

Computerlinguistische Beschreibung
der japanischen Verbmorphologie
und ihrer semantischen Entsprechungen

Diplomarbeit
bei
FLP/KIT

Aufgabensteller:
Prof. Dr. Bernd Mahr
(TU Berlin)

Betreuer:
Dr. Melanie Siegel
(DFKI Saarbrücken)
Dr. Manfred Stede
(TU Berlin)

Dietrich Bollmann
Matrikelnummer: 154540
Stromstraße 40
10551 Berlin

15. Dezember 1999

Inhaltsverzeichnis

Vorwort	1
1 Formale Linguistik und Constraintsysteme	5
1.1 Sprachliche Zeichen als Merkmalstrukturen	7
1.1.1 Informelle Einführung	8
1.1.2 Aufbau und Nomenklatur von Merkmalstrukturen . . .	8
1.2 Typenhierarchien	13
1.3 Merkmalstrukturen	18
1.4 Constraints	33
1.5 Resolution	46
1.6 Beispiele zur Anwendung von Constraintsystemen	47
1.6.1 Listen und Listen-Konkatenation	47
1.6.2 Syntax und Semantik	52
1.7 Defaults und Defaultvererbung	58
2 Japanische Verbmorphologie	61
2.1 Verben und japanische Grammatik	61
2.1.1 Agglutination	62
2.1.2 Verben und Adjektive	65
2.2 Die Verbmorphologie der japanischen Schulgrammatik	70
2.2.1 Flexionsklassen	72
2.2.2 Verbstämme	73
2.2.3 Hilfsverben (助動詞, <i>zyodousi</i>)	74
2.2.4 Phonologische Transformationen	77
2.2.5 Ausnahmen	78
2.3 Die überarbeitete Verbmorphologie	78
2.3.1 Kana-Schreibweise vs. phonologische Umschrift	79
2.3.2 Eliminierung der phonologischen Transformationen . . .	79
2.3.3 Eliminierung der Ausnahmen zugunsten einer komplizierteren Systematik	80

2.3.4	Von der traditionellen Grammatik zum überarbeiteten System	82
2.3.4.1	Verbkategorien	83
2.3.4.2	Verbstämme	89
2.3.4.3	Ausnahmen	95
2.4	Morphologie und Grammatik	96
3	Syntax und Semantik als Rahmen der Morphologie	99
3.1	Einleitung	99
3.2	Syntax	107
3.2.1	Komplemente	112
3.2.2	Adjunkte	130
3.2.3	Nichtlokale Abhängigkeiten	140
3.2.4	Relativsätze und Attributsätze	144
3.3	Semantik	147
3.3.1	Lexikalische Einträge	149
3.3.2	Dominanzschemata	153
3.3.3	Relativsätze	156
4	Struktur der morphologischen Beschreibung	161
4.1	Einleitung	161
4.2	Agglutination	162
4.2.1	Das Agglutinationsschema	170
4.2.2	Das morphologische Head-Merkmal-Prinzip	171
4.2.3	Das morphologische PHON-Merkmal-Prinzip	172
4.2.4	Das MORPHOLOGY-Merkmal-Prinzip (MMP)	172
4.2.5	Das STEM-Merkmal-Default	173
4.3	Lexikalisch-morphologische Regeln	179
4.3.1	Adjunkteinführung	184
4.4	Klassifizierung morphologischer Formen	184

5 Ein Fragment der japanischen Verbmorphologie	187
5.1 Organisation des morphologischen Lexikons	188
5.1.1 Lexikalische Einträge	188
5.1.2 Lexikalische Defaults	194
5.1.2.1 Lexikalische Defaultvererbungshierarchie . . .	194
5.1.2.2 Morphologische Lexikoneinträge (<i>m-form</i> [◇])	195
5.1.2.3 Verbstämme (<i>v-stem</i> [◇])	195
5.1.2.4 Suffixe im weiteren Sinne (<i>g-suffix</i> [◇])	197
5.1.2.5 Clitics (<i>clitic</i> [◇])	198
5.1.2.6 Suffixe im engeren Sinne (<i>m-suffix</i> [◇])	200
5.1.2.7 Flexive (<i>flexive</i> [◇])	200
5.1.2.8 Suffixverben und -adjektive (<i>suffix</i> [◇])	202
5.2 Das morphologische Lexikon	205
5.2.1 Verbstämme	205
5.2.1.1 v_v : いる – <i>i-ru</i> – sein	205
5.2.1.2 v_v : 見る – <i>mi-ru</i> – sehen	206
5.2.1.3 v_v : 寝る – <i>ne-ru</i> – schlafen	206
5.2.1.4 v_v : 食べる – <i>tabe-ru</i> – essen	207
5.2.1.5 v_{c_k} : 書く – <i>ka.k-u</i> – schreiben	207
5.2.1.6 $v_{c_{iku}}$: 行く – <i>i.k-u</i> – gehen, fahren	208
5.2.1.7 v_{c_g} : 泳ぐ – <i>oyo.g-u</i> – schwimmen	208
5.2.1.8 v_{c_s} : 話す – <i>hana.s-u</i> – sprechen	209
5.2.1.9 v_{c_t} : 待つ – <i>ma.t-u</i> – warten	209
5.2.1.10 v_{c_n} : 死ぬ – <i>si.n-u</i> – sterben	210
5.2.1.11 v_{c_b} : 呼ぶ – <i>yo.b-u</i> – rufen	210
5.2.1.12 v_{c_m} : 飲む – <i>no.m-u</i> – trinken	211
5.2.1.13 v_{c_r} : 降る – <i>fu.r-u</i> – 〈Regen〉 fällt	211
5.2.1.14 $v_{c_{aru}}$: ある – <i>a.r-u</i> – sein	212
5.2.1.15 $v_{c_{da}}$: だ – <i>da</i> – sein	213
5.2.1.16 $v_{c_{nasaru}}$: くださる – <i>kudasa.r-u</i> – bitte . . .	215

5.2.1.17	v_{cw} : 買う – <i>ka.-u</i> – kaufen	216
5.2.1.18	v_{ikuru} : 来る – <i>ku-ru</i> – kommen	216
5.2.1.19	v_{isuru} : する – <i>s.-uru</i> – tun, machen . . .	218
5.2.1.20	v_{imasu} : -ます – <i>-mas.-u</i> – Honorativ . . .	221
5.2.1.21	<i>adj</i> : 高い – <i>taka-i</i> – hoch, teuer	221
5.2.2	Clitics	222
5.2.2.1	Konsonantischer Stamm (<i>cons-stem</i>)	223
5.2.2.2	i-Stamm (<i>i-stem</i>)	224
5.2.2.3	a-Stamm (<i>a-stem</i>)	224
5.2.2.4	Onbin-Stamm (<i>onbin-stem</i>)	224
5.2.3	Flexive	226
5.2.3.1	Präsens (-FLEX)	226
5.2.3.2	Perfekt (-PAST)	231
5.2.3.3	Futur / Volitional (-VOLI)	233
5.2.3.4	Imperativ (-IMP)	235
5.2.3.5	Partizip (-PART)	237
5.2.3.6	Negiertes Partizip (-NEGPART)	240
5.2.3.7	Konditional (-COND)	242
5.2.3.8	Perfekt-Konditional (-PCOND)	244
5.2.3.9	Adverbial (-ADV)	245
5.2.4	Suffixverben und -adjektive (<i>suffix[◇]</i>)	248
5.2.4.1	Honorativ (-HON)	248
5.2.4.2	Negation (-NEG)	250
5.2.4.3	Voluntativ (-VOLU)	252
5.2.4.4	Kausativ (-CAUS)	254
5.2.4.5	Passiv (-PASS)	259
5.2.4.6	Kausativ-Passiv (-CAUS-PASS)	264
5.2.4.7	Potential (-POT)	264
5.2.4.8	Hilfssuffix ‘-Karu’ (- ϕ)	267

6	Entwicklung und Implementierung des vorgestellten Modells	269
6.1	TFS als experimentelles System	269
6.2	Strategien für die Implementierung komplexer Grammatiken .	270
6.2.1	Vereinfachung der Anfrage	270
6.2.2	Vereinfachung der Constraintbeschreibung	271
6.2.3	Entwicklung von Subsystemen	272
6.3	Entwicklung der japanischen Grammatik	273
6.3.1	Verbmorphologie	273
6.3.2	Syntax	274
6.3.3	Semantik	275
6.3.4	Integration von Morphologie, Syntax und Semantik . .	275
6.4	Kritik des benutzten Verfahrens	276
7	Ergebnisse und Ausblick	281
A	Beispielsätze zum Fragment der japanischen Verbmorphologie	287
A.1	Flexive	290
A.1.1	Präsens	290
A.1.2	Perfekt	291
A.1.3	Futur / Volitional	294
A.1.4	Imperativ	297
A.1.5	Partizip	298
A.1.6	Negiertes Partizip	302
A.1.7	Konditional	304
A.1.8	Perfekt-Konditional	305
A.1.9	Adverbial	307
A.2	Suffixe	311
A.2.1	Honorativ	311
A.2.2	Negation	313
A.2.3	Voluntativ	315

A.2.4	Kausativ	318
A.2.5	Passiv	324
A.2.6	Potential	332
A.2.7	Hilfssuffix ‘Karu’	334
B	Verbtafeln	337
	Zur Ableitung der Verbtafeln	338
B.1	v_v : いる – <i>i-ru</i> – sein, <Kontinuative>	339
B.2	v_v : 見る – <i>mi-ru</i> – sehen	340
B.3	v_v : 寝る – <i>ne-ru</i> – schlafen	341
B.4	v_v : 食べる – <i>tabe-ru</i> – essen	342
B.5	v_{c_k} : 書く – <i>ka.k-u</i> – schreiben	343
B.6	$v_{c_{iku}}$: 行く – <i>i.k-u</i> – gehen, fahren	344
B.7	v_{c_g} : 泳ぐ – <i>oyo.g-u</i> – schwimmen	345
B.8	v_{c_s} : 話す – <i>hana.s-u</i> – sprechen	346
B.9	v_{c_t} : 待つ – <i>ma.t-u</i> – warten	347
B.10	v_{c_n} : 死ぬ – <i>si.n-u</i> – sterben	348
B.11	v_{c_b} : 呼ぶ – <i>yo.b-u</i> – rufen	349
B.12	v_{c_m} : 飲む – <i>no.m-u</i> – trinken	350
B.13	v_{c_r} : 乗る – <i>no.r-u</i> – fahren, einsteigen	351
B.14	v_{c_r} : 取る – <i>to.r-u</i> – nehmen	352
B.15	$v_{c_{aru}}$: ある – <i>a.r-u</i> – sein	353
B.16	$v_{c_{da}}$: だ – <i>da</i> – sein (<Partikelverb>)	354
B.17	$v_{c_{nasaru}}$: なさる – <i>nasa.r-u</i> – tun, machen	355
B.18	v_{c_w} : 買う – <i>ka.-u</i> – kaufen	356
B.19	$v_{i_{kuru}}$: 来る – <i>ku-ru</i> – kommen	357
B.20	$v_{i_{suru}}$: する – <i>s.-uru</i> – tun, machen	358
	Bibliographie	359

Vorwort

Morphologische Fragestellungen spielen in vielen computerlinguistischen Systemen eine nur untergeordnete Rolle. Der Grund dafür ist jedoch nicht, daß die Morphologie für sprachverarbeitende Systeme keine Wichtigkeit hätte; vielmehr resultiert diese Unterbewertung aus der Komplexität des zu beschreibenden Gegenstandes und der Struktur der in der modernen Linguistik häufig betrachteten indoeuropäischen Sprachen. Der Modellierung von Sprache muß eine Eingrenzung der betrachteten Phänomene vorausgehen. Bei indoeuropäischen Sprachen bietet sich die Ausgrenzung der Morphologie an. Die begrenzte Anzahl ihrer morphologischen Formen legt eine Behandlung durch das Lexikon nahe; an die Stelle systematischer Beschreibungen treten Aufzählungen lexikalischer Einträge.

