I} \models A_1 \wedge \dots \wedge A_n$ den Fall $\mathcal{I} \not\models B$ aus, da sonst $\not\models A_1 \wedge \dots \wedge A_n \rightarrow B$. Also muss $\mathcal{I} \models B$ gelten. Somit gilt $A_1, \dots, A_n \models B$. QED

Insbesondere gilt für $n = 1$:

Korollar 3.9 $A \models B$ genau dann, wenn $\models A \rightarrow B$.

Bemerkung. Der Nachweis, dass eine Folgerungsbehauptung $A \models B$ gilt, kann also durch den Nachweis der Allgemeingültigkeit von $A \rightarrow B$ erfolgen und umgekehrt.

Man beachte, dass hier *nicht* etwa gesagt wird, dass die logische Folgerung $A \models B$ der Implikation $A \rightarrow B$ entspricht und umgekehrt; was gesagt wird, ist, dass $A \models B$ genau dann gilt, wenn $A \rightarrow B$ allgemeingültig ist.

Dass Implikation und logische Folgerung nicht dasselbe ist, macht man sich leicht an einem Beispiel klar: Angenommen, man weiß, dass $\mathcal{I} \not\models A$ und $\mathcal{I} \models B$. Dann weiß man, dass $\mathcal{I} \models A \rightarrow B$, d. h. dass $A \rightarrow B$ in diesem Fall wahr ist. Hingegen weiß man nicht, ob $A \models B$ gilt.

Bemerkung. Die *klassische Aussagenlogik* ist durch die Relation der logischen Folgerung, bzw. durch die Menge der allgemeingültigen aussagenlogischen Formeln, gegeben.

*klassische
Aussagenlogik*

Theorem 3.10 (Entscheidbarkeit der Aussagenlogik) Die klassische Aussagenlogik ist entscheidbar. Das heißt, es gibt ein Verfahren, das für jede aussagenlogische Formel in endlich vielen Schritten entscheidet, ob die Formel allgemeingültig ist oder nicht.

Beweis. Das Wahrheitstafelverfahren ist ein solches Verfahren.

QED

Bemerkung. Für die auf endliche Prämissenmengen eingeschränkte aussagenlogische Folgerung stellt das Wahrheitstafelverfahren ebenfalls ein Entscheidungsverfahren dar.

Bemerkung. Ein *Formelschema*, wie z. B. $(A \vee B) \leftrightarrow C$, steht für beliebige Formeln von bestimmter Form. Das Wahrheitstafelverfahren kann auch auf Formelschemata angewendet werden. Dabei ist jedoch der Unterschied zwischen Formel, z. B. $(p \vee q) \leftrightarrow r$, und *Formelschema*, z. B. $(A \vee B) \leftrightarrow C$, zu beachten. Zum Beispiel ist

Formelschema

A	B	$A \rightarrow B$
0	0	1
0	1	1
1	0	0
1	1	1

eine Wahrheitstafel für das Formelschema $A \rightarrow B$, das für beliebige Formeln mit dem

Hauptkonnektiv $\rightarrow$ steht. Anders als bei konkreten Formeln stehen hier links keine Bewertungen, sondern die möglichen Wahrheitswerte von A und B . Die für Formeln über Bewertungen definierten Begriffe *allgemeingültig*, *erfüllbar*, *unerfüllbar* und *kontingent* sind deshalb nicht ohne Weiteres auf Formelschemata anwendbar.

So bedeutet z. B. die Behauptung $\models A \vee \neg A$ strenggenommen nicht, dass das Formelschema $A \vee \neg A$ allgemeingültig ist, sondern dass alle Instanzen von $A \vee \neg A$, wie z. B. $(p \wedge q) \vee \neg(p \wedge q)$ oder $p \vee \neg p$, allgemeingültige Formeln sind.

Kommt in der Hauptspalte der Wahrheitstafel eines Formelschemas sowohl 1 als auch 0 vor, so schließt dies nicht aus, dass Instanzen des Schemas allgemeingültige oder unerfüllbare Formeln sein können. So ist z. B. die Formel $p \rightarrow p$ eine allgemeingültige Instanz des Schemas $A \rightarrow B$, während die Formel $\top \rightarrow \perp$ eine unerfüllbare Instanz ist. Das Formelschema $A \rightarrow B$ mit Blick auf die Hauptspalte als kontingent zu bezeichnen, wäre irreführend, da nicht alle Instanzen des Schemas kontingente Formeln sind.

Falls in der Hauptspalte der Wahrheitstafel eines Formelschemas nur 1 vorkommt, sind jedoch alle Instanzen allgemeingültige Formeln. Kommt nur 0 vor, sind alle Instanzen unerfüllbare Formeln. In diesen beiden Fällen ist die Bezeichnung des jeweiligen Schemas als allgemeingültig, bzw. unerfüllbar, unproblematisch.

3.3 Logische Äquivalenz

Definition 3.11 Zwei Formeln A und B heißen *logisch äquivalent*, falls $A \models B$ und $B \models A$. *logisch äquivalent*
Notation: $A \models B$.

Theorem 3.12 *Logische Äquivalenz $\models$ ist eine Äquivalenzrelation, d. h. es gilt:*

- (i) *Reflexivität:* $A \models A$.
- (ii) *Symmetrie:* Wenn $A \models B$, dann $B \models A$.
- (iii) *Transitivität:* Wenn $A \models B$ und $B \models C$, dann $A \models C$.

Beweis. Einfache Überlegungen zu Wahrheitstafeln.

QED

Beispiel. Es sei A die Formel $p \vee \top$, B die Formel $\top$, und C die Formel $p \vee \neg p$. Es gilt $p \vee \top \models \top$ und $p \vee \neg p \models \top$. Mit (ii) gilt $\top \models p \vee \neg p$, und mit (iii) auch $p \vee \top \models p \vee \neg p$.

Aufgrund von Theorem 3.12 können für schon gegebene logische Äquivalenzen weitere Äquivalenzen angegeben werden. Dabei werden immer ganze Formeln in Relation zueinander gesetzt. Um in einer gegebenen Formel auch beliebige Teilformeln durch zu diesen äquivalente Formeln ersetzen zu können, benötigen wir das folgende Theorem.

Theorem 3.13 (Ersetzungstheorem) *Es sei $A \models A'$, A eine Teilformel von B , und es sei B' das Resultat der Ersetzung von A durch A' in B . Dann gilt $B \models B'$.*

Beweis. Wegen $A \models A'$ ist die Hauptspalte der Wahrheitstafel für A identisch mit der Hauptspalte für A' . Durch Ersetzung von A durch A' in B ändert sich somit an der Hauptspalte von B' nichts. Also gilt $B \models B'$. QED

Beispiel. Sei A die Formel $p \vee q$ und B die Formel $(p \vee q) \wedge r$.

Da $\overbrace{p \vee q}^A \models \overbrace{q \vee p}^{A'}$, gilt auch $\underbrace{(p \vee q) \wedge r}_B \models \underbrace{(q \vee p) \wedge r}_{B'}$.

Wichtige Äquivalenzen und Folgerungen

Für beliebige Formeln A, B, C und D gilt:

Kommutativität von $\wedge$ und $\vee$

$$A \wedge B \models B \wedge A$$

$$A \vee B \models B \vee A$$

Distributivität von $\vee$ und $\wedge$

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

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

Verschmelzung

$$(A \wedge B) \vee A \models A$$

$$(A \vee B) \wedge A \models A$$

Extremalität von $\top$ und $\perp$

$$A \vee \top \models \top$$

$$A \wedge \perp \models \perp$$

Satz vom ausgeschlossenen Widerspruch

$$A \wedge \neg A \models \perp$$

$$\models \neg(A \wedge \neg A)$$

doppelte Negation

$$\neg\neg A \models A$$

Implikationsgesetze

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

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

$$A \rightarrow B \wedge C \models (A \rightarrow B) \wedge (A \rightarrow C)$$

$$A \rightarrow B \vee C \models (A \rightarrow B) \vee (A \rightarrow C)$$

modus (ponendo) ponens

$$A \rightarrow B, A \models B$$

Kontraposition

$$A \rightarrow B \models \neg B \rightarrow \neg A$$

praeclarum theorema

$$(A \rightarrow B) \wedge (D \rightarrow C) \models (A \wedge D) \rightarrow (B \wedge C)$$

Assoziativität von $\wedge$ und $\vee$

$$(A \wedge B) \wedge C \models A \wedge (B \wedge C)$$

$$(A \vee B) \vee C \models A \vee (B \vee C)$$

De Morgansche Gesetze

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

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

Idempotenz

$$A \wedge A \models A$$

$$A \vee A \models A$$

Neutralität von $\top$ und $\perp$

$$A \wedge \top \models A$$

$$A \vee \perp \models A$$

*Satz vom ausgeschlossenen Dritten
(tertium non datur)*

$$A \vee \neg A \models \top$$

$$\models A \vee \neg A$$

Transitivität von $\rightarrow$

$$A \rightarrow B, B \rightarrow C \models A \rightarrow C$$

weitere Implikationsgesetze

$$A \rightarrow \perp \models \neg A$$

$$\top \rightarrow A \models A$$

$$\neg A \rightarrow \perp \models A$$

$$A \rightarrow B \models \neg A \vee B$$

modus (tollendo) tollens

$$A \rightarrow B, \neg B \models \neg A$$

Importation (" $\models$ "), Exportation (" $\models$ ")

$$A \rightarrow (B \rightarrow C) \models A \wedge B \rightarrow C$$

Biimplikationsgesetz

$$A \leftrightarrow B \models (A \rightarrow B) \wedge (B \rightarrow A)$$

(Zur Bedeutung der Bezeichnungen *modus (ponendo) ponens* und *modus (tollendo) tollens* siehe J. Mittelstraß (Hrsg.), *Enzyklopädie Philosophie und Wissenschaftstheorie*, J. B. Metzler, 2005. Das *praeclarum theorema* geht auf Leibniz zurück.)

Unter Verwendung von Theorem 3.12, Ersetzungstheorem und schon gezeigter logischer Äquivalenzen können durch *Äquivalenzumformungen* weitere logische Äquivalenzen gezeigt werden.

Äquivalenzumformung

Beispiel. Wir zeigen durch Äquivalenzumformung $p \vee q \rightarrow r \models (p \rightarrow r) \wedge (q \rightarrow r)$:

$$\begin{aligned}
 p \vee q \rightarrow r &\models \neg(p \vee q) \vee r && (A \rightarrow B \models \neg A \vee B) \\
 &\models (\neg p \wedge \neg q) \vee r && (\text{De Morgan}) \\
 &\models (\neg p \vee r) \wedge (\neg q \vee r) && (\text{Distributivität von } \vee) \\
 &\models (p \rightarrow r) \wedge (\neg q \vee r) && (A \rightarrow B \models \neg A \vee B) \\
 &\models (p \rightarrow r) \wedge (q \rightarrow r) && (A \rightarrow B \models \neg A \vee B)
 \end{aligned}$$

Im 2., 4. und 5. Umformungsschritt wird bei der Verwendung der im Kommentar vermerkten Äquivalenzen auch das Ersetzungstheorem verwendet. Zum Beispiel wird im 2. Schritt zur Umformung der Teilformel $\neg(p \vee q)$ die Äquivalenz

$$\underbrace{\neg(p \vee q)}_A \models \underbrace{\neg p \wedge \neg q}_{A'}$$

verwendet (dies ist eine Instanz des De Morganschen Gesetzes $\neg(C \vee D) \models \neg C \wedge \neg D$), um aufgrund des Ersetzungstheorems die Äquivalenz

$$\underbrace{\underbrace{\neg(p \vee q)}_A \vee r}_B \models \underbrace{(\underbrace{\neg p \wedge \neg q}_{A'}) \vee r}_{B'}$$

zu erhalten.

Im 1. und 3. Schritt wird das Ersetzungstheorem nicht benötigt, da die im Kommentar vermerkten Äquivalenzen hier nicht auf echte Teilformeln angewendet werden, sondern auf ganze Formeln. Es ist

$$\underbrace{p \vee q \rightarrow r}_A \models \underbrace{\neg(p \vee q) \vee r}_B$$

eine Instanz von $A \rightarrow B \models \neg A \vee B$, und

$$\underbrace{(\neg p \wedge \neg q) \vee r}_A \models \underbrace{(\neg p \vee r) \wedge (\neg q \vee r)}_B$$

ist eine Instanz von $(A \wedge B) \vee C \models (A \vee C) \wedge (B \vee C)$. Die beiden Schritte können damit unter impliziter Verwendung von Reflexivität (nur im 1. Schritt) und Transitivität (im 1. und 3. Schritt) ausgeführt werden. Alternativ kann auch hier das Ersetzungstheorem verwendet werden, da es den Fall einschließt, dass die im Theorem genannte Teilformel A von B identisch mit B ist.

4 Normalformen und funktionale Vollständigkeit

Bisher haben wir uns auf wenige logische Konstanten beschränkt:

- 0-stellige: $\top, \perp$,
- 1-stellige: $\neg$,
- 2-stellige: $\wedge, \vee, \rightarrow, \leftrightarrow$.

Definition 4.1 Die *Stelligkeit* einer logischen Konstante ist die Anzahl der unmittelbaren Teilformeln, die durch die logische Konstante miteinander verbunden werden. (Im Fall natürlichsprachlicher Konnektive ist dies die Anzahl der unmittelbaren Teilaussagen.)

Stelligkeit

Bemerkung. Für logische Konstanten haben wir *Infixnotation* verwendet, und z. B. $(A \wedge B)$ oder $\neg A$ geschrieben. Stattdessen kann auch *Präfixnotation* verwendet werden, bei der man z. B. $\wedge(A, B)$ statt $(A \wedge B)$ und $\neg(A)$ statt $\neg A$ schreibt. Logische Konstanten K beliebiger Stelligkeit n können in Präfixnotation als $K(A_1, \dots, A_n)$ geschrieben werden.

Infixnotation
Präfixnotation

Wir können die Sprache der Aussagenlogik durch Hinzunahme von logischen Konstanten erweitern, deren Semantik wir z. B. durch Wahrheitstabeln festlegen.

Beispiele. (i) Die 2-stellige logische Konstante der *ausschließenden Disjunktion* (auch: *exklusives Oder, XOR*) ist definiert durch

A	B	$A \oplus B$
0	0	0
0	1	1
1	0	1
1	1	0

Sie entspricht dem natürlichsprachlichen “entweder ... oder ...”.

(ii) Das durch

A_1	A_2	A_3	$K^*(A_1, A_2, A_3)$
0	0	0	1
0	0	1	1
0	1	0	1
0	1	1	0
1	0	0	1
1	0	1	0
1	1	0	0
1	1	1	0

definierte 3-stellige Konnektiv K^* liefert den Wahrheitswert 1 genau dann, wenn mindestens zwei der Teilformeln A_1, A_2, A_3 falsch sind.

Auch können bestimmte logische Konstanten mittels anderer logischer Konstanten ausgedrückt werden; zum Beispiel kann die Biimplikation $\leftrightarrow$ aufgrund der Äquivalenz

$$A \leftrightarrow B \models (A \rightarrow B) \wedge (B \rightarrow A)$$

mittels $\wedge$ und $\rightarrow$ ausgedrückt werden. Des Weiteren können wir feststellen, dass z. B. das 2-stellige natürlichsprachliche Konnektiv “weder . . . noch . . .”, dem keine der bisher behandelten logischen Konstanten direkt entspricht, mittels $\neg$ und $\wedge$ durch $\neg A \wedge \neg B$ ausgedrückt werden kann.

Diese Beobachtungen führen zu folgenden Fragestellungen:

- (1) Welche weiteren logischen Konstanten können hinzugenommen werden?
- (2) Können prinzipiell alle wahrheitsfunktionalen Konnektive (beliebiger Stelligkeit) im Rahmen der Aussagenlogik behandelt werden?
- (3) Mit welchen logischen Konstanten können welche anderen logischen Konstanten ausgedrückt werden?
- (4) Gibt es überschaubare Mengen logischer Konstanten, mit denen alle logischen Konstanten ausgedrückt werden können?

Zur Untersuchung dieser Fragestellungen behandeln wir zunächst disjunktive und konjunktive Normalformen, die auch für Anwendungen wie z. B. Resolutionverfahren von großer Bedeutung sind.

4.1 Normalformen

Unter einer Konjunktion bzw. einer Disjunktion verstehen wir im Folgenden verallgemeinerte Konjunktionen bzw. Disjunktionen für $n \geq 0$ Konjunktions- bzw. Disjunktionsglieder.

Definition 4.2 (i) Ein *Literal* ist ein Aussagesymbol oder dessen Negation. Ist ein Literal ein Aussagesymbol, heißt es auch *positives Literal*; ist es ein negiertes Aussagesymbol, heißt es auch *negatives Literal*. Literale A und $\neg A$ sind zueinander *komplementär*.

Literal

(ii) Eine *Elementarkonjunktion* ist eine Konjunktion von Literalen.

(iii) Eine *Elementardisjunktion* ist eine Disjunktion von Literalen.

(iv) Eine *disjunktive Normalform* (kurz: *DNF*) ist eine Disjunktion von Elementarkonjunktionen.

disjunktive Normalform

(v) Eine *konjunktive Normalform* (kurz: *KNF*) ist eine Konjunktion von Elementardisjunktionen.

konjunktive Normalform

Bemerkungen. (i) Für Literale L_i (für $1 \leq i \leq k$) haben Elementarkonjunktionen die Form $L_1 \wedge \dots \wedge L_k$. Als Grenzfälle lassen wir für $k = 1$ das Literal L_1 , und für $k = 0$ die leere Konjunktion (Notation: $\top$), als Elementarkonjunktionen zu.

(ii) Bei Elementardisjunktionen der Form $L_1 \vee \dots \vee L_k$ lassen wir als Grenzfälle für $k = 1$ das Literal L_1 , und für $k = 0$ die leere Disjunktion (Notation: $\perp$), als Elementardisjunktionen zu.

(iii) Für die leere Konjunktion kann auch $p \vee \neg p$, und für die leere Disjunktion $p \wedge \neg p$, geschrieben werden, für ein beliebiges Aussagesymbol p .

Beispiele. (i) p ist ein positives Literal, $\neg p$ ein negatives. Die beiden Literale sind zueinander komplementär.

(ii) $p \wedge p \wedge \neg r$ ist eine Elementarkonjunktion aus 3 Literalen.

- (iii) $p \vee p_1 \vee \neg r \vee r$ ist eine Elementardisjunktion aus 4 Literalen.
- (iv) $\neg q$ ist sowohl eine Elementarkonjunktion als auch eine Elementardisjunktion.
- (v) $(p \wedge p \wedge \neg r) \vee (p \wedge \neg p_1)$ ist eine DNF.
- (vi) $(p_3 \vee \neg q) \wedge (p \vee p_1 \vee \neg r \vee r)$ ist eine KNF.
- (vii) Sowohl $p \wedge \neg q$ als auch $p \vee \neg q$ sind zugleich DNF und KNF.

Für jede Formel A kann eine zu ihr logisch äquivalente disjunktive Normalform anhand der Wahrheitstafel für A konstruiert werden. Wir führen das *Verfahren zur Konstruktion einer disjunktiven Normalform* anhand des folgenden Beispiels ein.

Konstruktion einer DNF

Beispiel. Konstruktion einer disjunktiven Normalform zu $p \vee q \rightarrow r$. Wir betrachten die Wahrheitstafel für $p \vee q \rightarrow r$ (wobei wir hier zur Erläuterung die 8 relevanten Bewertungen mit $\mathcal{I}_1$ bis $\mathcal{I}_8$ benennen):

	p	q	r	$p \vee q \rightarrow r$	
$\mathcal{I}_1$:	0	0	0	0	1
$\mathcal{I}_2$:	0	0	1	0	1
$\mathcal{I}_3$:	0	1	0	1	0
$\mathcal{I}_4$:	0	1	1	1	1
$\mathcal{I}_5$:	1	0	0	1	0
$\mathcal{I}_6$:	1	0	1	1	1
$\mathcal{I}_7$:	1	1	0	1	0
$\mathcal{I}_8$:	1	1	1	1	1

Der Hauptspalte entnimmt man, dass die Formel $p \vee q \rightarrow r$ für die 5 Bewertungen $\mathcal{I}_1$, $\mathcal{I}_2$, $\mathcal{I}_4$, $\mathcal{I}_6$ und $\mathcal{I}_8$ den Wahrheitswert 1 hat.

Wir betrachten die Bewertung $\mathcal{I}_1$. Es ist $\llbracket p \rrbracket^{\mathcal{I}_1} = \llbracket q \rrbracket^{\mathcal{I}_1} = \llbracket r \rrbracket^{\mathcal{I}_1} = 0$. Wir verwenden nun statt der positiven Literale p, q, r die negativen Literale $\neg p, \neg q$ und $\neg r$. Für diese gilt $\llbracket \neg p \rrbracket^{\mathcal{I}_1} = \llbracket \neg q \rrbracket^{\mathcal{I}_1} = \llbracket \neg r \rrbracket^{\mathcal{I}_1} = 1$, und somit nach Definition der Konjunktion $\llbracket \neg p \wedge \neg q \wedge \neg r \rrbracket^{\mathcal{I}_1} = 1$. Da auch $\llbracket p \vee q \rightarrow r \rrbracket^{\mathcal{I}_1} = 1$, gilt $\llbracket \neg p \wedge \neg q \wedge \neg r \rrbracket^{\mathcal{I}_1} = \llbracket p \vee q \rightarrow r \rrbracket^{\mathcal{I}_1}$. Für die Bewertung $\mathcal{I}_1$ ist der Wahrheitswert von $p \vee q \rightarrow r$ also durch $\neg p \wedge \neg q \wedge \neg r$ bestimmt.

Die Formel $\neg p \wedge \neg q \wedge \neg r$ ist die erste Elementarkonjunktion der gesuchten disjunktiven Normalform; wir notieren sie in der Zeile von $\mathcal{I}_1$:

	p	q	r	$p \vee q \rightarrow r$		
$\mathcal{I}_1$:	0	0	0	0	1	$(\neg p \wedge \neg q \wedge \neg r)$
$\mathcal{I}_2$:	0	0	1	0	1	
$\mathcal{I}_3$:	0	1	0	1	0	
$\mathcal{I}_4$:	0	1	1	1	1	
$\mathcal{I}_5$:	1	0	0	1	0	
$\mathcal{I}_6$:	1	0	1	1	1	
$\mathcal{I}_7$:	1	1	0	1	0	
$\mathcal{I}_8$:	1	1	1	1	1	

Nun betrachten wir die Bewertung $\mathcal{I}_2$, für die $p \vee q \rightarrow r$ ebenfalls wahr ist. Nun ist $\llbracket p \rrbracket^{\mathcal{I}_2} = \llbracket q \rrbracket^{\mathcal{I}_2} = 0$ und $\llbracket r \rrbracket^{\mathcal{I}_2} = 1$. Statt der positiven Literale p und q verwenden wir deshalb die negativen Literale $\neg p$ und $\neg q$, so dass $\llbracket \neg p \rrbracket^{\mathcal{I}_2} = \llbracket \neg q \rrbracket^{\mathcal{I}_2} = \llbracket r \rrbracket^{\mathcal{I}_2} = 1$,

und somit $\llbracket \neg p \wedge \neg q \wedge r \rrbracket^{\mathcal{I}_2} = 1$. Da $\llbracket \neg p \wedge \neg q \wedge r \rrbracket^{\mathcal{I}_2} = \llbracket p \vee q \rightarrow r \rrbracket^{\mathcal{I}_2}$, ist für $\mathcal{I}_2$ der Wahrheitswert von $p \vee q \rightarrow r$ durch die Elementarkonjunktion $\neg p \wedge \neg q \wedge r$ bestimmt. Da $p \vee q \rightarrow r$ (zumindest) für $\mathcal{I}_1$ oder $\mathcal{I}_2$ wahr ist, fügen wir diese zweite Elementarkonjunktion mit einer *Disjunktion* in der Zeile von $\mathcal{I}_2$ hinzu; man erhält also die Formel $(\neg p \wedge \neg q \wedge r) \vee (\neg p \wedge \neg q \wedge r)$:

	p	q	r	$p \vee q \rightarrow r$		
$\mathcal{I}_1$:	0	0	0	0	1	$(\neg p \wedge \neg q \wedge \neg r)$
$\mathcal{I}_2$:	0	0	1	0	1	$\vee (\neg p \wedge \neg q \wedge r)$
$\mathcal{I}_3$:	0	1	0	1	0	
$\mathcal{I}_4$:	0	1	1	1	1	
$\mathcal{I}_5$:	1	0	0	1	0	
$\mathcal{I}_6$:	1	0	1	1	1	
$\mathcal{I}_7$:	1	1	0	1	0	
$\mathcal{I}_8$:	1	1	1	1	1	

Entsprechend erzeugt man für die restlichen Bewertungen $\mathcal{I}_4$, $\mathcal{I}_6$ und $\mathcal{I}_8$, für die $p \vee q \rightarrow r$ wahr ist, die Elementarkonjunktionen $(\neg p \wedge q \wedge r)$, $(p \wedge \neg q \wedge r)$ und $(p \wedge q \wedge r)$, die jeweils mit einer Disjunktion hinzugefügt werden:

	p	q	r	$p \vee q \rightarrow r$		} DNF	
$\mathcal{I}_1$:	0	0	0	0	1		$(\neg p \wedge \neg q \wedge \neg r)$
$\mathcal{I}_2$:	0	0	1	0	1		$\vee (\neg p \wedge \neg q \wedge r)$
$\mathcal{I}_3$:	0	1	0	1	0		
$\mathcal{I}_4$:	0	1	1	1	1		$\vee (\neg p \wedge q \wedge r)$
$\mathcal{I}_5$:	1	0	0	1	0		
$\mathcal{I}_6$:	1	0	1	1	1		$\vee (p \wedge \neg q \wedge r)$
$\mathcal{I}_7$:	1	1	0	1	0		
$\mathcal{I}_8$:	1	1	1	1	1		$\vee (p \wedge q \wedge r)$

Die resultierende Formel ist eine disjunktive Normalform. Nun gilt Folgendes:

- (i) Für jede der 5 Bewertungen $\mathcal{I}_1$, $\mathcal{I}_2$, $\mathcal{I}_4$, $\mathcal{I}_6$ und $\mathcal{I}_8$, für die $p \vee q \rightarrow r$ wahr ist, wurde eine Elementarkonjunktion angegeben, die den Wahrheitswert von $p \vee q \rightarrow r$ für die jeweilige Bewertung zu wahr bestimmt. Für jede dieser 5 Bewertungen ist somit (nach Definition der Disjunktion) auch die DNF wahr.
- (ii) Für jede der übrigen 3 Bewertungen $\mathcal{I}_3$, $\mathcal{I}_5$ und $\mathcal{I}_7$ hat die DNF den Wahrheitswert 0, da keine der in ihr enthaltenen Elementarkonjunktionen unter diesen Bewertungen wahr sein kann.

(Zum Beispiel kann für $\mathcal{I}_7$ wegen $\llbracket r \rrbracket^{\mathcal{I}_7} = 0$ keine der 4 Elementarkonjunktionen wahr sein, in der das positive Literal r vorkommt; und die erste Elementarkonjunktion kann aufgrund der negativen Literale $\neg p$ und $\neg q$ nicht wahr sein, da $\llbracket p \rrbracket^{\mathcal{I}_7} = \llbracket q \rrbracket^{\mathcal{I}_7} = 1$, also $\llbracket \neg p \rrbracket^{\mathcal{I}_7} = \llbracket \neg q \rrbracket^{\mathcal{I}_7} = 0$, und somit $\llbracket \neg p \wedge \neg q \wedge r \rrbracket^{\mathcal{I}_7} = 0$.)

Folglich stimmt der Wahrheitswert von $p \vee q \rightarrow r$ mit dem der DNF für jede Bewertung überein. Es gilt also

$$p \vee q \rightarrow r \models \underbrace{(\neg p \wedge \neg q \wedge r) \vee (\neg p \wedge \neg q \wedge r) \vee (\neg p \wedge q \wedge r) \vee (p \wedge \neg q \wedge r) \vee (p \wedge q \wedge r)}_{\text{DNF zu } p \vee q \rightarrow r}$$

Man macht sich leicht klar, dass mit dem im Beispiel illustrierten Verfahren zu jeder Formel eine zu ihr logisch äquivalente disjunktive Normalform konstruiert werden kann. Damit haben wir schon eine Antwort zu Frage (3): Die bisher betrachteten logischen Konstanten einschließlich $\rightarrow$, $\leftrightarrow$, $\oplus$ und K^* können allein mit $\top$, $\perp$, $\neg$, $\wedge$ und $\vee$ ausgedrückt werden.

Verum $\top$ und Falsum $\perp$ werden nur für die Grenzfälle leerer Konjunktionen, beziehungsweise leerer Disjunktionen benötigt. Sofern man nicht verlangt, dass in der DNF zu einer Formel A ausschließlich die in A vorkommenden Aussagesymbole vorkommen, kann auf $\top$ und $\perp$ verzichtet werden, indem man $p \vee \neg p$ statt $\top$ und $p \wedge \neg p$ statt $\perp$ schreibt. Das Aussagesymbol p ist hierbei beliebig, muss also in der Formel A nicht vorkommen.

Beispiele. (i) Im Grenzfall nicht vorhandener Aussagesymbole erhält man z. B. für $\perp \rightarrow \top$ anhand von

$\perp \rightarrow \top$	
0	1

die leere Konjunktion $\top$ als DNF sowie als KNF. Statt $\top$ kann die Formel $p \vee \neg p$ geschrieben werden, die ebenfalls sowohl eine DNF als auch eine KNF zu $\perp \rightarrow \top$ ist.

(ii) Für $\top \rightarrow \perp$ erhält man anhand

$\top \rightarrow \perp$	
1	0

die leere Disjunktion $\perp$ als DNF sowie als KNF. Auch $p \wedge \neg p$ ist zugleich eine DNF und KNF zu $\top \rightarrow \perp$.

Damit kann eine weitere Antwort auf Frage (3) gegeben werden: Die bisher betrachteten logischen Konstanten können allein mit $\neg$, $\wedge$ und $\vee$ ausgedrückt werden, sofern ein beliebiges Aussagesymbol zur Verfügung steht.

Neben einer DNF kann immer auch eine KNF zu einer gegebenen Formel konstruiert werden.

Theorem 4.3 *Zu jeder Formel A kann sowohl*

- (i) *eine logisch äquivalente disjunktive Normalform als auch*
- (ii) *eine logisch äquivalente konjunktive Normalform angegeben werden.*

Beweis. (i) Die anhand der Wahrheitstafel für A konstruierte disjunktive Normalform A' hat dieselbe Hauptspalte wie A . Also gilt: $A \models A'$.

(ii) Eine konjunktive Normalform zu A kann wie folgt konstruiert werden:

- (1) Zunächst bildet man eine disjunktive Normalform zu $\neg A$. Es gilt:

$$\neg A \models \underbrace{(L_{11} \wedge \dots \wedge L_{1m_1}) \vee \dots \vee (L_{j1} \wedge \dots \wedge L_{jn_j})}_{\text{DNF zu } \neg A}$$

mit Literalen L_{ij} .

- (2) Nun negiert man beide Seiten der logischen Äquivalenz, und formt die rechte Seite in eine konjunktive Normalform um.

Durch Negation auf beiden Seiten erhält man

$$\neg\neg A \models \neg((L_{11} \wedge \dots \wedge L_{1n_1}) \vee \dots \vee (L_{j1} \wedge \dots \wedge L_{jn_j}))$$

Nach Anwendung von $\neg\neg A \models A$ auf der linken Seite erhält man

$$A \models \neg((L_{11} \wedge \dots \wedge L_{1n_1}) \vee \dots \vee (L_{j1} \wedge \dots \wedge L_{jn_j}))$$

Die rechte Seite wird mit De Morgan umgeformt; man erhält

$$A \models \neg(L_{11} \wedge \dots \wedge L_{1n_1}) \wedge \dots \wedge \neg(L_{j1} \wedge \dots \wedge L_{jn_j})$$

und nach weiterer Umformung mit De Morgan

$$A \models (\neg L_{11} \vee \dots \vee \neg L_{1n_1}) \wedge \dots \wedge (\neg L_{j1} \vee \dots \vee \neg L_{jn_j})$$

Nach Beseitigung doppelter Negationen bei negierten negativen Literalen L_{ij} liegt dann auf der rechten Seite die konjunktive Normalform zu A vor. (Hierbei wird auch das Ersetzungstheorem verwendet.) QED

Beispiel. Konstruktion einer konjunktiven Normalform zu $p \vee q \rightarrow r$:

- (1) Disjunktive Normalform zu $\neg(p \vee q \rightarrow r)$ konstruieren:

p	q	r	$\neg(p \vee q \rightarrow r)$	
0	0	0	0	
0	0	1	0	
0	1	0	1	$(\neg p \wedge q \wedge \neg r)$
0	1	1	0	
1	0	0	1	$\vee (p \wedge \neg q \wedge \neg r)$
1	0	1	0	
1	1	0	1	$\vee (p \wedge q \wedge \neg r)$
1	1	1	0	

Es ist $\neg(p \vee q \rightarrow r) \models \underbrace{(\neg p \wedge q \wedge \neg r) \vee (p \wedge \neg q \wedge \neg r) \vee (p \wedge q \wedge \neg r)}_{\text{DNF zu } \neg(p \vee q \rightarrow r)}$.

- (2) Negation von $\neg(p \vee q \rightarrow r)$ auf der linken Seite, und Negation der DNF zu $\neg(p \vee q \rightarrow r)$ auf der rechten Seite:

$$\neg\neg(p \vee q \rightarrow r) \models \neg((\neg p \wedge q \wedge \neg r) \vee (p \wedge \neg q \wedge \neg r) \vee (p \wedge q \wedge \neg r))$$

Beseitigung der doppelten Negation auf der linken Seite, und Umformung der rechten Seite in eine konjunktive Normalform durch Verwendung von De Morgan und Beseitigung doppelter Negationen:

$$\begin{aligned} p \vee q \rightarrow r &\models \neg((\neg p \wedge q \wedge \neg r) \vee (p \wedge \neg q \wedge \neg r) \vee (p \wedge q \wedge \neg r)) \\ &\models \neg(\neg p \wedge q \wedge \neg r) \wedge \neg(p \wedge \neg q \wedge \neg r) \wedge \neg(p \wedge q \wedge \neg r) \\ &\models (\neg\neg p \vee \neg q \vee \neg\neg r) \wedge (\neg p \vee \neg\neg q \vee \neg\neg r) \wedge (\neg p \vee \neg q \vee \neg\neg r) \\ &\models \underbrace{(p \vee \neg q \vee r) \wedge (\neg p \vee q \vee r) \wedge (\neg p \vee \neg q \vee r)}_{\text{KNF zu } p \vee q \rightarrow r} \end{aligned}$$

Bemerkung. Als leichte Übungsaufgabe überlegt man sich, wie eine KNF zu einer Formel A direkt aus der Wahrheitstafel für A konstruiert werden kann.

Alternativ kann eine KNF zu einer Formel A auch ohne die Verwendung einer Wahrheitstafel erzeugt werden, indem man Äquivalenzumformungen verwendet.

Definition 4.4 Das folgende *Umformungsverfahren zur Erzeugung einer KNF* beruht auf logischen Äquivalenzen und terminiert, sobald eine zur Ausgangsformel logisch äquivalente KNF vorliegt.

Umformungsverfahren zur Erzeugung einer KNF

(1) Eliminiere Vorkommen von $\top$ und $\perp$ durch folgende Umformungen:

$$\top \rightsquigarrow p \vee \neg p$$

$$\perp \rightsquigarrow p \wedge \neg p$$

(2) Eliminiere Vorkommen von $\leftrightarrow$:

$$(A \leftrightarrow B) \rightsquigarrow (A \rightarrow B) \wedge (B \rightarrow A)$$

(3) Eliminiere Vorkommen von $\rightarrow$:

$$(A \rightarrow B) \rightsquigarrow (\neg A \vee B)$$

(4) Ziehe Negationen nach innen (De Morgan) und beseitige doppelte Negationen:

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

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

$$\neg\neg A \rightsquigarrow A$$

(5) Ziehe $\wedge$ nach außen mit Hilfe der Distributivität:

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

Assoziativität von $\wedge$ und $\vee$ verwenden wir implizit.

4.2 Funktionale Vollständigkeit

Die bisher behandelten Konnektive sind maximal 3-stellig. Nun betrachten wir beliebige, n -stellige Konnektive. Ein n -stelliges Konnektiv K notieren wir in Präfixnotation als $K(A_1, \dots, A_n)$. Die Definition von *Formeln* kann damit wie folgt erweitert werden:

Formeln

Definition 4.5 Sind $A_1, \dots, A_n$ Formeln, und ist K ein n -stelliges Konnektiv, dann ist auch $K(A_1, \dots, A_n)$ eine Formel, deren Hauptkonnektiv K ist.

Wie schon bei 0-, 1-, 2- und 3-stelligen Konnektiven, kann auch die Bedeutung beliebiger n -stelliger Konnektive $K(A_1, \dots, A_n)$ durch eine Wahrheitstafel festgelegt werden. Da der

Wahrheitswert eines n -stelligen Konnektivs $K(A_1, \dots, A_n)$ von den n Wahrheitswerten von $A_1, \dots, A_n$ abhängt, hat eine solche Wahrheitstafel $n + 1$ Spalten mit 2^n Zeilen:

A_1	A_2	$\dots$	A_n	$K(A_1, \dots, A_n)$
0	0	$\dots$	0	x_1
$\vdots$	$\vdots$		1	x_2
$\vdots$	0	$\dots$	$\vdots$	$\vdots$
$\vdots$	1		$\vdots$	$\vdots$
0	1	$\vdots$	$\vdots$	$x_{2^n/2}$
1	0	$\vdots$	$\vdots$	$x_{2^n/2+1}$
$\vdots$	$\vdots$		$\vdots$	$\vdots$
$\vdots$	0	$\dots$	$\vdots$	$\vdots$
$\vdots$	1		$\vdots$	$\vdots$
$\vdots$	$\vdots$		0	x_{2^n-1}
1	1	$\dots$	1	x_{2^n}

Hierbei stehen die Variablen $x_1, \dots, x_{2^n}$ jeweils für einen der beiden Wahrheitswerte 0 und 1. (Die Anzahl aller n -stelligen Konnektive beträgt damit 2^{2^n} .)

Beispiel. Die Bedeutung des 3-stelligen Konnektivs K' sei durch folgende Wahrheitstafel festgelegt:

A_1	A_2	A_3	$K'(A_1, A_2, A_3)$
0	0	0	0
0	0	1	1
0	1	0	0
0	1	1	0
1	0	0	0
1	0	1	0
1	1	0	0
1	1	1	1

Mit dem Verfahren zur Konstruktion einer DNF kann K' durch die Formel

$$(\neg A_1 \wedge \neg A_2 \wedge A_3) \vee (A_1 \wedge A_2 \wedge A_3)$$

ausgedrückt werden, da $K'(A_1, A_2, A_3) \models (\neg A_1 \wedge \neg A_2 \wedge A_3) \vee (A_1 \wedge A_2 \wedge A_3)$ gilt. (Die Formel auf der rechten Seite hat lediglich die *Form* einer DNF, da A_1 , A_2 und A_3 nicht nur für Aussagesymbole, sondern für beliebige Formeln stehen. Bei den beiden Disjunktionsgliedern handelt es sich also nicht unbedingt um Elementarkonjunktionen, was bei einer DNF verlangt wird.)

Definition 4.6 Eine Menge von Konnektiven heißt *wahrheitsfunktional vollständig* (kurz: *funktional vollständig*), falls jedes Konnektiv durch die in der Menge enthaltenen Konnektive ausgedrückt werden kann. *wahrheitsfunktional vollständig*

Theorem 4.7 Die Menge $\{\top, \perp, \neg, \wedge, \vee\}$ ist funktional vollständig.

Beweis. Mit den Verfahren zur Konstruktion einer DNF oder KNF kann zu jeder Formel $K(A_1, \dots, A_n)$ mit Hauptkonnektiv K eine zu ihr logisch äquivalente Formel angegeben werden, welche die Form einer DNF oder KNF hat, d. h. in der nur die

Konnektive $\neg, \wedge, \vee$ und als Grenzfälle $\top, \perp$ vorkommen. Somit kann jedes Konnektiv K mit diesen Konnektiven ausgedrückt werden. QED

Theorem 4.8 *Die folgenden Mengen sind funktional vollständig:*

- (i) $\{\neg, \wedge\}$, (ii) $\{\neg, \vee\}$, (iii) $\{\neg, \rightarrow\}$, (iv) $\{\perp, \rightarrow\}$.

Beweis. (i) Es gilt $A \vee B \models \neg(\neg A \wedge \neg B)$, $\top \models \neg(A \wedge \neg A)$ und $\perp \models A \wedge \neg A$. Damit folgt die Behauptung aus Theorem 4.7.

(ii) Es gilt $A \wedge B \models \neg(\neg A \vee \neg B)$, $\top \models A \vee \neg A$ und $\perp \models \neg(A \vee \neg A)$. Damit folgt die Behauptung aus Theorem 4.7.

(iii) Es gilt $A \vee B \models \neg A \rightarrow B$, $A \wedge B \models \neg(A \rightarrow \neg B)$, $\top \models A \rightarrow A$ und $\perp \models \neg(A \rightarrow A)$. Damit folgt die Behauptung aus Theorem 4.7.

(iv) Es gilt $\neg A \models A \rightarrow \perp$. Damit folgt die Behauptung unter Verwendung von (iii). QED

Bemerkungen. (i) In den Fällen (i)-(iii) ist zu beachten, dass bei $\top$ und $\perp$ zumindest ein beliebiges Aussagesymbol zur Verfügung stehen muss, um die beiden Konnektive auszudrücken.

(ii) Andernfalls können $\top$ und $\perp$ nicht durch $\{\neg, \wedge\}$, $\{\neg, \vee\}$ oder $\{\neg, \rightarrow\}$ ausgedrückt werden. Um funktional vollständige Mengen zu erhalten, muss wenigstens eines der beiden Konnektive $\top$ und $\perp$ hinzugenommen werden; das jeweils andere kann dann aufgrund $\top \models \neg \perp$ bzw. $\perp \models \neg \top$ ausgedrückt werden.

(iii) Wir gehen im Folgenden davon aus, dass immer auch ein beliebiges Aussagesymbol verwendet werden kann, um ein Konnektiv durch andere auszudrücken.

Definition 4.9 Die *Negatkonjunktion* $\downarrow$ (auch: *Peircescher Pfeil*) ist durch folgende Wahrheitstafel definiert:

Negatkonjunktion

A	B	$A \downarrow B$
0	0	1
0	1	0
1	0	0
1	1	0

Bemerkungen. (i) Da $A \downarrow B \models \neg(A \vee B)$, wird $\downarrow$ auch als NOR (“not or”) bezeichnet. Die Bezeichnung als Negatkonjunktion rührt von $A \downarrow B \models \neg A \wedge \neg B$ (“Konjunktion der Negate”) her.

(ii) Die Negatkonjunktion entspricht dem natürlichsprachlichen Konnektiv “weder . . . noch . . .” bzw. “nicht eines von beiden”.

Theorem 4.10 *Die Menge $\{\downarrow\}$ ist funktional vollständig.*

Beweis. Es gilt:

(i) $\neg A \models A \downarrow A$,

(ii) $A \wedge B \models \neg A \downarrow \neg B \models (A \downarrow A) \downarrow (B \downarrow B)$,

(iii) $A \vee B \models \neg(A \downarrow B) \models (A \downarrow B) \downarrow (A \downarrow B)$.

Damit folgt die Behauptung aus Theorem 4.8(i) bzw. (ii) für die funktional vollständigen Mengen $\{\neg, \wedge\}$ bzw. $\{\neg, \vee\}$. QED

Bemerkung. Für $\top$ und $\perp$ gilt dies nur mit der in der Bemerkung nach Theorem 4.8 genannten Einschränkung.

Beispiel. Darstellung des Beispielkonnektivs $K'(A_1, A_2, A_3)$ durch $\downarrow$:

$$\begin{aligned}
K'(A_1, A_2, A_3) &\models (A_1 \wedge A_2 \wedge A_3) \vee (\neg A_1 \wedge \neg A_2 \wedge A_3) \\
&\models ((\neg A_1 \downarrow \neg A_2) \wedge A_3) \vee (\neg A_1 \wedge \neg A_2 \wedge A_3) \\
&\models (((A_1 \downarrow A_1) \downarrow (A_2 \downarrow A_2)) \wedge A_3) \vee (\neg A_1 \wedge \neg A_2 \wedge A_3) \\
&\models \underbrace{(((A_1 \downarrow A_1) \downarrow (A_2 \downarrow A_2)) \downarrow ((A_1 \downarrow A_1) \downarrow (A_2 \downarrow A_2))) \downarrow (A_3 \downarrow A_3)}_B \vee (\neg A_1 \wedge \neg A_2 \wedge A_3) \\
&\models B \vee ((\neg \neg A_1 \downarrow \neg \neg A_2) \wedge A_3) \\
&\models B \vee ((A_1 \downarrow A_2) \wedge A_3) \\
&\models B \vee (\neg(A_1 \downarrow A_2) \downarrow \neg A_3) \\
&\models B \vee \underbrace{(((A_1 \downarrow A_2) \downarrow (A_1 \downarrow A_2)) \downarrow (A_3 \downarrow A_3))}_C \\
&\models \neg(B \downarrow C) \\
&\models (B \downarrow C) \downarrow (B \downarrow C)
\end{aligned}$$

Bemerkungen. (i) Man beachte, dass $\downarrow$ kein assoziatives Konnektiv ist, das heißt i. A. gilt *nicht* $(A_1 \downarrow A_2) \downarrow A_3 \models A_1 \downarrow (A_2 \downarrow A_3)$.

(ii) Neben $\downarrow$ gibt es mit der Negatdisjunktion $|$ genau ein weiteres 2-stelliges Konnektiv, mit dem jedes Konnektiv ausgedrückt werden kann. Die *Negatdisjunktion* (auch: *Shefferscher Strich*) ist durch folgende Wahrheitstafel definiert: *Negatdisjunktion*

A	B	$A B$
0	0	1
0	1	1
1	0	1
1	1	0

Da $A | B \models \neg(A \wedge B)$, wird $|$ auch als NAND (“*not and*”) bezeichnet. Natürlichsprachlich kann $|$ als “höchstens eines von beiden” wiedergegeben werden.

Zu den eingangs aufgeführten Fragestellungen können wir nun Folgendes sagen (wobei wir unter logischen Konstanten allein die *wahrheitsfunktionalen* verstehen wollen):

(1) Welche logischen Konstanten können zu den bisher betrachteten hinzugenommen werden?

Es können beliebige n -stellige wahrheitsfunktionale Konnektive hinzugenommen werden. Dadurch erhöht sich die Ausdruckstärke der formalen Sprache jedoch nur scheinbar. Denn da die Menge der bisher betrachteten logischen Konstanten schon wahrheitsfunktional vollständig ist, kann durch weitere hinzugenommene Konnektive nichts ausgedrückt werden, was nicht schon durch die in der funktional vollständigen Menge enthaltenen Konnektive ausgedrückt werden könnte.

- (2) Können prinzipiell alle wahrheitsfunktionalen Konnektive (beliebiger Stelligkeit) im Rahmen der Aussagenlogik behandelt werden?

Ja. Zwar gibt es unendlich viele wahrheitsfunktionale Konnektive, die jedoch alle durch endliche, wahrheitsfunktional vollständige Mengen ausgedrückt werden können.

- (3) Mit welchen logischen Konstanten können welche anderen logischen Konstanten ausgedrückt werden?

Wir haben gesehen, dass einzelne logische Konstanten durch eine andere Konstante (z. B. $\top$ durch $\rightarrow$) oder durch mehrere andere Konstanten ausgedrückt werden können. Anhand von Normalformen und weiterer Überlegungen konnten wir feststellen, dass es verschiedene wahrheitsfunktional vollständige Mengen gibt.

- (4) Gibt es überschaubare Mengen logischer Konstanten, mit denen alle logischen Konstanten ausgedrückt werden können?

Ja. Ein Beispiel ist die funktional vollständige Menge $\{\top, \neg, \wedge, \vee\}$. Weitere Beispiele derartiger Mengen sind $\{\neg, \wedge, \vee\}$ sowie $\{\downarrow\}$ und $\{\mid\}$, wobei bei diesen zusätzlich die Verwendung eines beliebigen Aussagesymbols zugelassen sein muss.

5 Resolutionsverfahren

Die bisher behandelten Begriffe wie z. B. Erfüllbarkeit, Allgemeingültigkeit und logische Folgerung beruhen auf einer zweiwertigen, wahrheitsfunktionalen Semantik. Natürlichsprachliche Argumente können im Rahmen von Formalisierungen als Folgerungsbehauptungen ausgedrückt werden, die dann z. B. unter Verwendung des Wahrheitstafelverfahrens auf Gültigkeit überprüft werden können.

Dieses Vorgehen vernachlässigt allerdings wesentliche Aspekte von Argumenten, wie z. B. die Verwendung zusätzlicher Annahmen während einer Argumentation, oder das schrittweise Schließen von Aussage zu Aussage. Diese Aspekte werden in Kalkülen berücksichtigt, in denen logisches Schließen rein syntaktisch erfolgt. Durch Kalküle werden die für die Logik zentralen Begriffe der *Ableitbarkeit* und *Beweisbarkeit* bereitgestellt. Wesentlicher Untersuchungsgegenstand der Logik ist dann das Verhältnis dieser kalkülbasierten Begriffe zu den semantischen Begriffen der logischen Folgerung und Allgemeingültigkeit: Ist ein Kalkül *korrekt* in Bezug auf die gegebene Semantik, d. h., ist jede im Kalkül beweisbare Aussage allgemeingültig? Ist der Kalkül auch *vollständig*, d. h., ist jede allgemeingültige Aussage im Kalkül beweisbar?

Wir führen zunächst den Resolutionskalkül ein, der aus lediglich einem Axiomenschema und einer Regel besteht. Weder in der Regel noch im Axiomenschema kommen logische Konstanten vor. Es werden ausschließlich sogenannte Klauseln als besonders einfache syntaktische Objekte verwendet. Die Anwendung des Resolutionskalküls auf beliebige Formeln erfordert deshalb eine vorgeschaltete Übersetzung von Formeln in Klauseln. Anschließend behandeln wir Resolutionswiderlegungen. Die Grundidee besteht darin, die Allgemeingültigkeit einer Formel A nachzuweisen, indem mittels Resolution gezeigt wird, dass $\neg A$ unerfüllbar ist. Für Resolutionswiderlegungen besteht der Resolutionskalkül aus nur einer Regel, und ist deshalb sehr einfach zu implementieren. Das Verfahren setzt jedoch Klauselmengen voraus, was für Formeln bedeutet, dass diese erst in konjunktive Normalform überführt werden müssen; der KNF entspricht dann unmittelbar die benötigte Klauselmenge.

Die vorgestellten Resolutionsverfahren sind korrekt und vollständig für die Semantik der Aussagenlogik.

Literatur

- M. Ben-Ari, *Mathematical Logic for Computer Science. Third Edition*, Springer, 2012, § 4.
- J. H. Gallier, *Logic for Computer Science. Foundations of Automatic Theorem Proving*, 2nd edition, Dover Publications, 2015, § 4.
- A. Leitsch, *The Resolution Calculus*, Springer, 1997, § 2.5.
- U. Schöning, *Logik für Informatiker*, 5. Auflage, Spektrum Akademischer Verlag, 2000, § 1.5.

Definition 5.1 Es seien $A_1, \dots, A_n, B_1, \dots, B_m$ Aussagesymbole.

- (i) Eine *Klausel* ist eine Disjunktion

Klausel

$$\neg A_1 \vee \dots \vee \neg A_n \vee B_1 \vee \dots \vee B_m$$

von Literalen.

Eine Klausel wird auch als Menge $\{\neg A_1, \dots, \neg A_n, B_1, \dots, B_m\}$ von Literalen aufgefasst.

- (ii) Wir notieren Klauseln auch als *Sequenzen* der Form

Sequenz

$$A_1, \dots, A_n \vdash B_1, \dots, B_m$$

- (iii) Links vom *Sequenzenzeichen* $\vdash$ steht das *Antezedens*, rechts davon das *Sukzedens*. Antezedens und Sukzedens werden als Mengen aufgefasst; d. h., Multiplizität und Anordnung der Listenelemente spielen keine Rolle.

- (iv) Die *leere Klausel* enthält keine Aussagesymbole. Notation: $\vdash$ bzw. $\square$.

Die leere Klausel ist unerfüllbar.

Bemerkungen. (i) Metasprachliche Variablen für endliche Mengen von Aussagesymbolen sind X, Y, Z , ggf. mit Indizes; für Klauseln: $S, S_1, S_2, \dots$; und für Mengen von Klauseln: $\Gamma, \Delta, \dots$

- (ii) Da Sequenzen Formeln entsprechen, können wir die für Formeln eingeführte logische Folgerung auch für Klauseln notieren: $\Gamma \models S$.

- (iii) $X, A \vdash Y, B$ steht für $X \cup \{A\} \vdash Y \cup \{B\}$, wobei A und B hier auch Element von X bzw. Y sein können.

5.1 Resolutionskalkül

Definition 5.2 Für Aussagesymbole A und möglicherweise leere Mengen von Aussagesymbolen $X, Y, \dots$ ist der *aussagenlogische Resolutionskalkül* gegeben durch das Axiom (bzw. Axiomenschema)

Resolutionskalkül

$$(Ax) \frac{}{X, A \vdash A, Y}$$

und die *aussagenlogische Resolutionsregel*:

Resolutionsregel

$$\frac{X_1 \vdash Y_1, A \quad A, X_2 \vdash Y_2}{X_1, X_2 \vdash Y_1, Y_2} (\mathcal{R})$$

Die Konklusion von $(\mathcal{R})$ heißt *Resolvente* (bzgl. A) der Prämissen.

Resolvente

Definition 5.3 Eine *Ableitung* im aussagenlogischen Resolutionskalkül ist ein (nach oben verzweigender) Baum, der induktiv wie folgt definiert ist (wobei wir $\mathcal{D}$ für die mit der Klausel S endende Ableitung $\mathcal{D}$ schreiben):

Ableitung

- (i) $X \vdash Y$ ist eine Ableitung (von $X \vdash Y$ aus $X \vdash Y$).

- (ii) $(Ax) \frac{}{X, A \vdash A, Y}$ ist eine Ableitung.

- (iii) Sind $\mathcal{D}_1$ $X_1 \vdash Y_1, A$ und $\mathcal{D}_2$ $A, X_2 \vdash Y_2$ Ableitungen, dann ist auch

$$\frac{\mathcal{D}_1 \quad \mathcal{D}_2}{X_1 \vdash Y_1, A \quad A, X_2 \vdash Y_2} (\mathcal{R})$$

eine Ableitung.

Definition 5.4 Die Menge der *Annahmen* (oder *Hypothesen*) einer Ableitung ist rekursiv definiert durch: *Annahmen*

- (i) $Hyp(X \vdash Y) := \{X \vdash Y\}.$
- (ii) $Hyp\left(\frac{(Ax)}{X, A \vdash A, Y}\right) := \emptyset.$
- (iii) $Hyp\left(\frac{\frac{\mathcal{D}_1}{X_1 \vdash Y_1, A} \quad \frac{\mathcal{D}_2}{A, X_2 \vdash Y_2}}{X_1, X_2 \vdash Y_1, Y_2} (\mathcal{R})\right) := Hyp\left(\frac{\mathcal{D}_1}{X_1 \vdash Y_1, A}\right) \cup Hyp\left(\frac{\mathcal{D}_2}{A, X_2 \vdash Y_2}\right).$

Bemerkung. Eine Ableitung ist also ein (nach oben verzweigender) binärer Baum, dessen Blätter Annahmen sind. Im Fall des Axioms (Ax) ist dessen leere Prämisse als leere Annahme ein Blatt.

Definition 5.5 Es gilt $\Gamma \vdash_{Res} S$ genau dann, wenn es eine mit der Klausel S endende Ableitung $\frac{\mathcal{D}}{S}$ aus Annahmen $Hyp\left(\frac{\mathcal{D}}{S}\right) \subseteq \Gamma$ gibt. $\Gamma \vdash_{Res} S$

$\Gamma \vdash_{Res} S$ bedeutet also, dass die Klausel S im aussagenlogischen Resolutionskalkül aus der Menge von Klauseln Γ *ableitbar* ist. Es ist $\vdash_{Res}$ die *Ableitbarkeitsrelation* (sie darf nicht mit dem Sequenzenzeichen $\vdash$ verwechselt werden). *ableitbar*

Bemerkungen. (i) Das Axiom erzeugt *tautologische Klauseln*. Eine Klausel *tautologische Klausel*

$$\neg A_1 \vee \dots \vee \neg A_n \vee B_1 \vee \dots \vee B_m$$

ist genau dann tautologisch, wenn $A_i = B_j$ für mindestens ein Paar i, j , d. h. wenn sie zwei komplementäre Literale enthält. Denn dann enthält die Klausel die Teilformel $\neg A \vee A$, was wiederum genau dann der Fall ist, wenn die Klausel ein Axiom ist.

- (ii) Mithilfe des Axioms können Klauseln abgeschwächt (verdünnt) werden. Die Klausel $A, X_2 \vdash Y_2$ kann mit dem Axiom $X_1, A \vdash A, Y_1$ über

$$\frac{(Ax) \frac{X_1, A \vdash A, Y_1}{X_1, A, X_2 \vdash Y_1, Y_2} \quad A, X_2 \vdash Y_2}{X_1, A, X_2 \vdash Y_1, Y_2} (\mathcal{R})$$

zu $X_1, A, X_2 \vdash Y_1, Y_2$ verdünnt werden. Entsprechend können Klauseln der Form $X_2 \vdash A, Y_2$ zu $X_1, X_2 \vdash A, Y_1, Y_2$ verdünnt werden.

Theorem 5.6 (Korrektheit) Wenn $\Gamma \vdash_{Res} S$, dann $\Gamma \models S$. Das heißt, der aussagenlogische Resolutionskalkül ist korrekt.

5.2 Resolutionswiderlegung

Im Folgenden behandeln wir den Resolutionskalkül als *Widerlegungskalkül*. Auf das Axiom kann dann verzichtet werden, obgleich die jeweils zu widerlegenden Klauselmengen auch Klauseln von der Form des Axioms enthalten dürfen. Der Resolutionskalkül besteht dann allein aus der Resolutionsregel

$$\frac{X_1 \vdash Y_1, A \quad A, X_2 \vdash Y_2}{X_1, X_2 \vdash Y_1, Y_2} (\mathcal{R})$$

wobei wir im Folgenden $A \notin Y_1$ und $A \notin X_2$ annehmen.

Definition 5.7 Eine *Resolutionswiderlegung* einer Menge Γ von Klauseln ist eine Ableitung der leeren Klausel $\vdash$ (bzw. $\square$) aus Klauseln in Γ .

Resolutionswiderlegung

Da die Resolutionsregel ($\mathcal{R}$) korrekt ist, muss für jede Bewertung mindestens eine der Klauseln in Γ falsch sein, falls eine Resolutionswiderlegung von Γ vorliegt. Mit anderen Worten: Eine Resolutionswiderlegung von Γ bedeutet $\Gamma \vdash_{\text{Res}} \square$ und ist aufgrund Korrektheit der Resolutionsregel ein Beweis für $\Gamma \models \square$, d. h. für die Unerfüllbarkeit von Γ . Betrachtet man eine Formel A , so muss diese zunächst in eine *Klauselmenge* $Kl(A)$ übersetzt werden, damit Resolution anwendbar wird.

Definition 5.8 Ein *Resolutionsbeweis* für eine Formel A ist eine Resolutionswiderlegung von $Kl(\neg A)$.

Resolutionsbeweis

Aus einem Resolutionsbeweis für A folgt $\models A$. Denn es gilt allgemein

$$\begin{aligned} \Gamma \models A &\iff A \text{ ist wahr in allen Modellen von } \Gamma \\ &\iff \Gamma \cup \{\neg A\} \text{ hat kein Modell} \\ &\iff \Gamma \cup \{\neg A\} \text{ ist unerfüllbar} \end{aligned}$$

und für $\Gamma = \emptyset$ somit insbesondere

$$\models A \iff \neg A \text{ ist unerfüllbar.}$$

Aufgrund Korrektheit der Resolutionsregel folgt aus einer Ableitung der unerfüllbaren leeren Klausel $\square$ die Unerfüllbarkeit von $Kl(\neg A)$, und somit auch die Unerfüllbarkeit von $\neg A$.

Allerdings muss für einen Resolutionsbeweis für eine Formel A zunächst die konjunktive Normalform von $\neg A$ gebildet werden, um die erforderliche Klauselmenge $Kl(\neg A)$ angeben zu können. Zur Bildung der KNF wird kein semantisches Verfahren (z. B. Wahrheitstabellen) verwendet, da dieses die Formel entscheidet, und somit das Resolutionsverfahren überflüssig macht. Stattdessen verwenden wir das Umformungsverfahren zur Erzeugung einer KNF aus Definition 4.4, welches durch die Ersetzung möglicher Vorkommen von $\top$ und $\perp$ in der Ausgangsformel durch Formeln $p \vee \neg p$ bzw. $p \wedge \neg p$ für ein neues Aussagesymbol p zudem sicherstellt, dass alle in der KNF vorkommenden Atome Aussagesymbole sind.

Jedem Konjunktionsglied der resultierenden KNF entspricht eine Klausel. Der KNF

$$(\neg A_{11} \vee \dots \vee \neg A_{1k_1} \vee B_{11} \vee \dots \vee B_{1l_1}) \wedge \dots \wedge (\neg A_{j1} \vee \dots \vee \neg A_{jk_j} \vee B_{j1} \vee \dots \vee B_{jl_j})$$

einer Formel A entspricht somit eine *Klauselmenge*

Klauselmenge

$$Kl(A) = \{A_{11}, \dots, A_{1k_1} \vdash B_{11}, \dots, B_{1l_1} ; \dots ; A_{j1}, \dots, A_{jk_j} \vdash B_{j1}, \dots, B_{jl_j}\}$$

die per Resolution behandelbar ist. (Wir grenzen die Elemente von Klauselmengen durch ‘;’ ab.)

Bemerkung. In Disjunktionen von Literalen $L_1 \vee \dots \vee L_n$ können gleiche Literale mehrfach vorkommen. Sind L_i und L_j (für $1 \leq i < j \leq n$) zwei gleiche Literale, dann bezeichnet man die Disjunktion

$$L_1 \vee \dots \vee L_i \vee L_{i+1} \vee \dots \vee L_{j-1} \vee L_{j+1} \vee \dots \vee L_n$$

in der das Vorkommen L_j beseitigt wurde, als *Faktor* der ursprünglichen Disjunktion. Falls keine gleichen Literale mehrfach vorkommen, heißt die Disjunktion von Literalen *faktorfrei*.

Faktor

faktorfrei

Da Antezedens und Sukzedens einer Klausel als Mengen aufgefasst werden, entspricht z. B. der Formel $\neg A \vee \neg A \vee B \vee B$ (für Aussagesymbole A, B) die Klausel $A \vdash B$, der die faktorfreie Formel $\neg A \vee B$ entspricht.

Bemerkungen. (i) Statt $Kl(\neg(A \rightarrow B))$ kann auch $Kl(A) \cup Kl(\neg B)$ betrachtet werden.

(ii) Zur Behandlung von Folgerungsbehauptungen $A_1, \dots, A_n \models B$ kann

$$Kl(\neg(A_1 \wedge \dots \wedge A_n \rightarrow B)) \quad \text{oder} \quad \bigcup_{1 \leq i \leq n} Kl(A_i) \cup Kl(\neg B)$$

betrachtet werden.

Beispiele. Seien A, B, C Aussagesymbole. (Wir verzichten auf die Angabe aller Zwischenschritte bei der Erzeugung der jeweiligen KNF.)

(i) $\models A \rightarrow A \vee B$

KNF von $\neg(A \rightarrow A \vee B)$:

$$\begin{aligned} \neg(A \rightarrow A \vee B) &\rightsquigarrow \neg(\neg A \vee A \vee B) \\ &\rightsquigarrow \neg\neg A \wedge \neg A \wedge \neg B \\ &\rightsquigarrow A \wedge \neg A \wedge \neg B \\ &\rightsquigarrow \{ \vdash A ; A \vdash ; B \vdash \} = Kl(\neg(A \rightarrow A \vee B)) \end{aligned}$$

Resolutionswiderlegung: $\frac{\vdash A \quad A \vdash}{\vdash} (\mathcal{R})$

Es wurde

$$Kl(\neg(A \rightarrow A \vee B)) \vdash_{\text{Res}} \square$$

gezeigt. Mit Korrektheit folgt

$$Kl(\neg(A \rightarrow A \vee B)) \models \square$$

Da $\square$ unerfüllbar ist, muss $Kl(\neg(A \rightarrow A \vee B))$ und damit $\neg(A \rightarrow A \vee B)$ unerfüllbar sein, d. h., $\models A \rightarrow A \vee B$.

(ii) $\models A \vee (B \vee C) \rightarrow (A \vee B) \vee C$

KNF von $\neg(A \vee (B \vee C) \rightarrow (A \vee B) \vee C)$:

$$\begin{aligned} \neg(A \vee (B \vee C) \rightarrow (A \vee B) \vee C) &\rightsquigarrow \neg(\neg(A \vee (B \vee C)) \vee ((A \vee B) \vee C)) \\ &\rightsquigarrow (A \vee (B \vee C)) \wedge \neg((A \vee B) \vee C) \\ &\rightsquigarrow (A \vee B \vee C) \wedge \neg A \wedge \neg B \wedge \neg C \\ &\rightsquigarrow \{ \vdash A, B, C ; A \vdash ; B \vdash ; C \vdash \} \end{aligned}$$

Resolutionswiderlegung: $\frac{\vdash A, B, C \quad A \vdash}{\vdash B, C} (\mathcal{R}) \quad \frac{B \vdash}{\vdash C} (\mathcal{R}) \quad \frac{\vdash C \quad C \vdash}{\vdash} (\mathcal{R})$

Aufgrund des Import-Export-Theorems können wir alternativ auch die Folgerungsbehauptung $A \vee (B \vee C) \models (A \vee B) \vee C$ betrachten:

KNF von $A \vee (B \vee C)$: $A \vee B \vee C \rightsquigarrow \{ \vdash A, B, C \}$.

KNF von $\neg((A \vee B) \vee C)$: $\neg((A \vee B) \vee C) \rightsquigarrow \neg A \wedge \neg B \wedge \neg C \rightsquigarrow \{ A \vdash ; B \vdash ; C \vdash \}$.

Man erhält also ebenfalls die Klauselmenge $\{ \vdash A, B, C ; A \vdash ; A \vdash ; B \vdash ; C \vdash \}$ und die gezeigte Resolutionswiderlegung.

(iii) $\models (\neg A \rightarrow \neg B) \rightarrow (B \rightarrow A)$

KNF von $\neg((\neg A \rightarrow \neg B) \rightarrow (B \rightarrow A))$:

$$\begin{aligned} \neg((\neg A \rightarrow \neg B) \rightarrow (B \rightarrow A)) &\rightsquigarrow \neg((A \vee \neg B) \rightarrow (\neg B \vee A)) \\ &\rightsquigarrow \neg(\neg(A \vee \neg B) \vee \neg B \vee A) \\ &\rightsquigarrow \neg((\neg A \wedge B) \vee \neg B \vee A) \\ &\rightsquigarrow \neg(\neg A \wedge B) \wedge B \wedge \neg A \\ &\rightsquigarrow (A \vee \neg B) \wedge B \wedge \neg A \\ &\rightsquigarrow \{ B \vdash A ; \vdash B ; A \vdash \} \end{aligned}$$

Resolutionswiderlegung:

$$\frac{\vdash B \quad \frac{B \vdash A \quad A \vdash}{B \vdash} (\mathcal{R})}{\vdash} (\mathcal{R})$$

Alternativ für $\neg A \rightarrow \neg B \models B \rightarrow A$:

KNF von $\neg A \rightarrow \neg B$: $A \vee \neg B \rightsquigarrow \{ B \vdash A \}$,

KNF von $\neg(B \rightarrow A)$: $B \wedge \neg A \rightsquigarrow \{ \vdash B ; A \vdash \}$,

Klauselmenge (wie oben): $\{ B \vdash A ; \vdash B ; A \vdash \}$; mit der obigen Resolutionswiderlegung.

- Bemerkungen.** (i) Ein großer Teil des Aufwands geht hierbei in die Bestimmung der KNF.
- (ii) Die KNF muss nicht logisch äquivalent zur Ausgangsformel sein. Für den Widerlegungskalkül genügt eine *erfüllbarkeitsäquivalente* KNF, d. h., es muss lediglich gelten: KNF erfüllbar genau dann, wenn Ausgangsformel erfüllbar.
- (iii) Erfüllbarkeitsäquivalenz ist eine wesentlich schwächere Eigenschaft als logische Äquivalenz: Sei A eine kontingente Formel, dann gilt A erfüllbar gdw. $\neg A$ erfüllbar. Hingegen können A und $\neg A$ natürlich nicht logisch äquivalent sein.
- (iv) Die Anzahl der Schritte zur Erzeugung einer *logisch äquivalenten* KNF (oder einer DNF) ist exponentiell in der Anzahl der Aussagesymbole.
- (v) Für die Erzeugung einer *erfüllbarkeitsäquivalenten* KNF gibt es aber Verfahren mit polynomialer Laufzeit. Dies ist von Vorteil für das Resolutionsverfahren. (Für die Erzeugung einer *allgemeingültigkeitsäquivalenten* KNF gibt es kein Verfahren mit polynomialer Laufzeit.)
- (vi) Für das Auffinden von Resolutionswiderlegungen gibt es Verfahren, die für jede gegebene endliche Klauselmenge Γ entscheiden, ob es eine Resolutionswiderlegung für Γ gibt oder nicht (siehe z. B. [Goltz & Herre, 1990](#)).

6 Kalkül des natürlichen Schließens

Der Resolutionskalkül zeichnet sich durch seine Einfachheit aus, die im Wesentlichen dadurch erreicht wird, dass logische Konstanten im Kalkül nicht vorkommen. Dies schließt allerdings die direkte Behandlung beliebiger Aussagen aus, da diese zunächst in Klauselform gebracht werden müssen.

Hingegen werden in “natürlichen” Argumentationen und Beweisen Aussagen beliebiger logischer Form direkt verwendet, und es wird gemäß logischen Regeln geschlossen, die spezifisch für die jeweils vorkommenden logischen Konstanten sind. Des Weiteren spielt die Verwendung von Annahmen in Argumenten und Beweisen eine Rolle. Betrachtet man etwa die Argumentation im Beweis des Import-Export-Theorems 3.8, so stellt man fest, dass dort zunächst Annahmen eingeführt werden, unter denen dann mit logischen Schlüssen sowie mit Schlüssen, die auf Definitionen zurückgreifen, schrittweise von Aussage zu Aussage übergegangen wird. Im letzten Schritt wird auf das Theorem geschlossen (das hier als Konklusion nicht am Schluss des Beweises steht, sondern darüber), welches von keiner der zwischendurch zu Argumentationszwecken eingeführten Annahmen mehr abhängt.

Nun wollen wir auch diese Aspekte von Argumenten untersuchen. Das Mittel dazu ist der Kalkül des natürlichen Schließens, den wir im Anschluss an das folgende, einführende Beispiel behandeln.

Literatur

- G. Gentzen, *Untersuchungen über das logische Schließen*, Mathematische Zeitschrift **39** (1935), 176–210, 405–431.
- D. Prawitz, *Natural Deduction. A Proof-Theoretical Study*, Almqvist & Wiksell, 1965. Neuauflage 2006, Dover Publications.
- R. Bornat, *Proof and Disproof in Formal Logic. An Introduction for Programmers*, Oxford University Press, 2005.

Beispiel. Wir wollen zeigen, dass $A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$ gilt.

Dazu machen wir zunächst die Annahme, dass $A \vee (B \wedge C)$ gilt. Aufgrund dieser Annahme muss A oder $B \wedge C$ gelten. Als nächstes überlegen wir, was in jedem dieser beiden Fälle gilt.

1. Fall: A gilt. Dann gilt auch $A \vee B$. Dieser Schluss kann als Regelanwendung durch

$$\frac{A}{A \vee B}$$

ausgedrückt werden. Ebenso für $A \vee C$:

$$\frac{A}{A \vee C}$$

Damit haben wir unter der Annahme A gezeigt, dass $A \vee B$ und $A \vee C$ gilt. Damit können wir mit

$$\frac{A \vee B \quad A \vee C}{(A \vee B) \wedge (A \vee C)}$$

auf $(A \vee B) \wedge (A \vee C)$ schließen.

Das gesamte Argument lässt sich so darstellen:

$$\frac{\frac{A}{A \vee B} \quad \frac{A}{A \vee C}}{(A \vee B) \wedge (A \vee C)}$$

Die Formel $(A \vee B) \wedge (A \vee C)$ hängt noch von der (zweifach verwendeten) Annahme A ab.

2. Fall: $B \wedge C$ gilt. Somit gilt sowohl B als auch C ; als Regelanwendungen:

$$\frac{B \wedge C}{B} \quad \frac{B \wedge C}{C}$$

Nun können wir von B weiter auf $A \vee B$ und von C auf $A \vee C$ schließen; als Regelanwendungen:

$$\frac{B}{A \vee B} \quad \frac{C}{A \vee C}$$

Damit gilt mit

$$\frac{A \vee B \quad A \vee C}{(A \vee B) \wedge (A \vee C)}$$

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

Das gesamte Argument ist

$$\frac{\frac{B \wedge C}{B} \quad \frac{B \wedge C}{C}}{\frac{A \vee B \quad A \vee C}{(A \vee B) \wedge (A \vee C)}}$$

Die Formel $(A \vee B) \wedge (A \vee C)$ hängt hier noch von der (zweifach verwendeten) Annahme $B \wedge C$ ab.

Somit haben wir gezeigt, dass unter der Annahme $A \vee (B \wedge C)$ in beiden Fällen $(A \vee B) \wedge (A \vee C)$ gilt. Das Argument hat grob die Struktur

$$\frac{A \vee (B \wedge C)}{\vdots} \\ (A \vee B) \wedge (A \vee C)$$

Die Formel $(A \vee B) \wedge (A \vee C)$ hängt also noch von der Annahme $A \vee (B \wedge C)$ ab. Diese Abhängigkeit lösen wir durch die Einführung einer Implikation auf, deren Vorderglied gerade die angenommene Formel ist:

$$\frac{\frac{[A \vee (B \wedge C)]^1}{\vdots} \quad (A \vee B) \wedge (A \vee C)}{A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)} 1$$

(Die Tatsache, dass $A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$ nicht mehr von der Annahme $A \vee (B \wedge C)$ abhängt, drücken wir durch die Markierung der Annahme mit eckigen Klammern und einer Ziffer aus, die wir auch neben der für die Auflösung der Abhängigkeit verantwortlichen Regel notieren.)

Somit gilt $A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$.

Das Beispiel veranschaulicht Folgendes:

- (i) Argumente können aus elementaren Schritten baumförmig aufgebaut werden.
- (ii) Die elementaren Schritte erfolgen gemäß Regeln.
- (iii) Jede Regel bezieht sich auf nur eine logische Konstante.
(Etwaige Regeln, die sich aus Definitionen oder empirischen Zusammenhängen ergeben, lassen wir hier außer Acht.)
- (iv) Es gibt Regeln, bei deren Anwendung eine logische Konstante eingeführt wird (z. B. bei $\frac{A}{A \vee B}$), und es gibt Regeln, bei deren Anwendung eine logische Konstante beseitigt wird (z. B. bei $\frac{B \wedge C}{B}$).
- (v) Es dürfen Annahmen gemacht und verwendet werden.
- (vi) Annahmen ermöglichen weitere Regelanwendungen, wobei Aussagen, auf die unter Annahmen geschlossen wurde, noch von diesen abhängen.
- (vii) Es gibt Regeln, bei denen die Abhängigkeit von Annahmen aufgelöst werden kann (z. B. bei der Einführung einer Implikation).

Diese Eigenschaften werden im Kalkül des natürlichen Schließens umgesetzt. Zunächst geben wir die Regeln des Kalküls an. Danach erklären wir, wie die Regeln zu Argumenten (bzw. zu Ableitungen im Kalkül) zusammengesetzt werden können, und wie der Umgang mit Annahmen im Kalkül behandelt wird. Dabei beschränken wir uns auf die logischen Konstanten $\perp$, $\neg$, $\wedge$, $\vee$ und $\rightarrow$, wobei wir $\neg A$ als Abkürzung für $A \rightarrow \perp$ auffassen.

Definition 6.1 (i) Der *Kalkül NK des natürlichen Schließens (für klassische Logik)* ist durch folgende Regeln gegeben: *Kalkül NK*

<i>Einführungsregel</i>	<i>Beseitigungsregel</i>
$\frac{A_1 \quad A_2}{A_1 \wedge A_2} (\wedge E)$	$\frac{A_1 \wedge A_2}{A_i} (\wedge B) (i = 1 \text{ oder } 2)$
$\frac{A_i}{A_1 \vee A_2} (\vee E) (i = 1 \text{ oder } 2)$	$\frac{A_1 \vee A_2 \quad \begin{array}{c} [A_1] \\ C \end{array} \quad \begin{array}{c} [A_2] \\ C \end{array}}{C} (\vee B)$
$\frac{\begin{array}{c} [A] \\ B \end{array}}{A \rightarrow B} (\rightarrow E)$	$\frac{A \rightarrow B \quad A}{B} (\rightarrow B)$
<i>Widerspruchsregel</i>	
$\frac{\begin{array}{c} [\neg A] \\ \perp \\ A \end{array}}{\perp} (\perp)$	

- (ii) Rechts vom Regelstrich steht der *Regelname*; zum Beispiel $(\rightarrow E)$ für *Implikations-einführungsregel* oder $(\vee B)$ für *Disjunktionbeseitigungsregel*.
- (iii) Die jeweils unter dem Regelstrich stehende Formel heißt *Konklusion* der Regel. Die direkt über dem Regelstrich stehenden Formeln heißen *Prämissen* der Regel.

*Konklusion
Prämissen*

(iv) Die Prämisse $A_1 \vee A_2$ der Regel $(\vee B)$ heißt auch *Hauptprämisse*, die Prämissen C auch *Nebenprämissen*. Bei $(\rightarrow B)$ heißt die Prämisse $A \rightarrow B$ *Hauptprämisse* und die Prämisse A *Nebenprämisse*. *Hauptprämisse*
Nebenprämisse

(v) Die bei den Regeln $(\vee B)$, $(\rightarrow E)$ und $(\perp)$ in eckigen Klammern notierten Formeln $[A]$ zeigen an, dass die Konklusion der Regel nicht mehr von Annahmen A abhängt.

Bemerkungen. (i) Für jede der logischen Konstanten $\wedge$, $\vee$ und $\rightarrow$ gibt es jeweils eine Einführungs- und eine Beseitigungsregel.

Bei einer Einführungsregel kommt die jeweilige logische Konstante in der Konklusion der Regel hinzu. Bei einer Beseitigungsregel wird die in der (Haupt-)Prämisse vorkommende Konstante beim Übergang zur Konklusion eliminiert.

(ii) Da wir $\neg A$ als Abkürzung für $A \rightarrow \perp$ auffassen, können wir auf ein zusätzliches Regelpaar

$$\frac{[A]}{\frac{\perp}{\neg A}} (\neg E) \qquad \frac{\neg A \quad A}{\perp} (\neg B)$$

für die Negation $\neg$ verzichten. Denn mit $\neg A := A \rightarrow \perp$ ist

$$\frac{[A]}{\frac{\perp}{\neg A}} (\neg E)$$

die Instanz

$$\frac{[A]}{\frac{\perp}{A \rightarrow \perp}}$$

der Implikationseinführungsregel $(\rightarrow E)$, und

$$\frac{\neg A \quad A}{\perp} (\neg B)$$

ist die Instanz

$$\frac{A \rightarrow \perp \quad A}{\perp}$$

der Implikationsbeseitigungsregel $(\rightarrow B)$.

(iii) Jede der 6 Einführungs- und Beseitigungsregeln behandelt genau eine logische Konstante.

(iv) Bei der Widerspruchsregel

$$\frac{[\neg A]}{\frac{\perp}{A}} (\perp)$$

ist dies nicht der Fall, da neben $\perp$ auch $\neg$ vorkommt. Des Weiteren gibt es für $\perp$ kein Regelpaar aus Einführungs- und Beseitigungsregel. Die Widerspruchsregel kann zwar als Beseitigungsregel für $\perp$ aufgefasst werden; es gibt jedoch keine Einführungsregel.

Die Widerspruchsregel $(\perp)$ drückt aus, dass von der Prämisse $\perp$ zu einer beliebigen Formel A als Konklusion übergegangen werden darf, wobei letztere nicht mehr von Annahmen $\neg A$ abhängt.

(v) Die Konjunktionseinführungsregel

$$\frac{A_1 \quad A_2}{A_1 \wedge A_2} (\wedge E)$$

besagt, dass von den beiden Prämissen A_1 und A_2 zur Konklusion $A_1 \wedge A_2$ übergegangen werden darf; die Konjunktionsbeseitigungsregel

$$\frac{A_1 \wedge A_2}{A_i} (\wedge B) \quad (i = 1 \text{ oder } 2)$$

besagt, dass von der Prämisse $A_1 \wedge A_2$ zu A_1 oder zu A_2 als Konklusion übergegangen werden darf.

(vi) Bei der Disjunktionseinführungsregel

$$\frac{A_i}{A_1 \vee A_2} (\vee E) \quad (i = 1 \text{ oder } 2)$$

darf von der Prämisse A_1 ebenso wie von der Prämisse A_2 zur Konklusion $A_1 \vee A_2$ übergegangen werden. Die Regel

$$\frac{A_1 \vee A_2 \quad \begin{array}{c} [A_1] \\ C \end{array} \quad \begin{array}{c} [A_2] \\ C \end{array}}{C} (\vee B)$$

besagt, dass von der Hauptprämisse $A_1 \vee A_2$ und den beiden Nebenprämissen C zur Konklusion C übergegangen werden darf, wobei die Konklusion nicht mehr abhängt von Annahmen A_1 (von denen die linke Nebenprämisse noch abhängen kann) und A_2 (von denen die rechte Nebenprämisse noch abhängen kann).

(vii) Die Regel

$$\frac{\begin{array}{c} [A] \\ B \end{array}}{A \rightarrow B} (\rightarrow E)$$

drückt aus, dass von der Prämisse B zur Konklusion $A \rightarrow B$ übergegangen werden darf, die nicht mehr von Annahmen A abhängt. Die Implikationsbeseitigungsregel

$$\frac{A \rightarrow B \quad A}{B} (\rightarrow B)$$

(modus ponens) besagt, dass von der Hauptprämisse $A \rightarrow B$ mit der Nebenprämisse A zu B übergegangen werden darf.

(viii) Die Regeln $(\wedge B)$ und $(\vee E)$ könnten auch explizit als Regelpaare

$$\frac{A \wedge B}{A} (\wedge B) \quad \frac{A \wedge B}{B} (\wedge B)$$

bzw.

$$\frac{A}{A \vee B} (\vee E) \quad \frac{B}{A \vee B} (\vee E)$$

angegeben werden.

(ix) Auf eine weitere Differenzierung durch verschiedene Regelnamen wie z. B. $(\wedge B)_{\text{links}}$ und $(\wedge B)_{\text{rechts}}$, die eine Unterscheidung der beiden Regelanwendungen

$$\frac{A \wedge A}{A} (\wedge B)_{\text{links}} \quad \text{und} \quad \frac{A \wedge A}{A} (\wedge B)_{\text{rechts}}$$

ermöglichen würden, verzichten wir hier.

Definition 6.2 Eine *Ableitung* im Kalkül des natürlichen Schließens NK ist ein (nach oben verzweigender) Baum, der wie folgt induktiv definiert ist:

Ableitung

- (i) Für jede Formel A ist der Knoten

$$A$$

eine Ableitung.

Blatt und Wurzelknoten des Baums fallen hier zusammen. Der Knoten A ist eine Ableitung der *Endformel* A (Wurzelknoten) aus der *Annahme* A (Blatt).