Die Betrachtung einer anderen Sprachfamilie kann eine Neubestimmung der zu beschreibenden sprachlichen Bereiche notwendig machen. Für das Japanische hat dieses eine höhere Gewichtung der Morphologie zur Folge: Schon die Modellierung eines kleinen Ausschnittes dieser Sprache macht die Einbeziehung morphologischer Regularitäten notwendig. Insbesondere gilt dieses für die Morphologie der japanischen Verben und Adjektive. Damit ist die Motivation für die folgenden Betrachtungen umrissen: Sie widmen sich der Frage, wie sich die Verb- und Adjektivmorphologie im Rahmen einer computerlinguistischen Beschreibung des Japanischen adäquat darstellen läßt.

Das Japanische wird hinsichtlich der Morphologie seiner Verben und Adjektive zu den agglutinierenden Sprachen gezählt. Die Verschleifung der Morpheme erinnert jedoch gleichzeitig an Mechanismen, die charakteristisch für den flektierenden Sprachbau sind. Die Spezifikation adäquater und leicht zu handhabender formaler Mittel zur Beschreibung dieser Charakteristika der japanischer Grammatik bildet zusammen mit der Modellierung eines relevanten Ausschnittes japanischer Verbformen den Schwerpunkt dieser Arbeit.

Die Modifikation der Form eines Verbes hat im Japanischen oft auch eine Veränderung der Valenz zur Folge. Die Modellierung dieses Phänomens erfolgt durch die Integration des Subkategorisierungsrahmens in die lexikalische Beschreibung der Verbstämme und die Spezifikation von valenzmodifizierenden Funktionen im Lexikoneintrag der Verbsuffixe. Eine minimale Syntax, die auf dem so erzeugten Subkategorisierungsrahmen aufsetzt, veranschaulicht die Wechselwirkung von Morphologie und Syntax.

Die morphologische Modifikation von Verben und die Veränderungen, die sich dadurch in ihrer Syntax ergeben, spiegeln sich auch in ihrer semantischen Beschreibung. Diese muß so angelegt sein, daß die morphologisch verursachten Veränderungen der Verbbedeutung und der semantischen Argumentstruktur dargestellt werden können. Diese Aspekte der semantischen Beschreibung

lassen sich am besten anhand eines Formalismus veranschaulichen, der von allen Problemen abstrahiert, die durch die Verbmorphologie nicht direkt betroffen sind. Eine an der sprachlichen Oberfläche orientierte Semantik erfüllt diese Bedingung. Die Darstellung der Wechselwirkungen zwischen Morphologie, Syntax und Semantik anhand eines solchen Formalismus bildet daher einen weiteren wesentlichen Aspekt unserer Betrachtungen.

Die formale Ausrichtung dieser Arbeit und die drei Schwerpunkte Morphologie, Syntax und Semantik spiegeln sich in der Gliederung der folgenden Betrachtungen:

Kapitel 1 bietet einen Überblick über Strategie und Methodik formaler Beschreibungen in der Linguistik: Der benutzte Formalismus TFS (Emele 1993) wird anhand von Beispielen beschrieben und mathematisch fundiert.

Kapitel 2 führt in die japanische Verbmorphologie ein. Ausgehend von traditionellen japanischen Beschreibungen wird ein System entwickelt, daß sich zur formalen Modellierung eignet. In Teilen, wie der Subklassifizierung der morphologischen Primitiven und der Anzahl und Struktur der Stammformen, baut dieses Kapitel auf meiner Studienarbeit (Bollmann 1997) auf. Im Gegensatz zu dem dort entwickelten, defaultlogischen Modell, wird jedoch eine Beschreibung angestrebt, die sich — ohne große Einbußen in Eleganz und Ökonomie — mit monotonen Mitteln umsetzen läßt.

Kapitel 3 beschreibt den syntaktisch-semantischen Rahmen der Morphologie. Auf vorhandenen Beschreibungen des Japanischen — wie der Verbmobilgrammatik (Siegel 1996b), der JPSG (Gunji 1987, Gunji 1991, Gunji 1999) und der HPSG (Manning et al. 1999) — aufbauend, wird eine minimale Syntax und Semantik des Japanischen beschrieben. Dargestellt werden nur die Aspekte der japanischen Grammatik, die mit der Morphologie in Wechselwirkung treten und daher zu ihrer Beschreibung vorausgesetzt werden müssen.

In *Kapitel 4* wird die morphologische Beschreibung, die in Kapitel 2 entwickelt wurde, formalisiert und in die syntaktisch-semantische Beschreibung aus Kapitel 3 integriert. Auch wenn Ableitung und Struktur des morphologischen Systems, das in Kapitel 2 vorgestellt wurde, für den menschlichen Betrachter auf den ersten Blick kompliziert erscheinen, zeigen sich hier dessen Vorzüge aus formaler Perspektive: Auf zwei einfachen Typenhierarchien aufbauend lassen sich die Formen japanischer Verben mit nur einem morphologischen Schema einheitlich und elegant erklären.

In *Kapitel 5* wird das morphologische System auf ein repräsentatives Fragment der japanischen Verbmorphologie angewandt. Abgesehen von den speziellen Formen der Höflichkeitssprache (敬語, *keigo*) und einigen altertümlichen Formen werden die meisten Verbformen des modernen Japanischen abgedeckt. Dieses zeigt sich auch in Anhang B, in dem für alle Verbklassen

— von den oben erwähnten Höflichkeitsformen abgesehen — die Ableitung sämtlicher Verbformen gezeigt wird, die in den Verbtafeln des *“The Complete Japanese Verb Guide”* (Ishii et al. 1989) beschrieben werden. Im Gegensatz zu dem *“The Complete Japanese Verb Guide”* erlaubt unser System jedoch ebenso die Ableitung beliebiger Kombinationen dieser Formen.

In *Kapitel 6* wird auf die Vorteile und Probleme eingegangen, die sich aus einem deklarativen Ansatz bei der Operationalisierung der Grammatik ergeben: Aufgrund der Verwendung eines Kalküls zur Interpretation von Constraintbeschreibungen eignet sich TFS zwar gut als experimentelle Umgebung für die Spezifikation von Grammatikausschnitten; für konkrete computerlinguistische Anwendungen hingegen ist es nur begrenzt einsetzbar. Es werden Möglichkeiten zur Bewältigung der auftretenden Schwierigkeiten skizziert und ihre Anwendung anhand der Entwicklungsgeschichte der erarbeiteten japanischen Grammatik illustriert.

Den Schluß dieser Arbeit bildet *Kapitel 7*: Hier werden die Ergebnisse der Arbeit noch einmal zusammengefaßt und mögliche Fortsetzungen angesprochen.

Anhang A beschreibt den Korpus, der in dieser Arbeit verwendet wurde: Zu allen Verbformen werden die benutzten Beispielsätze mit ihren semantischen Entsprechungen aufgelistet.

Anhang B, auf den schon oben hingewiesen wurde, zeigt anhand einer Auswahl repräsentativer Verben, daß die hier beschriebene Morphologie alle Formen abzuleiten gestattet, die in Ishii et al. 1989 aufgelistet werden. Eine Ausnahme bilden die speziellen Höflichkeitsformen, deren Beschreibung mit den vorgestellten Mitteln aber keine Schwierigkeiten verursachen dürfte. Bei der Benennung der Formen wird in *Anhang B*, der besseren Vergleichbarkeit halber, ausnahmsweise Ishii et al. 1989 gefolgt.

Die Beispielsätze, die in dieser Arbeit verwendet werden, sind zum größten Teil dem Werk *“A Dictionary of Basic Japanese Grammar”* von Seiichi Makino und Michio Tsutsui (Makino und Tsutsui 1986) und der *“Japanische[n] Morphosyntax”* von Jens Rickmeyer (Rickmeyer 1995) entnommen. Da die Sätze auch in diesen Werken meist nur zitiert werden, läßt sich dort nachlesen, aus welchen Korpora, Grammatiken, Aufsätzen etc. diese ursprünglich stammen.

Danksagung

Zuallererst möchte ich mich bei Michelle Brehm, Ulrike Meibohm und Gitta-Miriam Oellinger bedanken. Sodann bei den Mitgliedern der Examensgruppe — Susanne Nowodnick, Ralf Behler, Barbara Kröber, Phillip Bartels, Christiane und Joanna Blümel, Roman Czyborra, Annette Kirsch und Giu-

seppe Pitronaci — und meinen Freunden in und außerhalb der Informatik: Obwohl — und oft *weil* — fachfremd, haben sie mir durch Anteilnahme, Unterstützung und Interesse über alle auftretenden Schwierigkeiten hinweggeholfen! Ebensolcher Dank gilt meinen Betreuern Melanie Siegel und Manfred Stede. Beide möchte ich für die Geduld und Ausdauer loben, mit der sie meinen immer neuen Vorstellungen und den daraus folgenden Verzögerungen begegnet sind: Für die Zukunft wünsche ich ihnen weniger “aufwendige” Diplomanden. Melanie Siegels Grammatik des Japanischen hat mir große Dienste erwiesen: Sie bildet Basis und Rahmen meines morphologischen Modells. Bei Fragen zur japanischen Sprache und Grammatik standen mir 中村隆治, 押元有治 und 鬼頭直 zur Seite. Besonders dankbar bin ich ihnen auch für die Hilfe bei der Auswahl und Bildung der japanischen Beispielsätze. Marianne Steffens, Wolfram Rauch und Reimund Noll haben mir mit Geduld und Kritik bei der Ausformulierung der Texte geholfen: Vielen Dank! Elizabeth, Sigrid, Leonard und Sebastian gilt mein Dank für das Vorleben der glücklichen Beispielsituationen aus Kapitel 1. Professor Bernd Mahr danke ich für die weite Auslegung des Begriffes *Informatik*: Seine Aufgeschlossenheit und Freude an interdisziplinären Thematiken machte es mir möglich, meine Interessen in dieser Weise zu kombinieren. In diesem Zusammenhang möchte ich auch Dirk Lüdtke danken, dessen Begeisterung für die Verbindung von Informatik und Japanologie manches Motivationsloch stopfte. Meine Eltern haben mich immer wieder ermutigt, im Studium meinen Interessen und Neigungen zu folgen. Ohne ihr Vertrauen und den durch sie gewährleisteten zeitlichen und finanziellen Freiraum hätte ich mich auf einen meiner Interessensschwerpunkte — Japanisch, Linguistik und Informatik — beschränken müssen!

1 Formale Linguistik und Constraintsysteme

In Anlehnung an die Methoden der exakten Wissenschaften fordert Noam Chomsky 1957 in seinem Buch “Syntactic Structures” (Chomsky 1957) die Benutzung *formaler Modelle* für die linguistische Theoriebildung:

Precisely constructed models for linguistic structure can play an important role, both negative and positive, in the process of discovery itself. By pushing a precise but inadequate formulation to an unacceptable conclusion, we can often expose the exact source of this inadequacy and, consequently, gain a deeper understanding of the linguistic data. More positively, a formalized theory may automatically provide solutions for many problems other than those for which it was explicitly designed. Obscure and intuition-bound notions can neither lead to absurd conclusions nor provide new and correct ones, and hence they fail to be useful in two important respects.¹

An späterer Stelle führt Chomsky weiter aus:

A grammar of the language L is essentially a theory of L . Any scientific theory is based on a finite number of observations, and it seeks to relate the observed phenomena and to predict new phenomena by constructing general laws in terms of hypothetical constructs such as (in physics, for example) “mass” and “electron.” Similarly, a grammar of English is based on a finite corpus of utterances (observations), and it will contain certain grammatical rules (laws) stated in terms of the particular phonemes, phrases, etc., of English (hypothetical constructs). These rules express structural relations among the sentences of the corpus and the indefinite number of sentences generated by the grammar beyond the corpus (predictions).²

Chomsky begründet damit nicht nur eine neue Schule innerhalb der Linguistik — die sogenannte *Generative Transformations-Grammatik*³ —, sondern

¹ Chomsky 1957, S. 5 (*preface*).

² Chomsky 1957, S. 49 (*On the goals of linguistic theory*, 6.1).

³ Die Bezeichnung *Generative Transformations-Grammatik* wird als Sammelbegriff für Theorien verwendet, die in der Tradition stehen, die von Noam Chomsky durch das Buch “*Syntactic Structures*” (Chomsky 1957) begründet wurde. Neuere Versionen dieser Theorie sind unter den Namen *Government and Binding* (Chomsky 1981) und *The Minimalist Program* (Chomsky 1995) bekannt.

prägt mit seinen Forderungen die gesamte moderne Linguistik in entscheidender Weise. Notwendige Folge dieser neuen Tendenz ist die Entwicklung einer Vielzahl mathematischer Formalismen, die speziell auf die Darstellung linguistischer Sachverhalte ausgerichtet sind.

Neben der ‘Generativen Transformationsgrammatik’ ist die ‘Head-Driven Phrase Structure Grammar’ (HPSG, Pollard und Sag 1987, Pollard und Sag 1994) unter den formalen linguistischen Theorien heute wohl die bekannteste.

Merkmalstrukturen und Constraints — der formale Rahmen der HPSG

Diese Arbeit orientiert sich weitgehend an der HPSG und benutzt daher auch ähnliche formale Mittel:

Merkmalstrukturen dienen der formalen Darstellung sprachlicher Entitäten im Modell: Morphe, Wörter, Phrasen, Sätze etc. werden durch Merkmalstrukturen repräsentiert.

Constraints sind das eigentliche Mittel der linguistischen Modellierung: Die Regularitäten aller linguistischen Beschreibungsebenen — morphologische, lexikalische, syntaktische, semantische etc. — werden durch Constraints formalisiert. Constraints schränken die Menge der kombinatorisch möglichen Merkmalstrukturen ein: Nur Strukturen, die allen Constraints genügen, sind im Sinne des linguistischen Modells korrekt.

Resolution — Mit Hilfe eines *Resolutionsverfahrens* lassen sich aus Sätzen linguistische Strukturbeschreibungen, und aus Strukturbeschreibungen Sätze ableiten. Beide Ableitungen entsprechen einer allgemeineren Form der Berechnung: Aus *unterspezifizierten* Merkmalstrukturen wird die Menge der *vollspezifizierten* Merkmalstrukturen errechnet.⁴

Merkmalstrukturen und Constraints zusammen bilden das linguistische Modell, die Resolution dient ihrer Operationalisierung.

Die mathematischen Grundlagen der HPSG — und damit auch unseres Modells — sind Thema der folgenden Abschnitte: Merkmalstrukturen, Constraints und Resolution werden formal eingeführt. Unsere Darstellung orientiert sich an der formalen Spezifikation des Systems TFS (*Typed Feature Structures Representation Formalism*, Emele und Zajac 1990a, Emele und Zajac 1990b, Emele 1991, Emele 1993, Emele 1994, und Zajac 1992), das zur Implementierung verwendet wurde, und an Carpenter 1992, dem “Standardwerk” zur Logik getypter Merkmalstrukturen. Aspekte beider Darstellungen

⁴ Was unter einer *vollspezifizierten* Merkmalstruktur zu verstehen ist, hängt von dem verwendeten Resolutionsverfahren ab (vgl. Absatz 1.5 (Resolution), S. 46).

wurden — zum Teil in veränderter Form — zu einer möglichst einfachen und verständlichen Darstellung integriert.⁵

1.1 Sprachliche Zeichen als Merkmalstrukturen

Form und Eigenschaften sprachlicher Zeichen, also Morphe, Wörter, Sätze etc., werden in Theorien wie der HPSG durch Merkmalstrukturen dargestellt. Der Einführung und Definition dieser *Merkmalstrukturen* dient der folgende Abschnitt.

Eine informelle Vorstellung von Merkmalstrukturen erleichtert das Verständnis der Definitionen. Es werden deshalb, anhand einiger Beispiele, die wichtigsten Begriffe für Merkmalstrukturen zuerst anschaulich vorgestellt. Anschließend wird die AVM⁶-Notation (Merkmal-Wert-Matrizen-Notation) für Merkmalstrukturen eingeführt. Der größeren Anschaulichkeit halber wird sie an vielen Stellen der Graphennotation vorgezogen.

Für die Spezifikation von Merkmalstrukturen spielen *Typen* und *Typenhierarchien* eine zentrale Rolle. Nach der informellen Einführung von Merkmalstrukturen werden daher die Typen und Typenhierarchien formalisiert.

Die Definition der Merkmalstrukturen als Graphen folgt am Ende des Abschnittes.

5 Folge der Vereinfachung sind unter anderem:

1. Merkmalstrukturen werden als gerichtete Graphen modelliert und die AVM-Schreibweise als eine Art Repräsentantensystem für Äquivalenzklassen von Merkmalstrukturen aufgefaßt.
2. Die in der formalen, modelltheoretischen Logik übliche Unterscheidung zwischen *Beschreibungssprache* und *Modelltheoretischer Semantik* wurde aufgegeben. Eine Merkmalstruktur steht — je nach Kontext — entweder für sich selbst, oder für die Menge der Strukturen, die von ihr subsumiert werden.
3. Die Unifikation von Typen — und damit auch von Merkmalstrukturen — wurde nicht-deterministisch definiert: Die Unifikation zweier Typen ergibt eine (evtl. leere) Menge von Typen; die Unifikation von Merkmalstrukturen ergibt eine (evtl. leere) Menge von Merkmalstrukturen.
4. Aus der nicht-deterministischen Definition der Unifikation ergibt sich ein vereinfachtes Resolutionsverfahren.

Die Synthese aus Carpenters und Emeles System wird ebenfalls als TFS bezeichnet. Oft wird auch allgemeiner von *Constraintsystemen* oder *Merkmalstlogiken* gesprochen.

6 Das Kürzel AVM steht für die Initialen der englischen Bezeichnung *Attribute-Value-Matrix*.

1.1.1 Informelle Einführung

Sprachliche Zeichen lassen sich unter verschiedenen Aspekten beschreiben: sie haben eine *Form* und eine *Bedeutung*; sie sind gekennzeichnet durch Eigenschaften, die ihr *morphologisches*, *syntaktisches* und *semantisches* Verhalten bestimmen; sie haben bestimmte *Funktionen im Text* usw. Dies gilt nicht nur für die Wörter einer Sprache, sondern ebenso für Syntagmen, ganze Sätze oder Texte, sowie für morphologische und phonologische Bestandteile von Wörtern.

Sprachliche Ausdrücke aller linguistischen Beschreibungsebenen müssen folglich durch Mengen von Eigenschaften unterschiedlicher Qualität beschrieben werden. Ihr komplexes Verhalten erfordert das Zusammenspiel verschiedener, weitgehend unabhängiger Theorien, wie Syntax, Semantik und Morphologie.

Sollen die sprachlichen Ausdrücke formal beschrieben werden, liegt es damit nahe, ihre unterschiedlichen Eigenschaften hierarchisch zu ordnen. Die Eigenschaften einer Verbalphrase beispielsweise lassen sich in solche unterteilen, die ihre innere Struktur beschreiben, und solche, die sie als Gesamtheit charakterisieren. Letztere Gruppe kann weiter in die Klasse der semantischen und der syntaktischen Eigenschaften unterteilt werden usw. Dadurch ergibt sich eine hierarchische Struktur, die vieles gemein hat mit den Datenstrukturen moderner objektorientierter Programmiersprachen.

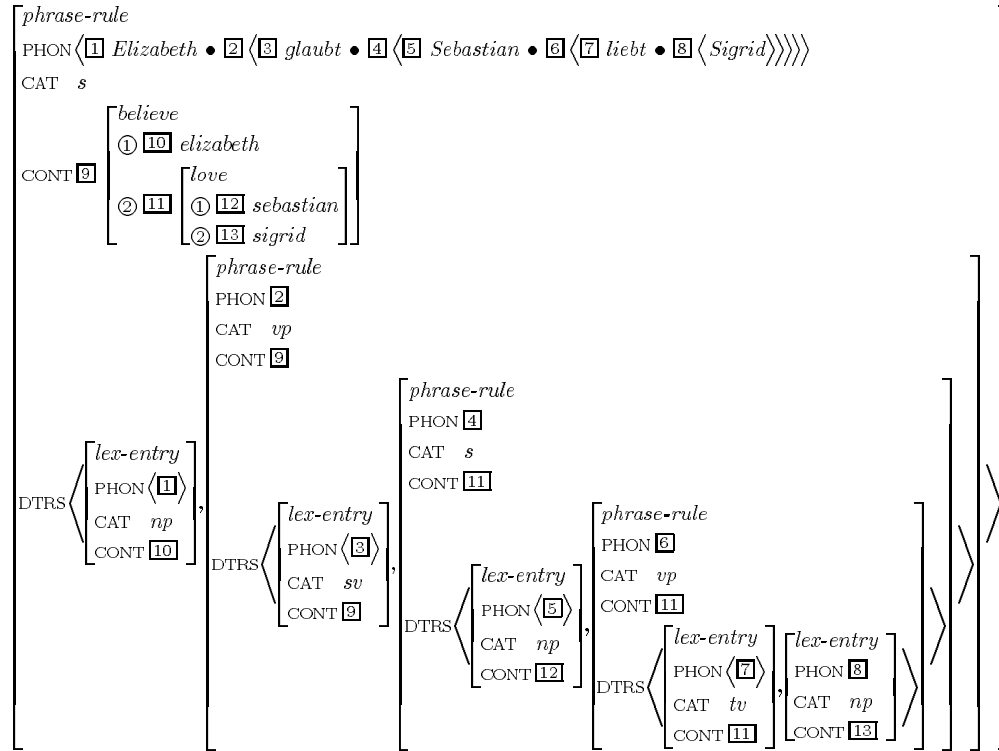
In der Linguistik — und insbesondere im Paradigma der Unifikationsgrammatiken — hat sich der Begriff *Merkmal-Wert-Matrix* (AVM)⁷ für diese Art von Datenstrukturen eingebürgert. Eher mathematische Darstellungen bevorzugen den Begriff *Merkmalstruktur*. In diesem Kapitel werden Merkmal-Wert-Matrizen als Notation für Merkmalstrukturen eingeführt. Im linguistischen Teil der Arbeit ist diese Unterscheidung nicht wichtig, so daß beide Begriffe dort synonym verwendet werden.

1.1.2 Aufbau und Nomenklatur von Merkmalstrukturen

Merkmalstrukturen lassen sich nicht nur zur Darstellung und Gliederung von sprachlichen Zeichen (Morphen, Wörtern, Phrasen, Sätzen, Texten etc.) verwenden; auch die Bestandteile von Zeichen (Form, Bedeutung etc.) und die Regularitäten, die diese betreffen (Syntaktische Schemata, Prinzipien etc.), werden durch Merkmalstrukturen dargestellt. Um die Universalität von Merkmalstrukturen besser verstehen zu können, ist es notwendig, ihren Aufbau genauer zu betrachten. Sehen wir uns dazu drei unterschiedlich komplexe Beispiele an:

⁷ engl.: *Attribute-Value-Matrix*.

(1) Merkmalstruktur für den Satz “Elizabeth glaubt, Sebastian liebt Sigrid”:



Diese Merkmalstruktur beschreibt — gegeben eine einfache formale Grammatik — Form, Syntax und Semantik des Satzes *Elizabeth glaubt, Sebastian liebt Sigrid*. Merkmalstrukturen können aber auch einfacher aussehen:

$$(2) \quad \left[\begin{array}{l} \text{phrase-rule} \\ \text{PHON} \langle \boxed{1} \text{ Elizabeth} \bullet \boxed{2} \langle \text{schläft} \rangle \rangle \\ \text{DTRS} \left\langle \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON} \langle \boxed{1} \rangle \end{array} \right], \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON} \quad \boxed{2} \end{array} \right] \right\rangle \end{array} \right]$$

(2) beschreibt Form und Syntax des Satzes *Elizabeth schläft*. Im einfachsten Fall besteht eine Merkmalstruktur nur aus einem Symbol:

(3) *Leonard*

Diese Struktur kann z.B. zur Beschreibung der phonologischen Form des Wortes *Leonard* verwendet werden.