Endformel
Annahme

- (ii) Sind $\mathcal{D}$, $\mathcal{D}_1$ und $\mathcal{D}_2$ Ableitungen, dann sind auch die folgenden Bäume Ableitungen:

$$\begin{array}{c} \mathcal{D}_1 \quad \mathcal{D}_2 \\ \frac{A_1 \quad A_2}{A_1 \wedge A_2} (\wedge E) \end{array} \qquad \begin{array}{c} \mathcal{D}_1 \\ \frac{A_1 \wedge A_2}{A_i} (\wedge B) \quad (i = 1 \text{ oder } 2) \end{array}$$

$$\begin{array}{c} \mathcal{D}_i \\ \frac{A_i}{A_1 \vee A_2} (\vee E) \quad (i = 1 \text{ oder } 2) \end{array} \qquad \begin{array}{c} [A_1]^n \quad [A_2]^n \\ \mathcal{D} \quad \mathcal{D}_1 \quad \mathcal{D}_2 \\ \frac{A_1 \vee A_2 \quad C \quad C}{C} (\vee B)^n \end{array}$$

$$\begin{array}{c} [A]^n \\ \mathcal{D} \\ \frac{B}{A \rightarrow B} (\rightarrow E)^n \end{array} \qquad \begin{array}{c} \mathcal{D}_1 \quad \mathcal{D}_2 \\ \frac{A \rightarrow B \quad A}{B} (\rightarrow B) \end{array}$$

$$\begin{array}{c} [\neg A]^n \\ \mathcal{D} \\ \frac{\perp}{A} (\perp)^n \end{array}$$

Definition 6.3 (i) Die Formeln unter dem Regelstrich sind die Wurzelknoten des Baums, die in Ableitungen als *Endformeln* bezeichnet werden. Die Blätter des Baums heißen *Annahmen*.

Endformeln
Annahmen

- (ii) Die Notation $\frac{[A]}{\mathcal{D}} \frac{B}{A \rightarrow B}$ drückt aus, dass in der Ableitung $\mathcal{D}$ mit Endformel B Annahmen A vorkommen können. Statt $\mathcal{D}$ schreiben wir auch $\dot{}$ als Platzhalter für eine Ableitung.

- (iii) Die durch eckige Klammern gekennzeichneten Annahmen müssen dabei in den jeweiligen Ableitungen nicht vorkommen. Falls sie vorkommen, können sie beim Übergang zur Endformel (d. h. bei Anwendung der angegebenen Regel) gelöscht werden. Die *Löschung* von Annahmen A besteht in der Markierung von Blättern A durch eckige Klammern und einer (noch nicht zu einer Löschung verwendeten) Ziffer n , wie folgt: $[A]^n$. Zusätzlich wird die Ziffer n bei der angewandten Regel notiert.

Löschung

Beispiel.

$$\begin{array}{c} [A]^1 \\ \mathcal{D} \\ \frac{B}{A \rightarrow B} (\rightarrow E)^1 \end{array}$$

(sofern die Ziffer 1 in $\mathcal{D}$ noch nicht zur Löschung einer Annahme verwendet wurde).

- (iv) Eine gelöschte Annahme heißt auch *geschlossen*. Eine nicht geschlossene (bzw. nicht gelöschte) Annahme heißt *offen*. *geschlossen*
offen

Beispiel. Sowohl $\frac{B}{A \rightarrow B} (\rightarrow E)$, $\frac{A}{A \rightarrow A} (\rightarrow E)$ als auch $\frac{[A]^1}{A \rightarrow A} (\rightarrow E)^1$ sind korrekte Ableitungen, bzw. korrekte Anwendungen der Regel $(\rightarrow E)$. In der ersten Ableitung ist die Annahme B offen, in der zweiten A ; in der dritten wurde die Annahme A (die zugleich Prämisse der Regel ist) gelöscht (bzw. geschlossen).

Definition 6.4 Eine *Teilableitung* $\mathcal{D}'$ einer Ableitung $\mathcal{D}$ ist ein Teilbaum von $\mathcal{D}$. Gelöschte Annahmen der Teilableitung $\mathcal{D}'$ sind nur jene gelöschten Annahmen von $\mathcal{D}$, deren bei ihrer Markierung verwendete Ziffern n bei einer Regelanwendung in $\mathcal{D}'$ stehen. *Teilableitung*

- Bemerkungen.** (i) Die Löschung von Annahmen ist nur bei Anwendung der Regeln $(\vee B)$, $(\rightarrow E)$ und $(\perp)$ möglich. Es dürfen ausschließlich die in den Regeln in eckigen Klammern angegebenen Formeln als Annahmen gelöscht werden. Bei $(\vee B)$ darf die Löschung von A_1 nur in Teilableitung $\mathcal{D}_1$, die von A_2 nur in $\mathcal{D}_2$ vorgenommen werden.
- (ii) Ist $\mathcal{D}'$ eine Teilableitung von $\mathcal{D}$, und kommt die Ziffer n einer in $\mathcal{D}$ geschlossenen Annahme A bei einer Regelanwendung in $\mathcal{D}$ vor, die nicht Bestandteil von $\mathcal{D}'$ ist, so ist A in $\mathcal{D}'$ offen.

Beispiele. (i) Es ist

$$\frac{(A \wedge B) \vee B \quad \frac{\frac{[A \wedge B]^1}{B} (\wedge B) \quad [B]^1}{(\vee B)^1}}{B}$$

eine korrekte Ableitung von B aus $(A \wedge B) \vee B$.

Die Annahme $A \wedge B$ ist in der Teilableitung $\frac{A \wedge B}{B} (\wedge B)$ offen, in der gesamten Ableitung ist sie geschlossen.

(ii) *Keine* korrekte Ableitung ist

$$\not\vdash \frac{A_1 \vee A_2 \quad \frac{\frac{[A_1]^1}{A_1 \wedge A_2} (\wedge E) \quad \frac{[A_2]^{1\cancel{z}}}{A_1 \wedge A_2} (\wedge E)}{A_1 \wedge A_2} (\vee B)^{1\cancel{z}}}{A_1 \wedge A_2} \not\vdash$$

Die Disjunktionsbeseitigungsregel $(\vee B)$ erlaubt nur die Löschung von A_1 (linkes Disjunktionsglied) in der Teilableitung der linken Nebenprämisse und von A_2 (rechtes Disjunktionsglied) in der Teilableitung der rechten Nebenprämisse. Die Löschung der beiden mit $\cancel{z}$ gekennzeichneten Annahmen bei $(\vee B)$ ist *inkorrekt*.

(iii) Eine korrekte Ableitung ist

$$\frac{A_1 \vee A_2 \quad \frac{\frac{[A_1]^1}{A_1 \wedge A_2} (\wedge E) \quad A_2}{A_1 \wedge A_2} (\wedge E) \quad \frac{A_1 \quad \frac{[A_2]^1}{A_1 \wedge A_2} (\wedge E)}{A_1 \wedge A_2} (\wedge E)}{A_1 \wedge A_2} (\vee B)^1$$

Die Endformel $A_1 \wedge A_2$ hängt von den offenen Annahmen $A_1 \vee A_2$, A_2 und A_1 ab. (Die Formel $A_1 \wedge A_2$ kann aber auch direkt mit $(\wedge E)$ aus den beiden Annahmen A_1 und A_2 allein abgeleitet werden.)

Definition 6.5 Die *Annahmenmenge* (oder *Hypothesenmenge*) $Hyp(\mathcal{D})$ einer Ableitung $\mathcal{D}$ ist rekursiv definiert durch:

- (i) $Hyp(A) := \{A\}$.

Das heißt, die Annahmenmenge einer aus einem Blatt A bestehenden Ableitung A ist die Menge $\{A\}$ (deren einziges Element die Formel A ist).

$$(ii) \quad Hyp\left(\frac{\mathcal{D}_1 \quad \mathcal{D}_2}{\frac{A_1 \quad A_2}{A_1 \wedge A_2} (\wedge E)}\right) := Hyp\left(\frac{\mathcal{D}_1}{A_1}\right) \cup Hyp\left(\frac{\mathcal{D}_2}{A_2}\right).$$

Entsprechend für $(\wedge B)$, $(\vee E)$ und $(\rightarrow B)$.

- (iii) Sofern bei Anwendung von $(\rightarrow E)$ Annahmen A nicht gelöscht wurden, gilt:

$$Hyp\left(\frac{\begin{array}{c} [A] \\ \mathcal{D}_1 \\ B \\ \hline A \rightarrow B \end{array} (\rightarrow E)}{A \rightarrow B} (\rightarrow E)\right) := Hyp\left(\frac{[A] \quad \mathcal{D}_1}{B}\right).$$

Entsprechend für $(\vee B)$ und $(\perp)$.

Dies soll den Fall einschließen, dass einige, aber nicht alle Annahmen gelöscht wurden. Da wir Mengen betrachten, ist die Anzahl nicht gelöschter Annahmen irrelevant, sofern mindestens eine Annahme nicht gelöscht wurde.

- (iv) Sofern bei Anwendung von $(\rightarrow E)$ alle Annahmen A gelöscht wurden, gilt:

$$Hyp\left(\frac{\begin{array}{c} [A]^n \\ \mathcal{D}_1 \\ B \\ \hline A \rightarrow B \end{array} (\rightarrow E)^n}{A \rightarrow B} (\rightarrow E)^n\right) := Hyp\left(\frac{[A] \quad \mathcal{D}_1}{B}\right) \setminus \{A\}.$$

Entsprechend für $(\vee B)$ und $(\perp)$.

$Hyp(\mathcal{D})$ ist also die Menge der offenen Annahmen in $\mathcal{D}$.

Beispiel. Wir betrachten die Ableitung $\mathcal{D}$ mit $Hyp(\mathcal{D}) = \{A \rightarrow (B \rightarrow C)\}$:

$$\mathcal{D} \left\{ \frac{\frac{\frac{A \rightarrow (B \rightarrow C) \quad [A]^1 (\rightarrow B)}{B \rightarrow C} \quad \frac{[A \rightarrow B]^2 \quad [A]^1 (\rightarrow B)}{B} (\rightarrow B)}{C} (\rightarrow E)^1}{A \rightarrow C} (\rightarrow E)^2}{(A \rightarrow B) \rightarrow (A \rightarrow C)} (\rightarrow E)^2 \right\} \mathcal{D}'$$

$\mathcal{D}'$ ist eine Teildableitung von $\mathcal{D}$. Es ist $Hyp(\mathcal{D}') = \{A \rightarrow (B \rightarrow C), A \rightarrow B\}$. Für die Teildableitung $\mathcal{D}''$ von $\mathcal{D}'$ ist $Hyp(\mathcal{D}'') = \{A \rightarrow (B \rightarrow C), A \rightarrow B, A\}$.

- Definition 6.6** (i) Aus Γ ist A *ableitbar* (Notation: $\Gamma \vdash A$), falls es eine Ableitung $\mathcal{D}$ mit Endformel A und $Hyp(\mathcal{D}) \subseteq \Gamma$ gibt. *ableitbar*
- (ii) Eine Formel A heißt *beweisbar* (Notation: $\vdash A$), falls es eine Ableitung $\mathcal{D}$ mit Endformel A und $Hyp(\mathcal{D}) = \emptyset$ gibt. Die Ableitung $\mathcal{D}$ heißt in diesem Fall auch (*formaler*) *Beweis* von A . *beweisbar*
Beweis

Bemerkungen. (i) Auf Zusätze wie in “ableitbar in NK”, “in NK beweisbar” oder “ $\Gamma \vdash_{NK} A$ ” verzichten wir an dieser Stelle, da wir uns bis auf Weiteres stets auf den Kalkül NK beziehen.

- (ii) Für die Notation von Mengen von Formeln gilt das bei der logischen Folgerung Gesagte (siehe Bemerkung (ii) nach Definition 3.7).

Beispiel. Wir zeigen die Beweisbarkeitsbehauptung

$$\vdash (A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)$$

Die *Beweissuche* kann von unten nach oben erfolgen:

- (i) Wir beginnen mit der zu beweisenden Endformel

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

die mit ($\rightarrow$ E) abgeleitet werden kann:

$$\frac{A \rightarrow B \wedge C}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)$$

Für die Beweissuche ist dieser Ansatz von Vorteil, da es jeweils nur eine Einführungsregel gibt, mit der eine Formel abgeleitet werden kann (die Wahl der Einführungsregel ist durch das Hauptkonjunktiv der Formel festgelegt); bei Beseitigungsregeln ist dies nicht der Fall. Zudem ist bei Einführungsregeln die Wahl der Prämisse(n) sehr viel stärker beschränkt als bei Beseitigungsregeln; im Fall ($\rightarrow$ E) gibt es überhaupt keine Wahlmöglichkeit. Bei der hier vorgenommenen Anwendung von ($\rightarrow$ E) wissen wir zusätzlich, dass Annahmen $(A \rightarrow B) \wedge (A \rightarrow C)$ gelöscht werden können. Diese Annahmen können wir später verwenden, um die Beweissuche von oben nach unten fortzusetzen.

- (ii) Wir setzen die Suche zunächst von unten nach oben fort. Die Formel $A \rightarrow B \wedge C$ kann mit ($\rightarrow$ E) abgeleitet werden:

$$\frac{\frac{B \wedge C}{A \rightarrow B \wedge C} (\rightarrow E)}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)$$

Bei dieser Anwendung von ($\rightarrow$ E) können Annahmen A gelöscht werden. Auch dies merken wir uns für eine spätere Verwendung.

- (iii) Die Formel $B \wedge C$ leiten wir mit ($\wedge$ E) ab:

$$\frac{\frac{\frac{B}{B \wedge C} (\wedge E) \quad \frac{C}{B \wedge C} (\wedge E)}{A \rightarrow B \wedge C} (\rightarrow E)}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)$$

- (iv) Die Prämissen B und C können nicht Konklusion einer Einführungsregel sein. Die Suche von unten nach oben muss also mit Beseitigungsregeln oder mit der Widerspruchregel fortgesetzt werden. Um den Suchraum einzuschränken, überlegen wir uns, welche Formeln wir mit Beseitigungsregeln aus den gemerkten Annahmen ableiten können. Aus $(A \rightarrow B) \wedge (A \rightarrow C)$ leiten wir mit ($\wedge$ B) die Formeln $A \rightarrow B$ und $A \rightarrow C$ ab:

$$\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow B} (\wedge B) \qquad \frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow C} (\wedge B)$$

- (v) Unter Verwendung der Annahme A erhalten wir mit ($\rightarrow$ B) die beiden Ableitungen:

$$\frac{\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow B} (\wedge B) \quad A}{B} (\rightarrow B) \qquad \frac{\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow C} (\wedge B) \quad A}{C} (\rightarrow B)$$

(vi) Damit haben wir die gesuchten (Teil-)Ableitungen für B und C gefunden, mit denen wir die Ableitung

$$\frac{\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow B} (\wedge B) \quad \frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow C} (\wedge B)}{\frac{B \quad C}{B \wedge C} (\wedge E)} (\rightarrow B) \quad \frac{B \wedge C}{A \rightarrow B \wedge C} (\rightarrow E) \quad \frac{A \rightarrow B \wedge C}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)$$

erhalten. Dies ist noch nicht der gesuchte Beweis für $(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)$, da die Endformel noch von den offenen Annahmen $(A \rightarrow B) \wedge (A \rightarrow C)$ und A abhängt.

(vii) Die noch offenen Annahmen können bei den Anwendungen von $(\rightarrow E)$ gelöscht werden. In der Ableitung

$$\frac{\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow B} (\wedge B) \quad [A]^1}{B} (\rightarrow B) \quad \frac{\frac{(A \rightarrow B) \wedge (A \rightarrow C)}{A \rightarrow C} (\wedge B) \quad [A]^1}{C} (\rightarrow B) \quad \frac{B \wedge C}{A \rightarrow B \wedge C} (\rightarrow E)^1 \quad \frac{A \rightarrow B \wedge C}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)$$

markiert die Ziffer 1 die Löschung der Annahmen A . Die Endformel $(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)$ hängt nun nur noch von den offenen Annahmen $(A \rightarrow B) \wedge (A \rightarrow C)$ ab.

(viii) Schließlich erhalten wir durch die Löschung der Annahmen $(A \rightarrow B) \wedge (A \rightarrow C)$ bei der untersten Anwendung von $(\rightarrow E)$, markiert durch die Ziffer 2, den Beweis

$$\frac{\frac{[(A \rightarrow B) \wedge (A \rightarrow C)]^2}{A \rightarrow B} (\wedge B) \quad [A]^1}{B} (\rightarrow B) \quad \frac{\frac{[(A \rightarrow B) \wedge (A \rightarrow C)]^2}{A \rightarrow C} (\wedge B) \quad [A]^1}{C} (\rightarrow B) \quad \frac{B \wedge C}{A \rightarrow B \wedge C} (\rightarrow E)^1 \quad \frac{A \rightarrow B \wedge C}{(A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)} (\rightarrow E)^2$$

womit $\vdash (A \rightarrow B) \wedge (A \rightarrow C) \rightarrow (A \rightarrow B \wedge C)$ gezeigt ist.

Beispiel. Wir zeigen die Ableitbarkeitsbehauptung $\neg\neg A \vdash A$. Diese ist gleichbedeutend mit der Behauptung $(A \rightarrow \perp) \rightarrow \perp \vdash A$, da wir $\neg A$ als Abkürzung für $A \rightarrow \perp$ auffassen.

$$\frac{(A \rightarrow \perp) \rightarrow \perp \quad [A \rightarrow \perp]^1}{\perp} (\rightarrow B) \quad \frac{\perp}{A} (\perp)^1$$

Die Ziffer 1 markiert die Löschung der zusätzlichen Annahme $A \rightarrow \perp$. Die Endformel A hängt nur noch von der Annahme $(A \rightarrow \perp) \rightarrow \perp$ ab.

Die Abkürzung $\neg A := A \rightarrow \perp$ muss in Ableitungen nicht aufgelöst werden; die Ableitung

$$\frac{\neg\neg A \quad [\neg A]^1}{\perp} (\rightarrow B) \quad \frac{\perp}{A} (\perp)^1$$

ist also ebenfalls korrekt.

Die Widerspruchsregel $(\perp)$ entspricht somit der Beseitigung der doppelten Negation.

Theorem 6.7 Sei Γ eine beliebige (möglicherweise unendliche) Menge von Formeln, und A eine beliebige Formel. Falls $\Gamma \vdash A$ gilt, gibt es eine endliche Teilmenge Γ' von Γ , so dass $\Gamma' \vdash A$ gilt.

Beweis. Wenn $\Gamma \vdash A$ gilt, muss es eine Ableitung $\mathcal{D}$ mit $\text{Hyp}(\mathcal{D}) \subseteq \Gamma$ geben. Da die Annahmenmenge einer Ableitung stets endlich ist, ist mit $\text{Hyp}(\mathcal{D})$ eine endliche Annahmenmenge $\Gamma' \subseteq \Gamma$ gefunden, so dass $\Gamma' \vdash A$. QED

Zum Auffinden einer Ableitung von A aus Annahmen Γ kann die folgende Strategie nützlich sein, bei der eine Ableitung von unten nach oben konstruiert wird:

- (1) Falls die Formel A Konklusion einer Einführungsregel sein kann, erweitert man A um die entsprechende Einführungsregel nach oben. Die zu verwendende Einführungsregel ist durch das Hauptkonnektiv von A eindeutig bestimmt.

Dies liefert die Prämisse(n) der Einführungsregel, wobei im Fall $(\vee E)$ eines der beiden Disjunktionsglieder als Prämisse gewählt werden muss. Falls diese Wahl am Ende nicht zum Ziel führt, kann an dieser Stelle das andere Disjunktionsglied als Prämisse gewählt werden.

Falls die gesuchte Einführungsregel eine Implikationseinführung $(\rightarrow E)$ ist, merkt man sich die Annahme, die bei dieser Regelanwendung gelöscht werden kann. Diese Annahme kann man sich später zunutze machen.

Nun geht man zu Schritt (2), oder man wiederholt dieses Vorgehen für jede der neu hinzugekommenen Prämissen, bis man nur noch Prämissen hat, die nicht mehr Konklusion einer Einführungsregel sein können.

- (2) Für die nun zuoberst stehenden Formeln probiert man eine Beseitigungsregel, mit der die jeweilige Formel abgeleitet werden kann. Hier gibt es mehrere Möglichkeiten, sowohl bei der Wahl der Regel als auch bei der Wahl der Prämissen. Einen Anhaltspunkt für die Wahl können die Annahmen Γ sowie die in Schritt (1) gemerkten Annahmen darstellen.
- (3) Für die gewählte Beseitigungsregel beschränkt man sich bei der Wahl der Prämisse(n) auf Formeln, in denen neben der Konklusion der Beseitigungsregel nur Teilformeln der Annahmen als Teilformeln vorkommen.

Ist die gewählte Beseitigungsregel $(\vee B)$, sind die Nebenprämissen schon durch die Konklusion festgelegt. Zur späteren Verwendung merkt man sich auch hier jene beiden Annahmen, die aufgrund der gewählten Hauptprämisse bei dieser Regelanwendung gelöscht werden können.

Gehe zu (1), wobei A nun eine der gewählten Prämissen sei.

Falls dieser Ansatz nicht zum Ziel führt, verwendet man in Schritt (2) statt einer Beseitigungsregel die Widerspruchsregel $(\perp)$. Schlägt auch dies fehl, verwendet man die Widerspruchsregel schon in Schritt (1) anstelle einer Einführungsregel. Zudem kann es hilfreich sein, zwischendurch Teildableitungen aus den gemerkten und den in Γ gegebenen Annahmen von oben nach unten zu entwickeln.

Beispiele. (i) Wir greifen das eingangs behandelte Beispiel (vgl. S. 47f.) auf, und zeigen

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

durch Angabe einer Ableitung:

$$\frac{[A \vee (B \wedge C)]^2 \quad \frac{\frac{[A]^1}{A \vee B} (\vee E) \quad \frac{[A]^1}{A \vee C} (\vee E)}{(A \vee B) \wedge (A \vee C)} (\wedge E) \quad \frac{\frac{[B \wedge C]^1}{B} (\wedge B) \quad \frac{[B \wedge C]^1}{C} (\wedge B)}{A \vee B} (\vee E) \quad \frac{[B \wedge C]^1}{A \vee C} (\vee E)}{(A \vee B) \wedge (A \vee C)} (\wedge E)}{(A \vee B) \wedge (A \vee C)} (\vee B)^1 \quad \frac{(A \vee B) \wedge (A \vee C)}{A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)} (\rightarrow E)^2$$

Die Ableitung ist ein Beweis der Formel $A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$.

(ii) Wir zeigen das De Morgansche Gesetz $\neg(A \wedge B) \vdash \neg A \vee \neg B$:

$$\frac{\neg(A \wedge B) \quad \frac{\frac{[\neg(\neg A \vee \neg B)]^3}{\perp} (\perp)^1 \quad \frac{[\neg A]^1}{\neg A \vee \neg B} (\vee E)}{\neg A \vee \neg B} (\rightarrow B) \quad \frac{\frac{[\neg(\neg A \vee \neg B)]^3}{\perp} (\perp)^2 \quad \frac{[\neg B]^2}{\neg A \vee \neg B} (\vee E)}{B} (\wedge E)}{A \wedge B} (\rightarrow B) \quad \frac{\perp}{\neg A \vee \neg B} (\perp)^3$$

6.1 Strukturregeln

Theorem 6.8 Es gelten die folgenden Strukturregeln:

Strukturregeln

- (i) *Identität:* $A \vdash A$.
- (ii) *Verdünnung:* Wenn $\Gamma \vdash A$, dann $\Gamma, \Delta \vdash A$.
- (iii) *Kontraktion:* Wenn $\Gamma, A, A \vdash B$, dann $\Gamma, A \vdash B$.
- (iv) *Schnitt:* Wenn $\Gamma \vdash A$ und $\Delta, A \vdash B$, dann $\Gamma, \Delta \vdash B$.

Beweis. (i) Es ist A eine Ableitung $\mathcal{D}$ mit Endformel A und $Hyp(\mathcal{D}) = A$. Somit gilt $A \vdash A$.

(ii) Angenommen $\Gamma \vdash A$. Dann muss es eine Ableitung $\mathcal{D}$ mit Endformel A und $Hyp(\mathcal{D}) \subseteq \Gamma$ geben. Wegen $Hyp(\mathcal{D}) \subseteq \Gamma \subseteq \Gamma \cup \Delta$ gibt es dann auch eine Ableitung $\mathcal{D}'$ mit Endformel A und $Hyp(\mathcal{D}') \subseteq \Gamma \cup \Delta$. Also gilt $\Gamma, \Delta \vdash A$.

(iii) Angenommen $\Gamma, A, A \vdash B$. Da Γ, A, A für die Menge $\{\Gamma, A, A\} = \{\Gamma, A\}$ steht, gilt auch $\Gamma, A \vdash B$.

(iv) Angenommen, es gilt $\Gamma \vdash A$ und $\Delta, A \vdash B$. Dann muss es eine Ableitung $\mathcal{D}_1$ mit $Hyp(\mathcal{D}_1) \subseteq \Gamma$ und eine Ableitung $\mathcal{D}_2$ mit $Hyp(\mathcal{D}_2) \subseteq \Delta \cup \{A\}$ geben. Nun erweitern wir jede offene Annahme A in $\mathcal{D}_2$ zu einer Teilableitung $\mathcal{D}_1$ von $\mathcal{D}_2$. Man erhält eine neue Ableitung $\mathcal{D}_2'$ mit $Hyp(\mathcal{D}_2') \subseteq (\Delta \cup \{A\}) \setminus \{A\} \cup Hyp(\mathcal{D}_1) = \Gamma \cup \Delta$. Somit gibt es eine Ableitung $\mathcal{D}_2'$ mit Endformel B und $Hyp(\mathcal{D}_2') \subseteq \Gamma \cup \Delta$. Also gilt $\Gamma, \Delta \vdash B$. QED

Bemerkung. Die Schnittregel drückt die Transitivität der Ableitbarkeitsrelation $\vdash$ aus. Die “weggeschnittene” Formel A bezeichnet man auch als *Schnittformel*.

Schnittformel

Definition 6.9 Eine Regel

ableitbare Regel

$$\frac{A_1 \quad \dots \quad A_n}{B}$$

mit Prämissenmenge $\Gamma = \{A_1, \dots, A_n\}$ und Konklusion B heißt *ableitbar*, falls $\Gamma \vdash B$.

Die Schnittregel rechtfertigt die Anwendung abgeleiteter Regeln.

Beispiele. (i) Im Beispiel auf S. 56 wurde $\neg\neg A \vdash A$ (doppelte Negationsbeseitigung) gezeigt. Somit ist

$$\frac{\neg\neg A}{A} \text{ (DNB)}$$

eine abgeleitete Regel.

Ableitungen der Form $\frac{\mathcal{D}_1}{\frac{\neg\neg A}{\mathcal{D}_2}} \frac{A}{A}$ können damit abgekürzt werden zu $\frac{\mathcal{D}_1}{A} \text{ (DNB)}$.

(ii) Aufgrund

$$\frac{\neg B \quad \frac{A \rightarrow B \quad [A]^1}{B} (\rightarrow B)}{\frac{\perp}{\neg A} (\rightarrow E)^1} (\rightarrow B)$$

gilt $\neg B, A \rightarrow B \vdash \neg A$. Somit ist

$$\frac{\neg B \quad A \rightarrow B}{\neg A} \text{ (modus tollens)}$$

eine abgeleitete Regel.

$\frac{\mathcal{D}_1}{\neg B, A \rightarrow B} \frac{\mathcal{D}_2}{\neg A}$ kann damit abgekürzt werden zu $\frac{\mathcal{D}_1 \quad \mathcal{D}_2}{\neg A} \text{ (modus tollens)}$.

(iii) Es gilt $B \vee A, \neg A \vee C \vdash B \vee C$ aufgrund

$$\frac{B \vee A \quad \frac{[B]^2}{B \vee C} (\vee E) \quad \frac{\neg A \vee C \quad \frac{\frac{[\neg A]^1 \quad [A]^2}{\perp} (\rightarrow B) \quad \frac{[C]^1}{B \vee C} (\vee E)}{B \vee C} (\perp)}{B \vee C} (\vee B)^2}{B \vee C} (\vee B)^1$$

Die Regel

$$\frac{B \vee A \quad \neg A \vee C}{B \vee C}$$

ist also ableitbar. (Man vergleiche diese Regel mit der in § 5 behandelten Resolutionsregel.)

Bemerkungen. (i) Annahmen, die bei einer Regelanwendung gelöscht werden dürfen, müssen nicht gelöscht werden; dies entspricht der strukturellen Operation der *Verdünnung*.

Verdünnung

Im Beweis von $\neg A \rightarrow (A \rightarrow B)$ ist Verdünnung wesentlich:

$$\frac{\frac{\frac{[\neg A]^2 \quad [A]^1}{\perp} (\perp) \quad \frac{A \rightarrow B}{\neg A \rightarrow (A \rightarrow B)} (\rightarrow E)^1}{\neg A \rightarrow (A \rightarrow B)} (\rightarrow E)^2}$$

Die Anwendung von $(\perp)$ würde die Löschung von Annahmen $\neg B$ erlauben, die hier aber nicht vorgenommen wird.

(ii) Mehrere Vorkommen derselben Annahme können bei einer einzigen Regelanwendung gelöscht werden; dies entspricht der strukturellen Operation der *Kontraktion*.

Kontraktion

Im Beweis des Satzes vom ausgeschlossenen Widerspruch $\neg(A \wedge \neg A)$ ist Kontraktion wesentlich:

$$\frac{\frac{[A \wedge \neg A]^1}{\neg A} (\wedge B) \quad \frac{[A \wedge \neg A]^1}{A} (\wedge B)}{\frac{\perp}{\neg(A \wedge \neg A)} (\rightarrow E)^1}$$

(iii) Neben Einführung und Beseitigung logischer Konstanten spielen also auch strukturelle Operationen bezüglich Annahmen eine Rolle. Im Kalkül des natürlichen Schließens sind diese Operationen nicht explizit durch eigene Regeln repräsentiert, sondern nur implizit durch die Festlegungen zur Löschung von Annahmen gegeben.

Es gibt auch sogenannte *Sequenzkalküle*, in denen eigene Regeln für strukturelle Operationen zur Verfügung stehen. (Siehe z. B. [Kleene, 2000](#) und [Troelstra & Schwichtenberg, 2000](#).)

6.2 Widerspruchsregel und reductio ad absurdum

Anwendungen der Widerspruchsregel bezeichnet man auch als *reductio ad absurdum*. Jedoch werden in der Literatur sowohl Argumentationen der Form

reductio ad absurdum

$$\frac{[A] \quad \vdots \quad \perp}{\neg A} \quad \text{als auch} \quad \frac{[\neg A] \quad \vdots \quad \perp}{A}$$

als *reductio ad absurdum* bezeichnet. Die linke Argumentation entspricht einer Anwendung der Implikationseinführung

$$\frac{[A] \quad \vdots \quad B}{A \rightarrow B} (\rightarrow E)$$

bei der B die Formel $\perp$ ist. Die rechte Argumentation stellt hingegen eine Anwendung der Widerspruchsregel $(\perp)$ dar. Würde man statt der Widerspruchsregel $(\perp)$ nur $(\rightarrow E)$

zulassen, erhielte man einen schwächeren Kalkül (die resultierende Logik bezeichnet man als *minimale Logik*), in dem die Widerspruchsregel keine ableitbare Regel ist. Man hat in diesem Fall zwar

$$\begin{array}{c} [\neg A] \\ \vdots \\ \frac{\perp}{\neg\neg A} (\rightarrow E) \end{array}$$

kann aber von $\neg\neg A$ nicht auf A schließen, da dies gerade die Widerspruchsregel erfordert (vgl. das obige Beispiel 20 auf S. 56 zu $\neg\neg A \vdash A$).

Beispiele. (i) Wir beweisen den Satz vom ausgeschlossenen Dritten (tertium non datur) $A \vee \neg A$, indem wir dessen Negation zum Widerspruch führen:

$$\frac{\frac{[\neg(A \vee \neg A)]^3 \quad \frac{[A]^1}{A \vee \neg A} (\vee E)}{\frac{\perp}{\neg A} (\rightarrow E)^1} (\rightarrow B) \quad \frac{[\neg(A \vee \neg A)]^3 \quad \frac{[\neg A]^2}{A \vee \neg A} (\vee E)}{\frac{\perp}{A} (\perp)^2} (\rightarrow B)}{\frac{\perp}{A \vee \neg A} (\perp)^3}$$

Hier wird in der linken Teildableitung unter der Annahme A ein Widerspruch $\perp$ abgeleitet, um dann mit Implikationseinführung ($\rightarrow E$) von $\perp$ auf $\neg A$ zu schließen. In der rechten Teildableitung wird der Widerspruch hingegen unter Annahme von $\neg A$ abgeleitet, um dann mit der Widerspruchsregel ($\perp$) auf A zu schließen. In beiden Teildableitungen wurde zusätzlich die Annahme $\neg(A \vee \neg A)$ gemacht. Beide Vorkommen der Annahme werden bei Anwendung der Widerspruchsregel im letzten Schritt gelöscht.

(ii) Die folgende Ableitung von $A \vee \neg A$ kommt mit nur einer Anwendung der Widerspruchsregel aus:

$$\frac{[\neg(A \vee \neg A)]^2 \quad \frac{[A]^1}{A \vee \neg A} (\vee E)}{\frac{\perp}{\neg A} (\rightarrow E)^1} (\rightarrow B) \quad \frac{[\neg(A \vee \neg A)]^2 \quad \frac{\perp}{A \vee \neg A} (\vee E)}{\frac{\perp}{A \vee \neg A} (\perp)^2} (\rightarrow B)$$

($\rightarrow E$) und ($\perp$) werden wie in (i) verwendet.

6.3 Korrektheit und Vollständigkeit

In der Semantik haben wir als zentrale Begriffe der Logik *logische Folgerung* ($\Gamma \models A$) und *Allgemeingültigkeit* ($\models A$) behandelt. Mit dem Kalkül des natürlichen Schließens haben wir als weitere zentrale Begriffe *Ableitbarkeit* ($\Gamma \vdash A$) und *Beweisbarkeit* ($\vdash A$) erhalten. Für das Verhältnis dieser Begriffe gilt Folgendes:

Theorem 6.10 (Korrektheit) Wenn $\Gamma \vdash A$, dann $\Gamma \models A$. Das heißt, der Kalkül NK ist korrekt.

Bemerkung. Um nachzuweisen, dass eine Formel A aus Annahmen Γ nicht ableitbar ist, genügt also der Nachweis, dass A aus Γ nicht logisch folgt. (Verwendung von Korrektheit per Kontraposition.)

Theorem 6.11 (Vollständigkeit) Wenn $\Gamma \models A$, dann $\Gamma \vdash A$. Das heißt, der Kalkül NK ist vollständig.

Für die Beweise von Korrektheit und Vollständigkeit siehe z. B. [van Dalen \(2013, § 2.5\)](#).

Korollar 6.12 Insbesondere gilt für $\Gamma = \emptyset$: $\vdash A$ genau dann, wenn $\models A$. Das heißt, eine Formel ist genau dann beweisbar, wenn sie allgemeingültig ist.

Bemerkung. Die Begriffe *logische Folgerung* und *Ableitbarkeit* sind zwar extensional gleich, da die Mengen der durch sie ausgezeichneten Gegenstände (geordnete Paare $\langle \Gamma, A \rangle$) gleich sind: $\{\langle \Gamma, A \rangle \mid \Gamma \models A\} = \{\langle \Gamma, A \rangle \mid \Gamma \vdash A\}$ (d. h. die Paare $\langle \Gamma, A \rangle$ mit der Eigenschaft $\Gamma \models A$ sind genau jene Paare $\langle \Gamma, A \rangle$ mit der Eigenschaft $\Gamma \vdash A$). Die Begriffe sind jedoch intensional verschieden, da die durch sie ausgezeichneten Gegenstände $\langle \Gamma, A \rangle$ auf verschiedene Weise gegeben sind. Gleiches gilt für *Allgemeingültigkeit* und *Beweisbarkeit*.

Korollar 6.13 Sei Γ eine beliebige (möglicherweise unendliche) Menge von Formeln, und A eine beliebige Formel. Falls $\Gamma \models A$ gilt, gibt es eine endliche Teilmenge Γ' von Γ , so dass $\Gamma' \models A$ gilt.

Beweis. Aufgrund Vollständigkeit gilt mit $\Gamma \models A$ auch $\Gamma \vdash A$. Nach Theorem 6.7 gibt es dann eine endliche Teilmenge $\Gamma' \subseteq \Gamma$, so dass $\Gamma' \vdash A$. Aufgrund Korrektheit gilt dann auch $\Gamma' \models A$. QED

Der Kalkül des natürlichen Schließens eignet sich in besonderer Weise zur Untersuchung der Struktur natürlichsprachlicher Argumente und mathematischer Beweise. Dies ist Gegenstand der *Beweistheorie*. Da der Kalkül korrekt und vollständig bezüglich der Semantik der Aussagenlogik ist, kann er zum Nachweis von logischen Folgerungen oder der Allgemeingültigkeit von Formeln verwendet werden.

Beweistheorie

7 Die formale Sprache der Quantorenlogik

7.1 Einleitende Bemerkungen

Ganz zu Beginn (vgl. S. 7) hatten wir als Beispiel für einen gültigen Schluss angegeben:

Alle Menschen sind sterblich.
Sokrates ist ein Mensch.
Sokrates ist sterblich.

Mit den Mitteln der Aussagenlogik erhält man eine Formalisierung

$$p, q \models r$$

wobei p : Alle Menschen sind sterblich; q : Sokrates ist ein Mensch; r : Sokrates ist sterblich. Obgleich die Formalisierung aussagenlogisch adäquat ist, gilt die Folgerungsbehauptung $p, q \models r$ nicht (ein Gegenbeispiel ist die Bewertung $\mathcal{I}$ mit $\llbracket p \rrbracket^{\mathcal{I}} = \llbracket q \rrbracket^{\mathcal{I}} = 1$ und $\llbracket r \rrbracket^{\mathcal{I}} = 0$). Aus Sicht der Aussagenlogik könnte man daraus folgern, dass auch das natürlichsprachliche Ausgangsargument nicht gültig ist.

Hingegen würde man aus Sicht des intuitiven Begriffs von Gültigkeit natürlichsprachlicher Argumente den Schluss ziehen, dass die formale Aussagenlogik nicht ausreicht, um alle intuitiv gültigen Argumente auch als formal gültige Schlüsse zu erfassen.

Der Grund dafür besteht in der mangelnden Ausdruckstärke der formalen Sprache der Aussagenlogik. Diese ermöglicht bestenfalls eine adäquate Formalisierung natürlichsprachlicher Aussagen auf der Ebene ganzer Aussagen, bzw. von mit Konnektiven aus Aussagen zusammengesetzten Aussagen. In unserem Beispiel kommen jedoch nur Aussagen ohne Konnektive vor, so dass als Formalisierung nur die angegebene bleibt. Eine Betrachtung des intuitiv gleichwertigen, gültigen Schlusses

Wenn alle Menschen sterblich sind, und Sokrates ein Mensch ist,
dann ist Sokrates sterblich.

führt auch nicht weiter. Zwar kommen nun Konnektive vor, doch die aussagenlogisch adäquate Formalisierung dieses Schlusses als

$$\models (p \wedge q) \rightarrow r$$

stellt eine ungültige Behauptung der Allgemeingültigkeit von $(p \wedge q) \rightarrow r$ dar.

Als eine erste Präzisierung des intuitiven Begriffs von Gültigkeit hatten wir vorgeschlagen:

Ein Schluss ist *gültig*, falls es keine Interpretation gibt, für die zwar die Prämissen wahr sind, die Konklusion aber falsch ist.

Zur besseren Vergleichbarkeit mit dem Begriff der aussagenlogischen Folgerung können wir diese Präzisierung auch so fassen:

Ein Schluss von Prämissen Γ auf eine Konklusion A ist *gültig* genau dann, wenn für alle Interpretationen $\mathcal{I}$ gilt: Wenn die Prämissen Γ unter $\mathcal{I}$ wahr sind, dann ist die Konklusion A unter $\mathcal{I}$ wahr.

Dabei wollen wir die zulässigen Interpretationen aber nicht auf aussagenlogische Bewertungen beschränken. Der Begriff ist dann offensichtlich weiter gefasst als der aussagenlogische Begriff der logischen Folgerung. Dies ermöglicht es uns, auch innerhalb

von aus aussagenlogischer Sicht unzerlegbaren Aussagen für bestimmte Komponenten verschiedene Interpretationen zuzulassen. Diese Komponenten hatten wir im Beispiel durch Zeichen P, Q, C repräsentiert (vgl. S. 8):

$$\frac{\begin{array}{l} \text{Alle } P \text{ sind } Q. \\ C \text{ ist } P. \end{array}}{C \text{ ist } Q.}$$

Eine Interpretation besteht nun in der Angabe von Bedeutungen für P, Q und C . Mit Blick auf den ursprünglichen Schluss stellt man fest, dass die möglichen Bedeutungen von P und Q von verschiedener Art wie die möglichen Bedeutungen von C sein müssen. Zum Beispiel steht Q anstelle des Eigenschaftswortes “sterblich”, während C anstelle des Namens (Designators) “Sokrates” steht. Für Q wird man deshalb als Bedeutungen nur Eigenschaften (d. h. Attribute von Gegenständen oder Relationen zwischen Gegenständen) zulassen, während man für C nur Gegenstände (Individuen) als Bedeutungen zulässt. Für die Kopula (“ist”, “sind”) haben wir keine verschiedenen Interpretationen zugelassen; das ist jedoch unwesentlich, da wir die Kopula immer auch als Bestandteil der Eigenschaftswörter auffassen können. Wesentlich ist aber, dass wir im Fall des Quantors “alle” keine verschiedenen Interpretationen zulassen, da wir diesen als logische Konstante behandeln wollen.

Um diese zusätzlichen Interpretationsmöglichkeiten im Rahmen der formalen Logik behandeln zu können, müssen wir die formale Sprache der Aussagenlogik in geeigneter Weise erweitern. Als Resultat werden wir die formale Sprache der Quantorenlogik erhalten.