Die wichtigsten Begriffe und Charakteristika von Merkmalstrukturen lassen sich anhand dieser Beispiele veranschaulichen:

Merkmale⁸ — In den Merkmalstrukturen (1) und (2) befinden sich rechts neben den sich öffnenden eckigen Klammern Symbole, die in KAPITÄLCHEN gesetzt sind: Diese werden *Merkmale* genannt. Auch mit kleinen Kreisen umrandete Zahlen wie ① und ② sind Merkmale.

In (1) finden sich die Merkmale PHON, CAT, CONT, ①, ② und DTRS. Merkmalstruktur (2) enthält die Merkmale PHON und DTRS.

Typen — In allen drei Merkmalstrukturen befinden sich kursiv geschriebene Symbole, die sogenannten *Typen*:

- Das Zeichen (3) beispielsweise besteht nur aus dem Typ *Leonard*.
- In den Zeichen (1) und (2) befinden sich an der linken oberen Ecke sich öffnender, eckiger Klammern ebenfalls Typen: In (1) beispielsweise *phrase-rule*, *believe*, *love* und *lex-entry*; In (2) befinden sich die Typen *phrase-rule* und *lex-entry*.
- Hinter Merkmalen können ebenfalls Typen stehen: in (1) z.B. unter anderem *s*, *elizabeth*, *sebastian*, *sigrid*, *np*, *vp* usw.

Merkmal-Wert-Paare — Typisch für (1) und (2) sind Paare aus Merkmalen und Typen bzw. Klammersausdrücken. Erstere heißen *Merkmale*, letztere *Merkmalwerte*. Paare aus einem Merkmal und einem Merkmalwert heißen entsprechend *Merkmal-Wert-Paare*.

In (1) beispielsweise finden sich in der ersten Klammerebene folgende Merkmal-Wert-Paare:

(4)	Merkmal:	Merkmalwert:
	PHON	$\langle \boxed{1} \textit{Elizabeth} \bullet \boxed{2} \langle \boxed{3} \textit{glaubt} \dots \rangle \rangle$
	CAT	<i>s</i>
	CONT	$\boxed{9} \left[\begin{array}{l} \textit{believe} \\ \textcircled{1} \boxed{10} \textit{elizabeth} \\ \textcircled{2} \boxed{11} \left[\begin{array}{l} \textit{love} \\ \textcircled{1} \boxed{12} \textit{sebastian} \\ \textcircled{2} \boxed{13} \textit{sigrid} \end{array} \right] \end{array} \right]$
	DTRS	$\langle \left[\begin{array}{l} \textit{lex-entry} \\ \text{PHON} \langle \boxed{1} \rangle \\ \text{CAT} \textit{np} \\ \text{CONT} \boxed{10} \end{array} \right], \dots \rangle$

8 engl.: *attribute* oder *feature*.

Atomare und komplexe Merkmalstrukturen — Die Menge der Merkmalstrukturen zerfällt in zwei Untermengen: die *atomaren* und die *komplexen* Merkmalstrukturen. Im Gegensatz zu den atomaren Merkmalstrukturen, die nur aus einem Typ bestehen,⁹ setzen sich die komplexen Merkmalstrukturen aus einem Paar eckiger Klammern zusammen, die durch einen Typ gekennzeichnet werden und eine Menge von Merkmal-Wert-Paaren umschließen.

Beispiel (3) zeigt eine *atomare Merkmalstruktur*; (1) und (2) hingegen repräsentieren komplexe Merkmalstrukturen.

Listen — Neben den eckigen Klammern gibt es in (1) und (2) auch spitze Klammern, die ein oder mehrere durch Kommata getrennte Elemente enthalten. Diese Strukturen heißen *Listen*. Wie wir später sehen werden, sind die spitzen Klammern nur eine Abkürzung für spezielle, rekursive Strukturen aus eckigen Klammern.

Variablen — Zwischen Merkmalen und Werten oder an Stelle der Werte stehen in manchen Fällen rechteckig umrandete Zahlen: die *Variablen*. Steht hinter einer Variablen ein Wert, wird die Variable mit diesem belegt. Merkmale mit gleicher Variable haben (extensional) gleiche Werte.

Durch Merkmalstrukturen lassen sich damit, dank der Variablen, auch Graphen darstellen.

Rekursivität — Listen und Variablen sind nur Kürzel für eckig geklammerter Strukturen. Die *Rekursivität* kann daher als wichtigstes Strukturprinzip der Merkmalstrukturen festgestellt werden:

Merkmalstrukturen sind entweder:

atomar — d.h. sie bestehen nur aus einem Typ (vgl. (3)), oder

komplex — d.h. sie bestehen aus einem eckigen Klammerpaar, einem Typ und einer Menge von Merkmal-Wert-Paaren (vgl. (1) und (2)). Die Werte sind rekursiv nach dem gleichen Prinzip aufgebaut: Sie sind also selbst wiederum atomar oder komplex.

Merkmalstrukturen sind damit *rekursive Datenstrukturen*.

Merkmal-Wert-Matrizen (AVM)

Merkmalstrukturen sind formal gesehen Graphen. Wenn sie mit der Hilfe von Variablen trotzdem als Klammerstrukturen geschrieben werden, gibt es dafür einen einzigen, einfachen Grund: Klammerstrukturen können graphisch leichter dargestellt werden und sind leichter lesbar.

⁹ Atomare Merkmalstrukturen können mit oder ohne eckige Klammern dargestellt werden: [*Leonard*] und *Leonard* stehen also für die gleiche atomare Merkmalstruktur.

Für die geklammerte Form der Darstellung von Merkmalstrukturen soll von nun an der Begriff *Merkmal-Wert-Matrix (AVM)* Verwendung finden. Auch wenn die Begriffe *Merkmalstruktur* und *Merkmal-Wert-Matrix* in dem linguistischen Teil dieser Arbeit synonym gebraucht werden, sind AVMs unserer Definition nach nur eine *Notation* für Merkmalstrukturen. Merkmalstrukturen werden im Abschnitt 1.3 formal als Graphen definiert, da sich Operationen, wie die Unifikation, leichter für Graphen definieren lassen.

Die Beobachtungen des letzten Abschnittes lassen sich in einer einfachen Definition von AVMs zusammenfassen:

Definition 1.1 (Merkmal-Wert-Matrizen (AVM)¹⁰)

Sei **Type** eine endliche Menge von Typen, **Feat** eine endliche Menge von Merkmalen und **Var** = $\{\mathbb{1}, \mathbb{2}, \mathbb{3}, \dots\}$ eine Menge von Variablen. Die Menge der Merkmal-Wert-Matrizen (AVM) über **Type** und **Feat** ist die kleinste Menge AVM, die die folgenden Bedingungen erfüllt:

$$1. \begin{bmatrix} \sigma \\ f_1 \phi_1 \\ f_2 \phi_2 \\ \vdots \\ f_n \phi_n \end{bmatrix} \in \text{AVM} \quad \text{falls} \quad \begin{array}{l} \sigma \in \text{Type}, \\ n \in \mathbb{N}_0, \\ f_1, f_2, \dots, f_n \in \text{Feat}, \\ \phi_1, \phi_2, \dots, \phi_n \in \text{AVM}; \end{array}$$

Für $n = 0$ erhalten wir folglich den Spezialfall $[\sigma]$, für den wir auch einfach σ schreiben.

$$2. v \in \text{AVM} \quad \text{falls} \quad v \in \text{Var}$$

$$3. v\phi \in \text{AVM} \quad \text{falls} \quad v \in \text{Var} \quad \text{und} \quad \phi \in \text{AVM}$$

10 Eine andere Notation wird in Carpenter 1992 (Def. 4.1 (Descriptions)) vorgestellt: Die Merkmal-Wert-Paare, die sich in unserer Darstellung in einer Klammerebene befinden, werden als logische Konjunktionen beschrieben, und Merkmale mit gleichen Werten werden als Gleichungen dargestellt:

Die Struktur

$$\begin{bmatrix} a \\ x \mathbb{1} b \\ y \begin{bmatrix} c \\ z \mathbb{1} \end{bmatrix} \end{bmatrix}$$

beispielsweise wird in seiner Notation durch

$$a \wedge (x : b) \wedge (y : c) \wedge (x \doteq y \cdot z)$$

dargestellt.

- Der Typ *Person* kann die “Werte” *Frau*, *Feinschmecker*, *Mann*, *Elizabeth*, *Sigrid* und *Sebastian* annehmen.
- *Elizabeth*, *Sigrid* und *Sebastian* sind *minimale Typen*¹¹ und können nicht durch speziellere Typen ersetzt werden.
- *Frau*, *Feinschmecker* und *Mann* hingegen sind spezieller und damit “informativer” als *Person*, können aber trotzdem noch weiter präzisiert werden:
 - *Frau* durch *Elizabeth* oder *Sigrid*,
 - *Feinschmecker* durch *Sigrid* oder *Sebastian*,
 - *Mann* durch *Sebastian*.

größter Typ — Typenhierarchien enthalten einen *größten Typ*¹²: Dieser dominiert alle anderen Typen, ohne selber dominiert zu werden. Er wird durch das Symbol ‘ $\top$ ’ (gelesen: “top”) bezeichnet.

In der Hierarchie (7) auf S. 13 beispielsweise ist *Person* ($= \top$) der größte Typ.

Aus den zwei Eigenschaften ergibt sich die Definition von Typenhierarchien:

Definition 1.2 (Typenhierarchie¹³)

Eine Typenhierarchie ist ein 3-Tupel $\mathcal{H} = \langle \text{Type}, \sqsubseteq, \top \rangle$, das die folgenden Bedingungen erfüllt:

1. *Type* ist eine endliche Menge von Typen mit $\top \in \text{Type}$;
2. $\langle \text{Type}, \sqsubseteq \rangle$ ist eine partielle Ordnung;
3. $\top$ (gelesen: “top”) ist das größte Element¹⁴ hinsichtlich der partiellen Ordnung $\langle \text{Type}, \sqsubseteq \rangle$,
d.h. $\forall \sigma \in \text{Type} : \top \sqsubseteq \sigma$.

¹¹ *minimale* und *maximale Elemente* einer geordneten Menge M sind folgendermaßen definiert:

$$m = \text{maximales Element}(M) : \Longleftrightarrow m \in M \wedge \forall x \in M [m \sqsubseteq x \Rightarrow x = m]$$

$$m = \text{minimales Element}(M) : \Longleftrightarrow m \in M \wedge \forall x \in M [m \sqsupseteq x \Rightarrow x = m]$$

¹² *größtes* und *kleinstes Element* einer geordneten Menge M sind folgendermaßen definiert:

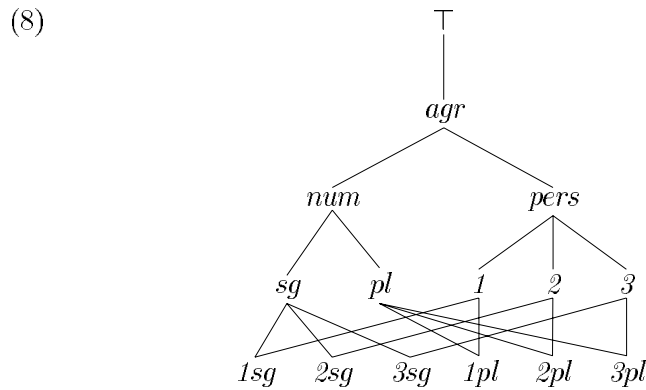
$$g = \text{größtes Element}(M) : \Longleftrightarrow g \in M \wedge \forall x \in M [x \sqsubseteq g]$$

$$k = \text{kleinstes Element}(M) : \Longleftrightarrow k \in M \wedge \forall x \in M [x \sqsupseteq k]$$

Subsumtion von Typen

Die partielle Ordnung $\sqsubseteq$ wird bei Typenhierarchien auch als *Subsumtionsordnung* bezeichnet. Für $\sigma \sqsubseteq \tau$ sagt man daher: “ σ subsumiert τ ” bzw. “ τ wird durch σ subsumiert”.