Dazu führen wir zunächst Prädikatsymbole $P, Q, R, \dots$ und Individuenkonstanten $k, k', l, \dots$ ein. Prädikatsymbole werden anstelle von Eigenschaftswörtern verwendet. Individuenkonstanten übernehmen die Funktion von Designatoren, d. h. sie stehen anstelle von Ausdrücken, die einzelne Gegenstände bezeichnen. Um auch über beliebige Gegenstände etwas sagen zu können, werden wir Individuenvariablen $x, y, z, \dots$ verwenden. Schließlich führen wir den Allquantor $\forall$ und den Existenzquantor $\exists$ ein, mit denen quantifizierte Aussagen über beliebige Gegenstände gebildet werden können. Anhand von Beispielen werden diese Spracherweiterungen im Folgenden motiviert und deren Funktion erläutert.

7.1.1 Prädikatsymbole und Individuenkonstanten

Beispiel. Die Aussage

Sokrates ist ein Mensch

enthält den Namen “Sokrates”. Das, was übrigbleibt, wenn man den Namen entfernt, ist das Prädikat:

... ist ein Mensch

An die Stelle der Auslassung “...” können beliebige Namen treten. Für jede Einsetzung eines Namens erhält man eine neue Aussage; z. B.

Aristoteles ist ein Mensch

oder

Tübingen ist ein Mensch.

Man kann auch sagen, dass das Prädikat “. . . ist ein Mensch” eine Argumentstelle hat, was so geschrieben werden kann:

ist ein Mensch(...)

Setzt man den Namen “Sokrates” als Argument ein, erhält man den Ausdruck

ist ein Mensch(Sokrates)

In dieser Form wird deutlich gemacht, dass die Ausgangsaussage “Sokrates ist ein Mensch” so gelesen werden kann: Die Eigenschaft, ein Mensch zu sein, trifft auf das mit dem Namen “Sokrates” bezeichnete Individuum zu; bzw.: Sokrates hat die Eigenschaft, ein Mensch zu sein. Unter Verwendung des (1-stelligen) Prädikatsymbols P kann man entsprechend schreiben:

$P(...)$

Setzt man eine Individuenkonstante k ein, erhält man die Aussage

$P(k)$

welche besagt, dass die Eigenschaft P auf k zutrifft. Legt man zusätzlich fest, dass P für die Eigenschaft steht, ein Mensch zu sein, und dass k das Individuum Sokrates bezeichnet, so stellt $P(k)$ eine Formalisierung der Ausgangsaussage “Sokrates ist ein Mensch” dar. Entsprechend kann man festlegen, dass das (1-stellige) Prädikatsymbol Q für die Eigenschaft steht, sterblich zu sein. Dann ist $Q(k)$ eine Formalisierung der Aussage “Sokrates ist sterblich”.

Prädikate können auch mehr als eine Argumentstelle haben. Zu deren Formalisierung verwendet man entsprechend mehrstellige Prädikatsymbole.

Beispiel. Die Aussage

Der Neckar fließt durch Tübingen

enthält die beiden Namen “Neckar” und “Tübingen”. Entfernt man beide, bleibt das 2-stellige Prädikat

... fließt durch ...

bzw.

fließt durch(..., ...)

übrig. Bei mehrstelligen Prädikaten kommt es auf die Reihenfolge der Argumente an. Sei R ein 2-stelliges Prädikatsymbol (bzw. Relationszeichen), das für die Eigenschaft “fließt durch” stehe, und bezeichne die Individuenkonstante k den Neckar und l die Stadt Tübingen, so erhält man mit

$R(k, l)$

die der natürlichsprachlichen Ausgangsaussage “Der Neckar fließt durch Tübingen” entsprechende wahre Aussage. Hingegen ist

$R(l, k)$

eine falsche Aussage (sie entspricht “Tübingen fließt durch den Neckar”).

Bemerkung. Individuen müssen nicht immer durch Eigennamen wie “Tübingen” bezeichnet sein. Zum Beispiel wird in der Aussage

Der Verfasser der *Begriffsschrift* lehrte in Jena

ein Individuum (G. Frege) durch den Designator “der Verfasser der *Begriffsschrift*” bezeichnet. Als weiterer Designator kommt der Name “Jena” vor, der die Stadt Jena bezeichnet.

Durch Weglassung von Designatoren können hier folgende Prädikate übrigbleiben:

(i) Der Verfasser der *Begriffsschrift* lehrte in ...

Man erhält ein 1-stelliges Prädikat der Form $S(\dots)$.

(ii) ... lehrte in Jena

Man erhält ein anderes 1-stelliges Prädikat $T(\dots)$ derselben Form.

(iii) ... lehrte in ...

Man erhält ein 2-stelliges Prädikat der Form $R(\dots, \dots)$.

Lässt man auch die “leere” Weglassung zu, so bleibt als Prädikat die Aussage selbst:

(iv) Der Verfasser der *Begriffsschrift* lehrte in Jena.

Man erhält ein 0-stelliges Prädikat der Form P (d. h. ein Prädikat ohne Argumentstelle).

Wir lassen auch 0-stellige Prädikatsymbole wie P zu. Sie haben dieselbe Funktion wie die Aussagesymbole $p, q, r, \dots$ in der Aussagenlogik.

7.1.2 Individuenvariablen

Um auch Aussagen über beliebige Individuen machen zu können, erweitern wir unsere Sprache nun um Individuenvariablen $x, y, z, \dots$.

Beispiel. Statt der Aussage

Alle Menschen sind sterblich

kann man auch sagen

Für alle Individuen gilt: Wenn ein Individuum ein Mensch ist, dann ist dieses Individuum sterblich.

“Ein Individuum” und “dieses Individuum” stehen hier für ein beliebiges, aber in beiden Fällen identisches Individuum. Unter Verwendung der Individuenvariable x kann man sagen:

Für alle x gilt: Wenn x ein Mensch ist, dann ist x sterblich.

Auf das “dieses” kann also verzichtet werden, da wir für beide Fälle dieselbe Variable x verwendet haben. Unter Verwendung der 1-stelligen Prädikatsymbole P und Q können wir dies weiter umschreiben zu

Für alle x gilt: Wenn $P(x)$, dann $Q(x)$.

Den hinteren Teil “Wenn $P(x)$, dann $Q(x)$ ” können wir aussagenlogisch formalisieren:

Für alle x gilt: $(P(x) \rightarrow Q(x))$

wobei $P(x)$: x ist ein Mensch; $Q(x)$: x ist sterblich.

7.1.3 Quantoren

Als weitere Spracherweiterung führen wir jetzt das Symbol $\forall$ für den Allquantor “für alle” ein. Damit können wir die Formalisierung aus dem vorherigen Beispiel fortsetzen:

Beispiel. Aus “Für alle x gilt: $(P(x) \rightarrow Q(x))$ ” wird $\forall x(P(x) \rightarrow Q(x))$.

Eine Formalisierung des Schlusses

Alle Menschen sind sterblich.
Sokrates ist ein Mensch.
Sokrates ist sterblich.

kann damit wie folgt lauten: $\forall x(P(x) \rightarrow Q(x)), P(k) \models Q(k)$, wobei

$P(x)$: x ist ein Mensch

$Q(x)$: x ist sterblich

k : Sokrates

Diese Formalisierung weist eine wesentlich feinere Struktur als die eingangs betrachtete aussagenlogische Formalisierung auf. Dadurch sind weitaus vielfältigere Interpretationen möglich, die wir in der Semantik der Quantorenlogik untersuchen werden.

Als weiteren Quantor betrachten wir den Existenzquantor “es gibt (mindestens) ein”, für den wir das Symbol $\exists$ verwenden.

Beispiel. Die Aussage

Es gibt einen Menschen

kann schrittweise über

Es gibt ein x , für das gilt: x ist ein Mensch

Es gibt ein x , für das gilt: $P(x)$

zu

$\exists x P(x)$

umgeformt werden, wobei P wieder für die Eigenschaft, ein Mensch zu sein, stehe.

Beispiel. Eine Aussage wie

Einige Menschen sind Logiker

kann schrittweise umgeschrieben werden zu

Es gibt ein x , für das gilt: x ist ein Mensch, und x ist ein Logiker

Es gibt ein x , für das gilt: $(P(x) \wedge S(x))$

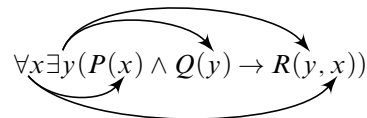
wobei das Prädikatsymbol S für die Eigenschaft stehe, ein Logiker zu sein. Unter Verwendung des Symbols $\exists$ wird daraus schließlich

$\exists x(P(x) \wedge S(x))$

als Formalisierung der Ausgangsaussage “Einige Menschen sind Logiker”.

- Bemerkungen.** (i) Eine Formalisierung $\exists x(P(x) \rightarrow S(x))$ wäre nicht adäquat, da diese aufgrund der Bedeutung der Implikation schon dann wahr wäre, wenn es ein Individuum gibt, für das $P(x)$ nicht gilt.
- (ii) Bei der Formalisierung von “Alle Menschen sind sterblich” wäre hingegen $\forall x(P(x) \wedge Q(x))$ nicht adäquat, da Letzteres nur dann wahr ist, wenn alle Individuen Menschen und sterblich sind.
- (iii) Bei den Quantoren handelt es sich um *variablenbindende* Operatoren. Sie unterscheiden sich dadurch wesentlich von den aussagenlogischen Operatoren, den Konnektiven. Konnektive verbinden Formeln zu einer neuen Formel, wobei erstere dann unmittelbare Teilformeln letzterer sind. In Formeln der Form $\forall xA$ und $\exists xA$ stellen die Quantoren $\forall x$ und $\exists x$ zusätzlich eine Verbindung zu allen Vorkommen der Variablen x in A her, d. h. einschließlich der Vorkommen von x in Teilformeln der unmittelbaren Teilformel A .

Zur Veranschaulichung:



Die Pfeile drücken die Bindung der Variablen durch den jeweiligen Quantor aus.

Die durch die Spracherweiterungen gewonnene Ausdrucksstärke kommt insbesondere dann zum Tragen, wenn mehrere Quantoren in einer Aussage verwendet werden müssen.

Beispiele. Es sei

k : Sokrates

$S(x, y)$: x sieht y

Die Variablen x und y stehen hier für beliebige Individuenvariablen $x, y, z, \dots$. Wichtig ist, dass x und y prinzipiell verschiedene Variablen sein können, aber nicht müssen; d. h. es darf anstelle von x und y dasselbe oder Verschiedenes stehen. Hingegen dürfte bei einer Festlegung wie “ $S(x, y)$: x sieht x ” an beiden Argumentstellen von S nur dasselbe stehen.

Den Individuenbereich (d. h. den Bereich der Gegenstände, über den quantifiziert wird) beschränken wir auf Menschen.

Formalisierungen:

- (i) Es gibt jemanden, der Sokrates sieht. $\exists xS(x, k)$
- (ii) Sokrates sieht jemanden. (Es gibt jemanden, den Sokrates sieht.) $\exists xS(k, x)$
- (iii) Jeder sieht jeden. $\forall x\forall yS(x, y)$
- (iv) Jeder sieht jemanden. $\forall x\exists yS(x, y)$
- (v) Einer sieht jeden. (Es gibt jemanden, der jeden sieht.) $\exists x\forall yS(x, y)$
- (vi) Jeder wird von jemandem gesehen. $\forall y\exists xS(x, y)$
- (vii) Es gibt jemanden, der von jedem gesehen wird. $\exists y\forall xS(x, y)$

- (viii) Es gibt jemanden, der jemanden sieht. $\exists x \exists y S(x, y)$
- (ix) Jeder sieht sich selbst. $\forall x S(x, x)$
- (x) Es gibt jemanden, der sich selbst sieht. $\exists x S(x, x)$

Die vorgenommenen Spracherweiterungen bilden die Grundlage der Quantorenlogik. Da hierbei neben Quantoren Prädikate eine herausragende Rolle spielen, bezeichnet man die Quantorenlogik auch als *Prädikatenlogik*. Eine weitere Bezeichnung ist *Logik erster Stufe* (engl. *first-order logic*), die zum Ausdruck bringt, dass ausschließlich über Individuen eines Gegenstandsbereichs quantifiziert wird, nicht aber über höherstufige Objekte wie Relationen, Eigenschaften von Relationen usw.

Nach diesen einleitenden Bemerkungen legen wir nun die formale Sprache der Quantorenlogik noch etwas präziser fest.

7.2 Syntax der Quantorenlogik

Definition 7.1 Wir erweitern das *Alphabet* um folgende Zeichen:

Alphabet

- (i) *Quantoren*: $\forall$ (Allquantor, “für alle”) und $\exists$ (Existenzquantor, “es gibt”).
- (ii) *Terme*:
 - *Individuenvariablen* (kurz: *Variablen*): $x, y, z, u, v, w, x_1, x_2, \dots$
 - *Individuenkonstanten* (kurz: *Konstanten*): $k, k', k_1, k_2, \dots, l, \dots$
- (iii) *Relationszeichen* (bzw. *Prädikatsymbole*): $P, Q, R, \dots$
 (Gegebenenfalls mit fester Stelligkeit, so dass z. B. R^1 ein 1-stelliges und R^2 ein 2-stelliges Relationszeichen ist. 0-stellige Relationszeichen entsprechen Aussagesymbolen.)
- (iv) Komma: ,

Bemerkung. Als metasprachliche Variablen für Terme verwenden wir im Folgenden Buchstaben $r, s, t, t_1, t_2, \dots$.

Wir verzichten aber auf gesonderte metasprachliche Variablen wie z. B. $\xi, \xi_1, \xi_2, \dots$ für objektsprachliche Variablen $x, y, z, \dots$ und auf metasprachliche Variablen wie z. B. $\alpha, \alpha_1, \alpha_2, \dots$ für objektsprachliche Konstanten $k, k', l, \dots$; ebenso verzichten wir auf metasprachliche Variablen wie z. B. $\Pi, \Pi_1, \Pi_2, \dots$ für objektsprachliche Relationszeichen $P, Q, R, \dots$.

Stattdessen verwenden wir objektsprachliche Variablen und Relationszeichen gelegentlich auch als metasprachliche Variablen für objektsprachliche Zeichen der jeweiligen Art. Durch den Kontext wird immer klar sein, was gemeint ist.

Zum Beispiel wird in der folgenden Definition von Formeln das Zeichen R als metasprachliche Variable für beliebige objektsprachliche Relationszeichen verwendet, und das Zeichen x wird als metasprachliche Variable für beliebige objektsprachliche Individuenvariablen verwendet.

Definition 7.2 Wir erweitern die Definition von *Formeln* wie folgt:

Formel

- (i) Ist R ein n -stelliges Relationszeichen und sind $t_1, \dots, t_n$ Terme, dann ist $R(t_1, \dots, t_n)$ eine Formel, die wir *atomare Formel* bzw. *Atom* nennen.
 (Abkürzend kann man auch $Rt_1 \dots t_n$ schreiben. Im Fall 0-stelliger Relationszeichen R schreiben wir statt $R()$ stets R .)

(ii) Ist A eine Formel und ist x eine Variable, dann sind $\forall xA$ und $\exists xA$ ebenfalls Formeln.

Hierdurch ist die *Sprache der Quantorenlogik* als Menge aller (quantorenlogischen) Formeln definiert, die wir auch wieder mit FORMEL bezeichnen.

Sprache der Quantorenlogik

Bemerkung. Zur Sprache der Quantorenlogik können noch mit Funktionszeichen f gebildete Funktionsterme $f(t_1, \dots, t_n)$ und die Identität $=$ als logische Konstante hinzugefügt werden. Auf die Behandlung der Identität verzichten wir in dieser Einführung; auf Funktionsterme werden wir in § 9.3 eingehen.

Definition 7.3 Wir ergänzen die Definition von (unmittelbare) Teilformel um die Klausel: A ist *unmittelbare Teilformel* von $\forall xA$ und $\exists xA$.

unmittelbare Teilformel

Bemerkungen. (i) Die Quantoren binden stärker als die übrigen logischen Konstanten. Klammerersparnis entsprechend.

Bindungsstärke

(ii) Das *Vorkommen* eines Zeichens sei nur anhand eines Beispiels erläutert: In der Formel $\forall x(P(x) \wedge Q(x, y))$ kommt x dreimal und y einmal vor.

Vorkommen

Definition 7.4 Wir definieren *freie* bzw. *gebundene* Variablenvorkommen:

frei/gebunden

(i) Jedes Variablenvorkommen in einer atomaren Formel ist frei.

(ii) Die freien bzw. gebundenen Variablenvorkommen in A und B sind auch freie bzw. gebundene Variablenvorkommen in $\neg A$, $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$ und $(A \leftrightarrow B)$.

(iii) Jedes freie Vorkommen von x in A ist ein gebundenes Vorkommen von x in $\forall xA$ und $\exists xA$.

Alle anderen freien bzw. gebundenen Variablenvorkommen in A sind auch freie bzw. gebundene Variablenvorkommen in $\forall xA$ und $\exists xA$.

Die freien Vorkommen von x in A liegen in den Formeln $\forall xA$ bzw. $\exists xA$ im *Wirkungsbereich* des Allquantors $\forall x$ bzw. des Existenzquantors $\exists x$.

Wirkungsbereich

Es sei $FV(A)$ die Menge der Variablen, die in einer Formel A frei vorkommen. Ist $FV(A) \neq \emptyset$, so heißt die Formel A *offen*. Ist $FV(A) = \emptyset$, dann heißt A *geschlossen* oder *Aussage*.

offen/geschlossen

Aussage

Bemerkung. Bei metasprachlichen Variablen für Formeln A schreiben wir z. B. $A(x)$, um auszudrücken, dass x in A frei vorkommen kann. (Der Ausdruck $A(x)$ steht also i. A. nicht für ein Atom wie z. B. $P(x)$, sondern für eine beliebige Formel.)

Beispiele. Für die Variablenvorkommen gilt:

(i) $R(x, y, y)$ – x, y sind frei.

(ii) $\exists z(Q(z, k) \rightarrow \neg R(x, y, y))$ – x, y sind frei; z ist gebunden. Der Wirkungsbereich des Existenzquantors $\exists z$ ist die unmittelbare Teilformel $(Q(z, k) \rightarrow \neg R(x, y, y))$, in der z frei vorkommt.

(In $\exists z$ ist z weder frei noch gebunden.)

Es ist $FV = \{x, y\}$. Die Formel ist also offen.

(iii) $\exists y \overbrace{(P(x, y) \vee \underbrace{\exists z (Q(z, k) \rightarrow \neg R(x, y, y))}_{\text{Wirkungsbereich von } \exists z})}_{\text{Wirkungsbereich von } \exists y})$ – x ist frei; y, z sind gebunden.

(In $\exists y$ ist y weder frei noch gebunden.)

Es ist $FV = \{x\}$. Die Formel ist also offen.

(iv) $\forall x P(x) \wedge Q(x)$ – Das erste Vorkommen von x ist weder frei noch gebunden; das zweite ist gebunden, das dritte frei (nur $P(x)$ steht im Wirkungsbereich des Allquantors $\forall x$, nicht aber $Q(x)$).

Es ist $FV = \{x\}$. Die Formel ist also offen.

(v) $\forall x \exists y (P(x) \vee Q(y, x))$ – x, y sind gebunden.

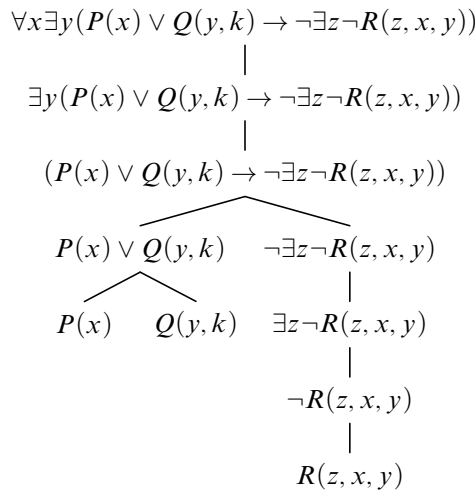
Es ist $FV = \emptyset$. Es handelt sich also um eine geschlossene Formel, d. h. um eine Aussage.

Definition 7.5 Die Definition von *Strukturbaum* erweitern wir um die folgende Klausel: *Strukturbaum*

(iv) Die Strukturbäume $Sb(\forall x A)$ bzw. $Sb(\exists x A)$ von Formeln $\forall x A$ bzw. $\exists x A$ sind

$$\begin{array}{ccc} \forall x A & & \exists x A \\ | & \text{bzw.} & | \\ Sb(A) & & Sb(A) \end{array}$$

Beispiel. Der Strukturbaum der Formel $\forall x \exists y (P(x) \vee Q(y, k) \rightarrow \neg \exists z \neg R(z, x, y))$ ist:



Die Formel enthält die drei atomaren Formeln $P(x)$, $Q(y, k)$ und $R(z, x, y)$.

7.3 Strukturelle Induktion

Wir haben Formeln durch eine induktive Definition eingeführt, die im Wesentlichen aus zwei Teilen besteht. Wir haben zunächst für bestimmte Grundelemente, nämlich für Aussagesymbole, $\top$ und $\perp$, festgelegt, dass diese Formeln sind. Dann haben wir unter der Annahme, dass gewisse Gegenstände Formeln oder bestimmte Zeichen sind, gesagt,

welche unter Verwendung dieser und ggf. weiterer Gegenstände zusammengesetzten Gegenstände ebenfalls Formeln sind. Der Gegenstandsbereich war hierbei durch die Zeichen eines vorher festgelegten Alphabets gegeben.

Im Allgemeinen haben *induktive Definitionen* die folgende Form:

induktive Definition

- (i) *Basisklauseln (Induktionsbasis):* Der Gegenstand g_1 ist ein X ; der Gegenstand g_2 ist ein X ; ...; der Gegenstand g_n ist ein X .
- (ii) *Induktive Klauseln:* Wenn der Gegenstand g ein X ist, dann ist auch g' ein X .

Dabei wird zusätzlich vorausgesetzt, dass nichts sonst ein X sei. Diese Abschlussbedingung kann auch explizit vorgenommen werden, z. B. indem man die induktive Definition wie folgt formuliert: Sei X die kleinste Menge, so dass (i) und (ii) gilt.

Die Klauseln in induktiven Definitionen stellen Erzeugungsregeln für die Elemente eines Gegenstandsbereichs X dar, und da nichts sonst Element von X sein soll, ist X definiert. Jeder induktiven Definition entspricht ein Induktionsprinzip, das in einem Beweis per Induktion über dem Aufbau der induktiv definierten Gegenstände verwendet werden kann, um Eigenschaften E nachzuweisen, die allen Gegenständen in X zukommen. Das jeweilige *Induktionsprinzip* besagt:

Induktionsprinzip

Sei E eine Eigenschaft, so dass für $g \in X$ gilt:

- (1) Für die in den Basisklauseln angegebenen Gegenstände $g_1, \dots, g_n$ gilt $E(g_1), \dots, E(g_n)$, d. h. diese Gegenstände haben die Eigenschaft E .
- (2) Wenn $E(g)$, dann auch $E(g')$. (Analog zu den induktiven Klauseln.)

Dann gilt $E(g)$ für jedes $g \in X$.

Ein Beweis per *struktureller Induktion* (kurz: *Induktionsbeweis*) hat dann die Form:

*Beweis per
struktureller
Induktion*

Induktionsanfang: Seien $g_1, \dots, g_n \in X$. Zu zeigen: $E(g_1), \dots, E(g_n)$.

Induktionsannahme: Es gelte $E(g)$.

Induktionsschritt: Unter der Induktionsannahme ist zu zeigen, dass $E(g')$ gilt. (Hier ist ggf. eine Fallunterscheidung bezüglich der induktiven Klauseln vorzunehmen.)

Man spricht von *struktureller Induktion*, da der Beweis per Induktion über der Struktur von Gegenständen erfolgt (die durch eine induktive Definition gegeben sind). In Induktionsbeweisen verwendet man auch die Abkürzungen IA für Induktionsanfang, IH für Induktionshypothese (= Induktionsannahme; oder IV für Induktionsvoraussetzung) und IS für Induktionsschritt.

Als Beispiel betrachten wir die induktive Definition der Menge FORMEL, wobei wir uns auf den aussagenlogischen Teil beschränken. Wir verwenden dabei $\circ$ als metasprachliche Variable für die 2-stelligen Konnektive $\wedge, \vee, \rightarrow$ und $\leftrightarrow$.

Induktive Definition Sei FORMEL die kleinste Menge, so dass Folgendes gilt:

- (i) (a) Sei $A \in \text{AS}$. Dann ist $A \in \text{FORMEL}$.
- (b) $\top, \perp \in \text{FORMEL}$.
- (ii) (a) Wenn $A \in \text{FORMEL}$, dann $\neg A \in \text{FORMEL}$.
- (b) Wenn $A, B \in \text{FORMEL}$, dann $(A \circ B) \in \text{FORMEL}$.

In (i) stehen die Basisklauseln für die atomaren Formeln, in (ii) die induktiven Klauseln für die komplexen Formeln (für jedes Konnektiv eine Klausel). Nun beweisen wir das zugehörige Induktionsprinzip.

Induktionsprinzip Sei E eine Eigenschaft, so dass Folgendes gilt:

- (1) (a) Sei $A \in \text{AS}$. Dann gilt $E(A)$.
 (b) $E(\top)$ und $E(\perp)$.
- (2) (a) Wenn $E(A)$, dann $E(\neg A)$.
 (b) Wenn $E(A)$ und $E(B)$, dann $E((A \circ B))$.

Dann gilt $E(A)$ für alle $A \in \text{FORMEL}$.

Beweis. Wir betrachten die Menge X der Formeln mit der Eigenschaft E , d. h. $X := \{A \in \text{FORMEL} \mid E(A)\}$. Offensichtlich gilt $X \subseteq \text{FORMEL}$. Dann erfüllt X die Eigenschaften (i)(a,b) und (ii)(a,b) der induktiven Definition von FORMEL. Da FORMEL die kleinste Menge mit diesen Eigenschaften ist, gilt $\text{FORMEL} \subseteq X$. Damit ist $\text{FORMEL} = X$, d. h. $E(A)$ gilt für alle $A \in \text{FORMEL}$. QED

Als einfaches Beispiel für die Anwendung dieses Induktionsprinzips zeigen wir per Induktion über dem Formelaufbau die folgende Aussage.

Theorem Sei E die Eigenschaft "hat gerade Anzahl an Klammern". Es gilt $E(A)$ für alle $A \in \text{FORMEL}$.

Beweis. Per Induktion über dem Aufbau von A . Es sei $\#A$ die Anzahl der Klammern von A .

Induktionsanfang: Zeige die Aussage für atomare Formeln.

Sei $A \in \text{AS}$. Es ist $\#A = 0$ und $0 = 2 \cdot 0$. Also $E(A)$.

Es ist $\#\top = \#\perp = 0$. Also $E(\top)$ und $E(\perp)$.

Induktionsannahme: $E(B)$ und $E(C)$ gilt für $B, C \in \text{FORMEL}$, wobei $\#B = 2n$ und $\#C = 2m$ sei.

Induktionsschritt: Zeige die Aussage für komplexe Formeln.

Betrachte $\neg B$: Es ist $\#\neg B = \#B + 0 \stackrel{\text{IH}}{=} 2n + 0$ gerade. Also $E(\neg B)$.

Betrachte $(B \circ C)$: Es ist $\#(B \circ C) = \#B + \#C + 2 \stackrel{\text{IH}}{=} 2n + 2m + 2 = 2(n + m + 1)$, d. h. $E((B \circ C))$.

(Hier konnten wir vier Fälle in einem behandeln, da keine Fallunterscheidung für die 2-stelligen Konnektive nötig war. Die Verwendung der Induktionsannahme haben wir jeweils durch "IH" angezeigt.)

Somit gilt $E(A)$ für alle $A \in \text{FORMEL}$. QED

7.4 Mathematische Induktion

Neben struktureller Induktion können Induktionsbeweise auch unter Verwendung eines geeigneten *Induktionsmaßes* geführt werden, nach dem die Gegenstände, über denen Induktion geführt werden soll, wie die natürlichen Zahlen geordnet werden können. Dies garantiert die Gültigkeit eines *Induktionsprinzips* für diese Gegenstände, denn für die natürlichen Zahlen gilt das *Prinzip der mathematischen Induktion*:

Induktionsmaß

Angenommen, für eine Eigenschaft E gelten die Bedingungen

- (1) $E(0)$, d. h. 0 hat die Eigenschaft E , und

(2) für jede natürliche Zahl n : wenn $E(n)$, dann $E(n + 1)$.

Dann hat jede natürliche Zahl die Eigenschaft E .

Als (nicht-logische) Regel hat das Prinzip der mathematischen Induktion im natürlichen Schließen die Form

$$\frac{n \in \mathbb{N} \quad E(0) \quad E(n+1)}{\forall n E(n)} \quad [E(n)]$$

wobei n in keiner Annahme außer $E(n)$ vorkommen darf, von der die Prämisse $E(n + 1)$ abhängt.

Bei einem Beweis per Induktion über dem Formelaufbau kann man ausnutzen, dass Formeln gemäß ihrer Komplexität wie die natürlichen Zahlen geordnet werden können. Die Komplexität einer Formel kann man z. B. als Anzahl der in ihr vorkommenden logischen Konstanten festlegen.

Definition 7.6 Der *Grad* g einer Formel ist induktiv definiert wie folgt:

Grad

- (i) $g(A) := 0$, falls A atomar.
- (ii) $g(*A) := g(A) + 1$, falls $*$ eine 1-stellige logische Konstante $\neg$, $\forall$ oder $\exists$ ist.
- (iii) $g(A \circ B) := g(A) + g(B) + 1$, falls $\circ$ ein 2-stelliges Konnektiv $\wedge$, $\vee$, $\rightarrow$ oder $\leftrightarrow$ ist.

Verschiedene Formeln können denselben Grad haben; z. B. ist

$$\begin{aligned} g(\forall x P(x) \wedge \perp) &= g(\forall x P(x)) + g(\perp) + 1 \\ &= g(\forall x P(x)) + 0 + 1 \\ &= g(P(x)) + 1 + 0 + 1 \\ &= 0 + 1 + 0 + 1 = 2 = g(p \wedge q \rightarrow r) \end{aligned}$$

Eine Induktion über dem Grad einer Formel erfordert deshalb eine Fallunterscheidung, bei der Bedingung (2) des Prinzips der mathematischen Induktion für jede logische Konstante nachgewiesen werden muss.

Bei einer Induktion über dem Formelaufbau zum Nachweis einer Eigenschaft E für alle Formeln nimmt man an, dass die unmittelbaren Teilformeln einer beliebigen Formel A die Eigenschaft E haben. Dann weist man nach, dass unter dieser Annahme auch die aus diesen Teilformeln zusammengesetzte Formel A die Eigenschaft E hat (vgl. Bedingung (2)). Dieser Nachweis ist der *Induktionsschritt*, die dabei gemachte Annahme ist die *Induktionsannahme*. Zusätzlich muss gezeigt werden, dass auch die Formeln kleinster Komplexität (also die atomaren Formeln) die Eigenschaft E haben (vgl. Bedingung (1) im Prinzip der mathematischen Induktion). Dieser Nachweis ist der *Induktionsanfang*. Sind Induktionsanfang und Induktionsschritt (für jede logische Konstante) gezeigt, kann unter Verwendung des Induktionsprinzips geschlossen werden, dass alle Formeln die Eigenschaft E haben.

Zur Veranschaulichung: Seien $UT(A)$ die unmittelbaren Teilformeln von A . Dann hat ein solcher Induktionsbeweis im natürlichen Schließen die Form

$$\frac{\begin{array}{cc} \text{IA} & \text{IS} \\ E(A), \text{ für alle Atome} & E(A) \end{array}}{\text{Für alle Formeln } A: E(A)} \quad \frac{[E(UT(A))]^k}{(\text{Induktionsprinzip})^k}$$

Hierbei steht IA für den Beweis des Induktionsanfangs, und IS steht für die Ableitung von $E(A)$ aus Annahmen $E(UT(A))$, d. h. für den Induktionsschritt.

Literatur

- D. Makinson, *Sets, Logic and Maths for Computing. Third Edition*, Springer, 2020, § 4.6.2.
- R. M. Smullyan, *Logical Labyrinths*, A K Peters, 2009, § 15.

8 Semantik der Quantorenlogik

Zur Bestimmung des Wahrheitswerts einer aussagenlogischen Formel genügt die Betrachtung von Bewertungen $\mathcal{I}$, d. h. von Interpretationen der in der Formel vorkommenden Aussagesymbole durch Wahrheitswerte. Bei quantorenlogischen Formeln müssen zusätzlich Interpretationen der in der Formel vorkommenden Individuenkonstanten und Relationszeichen betrachtet werden. Individuenkonstanten werden durch Gegenstände interpretiert, und n -stellige Relationszeichen werden durch n -stellige Relationen zwischen Gegenständen interpretiert (bzw. im Fall 1-stelliger Relationszeichen durch Eigenschaften von Gegenständen, und im Fall 0-stelliger Relationszeichen durch Wahrheitswerte). Interpretationen ordnen diesen nicht-logischen Zeichen der formalen Sprache also Bedeutungen zu wie folgt:

<i>formale Sprache</i>	<i>zugeordnete Bedeutung, dargestellt als</i>	
Aussagesymbol	Wahrheitswert	leere bzw. nichtleere Menge
0-stelliges Relationszeichen	Wahrheitswert	leere bzw. nichtleere Menge
Individuenkonstante	Gegenstand	Element einer Menge
1-stelliges Relationszeichen	1-stellige Relation	Menge von Gegenständen
2-stelliges Relationszeichen	2-stellige Relation	Menge geordneter Paare
3-stelliges Relationszeichen	3-stellige Relation	Menge von Tripeln
$\vdots$	$\vdots$	$\vdots$
n -stelliges Relationszeichen	n -stellige Relation	Menge von n -Tupeln

Für die Interpretationen können beliebige Mengen von Gegenständen zugelassen werden. Insbesondere kann auch eine Beschränkung auf einen bestimmten Gegenstandsbereich vorgenommen werden. Wählt man zum Beispiel den Bereich der natürlichen Zahlen (d. h. die Menge $\mathbb{N} = \{0, 1, 2, 3, \dots\}$), so muss jeder Konstante eine natürliche Zahl, und jedem Relationszeichen eine Relation zwischen natürlichen Zahlen zugeordnet werden. Dabei sind jedoch nur Interpretationen für jene Konstanten und Relationszeichen anzugeben, die auch tatsächlich verwendet werden. Im Folgenden betrachten wir deshalb quantorenlogische formale Sprachen $\mathcal{L}$, in denen (neben den logischen Konstanten und Hilfszeichen) nur bestimmte Individuenkonstanten und Relationszeichen vorkommen. Wir schreiben $A \in \mathcal{L}$ bzw. $\Gamma \subseteq \mathcal{L}$, um auszudrücken, dass die Formel A bzw. jede Formel in Γ ein Ausdruck in $\mathcal{L}$ ist. Die Kombination von Gegenstandsbereich und entsprechender Interpretation der nicht-logischen Zeichen einer Sprache $\mathcal{L}$ wird als (*mengentheoretische*) *Struktur* für $\mathcal{L}$ bezeichnet.

8.1 Strukturen

Definition 8.1 Eine *Struktur* für eine Sprache $\mathcal{L}$ ist ein geordnetes Paar $\mathfrak{M} = \langle M, \mathcal{I} \rangle$, wobei M eine nichtleere Menge ist und $\mathcal{I}$ eine Funktion, die jedem Symbol aus $\mathcal{L}$ eine *Interpretation* zuordnet wie folgt:

- (i) $\mathcal{I}(k) \in M$ für Konstanten $k \in \mathcal{L}$,
- (ii) $\mathcal{I}(R) \subseteq M^n$, falls $R \in \mathcal{L}$ ein n -stelliges Relationszeichen ist ($n \geq 1$),
- (iii) $\mathcal{I}(R) \in \{0, 1\}$, falls $R \in \mathcal{L}$ ein 0-stelliges Relationszeichen, d. h. ein Aussagesymbol ist, wobei die Wahrheitswerte 0 und 1 wie folgt definiert werden können: $0 := \emptyset$ und $1 := M$.

Die Menge M heißt *Gegenstandsbereich* (oder auch *Individuenbereich*, *Universum*).

Struktur

Interpretation

Gegenstandsbereich

Gemäß Klausel (i) ist jeder Konstante $k \in \mathcal{L}$ ein Gegenstand aus M zuzuordnen.

Beispiel. Ist M die Menge der Menschen und $k_1, l \in \mathcal{L}$, dann ist $\mathcal{I}$ mit

$$\begin{aligned}\mathcal{I}(k_1) &= \text{Sokrates} \\ \mathcal{I}(l) &= \text{Aristoteles}\end{aligned}$$

eine Interpretation der Konstanten in $\mathcal{L}$.

Gemäß Klausel (ii) ist für $n \geq 1$ jedem n -stelligen Relationszeichen $R \in \mathcal{L}$ eine n -stellige Relation auf M^n zuzuordnen. (M^n ist die Menge aller n -Tupel, die mit Elementen aus M gebildet werden können.)

Beispiel. Ist $M = \{\text{Sokrates, Alexander der Große, Aristoteles, Platon}\}$, dann kann die Interpretation $\mathcal{I}(R)$ des 2-stelligen Relationszeichens R als Relation "... ist Schüler von ..." wie folgt als Menge von geordneten Paaren notiert werden (d. h. durch explizite Angabe aller Elemente; hier: Paare $\langle m_1, m_2 \rangle$):

$$\mathcal{I}(R) = \{\langle \text{Aristoteles, Platon} \rangle, \langle \text{Platon, Sokrates} \rangle, \langle \text{Alexander der Große, Aristoteles} \rangle\}$$

Alternativ kann die Menge durch Angabe der Eigenschaft, die auf alle Elemente zutreffen muss, notiert werden:

$$\mathcal{I}(R) = \{\langle m_1, m_2 \rangle \mid m_1 \text{ ist Schüler von } m_2\}$$

wobei m_1 und m_2 für Gegenstände in M stehen.

Allgemein werden Interpretationen n -stelliger Relationszeichen für $n \geq 2$ als Mengen von n -Tupeln $\langle m_1, m_2, \dots, m_n \rangle$ angegeben, d. h. je nach Stelligkeit als Paare $\langle m_1, m_2 \rangle$, Tripel $\langle m_1, m_2, m_3 \rangle$, Quadrupel $\langle m_1, m_2, m_3, m_4 \rangle$ usw. Ein 1-Tupel $\langle m_1 \rangle$ sei dasselbe wie der Gegenstand m_1 . Da Paare $\langle m_1, m_2 \rangle$ als Mengen von Mengen $\{\{m_1\}, \{m_1, m_2\}\}$ definiert werden können, und n -Tupel $\langle m_1, m_2, \dots, m_n \rangle$ auf geordnete Paare $\langle \langle m_1, m_2, \dots, m_{n-1} \rangle, m_n \rangle$ zurückgeführt werden können, müssen letztlich nur Mengen vorausgesetzt werden.

Interpretationen 0-stelliger Relationszeichen $R \in \mathcal{L}$ gemäß Klausel (iii) entsprechen aussagenlogischen Bewertungen.

Die mengentheoretische Festlegung der beiden Wahrheitswerte durch $0 := \emptyset$ und $1 := M$ in Strukturen $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ ist möglich, da für jede Struktur $M \neq \emptyset$ gefordert wird (und stets $\emptyset \subset M$ gilt). Auch für 0-stellige Relationszeichen R (bzw. Aussagesymbole) gilt damit immer $\mathcal{I}(R) \subseteq M$, so dass sich Interpretationen $\mathcal{I}$ in jedem Fall ausschließlich auf den Gegenstandsbereich M beziehen.

Beispiele. Die Sprache $\mathcal{L}_s$ enthalte lediglich die beiden 1-stelligen Relationszeichen P und Q , sowie die Konstante k .

(i) Sei der Gegenstandsbereich $M = \{\text{Neckar, Tübingen}\}$ und die Interpretation $\mathcal{I}$:

$$\begin{aligned}\mathcal{I}(P) &= \{\text{Neckar}\} \\ \mathcal{I}(Q) &= \{\text{Neckar, Tübingen}\} \\ \mathcal{I}(k) &= \text{Neckar}\end{aligned}$$

Dann ist $\langle M, \mathcal{I} \rangle$ eine Struktur für $\mathcal{L}_s$.

- (ii) Nun umfasse der Gegenstandsbereich M_s beliebige Gegenstände, und die Interpretation $\mathcal{I}_s$ sei gegeben durch:

$$\mathcal{I}_s(P) = \{m \mid m \text{ ist ein Mensch}\}$$

$$\mathcal{I}_s(Q) = \{m \mid m \text{ ist sterblich}\}$$

$$\mathcal{I}_s(k) = \text{Sokrates}$$

d. h. durch $\mathcal{I}_s$ wird dem Prädikatsymbol P die *Eigenschaft* (1-stellige Relation) zugeordnet, ein Mensch zu sein, dem Prädikatsymbol Q wird die *Eigenschaft* zugeordnet, sterblich zu sein, und der Individuenkonstante (d. h. dem *Term*) k wird der *Gegenstand* Sokrates zugeordnet.

Dann ist $\mathfrak{G} = \langle M_s, \mathcal{I}_s \rangle$ ebenfalls eine Struktur für $\mathcal{L}_s$.

Beispiel. $\mathcal{L}$ enthalte lediglich die Konstanten $k', k_1, k_2, \dots$ und ein 2-stelliges Relationszeichen Q .

- Der Gegenstandsbereich M sei $\mathbb{N} = \{0, 1, 2, 3, \dots\}$,
- die Interpretation der Konstanten sei durch $\mathcal{I}(k') = 5$ und $\mathcal{I}(k_n) = n$ gegeben (d. h. $\mathcal{I}$ ordnet der Individuenkonstante (d. h. dem *Term*) $k' \in \mathcal{L}$ als *Gegenstand* die Zahl $5 \in \mathbb{N}$ zu; den Individuenkonstanten (d. h. den *Termen*) $k_n \in \mathcal{L}$ ordnet $\mathcal{I}$ als *Gegenstände* die Zahlen $n \in \mathbb{N}$ zu, z. B. $\mathcal{I}(k_7) = 7$),
- und die Interpretation von Q sei $\mathcal{I}(Q) = \{\langle n, n' \rangle \mid n < n'\} \subseteq \mathbb{N}^2$ (d. h. das 2-stellige Relationszeichen wird durch $\mathcal{I}$ als die 2-stellige kleiner-als-Relation auf $\mathbb{N}^2$ interpretiert; es ist z. B. $\langle 5, 3 \rangle \notin \mathcal{I}(Q)$, da $5 \not< 3$).