Das folgende Beispiel dient der Modellierung der *Kongruenzmerkmale*:



In (8) gelten unter anderem die folgenden Subsumtionsbeziehungen:

- (9)
- a. $\forall_{\sigma \in \text{Type}} : \top \sqsubseteq \sigma$;
 - b. $\text{agr} \sqsubseteq \top$;
 - c. $\text{num} \sqsubseteq \text{agr}, \text{num} \sqsubseteq \top$;
 - d. $\text{sg} \sqsubseteq \text{num}, \text{sg} \sqsubseteq \text{agr}, \text{sg} \sqsubseteq \top$;
 - e. $1\text{sg} \sqsubseteq \text{sg}, 1\text{sg} \sqsubseteq \text{num}, 1\text{sg} \sqsubseteq \text{agr}, 1\text{sg} \sqsubseteq \top, 1\text{sg} \sqsubseteq 1, 1\text{sg} \sqsubseteq \text{pers}$;

Der Einfachheit halber schreiben wir für $\sigma \sqsubseteq \sigma_1, \sigma \sqsubseteq \sigma_2, \dots, \sigma \sqsubseteq \sigma_n$ auch $\sigma \sqsubseteq \{\sigma_1, \sigma_2, \dots, \sigma_n\}$

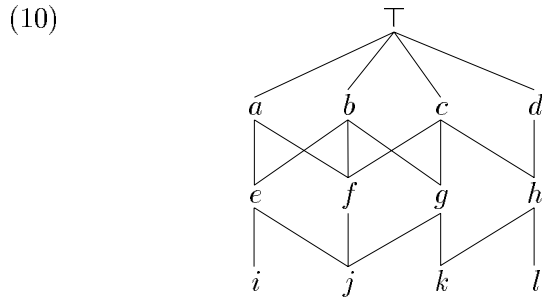
13 Carpenter 1992 definiert Typenhierarchien als *finite bounded complete partial order* $\langle \text{Type}, \sqsubseteq \rangle$, vgl. Carpenter 1992, Def. 2.1 (*Inheritance Hierarchie*).

Dank dieser Definition läßt sich die Unifikation als partielle Funktion in die Menge der Typen definieren: Typen haben in seinem System also entweder einen Typ als eindeutiges Unifikat oder sie sind nicht unifizierbar.

14 Carpenter 1992 betrachtet Typenhierarchien als Träger von “Information”: der Typ *Person* aus (7) enthält nach dieser Auffassung weniger Information als ein speziellerer Typ wie z.B. *Frau* oder *Sigrid*. Diese Interpretation der Typenhierarchie hat zur Folge, daß der allgemeinste Typ — d.h. der Typ, der von allen Typen am wenigsten Information trägt — nicht, wie in unserem Fall, *oberhalb*, sondern *unterhalb* aller anderen Symbole dargestellt wird. Carpenter verwendet daher für diesen Typ auch nicht das Symbol $\top$, sondern das Symbol $\perp$. Die Bezeichnungen *maximaler Typ* und *minimaler Typ* sind in entsprechender Weise vertauscht. Auch die Subsumtionsordnung $\sqsubseteq$ wird von Carpenter genau spiegelverkehrt definiert.

Unifikation von Typen

Es gibt *unifizierbare*, d.h. “verträgliche”, und *nicht unifizierbare*, d.h. “unverträgliche”, Typen. *Unifizierbare Typen* besitzen gemeinsame Untertypen, *nicht unifizierbare Typen* haben keine gemeinsamen Untertypen:



In dieser Typenhierarchie gelten unter anderem die folgenden Beziehungen:

- (11)
- | | |
|----|------------------------------------|
| a. | $a \supseteq \{e, f, i, j\}$ |
| b. | $b \supseteq \{e, f, g, i, j, k\}$ |
| c. | $c \supseteq \{f, g, h, j, k, l\}$ |
| d. | $d \supseteq \{h, k, l\}$ |

b und *c* haben gemeinsame Untertypen und sind damit unifizierbar:

$$(12) \quad \{e, f, g, i, j, k\} \cap \{f, g, h, j, k, l\} = \{f, g, j, k\}$$

a und *d* hingegen haben keine gemeinsamen Untertypen, sind also auch nicht unifizierbar:

$$(13) \quad \{e, f, i, j\} \cap \{h, k, l\} = \emptyset$$

Die Schnittmenge der Untertypen enthält zu jedem ihrer Elemente auch alle seine Untertypen. Sie läßt sich daher durch die Menge ihrer maximalen Typen charakterisieren: Die Menge der Maximalen Elemente aus $\{f, g, j, k\}$ relativ zu (10) beispielsweise ist $\{f, g\}$.

Diese Menge — also die Menge der maximalen Typen aus der Menge der gemeinsamen Untertypen — definieren wir als *Unifikation*. Noch einfacher

läßt sich die Unifikation mit dem Begriff der *maximalen unteren Schranke*¹⁵ charakterisieren:

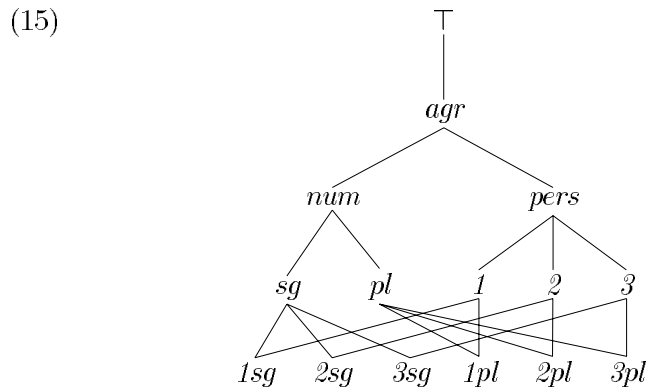
Definition 1.3 (Unifikation von Typen)

Gegeben sei eine Typenhierarchie $\langle \mathbf{Type}, \sqsubseteq, \top \rangle$. Unifikation heißt die Funktion $\sqcup : \mathbf{Type} \times \mathbf{Type} \longrightarrow 2^{\mathbf{Type}}$, die jedem Tupel $\langle \sigma, \tau \rangle \in \mathbf{Type} \times \mathbf{Type}$ die Menge seiner maximalen unteren Schranken hinsichtlich der partiellen Ordnung $\langle \mathbf{Type}, \sqsubseteq \rangle$ zuweist.

Für zwei Elemente $\sigma, \tau \in \mathbf{Type}$ ist $\sigma \sqcup \tau$ folglich nur dann ungleich der leeren Menge, wenn zumindest ein Element $\varphi \in \mathbf{Type}$ existiert, mit $\varphi \sqsubseteq \sigma$ und $\varphi \sqsubseteq \tau$. Für die Typenhierarchie (10) beispielsweise erhalten wir:

- (14) a. $a \sqcup d = \emptyset$;
 b. $b \sqcup c = \{f, g\}$;
 c. $b \sqcup f = \{f\}$;
 d. $b \sqcup b = \{b\}$;
 e. $i \sqcup j = \emptyset$;

Warum diese Definition der Unifikation sinnvoll ist, läßt sich leichter anhand der Hierarchie der *Kongruenzmerkmale* verstehen, die hier noch einmal dargestellt wird:



¹⁵ Sei M eine geordnete Menge und $T \subseteq M$, so sind die Begriffe *untere* und *obere Schranke* folgendermaßen definiert:

$$\begin{aligned} s &= \text{obere Schranke}(T) :\iff s \in M \wedge \forall_{x \in T} [x \sqsubseteq s] \\ s &= \text{untere Schranke}(T) :\iff s \in M \wedge \forall_{x \in T} [x \sqsupseteq s] \end{aligned}$$

Für diese Hierarchie gilt unter anderem:

- (16) a. $sg \sqcup pl = \emptyset$;
 b. $sg \sqcup sg = \{sg\}$;
 c. $sg \sqcup agr = \{sg\}$;
 d. $sg \sqcup 1 = \{1sg\}$;
 e. $sg \sqcup pers = \{1sg, 2sg, 3sg\}$;

1.3 Merkmalstrukturen

Nachdem Typenhierarchien und die Begriffe Subsumtion und Unifikation von Typen formal eingeführt wurden, soll nun der Begriff *Merkmalstruktur* präzisiert werden.

Am Anfang dieses Kapitels wurden Merkmalstrukturen informell beschrieben und ihre AVM-Schreibweise vorgestellt. Für die formale Behandlung von Merkmalstrukturen ist es jedoch sinnvoll, sie als Graphen zu beschreiben. Zu diesem Zweck wollen wir die Struktur (2) wieder aufgreifen und schrittweise den ihr entsprechenden Graphen ableiten:

$$(17) \left[\begin{array}{l} \text{phrase-rule} \\ \text{PHON} \left\langle \boxed{1} \text{ Elizabeth} \bullet \boxed{2} \langle \text{schläft} \rangle \right\rangle \\ \text{DTRS} \left\langle \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON} \langle \boxed{1} \rangle \end{array} \right], \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON} \boxed{2} \end{array} \right] \right\rangle \end{array} \right]$$

Auflösung der Abkürzungen:

Listen werden in AVMs meistens durch spitze Klammern dargestellt. Diese sind jedoch nur Abkürzungen für rekursive Merkmalstrukturen:

elist steht für die leere Liste:

- (18) Abkürzung: abgekürzte Merkmalstruktur:

$$\langle \rangle \qquad \qquad \qquad elist$$

Nicht-leere Listen werden durch Merkmalstrukturen des Types *nelist* dargestellt. HEAD beschreibt das erste Element der Liste, TAIL die Restliste:

(19) Abkürzung: abgekürzte Merkmalstruktur:

$$\langle a, b, c \rangle \quad \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } a \\ \text{TAIL } \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } b \\ \text{TAIL } \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } c \\ \text{TAIL } \textit{elist} \end{array} \right] \end{array} \right] \end{array} \right]$$

Das Zeichen ‘•’ verbindet die ersten Elemente einer Liste mit der Restliste:

(20) Abkürzung: abgekürzte Merkmalstruktur:

$$\langle a, b \bullet \textit{list} \rangle \quad \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } a \\ \text{TAIL } \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } b \\ \text{TAIL } \textit{list} \end{array} \right] \end{array} \right]$$

Ersetzen wir die Listen in (17) durch ihre Merkmalstrukturen und schreiben atomare Merkmalstrukturen in geklammerter Schreibweise (also [*Elizabeth*] statt *Elizabeth*, vgl. Def. 1.1 (Merkmal-Wert-Matrizen)) erhalten wir die folgende Darstellung:

$$(21) \left[\begin{array}{l} \textit{phrase-rule} \\ \text{PHON} \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } \boxed{1} [\textit{Elizabeth}] \\ \text{TAIL } \boxed{2} \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } [\textit{schl\"{u}ft}] \\ \text{TAIL } [\textit{elist}] \end{array} \right] \end{array} \right] \\ \text{DTRS} \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD} \left[\begin{array}{l} \textit{lex-entry} \\ \text{PHON} \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD } \boxed{1} \\ \text{TAIL } [\textit{elist}] \end{array} \right] \end{array} \right] \\ \text{TAIL} \left[\begin{array}{l} \textit{nelist} \\ \text{HEAD} \left[\begin{array}{l} \textit{lex-entry} \\ \text{PHON } \boxed{2} \end{array} \right] \\ \text{TAIL } [\textit{elist}] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

Umsetzung der AVM in einen Graphen:

Merkmalstrukturen wie (21) werden folgendermaßen in Graphen umgesetzt:

1. *Klammern* und *Typen* werden durch Knoten dargestellt, an denen ihr Typ angetragen wird:

$$(22) \left[\begin{array}{l} \textit{lex-entry} \\ \dots \end{array} \right] \longrightarrow \boxed{\textit{lex-entry}}$$

Die äußere Klammerenebene entspricht dem Wurzelknoten (graphisch durch eine stärkere Umrandung hervorgehoben):

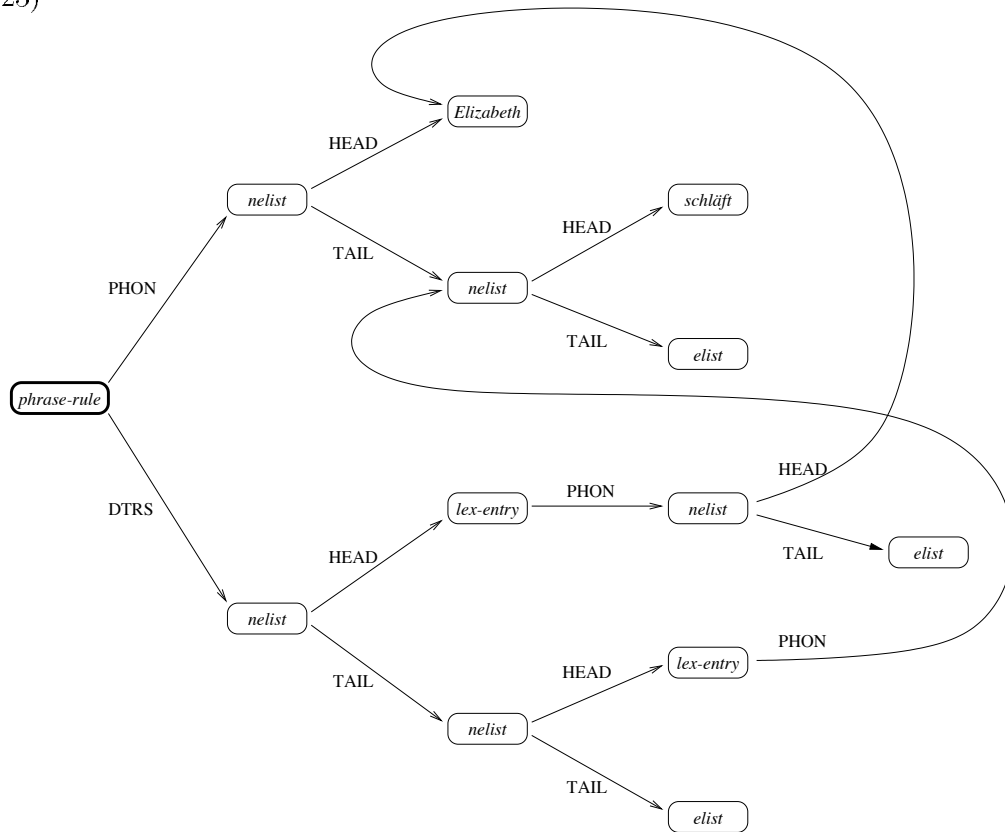
$$(23) \left[\begin{array}{l} \textit{phrase-rule} \\ \dots \end{array} \right] \longrightarrow \boxed{\boxed{\textit{phrase-rule}}}$$

2. *Merkmal-Wert-Paare* werden durch annotierte Kanten dargestellt:

$$(24) \left[\begin{array}{l} \textit{lex-entry} \\ \text{PHON} \left[\begin{array}{l} \textit{nelist} \\ \dots \end{array} \right] \\ \dots \end{array} \right] \longrightarrow \boxed{\textit{lex-entry}} \xrightarrow{\text{PHON}} \boxed{\textit{nelist}}$$

Damit erhalten wir aus (21) den folgenden Graphen:

(25)



Formal läßt sich dieser Graph folgendermaßen beschreiben:

- (26) 1. Q ist die Menge der Knoten:
 $Q = \{q_0, q_1, q_2, \dots, q_{12}\}$
2. $q_0 \in Q$ ist der Wurzelknoten
3. Die Funktion $\theta : Q \rightarrow \text{Type}$ ordnet jedem Knoten des Graphen einen Typ zu:
- | | | |
|-------------------------------------|-----------------------------------|--------------------------------------|
| $\theta(q_0) = \text{phrase-rule},$ | $\theta(q_1) = \text{nelist},$ | $\theta(q_2) = \text{Elizabeth},$ |
| $\theta(q_3) = \text{nelist},$ | $\theta(q_4) = \text{schläft},$ | $\theta(q_5) = \text{elist},$ |
| $\theta(q_6) = \text{nelist},$ | $\theta(q_7) = \text{lex-entry},$ | $\theta(q_8) = \text{nelist},$ |
| $\theta(q_9) = \text{elist},$ | $\theta(q_{10}) = \text{nelist},$ | $\theta(q_{11}) = \text{lex-entry},$ |
| $\theta(q_{12}) = \text{elist}.$ | | |

4. Die partielle Funktion $\delta : \mathbf{Feat} \times Q \longrightarrow Q$ beschreibt die gerichteten Kanten mit den zugeordneten Merkmalen:

$$\begin{aligned} \delta(\text{PHON}, q_0) &= q_1, & \delta(\text{DTRS}, q_0) &= q_6, \\ \delta(\text{HEAD}, q_1) &= q_2, & \delta(\text{TAIL}, q_1) &= q_3, \\ \delta(\text{HEAD}, q_3) &= q_4, & \delta(\text{TAIL}, q_3) &= q_5, \\ \delta(\text{HEAD}, q_6) &= q_7, & \delta(\text{TAIL}, q_6) &= q_{10}, \\ \delta(\text{PHON}, q_7) &= q_8, \\ \delta(\text{HEAD}, q_8) &= q_2, & \delta(\text{TAIL}, q_8) &= q_9, \\ \delta(\text{HEAD}, q_{10}) &= q_{11}, & \delta(\text{TAIL}, q_{10}) &= q_{12}, \\ \delta(\text{PHON}, q_{11}) &= q_3. \end{aligned}$$

Eine Merkmalstruktur F ist also ein gerichteter, evt. zyklischer Graph $F = \langle Q, \bar{q}, \theta, \delta \rangle$ mit einem ausgezeichneten Wurzelknoten $q_0 \in Q$.

Ein Charakteristikum von Merkmalstrukturen ist, daß alle Kanten vom Wurzelknoten ausgehend mittels der gerichteten Kanten erreicht werden können. Man sagt: F ist in $\bar{q}$ *verwurzelt*.

Von unserem Beispiel ausgehend können wir nun den Begriff *Merkmalstruktur* definieren:

Definition 1.4 (Merkmalstruktur (MS)¹⁶)

Sei **Type** eine endliche Menge von Typen und **Feat** eine endliche Menge von Merkmalen. Ein Tupel $F = \langle Q, \bar{q}, \theta, \delta \rangle$ ist eine Merkmalstruktur (MS), gdw. gilt:

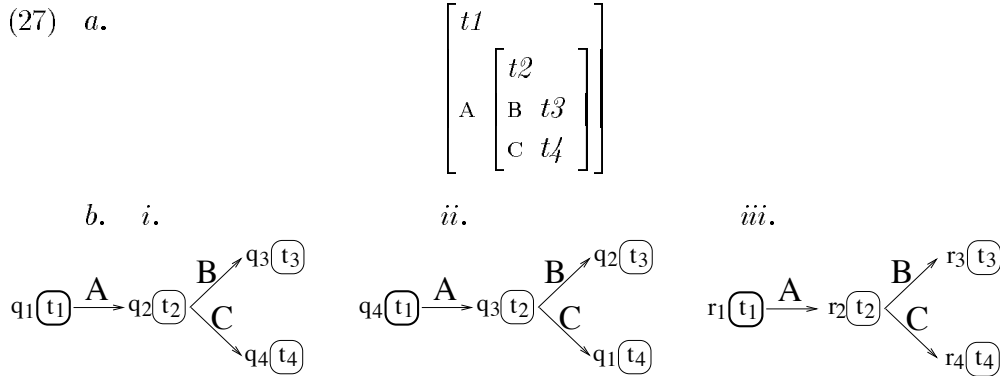
1. Q ist eine endliche Menge von Knoten;
2. Es gibt einen ausgezeichneten Knoten $\bar{q} \in Q$, genannt Wurzelknoten;
3. $\theta : Q \longrightarrow \mathbf{Type}$ ist eine Funktion, die jedem Knoten des Graphen einen Typ zuordnet;
4. $\delta : \mathbf{Feat} \times Q \longrightarrow Q$ ist eine partielle Funktion, durch welche die gerichteten Kanten mit den zugeordneten Merkmalen beschrieben werden.
5. F ist verwurzelt in $\bar{q}$.

Mit $\mathbf{MS}^{\mathbf{Type}, \mathbf{Feat}}$ sei die Menge der Merkmalstrukturen über **Type** und **Feat** bezeichnet.

¹⁶ vgl. Carpenter 1992, Def. 3.1 (Feature Structure)

Alphabetische Varianten

Die AVM-Notation von Merkmalstrukturen unterscheidet sich in einem wesentlichen Punkt von der Graphendarstellung: Graphen, die bezüglich der Kanten und Typen isomorph sind, können sich hinsichtlich der Knoten unterscheiden; Die AVM-Notation hingegen abstrahiert von diesem Unterschied.



In linguistischen Abhandlungen findet normalerweise ausschließlich die Matrixschreibweise Verwendung: Aus linguistischer Sicht ist nur die Struktur der Kanten und Typen interessant, nicht die Identität der Knoten.

Um die Abstraktion von den Knoten formal ausdrücken zu können, führen wir den Begriff *alphabetische Variante* ein:

Definition 1.5 (Alphabetische Variante¹⁷)

Zwei Merkmalstrukturen F_1 und F_2 sind alphabetische Varianten, wenn sie bezüglich ihrer Kanten und Typen isomorph sind.

Wir schreiben in diesem Fall: $F_1 \sim F_2$

Die Isomorphie ist eine Äquivalenzbeziehung¹⁸. Die Menge der Merkmalstrukturen MS zerfällt damit in Äquivalenzklassen (Mengen alphabetischer Varianten) über die sich die Quotientenmenge (die Menge der Äquivalenzklassen) konstruieren läßt. Diese Quotientenmenge ist für uns die eigentlich interessante. Die Matrixschreibweise läßt sich als eine Art Repräsentantensystem¹⁹ bezüglich der alphabetischen Varianten auffassen.

17 vgl. Carpenter 1992, Def. 3.6 (Alphabetic Variants): Carpenter definiert Alphabetische Varianten über die Subsumtionsbeziehung: *If F and F' are feature structures such that $F \sqsubseteq F'$ and $F' \sqsubseteq F$, then we write $F \sim F'$ and say that they are alphabetic variants.*

18 Äquivalenzrelationen sind Relationen, die reflexiv, symmetrisch und transitiv sind.

19 Auch die Matrixschreibweise ist nicht eindeutig, wie dieses für Repräsentantensysteme gefordert ist. Das folgende Beispiel beweist dieses:

$$\begin{bmatrix} A & \boxed{a} \\ B & \boxed{a} \end{bmatrix} \sim \begin{bmatrix} A & \boxed{a} \\ B & \boxed{a} \end{bmatrix} \sim \begin{bmatrix} B & \boxed{a} \\ A & \boxed{a} \end{bmatrix} \sim \begin{bmatrix} B & \boxed{a} \\ A & \boxed{a} \end{bmatrix}$$

Hilfsdefinitionen

Um die Ausweitung der Begriffe *Subsumtion* und *Unifikation* auf die Menge der Merkmalstrukturen **MS** im nächsten Kapitel zu vereinfachen, sollen an dieser Stelle einige Begriffe eingeführt werden, die die formale Behandlung von Merkmalstrukturen vereinfachen.