Dann ist $\mathfrak{N} = \langle \mathbb{N}, \mathcal{I} \rangle$ eine Struktur für die Sprache $\mathcal{L}$.

8.2 Variablenbelegungen

Durch Strukturen $\langle M, \mathcal{I} \rangle$ für eine Sprache $\mathcal{L}$ können wir schon jeder variablenfreien Formel $R(k_1, \dots, k_n)$ in $\mathcal{L}$ eine Bedeutung zuordnen. Um beliebigen Formeln $R(t_1, \dots, t_n)$, in denen auch Variablen vorkommen können, ebenfalls eine Bedeutung zuordnen zu können, benötigen wir zusätzlich Variablenbelegungen. Diese ordnen jeder Variable einen Gegenstand in M zu.

Definition 8.2 Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine Struktur. Eine *Variablenbelegung* in $\mathfrak{M}$ ist eine Funktion v , die jeder Variable x einen Wert $v(x) \in M$ zuordnet. (Also $v : \text{VAR} \rightarrow M$, wobei VAR die Menge aller Variablen sei.)

Variablenbelegung

Unterscheidet sich eine Variablenbelegung v' höchstens durch die Zuordnung $v'(x) = m \in M$ von einer Variablenbelegung v , so schreiben wir $v[x \mapsto m]$ für v' . Es gilt

$$v[x \mapsto m](y) = \begin{cases} m & \text{falls } x = y, \\ v(y) & \text{falls } x \neq y. \end{cases}$$

Die Variablenbelegung $v' = v[x \mapsto m]$ heißt *Variante* von v . Da v' sich von v nur durch die Belegung von x unterscheidet, bezeichnet man v' auch als x -*Variante* von v .

Variante

Beispiele. Wir betrachten die Struktur $\mathfrak{N} = \langle \mathbb{N}, \mathcal{I} \rangle$ für $\mathcal{L}$ aus dem vorherigen Beispiel.

(i) Dann ist durch

$$v(x) = 4, \quad v(y) = 7, \quad v(z) = 0, \quad v(z_1) = 7, \quad \dots$$

eine Variablenbelegung v in $\mathfrak{M}$ gegeben.

(ii) Nun ändern wir v zu einer neuen Variablenbelegung v' ab, indem wir der Variable y statt 7 den Gegenstand 11 zuordnen:

$$v'(x) = 4, \quad v'(y) = 11, \quad v'(z) = 0, \quad v'(z_1) = 7, \quad \dots$$

Die Variablenbelegung v' ist eine Variante (genauer: eine y -Variante) von v , da $v' = v[y \mapsto 11]$.

(iii) Des Weiteren ist z. B. $v[z \mapsto 0]$, also die Variablenbelegung v selbst, eine (triviale) z -Variante von v .

Interpretationen $\mathcal{I}(k)$ für Konstanten $k \in \mathcal{L}$ und Variablenbelegungen v (für alle Variablen) können zu einer Termbelegung zusammengefasst werden:

Definition 8.3 Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine Struktur für $\mathcal{L}$, v eine Variablenbelegung in $\mathfrak{M}$ und TERM die Menge aller Terme. Dann ist die *Termbelegung* als Funktion

Termbelegung

$$\llbracket \cdot \rrbracket_v^{\mathfrak{M}} : \text{TERM} \rightarrow M$$

definiert durch:

$$\begin{aligned} \llbracket x \rrbracket_v^{\mathfrak{M}} &:= v(x), \text{ falls } x \text{ eine Variable ist,} \\ \llbracket k \rrbracket_v^{\mathfrak{M}} &:= \mathcal{I}(k), \text{ falls } k \text{ eine Konstante in } \mathcal{L} \text{ ist.} \end{aligned}$$

Für Konstanten $k \notin \mathcal{L}$ ist $\llbracket k \rrbracket_v^{\mathfrak{M}}$ nicht definiert.

Beispiel. Für die Struktur $\mathfrak{M} = \langle \mathbb{N}, \mathcal{I} \rangle$ und die Variablenbelegung v' erhalten wir beispielsweise die Termbelegungen $\llbracket k_{11} \rrbracket_{v'}^{\mathfrak{M}} = 11$ und $\llbracket y \rrbracket_{v'}^{\mathfrak{M}} = 11$. Für die Variablenbelegung v ist $\llbracket y \rrbracket_v^{\mathfrak{M}} = \llbracket z_1 \rrbracket_v^{\mathfrak{M}} = \llbracket k_7 \rrbracket_v^{\mathfrak{M}} = 7$. Für die Konstante $l \notin \mathcal{L}$ ist $\llbracket l \rrbracket_v^{\mathfrak{M}}$ nicht definiert.

8.3 Gültigkeit

Nachdem wir nun in der Lage sind, sowohl den nicht-logischen Zeichen (Individuenkonstanten und Relationszeichen) als auch den Variablen Bedeutungen zuzuordnen, können wir schließlich auch die Bedeutung der logischen Konstanten (insbesondere von $\forall$ und $\exists$) festlegen. Die *Semantik der Quantorenlogik* geben wir durch die induktive Definition einer Relation $\mathfrak{M} \models_v A$ zwischen Strukturen $\mathfrak{M}$ und Formeln A unter Variablenbelegungen v an:

Semantik der Quantorenlogik

Definition 8.4 Sei $\mathfrak{M}$ eine Struktur $\langle M, \mathcal{I} \rangle$, v eine Variablenbelegung und $\llbracket \cdot \rrbracket_v^{\mathfrak{M}}$ eine Termbelegung. Wir definieren die Relation $\mathfrak{M} \models_v A$ (“ A gilt in der Struktur $\mathfrak{M}$ unter der Variablenbelegung v in $\mathfrak{M}$ ”) wie folgt:

A gilt in $\mathfrak{M}$ unter v

Für n -stellige Relationszeichen R mit $n \geq 1$:

$$\mathfrak{M} \models_v R(t_1, \dots, t_n) :\iff \langle \llbracket t_1 \rrbracket_v^{\mathfrak{M}}, \dots, \llbracket t_n \rrbracket_v^{\mathfrak{M}} \rangle \in \mathcal{I}(R)$$

(Entsprechend sei $\mathfrak{M} \not\models_v R(t_1, \dots, t_n) :\iff \langle \llbracket t_1 \rrbracket_v^{\mathfrak{M}}, \dots, \llbracket t_n \rrbracket_v^{\mathfrak{M}} \rangle \notin \mathcal{I}(R)$.)

Für 0-stellige Relationszeichen R (bzw. Aussagesymbole):

$$\mathfrak{M} \models_v R :\iff \mathcal{I}(R) = 1$$

(Entsprechend sei $\mathfrak{M} \not\models_v R :\iff \mathcal{I}(R) = 0$.)

Für die mit logischen Konstanten gebildeten Formeln:

$$\mathfrak{M} \models_v \top$$

$$\mathfrak{M} \not\models_v \perp$$

$$\mathfrak{M} \models_v \neg A :\iff \mathfrak{M} \not\models_v A$$

$$\mathfrak{M} \models_v A \wedge B :\iff \mathfrak{M} \models_v A \text{ und } \mathfrak{M} \models_v B$$

$$\mathfrak{M} \models_v A \vee B :\iff \mathfrak{M} \models_v A \text{ oder } \mathfrak{M} \models_v B$$

$$\mathfrak{M} \models_v A \rightarrow B :\iff \mathfrak{M} \not\models_v A \text{ oder } \mathfrak{M} \models_v B$$

$$\iff \text{Wenn } \mathfrak{M} \models_v A, \text{ dann } \mathfrak{M} \models_v B$$

$$\mathfrak{M} \models_v A \leftrightarrow B :\iff \mathfrak{M} \models_v A \text{ genau dann, wenn } \mathfrak{M} \models_v B$$

$$\mathfrak{M} \models_v \forall x A(x) :\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} A(x)$$

$$\mathfrak{M} \models_v \exists x A(x) :\iff \text{Es gibt eine } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} A(x)$$

Bemerkung. Die Klauseln für die Quantoren können auch wie folgt formuliert werden (wobei M der Gegenstandsbereich von $\mathfrak{M}$ ist):

$$\mathfrak{M} \models_v \forall x A(x) :\iff \text{Für jedes } m \in M \text{ gilt } \mathfrak{M} \models_{v[x \mapsto m]} A(x)$$

$$\mathfrak{M} \models_v \exists x A(x) :\iff \text{Es gibt ein } m \in M, \text{ so dass } \mathfrak{M} \models_{v[x \mapsto m]} A(x)$$

Definition 8.5 Sei $A \in \mathcal{L}$ und $\mathfrak{M}$ eine Struktur für $\mathcal{L}$.

(i) Wir definieren *Gültigkeit in einer Struktur* (“ A gilt in der Struktur $\mathfrak{M}$ ”):

Gültigkeit in einer Struktur

$$\mathfrak{M} \models A :\iff \mathfrak{M} \models_v A \text{ für alle Variablenbelegungen } v \text{ in } \mathfrak{M}.$$

(ii) Falls A in $\mathfrak{M}$ gilt (d. h. falls $\mathfrak{M} \models A$), heißt die Struktur $\mathfrak{M}$ *Modell von A* .

Modell von A

(iii) Falls $\mathfrak{M}$ kein Modell von A ist, schreiben wir auch $\mathfrak{M} \not\models A$. In diesem Fall heißt $\mathfrak{M}$ *Gegenmodell* (oder *Gegenbeispiel*) zu A .

Gegenmodell

Bemerkungen. (i) Im Fall aussagenlogischer Formeln A waren lediglich Aussagesymbole durch Wahrheitswerte zu interpretieren. Dies geschah durch Bewertungen $\mathcal{I}$, für die wir die Relation $\mathcal{I} \models A$ (“ A gilt in $\mathcal{I}$ ”, “ A ist wahr unter $\mathcal{I}$ ”) definiert hatten.

(ii) Die Relation $\mathfrak{M} \models A$ (“ A gilt in der Struktur $\mathfrak{M}$ ”) ist der dazu analoge quantorenlogische Begriff. Statt “ A gilt in $\mathfrak{M}$ ” sagt man auch “ A ist in $\mathfrak{M}$ wahr”, bzw. statt “ A gilt nicht in $\mathfrak{M}$ ” ($\mathfrak{M} \not\models A$) auch “ A ist in $\mathfrak{M}$ falsch”.

(iii) Diese Art von Semantik bezeichnet man als *modelltheoretische Semantik*. (Sie geht zurück auf [Tarski \(1935\)](#).)

Beispiel. Die in Beispiel (ii) auf S. 78 angegebene Struktur $\mathfrak{G} = \langle M_s, \mathcal{I}_s \rangle$ mit uneingeschränktem Gegenstandsbereich M_s und der durch

$$\mathcal{I}_s(P) = \{m \mid m \text{ ist ein Mensch}\}$$

$$\mathcal{I}_s(Q) = \{m \mid m \text{ ist sterblich}\}$$

$$\mathcal{I}_s(k) = \text{Sokrates}$$

gegebenen Interpretation $\mathcal{I}_s$ für die Sprache $\mathcal{L}_s$ ist ein Modell der Formeln $P(k)$, $Q(k)$ und $\forall x(P(x) \rightarrow Q(x))$.

Sei v eine beliebige Variablenbelegung in $\mathfrak{S}$.

– Für $P(k)$ gilt

$$\mathfrak{S} \models_v P(k) \iff \llbracket k \rrbracket_v^{\mathfrak{S}} \in \mathcal{I}_s(P)$$

Die rechte Seite gilt, da $\llbracket k \rrbracket_v^{\mathfrak{S}} = \text{Sokrates}$, und $\text{Sokrates} \in \mathcal{I}_s(P)$. Damit gilt auch $\mathfrak{S} \models_v P(k)$. Da v beliebig ist, gilt $\mathfrak{S} \models P(k)$, d. h. $\mathfrak{S}$ ist ein Modell von $P(k)$.

– Entsprechend gilt auch $\mathfrak{S} \models Q(k)$, da $\text{Sokrates} \in \mathcal{I}_s(Q)$.

– Für $\forall x(P(x) \rightarrow Q(x))$ gilt

$$\mathfrak{S} \models_v \forall x(P(x) \rightarrow Q(x))$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{S}: \mathfrak{S} \models_{v'} P(x) \rightarrow Q(x)$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{S}: \text{Wenn } \mathfrak{S} \models_{v'} P(x), \text{ dann } \mathfrak{S} \models_{v'} Q(x)$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{S}: \text{Wenn } \llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(P), \text{ dann } \llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(Q)$$

Sei $v' = v[x \mapsto m]$ für ein $m \in \mathcal{I}_s(P)$. In diesem Fall gilt sowohl $\mathfrak{S} \models_{v'} P(x)$, da $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} = m \in \mathcal{I}_s(P)$, als auch $\mathfrak{S} \models_{v'} Q(x)$, da $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} = m \in \mathcal{I}_s(Q)$. Letzteres folgt wegen $\mathcal{I}_s(P) \subset \mathcal{I}_s(Q)$; dies sieht man den mittels Eigenschaften notierten Mengen $\mathcal{I}_s(P)$ und $\mathcal{I}_s(Q)$ nicht an, aber wir gehen davon aus, dass alle Menschen sterblich sind. Somit gilt: Wenn $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(P)$, dann $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(Q)$.

Ist hingegen $v' = v[x \mapsto m]$ für ein $m \notin \mathcal{I}_s(P)$, so ist $\mathfrak{S} \not\models_{v'} P(x)$. Damit gilt ebenfalls: Wenn $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(P)$, dann $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(Q)$.

Somit gilt für jede x -Variante v' von v in $\mathfrak{S}$: Wenn $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(P)$, dann $\llbracket x \rrbracket_{v'}^{\mathfrak{S}} \in \mathcal{I}_s(Q)$.

Da v beliebig ist, gilt $\mathfrak{S} \models \forall x(P(x) \rightarrow Q(x))$. Die Struktur $\mathfrak{S}$ ist also auch ein Modell von $\forall x(P(x) \rightarrow Q(x))$.

Beispiele. Wir betrachten die Formel $\forall x \exists y(S(y) \wedge R(x, y))$.

(i) Die Struktur $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ mit

$$M = \{m_1, m_2\}$$

$$\mathcal{I}(S) = \{m_1\}$$

$$\mathcal{I}(R) = \{\langle m_1, m_1 \rangle, \langle m_2, m_1 \rangle\}$$

ist ein Modell der Formel.

Sei v eine beliebige Variablenbelegung in $\mathfrak{M}$. Es gilt

$$\mathfrak{M} \models_v \forall x \exists y(S(y) \wedge R(x, y))$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} \exists y(S(y) \wedge R(x, y))$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M} \text{ gibt es eine } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M}: \\ \mathfrak{M} \models_{v''} S(y) \wedge R(x, y)$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M} \text{ gibt es eine } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M}: \\ \mathfrak{M} \models_{v''} S(y) \text{ und } \mathfrak{M} \models_{v''} R(x, y)$$

$$\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M} \text{ gibt es eine } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M}: \\ \llbracket y \rrbracket_{v''}^{\mathfrak{M}} \in \mathcal{I}(S) \text{ und } \langle \llbracket x \rrbracket_{v'}^{\mathfrak{M}}, \llbracket y \rrbracket_{v''}^{\mathfrak{M}} \rangle \in \mathcal{I}(R)$$

Da M zwei Gegenstände enthält, gibt es zwei x -Varianten v' von v in $\mathfrak{M}$: $v[x \mapsto m_1]$ und $v[x \mapsto m_2]$. Mit der y -Variante $v'' = v'[y \mapsto m_1]$ von v' gilt für die erste x -Variante $\llbracket y \rrbracket_{v''}^{\mathfrak{M}} = m_1 \in \mathcal{I}(S)$ und $\langle \llbracket x \rrbracket_{v''}^{\mathfrak{M}}, \llbracket y \rrbracket_{v''}^{\mathfrak{M}} \rangle = \langle m_1, m_1 \rangle \in \mathcal{I}(R)$; für die zweite x -Variante gilt $\llbracket y \rrbracket_{v''}^{\mathfrak{M}} = m_1 \in \mathcal{I}(S)$ und $\langle \llbracket x \rrbracket_{v''}^{\mathfrak{M}}, \llbracket y \rrbracket_{v''}^{\mathfrak{M}} \rangle = \langle m_2, m_1 \rangle \in \mathcal{I}(R)$.

Damit ist für jede x -Variante v' von v in $\mathfrak{M}$ eine y -Variante v'' von v' in $\mathfrak{M}$ gefunden, so dass $\mathfrak{M} \models_{v''} S(y) \wedge R(x, y)$. Also gilt $\mathfrak{M} \models_v \forall x \exists y (S(y) \wedge R(x, y))$. Da v beliebig ist, gilt $\mathfrak{M} \models \forall x \exists y (S(y) \wedge R(x, y))$, d. h. $\mathfrak{M}$ ist ein Modell von $\forall x \exists y (S(y) \wedge R(x, y))$.

- (ii) Hingegen ist die Struktur $\mathfrak{M}' = \langle \{m_1, m_2\}, \mathcal{I}' \rangle$ mit

$$\begin{aligned}\mathcal{I}'(S) &= \{m_2\} \\ \mathcal{I}'(R) &= \{\langle m_1, m_1 \rangle, \langle m_2, m_1 \rangle\}\end{aligned}$$

ein Gegenmodell für die Formel.

Denn für die x -Variante $v' = v[x \mapsto m_1]$ in $\mathfrak{M}'$ gibt es nun keine y -Variante v'' von v' in $\mathfrak{M}'$, so dass $\mathfrak{M}' \models_{v''} S(y)$ und $\mathfrak{M}' \models_{v''} R(x, y)$: Für $v'' = v'[y \mapsto m_1]$ ist $\mathfrak{M}' \not\models_{v''} S(y)$, da $\llbracket y \rrbracket_{v''}^{\mathfrak{M}'} = m_1 \notin \mathcal{I}'(S)$; und für $v'' = v'[y \mapsto m_2]$ ist $\mathfrak{M}' \not\models_{v''} R(x, y)$, da $\langle \llbracket x \rrbracket_{v''}^{\mathfrak{M}'}, \llbracket y \rrbracket_{v''}^{\mathfrak{M}'} \rangle = \langle m_1, m_2 \rangle \notin \mathcal{I}'(R)$.

Somit gibt es eine Variablenbelegung v in $\mathfrak{M}'$ (nämlich v mit $v(x) = m_1$ und $v(y)$ beliebig), für die $\mathfrak{M}' \not\models_v \forall x \exists y (S(y) \wedge R(x, y))$; also $\mathfrak{M}' \not\models \forall x \exists y (S(y) \wedge R(x, y))$, d. h. $\mathfrak{M}'$ ist ein Gegenmodell zu $\forall x \exists y (S(y) \wedge R(x, y))$.

Die Berücksichtigung von Variablenbelegungen ist nötig, um auch bei offenen Formeln eine Aussage über deren Gültigkeit machen zu können.

Sei z. B. $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ mit $M = \{m, m'\}$ und $\mathcal{I}(P) = \{m, m'\}$ für das 1-stellige Relationszeichen P . Die Formel $\forall x P(x)$ soll in $\mathfrak{M}$ unter einer Variablenbelegung v genau dann gelten, wenn die offene Formel $P(x)$ für alle Gegenstände in M gilt. Um Letzteres auszudrücken, können wir jedoch *nicht* einfach schreiben, dass $P(m)$ und $P(m')$ gilt. Denn m und m' sind keine Zeichen unserer Objektsprache, sondern Gegenstände in M ; die Ausdrücke “ $P(m)$ ” und “ $P(m')$ ” sind somit keine Formeln. Hingegen können wir sagen, dass die offene Formel $P(x)$ für jede Zuordnung eines Gegenstands in M zur freien Variable x , also hier für die beiden x -Varianten $v' = v[x \mapsto m]$ und $v'' = v[x \mapsto m']$, gilt. Es ist sowohl $\mathfrak{M} \models_{v'} P(x)$ als auch $\mathfrak{M} \models_{v''} P(x)$, da sowohl $v'(x) = m \in \mathcal{I}(P)$ als auch $v''(x) = m' \in \mathcal{I}(P)$.

Das begriffliche Hilfsmittel der Varianten benötigen wir auch, um eine Aussage über die Gültigkeit von z. B. $\forall x Q(x, y)$ in $\mathfrak{M}$ unter einer Variablenbelegung v machen zu können. Sei wieder $M = \{m, m'\}$, $\mathcal{I}(Q) = \{\langle m, m \rangle, \langle m', m \rangle\}$ für das 2-stellige Relationszeichen Q , und $v(y) = m$. Sowohl für die x -Variante $v' = v[x \mapsto m]$ als auch für die x -Variante $v'' = v[x \mapsto m']$ gilt $Q(x, y)$, denn $\langle \llbracket x \rrbracket_{v'}^{\mathfrak{M}}, \llbracket y \rrbracket_{v'}^{\mathfrak{M}} \rangle = \langle m, m \rangle \in \mathcal{I}(Q)$ und $\langle \llbracket x \rrbracket_{v''}^{\mathfrak{M}}, \llbracket y \rrbracket_{v''}^{\mathfrak{M}} \rangle = \langle m', m \rangle \in \mathcal{I}(Q)$. Die offene Formel $Q(x, y)$ gilt also für jede x -Variante von v in $\mathfrak{M}$. Somit gilt auch $\forall x Q(x, y)$ in $\mathfrak{M}$ unter v .

Für die Gültigkeit einer Formel A in einer Struktur $\mathfrak{M}$ unter einer Variablenbelegung v sind jedoch nur jene Variablen relevant, die in A frei vorkommen. Denn es gilt:

Theorem 8.6 (Koinzidenz) *Sei A eine beliebige Formel und $\mathfrak{M}$ eine beliebige Struktur, und seien v und v' zwei Variablenbelegungen, so dass für jede freie Variable x in A gilt: $v(x) = v'(x)$. Dann gilt $\mathfrak{M} \models_v A \iff \mathfrak{M} \models_{v'} A$.*

8.4 Erfüllbarkeit und Allgemeingültigkeit

Ganz analog zur Aussagenlogik interessieren wir uns in der Quantorenlogik für folgende Eigenschaften von Formeln:

Definition 8.7 Sei $A \in \mathcal{L}$ und $\mathfrak{M}$ eine Struktur für $\mathcal{L}$.

- (i) A heißt *erfüllbar* (oder *konsistent*), falls es eine Struktur $\mathfrak{M}$ und eine Variablenbelegung v gibt, so dass $\mathfrak{M} \models_v A$. *Erfüllbarkeit*
 Ist A eine geschlossene Formel, dann heißt A *erfüllbar*, falls es eine Struktur $\mathfrak{M}$ gibt, so dass $\mathfrak{M} \models A$, d. h. falls A ein Modell hat.
 Andernfalls heißt A *unerfüllbar* (oder *inkonsistent* oder *kontradiktorisch*).
- (ii) *Allgemeingültigkeit*: $\models A : \iff \mathfrak{M} \models A$ für alle Strukturen $\mathfrak{M}$. *Allgemeingültigkeit*
 (Mit anderen Worten: Eine Formel $A \in \mathcal{L}$ ist *allgemeingültig* genau dann, wenn jede Struktur $\mathfrak{M}$ für $\mathcal{L}$ ein Modell für A ist.)
- (iii) A heißt *kontingent*, falls A erfüllbar, aber nicht allgemeingültig ist.
 Für geschlossene Formeln ist dies genau dann der Fall, wenn es sowohl ein Modell für A als auch ein Gegenmodell zu A gibt.

Beispiele. Wir betrachten wieder die Struktur $\mathfrak{N} = \langle \mathbb{N}, \mathcal{I} \rangle$ mit

$$\begin{aligned} \mathcal{I}(k') &= 5 \\ \mathcal{I}(k_n) &= n \quad (\text{d. h. } \mathcal{I}(k_0) = 0, \mathcal{I}(k_1) = 1, \mathcal{I}(k_2) = 2, \dots) \\ \mathcal{I}(Q) &= \{ \langle n, n' \rangle \mid n < n' \} \subseteq \mathbb{N}^2 \end{aligned}$$

- (i) Es ist $\mathfrak{N} \not\models_v Q(k', k_3)$, da $\langle \llbracket k' \rrbracket_v^{\mathfrak{N}}, \llbracket k_3 \rrbracket_v^{\mathfrak{N}} \rangle = \langle 5, 3 \rangle \notin \mathcal{I}(Q)$. (Es ist $5 \not< 3$.)
 Die Variablenbelegung v spielt keine Rolle, da $Q(k', k_3)$ eine geschlossene Formel ist. Folglich gilt $\mathfrak{N} \not\models_v Q(k', k_3)$ für alle v . Also $\mathfrak{N} \not\models Q(k', k_3)$.
 Die Formel $Q(k', k_3)$ ist jedoch erfüllbar, da z. B. die Struktur $\mathfrak{N}' = \langle \mathbb{N}, \mathcal{I}' \rangle$ mit $\mathcal{I}'(Q) = \{ \langle n, n' \rangle \mid n > n' \} \subseteq \mathbb{N}^2$, und $\mathcal{I}'(k') = 5$, $\mathcal{I}'(k_3) = 3$, ein Modell von $Q(k', k_3)$ ist, d. h. $\mathfrak{N}' \models Q(k', k_3)$.
- (ii) Sei v' die Variablenbelegung mit $v'(y) = 11$ und $v'(z_1) = 7$. Dann ist $\mathfrak{N} \models_{v'} Q(z_1, y)$, da $\langle \llbracket z_1 \rrbracket_{v'}^{\mathfrak{N}}, \llbracket y \rrbracket_{v'}^{\mathfrak{N}} \rangle = \langle 7, 11 \rangle \in \mathcal{I}(Q)$. (Es ist $7 < 11$.)
 Die Formel $Q(z_1, y)$ ist also erfüllbar. Die Struktur $\mathfrak{N}$ ist jedoch kein Modell von $Q(z_1, y)$, da es eine Variablenbelegung v in $\mathfrak{N}$ gibt (z. B. $v(y) = v(z_1) = 7$), so dass $\mathfrak{N} \not\models_v Q(z_1, y)$.
- (iii) Für beliebige Variablenbelegungen v gilt $\mathfrak{N} \models_v \forall x \exists y Q(x, y)$.

$$\begin{aligned} \mathfrak{N} \models_v \forall x \exists y Q(x, y) &\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{N}: \mathfrak{N} \models_{v'} \exists y Q(x, y) \\ &\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{N} \\ &\quad \text{gibt es eine } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{N}: \mathfrak{N} \models_{v''} Q(x, y) \end{aligned}$$

Sei v beliebig, und $v' = v[x \mapsto n]$ für ein beliebiges $n \in \mathbb{N}$. Setze $v'' = v'[y \mapsto n+1]$. Dann ist $\mathfrak{N} \models_{v''} Q(x, y)$, da $\langle \llbracket x \rrbracket_{v''}^{\mathfrak{N}}, \llbracket y \rrbracket_{v''}^{\mathfrak{N}} \rangle = \langle n, n+1 \rangle \in \mathcal{I}(Q)$.

Da v beliebig, gilt $\mathfrak{N} \models \forall x \exists y Q(x, y)$. Die Struktur $\mathfrak{N}$ ist also ein Modell von $\forall x \exists y Q(x, y)$.

Keine der betrachteten Formeln ist allgemeingültig. Ein Gegenmodell ist z. B. die Struktur $\mathfrak{M} = \langle \{m\}, \mathcal{I}'' \rangle$ mit $\mathcal{I}''(Q) = \{ \langle m, m' \rangle \mid m \neq m' \}$ (und $\mathcal{I}''(k) = m$ für alle Konstanten $k \in \mathcal{L}$).

Da es für $Q(k', k_3)$ und $\forall x \exists y Q(x, y)$ mit $\mathfrak{M}$ auch ein Modell gibt, sind diese Formeln kontingent. Für $Q(z_1, y)$ lässt sich ebenfalls ein Modell angeben: Für $\mathfrak{M}^* = \langle \{m\}, \mathcal{I}^* \rangle$ mit $\mathcal{I}^*(Q) = \{ \langle m, m \rangle \}$, und $\mathcal{I}^*(k) = m$ für alle Konstanten $k \in \mathcal{L}$, gilt für jede Variablenbelegung v in $\mathfrak{M}^*$ (es gibt nur eine, mit $v(z_1) = v(y) = m$) $\mathfrak{M}^* \models_v Q(z_1, y)$. Damit ist auch $Q(z_1, y)$ kontingent.

Für das Verhältnis von Quantorenlogik und Aussagenlogik gilt offensichtlich:

Theorem 8.8 (konservative Erweiterung) *Die Quantorenlogik ist eine konservative Erweiterung der Aussagenlogik. Das heißt, jede allgemeingültige aussagenlogische Formel ist auch in der Quantorenlogik allgemeingültig.*

Theorem 8.9 (Permanenz) *Hat eine quantorenlogische Formel A die Form einer allgemeingültigen aussagenlogischen Formel B , so ist auch A allgemeingültig. Diese Eigenschaft bezeichnet man als Permanenz der Aussagenlogik in der Quantorenlogik.*

Beispiel. Sei A die quantorenlogische Formel

$$((\forall x \neg P(x) \rightarrow \exists y \forall z R(y, z)) \rightarrow \forall x \neg P(x)) \rightarrow \forall x \neg P(x)$$

Die Formel A hat die Form der aussagenlogisch allgemeingültigen Formel

$$((p \rightarrow q) \rightarrow p) \rightarrow p$$

Damit ist auch A allgemeingültig.

Bemerkung. Für beliebige (offene oder geschlossene) Formeln A und B gilt:

$$\begin{aligned} \mathfrak{M} \models \neg A &\implies \mathfrak{M} \not\models A \\ \mathfrak{M} \models A \rightarrow B &\implies \text{Wenn } \mathfrak{M} \models A, \text{ dann } \mathfrak{M} \models B \\ \mathfrak{M} \models A \text{ oder } \mathfrak{M} \models B &\implies \mathfrak{M} \models A \vee B \end{aligned}$$

Die umgekehrte Richtung “ $\Leftarrow$ ” gilt i. A. jeweils *nicht*. Es gilt aber

$$\mathfrak{M} \models A \wedge B \iff \mathfrak{M} \models A \text{ und } \mathfrak{M} \models B$$

Ein Gegenbeispiel für $\mathfrak{M} \not\models A \implies \mathfrak{M} \models \neg A$ ist die offene Formel $R(x, y)$ zusammen mit der Struktur $\mathfrak{M} = \langle M, \mathcal{I} \rangle$, wobei $M = \{0, 1\}$ und $\mathcal{I}(R) = \{ \langle n, n' \rangle \mid n = n' \} \subseteq M^2$. Es gilt $\mathfrak{M} \not\models R(x, y)$, da es eine Variablenbelegung v in $\mathfrak{M}$ gibt, so dass $\mathfrak{M} \not\models_v R(x, y)$. Ist z. B. $v(x) = 0$ und $v(y) = 1$, so ist $\langle \llbracket x \rrbracket_v^{\mathfrak{M}}, \llbracket y \rrbracket_v^{\mathfrak{M}} \rangle = \langle 0, 1 \rangle \notin \mathcal{I}(R)$, da $0 \neq 1$. Hingegen gilt $\mathfrak{M} \models \neg R(x, y)$ *nicht*, da es eine Variablenbelegung v' in $\mathfrak{M}$ gibt, z. B. $v'(x) = v'(y) = 1$, so dass $\mathfrak{M} \models_{v'} \neg R(x, y)$, und somit $\mathfrak{M} \not\models \neg R(x, y)$.

Für *geschlossene* Formeln A und B gelten jeweils beide Richtungen, d. h. es ist

$$\begin{aligned} \mathfrak{M} \models \neg A &\iff \mathfrak{M} \not\models A \\ \mathfrak{M} \models A \rightarrow B &\iff \text{Wenn } \mathfrak{M} \models A, \text{ dann } \mathfrak{M} \models B \\ \mathfrak{M} \models A \vee B &\iff \mathfrak{M} \models A \text{ oder } \mathfrak{M} \models B \end{aligned}$$

(und, wie schon festgestellt, $\mathfrak{M} \models A \wedge B \iff \mathfrak{M} \models A \text{ und } \mathfrak{M} \models B$).

8.5 Logische Folgerung

Nun geben wir den Begriff der logischen Folgerung für die Quantorenlogik an. Für Formelmengen Γ schreiben wir $\mathfrak{M} \models_v \Gamma$, falls $\mathfrak{M} \models_v A$ für alle Formeln $A \in \Gamma$, bzw. $\mathfrak{M} \models \Gamma$, falls Γ eine Menge von geschlossenen Formeln (d. h. von Aussagen) ist.

Definition 8.10 Sei $\Gamma \subseteq \mathcal{L}$ eine Formelmenge und $A \in \mathcal{L}$ eine Formel. Aus Γ folgt (quantoren-)logisch A (Notation: $\Gamma \models A$), falls für alle Strukturen $\mathfrak{M}$ für $\mathcal{L}$ und alle Variablenbelegungen v in $\mathfrak{M}$ gilt: Wenn $\mathfrak{M} \models_v \Gamma$, dann $\mathfrak{M} \models_v A$. Das heißt *logische Folgerung*

$$\Gamma \models A :\iff \text{Für alle } \mathfrak{M} \text{ für } \mathcal{L} \text{ und alle } v \text{ in } \mathfrak{M}: \text{ Wenn } \mathfrak{M} \models_v \Gamma, \text{ dann } \mathfrak{M} \models_v A.$$

Bemerkungen. (i) Für Formeln mit freien Variablen ist die logische Folgerung hier so definiert, dass $\Gamma \models A$ für jede Variablenbelegung gilt.

(ii) Für geschlossene Formeln (d. h. für Aussagen) vereinfacht sich der Begriff der logischen Folgerung. Sei $\Gamma \subseteq \mathcal{L}$ eine Menge von Aussagen und $A \in \mathcal{L}$ eine Aussage. Dann ist *logische Folgerung* definiert durch:

$$\Gamma \models A :\iff \text{Für alle } \mathfrak{M} \text{ für } \mathcal{L}: \text{ Wenn } \mathfrak{M} \models \Gamma, \text{ dann } \mathfrak{M} \models A.$$

Die rechte Seite bedeutet (nach Definition 8.5 (i)):

Für alle $\mathfrak{M}$ für $\mathcal{L}$: Wenn für alle v in $\mathfrak{M}$: $\mathfrak{M} \models_v \Gamma$, dann für alle v in $\mathfrak{M}$: $\mathfrak{M} \models_v A$.

(iii) Der Begriff der logischen Folgerung beruht (wie auch schon in der Aussagenlogik) auf der Idee der Wahrheitskonservierung: Eine Folgerungsbehauptung $\Gamma \models A$ gilt genau dann, wenn stets für gültige Prämissen Γ auch die Konklusion A gültig ist.

Beispiel. Nun können wir die Frage beantworten, ob der folgende natürlichsprachliche Schluss gültig ist:

$$\begin{array}{l} \text{Alle Menschen sind sterblich.} \\ \text{Sokrates ist ein Mensch.} \\ \hline \text{Sokrates ist sterblich.} \end{array}$$

Die Formalisierung der Prämissen ergab $\forall x(P(x) \rightarrow Q(x))$ bzw. $P(k)$, die der Konklusion $Q(k)$. Damit lautet die Frage:

$$\text{Gilt } \forall x(P(x) \rightarrow Q(x)), P(k) \models Q(k)?$$

Da die Folgerungsbehauptung $\forall x(P(x) \rightarrow Q(x)), P(k) \models Q(k)$ nur Aussagen enthält, genügt es, zu zeigen, dass

$$\text{Wenn } \mathfrak{M} \models \forall x(P(x) \rightarrow Q(x)) \text{ und } \mathfrak{M} \models P(k), \text{ dann } \mathfrak{M} \models Q(k).$$

Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine beliebige Struktur für die Sprache $\mathcal{L}_s$, welche die beiden 1-stelligen Prädikatsymbole P und Q , sowie die Konstante k enthält. Es ist $\{\forall x(P(x) \rightarrow Q(x)), P(k)\} \subseteq \mathcal{L}_s$ und $Q(k) \in \mathcal{L}_s$.

Angenommen, es gilt $\mathfrak{M} \models \forall x(P(x) \rightarrow Q(x))$ und $\mathfrak{M} \models P(k)$. Da

$$\mathfrak{M} \models \forall x(P(x) \rightarrow Q(x))$$

$$\iff \text{Für jede Variablenbelegung } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_v P(x) \rightarrow Q(x)$$

- $\Longleftrightarrow$ Für jede Variablenbelegung v in $\mathfrak{M}$: Wenn $\mathfrak{M} \models_v P(x)$, dann $\mathfrak{M} \models_v Q(x)$
 $\Longleftrightarrow$ Für jede Variablenbelegung v in $\mathfrak{M}$: Wenn $\llbracket x \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(P)$, dann $\llbracket x \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(Q)$

gilt für jede Variablenbelegung v in $\mathfrak{M}$:

$$\text{Wenn } \llbracket x \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(P), \text{ dann } \llbracket x \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(Q); \quad (1)$$

und da

$$\mathfrak{M} \models P(k) \Longleftrightarrow \llbracket k \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(P)$$

gilt

$$\llbracket k \rrbracket_v^{\mathfrak{M}} \in \mathcal{I}(P) \quad (2)$$

Betrachte die Variablenbelegung v mit $v(x) = m$ für einen Gegenstand $m \in M$, und sei $\llbracket k \rrbracket_v^{\mathfrak{M}} = m$; da der Gegenstandsbereich einer Struktur stets nichtleer ist, muss es mindestens einen solchen Gegenstand m geben. Dann gilt $m \in \mathcal{I}(P)$ aufgrund von (2), und somit wegen (1) auch $m \in \mathcal{I}(Q)$. Damit gilt $\mathfrak{M} \models Q(k)$. Da $\mathfrak{M}$ für $\mathcal{L}_s$ beliebig ist, gilt für alle $\mathfrak{M}$ für $\mathcal{L}_s$: Wenn $\mathfrak{M} \models \forall x(P(x) \rightarrow Q(x))$ und $\mathfrak{M} \models P(k)$, dann $\mathfrak{M} \models Q(k)$. Das heißt, es gilt

$$\forall x(P(x) \rightarrow Q(x)), P(k) \models Q(k)$$

Da die Formalisierung adäquat ist, ist der natürlichsprachliche Schluss gültig.

Im Unterschied zur Aussagenlogik ist die Quantorenlogik nicht entscheidbar:

Theorem 8.11 (Unentscheidbarkeit der Quantorenlogik)

- (i) *Es gibt kein Verfahren, das für beliebige Formelmengen Γ und beliebige Formeln A entscheidet, ob $\Gamma \models A$.*
(ii) *Erfüllbarkeit ist ebenfalls unentscheidbar.*

Beweis. Siehe Church (1936) und Turing (1936).

QED

9 Gesetze der Quantorenlogik, pränexe Normalform und Skolemisierung

Im Folgenden behandeln wir einige grundlegende Gesetze der Quantorenlogik. Eine besondere Anwendung dieser Gesetze besteht in der Konstruktion von Formeln in sogenannter pränexer Normalform, die wir im Anschluss betrachten.

9.1 Gesetze der Quantorenlogik

Logische Äquivalenz ist unter Verwendung der quantorenlogischen Folgerung “ $\models$ ” als direkte Erweiterung der aussagenlogischen Äquivalenz definiert (und wie in der Aussagenlogik eine Äquivalenzrelation):

Definition 9.1 Zwei quantorenlogische Formeln A und B heißen *logisch äquivalent*, falls $A \models B$ und $B \models A$. Notation: $A \models B$. *logisch äquivalent*

Auch in der Quantorenlogik gilt:

Theorem 9.2 (Import-Export) $A_1, \dots, A_n \models B$ genau dann, wenn $\models A_1 \wedge \dots \wedge A_n \rightarrow B$.

Beweis. Analog zum Beweis in der Aussagenlogik. QED

Korollar 9.3 Für quantorenlogische Formeln A und B gilt $A \models B$ genau dann, wenn $\models A \leftrightarrow B$.

Bemerkung. Wir verwenden hier den in Definition 8.10 angegebenen Begriff der logischen Folgerung. Würde man stattdessen den in Bemerkung (ii) zu Definition 8.10 für geschlossene Formeln angegebenen vereinfachten Begriff auch für offene Formeln verwenden, so würde das Import-Export-Theorem nur noch für geschlossene Formeln gelten. Für beliebige (offene oder geschlossene Formeln) gilt dann zwar noch Export (d. h., wenn $\models A_1 \wedge \dots \wedge A_n \rightarrow B$, dann $A_1, \dots, A_n \models B$), aber Import (d. h., wenn $A_1, \dots, A_n \models B$, dann $\models A_1 \wedge \dots \wedge A_n \rightarrow B$) gilt nicht mehr allgemein. Es ist z. B. $P(x) \models \forall y P(y)$, aber $\not\models P(x) \rightarrow \forall y P(y)$. Sei $\{P(x), \forall y P(y)\} \subseteq \mathcal{L}, \mathfrak{M}$ für $\mathcal{L}$ beliebig und $\mathfrak{M} \models P(x)$. Es gilt

$$\begin{aligned} \mathfrak{M} \models P(x) &\iff \text{Für alle } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_v P(x) \\ &\iff \mathfrak{M} \models \forall x P(x) \\ &\iff \mathfrak{M} \models \forall y P(y) \end{aligned}$$

Also insbesondere: $\mathfrak{M} \models P(x) \implies \mathfrak{M} \models \forall y P(y)$, d. h. $P(x) \models \forall y P(y)$, da $\mathfrak{M}$ für $\mathcal{L}$ beliebig. Jedoch gilt *nicht* $\models P(x) \rightarrow \forall y P(y)$. Es ist

$$\begin{aligned} \mathfrak{M} \models P(x) \rightarrow \forall y P(y) &\iff \text{Für alle } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_v P(x) \rightarrow \forall y P(y) \\ &\iff \text{Für alle } v \text{ in } \mathfrak{M}: \text{Wenn } \mathfrak{M} \models_v P(x), \text{ dann } \mathfrak{M} \models_v \forall y P(y) \\ &\iff \text{Für alle } v \text{ in } \mathfrak{M}: \text{Wenn } \mathfrak{M} \models_v P(x), \text{ dann } \mathfrak{M} \models_{v'} P(y) \\ &\quad \text{für jede } y\text{-Variante } v' \text{ von } v \end{aligned}$$

Sei $\mathfrak{M} = \langle \{m_1, m_2\}, \mathcal{I} \rangle$ mit $\mathcal{I}(P) = \{m_1\}$. Dann ist $\mathfrak{M} \models_v P(x)$ für $v(x) = m_1$, aber für die y -Variante $v' = v[y \mapsto m_2]$ gilt $\mathfrak{M} \not\models_{v'} P(y)$. Somit $\mathfrak{M} \not\models P(x) \rightarrow \forall y P(y)$, also auch $\not\models P(x) \rightarrow \forall y P(y)$.