Merkmale verbinden als gerichtete Kanten die Knoten von Merkmalstrukturen. Oft ist es interessant, von einem Knoten ausgehend nacheinander mehreren Kanten zu folgen. Eine solche Folge von Kanten wird als *Pfad* bezeichnet:

Definition 1.6 (Pfad)

Ein Pfad $\pi \in \mathbf{Path}$ ist eine Folge von Merkmalen, verbunden durch das Zeichen ‘|’:

1. $m \in \mathbf{Path}$ falls $m \in \mathbf{Feat}$;
2. $m|\pi \in \mathbf{Path}$ falls $m \in \mathbf{Feat}$, $\pi \in \mathbf{Path}$;
3. $\mathbf{Path}$ ist die kleinste Menge, die 1. und 2. erfüllt.

Oft wird auch ϵ — der leere Pfad bzw. Pfad der Länge 0 — zu $\mathbf{Path}$ gerechnet. Pfade werden als Abkürzungen auch in der AVM-Notation verwendet:

(28) Abkürzung: abgekürzte Merkmalstruktur:

$$\left[\begin{array}{c} \dots \\ \text{SYNSEM} | \text{LOCAL} | \text{CAT} \text{ value} \\ \dots \end{array} \right] \quad \left[\begin{array}{c} \dots \\ \text{SYNSEM} \left[\begin{array}{c} \top \\ \text{LOCAL} \left[\begin{array}{c} \top \\ \text{CAT} \text{ value} \end{array} \right] \end{array} \right] \\ \dots \end{array} \right]$$

Die Gleichung $\delta(\mathbf{M}, q) = q'$ beschreibt den Übergang von q zu q' entlang des Merkmals $\mathbf{M}$ (vgl. Def. 1.4 (Merkmalstruktur)):

$$(29) \quad q \xrightarrow{\mathbf{M}} q'$$

Oft ist es interessant, zu wissen, welcher Knoten q' von q ausgehend entlang eines Pfades $\pi = \mathbf{M}_1|\mathbf{M}_2|\dots|\mathbf{M}_i$ erreicht wird:

$$(30) \quad q \xrightarrow{\mathbf{M}_1} q_1 \xrightarrow{\mathbf{M}_2} q_2 \quad \dots \quad q_{i-1} \xrightarrow{\mathbf{M}_i} q'$$

Zu diesem Zweck weiten wir δ von der Menge der Merkmale **Feat** auf die Menge der Pfade **Path** aus:

- $$(31) \quad \begin{array}{l} 1. \quad \delta(\epsilon, q) = q \\ 2. \quad \delta(m|\pi, q) = \delta(\pi, \delta(m, q)) \end{array}$$

Merkmalstrukturen sind rekursive Datenstrukturen: Jeder Wert eines Pfades entspricht wiederum einer Merkmalstruktur.

Mit der partiellen Funktion $@ : \mathbf{MS} \times \mathbf{Path} \longrightarrow \mathbf{MS}$ lassen sich diese Unterstrukturen gewinnen:

Definition 1.7 (Wert eines Pfades²⁰)

Sei $F = \langle Q, \bar{q}, \theta, \delta \rangle$. Wenn $\delta(\pi, \bar{q})$ für F definiert ist, so wird die Unterstruktur $F@ \pi = \langle Q', \bar{q}', \theta', \delta' \rangle$ durch die folgenden Bedingungen bestimmt:

1. $\bar{q}' = \delta(\pi, \bar{q})$
2. $Q' = \{\delta(\pi', \bar{q}') \mid \pi' \in \mathbf{Path}\} = \{\delta(\pi|\pi', \bar{q}) \mid \pi' \in \mathbf{Path}\}$
3. $\delta'(f, q) = \delta(f, q) \quad \text{falls} \quad q \in Q'$
4. $\theta'(q) = \theta(q) \quad \text{falls} \quad q \in Q'$

Subsumtion von Merkmalstrukturen

Genau wie Typen lassen sich Merkmalstrukturen in Subsumtionshierarchien anordnen. Wiederum spielt dabei der “Informationsgehalt” die entscheidende Rolle: In der Typenhierarchie (7) enthält der Typ *Mann* mehr Informationen als der Typ *Person*. Die Extension von *Mann* ist daher kleiner als die von *Person*: Sie entspricht der Menge der “Dinge”, die nicht nur *Personen* sind, sondern auch *männlich*. In der Hierarchie (7) spiegelt sich dieser Zusammenhang wieder: Der Typ *Person* subsumiert den Typen *Mann*.

Die Anordnung über den Informationsgehalt bzw. die Extension läßt sich in gleicher Weise auf Merkmalstrukturen ausdehnen:

²⁰ vgl. Carpenter 1992, Def. 3.3 (Path Value), S.39.

1. Ausweitung der Typensubsumtion auf Merkmalstrukturen

Da Typen auch Merkmalstrukturen sind, gilt natürlich auch:

$$(32) \quad Person \sqsubseteq Mann. \quad \text{bzw.} \quad (33) \quad [Person] \sqsubseteq [Mann].$$

Diese Beziehung setzt sich auch auf komplexe Merkmalstrukturen fort:

$$(34) \quad \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \left[\begin{array}{c} sign \\ \text{PHON} \langle Person \rangle \end{array} \right] \bullet \top \right\rangle \end{array} \right] \sqsubseteq \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \left[\begin{array}{c} sign \\ \text{PHON} \langle Mann \rangle \end{array} \right] \bullet \top \right\rangle \end{array} \right]$$

Aus der Perspektive des Informationsgehaltes, ist das erste Zeichen allgemeiner als das zweite: Die Aussage “eine Phrase, deren erstes Zeichen eine Person beschreibt” ist weniger informativ als die Aussage “eine Phrase, deren erstes Zeichen eine männliche Person (Mann) beschreibt”. Vom Standpunkt der Extension aus gesehen, enthält die Menge der Phrasen, die mit einem Zeichen beginnen, das eine *Person* beschreibt, die Menge der Phrasen, die mit einem Zeichen beginnen, das eine *männliche Person* beschreibt. Das erste Zeichen subsumiert also das zweite.

2. Stärker strukturierte Merkmalstrukturen werden von weniger stark strukturierten subsumiert

Stärker strukturierte Merkmalstrukturen enthalten mehr Information als weniger stark strukturierte Strukturen: Ein zusätzliches Merkmal-Wert-Paar beispielsweise fügt einer Merkmalstruktur weitere Information hinzu:

$$(35) \quad \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \left[\begin{array}{c} sign \\ \text{PHON} \langle Mann \rangle \end{array} \right] \bullet \top \right\rangle \end{array} \right] \sqsubseteq \left[\begin{array}{c} phrase \\ \text{CONT } love \\ \text{DTRS} \left\langle \left[\begin{array}{c} sign \\ \text{PHON} \langle Mann \rangle \end{array} \right] \bullet \top \right\rangle \end{array} \right]$$

Während die linke Merkmalstruktur in (35) alle Phrasen beschreibt, die mit einer männlichen Person beginnen, beschreibt die rechte nur solche Phrasen, die zusätzlich semantisch durch die Struktur *love* beschrieben werden.

Auch ein “Zusammenfallen” von Unterstrukturen bedeutet ein Zuwachs von Information:

$$(36) \quad \langle \top, \top \rangle \sqsubseteq \langle \sqcup \top, \sqcup \rangle$$

Natürlich gibt es auch Merkmalstrukturen, die sich nicht subsumieren:

$$(37) \quad a. \quad \left[\begin{array}{c} phrase \\ \text{PHON} \langle Mann \rangle \end{array} \right] \quad \text{und} \quad \left[\begin{array}{c} phrase \\ \text{CONT} love \end{array} \right]$$

$$b. \quad [Mann] \quad \text{und} \quad [Frau]$$

Die Subsumtionshierarchie der Merkmalstrukturen ist damit, wie die Subsumtionshierarchie von Typen (Typenhierarchie), eine partielle Ordnung mit einem größten Element: größtes Element ist die am wenigsten strukturierte Merkmalstruktur mit dem allgemeinsten Type, also wiederum $\top$, bzw. $[\top]$.

Formale Definition der Subsumtion von Merkmalstrukturen

Formal läßt sich die Subsumtionsbeziehung von Merkmalstrukturen über einen *Morphismus* definieren: Alle Knoten der subsumierenden Struktur müssen auf Knoten der subsumierten Struktur abgebildet werden, die gleich oder kleiner (spezieller) getypt sind. Die Kanten der subsumierenden Struktur müssen sich in der subsumierten Struktur wiederfinden.

Definition 1.8 (Subsumtion von Merkmalstrukturen²¹)

$F = \langle Q, \bar{q}, \theta, \delta \rangle$ subsumiert $F' = \langle Q', \bar{q}', \theta', \delta' \rangle$, geschrieben $F \sqsupseteq F'$, genau dann, wenn es eine totale Funktion $h : Q \longrightarrow Q'$, genannt Morphismus gibt, mit:

1. $h(\bar{q}) = \bar{q}'$
2. $\theta(q) \sqsupseteq \theta'(h(q))$ für alle $q \in Q$
3. $h(\delta(f, q)) = \delta'(f, h(q))$ für alle $q \in Q$ und alle Merkmale $f \in \mathbf{Feat}$,
für die $\delta(f, q)$ definiert ist

Unifikation von Merkmalstrukturen

So, wie sich die Subsumtionsbeziehung von Merkmalstrukturen aus der Subsumtionsbeziehung von Typen ableiten läßt, kann auch die Unifikation von Merkmalstrukturen aus der Unifikation von Typen entwickelt werden.

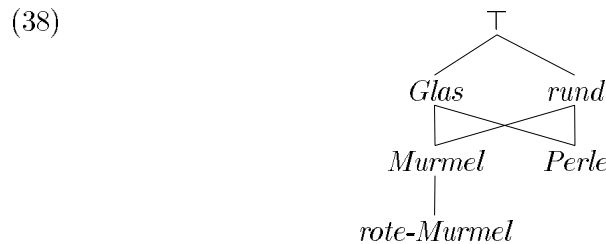
Die Unifikation zweier Typen ist die Menge ihrer *maximalen unteren Schranken* hinsichtlich der Typenhierarchie (vgl. Def. 1.3 (Unifikation von Typen)) Entsprechend ist die Unifikation zweier Merkmalstrukturen definiert als die

²¹ vgl. Carpenter 1992, Def. 3.4 (Subsumtion), S. 41.

Menge ihrer *maximalen unteren Schranken* hinsichtlich der Subsumtionshierarchie von Merkmalstrukturen.

Aus der informationstheoretischen Perspektive kann gesagt werden, daß jedes Element des Unifikates zweier Typen bzw. Merkmalstrukturen σ und τ für eine Möglichkeit steht, die Information beider Typen bzw. Merkmalstrukturen zu kombinieren, ohne weitere Information hinzuzufügen. Das Unifikat insgesamt bezeichnet alle Möglichkeiten dieser Kombination.

Betrachten wir zuerst ein Beispiel für die Unifikation von Typen:



Die “Welt”, die durch diese Typenhierarchie beschrieben wird, besteht aus Dingen aus Glas, runden Dingen, Murmeln, roten Murmeln und (Glas)perlen. Runde Dinge aus Glas können in diesem Universum nur Murmeln oder Perlen sein: Beide stehen hinsichtlich dieses speziellen Universums für eine allgemeinste Möglichkeit, die Informationen *Glas* und *rund* zu vereinen. Der Typ *rote-Murmel* hingegen ist nicht Element des Unifikates, da er zu den Informationen *Glas* und *rund* zusätzlich die Information *rot* hinzufügt.