Theorem 9.4 (Dualität)

- (i) $\neg\forall xA \models \exists x\neg A$, (iii) $\forall xA \models \neg\neg\exists x\neg A$,
(ii) $\neg\exists xA \models \forall x\neg A$, (iv) $\exists xA \models \neg\neg\forall x\neg A$.

Beweis. Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine beliebige Struktur mit $m \in M$ beliebig, und sei v eine beliebige Variablenbelegung in $\mathfrak{M}$.

- (i) Es ist zu zeigen, dass nicht sowohl $\mathfrak{M} \models_v \neg\forall xA$ als auch $\mathfrak{M} \models_v \exists x\neg A$ gilt. Aus $\mathfrak{M} \models_v \neg\forall xA$ folgt $\mathfrak{M} \not\models_v \forall xA$. Somit muss es eine x -Variante v' von v in $\mathfrak{M}$ geben, so dass $\mathfrak{M} \not\models_{v'} A$, d. h. $\mathfrak{M} \models_{v'} \neg A$. Also $\mathfrak{M} \models_v \exists x\neg A$.

Des Weiteren ist zu zeigen, dass nicht sowohl $\mathfrak{M} \models_v \exists x\neg A$ als auch $\mathfrak{M} \models_v \neg\forall xA$ gilt. Aus $\mathfrak{M} \models_v \exists x\neg A$ folgt, dass es eine x -Variante v' von v in $\mathfrak{M}$ geben muss, so dass $\mathfrak{M} \models_{v'} \neg A$, d. h. $\mathfrak{M} \not\models_{v'} A$. Folglich $\mathfrak{M} \not\models_v \forall xA$, d. h. $\mathfrak{M} \models_v \neg\forall xA$.

(ii) und (iii) als Übungsaufgabe.

- (iv) Angenommen $\mathfrak{M} \models_v \exists xA$. Dann gibt es eine x -Variante $v' = v[x \mapsto m]$ in $\mathfrak{M}$, so dass $\mathfrak{M} \models_{v'} A$, also $\mathfrak{M} \not\models_{v'} \neg A$, und folglich $\mathfrak{M} \models_v \neg\forall x\neg A$, d. h. $\mathfrak{M} \models_v \forall x\neg A$.

Angenommen $\mathfrak{M} \not\models_v \exists xA$. Dann gibt es keine x -Variante v' von v , so dass $\mathfrak{M} \models_{v'} A$. Damit gilt $\mathfrak{M} \models_{v'} \neg A$ für jede x -Variante v' von v in $\mathfrak{M}$. Folglich $\mathfrak{M} \models_v \neg\forall x\neg A$, d. h. $\mathfrak{M} \not\models_v \forall x\neg A$. QED

Theorem 9.5 (Vertauschung) (i) $\forall x\forall yA \models \forall y\forall xA$, (ii) $\exists x\exists yA \models \exists y\exists xA$.

Beweis. (i) Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine beliebige Struktur und v eine beliebige Variablenbelegung in $\mathfrak{M}$. Es gilt

$$\begin{aligned} \mathfrak{M} \models_v \forall x\forall yA &\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} \forall yA \\ &\iff \text{Für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M} \text{ und} \\ &\quad \text{für jede } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v''} A \end{aligned}$$

Seien m_1 und m_2 zwei beliebige, möglicherweise identische Gegenstände in M . Es gilt

$$v'' = v[x \mapsto m_1][y \mapsto m_2] = v[y \mapsto m_2][x \mapsto m_1] = v'$$

Somit gilt

$$\begin{aligned} \mathfrak{M} \models_v \forall x\forall yA &\iff \text{Für jede } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M} \text{ und} \\ &\quad \text{für jede } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} A \\ &\iff \text{Für jede } y\text{-Variante } v'' \text{ von } v' \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v''} \forall xA \\ &\iff \mathfrak{M} \models_v \forall y\forall xA \end{aligned}$$

(ii) analog. QED

Theorem 9.6 Es gilt die logische Folgerung $\exists x\forall yA \models \forall y\exists xA$.

(Die Umkehrung gilt jedoch nicht; d. h. $\forall x\exists yA \not\models \exists y\forall xA$.)

Beweis. Übungsaufgabe. QED

Theorem 9.7 (Leere Quantoren)

Falls x nicht frei in A vorkommt, gilt: (i) $\forall xA \models A$, (ii) $\exists xA \models A$.

Beweis. Übungsaufgabe.

QED

Bemerkung. Die Einschränkung, dass x nicht frei in A vorkommen darf, ist notwendig, da z. B. für $\mathfrak{M} = \langle \{m_1, m_2\}, \mathcal{I} \rangle$ mit $\mathcal{I}(P) = \{m_1\}$ und v beliebig zwar $\mathfrak{M} \models_{v[x \mapsto m_1]} P(x)$, aber $\mathfrak{M} \not\models_v \forall x P(x)$, also $P(x) \not\models \forall x P(x)$; und da zwar $\mathfrak{M} \models_v \exists x P(x)$, aber $\mathfrak{M} \not\models_{v[x \mapsto m_2]} P(x)$, ist $\exists x P(x) \not\models P(x)$.

Theorem 9.8 (Distributivität über $\wedge$ und $\vee$)

(i) $\exists x(A \vee B) \models \exists x A \vee \exists x B$, (ii) $\forall x(A \wedge B) \models \forall x A \wedge \forall x B$.

Beweis. (i) Sei $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ eine beliebige Struktur und v eine beliebige Variablenbelegung in $\mathfrak{M}$. Es gilt

$$\begin{aligned} \mathfrak{M} \models_v \exists x(A \vee B) &\iff \text{Es gibt eine } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} A \vee B \\ &\iff \text{Es gibt eine } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \\ &\quad \mathfrak{M} \models_{v'} A \text{ oder } \mathfrak{M} \models_{v'} B \\ &\iff \text{Es gibt eine } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} A \\ &\quad \text{oder es gibt eine } x\text{-Variante } v' \text{ von } v \text{ in } \mathfrak{M}: \mathfrak{M} \models_{v'} B \\ &\iff \mathfrak{M} \models_v \exists x A \text{ oder } \mathfrak{M} \models_v \exists x B \\ &\iff \mathfrak{M} \models_v \exists x A \vee \exists x B \end{aligned}$$

(ii) analog.

QED

Theorem 9.9 Es gelten die beiden logischen Folgerungen:

(i) $\exists x(A \wedge B) \models \exists x A \wedge \exists x B$, (ii) $\forall x A \vee \forall x B \models \forall x(A \vee B)$.

(Die Umkehrungen gelten jedoch nicht; d. h. $\exists x A \wedge \exists x B \not\models \exists x(A \wedge B)$ und $\forall x(A \vee B) \not\models \forall x A \vee \forall x B$.)

Beweis. Übungsaufgabe.

QED

Um ausdrücken zu können, dass Formeln, die sich lediglich durch ihre gebundenen Variablen unterscheiden, zueinander logisch äquivalent sind, benötigen wir das Hilfsmittel der Substitution.

Definition 9.10 (i) Eine *Substitution* ist eine Funktion, die Variablen auf Terme abbildet und für nur endlich viele Variablen nicht identisch ist. *Substitution*

(ii) Wir notieren Substitutionen als endliche Mengen (die wir mit eckigen Klammern $[\]$ schreiben) der Form

$$[x_1/t_1, \dots, x_n/t_n] \quad (\text{für } 0 \leq i \leq n),$$

für paarweise verschiedene Variablen x_i und Terme t_i , so dass $t_i \neq x_i$.

In der Notation $[x_1/t_1, \dots, x_n/t_n]$ werden also nur die (immer endlich vielen) nicht-identischen Zuordnungen von Termen zu Variablen einer Substitution angegeben.

(iii) Es heißt x_i/t_i *Bindung* für x_i .

Bindung

Wir sagen auch “ t_i wird für x_i substituiert” oder “ x_i wird durch t_i ersetzt”.

(iv) Substitutionen bezeichnen wir mit $\sigma, \rho, \tau, \vartheta, \dots$

(v) Ist $\sigma = \emptyset$, so heißt σ *leere Substitution* und wird mit ε bezeichnet.

leere Substitution

Da wir später auch Funktionsterme der Form $f(t_1, \dots, t_n)$ benötigen, schließen wir diese in der folgenden Definition von Anwendungen von Substitutionen schon mit ein.

Definition 9.11 Sei $\sigma = [x_1/t_1, \dots, x_n/t_n]$ eine Substitution und E ein Ausdruck, d. h. eine Formel oder ein Term.

(i) Die simultane Ersetzung jedes freien Vorkommens von x_i in E durch t_i für $0 \leq i \leq n$ heißt *Anwendung* von σ auf E .

Anwendung

Notation: $E\sigma$. Die Substitution wirkt also nach links.

Ist E eine Formel, wird zusätzlich gefordert, dass t_i in E *frei einsetzbar* ist für x_i ; d. h., x_i darf nur dann durch t_i ersetzt werden, wenn in t_i vorkommende Variablen in E nicht durch einen Quantor gebunden werden.

frei einsetzbar

(ii) Der resultierende Ausdruck $E\sigma$ heißt (*Substitutions-*) *Instanz* von E für σ .

Instanz

(iii) Ist $X = \{E_1, \dots, E_n\}$ eine endliche Menge von Ausdrücken, so steht $X\sigma$ für die Menge $\{E_1\sigma, \dots, E_n\sigma\}$.

Bemerkungen. (i) Da Substitutionen simultan ausgeführt werden, ist

$$P(x, y) [x/y, y/k] = P(y, k)$$

und *nicht* $P(k, k)$, was man durch die Hintereinanderausführung

$$(P(x, y) [x/y]) [y/k] = P(y, y) [y/k] = P(k, k)$$

erhalten würde.

(ii) Ist A eine Formel, x eine Variable und t ein Term, so ist $A [x/t]$ das Resultat der Ersetzung aller freien Vorkommen von x in A durch t .

(iii) Kommt x in A im Wirkungsbereich eines Quantors $\forall y$ oder $\exists y$ für eine Variable y vor, so ist $A [x/y]$ nicht definiert. Andernfalls ist $A [x/y]$ das Resultat der Ersetzung aller freien Vorkommen von x in A durch y .

(iv) Falls x in A nicht frei vorkommt, ist $A [x/y]$ identisch mit A .

(v) Die Operation $[x/t]$ wirkt global, d. h. auf die gesamte links von $[x/t]$ stehende Formel. Möchte man die Wirkung auf eine Teilformel einschränken, muss entsprechend geklammert werden.

(vi) Man beachte, dass das Vorkommen von x bei den Quantoren $\forall x$ und $\exists x$ weder frei noch gebunden ist. Deshalb ist z. B. $\forall x P(y) [x/z]$ die Formel $\forall x P(y)$, und *nicht* etwa $\forall z P(y)$.

Beispiele. (i) $\exists x (P(x) \vee Q(y)) [y/k]$ ist $\exists x (P(x) \vee Q(k))$.

(ii) $\forall x (P(y) \rightarrow R(x, y)) [y/k]$ ist $\forall x (P(k) \rightarrow R(x, k))$.

(iii) $\exists x R(x, y) [y/x]$ ist nicht definiert.

(iv) $P(x) \wedge (R(x, y) [x/z])$ ist $P(x) \wedge R(z, y)$.

(v) $P(x) \wedge R(x, y) [x/z]$ ist $P(z) \wedge R(z, y)$.

Definition 9.12 Die *Komposition* $\sigma\tau$ zweier Substitutionen

Komposition

$$\sigma = [x_1/s_1, \dots, x_n/s_n] \quad \text{und} \quad \tau = [y_1/t_1, \dots, y_m/t_m]$$

mit $0 \leq i \leq n$ und $0 \leq j \leq m$ ist wie folgt gegeben: Wir bilden die Folge von Bindungen

$$x_1/(s_1\tau), \dots, x_n/(s_n\tau), y_1/t_1, \dots, y_m/t_m.$$

Aus dieser Folge entfernen wir

- (i) alle Bindungen $x_i/(s_i\tau)$, für die $x_i = (s_i\tau)$,
- (ii) und alle Bindungen y_j/t_j , für die $y_j \in \{x_1, \dots, x_n\}$.

Die aus den Bindungen in der resultierenden Folge bestehende Substitution ist die *Komposition* $\sigma\tau$ von σ und τ .

Beispiel. Sei $\sigma = [x/y, y/z]$ und $\tau = [x/k, y/l, z/y]$. Dann ist die Folge von Bindungen

$$\underbrace{x/(y[x/k, y/l, z/y])}_{x/l}, \underbrace{y/(z[x/k, y/l, z/y])}_{y/y}, x/k, y/l, z/y.$$

Aus dieser sind die Bindungen y/y gemäß (i) und $x/k, y/l$ gemäß (ii) zu entfernen. Man erhält die Komposition $\sigma\tau = [x/l, z/y]$.

Lemma 9.13 Seien ρ, σ, ϑ Substitutionen und ε die leere Substitution. Dann gilt:

- (i) $\rho\varepsilon = \varepsilon\rho = \rho$.
- (ii) $(\rho\sigma)\vartheta = \rho(\sigma\vartheta)$.
- (iii) $(E\rho)\sigma = E(\rho\sigma)$, für beliebige Formeln und Terme E .

Beweis. Übungsaufgabe.

QED

Lemma 9.14 Die *Komposition* von Substitutionen ist nicht kommutativ.

Beweis. Übungsaufgabe.

QED

Theorem 9.15 (Gebundene Umbenennung)

Falls y nicht frei in A vorkommt, gilt:

- (i) $\forall x A \models \forall y (A[x/y])$, (ii) $\exists x A \models \exists y (A[x/y])$.

Beweis. Offensichtlich, da Substitution nicht induktiv, sondern global definiert wurde. QED

Beispiele. (i) $\forall x P(x) \models \forall y P(y)$.

(ii) $\exists x P(x) \wedge Q(y) \models \exists y P(y) \wedge Q(y)$.

(iii) $\forall x \exists y R(x, y) \models \forall x \exists z R(x, z)$.

(iv) Hingegen sind $\exists x R(x, y)$ und $\exists y R(y, y)$ nicht logisch äquivalent.

In Beispiel (ii) und (iii) haben wir auch das schon aus der Aussagenlogik bekannte Ersetzungstheorem verwendet, das für die Quantorenlogik erweitert werden kann:

Theorem 9.16 (Ersetzungstheorem) Es sei $A \models A'$, A eine Teilformel von B , und es sei B' das Resultat der Ersetzung von A durch A' in B . Dann gilt $B \models B'$.

Beweis. Per Induktion über dem Aufbau der Formel B :

Induktionsanfang: B sei atomar. Dann ist A mit B identisch. Mit $A \models A'$ gilt dann auch $B \models B'$.

Induktionsannahme: Die Behauptung (d. h. das Ersetzungstheorem) gelte für die unmittelbaren Teilformeln einer Formel B .

Induktionsschritt: Es sei $\mathfrak{M}$ eine beliebige Struktur und v eine beliebige Variablenbelegung in $\mathfrak{M}$.

1. Fall: B sei $\exists xA$. Es ist zu zeigen, dass $B \models B'$. Ist $\mathfrak{M} \models_v \exists xA$, so gibt es eine x -Variante v' von v in $\mathfrak{M}$, für die $\mathfrak{M} \models_{v'} A$. Das ist wegen $A \models A'$ genau dann der Fall, wenn es eine x -Variante v' von v in $\mathfrak{M}$ gibt, so dass $\mathfrak{M} \models_{v'} A'$. Somit gilt $\mathfrak{M} \models_v \exists xA'$, d. h. $\mathfrak{M} \models B'$.

2. Fall: B sei $\forall xA$. Es ist wieder zu zeigen, dass $B \models B'$. Ist $\mathfrak{M} \models_v \forall xA$, so gilt für jede x -Variante v' von v in $\mathfrak{M}$: $\mathfrak{M} \models_{v'} A$, was wegen $A \models A'$ genau dann der Fall ist, wenn für jede x -Variante v' von v in $\mathfrak{M}$ gilt: $\mathfrak{M} \models_{v'} A'$. Somit gilt $\mathfrak{M} \models_v \forall xA'$, d. h. $\mathfrak{M} \models B'$.

3. Fall: B beginne nicht mit einem Quantor, d. h. B sei $\neg A$ oder habe die Form $A \circ C$ bzw. $C \circ A$ für ein 2-stelliges Konnektiv $\circ$. Es ist wieder zu zeigen, dass $B \models B'$.

– B sei $\neg A$. Es gilt $\mathfrak{M} \models_v \neg A$ genau dann, wenn $\mathfrak{M} \not\models_v A$. Dies gilt wegen $A \models A'$ genau dann, wenn $\mathfrak{M} \not\models_v A'$, was genau dann gilt, wenn $\mathfrak{M} \models_v \neg A'$, d. h. $\mathfrak{M} \models_v B'$.

– B sei $A \wedge C$. Es gilt $\mathfrak{M} \models_v A \wedge C$ genau dann, wenn $\mathfrak{M} \models_v A$ und $\mathfrak{M} \models_v C$. Wegen $A \models A'$ gilt dann auch $\mathfrak{M} \models_v A'$, und damit $\mathfrak{M} \models_v A' \wedge C$, d. h. $\mathfrak{M} \models_v B'$. Entsprechend für $C \wedge A$.

Für die restlichen Konnektive zeigt man den Induktionsschritt analog.

QED

Theorem 9.17 (Eingeschränkte Distributivität über $\rightarrow$)

Falls x nicht frei in B vorkommt, gilt:

(i) $\forall xA \rightarrow B \models \exists x(A \rightarrow B)$, (iii) $B \rightarrow \forall xA \models \forall x(B \rightarrow A)$,

(ii) $\exists xA \rightarrow B \models \forall x(A \rightarrow B)$, (iv) $B \rightarrow \exists xA \models \exists x(B \rightarrow A)$.

Beweis. (i) $\forall xA \rightarrow B \models \neg \forall xA \vee B$ (Aussagenlogik)

$\models \exists x \neg A \vee B$ (Dualität, Ersetzung)

$\models \exists x \neg A \vee \exists xB$ (Leerer Quantor, Ersetzung)

$\models \exists x(\neg A \vee B)$ (Distributivität über $\vee$)

$\models \exists x(A \rightarrow B)$ (Aussagenlogik, Ersetzung)

(ii)-(iv) analog.

QED

Beispiel. Es gilt $\models \exists x(P(x) \rightarrow \forall yP(y))$. Denn

$\exists x(P(x) \rightarrow \forall yP(y)) \models \forall xP(x) \rightarrow \forall yP(y)$ (Eingeschr. Distrib. über $\rightarrow$; (i))

$\models \forall xP(x) \rightarrow \forall xP(x)$ (Geb. Umbenennung, Ersetzung)

$\models p \rightarrow p$ (Permanenz)

Da $\models p \rightarrow p$, muss auch $\models \exists x(P(x) \rightarrow \forall yP(y))$ gelten.

Theorem 9.18 (Eingeschränkte Distributivität über $\wedge$ und $\vee$)

Falls x nicht frei in B vorkommt, gilt:

- (i) $\exists x(A \wedge B) \models \exists xA \wedge B$, (iii) $\forall x(A \wedge B) \models \forall xA \wedge B$,
(ii) $\exists x(A \vee B) \models \exists xA \vee B$, (iv) $\forall x(A \vee B) \models \forall xA \vee B$.

Beweis. (i) $\exists x(A \wedge B) \models \neg \forall x \neg(A \wedge B)$ (Dualität)
 $\models \neg \forall x(\neg A \vee \neg B)$ (Aussagenlogik)
 $\models \neg \forall x(A \rightarrow \neg B)$ (Aussagenlogik)
 $\models \neg(\exists xA \rightarrow \neg B)$ (Eingeschr. Distrib. über $\rightarrow$; (ii))
 $\models \neg(\neg \exists xA \vee \neg B)$ (Aussagenlogik)
 $\models \exists xA \wedge B$ (Aussagenlogik)

Außer im 1. Umformungsschritt wurde überall auch das Ersetzungstheorem verwendet.

(ii)-(iv) analog.

QED

Wir fassen die Gesetze noch in einem Überblick zusammen. Die uneingeschränkten Gesetze gelten für beliebige Formeln A und B sowie für beliebige Variablen x und y ; bei den eingeschränkten Gesetzen ist die jeweils angegebene Variablenbedingung zu beachten.

<i>Uneingeschränkte Gesetze</i>	<i>Eingeschränkte Gesetze</i>
<i>Dualität</i> $\neg \forall xA \models \exists x \neg A$, $\forall xA \models \neg \exists x \neg A$ $\neg \exists xA \models \forall x \neg A$, $\exists xA \models \neg \forall x \neg A$ <i>Vertauschung</i> $\forall x \forall yA \models \forall y \forall xA$ $\exists x \exists yA \models \exists y \exists xA$ $\exists x \forall yA \models \forall y \exists xA$ <i>Distributivität über $\wedge$ und $\vee$</i> $\exists x(A \vee B) \models \exists xA \vee \exists xB$ $\forall x(A \wedge B) \models \forall xA \wedge \forall xB$ $\exists x(A \wedge B) \models \exists xA \wedge \exists xB$ $\forall xA \vee \forall xB \models \forall x(A \vee B)$	<i>Leere Quantoren</i> Falls x nicht frei in A vorkommt, gilt: $\forall xA \models A$, $\exists xA \models A$ <i>Gebundene Umbenennung</i> Falls y nicht frei in A vorkommt, gilt: $\forall xA \models \forall y(A[x/y])$, $\exists xA \models \exists y(A[x/y])$ <i>Eingeschränkte Distributivität über $\rightarrow$</i> Falls x nicht frei in B vorkommt, gilt: $\forall xA \rightarrow B \models \exists x(A \rightarrow B)$ $\exists xA \rightarrow B \models \forall x(A \rightarrow B)$ $B \rightarrow \forall xA \models \forall x(B \rightarrow A)$ $B \rightarrow \exists xA \models \exists x(B \rightarrow A)$ <i>Eingeschränkte Distributivität über $\wedge$ und $\vee$</i> Falls x nicht frei in B vorkommt, gilt: $\exists x(A \wedge B) \models \exists xA \wedge B$ $\exists x(A \vee B) \models \exists xA \vee B$ $\forall x(A \wedge B) \models \forall xA \wedge B$ $\forall x(A \vee B) \models \forall xA \vee B$

9.2 Pränexe Normalform

Definition 9.19 Eine Formel A einer Sprache $\mathcal{L}$ ist in *pränexer Normalform* (kurz: PNF), *pränexe Normalform* wenn sie die Form

$$Q_1 x_1 \dots Q_n x_n B$$

hat, wobei B quantorenfrei, $n \geq 0$, Q_i entweder $\forall$ oder $\exists$ ist, und alle x_i paarweise verschieden sind.

Es heißt $Q_1 x_1 \dots Q_n x_n$ *Präfix* der Formel A , und B heißt *Kern* oder *Matrix* von A . (Im Fall $n = 0$ ist A quantorenfrei und identisch mit B .)

Theorem 9.20 Sei A eine Formel. Dann gibt es eine Formel B in pränexer Normalform, so dass A und B logisch äquivalent sind.

Beweis. Wir geben ein Verfahren an, um die Formel A in eine logisch äquivalente Formel B in pränexer Normalform zu überführen. Jeder Umformungsschritt beruht auf einer logischen Äquivalenz. Die einzelnen Umformungsschritte werden am Beispiel der Formel

$$\forall y (\forall x (Px \rightarrow Qx) \vee \neg \forall x (Qx \vee Rxzk))$$

veranschaulicht. (Zur besseren Lesbarkeit verwenden wir die abkürzende Schreibweise bei Relationszeichen.)

(1) Streiche alle leeren Quantoren in A , d. h.

(a) $\forall x A_1 \rightsquigarrow A_1$, falls x nicht frei in A_1 ; (b) $\exists x A_1 \rightsquigarrow A_1$, falls x nicht frei in A_1 .

Es gilt $A \models A_1$.

Beispiel. $\forall y (\forall x (Px \rightarrow Qx) \vee \neg \forall x (Qx \vee Rxzk)) \rightsquigarrow$

$$\forall x (Px \rightarrow Qx) \vee \neg \forall x (Qx \vee Rxzk)$$

(2) Benenne gebundene Variablen in A_1 so um, dass alle Quantoren verschiedene Variablen haben, keine Variable sowohl frei als auch gebunden vorkommt, und freie Variablen nicht gebunden werden. (Letzteres ist durch die Forderung nach freier Einsetzbarkeit in Definition 9.11(i) gewährleistet.)

Die resultierende Formel sei A_2 . Es gilt $A_1 \models A_2$.

Beispiel. $\forall x (Px \rightarrow Qx) \vee \neg \forall x (Qx \vee Rxzk) \rightsquigarrow$

$$\forall x (Px \rightarrow Qx) \vee \neg \forall y (Qy \vee Ryzk)$$

(3) Ziehe Negationen nach innen, so dass diese nur noch vor Atomen vorkommen, und beseitige doppelte Negationen:

(a) $\neg \forall x C \rightsquigarrow \exists x \neg C$

(d) $\neg (C \vee D) \rightsquigarrow (\neg C \wedge \neg D)$

(b) $\neg \exists x C \rightsquigarrow \forall x \neg C$

(e) $\neg (C \rightarrow D) \rightsquigarrow (C \wedge \neg D)$

(c) $\neg (C \wedge D) \rightsquigarrow (\neg C \vee \neg D)$

(f) $\neg \neg C \rightsquigarrow C$

Die resultierende Formel sei A_3 . Es gilt $A_2 \models A_3$.

Beispiel. $\forall x (Px \rightarrow Qx) \vee \neg \forall y (Qy \vee Ryzk) \rightsquigarrow$

$$\forall x (Px \rightarrow Qx) \vee \exists y \neg (Qy \vee Ryzk) \rightsquigarrow$$

$$\forall x (Px \rightarrow Qx) \vee \exists y (\neg Qy \wedge \neg Ryzk)$$

(4) Ziehe Quantoren nach außen:

- | | |
|--|--|
| (a) $\forall x C \wedge D \rightsquigarrow \forall x (C \wedge D)$ | (g) $\exists x C \vee D \rightsquigarrow \exists x (C \vee D)$ |
| (b) $C \wedge \forall x D \rightsquigarrow \forall x (C \wedge D)$ | (h) $C \vee \exists x D \rightsquigarrow \exists x (C \vee D)$ |
| (c) $\exists x C \wedge D \rightsquigarrow \exists x (C \wedge D)$ | (i) $\forall x C \rightarrow D \rightsquigarrow \exists x (C \rightarrow D)$ |
| (d) $C \wedge \exists x D \rightsquigarrow \exists x (C \wedge D)$ | (j) $C \rightarrow \forall x D \rightsquigarrow \forall x (C \rightarrow D)$ |
| (e) $\forall x C \vee D \rightsquigarrow \forall x (C \vee D)$ | (k) $\exists x C \rightarrow D \rightsquigarrow \forall x (C \rightarrow D)$ |
| (f) $C \vee \forall x D \rightsquigarrow \forall x (C \vee D)$ | (l) $C \rightarrow \exists x D \rightsquigarrow \exists x (C \rightarrow D)$ |

Durch (2) ist sichergestellt, dass hierdurch keine freien Variablen gebunden werden.

Die resultierende Formel ist die gesuchte pränexe Normalform B der Ausgangsformel A . Es gilt $A_3 \models B$ und, da $A \models A_1 \models A_2 \models A_3 \models B$, auch $A \models B$. Die pränexe Normalform ist also logisch äquivalent zur Ausgangsformel.

Beispiel. $\forall x (Px \rightarrow Qx) \vee \exists y (\neg Qy \wedge \neg Ryzk) \rightsquigarrow$

$$\forall x ((Px \rightarrow Qx) \vee \exists y (\neg Qy \wedge \neg Ryzk)) \rightsquigarrow$$

$$\forall x \exists y ((Px \rightarrow Qx) \vee (\neg Qy \wedge \neg Ryzk))$$

oder auch $\forall x (Px \rightarrow Qx) \vee \exists y (\neg Qy \wedge \neg Ryzk) \rightsquigarrow$

$$\exists y (\forall x (Px \rightarrow Qx) \vee (\neg Qy \wedge \neg Ryzk)) \rightsquigarrow$$

$$\exists y \forall x ((Px \rightarrow Qx) \vee (\neg Qy \wedge \neg Ryzk))$$

QED

Für das Verfahren ist nur die Reihenfolge der Schritte (1)-(4) zu beachten. Die Reihenfolge der Teilschritte (1a)/(1b), (3a)-(3f) und (4a)-(4l) ist hingegen beliebig.

Beispiele. (i) $\forall z (\exists x Px \rightarrow \neg \exists x Q(x, y)) \xrightarrow{(1a)} \exists x Px \rightarrow \neg \exists x Q(x, y)$
 $\xrightarrow{(2)} \exists x Px \rightarrow \neg \exists z Qzy$
 $\xrightarrow{(3b)} \exists x Px \rightarrow \forall z \neg Qzy$
 $\xrightarrow{(4j)} \forall z (\exists x Px \rightarrow \neg Qzy)$
 $\xrightarrow{(4k)} \forall z \forall x (Px \rightarrow \neg Qzy)$

(ii) Man kann auch zuerst (4k) und dann (4j) anwenden:

$$\exists x Px \rightarrow \forall z \neg Qzy \xrightarrow{(4k)} \forall x (Px \rightarrow \forall z \neg Qzy) \xrightarrow{(4j)} \forall x \forall z (Px \rightarrow \neg Qzy)$$

Die pränexe Normalform einer Formel ist nicht eindeutig bestimmt, da

- (i) unterschiedliche gebundene Umbenennungen möglich sind,
- (ii) die Reihenfolge, in der Quantoren nach außen gezogen werden können, nicht festgelegt ist,
- (iii) und der Kern prinzipiell durch jede zu diesem logisch äquivalente quantorenfreie Formel ersetzt werden darf.

9.3 Skolemisierung

Zu einer Formel in pränexer Normalform in einer Sprache $\mathcal{L}$ kann mittels *Skolemisierung* (nach Thoralf Skolem, 1887–1963) eine quantorenfreie Formel in einer erweiterten Sprache $\mathcal{L}_s \supseteq \mathcal{L}$ gebildet werden. Allerdings ist diese Formel im Allgemeinen nicht mehr logisch äquivalent zur Ausgangsformel, jedoch in jedem Fall noch erfüllbarkeitsäquivalent zu derselben.

Zur Erläuterung der Idee bei Skolemisierung betrachten wir die (in der Standardinterpretation wahre) Aussage

Für alle $n \in \mathbb{N}$ gibt es ein $m \in \mathbb{N}$, so dass $n < m$.

Für die Funktion $f(n) = n + 1$ gilt $n < f(n)$ für alle n . Die durch den Existenzquantor gebundene Variable m kann also durch die Funktion $f(n)$ ersetzt werden, wodurch der Existenzquantor überflüssig wird. Fasst man freie Variablen zudem als universell auf, so kann auch der Allquantor weggelassen werden. Lässt man auch von der Standardinterpretation abweichende Interpretationen zu, so ist die resultierende Aussage zwar nicht mehr logisch äquivalent zur Ausgangsaussage, doch die resultierende Aussage ist erfüllbar genau dann, wenn die Ausgangsaussage erfüllbar ist.

Obige Aussage hat die Form

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

wobei wir hier annehmen wollen, dass dies eine Formel in einer Sprache $\mathcal{L}$ ist, die keine Funktionszeichen enthält. Nun ersetzt man in der Formel $\forall x \exists y P(x, y)$ die durch den Existenzquantor gebundene Variable y durch einen Funktionsterm $f(x)$. Dies stellt eine Spracherweiterung um das einstellige Funktionszeichen f dar. Lässt man jetzt noch den Allquantor weg, resultiert als Skolem-Normalform die Formel $P(x, f(x))$. Diese Formel ist erfüllbar genau dann, wenn die Ausgangsformel erfüllbar ist.

Wir nehmen zunächst die für die Verwendung von Funktionszeichen nötigen Anpassungen von Syntax und Semantik vor.

Definition 9.21 Wir erweitern das *Alphabet* um *Funktionszeichen* $f, g, h, \dots$ (ggf. mit Indizes) und erweitern unsere Sprache um *Funktionsterme* wie folgt: Ist f ein n -stelliges Funktionszeichen (für $n \geq 1$) und sind $t_1, \dots, t_n$ Terme, dann ist auch $f(t_1, \dots, t_n)$ ein Term.

*Funktionszeichen
Funktionsterme*

Definition 9.22 Den Begriff einer *Struktur* $\mathfrak{M} = \langle M, \mathcal{I} \rangle$ für eine Sprache $\mathcal{L}$ erweitern wir um *Interpretationen von Funktionszeichen* wie folgt:

*Struktur
Interpretation von
Funktionszeichen*

$$\mathcal{I}(f) : M^n \rightarrow M, \text{ falls } f \in \mathcal{L} \text{ ein } n\text{-stelliges Funktionszeichen ist } (n \geq 1).$$

Das heißt, n -stellige Funktionszeichen in $\mathcal{L}$ werden n -stellige Funktionen $M^n \rightarrow M$ zugeordnet.

Beispiel. Die Sprache $\mathcal{L}$ enthalte lediglich die Konstante k und das 1-stellige Funktionszeichen s . Wir betrachten die Struktur $\mathfrak{N} = \langle \mathbb{N}, \mathcal{I} \rangle$ mit

- (i) $\mathcal{I}(k) = 0$;
- (ii) $\mathcal{I}(s) : M^1 \rightarrow M, m \mapsto m + 1$, wobei ‘+’ hier die normale Addition ist; d. h., $\mathcal{I}$ ordnet dem 1-stelligen Funktionszeichen s die 1-stellige Nachfolger-Funktion $f(m) = m + 1$ zu.

In dieser Struktur bezeichnet k die Zahl 0, $s(k)$ die Zahl 1 (da $\mathcal{I}(s)(\mathcal{I}(k)) = f(0) = 1$), $s(s(k))$ die Zahl 2 (da $\mathcal{I}(s)(\mathcal{I}(s)(\mathcal{I}(k))) = f(f(0)) = f(1) = 2$), usw.

Definition 9.23 Den Begriff der *Termbelegung* (für Strukturen $\mathfrak{M}$ und Variablenbelegungen v) erweitern wir durch eine Klausel für Funktionsterme:

Termbelegung

$$\llbracket f(t_1, \dots, t_n) \rrbracket_v^{\mathfrak{M}} := \mathcal{I}(f)(\llbracket t_1 \rrbracket_v^{\mathfrak{M}}, \dots, \llbracket t_n \rrbracket_v^{\mathfrak{M}}).$$

Beispiel. Für obige Struktur $\mathfrak{N} = \langle \mathbb{N}, \mathcal{I} \rangle$ und beliebige Variablenbelegungen v erhalten wir z. B. die Termbelegung

$$\llbracket s(s(k)) \rrbracket_v^{\mathfrak{N}} = \mathcal{I}(s)(\mathcal{I}(s)(\llbracket k \rrbracket_v^{\mathfrak{N}})) = f(f(\llbracket k \rrbracket_v^{\mathfrak{N}})) = f(f(0)) = 2.$$

Die semantischen Begriffe wie Gültigkeit, Erfüllbarkeit, Allgemeingültigkeit usw. sind damit auch für Sprachen mit Funktionszeichen definiert.

Nun behandeln wir ein Verfahren zur Skolemisierung.

Definition 9.24 Sei $A \in \mathcal{L}$ eine Formel der Form $Q_1 x_1 \dots Q_n x_n B$ (für $n \geq 0$) in pränexer Normalform mit Kern B . Das folgende Verfahren liefert eine *Skolem-Normalform* A_S von A in einer geeigneten Sprache $\mathcal{L}_S \supseteq \mathcal{L}$.

Skolem-Normalform

Für i von 1 bis n :

- (1) Falls $n = 0$, setze $A_S := A$.
- (2) Falls Q_i ein Allquantor ist, setze $A := Q_{i+1} x_{i+1} \dots Q_n x_n B$.
- (3) Falls Q_i ein Existenzquantor ist, und $y_1, \dots, y_m$ die in A frei vorkommenden Variablen sind, setze

$$A := Q_{i+1} x_{i+1} \dots Q_n x_n B[x_i / f_i(y_1, \dots, y_m)]$$

wobei f_i ein neues (noch nicht in A vorkommendes) m -stelliges Funktionszeichen ist, bzw. eine neue Konstante k_i , falls $m = 0$.

Beispiele. Wir verwenden im Folgenden indizierte Variablen. Die Wahl der Individuenkonstanten und Funktionszeichen ist dann durch das Verfahren festgelegt. Bei Formeln mit nicht-indizierten Variablen ist darauf zu achten, dass in Schritt (3) jeweils neue Konstanten, bzw. neue Funktionszeichen eingeführt werden (vgl. Beispiel (v)).

- (i) $\exists x_1 P(x_1) \rightsquigarrow P(k_1)$
- (ii) $\exists x_1 \forall x_2 P(x_1, x_2) \rightsquigarrow \forall x_2 P(k_1, x_2) \rightsquigarrow P(k_1, x_2)$
- (iii) $\forall x_1 \exists x_2 P(x_1, f_1(x_2)) \rightsquigarrow \exists x_2 P(x_1, f_1(x_2)) \rightsquigarrow P(x_1, f_1(f_2(x_1)))$
- (iv) $\exists x_1 \forall x_2 \exists x_3 P(x_1, x_2, x_3) \rightsquigarrow \forall x_2 \exists x_3 P(k_1, x_2, x_3)$
 $\rightsquigarrow \exists x_3 P(k_1, x_2, x_3)$
 $\rightsquigarrow P(k_1, x_2, f_3(x_2))$
- (v) $\forall x_1 \exists x_2 \forall x_3 \exists x_4 P(x_1, x_2, x_3, x_4) \rightsquigarrow \exists x_2 \forall x_3 \exists x_4 P(x_1, x_2, x_3, x_4)$
 $\rightsquigarrow \forall x_3 \exists x_4 P(x_1, f_2(x_1), x_3, x_4)$
 $\rightsquigarrow \exists x_4 P(x_1, f_2(x_1), x_3, x_4)$
 $\rightsquigarrow P(x_1, f_2(x_1), x_3, f_4(x_1, x_3))$

Für nicht indizierte Variablen:

$$\begin{aligned}
\forall x \exists y \forall z \exists u P(x, y, z, u) &\rightsquigarrow \exists y \forall z \exists u P(x, y, z, u) \\
&\rightsquigarrow \forall z \exists u P(x, f(x), z, u) \\
&\rightsquigarrow \exists u P(x, f(x), z, u) \\
&\rightsquigarrow P(x, f(x), z, g(x, z))
\end{aligned}$$

Definition 9.25 Der *All-Abschluss* $\forall A$ einer Formel A ist die Formel

All-Abschluss

$$\forall x_1 \dots \forall x_n A [y_1/x_1, \dots, y_n/x_n]$$

wobei $y_1, \dots, y_n$ *alle* freien Variablen in A und $x_1, \dots, x_n$ Variablen sind, die in A nicht vorkommen. (Da die paarweise verschiedenen Variablen x_i ansonsten frei wählbar sind, ist dies eine nicht-deterministische Definition.)

Die Skolem-Normalform A_S einer Formel A ist im Allgemeinen *nicht* logisch äquivalent zu A . Es gilt zwar stets $\forall A_S \models A$, aber für gewisse Formeln A ist $A \not\models A_S$. Betrachte z. B. die Formel $\exists x P(x)$ mit Skolem-Normalform $P(k)$. Die Struktur $\langle \{m_1, m_2\}, \mathcal{I} \rangle$ mit $\mathcal{I}(P) = \{m_1\}$ und $\mathcal{I}(k) = m_2$ ist zwar ein Modell von $\exists x P(x)$, aber sie ist kein Modell von $P(k)$.

Da nicht jede Formel eine logisch äquivalente Skolem-Normalform hat, spricht man manchmal auch von Skolem-*Standardform*.

Es gilt aber Erfüllbarkeitsäquivalenz:

Theorem 9.26 *Es ist $\forall A$ erfüllbar unter einer Interpretation genau dann, wenn gilt: $\forall A_S$ ist erfüllbar in einer geeigneten erweiterten Interpretation, in der die durch Skolemisierung eingeführten Konstanten und Funktionszeichen interpretiert sind.*

Beweis. Siehe z. B. Ebbinghaus, Flum & Thomas (2007, Satz 8.4.6) oder Nienhuys-Cheng & de Wolf (1997, § 3). QED

10 Kalkül des natürlichen Schließens

Wir erweitern den aussagenlogischen Kalkül NK für die Quantorenlogik, indem wir Einführungs- und Beseitigungsregeln für die Quantoren $\forall$ und $\exists$ hinzufügen. Der resultierende Kalkül ist korrekt und vollständig für die Quantorenlogik.

Wir verzichten hier auf die Verwendung von Funktionstermen, d. h. Terme t sind ausschließlich Individuenkonstanten oder Variablen. Der Kalkül kann aber ohne Weiteres auch für die um Funktionsterme erweiterte Sprache verwendet werden.

10.1 Der quantorenlogische Kalkül NK

Zur Notation beliebiger Formeln A in Regeln und Ableitungen verwenden wir folgende Konventionen, die insbesondere auch der Formulierung bestimmter Variablenbedingungen für Regelanwendungen dienen.

Definition 10.1 (i) Wie gehabt bedeutet $A(x)$, dass die Variable x in der Formel A frei vorkommen kann.

(ii) Es bedeutet $A(t)$, dass der Term t in der Formel A vorkommen kann. Entsprechend auch für $A(t_1, \dots, t_n)$.

(iii) Im Kontext einer Formel $A(x)$ steht $A(t)$ für die Formel $A(x)[x/t]$.

Im Kontext einer Formel $A(t)$ bedeutet $A(x)$, dass alle Vorkommen des Terms t in A durch die Variable x ersetzt wurden.

Definition 10.2 (i) Der Kalkül NK des natürlichen Schließens (für klassische Quantorenlogik) besteht aus den in Definition 6.1 angegebenen aussagenlogischen Regeln zusammen mit den folgenden quantorenlogischen Regeln:

Kalkül NK

Einführungsregel	Beseitigungsregel
$\frac{A(y)}{\forall x A(x)} (\forall E)$ die Variable y darf in keiner Annahme frei vorkommen, von der $A(y)$ abhängt	$\frac{\forall x A(x)}{A(t)} (\forall B)$
$\frac{A(t)}{\exists x A(x)} (\exists E)$	$\frac{\begin{array}{c} [A(y)] \\ \exists x A(x) \end{array}}{C} (\exists B)$ die Variable y darf weder in C noch in einer Annahme außer $A(y)$ frei vorkommen, von der die Prämisse C abhängt

(ii) Die Variable y bei Anwendungen der Regeln $(\forall E)$ und $(\exists B)$ heißt *Eigenvariable*.

Eigenvariable

(iii) Die einschränkende Bedingung bezüglich der Eigenvariable y bei $(\forall E)$ und $(\exists B)$ heißt *Eigenvariablenbedingung*.

(iv) Die Prämisse $\exists x A(x)$ der Regel $(\exists B)$ heißt auch *Hauptprämisse*, die Prämisse C auch *Nebenprämisse*.

Wir erläutern zunächst die quantorenlogischen Regeln. Dabei geben wir diese zusätzlich in einer Form an, die ohne die in Definition 10.1 festgelegten Konventionen auskommt.

(i) Die All-Einführungsregel ($\forall E$) besagt, dass von der Prämisse $A(y)$ zur Konklusion $\forall x A(x)$ übergegangen werden darf, wobei jedes freie Vorkommen von y in A durch die Variable x ersetzt wurde.

Die Formel $A(y)$ kann man auch lesen als “ A trifft auf ein beliebiges y zu”, sofern y in A frei vorkommt. Dass y beliebig ist, wird gerade durch die Eigenvariablenbedingung ausgedrückt: Bei einer Anwendung von ($\forall E$) darf die Prämisse $A(y)$ von keiner Annahme abhängen, in der y frei vorkommt. (Andernfalls wäre y in $A(y)$ nicht beliebig, da zusätzliche Voraussetzungen bezüglich y bestünden.)

Man beachte, dass von einer offenen Formel $A(y)$, in der y frei vorkommt, als Annahme nie zu $\forall x A(x)$ übergegangen werden darf, da die Prämisse $A(y)$ als Annahme von sich selbst abhängt. In diesem Fall ist die Eigenvariablenbedingung also immer verletzt. (Vergleiche hierzu auch die auf Theorem 9.7 folgende **Bemerkung** zu $P(x) \not\models \forall x P(x)$.)

Alternative Formulierung der All-Einführungsregel:

$$\frac{A[x/y]}{\forall x A} (\forall E) \quad y = x \text{ oder } y \notin \text{FV}(A), \text{ und } y \text{ darf in keiner Annahme frei vorkommen, von der } A[x/y] \text{ abhängt.}$$

(ii) Die Allbeseitigungsregel ($\forall B$) besagt, dass von der Prämisse $\forall x A(x)$ zur Konklusion $A(t)$ übergegangen werden darf, wobei Vorkommen von x in A durch den Term t ersetzt werden.

Bei einer konkreten Anwendung ist die Konklusion entweder $A(y)$, für eine Variable y , oder $A(k)$, für eine Konstante k .

Alternative Formulierung der All-Beseitigungsregel:

$$\frac{\forall x A}{A[x/t]} (\forall B)$$

(iii) Die Existenz-einführungsregel ($\exists E$) besagt, dass von der Prämisse $A(t)$ zur Konklusion $\exists x A(x)$ übergegangen werden darf, wobei Vorkommen von t in A durch die Variable x ersetzt werden.

Bei einer konkreten Anwendung ist die Prämisse entweder $A(y)$, für eine Variable y , oder $A(k)$, für eine Konstante k .

Alternative Formulierung der Existenz-Einführungsregel:

$$\frac{A[x/t]}{\exists x A} (\exists E)$$

(iv) Die Existenzbeseitigungsregel ($\exists B$) besagt, dass von der Hauptprämisse $\exists x A(x)$ und der Nebenprämisse C zur Konklusion C übergegangen werden darf, wobei die Konklusion nicht mehr von Annahmen $A(y)$ abhängt, von denen die Nebenprämisse noch abhängen kann.

Die Eigenvariablenbedingung drückt wieder aus, dass y beliebig ist. Käme y in C frei vor, oder käme y in weiteren (d. h. von $A(y)$ verschiedenen) Annahmen frei vor, von denen die Nebenprämisse C abhängt, so würden über A hinaus zusätzliche Bedingungen bezüglich y bestehen. Mit anderen Worten, es würde dann nicht nur vorausgesetzt, dass

auf y die Eigenschaft A zutrifft, sondern dass auch jene Eigenschaften zutreffen, die in den zusätzlichen, von $A(y)$ verschiedenen Annahmen formuliert sind. Die Hauptprämisse sagt aber lediglich, dass es einen Gegenstand gibt, der die Eigenschaft A hat; über die Existenz von Gegenständen, auf die neben A noch weitere Eigenschaften zutreffen, sagt die Hauptprämisse nichts. Ein Übergang zur Konklusion wäre in diesem Fall also nicht gerechtfertigt.

Alternative Formulierung der Existenz-Beseitigungsregel:

$$\frac{\frac{[A[x/y]]}{\exists x A} \quad C}{C} (\exists B) \quad y = x \text{ oder } y \notin FV(A), \text{ und die Variable } y \text{ darf weder in } C \text{ noch in einer Annahme außer } A[x/y] \text{ frei vorkommen, von der die Nebenprämisse } C \text{ abhängt.}$$

Für die Definition von Ableitungen verwenden wir wieder die Konventionen gemäß Definition 10.1.

Definition 10.3 *Ableitungen* in NK definieren wir als Erweiterung von Definition 6.2 wie folgt. Sind $\mathcal{D}$ und $\mathcal{D}'$ Ableitungen, dann sind auch die folgenden Bäume Ableitungen:

Ableitung

$$\frac{\mathcal{D}}{A(y)} \quad \frac{\mathcal{D}}{\forall x A(x)} (\forall E) \quad \frac{\mathcal{D}}{\forall x A(x)} (\forall I) \quad \frac{\mathcal{D}}{\forall x A(x)} (\forall B) \quad \frac{\mathcal{D}}{A(t)} (\forall B)$$

die Variable y darf in keiner Annahme frei vorkommen, von der $A(y)$ abhängt

$$\frac{\mathcal{D}}{A(t)} (\exists E) \quad \frac{\mathcal{D}}{\exists x A(x)} (\exists E) \quad \frac{\mathcal{D}}{\exists x A(x)} (\exists I) \quad \frac{\mathcal{D}}{\exists x A(x)} (\exists B) \quad \frac{\mathcal{D}}{\exists x A(x)} (\exists B)^n$$

die Variable y darf weder in C noch in einer Annahme von $\mathcal{D}'$ außer $A(y)$ frei vorkommen, von der die Nebenprämisse C abhängt

Die in den Definitionen 6.3-6.6 festgelegten Begriffe und Notationen können direkt übernommen werden (wobei wir davon ausgehen, dass die rekursive Definition 6.5 von *Annahmenmenge* in offensichtlicher Weise für Ableitungen gemäß Definition 10.3 erweitert wurde). Die in Theorem 6.8 angegebenen Strukturregeln gelten entsprechend auch hier.

Beispiele. (i) Wir zeigen $\forall x(P(x) \rightarrow Q(x)), P(k) \vdash Q(k)$ durch Angabe einer Ableitung in NK:

$$\frac{\frac{\forall x(P(x) \rightarrow Q(x))}{P(k) \rightarrow Q(k)} (\forall B) \quad P(k)}{Q(k)} (\rightarrow B)$$

(Man vergleiche dies mit dem Nachweis von $\forall x(P(x) \rightarrow Q(x)), P(k) \models Q(k)$ in § 8.5.)

(ii) Wir zeigen $\exists x P(x) \vdash \exists y P(y)$:

$$\frac{\exists x P(x) \quad \frac{[P(y)]^1}{\exists y P(y)} (\exists E)}{\exists y P(y)} (\exists B)^1$$

(Vergleiche Theorem 9.15 (ii).)

(iii) Wir zeigen $\neg \forall x \neg A(x) \vdash \exists x A(x)$:

$$\frac{\neg \forall x \neg A(x) \quad \frac{\frac{\frac{[A(y)]^1}{\exists x A(x)} (\exists E)}{[\neg \exists x A(x)]^2} (\rightarrow B)}{\frac{\perp}{\neg A(y)} (\rightarrow E)^1} (\rightarrow B)}{\frac{\perp}{\forall x \neg A(x)} (\forall E)} (\rightarrow B)$$

$$\frac{\perp}{\exists x A(x)} (\perp)^2$$

(Vergleiche Theorem 9.4 (iv).)

Bemerkung. Die Ableitung in Beispiel (iii) zeigt, dass prinzipiell auch dann auf eine Existenzaussage $\exists x A(x)$ geschlossen werden kann, wenn überhaupt kein Gegenstand vorgewiesen wird, der die Eigenschaft A hat. Um die Existenz eines Gegenstands mit der Eigenschaft A zu zeigen, genügt es zu zeigen, dass es nicht der Fall ist, dass allen Gegenständen die Eigenschaft A nicht zukommt (hier als Annahme $\neg \forall x \neg A(x)$). Mit anderen Worten, $\exists x A(x)$ muss *nicht* (und kann auch nicht) in jedem Fall per Existenzführungsregel ($\exists E$) mit Prämisse $A(k)$ abgeleitet werden, d. h. aus der Aussage, dass der Gegenstand k die Eigenschaft A hat:

$$\frac{\mathcal{D} \quad A(k)}{\exists x A(x)} (\exists E)$$

(bzw. mit der Prämisse $A(y)$, die besagt, dass ein beliebiger Gegenstand die Eigenschaft A hat).

Dies ist ein Merkmal der *klassischen* Quantorenlogik. Hingegen gilt in der *konstruktiven* bzw. *intuitionistischen Logik* für jede geschlossene Formel $\exists x A(x)$ die sogenannte *Existenzeigenschaft*:

Wenn $\exists x A(x)$ beweisbar ist, dann ist auch $A(k)$ beweisbar, für eine Konstante k .

Einen Kalkül für die intuitionistische Logik erhält man durch Abschwächung der (klassischen) Widerspruchsregel ($\perp$) zur (intuitionistischen) *ex falso-Regel*

$$\frac{\perp}{A}$$

bei der im Unterschied zur Widerspruchsregel Annahmen der Form $\neg A$ nicht gelöscht werden können. Dieser Kalkül ist also schwächer als NK.

Weitere Beispiele. (iv) Wir zeigen $\exists x \forall y A(x, y) \vdash \forall y \exists x A(x, y)$ durch Angabe einer Ableitung in NK:

$$\frac{\frac{\frac{[\forall y A(x, y)]^1}{A(x, y)} (\forall B)}{\exists x A(x, y)} (\exists E)}{\frac{\exists x \forall y A(x, y)}{\exists x A(x, y)} (\exists B)^1} \frac{}{\forall y \exists x A(x, y)} (\forall E)$$

Im rechten Zweig könnte von $\forall y A(x, y)$ mit $(\forall B)$ zu $A(x, k)$ übergegangen werden, für eine Konstante k . Dann dürfte jedoch im letzten Schritt von $\exists x A(x, k)$ nicht mit $(\forall E)$ auf $\forall y \exists x A(x, y)$ geschlossen werden, da bei $(\forall E)$ eine freie Variable verlangt wird.

(Vergleiche Theorem 9.6.)

(v) Wir zeigen, dass $\exists x A(x) \rightarrow B \vdash \forall x (A(x) \rightarrow B)$, falls die Variable x nicht frei in B vorkommt:

$$\frac{\frac{\frac{[A(x)]^1}{\exists x A(x)} (\exists E)}{B} (\rightarrow B)}{\frac{A(x) \rightarrow B}{\forall x (A(x) \rightarrow B)} (\forall E)} (\rightarrow E)^1$$

Statt der Annahme $A(x)$ hätte man auch die (weniger allgemeine) Annahme $A(k)$, für eine Konstante k , als Prämisse von $(\exists E)$ wählen können. Allerdings könnte dann von $A(k) \rightarrow B$ nicht mit $(\forall E)$ auf die Endformel geschlossen werden, da $(\forall E)$ in der Prämisse eine freie Variable verlangt.

Die Einschränkung, dass die Variable x nicht in B vorkommt, ist nötig, da x Eigenvariable von $(\forall E)$ ist. Käme x in B frei vor, würde die Prämisse von $(\forall E)$ von einer Annahme der Form $\exists x A(x) \rightarrow B(x)$ abhängen, in der x frei vorkommt. Damit wäre die Eigenvariablenbedingung verletzt. (Die Annahme $A(x)$ ist hingegen unproblematisch, da diese bei der Anwendung von $(\rightarrow E)$ geschlossen wurde; die Prämisse von $(\forall E)$ hängt also nicht mehr von der Annahme $A(x)$ ab.)

(Vergleiche Theorem 9.17 (ii).)

10.2 Zur Notwendigkeit der Eigenvariablenbedingung bei $(\forall E)$ und $(\exists B)$

Um die Notwendigkeit der Eigenvariablenbedingung bei den Regeln $(\forall E)$ und $(\exists B)$ noch etwas deutlicher zu machen, betrachten wir folgende *inkorrekte* Ableitungen für die Ableitbarkeitsbehauptungen $\exists x P(x) \vdash \forall x P(x)$ und $\exists x P(x), \exists x \neg P(x) \vdash \perp$. Die Eigenvariablenbedingung wird dabei jeweils verletzt. (Es gibt keine korrekten Ableitungen, die diese Ableitbarkeitsbehauptungen beweisen würden, d. h. es ist $\exists x P(x) \not\vdash \forall x P(x)$ und $\exists x P(x), \exists x \neg P(x) \not\vdash \perp$.)

(i) *Inkorrekte* Ableitung für $\exists x P(x) \vdash \forall x P(x)$:

$$\frac{\frac{[P(y)]^1}{\exists x P(x)} (\forall E) \not\vdash}{\forall x P(x)} (\exists B)^1 \not\vdash$$

Die Anwendung von $(\forall E)$ ist *inkorrekt*, da an dieser Stelle $P(y)$ eine offene Annahme ist,

(ii) Weitere *inkorrekte* Ableitung für $\exists xP(x) \vdash \forall xP(x)$:

$$\frac{\exists x P(x) \quad [P(y)]^1}{\frac{P(y)}{\forall x P(x)} (\forall E)} (\exists B)^1 \downarrow$$

(iii) Man sieht hier auch, dass

$$\frac{\exists x A(x)}{A(t)} \downarrow$$

$$\frac{\frac{\exists x P(x)}{P(y)} \downarrow}{\forall x P(x)} (\forall E) \quad \downarrow$$

(iv) Eine weitere *inkorrekte* Ableitung für $\exists xP(x) \vdash \forall xP(x)$ ist:

$$\frac{\frac{\frac{[\neg P(y)]^2 \quad [P(y)]^1}{\exists x P(x)} (\rightarrow B) \quad \perp}{(\exists B)^1} \not\vdash \quad \perp}{\frac{\perp}{P(y)} (\bot)^2 \quad \perp}{\forall x P(x)} (\forall E)$$

(v) *Inkorrekte Ableitung* für $\exists xP(x), \exists x\neg P(x) \vdash \perp$:

$$\frac{\frac{\frac{\exists x \neg P(x)}{\exists x P(x)} \perp}{\exists x \neg P(x)} \perp}{\exists x P(x)} \perp \quad \frac{\frac{\frac{[\neg P(y)]^2}{[P(y)]^1} (\rightarrow B)}{\exists x P(x)} \perp}{(\exists B)^1} \perp \quad \frac{\perp}{(\exists B)^2} \perp$$

104

10.3 Definierbarkeit von $\forall$ bzw. $\exists$

Statt Regelpaare sowohl für $\forall$ als auch für $\exists$ einzuführen, genügt es, nur für einen der beiden Quantoren ein Regelpaar anzugeben.

Gibt man z. B. Regeln für den Allquantor $\forall$ an, dann kann $\exists xA(x)$ durch $\neg\forall x\neg A(x)$ definiert werden. Der resultierende Kalkül sei $NK_{\forall}$, und die durch ihn gegebene Ableitbarkeitsrelation sei $\vdash_{\forall}$. Damit $NK_{\forall}$ gleichwertig zu NK ist (d. h., damit $\vdash_{\forall} = \vdash$ gilt), muss Folgendes gelten:

Theorem 10.4 (i) $A(t) \vdash_{\forall} \exists xA(x)$, für beliebige Terme t .

(ii) Wenn $\Gamma, A(y) \vdash_{\forall} C$, dann $\Gamma, \exists xA(x) \vdash_{\forall} C$, wobei Γ eine Menge von Annahmen ist, und die Variable y nicht in C und in keiner Annahme in Γ frei vorkommt, von der C abhängt.

Beweis. (i) Sei t ein beliebiger Term. Dann ist

$$\frac{\frac{[\forall x\neg A(x)]^1}{\neg A(t)} (\forall B) \quad A(t)}{\perp} (\rightarrow B) \quad \frac{\perp}{\neg\forall x\neg A(x)} (\rightarrow E)^1$$

eine Ableitung in $NK_{\forall}$. Somit gilt mit $\exists xA(x) := \neg\forall x\neg A(x)$, dass $A(t) \vdash_{\forall} \exists xA(x)$.

Man erhält also die NK-Regel $\frac{A(t)}{\exists xA(x)} (\exists E)$.

(ii) Sei $\frac{\Gamma, A(y)}{C}$ eine Ableitung von C aus Γ und $A(y)$, wobei Γ eine Menge von Annahmen ist, und die Variable y nicht in C und in keiner Annahme in Γ frei vorkommt, von der C abhängt. Dann gilt wegen

$$\frac{\frac{\frac{\Gamma, [A(y)]^1}{C} (\rightarrow B) \quad [\neg C]^2}{\perp} (\rightarrow B) \quad \frac{\perp}{\neg A(y)} (\rightarrow E)^1}{\frac{\neg\forall x\neg A(x)}{\forall x\neg A(x)} (\forall E)} (\rightarrow B) \quad \frac{\perp}{C} (\perp)^2$$

unter Verwendung von $\exists xA(x) := \neg\forall x\neg A(x)$, dass $\Gamma, \exists xA(x) \vdash_{\forall} C$.

Man erhält also die NK-Regel

$$\frac{\frac{[A(y)]}{\exists xA(x)} C (\exists B)}{C} (\exists B)$$

mit entsprechender Eigenvariablenbedingung.

QED

Entsprechend kann man sich auch auf den Existenzquantor $\exists$ beschränken, und dann

$\forall x A(x)$ durch $\neg \exists x \neg A(x)$ definieren. Der resultierende Kalkül sei $NK_{\exists}$; seine Ableitbarkeitsrelation sei $\vdash_{\exists}$. Damit $NK_{\exists}$ gleichwertig zu NK ist (d. h., damit $\vdash_{\exists} = \vdash$ gilt), muss Folgendes gelten:

Theorem 10.5 (i) Wenn $\Gamma \vdash_{\exists} A(y)$, dann $\Gamma \vdash_{\exists} \forall x A(x)$, wobei Γ eine Menge von Annahmen ist, und die Variable y in keiner Annahme in Γ frei vorkommt, von der $A(y)$ abhängt.

(ii) $\forall x A(x) \vdash_{\exists} A(t)$, für beliebige Terme t .

Beweis.

(i) Sei $\mathcal{D}$ eine Ableitung von $A(y)$ aus Γ , wobei Γ eine Menge von Annahmen ist, und die Variable y in keiner Annahme in Γ frei vorkommt, von der $A(y)$ abhängt. Dann gilt wegen

$$\frac{\frac{\frac{\Gamma}{\mathcal{D}} \quad \frac{[\neg A(y)]^1}{A(y)} (\rightarrow B)}{[\exists x \neg A(x)]^2} \quad \frac{\perp}{(\exists B)^1}}{\frac{\perp}{\neg \exists x \neg A(x)} (\rightarrow E)^2}$$

unter Verwendung von $\forall x A(x) := \neg \exists x \neg A(x)$, dass $\Gamma \vdash_{\exists} \forall x A(x)$.

Man erhält also die NK-Regel

$$\frac{A(y)}{\forall x A(x)} (\forall E)$$

mit entsprechender Eigenvariablenbedingung.

(ii) Sei t ein beliebiger Term. Dann ist

$$\frac{\frac{\neg \exists x \neg A(x) \quad \frac{[\neg A(t)]^1}{\exists x \neg A(x)} (\exists E)}{(\rightarrow B)} \quad \frac{\perp}{A(t)} (\perp)^1}{\perp}$$

eine Ableitung in $NK_{\exists}$. Somit gilt mit $\forall x A(x) := \neg \exists x \neg A(x)$, dass $\forall x A(x) \vdash_{\exists} A(t)$.

Man erhält also die NK-Regel $\frac{\forall x A(x)}{A(t)} (\forall B)$.

QED

10.4 Korrektheit und Vollständigkeit

Der quantorenlogische Kalkül NK ist für den semantischen Begriff der quantorenlogischen Folgerung korrekt und vollständig. Es gilt:

Theorem 10.6 (Korrektheit und Vollständigkeit)

(i) Korrektheit von NK : Wenn $\Gamma \vdash A$, dann $\Gamma \models A$.

(ii) Vollständigkeit von NK : Wenn $\Gamma \models A$, dann $\Gamma \vdash A$.

Wie schon in der Aussagenlogik, gilt auch hier:

Korollar 10.7 (i) *Ist $\Gamma = \emptyset$, dann gilt $\vdash A$ genau dann, wenn $\models A$. Das heißt, eine Formel ist genau dann beweisbar, wenn sie allgemeingültig ist.*

(ii) *Sei Γ eine beliebige (möglicherweise unendliche) Menge von Formeln, und A eine beliebige Formel. Falls $\Gamma \models A$ gilt, gibt es eine endliche Teilmenge Γ' von Γ , so dass $\Gamma' \models A$ gilt.*

Bemerkungen. (i) Die Zusammenfassung von Korrektheit und Vollständigkeit in der Aussage “ $\Gamma \vdash A$ genau dann, wenn $\Gamma \models A$ ” bezeichnet man oft einfach als *Vollständigkeitssatz*.

(ii) Den ersten Beweis eines Vollständigkeitssatzes (für den sogenannten engeren Funktionenkalkül, der unserer Quantorenlogik entspricht) gab Gödel (1929).

(iii) Für einen Beweis des Vollständigkeitssatzes für den Kalkül des natürlichen Schließens siehe z. B. van Dalen (2013, § 4.1).

11 Quantorenlogische Resolution

Wir erweitern aussagenlogische Resolution für die Sprache der Quantorenlogik. Dabei lassen wir auch wieder Funktionszeichen in unserer Sprache zu.

Zunächst behandeln wir quantorenlogische Resolution als Widerlegungsverfahren für beliebige quantorenlogische Formeln, danach SLD-Resolution. Letztere stellt die Grundlage für die sogenannte Logikprogrammierung dar. Bei dieser werden Mengen logischer Formeln als Programme aufgefasst, die ein bestimmtes Problem spezifizieren. Mittels SLD-Resolution können Lösungen des Problems dann berechnet werden.

Ist A eine beliebige quantorenlogische Formel und A_S deren Skolem-Normalform, so gilt für $\forall A$ und $\forall A_S$ Erfüllbarkeitsäquivalenz (Theorem 9.26): Die Formel $\forall A$ ist erfüllbar unter einer Interpretation genau dann, wenn $\forall A_S$ in einer geeigneten erweiterten Interpretation erfüllbar ist, in der die durch Skolemisierung eingeführten Konstanten und Funktionszeichen interpretiert sind. Eine Widerlegung von $\forall A$ kann somit durch eine Widerlegung von $\forall A_S$ erfolgen:

Korollar 11.1 (i) Um zu zeigen, dass $\forall A$ unerfüllbar ist, d. h., um $\forall A$ zu widerlegen, reicht es aus, $\forall A_S$ zu widerlegen.

(ii) Sei $C_1 \wedge \dots \wedge C_m$ eine KNF von A_S , d. h. die konjunktive Normalform der Skolem-Normalform (kurz: konjunktive Skolem-Normalform) der Ausgangsformel A . Um zu zeigen, dass $\forall A$ unerfüllbar ist, reicht es aus, $\forall C_1 \wedge \dots \wedge \forall C_m$ zu widerlegen.

Die konjunktive Skolem-Normalform $C_1 \wedge \dots \wedge C_m$ von A kann man (wie in der Aussagenlogik) als *Klauselmenge* schreiben, die wir mit A_{SK} (*konjunktive Skolem-Normalform in Klauselform*) bezeichnen.

Literale seien nun quantorenlogische Atome $R(t_1, \dots, t_n)$ (für $n \geq 0$) oder deren Negation, womit *Klauseln* definiert sind wie bisher.

Definiert man ein *Widerlegungsverfahren* für Klauseln so, dass man dabei die freien Variablen in Klauseln universell versteht (siehe § 11.1), dann gilt: Um zu zeigen, dass $\forall A$ unerfüllbar ist, genügt es, A_{SK} nach diesem Verfahren zu widerlegen; zum Nachweis der Allgemeingültigkeit einer geschlossenen Formel A genügt es, $(\neg A)_{SK}$ nach diesem Verfahren zu widerlegen.

Für eine beliebige quantorenlogische Formel A sind also die folgenden drei Schritte auszuführen:

- (1) Bilde eine konjunktive Skolem-Normalform für die negierte Ausgangsformel $\neg A$. (Dazu überführt man $\neg A$ entweder in eine Skolem-Normalform, die anschließend in eine KNF umgeformt wird, oder man formt $\neg A$ zunächst in eine PNF mit Kern in KNF um und skolemisiert dann.)
- (2) Bilde die entsprechende Klauselmenge $(\neg A)_{SK}$.
- (3) Wende das Widerlegungsverfahren auf die Klauselmenge an.

Im Fall von Folgerungsbehauptungen $A_1, \dots, A_n \models B$ geht man wie folgt vor:

- (1) Bilde $(A_1)_{SK}, \dots, (A_n)_{SK}$ und $(\neg B)_{SK}$. Hierbei sind die durch die jeweilige Skolemisierung eingeführten Terme so zu wählen, dass die Mengen dieser Terme für $(A_1)_{SK}, \dots, (A_n)_{SK}$ und $(\neg B)_{SK}$ paarweise disjunkt sind.
- (2) Bilde die Klauselmenge $\bigcup_{1 \leq i \leq n} (A_i)_{SK} \cup (\neg B)_{SK}$.
- (3) Wende das Widerlegungsverfahren auf die Klauselmenge an.

Klauselmenge

Literal

Klausel

Widerlegungsverfahren

Beispiele. (i) $\models \exists x \forall y P(x, y) \rightarrow \forall y \exists x P(x, y)$

$$\begin{aligned}
& \neg(\exists x \forall y P(x, y) \rightarrow \forall y \exists x P(x, y)) && \text{(Negation der Ausgangsformel)} \\
& \rightsquigarrow \neg(\exists x \forall y P(x, y) \rightarrow \forall z \exists u P(u, z)) && \text{(gebundene Umbenennung)} \\
& \rightsquigarrow \exists x \forall y P(x, y) \wedge \exists z \forall u \neg P(u, z) && \text{(Negation nach innen gezogen)} \\
& \rightsquigarrow \exists x \forall y \exists z \forall u (P(x, y) \wedge \neg P(u, z)) && \text{(PNF mit Kern in KNF)} \\
& \rightsquigarrow \forall y \exists z \forall u (P(k, y) \wedge \neg P(u, z)) && \text{(Skolemisiere . . .)} \\
& \rightsquigarrow \exists z \forall u (P(k, y) \wedge \neg P(u, z)) \\
& \rightsquigarrow \forall u (P(k, y) \wedge \neg P(u, f(y))) \\
& \rightsquigarrow P(k, y) \wedge \neg P(u, f(y)) && \text{(Skolem-Normalform in KNF)} \\
& \rightsquigarrow \{ \vdash P(k, y) ; P(u, f(y)) \vdash \} && \text{(Klauselmenge)}
\end{aligned}$$

(ii) $\exists x \forall y P(x, y) \models \forall y \exists x P(x, y)$

$$\begin{aligned}
& \exists x \forall y P(x, y) \rightsquigarrow \forall y P(k, y) \rightsquigarrow P(k, y) \\
& \neg \forall y \exists x P(x, y) \rightsquigarrow \exists y \forall x \neg P(x, y) \rightsquigarrow \forall x \neg P(x, l) \rightsquigarrow \neg P(x, l)
\end{aligned}$$

Man erhält die Klauselmenge $\{ \vdash P(k, y) ; P(x, l) \vdash \}$.

Aufgrund des Import-Export-Theorems 9.2 gilt die Behauptung der Allgemeingültigkeit in (i) genau dann, wenn die Folgerungsbehauptung in (ii) gilt. Dennoch unterscheiden sich die resultierenden Klauselmengen: Bei (ii) steht die Konstante l anstelle des Funktionsterms $f(y)$. Das ist jedoch unproblematisch, da die freien Variablen in jeder Klausel universell verstanden werden.

Obgleich beide Klauselmengen unerfüllbar sind, ist eine Resolutionswiderlegung mit der aussagenlogischen Resolutionsregel nicht möglich, da $P(k, y)$ und $P(u, f(y))$, bzw. $P(k, y)$ und $P(x, l)$, jeweils zwei verschiedene Formeln sind. (Dass dies keine aussagenlogischen, sondern quantorenlogische Formeln sind, ist hier nebensächlich.)

11.1 Unifikation und quantorenlogische Resolution

Um auch für unerfüllbare Klauselmengen wie jene im letzten Beispiel eine Resolutionswiderlegung erhalten zu können, müssen Atome A und B , für die $\{A, \neg B\}$ unerfüllbar ist, in geeigneter Weise durch *Unifikation* vereinheitlicht werden. Dies erreicht man durch Angabe einer Instanz A' von A und einer Instanz B' von B , so dass A' und B' syntaktisch gleich sind.

- Definition 11.2** (i) Ist der Ausdruck (d. h. die Formel oder der Term) E' eine Instanz des Ausdrucks E (d. h. $E' = E\sigma$, für eine Substitution σ), dann heißt E *allgemeiner* als E' . *allgemeiner*
- (ii) Sind σ, τ zwei Substitutionen, dann heißt σ *allgemeiner* als τ , wenn es eine Substitution ϑ gibt, so dass $\sigma\vartheta = \tau$.
- (iii) Es ist $\text{dom}(\vartheta) := \{x \mid x\vartheta \neq x\}$.
- (iv) Eine Substitution ϑ heißt *Variablensubstitution*, falls alle $x_i\vartheta$ für $x_i \in \text{dom}(\vartheta)$ Variablen sind. *Variablensubstitution*
- (v) Eine Variablensubstitution ϑ ist eine *Umbenennung*, falls ϑ die Variablen eineindeutig aufeinander abbildet. (Eine Umbenennung ist also eine Permutation von Variablen.) *Umbenennung*

(vi) Ist ϑ eine Umbenennung, so heißt $E\vartheta$ *Variante* von E .

Variante

Beispiele. (i) Es ist x allgemeiner als $g(k, h(l))$, denn für $[x/g(k, h(l))]$ ist $g(k, h(l))$ eine Instanz von x .

(ii) Es ist $f(x, y)$ allgemeiner als $f(x, x)$, denn $f(x, y)[y/x] = f(x, x)$.

(iii) $[x/y]$ ist allgemeiner als $[x/k, y/k]$, da $[x/y][y/k] = [x/k, y/k]$.

(iv) Es ist σ allgemeiner als σ für jede Substitution σ , da $\sigma\varepsilon = \sigma$ für die leere Substitution ε .

(v) $[x/y]$ ist *nicht* allgemeiner als $[x/k]$.

Angenommen, es gäbe eine Substitution ϑ , so dass die Bindung x/k in $[x/y]\vartheta$ enthalten ist. Dann muss y/k in ϑ enthalten sein, und damit $y \in \text{dom}([x/y]\vartheta)$. Folglich kann es keine Substitution ϑ mit $[x/y]\vartheta = [x/k]$ geben.

(vi) $[x/z, y/z]$ ist eine Variablensubstitution.

(vii) $[x/z, z/y, y/x]$ ist eine Umbenennung.

Die leere Substitution ε ist ebenfalls eine Umbenennung.

(viii) Es ist $f(x, y)$ eine Variante von $f(y, x)$, denn $f(y, x)[y/x, x/y] = f(x, y)$.

(ix) Es ist $f(x, z)$ eine Variante von $f(x, y)$, denn $f(x, y)[y/z, z/y] = f(x, z)$.

Die Bindung z/y ist nötig, damit die angegebene Substitution eine Umbenennung gemäß unserer Definition ist.

(x) Es ist $f(x, x)$ *keine* Variante von $f(x, y)$.

Wäre es eine Variante, dann müsste $f(x, y)\vartheta = f(x, x)$ für eine Umbenennung ϑ mit Bindung y/x gelten. Da ϑ eine Umbenennung ist, muss ϑ aber außerdem eine Bindung x/y mit $x \neq y$ enthalten. Doch dann wäre $f(x, y)\vartheta = f(y, x) \neq f(x, x)$.

Definition 11.3 (i) Eine Substitution σ heißt *Unifikator* von zwei Atomen A und B , wenn gilt: $A\sigma = B\sigma$ (d. h., wenn die beiden Ausdrücke $A\sigma$ und $B\sigma$ syntaktisch gleich sind).

Unifikator

(ii) Gibt es einen Unifikator σ von A und B , dann heißen die Atome A und B *unifizierbar*. Ist $\Gamma = \{A, B\}$ für zwei Atome A und B mit Unifikator σ , so sagen wir auch, dass σ die Menge Γ *unifiziert*.

unifizierbar

Beispiel. Die beiden Atome $P(k, x)$ und $P(y, l)$ sind unifizierbar: die Substitution $\sigma = [x/l, y/k]$ ist ein Unifikator, da $P(k, x)\sigma = P(y, l)\sigma$.

Definition 11.4 Die *quantorenlogische Resolutionsregel* ist

quantorenlogische Resolutionsregel

$$\frac{X_1 \vdash Y_1, A \quad B, X_2 \vdash Y_2}{(X_1, X_2 \vdash Y_1, Y_2)\sigma} (\mathcal{R})$$

wobei σ ein Unifikator von A und B ist, und die Mengen der freien Variablen der beiden Prämissen disjunkt sein müssen.

Beispiel. Die beiden Klauseln $Q(x) \vdash P(k, x)$ und $P(y, l) \vdash R(y)$ haben keine gemeinsamen Variablen, und die Substitution $[x/l, y/k]$ unifiziert $P(k, x)$ und $P(y, l)$. Somit ist

$$[x/l, y/k] \frac{Q(x) \vdash P(k, x) \quad P(y, l) \vdash R(y)}{Q(l) \vdash R(k)} (\mathcal{R})$$

eine korrekte Anwendung der quantorenlogische Resolutionsregel. (Den verwendeten Unifikator haben wir neben dem Regelstrich notiert.)

Ist σ eine Substitution, die A und B unifiziert, so führt die Anwendung des Unifikators σ auf die Formeln in der quantorenlogischen Resolutionsregel der Form nach auf die aussagenlogische Resolutionsregel:

$$\frac{\begin{array}{c} X_1 \vdash Y_1, A \\ \downarrow \sigma \\ X_1\sigma \vdash Y_1\sigma, A\sigma (= B\sigma) \end{array} \quad \begin{array}{c} B, X_2 \vdash Y_2 \\ \downarrow \sigma \\ B\sigma (= A\sigma), X_2\sigma \vdash Y_2\sigma \end{array}}{X_1\sigma, X_2\sigma \vdash Y_1\sigma, Y_2\sigma}$$

(Weil die Formeln hier Formeln in der Sprache der Quantorenlogik sind, erhält man nicht die eigentliche aussagenlogische Resolutionsregel, da diese nur für aussagenlogische Formeln angegeben wurde. Das ist hier aber nicht wesentlich.)

Beispiel. Für $\sigma = [x/l, y/k]$ als Unifikator von $P(k, x)$ und $P(y, l)$ erhält man:

$$\frac{\begin{array}{c} Q(x) \vdash P(k, x) \\ \downarrow \sigma \\ Q(l) \vdash P(k, l) \end{array} \quad \begin{array}{c} P(y, l) \vdash R(y) \\ \downarrow \sigma \\ P(k, l) \vdash R(k) \end{array}}{Q(l) \vdash R(k)} (\mathcal{R})$$

Da freie Variablen universell verstanden werden, bedeutet die Forderung disjunkter Mengen freier Variablen der beiden Prämissen bei $(\mathcal{R})$ keine Einschränkung der Allgemeinheit. Die Forderung kann immer durch geeignete Ersetzungen der freien Variablen erfüllt werden, durch die freie Variablen in verschiedenen Klauseln *separiert* werden. (Vergleiche auch Korollar 11.1 (ii).)

Separierung freier Variablen

Betrachtet man z. B. die beiden Klauseln

$$\vdash P(k), Q(x) \quad \text{und} \quad P(k) \vdash R(x)$$

dann erhält man per (inkorrekt) Resolution

$$\frac{\vdash P(k), Q(x) \quad P(k) \vdash R(x)}{\vdash Q(x), R(x)} \not\vdash$$

die Klausel $\vdash Q(x), R(x)$ als Resolvente; der Resolutionsschritt ist inkorrekt, da x in beiden Prämissen als freie Variable vorkommt. Separiert man die freien Variablen stattdessen zuerst, z. B. durch die Umbenennung

$$P(k) \vdash R(x) [x/y] = P(k) \vdash R(y)$$

so erhält man als Resolvente die Klausel $\vdash Q(x), R(y)$. Die beiden Resolventen sind nicht logisch äquivalent, da

$$Q(x) \vee R(y) \models Q(x) \vee R(x) \quad \text{aber} \quad Q(x) \vee R(x) \not\models Q(x) \vee R(y).$$

Das Problem besteht darin, dass die Variable x in beiden Ausgangsklauseln vorkommt, obgleich die beiden Klauseln voneinander unabhängig sind. Es handelt sich deshalb eigentlich nicht um zwei Vorkommen derselben Variable x , sondern um zwei verschiedene Variablen mit demselben Namen. Diese Unterscheidung zweier Variablen geht aber im obigen (inkorrekten) Resolutionsschritt verloren.

Um stets die allgemeinsten Resolventen zu erhalten, werden Variablen in den Prämissen immer erst durch Umbenennung separiert.

Beispiele. (i) Jetzt kann das Beispiel (i) bzw. (ii) von S. 109 fortgesetzt werden:

$$\{ \vdash P(k, y) ; P(u, f(y)) \vdash \} \xrightarrow[\text{Separierung der Variablen}]{\sim} \{ \vdash P(k, y) ; P(u, f(v)) \vdash \}$$

Resolutionswiderlegung:

$$[y/f(v), u/k] \frac{\vdash P(k, y) \quad P(u, f(v)) \vdash}{\vdash} (\mathcal{R})$$

Beziehungsweise für $\{ \vdash P(k, y) ; P(x, l) \vdash \}$:

$$[x/k, y/l] \frac{\vdash P(k, y) \quad P(x, l) \vdash}{\vdash} (\mathcal{R})$$

(ii) $\exists x(P(x) \wedge \neg Q(x)), \forall x(P(x) \rightarrow R(x)) \models \exists x(R(x) \wedge \neg Q(x))$

$$\begin{array}{lll} \exists x(P(x) \wedge \neg Q(x)) & \forall x(P(x) \rightarrow R(x)) & \neg \exists x(R(x) \wedge \neg Q(x)) \\ \sim \{ Q(k) \vdash ; \vdash P(k) \} & \sim \forall x(\neg P(x) \vee R(x)) & \sim \forall x(\neg R(x) \vee Q(x)) \\ & \sim \{ P(x) \vdash R(x) \} & \sim \{ R(y) \vdash Q(y) \} \end{array}$$

Man kann separat skolemisieren, da die Variablen separiert werden können.

Als Klauselmenge erhält man $\{ Q(k) \vdash ; \vdash P(k) ; P(x) \vdash R(x) ; R(y) \vdash Q(y) \}$.

Resolutionswiderlegung:

$$\begin{array}{c} [y/k] \frac{\vdash P(k)}{\vdash Q(k)} \quad [x/y] \frac{P(x) \vdash R(x) \quad R(y) \vdash Q(y)}{P(y) \vdash Q(y)} (\mathcal{R}) \\ \varepsilon \frac{\vdash Q(k) \quad P(y) \vdash Q(y)}{\vdash} (\mathcal{R}) \end{array}$$

Die quantorenlogische Resolutionsregel ($\mathcal{R}$) allein reicht nicht aus, um z. B. aus den beiden Klauseln

$$\vdash P(x), P(y) \quad \text{und} \quad P(z), P(u) \vdash$$

die leere Klausel ableiten zu können, obgleich die entsprechende Menge

$$\{ P(x) \vee P(y), \neg P(z) \vee \neg P(u) \}$$

unerfüllbar ist, da freie Variablen hier universell verstanden werden. Denn jeder Resolutionsschritt erzeugt hier wieder eine Klausel mit zwei Atomen, z. B.

$$[y/z] \frac{\vdash P(x), P(y) \quad P(z), P(u) \vdash}{P(u) \vdash P(x)} (\mathcal{R})$$

Da in jedem Resolutionsschritt die Mengen der freien Variablen der beiden Prämissen disjunkt sein müssen, könnte man nur mit geeigneten Varianten fortfahren; z. B. so:

$$[y/u] \frac{\vdash P(v), P(y) \quad [y/z] \frac{\vdash P(x), P(y) \quad P(z), P(u) \vdash}{P(u) \vdash P(x)} (\mathcal{R})}{\vdash P(v), P(u)} (\mathcal{R})$$

wobei $\vdash P(v), P(u)$ eine Variante der ersten Klausel $\vdash P(x), P(y)$ ist. Dies führt jedoch ebenfalls nur zu einer Klausel mit zwei Atomen, die in diesem Fall die Form der ersten Ausgangsklausel hat. Entsprechend erhält man auch in allen anderen Fällen stets Klauseln mit zwei Atomen.

Um ein korrektes Resolutionsverfahren zu erhalten, müssen deshalb zusätzlich sogenannte *Faktoren* von Klauseln gebildet werden können.

Definition 11.5 Ein *Faktor* S' einer Klausel S ist das Ergebnis der Anwendung einer Substitution auf S , durch die mehrere Atome in S unifiziert werden. *Faktor*

Die Ableitung von Faktoren ermöglichen wir durch Hinzunahme folgender Regel.

Definition 11.6 Die *Faktorisierungsregel* ($\mathcal{F}$) ist als Regelpaar für faktorisierbare Atome A, B im Antezedens bzw. Sukzedens gegeben durch: *Faktorisierungsregel*

$$\frac{X, A, B \vdash Y}{(X, A \vdash Y)\sigma} (\mathcal{F}) \quad \frac{X \vdash A, B, Y}{(X \vdash A, Y)\sigma} (\mathcal{F})$$

wobei σ ein Unifikator von A und B ist.

Die Faktorisierungsregel ist eine Substitutionsregel, bei deren Anwendung A und B durch Unifikation identifiziert werden. Dies ist dadurch gerechtfertigt, dass Variablen in Klauseln universell aufgefasst werden.

Beispiel. Mithilfe der Faktorisierung ist es möglich, eine Resolutionswiderlegung von

$$\{ \vdash P(x), P(y) ; P(z), P(u) \vdash \}$$

zu erzeugen:

$$\frac{[y/x] \frac{\vdash P(x), P(y)}{\vdash P(x)} (\mathcal{F}) \quad [u/z] \frac{P(z), P(u) \vdash}{P(z) \vdash} (\mathcal{F})}{[x/z] \frac{\vdash P(x) \quad P(z) \vdash}{\vdash} (\mathcal{R})}$$

Es ist auch möglich, ($\mathcal{R}$) und ($\mathcal{F}$) in einer einzigen Regel zusammenzufassen, indem man erlaubt, in einem Schritt mehrere Paare von Formeln zu unifizieren und zu resolvieren.

Definition 11.7 Die *verallgemeinerten Resolutionsregel* ($\mathcal{R} + \mathcal{F}$) lautet:

$$\frac{X_1 \vdash Y_1, A, B \quad C, D, X_2 \vdash Y_2}{(X_1, X_2 \vdash Y_1, Y_2)\sigma} (\mathcal{R} + \mathcal{F})$$

*verallgemeinerte
Resolutionsregel*

falls σ die Menge $\{A, B, C, D\}$ unifiziert.

Beispiel. Für die Klauselmeng aus dem vorherigen Beispiel genügt eine Anwendung der verallgemeinerten Resolutionsregel, um die leere Klausel abzuleiten:

$$[y/x, u/z][z/x] \frac{\vdash P(x), P(y) \quad P(z), P(u) \vdash}{\vdash} (\mathcal{R} + \mathcal{F})$$

11.2 Allgemeine Unifikatoren, Unifikationsalgorithmus

Um Vollständigkeit zu erreichen, müssen *allgemeinste Unifikatoren* verwendet werden.

Definition 11.8 Es ist σ ein *allgemeinster Unifikator* (*most general unifier*, kurz: mgu) von A und B , wenn für alle Unifikatoren τ von A und B gilt: $\tau = \sigma\rho$ für eine Substitution ρ . *allgemeinster Unifikator*

Beispiele. (i) Für $P(f(x, g(k, y)))$ und $P(f(l, z))$ ist $[x/l, z/g(k, y)]$ ein allgemeinster Unifikator.

Für die angegebenen Atome ist $[x/l, y/k', z/g(k, k')]$ zwar ein Unifikator, aber *kein* allgemeinster Unifikator, da $[x/l, y/k', z/g(k, k')] = [x/l, z/g(k, y)][y/k']$.

(ii) Für die Menge $\{P(x), P(y)\}$ ist sowohl $[x/y]$ als auch $[y/x]$ ein mgu.

Allgemeinste Unifikatoren sind also nicht eindeutig bestimmt.

(iii) Es ist $\sigma = [x/z, y/z]$ *kein* mgu für $\{P(x), P(y)\}$, da es keine Substitution ρ gibt, so dass $[x/y] = \sigma\rho$.

Insbesondere ist auch $[z/y]$ keine derartige Substitution ρ , da $\sigma[z/y] = [x/y, z/y] \neq [x/y]$.

Theorem 11.9 Mehrere *allgemeinste Unifikatoren* von A und B unterscheiden sich nur durch Umbenennung.

Beweis. Siehe Lassez, Maher & Marriot (1987, S. 605).

QED

Definition 11.10 Sei Γ eine endliche Menge von Termen oder Atomen. Dann ist die *Unterschiedsmenge* U von Γ wie folgt definiert:

Unterschiedsmenge

Finde die linkeste Stelle, an der nicht alle Ausdrücke in Γ dasselbe Symbol haben, und wähle von jedem Ausdruck in Γ jenen Teilausdruck, der an dieser Stelle beginnt.

Die Menge aller dieser Teilausdrücke ist die Unterschiedsmenge.

Beispiele. (i) Sei $\Gamma = \{P(k, \underline{f(x, y)}, g(z)), P(k, \underline{z}, h(u)), P(k, \underline{x}, y)\}$.

(Die relevanten Teilausdrücke sind unterstrichen.)

Dann ist die Unterschiedsmenge $U = \{f(x, y), z, x\}$.

(ii) Die Unterschiedsmenge für $\Gamma = \{Q(x), R(x, y)\}$ ist $U = \{Q(x), R(x, y)\}$.

Wir können uns auf die Behandlung zweielementiger Mengen $\Gamma = \{A, B\}$ für Atome A, B beschränken. Für diese erzeugt der folgende *Unifikationsalgorithmus* einen mgu, sofern A und B unifizierbar sind.

Unifikationsalgorithmus

Unifikationsalgorithmus

Eingabe: Eine Menge $\Gamma = \{A, B\}$ zweier Atome A und B .

Ausgabe: Ein mgu für Γ , falls A und B unifizierbar sind.

- (1) Setze $n = 0$ und $\sigma_0 = \varepsilon$, wobei ε die leere Substitution ist.
- (2) Wenn $\Gamma\sigma_n$ einelementig ist, dann halte; σ_n ist ein mgu für Γ .
Andernfalls bilde die Unterschiedsmenge U_n von $\Gamma\sigma_n$.
- (3) Wenn es eine Variable x und einen Term t in U_n gibt derart, dass die Variable x nicht in t vorkommt, dann setze $\sigma_{n+1} = \sigma_n[x/t]$ (d. h., bilde die Komposition $\sigma_n[x/t]$ mit der bisher erhaltenen Substitution σ_n), erhöhe n um 1, und gehe zu (2).
Andernfalls halte mit der Ausgabe, dass Γ nicht unifizierbar ist.

Bei Anwendung des Unifikationsalgorithmus im Rahmen der Resolution gilt für die Eingabe $\Gamma = \{A, B\}$, dass $FV(A) \cap FV(B) = \emptyset$; denn im Resolutionsschritt wird für die Mengen der freien Variablen der beiden Prämissen Disjunktheit gefordert. Im Fall von Eingabemengen $\Gamma = \{A, B\}$ mit $FV(A) \cap FV(B) \neq \emptyset$ können die in den Atomen A und B vorkommenden Variablen erst separiert werden. Diese Variablenseparierung ist (wie bei Klauseln) dadurch gerechtfertigt, dass A und B stets als voneinander unabhängig aufgefasst werden können.

Bemerkung. Die Überprüfung in Schritt (3), ob x in t vorkommt, bezeichnet man als *occur check*. Ohne diesen würde der Algorithmus nicht in jedem Fall terminieren. Zum Beispiel würde der Algorithmus *ohne* occur check für die nicht unifizierbare Menge $\Gamma = \{P(\underline{x}, x), P(\underline{y}, f(x))\}$ Folgendes liefern:

occur check

- (1) $\sigma_0 = \varepsilon$
- (2) $U_0 = \{x, y\}$
- (3) $\sigma_1 = [x/y]$, $\Gamma\sigma_1 = \{P(y, \underline{y}), P(y, \underline{f(y)})\}$
- (2) $U_1 = \{y, f(y)\}$
- (3) Die Variable y kommt im Term $f(y)$ vor! Ohne occur check erhält man:
 $\sigma_2 = [x/y][y/f(y)] = [x/f(y), y/f(y)]$,
 $\Gamma\sigma_2 = \{P(f(y), \underline{f(y)}), P(f(y), \underline{f(f(y))})\}$
- (2) $U_2 = \{y, f(y)\}$
- (3) Die Variable y kommt im Term $f(y)$ vor! Ohne occur check erhält man:
 $\sigma_3 = [x/f(y), y/f(y)][y/f(y)] = [x/f(f(y)), y/f(f(y))]$,
 $\Gamma\sigma_3 = \{P(f(f(y)), \underline{f(f(y))}), P(f(f(y)), \underline{f(f(f(y))}))\}$
 $\vdots$

Der Algorithmus würde also nicht halten, sondern in Schritt (3) stets $\sigma_{n+1} = \sigma_n[y/f(y)]$ setzen.

Beispiele. (i) $\Gamma = \{P(\underline{f(x)}, g(y)), P(\underline{z}, z)\}$

- (1) $\sigma_0 = \varepsilon$
- (2) $U_0 = \{f(x), z\}$
- (3) $\sigma_1 = [z/f(x)]$, $\Gamma\sigma_1 = \{P(f(x), \underline{g(y)}), P(f(x), \underline{f(x)})\}$
- (2) $U_1 = \{g(y), f(x)\}$
- (3) Die Unterschiedsmenge U_1 enthält keine Variable als Element, also sind $P(f(x), g(y))$ und $P(z, z)$ nicht unifizierbar.

- (ii) $\Gamma = \{P(\underline{f(x)}, y), P(\underline{z}, k)\}$
- (1) $\sigma_0 = \varepsilon$
 - (2) $U_0 = \{f(x), z\}$
 - (3) $\sigma_1 = [z/f(x)], \Gamma\sigma_1 = \{P(f(x), \underline{y}), P(f(x), \underline{k})\}$
 - (2) $U_1 = \{y, k\}$
 - (3) $\sigma_2 = [z/f(x)][y/k], \Gamma\sigma_2 = \{P(f(x), k), P(f(x), k)\} = \{P(f(x), k)\}$
 - (2) $\Gamma\sigma_2$ einelementig, also Γ unifizierbar mit mgu $\sigma_2 = [z/f(x)][y/k]$.
- (iii) $\Gamma = \{P(\underline{x}, x), P(\underline{y}, f(y))\}$
- (1) $\sigma_0 = \varepsilon$
 - (2) $U_0 = \{x, y\}$
 - (3) $\sigma_1 = [x/y], \Gamma\sigma_1 = \{P(y, y), P(y, f(y))\}$
 - (2) $U_1 = \{y, f(y)\}$
 - (3) Es kommt y in $f(y)$ vor (occur check), also sind $P(x, x)$ und $P(y, f(y))$ nicht unifizierbar.

11.3 SLD-Resolution

Abschließend behandeln wir eine eingeschränkte Form der quantorenlogischen Resolution, die sogenannte *SLD-Resolution*. Diese stellt die theoretische Grundlage der *Logikprogrammierung* und der Programmiersprache *Prolog* dar.

Bei SLD-Resolution werden nicht mehr beliebige Klauseln zugelassen, sondern nur noch sogenannte *Hornklauseln* (benannt nach Alfred Horn, 1918–2001), die neben beliebig vielen negativen Literalen höchstens *ein* positives Literal enthalten dürfen. Zudem dürfen im Resolutionsschritt als Prämissen nur eine Hornklausel mit keinem positiven Literal und eine Hornklausel mit genau einem positiven Literal vorkommen. Letztere muss als Variante aus einer gegebenen Klauselmengende, dem *Logikprogramm*, gewählt werden. Diese Beschränkungen führen zu einfacheren Resolutionsableitungen. Das auf Hornklauseln eingeschränkte Fragment der Quantorenlogik ist jedoch wie diese unentscheidbar. Bei SLD-Resolution rücken zudem die in einer Ableitung berechneten Unifikatoren in den Vordergrund; sie stellen Antworten auf Anfragen dar, die an das jeweilige Logikprogramm gerichtet werden. Logikprogrammierung im Sinne von SLD-Resolution ist Turing-vollständig (vgl. [Lloyd, 1993](#), Theorem 9.6), d. h. jede durch eine Turing-Maschine berechenbare Funktion ist auch durch ein Logikprogramm berechenbar.

Definition 11.12 (i) Eine *Hornklausel* ist eine Klausel mit höchstens einem positiven Literal, d. h., eine Klausel der Form $X \vdash$ oder $X \vdash A$ (wobei A ein positives Literal ist, und die Menge X von Atomen auch leer sein darf).

Hornklausel

Eine Klausel mit genau einem positiven Literal heißt *definite Hornklausel*.

definite Hornklausel

Man schreibt mit der *reversen Implikation* ‘ $\leftarrow$ ’ auch $\leftarrow X$ bzw. $A \leftarrow X$.

reverse Implikation

Motivation für die Schreibweise: $X \vdash A$ kann man als „ A , falls X “ lesen, $X \vdash$ als „ X ist falsch“ (bzw. als „Falsum, falls X “).

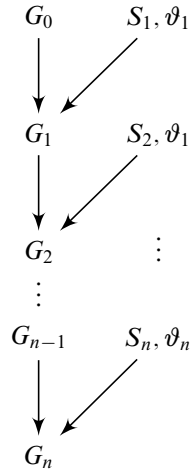
(ii) Eine Klausel der Form $\leftarrow X$ heißt auch *Zielklausel* (kurz: *Ziel*, oder auch *Anfrage*).

Zielklausel

Das leere Ziel wird mit $\leftarrow$ bezeichnet.

- (iii) Eine Klausel der Form $A \leftarrow X$ heißt auch *Regel*; A heißt *Kopf* und X heißt *Rumpf* von $A \leftarrow X$. Eine Klausel der Form $A \leftarrow$ heißt auch *Faktum*. *Regel*
- (iv) *Programmklauseln* sind Regeln oder Fakten (d. h. definite Hornklauseln). *Programmklausel*
- (v) Ein *Logikprogramm* (kurz: *Programm*) Π ist eine endliche Menge von Programmklauseln. *Logikprogramm*
- Definition 11.13** (i) Eine *SLD-Ableitung* eines Ziels G aus einem Programm Π besteht aus *SLD-Ableitung*
- einer Folge $G_0, G_1, \dots, G_n$ von Zielen,
 - einer Folge $S_1, \dots, S_n$ von Varianten von Klauseln aus Π ,
 - einer Folge $\vartheta_1, \dots, \vartheta_n$ von Substitutionen,
- so dass $G_0 = G$ und für alle $i < n$ Folgendes gilt (wobei die Listen von Atomen X_i , Y_i und Z_{i+1} auch leer sein können):
- G_i ist von der Form $\leftarrow X_i, A_i, Y_i$; dies schließt den Fall $\leftarrow$ ein.
 - S_{i+1} ist von der Form $B_{i+1} \leftarrow Z_{i+1}$.
 - S_{i+1} hat keine Variablen gemeinsam mit G_i oder $G_0\vartheta_1 \dots \vartheta_i$.
 - ϑ_{i+1} ist ein mgu von A_i und B_{i+1} .
 - $G_{i+1} = (\leftarrow X_i, Z_{i+1}, Y_i)\vartheta_{i+1}$.
- (ii) Die Zahl $n + 1$ heißt *Länge* der SLD-Ableitung. *Länge*
- (iii) Der Übergang von einem Ziel G_i zu einem neuen Ziel G_{i+1} ist ein *SLD-Resolutionsschritt*. *SLD-Resolutionsschritt*
- (iv) Die Klausel S_i heißt *i-te Inputklausel*. *Inputklausel*
- (v) Das Atom A_i heißt *ausgewähltes Atom*. *ausgewähltes Atom*

Eine SLD-Ableitung der Länge $n + 1$ hat folgende Form:



SLD-Ableitungen können aber auch unendlich sein.

Bemerkung. In einem Resolutionsschritt von G_i nach G_{i+1} geschieht Folgendes:

- (1) Im Ziel G_i wird ein Atom A_i ausgewählt.

- (2) Aus dem Programm Π wird eine Klausel gewählt.
- (3) Durch Separierung der Variablen (falls nötig) erhält man aus der gewählten Programmklausel die Variante S_{i+1} ; dies ist die neue Inputklausel $B_{i+1} \leftarrow Z_{i+1}$.
- (4) Falls es einen allgemeinsten Unifikator ϑ_{i+1} für den Kopf B_{i+1} der Klausel S_{i+1} und das ausgewählte Atom A_i gibt, dann wird A_i durch den Rumpf Z_{i+1} von S_{i+1} ersetzt, und ϑ_{i+1} wird auf die neue Zielklausel G_{i+1} angewendet.

Definition 11.14 Sei ein Logikprogramm Π und eine Zielklausel $\leftarrow X$ gegeben. Eine *SLD-Widerlegung* von $\leftarrow X$ relativ zu Π ist eine SLD-Ableitung, die mit $\leftarrow X$ als linker oberster Annahme beginnt, ansonsten nur (Varianten von) Klauseln aus Π als Annahmen benutzt, und mit der leeren Klausel $\leftarrow$ endet.

SLD-Widerlegung

Definition 11.15 (i) Eine SLD-Ableitung heißt *erfolgreich*, falls es sich um eine SLD-Widerlegung handelt, d. h., wenn G_n die leere Klausel $\leftarrow$ ist.

erfolgreich

(ii) Die auf $FV(G_0)$ eingeschränkte Komposition $\vartheta_1 \dots \vartheta_n$ der allgemeinsten Unifikatoren in einer erfolgreichen SLD-Ableitung heißt *berechnete Antwortsubstitution* für das Ziel G_0 , und $G_0\vartheta_1 \dots \vartheta_n$ heißt *berechnete Instanz* von G_0 .

*berechnete
Antwortsubstitution*

Ist V eine Menge von Variablen, so bezeichnet $\vartheta \upharpoonright V$ die *Einschränkung* von ϑ auf die Variablen in V .

*Einschränkung auf
Variablen*

(iii) Eine SLD-Ableitung heißt *gescheitert*, falls G_n nicht die leere Klausel ist, und keine (Variante einer) Klausel aus Π für das in G_n ausgewählte Atom anwendbar ist.

gescheitert

Definition 11.16 Ein *SLD-Beweis* für X aus Π ist eine SLD-Widerlegung von $\leftarrow X$ relativ zu Π .

SLD-Beweis

Bemerkungen. (i) Notiert man die *SLD-Resolutionsschritte* in der Form

SLD-Resolutionsschritt

$$\frac{\leftarrow X, A \quad B \leftarrow Y}{(\leftarrow X, Y)\sigma} \sigma$$

wobei σ ein allgemeinsten Unifikator von A und B ist, und die Mengen der freien Variablen der beiden Prämissen disjunkt sein müssen, so sieht ein *SLD-Beweis* für X_1 aus Π wie folgt aus:

$$\frac{\frac{\leftarrow X_1 \quad A_1 \leftarrow Y_1}{\leftarrow X_2} \sigma_1 \quad A_2 \leftarrow Y_2}{\leftarrow X_3} \sigma_2 \dots \leftarrow$$

Die Substitutionen $\sigma_1, \sigma_2, \dots$ sind hierbei die im jeweiligen Schritt durch den Unifikationsalgorithmus berechneten allgemeinsten Unifikatoren.

(ii) Die Abkürzung SLD steht für Folgendes:

S: Es wird eine *selection function* für die Auswahl des zu unifizierenden Atoms in der Zielklausel verwendet.

L: Die Form der Ableitung ist *linear*.

Dies wird besonders deutlich, wenn man die SLD-Ableitung wie folgt notiert:

$$G_0 \xrightarrow{S_1, \vartheta_1} G_1 \xrightarrow{S_2, \vartheta_2} G_2 \cdots G_{n-1} \xrightarrow{S_n, \vartheta_n} G_n \cdots$$

D: Logikprogramme bestehen aus *definiten* Hornklauseln.

Wir lassen in unseren formalen Sprachen $\mathcal{L}$ nun auch ‘sprechende’ Relationssymbole wie *add*, Funktionszeichen wie *s* (für *successor*) und Konstanten wie 0 zu, um deren intendierte Interpretation anzudeuten. Wir erweitern hierdurch lediglich das Alphabet; die neuen Zeichen werden aber mit keinerlei bestimmter Bedeutung versehen. Die Bedeutung dieser Zeichen ist operativ bezüglich SLD-Resolution gegeben.

Beispiel. Wir betrachten das Logikprogramm Π_{add} für die Addition, welches aus zwei Programmklauseln, nämlich einem Faktum S_1 und einer Regel S_2 , besteht:

$$\begin{aligned} S_1. \quad & \text{add}(x, 0, x) \leftarrow \\ S_2. \quad & \text{add}(x, s(y), s(z)) \leftarrow \text{add}(x, y, z) \end{aligned}$$

Eine SLD-Widerlegung der Zielklausel

$$\leftarrow \text{add}(s(s(0)), s(s(0)), z)$$

relativ zum Programm Π_{add} ist dann beispielsweise:

$$\begin{array}{ccc} \leftarrow \text{add}(s(s(0)), s(s(0)), z) & \text{add}(x_1, s(y_1), s(z_1)) \leftarrow \text{add}(x_1, y_1, z_1) & \\ \downarrow & \swarrow [x_1/s(s(0)), y_1/s(0), z/s(z_1)] = \vartheta_1 & \\ \leftarrow \text{add}(s(s(0)), s(0), z_1) & \text{add}(x_2, s(y_2), s(z_2)) \leftarrow \text{add}(x_2, y_2, z_2) & \\ \downarrow & \swarrow [x_2/s(s(0)), y_2/0, z_1/s(z_2)] = \vartheta_2 & \\ \leftarrow \text{add}(s(s(0)), 0, z_2) & \text{add}(x_3, 0, x_3) \leftarrow & \\ \downarrow & \swarrow [x_3/s(s(0)), z_2/s(s(0))] = \vartheta_3 & \\ \leftarrow & & \end{array}$$

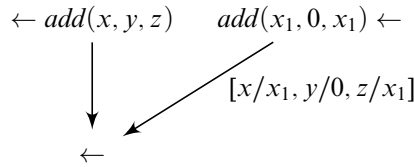
Die Komposition $\vartheta_1\vartheta_2\vartheta_3$ der allgemeinsten Unifikatoren schränken wir nun auf die im Ziel $\leftarrow \text{add}(s(s(0)), s(s(0)), z)$ vorkommenden Variablen (d. h. hier auf z) ein, um die berechnete Antwortsubstitution zu erhalten:

$$\begin{aligned} & \vartheta_1\vartheta_2\vartheta_3 \mid \text{FV}(\leftarrow \text{add}(s(s(0)), s(s(0)), z)) \\ &= \vartheta_1\vartheta_2\vartheta_3 \mid \{z\} \\ &= [x_1/s(s(0)), y_1/s(0), z/s(z_1)][x_2/s(s(0)), y_2/0, z_1/s(z_2)]\vartheta_3 \mid \{z\} \\ &= [x_1/s(s(0)), x_2/s(s(0)), y_1/s(0), y_2/0, z/s(z_2)][x_3/s(s(0)), z_2/s(s(0))] \mid \{z\} \\ &= [x_1/s(s(0)), x_2/s(s(0)), x_3/s(s(0)), y_1/s(0), y_2/0, z/s(s(s(0))))] \mid \{z\} \\ &= [z/s(s(s(0)))] \end{aligned}$$

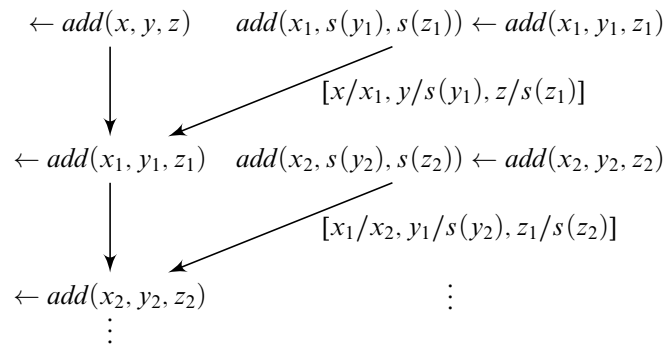
Die berechnete Antwortsubstitution ist also $[z/s(s(s(0)))]$; die berechnete Instanz der

Zielklausel ist $\leftarrow \text{add}(s(s(0)), s(s(0)), s(s(s(s(0)))))$. (Für die intendierte Interpretation ist also $2 + 2 = 4$.)

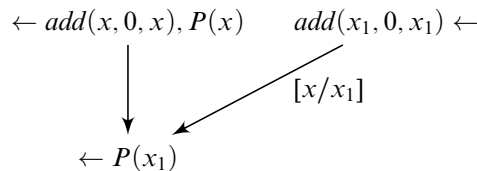
Für die Anfrage $\leftarrow \text{add}(x, y, z)$ gibt es relativ zu Π_{add} eine erfolgreiche Ableitung



aber auch unendliche SLD-Ableitungen wie z. B.



Für die Zielklausel $\leftarrow \text{add}(x, y, z), P(x)$ gibt es relativ zu Π_{add} z. B. die folgende gescheiterte Ableitung:



SLD-Ableitungen können also erfolgreich, unendlich oder gescheitert sein.

SLD-Ableitungen können auch als ‘Rückwärts-Schließen’ bzw. *zielgerichtetes Rasonieren* interpretiert werden:

- Gegeben als Ziel: X, A (notiert als $\leftarrow X, A$).
- Wir wissen: $B\sigma$ falls $Y\sigma$ für alle σ (notiert als $B \leftarrow Y$).
- Ferner wissen wir: $A\sigma = B\sigma$.
- Als neues Ziel erhält man: $X\sigma, Y\sigma$ (notiert als $(\leftarrow X, Y)\sigma$).

Wenn wir also X, A für die Substitution σ zeigen wollen, dann genügt es, $X\sigma, Y\sigma$ zu zeigen. Dies schließt die Abfrage von Daten ein. Das Ziel enthält freie Variablen für die zu suchenden Daten; das Programm enthält die Daten in Form von Fakten. Einen prinzipiellen Unterschied zwischen Daten und Programmen gibt es nicht.

Beispiel. Gegeben das Faktum $P(k, l) \leftarrow$, für welche x, y gilt $P(x, y)$?

SLD-Resolution:
$$\frac{\leftarrow P(x, y) \quad P(k, l) \leftarrow}{\leftarrow} [x/k, y/l]$$

Antwort: Für $x = k$ und $y = l$.

In einer SLD-Ableitung spielen verschiedene Arten von Nichtdeterminismus eine Rolle, da in jedem Resolutionsschritt vier Entscheidungen zu treffen sind:

- (A) Die Auswahl eines Atoms in der aktuellen Zielklausel.
- (B) Die Wahl einer Programmklausele für das ausgewählte Atom.
- (C) Die Umbenennung der in der gewählten Programmklausele vorkommenden Variablen; also die Auswahl einer Variante.
- (D) Die Wahl des allgemeinsten Unifikators.

Der Nichtdeterminismus durch (C) und (D) ist unproblematisch. Die Wahl des allgemeinsten Unifikators (D) ist durch die Verwendung des Unifikationsalgorithmus festgelegt. Des Weiteren spielt die konkrete Variante eines allgemeinsten Unifikators für die für ein Ziel G berechnete Antwortsubstitution keine Rolle, da aus dem vom Unifikationsalgorithmus gelieferten allgemeinsten Unifikator alle anderen allgemeinsten Unifikatoren durch Umbenennung erzeugt werden können (vgl. Theorem 11.9). Dasselbe gilt auch für die Umbenennung der in der gewählten Programmklausele vorkommenden Variablen (C), da verschiedene Umbenennungen lediglich zu Varianten der berechneten Antwortsubstitutionen führen. Durch bloße Umbenennung können also keine Antworten verlorengehen.

Hingegen ist es für die SLD-Resolution entscheidend, wie der durch (A) und (B) gegebene Nichtdeterminismus gehandhabt wird. Die Auswahl eines Atoms in der aktuellen Zielklausel (A) wird durch *Auswahlfunktionen* $\mathcal{A}$ festgelegt. Man kann zeigen, dass die *Existenz* einer erfolgreichen SLD-Ableitung unabhängig von der gegebenen Auswahlfunktion ist (d. h., der durch (A) gegebene Nichtdeterminismus in diesem Sinne unproblematisch ist). Die Wahl der Programmklausele (B) ist hingegen entscheidend für das *Auffinden* einer solchen SLD-Ableitung. Durch (B) wird ein *Suchraum* für SLD-Ableitungen gemäß $\mathcal{A}$, der sogenannte *SLD-Baum*, aufgespannt.

Definition 11.17 Ein *SLD-Baum* für ein Ziel G_0 relativ zu einem Programm Π gemäß einer Auswahlfunktion $\mathcal{A}$ ist ein wie folgt gegebener (nach unten verzweigender) Baum:

SLD-Baum

- (i) Jeder Knoten des Baums ist ein (möglicherweise leeres) Ziel.
- (ii) Der Wurzelknoten ist G_0 .
- (iii) Jeder Knoten G_i mit ausgewähltem Atom A_i hat genau einen Nachfolger für jede auf A_i anwendbare Klausel S aus Π .
Der Nachfolger ist eine Resolvente von G_i und S (bzw. von G_i und der gewählten Variante von S) bezüglich A_i .
- (iv) Knoten, die ein leeres Ziel sind, haben keine Nachfolger.

Zweige in SLD-Bäumen sind SLD-Ableitungen für G_0 relativ zu Π . Wir notieren diese in der Form

$$G_0 \xrightarrow{S_1, \vartheta_1} G_1 \xrightarrow{S_2, \vartheta_2} G_2 \cdots G_{n-1} \xrightarrow{S_n, \vartheta_n} G_n \cdots$$

von oben (d. h., vom Wurzelknoten G_0 aus) nach unten, wie in den folgenden Beispielen gezeigt.

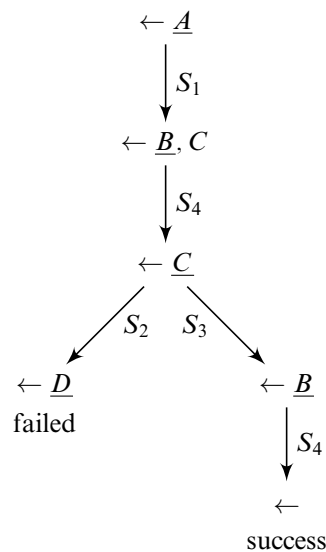
- Definition 11.18** (i) Ein SLD-Baum heißt *erfolgreich*, falls er die leere Klausel enthält. *erfolgreich*
 (Zweige, die mit der leeren Klausel enden, markieren wir mit ‘success’.)
- (ii) Ein SLD-Baum heißt *endlich gescheitert*, falls er endlich und nicht erfolgreich ist. *endlich gescheitert*
 Dies ist der Fall, wenn alle Zweige gescheiterte SLD-Ableitungen sind.
 (Zweige gescheiterter SLD-Ableitungen markieren wir mit ‘failed’.)

Beispiel. Wir betrachten zunächst das Logikprogramm

- $S_1.$ $A \leftarrow B, C$
 $S_2.$ $C \leftarrow D$
 $S_3.$ $C \leftarrow B$
 $S_4.$ $B \leftarrow$

(in dem nur 0-stellige Relationszeichen, bzw. Aussagesymbole vorkommen).

Sei $\mathcal{A}$ jene Auswahlfunktion, die stets das linkeste Atom (im Folgenden unterstrichen) im Rumpf einer Zielklausel auswählt. Für die Anfrage $\leftarrow A$ relativ zu unserem Programm ist der folgende Baum ein SLD-Baum gemäß $\mathcal{A}$:



Der linke Zweig ist eine gescheiterte SLD-Ableitung, da unser Logikprogramm keine auf die Zielklausel $\leftarrow \underline{D}$ anwendbare Klausel enthält. Der rechte Zweig endet mit der leeren Klausel, ist also eine erfolgreiche SLD-Ableitung, und der SLD-Baum ist somit erfolgreich.

Eine *Tiefensuche* könnte hier zuerst in den linken, gescheiterten Zweig geraten. Durch *backtracking*, d. h., durch das Zurückgehen zu einer früheren Verzweigung und Wahl einer anderen Alternative, kann aber der erfolgreiche Zweig dennoch aufgefunden werden.

backtracking

Beispiel. Nun betrachten wir das Logikprogramm

$$S_1. \quad P(x, z) \leftarrow Q(x, y), P(y, z)$$

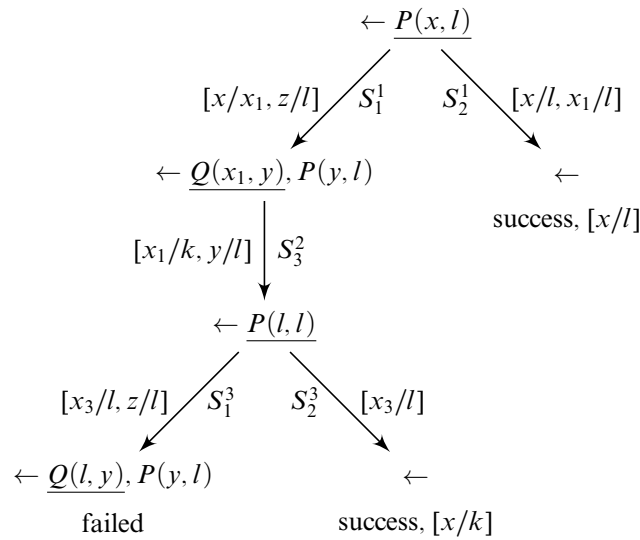
$$S_2. \quad P(x, x) \leftarrow$$

$$S_3. \quad Q(k, l) \leftarrow$$

und die Zielklausel

$$G. \quad \leftarrow P(x, l)$$

Wird durch die Auswahlfunktion stets das *linkeste Atom* gewählt, so erhält man einen SLD-Baum der folgenden Form, wobei die im jeweils i -ten Resolutionsschritt verwendeten Inputklauseln jene Varianten S_1^i , S_2^i und S_3^i der entsprechenden Programmklauseln S_1 , S_2 und S_3 sind, bei denen die zur Variablenseparierung umzubenennenden Variablen den Index i erhalten:



Der SLD-Baum ist endlich und erfolgreich. Die berechneten Antwortsubstitutionen sind

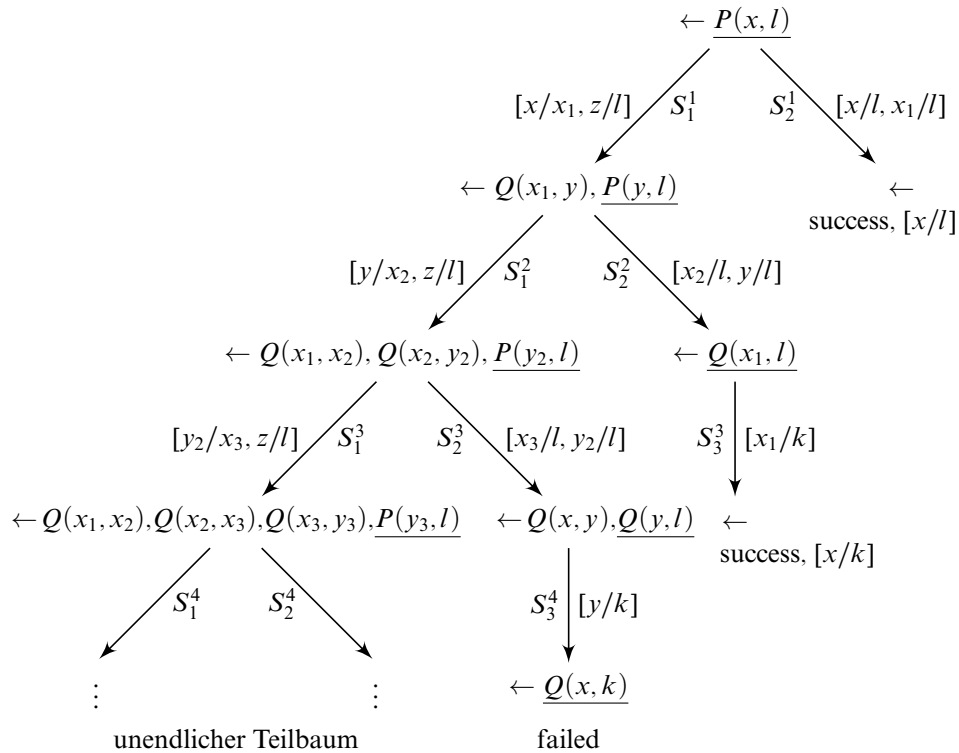
$$\begin{aligned} [x/x_1, z/l][x_1/k, y/l][x_3/l] \mid \{x\} &= [x/k, x_1/k, y/l, z/l][x_3/l] \mid \{x\} \\ &= [x/k, x_1/k, y/l, z/l, x_3/l] \mid \{x\} \\ &= [x/k] \quad (\text{mittlerer Zweig}) \end{aligned}$$

und

$$[x/l, x_1/l] \mid \{x\} = [x/l] \quad (\text{rechter Zweig}).$$

Andere SLD-Bäume unterscheiden sich von dem gezeigten SLD-Baum nur bezüglich Umbenennungen, sind aber der Form nach gleich.

Wird durch die Auswahlfunktion stets das *rechteste Atom* gewählt, so erhält man einen SLD-Baum der folgenden Form, wobei für die im i -ten Resolutionsschritt verwendeten Varianten S_1^i , S_2^i und S_3^i der entsprechenden Programmklauseln S_1 , S_2 und S_3 das oben Gesagte gilt.



Dieser SLD-Baum ist ebenfalls erfolgreich mit denselben berechneten Antwortsubstitutionen $[x/k]$ und $[x/l]$ wie beim ersten SLD-Baum. Im Unterschied zum ersten ist dieser zweite SLD-Baum jedoch unendlich; im linken Zweig können stets neue Zielklauseln erzeugt werden.

Startet eine Tiefensuche hier im linken Zweig, so können die erfolgreichen SLD-Ableitungen rechts nicht mehr aufgefunden werden, da der linke Zweig unendlich ist, und *backtracking* natürlich nur bei endlichen Zweigen möglich ist. Bei einer *Breitensuche* werden auch die erfolgreichen SLD-Ableitungen gefunden.

Literatur

- K. R. Apt, *From Logic Programming to Prolog*, Prentice Hall, 1997.
- P. Blackburn, J. Bos & K. Striegnitz, *Learn Prolog Now!*, College Publications, 2006, sowie online frei zugänglich unter <http://www.learnprolognow.org/>.
- J. W. Lloyd, *Foundations of Logic Programming*, 2nd edition, Springer-Verlag, 1993.

Literaturverzeichnis

- K. R. Apt, *From Logic Programming to Prolog*, Prentice Hall, 1997.
- M. Ben-Ari, *Mathematical Logic for Computer Science. Third Edition*, Springer, 2012.
- P. Blackburn, J. Bos & K. Striegnitz, *Learn Prolog Now!*, College Publications, 2006, sowie online frei zugänglich unter <http://www.learnprolognow.org/>.
- A. Church, *A Note on the Entscheidungsproblem*, The Journal of Symbolic Logic 1 (1936), 40–41.
- H.-D. Ebbinghaus, J. Flum & W. Thomas, *Einführung in die mathematische Logik*, 5. Auflage, Spektrum Akademischer Verlag, 2007.
- J. H. Gallier, *Logic for Computer Science. Foundations of Automatic Theorem Proving*, 2nd edition, Dover Publications, 2015.
- G. Gentzen, *Untersuchungen über das logische Schließen*, Mathematische Zeitschrift 39 (1935), 176–210, 405–431.
- H. J. Goltz & H. Herre, *Grundlagen der logischen Programmierung*, Wiley-VCH, 1990.
- K. Gödel, *Über die Vollständigkeit des Logikkalküls*, Dissertation, Universität Wien, 1929. (In: S. Feferman et al. (Hrsg.), *Kurt Gödel, Collected Works, Volume I, Publications 1929-1936*, Oxford University Press, 1986.)
- S. C. Kleene, *Introduction to Metamathematics*, Elsevier North-Holland, 2000.
- J.-L. Lassez, M. J. Maher & K. Marriot (1987), Unification Revisited. In: J. Minker (ed.), *Foundations of Deductive Databases and Logic Programming*, Morgan Kaufmann, S. 587–625.
- A. Leitsch, *The Resolution Calculus*, Springer, 1997.
- J. W. Lloyd, *Foundations of Logic Programming*, 2nd edition, Springer-Verlag, 1993.
- D. Makinson, *Sets, Logic and Maths for Computing. Third Edition*, Springer, 2020.
- J. Mittelstraß (Hrsg.), *Enzyklopädie Philosophie und Wissenschaftstheorie*, J. B. Metzler, 2005.
- S.-H. Nienhuys-Cheng & R. de Wolf, *Foundations of Inductive Logic Programming. Lecture Notes in Artificial Intelligence 1228*, Springer, 1997.
- D. Prawitz, *Natural Deduction. A Proof-Theoretical Study*, Almqvist & Wiksell, 1965. Neuauflage 2006, Dover Publications.
- U. Schöning, *Logik für Informatiker*, 5. Auflage, Spektrum Akademischer Verlag, 2000.
- P. Schroeder-Heister, *Einführung in die Logik*, Skript zur Vorlesung, 2008.
- R. M. Smullyan, *Logical Labyrinths*, A K Peters, 2009.
- R. M. Smullyan, *A Beginner's Guide to Mathematical Logic*, Dover Publications, 2014.

- A. Tarski, *Der Wahrheitsbegriff in den formalisierten Sprachen*, *Studia Philosophica* **1** (1935), 261–405.
- A. S. Troelstra & H. Schwichtenberg, *Basic Proof Theory*, Cambridge University Press, 2000.
- A. M. Turing, *On Computable Numbers, with an Application to the Entscheidungsproblem*, *Proceedings of the London Mathematical Society, Series 2*, **42** (1936), 230–265.
- D. van Dalen, *Logic and Structure*, 5th ed., Springer, 2013.