Die Unifikation der Merkmalstrukturen ergibt sich entsprechend:

1. Ausweitung der Typen-Unifikation auf Merkmalstrukturen

Die Typenunifikation findet sich in der Unifikation von Merkmalstrukturen wieder. Analog zu dem folgenden Beispiel:

$$(39) \quad \text{num} \sqcup 1 = \{1sg, 1pl\}^{22}$$

ergibt sich:

$$(40) \quad \left[\text{num} \right] \sqcup \left[1 \right] = \left\{ \left[1sg \right], \left[1pl \right] \right\}$$

Übertragen auf ein komplexeres Beispiel:

²² vgl. die Typenhierarchie (8) bzw. (15).

$$\begin{aligned}
(41) \quad & \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \begin{bmatrix} sign \\ \text{AGR } num \end{bmatrix} \bullet \top \right\rangle \end{array} \right] \sqcup \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \begin{bmatrix} sign \\ \text{AGR } 1 \end{bmatrix} \bullet \top \right\rangle \end{array} \right] \\
&= \left\{ \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \begin{bmatrix} sign \\ \text{AGR } 1sg \end{bmatrix} \bullet \top \right\rangle \right], \left[\begin{array}{c} phrase \\ \text{DTRS} \left\langle \begin{bmatrix} sign \\ \text{AGR } 1pl \end{bmatrix} \bullet \top \right\rangle \right] \right\}
\end{aligned}$$

1. Unifikation durch Kombination der “strukturellen Information”

Jedes Pfad-Merkmal-Paar aus einer der zu unifizierenden Merkmalstrukturen findet sich im Unifikat wieder:

$$(42) \quad \left[\begin{array}{c} sign \\ \text{AGR } num \\ \text{CAT } v \end{array} \right] \sqcup \left[\begin{array}{c} sign \\ \text{AGR } 1 \\ \text{CONT } sleep \end{array} \right] = \left\{ \left[\begin{array}{c} sign \\ \text{AGR } 1sg \\ \text{CAT } v \\ \text{CONT } sleep \end{array} \right], \left[\begin{array}{c} sign \\ \text{AGR } 1pf \\ \text{CAT } v \\ \text{CONT } sleep \end{array} \right] \right\}$$

Koindizierte Pfade, d.h. Pfade mit identischem Wert, finden sich natürlich auch im Unifikat wieder:

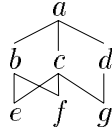
$$(43) \quad \langle \top, \top \rangle \sqcup \langle \sqcup \top, \sqcup \rangle = \left\{ \langle \sqcup \top, \sqcup \rangle \right\}$$

Anschaulich läßt sich die Unifikation von Merkmalstrukturen als *Projektion* der zu unifizierenden Strukturen *aufeinander* beschreiben:

1. Das Unifikat enthält alle Pfade, die in den Ausgangsstrukturen vorkommen;
2. Typen, die vom Wurzelknoten aus gesehen am Ende des gleichen Pfades liegen, und — im Sinne dieses Bildes gesprochen — aufeinander projiziert werden, werden durch ein Element ihrer Typen-Unifikation ersetzt.

Die verschiedenen Fälle, die bei einer Unifikation auftreten können, lassen sich am besten anhand von Beispielen veranschaulichen. Wir setzen die folgende Typenhierarchie voraus:

(44)



Für die Typenunifikation ergeben sich unter anderem folgende Fälle:

- (45) a. $a \sqcup a = \{a\}$
 b. $a \sqcup b = \{b\}$
 c. $b \sqcup d = \emptyset$
 d. $c \sqcup d = \{g\}$
 e. $b \sqcup c = \{e, f\}$

Mit den Merkmalen **Feat** = {1, 2} können unter anderem die folgenden Beispiele der Unifikation gebildet werden:

Unifikation ohne Veränderung der Struktur:

- (46) a. $\begin{bmatrix} a \\ 1 \ a \end{bmatrix} \sqcup \begin{bmatrix} a \\ 1 \ b \end{bmatrix} = \left\{ \begin{bmatrix} a \\ 1 \ b \end{bmatrix} \right\}$
 b. $\begin{bmatrix} a \\ 1 \ b \end{bmatrix} \sqcup \begin{bmatrix} a \\ 1 \ d \end{bmatrix} = \emptyset$
 c. $\begin{bmatrix} c \\ 1 \ b \end{bmatrix} \sqcup \begin{bmatrix} d \\ 1 \ c \end{bmatrix} = \left\{ \begin{bmatrix} g \\ 1 \ e \end{bmatrix}, \begin{bmatrix} g \\ 1 \ f \end{bmatrix} \right\}$
 d. $\begin{bmatrix} b \\ 1 \ b \end{bmatrix} \sqcup \begin{bmatrix} c \\ 1 \ c \end{bmatrix} = \left\{ \begin{bmatrix} e \\ 1 \ e \end{bmatrix}, \begin{bmatrix} e \\ 1 \ f \end{bmatrix}, \begin{bmatrix} f \\ 1 \ e \end{bmatrix}, \begin{bmatrix} f \\ 1 \ f \end{bmatrix} \right\}$

Kombination "struktureller Information":

- (47) a. $\begin{bmatrix} a \\ 1 \ a \end{bmatrix} \sqcup \begin{bmatrix} a \\ 2 \ a \end{bmatrix} = \left\{ \begin{bmatrix} a \\ 1 \ a \\ 2 \ a \end{bmatrix} \right\}$
 b. $\begin{bmatrix} a \\ 1 \ b \\ 2 \ e \end{bmatrix} \sqcup \begin{bmatrix} f \\ 2 \ \begin{bmatrix} a \\ 1 \ c \end{bmatrix} \\ 3 \ a \end{bmatrix} = \left\{ \begin{bmatrix} f \\ 1 \ b \\ 2 \ \begin{bmatrix} e \\ 1 \ c \end{bmatrix} \\ 3 \ a \end{bmatrix} \right\}$

Koindizierung von Merkmalwerten:

$$(48) \quad a. \quad \begin{bmatrix} b \\ 1 \ c \\ 2 \ a \end{bmatrix} \sqcup \begin{bmatrix} c \\ 1 \ \sqcup \ d \\ 2 \ \sqcup \end{bmatrix} = \left\{ \begin{bmatrix} e \\ 1 \ \sqcup \ g \\ 2 \ \sqcup \end{bmatrix}, \begin{bmatrix} f \\ 1 \ \sqcup \ g \\ 2 \ \sqcup \end{bmatrix} \right\}$$

$$b. \quad \begin{bmatrix} b \\ 1 \ c \\ 2 \ b \end{bmatrix} \sqcup \begin{bmatrix} c \\ 1 \ \sqcup \ d \\ 2 \ \sqcup \end{bmatrix} = \emptyset$$

Bei der Unifikation können aus *nicht-zyklischen* Merkmalstrukturen *zyklische*²³ entstehen:

$$(50) \quad \begin{bmatrix} a \\ 1 \ \sqcup \\ 2 \ \sqcup \end{bmatrix} \sqcup \begin{bmatrix} a \\ 1 \ \begin{bmatrix} a \\ 1 \ \sqcup \end{bmatrix} \\ 2 \ \sqcup \end{bmatrix} = \left\{ \begin{bmatrix} a \\ 1 \ \sqcup \ \begin{bmatrix} a \\ 1 \ \sqcup \end{bmatrix} \\ 2 \ \sqcup \end{bmatrix} \right\}$$

Operationale Beschreibung der Unifikation

Die Unifikation läßt sich auch konstruktiv beschreiben:

Von zwei Strukturen F_1 und F_2 ausgehend wird eine dritte Struktur F_u — Element des Unifikates von F_1 und F_2 — rekursiv mittels der folgenden Schritte konstruiert:

1. Als Typ des Wurzelknotens von F_u wird ein beliebiges Element aus der Unifikation der Typen der Wurzelknoten von F_1 und F_2 ausgewählt;
2. Die Menge der Merkmale, für die F_u am Wurzelknoten definiert ist, wird festgelegt als die Vereinigung der Merkmale, für die F_1 und F_2 am Wurzelknoten definiert sind;
3. Die Werte der Merkmale von F_u am Wurzelknoten ergeben sich durch Auswahl eines Elementes aus dem Unifikat der entsprechenden Merkmalwerte von F_1 und F_2 .

Diese drei Schritte werden rekursiv die Strukturen traversierend so lange wiederholt, bis das Ergebnis sich nicht mehr ändert.

²³ Beispiel einer einfachen *zyklischen Merkmalstruktur*:

$$(49) \quad \sqcup \begin{bmatrix} a \\ 1 \ \sqcup \end{bmatrix}$$

Definition der Unifikation

Bei der formalen Definition der Unifikation lehnen wir uns an Carpenter 1992 an. Im Gegensatz zur Definition in Carpenter 1992 ist die Unifikation von Merkmalstrukturen allerdings nicht als Funktion in die Menge der Merkmalstrukturen definiert, sondern als Funktion in die Potenzmenge der Merkmalstrukturen.

Definition 1.9 (Unifikation²⁴)

Gegeben seien zwei Merkmalstrukturen $F_1, F_2 \in \mathbf{MS}$. Mit $\langle Q_1, \bar{q}_1, \theta_1, \delta_1 \rangle \sim F_1$ und $\langle Q_2, \bar{q}_2, \theta_2, \delta_2 \rangle \sim F_2$ wählen wir zu diesen alphabetische Varianten, derart, daß $Q_1 \cap Q_2 = \emptyset$. Die Äquivalenz-Relation²⁵ $\bowtie$ auf $Q_1 \cup Q_2$ definieren wir als die kleinste Relation, die die folgenden Bedingungen erfüllt:

1. $\bar{q}_1 \bowtie \bar{q}_2$
2. $\delta(f, q_1) \bowtie \delta(f, q_2)$ falls $q_1 \bowtie q_2$
und $\delta(f, q_1)$ und $\delta(f, q_2)$ definiert

Die Unifikation $F_1 \sqcup F_2$ von F_1 und F_2 ist definiert durch:

$$F_1 \sqcup F_2 = \{ \langle (Q_1 \cup Q_2) / \bowtie, [\bar{q}_1]_{\bowtie}, \theta^\bowtie, \delta^\bowtie \rangle \mid \theta^\bowtie([q]_{\bowtie}) \in \sqcup \{ (\theta_1 \cup \theta_2)(q') \mid q' \bowtie q \} \}$$

wobei:

$$\delta^\bowtie(f, [q]_{\bowtie}) = \begin{cases} [(\delta_1 \cup \delta_2)(f, q)]_{\bowtie} & \text{falls } (\delta_1 \cup \delta_2)(f, q) \text{ definiert} \\ \text{undefiniert} & \text{ansonsten} \end{cases}$$

.

²⁴ vgl. Carpenter 1992, Def. 3.11 (Unification)

²⁵ Für eine Äquivalenz-Relation (also eine reflexive, symmetrische und transitive Relation) $\equiv \subseteq M \times M$ ist die Äquivalenzklasse von $x \in M$ folgendermaßen definiert:

$$(51) \quad [x]_{\equiv} = \{y \in M \mid y \equiv x\}$$

Die Quotientenmenge $M/\equiv$ von M modulo $\equiv$ ist die Menge der Äquivalenzklassen, in die M durch $\equiv$ zerfällt:

$$(52) \quad M/\equiv = \{[x]_{\equiv} \mid x \in M\}$$

1.4 Constraints

Merkmalslogiken sind Formalismen zur Modellierung von Gegenstandsbereichen. Die Objekte eines Gegenstandsbereiches werden im Modell durch Merkmalstrukturen repräsentiert. Ist der Gegenstandsbereich, wie in unserem Falle, ein Ausschnitt der japanischen Sprache, repräsentieren Merkmalstrukturen also Wörter, Phrasen, Sätze, semantische Strukturen etc.

Die Wahl der Merkmale und die Definition der Typenhierarchie bilden den ersten Teil der Modellierung. Durch die Menge der Typen und Merkmale wird, wie in Def. 1.4 (Merkmalstruktur) dargestellt, eine unendliche Menge von Merkmalstrukturen definiert. Die Form der möglichen Merkmalstrukturen über einer Menge von Typen und einer Menge von Merkmalen wird bisher jedoch allein durch die Kombinatorik bestimmt. Nur wenige der kombinatorisch möglichen Strukturen sind jedoch sinnvolle Beschreibungen von Objekten.

Gibt es beispielsweise die Merkmale PHON und CAT und die Typen *Elizabeth*, *Sigrid*, *np* und *lex-entry* lassen sich Strukturen bilden wie:

$$(53) \quad a. \quad \left[\begin{array}{c} Elizabeth \\ \text{PHON} \quad \boxed{} \\ \text{CAT} \quad \boxed{} \end{array} \right] \quad b. \quad \left[\begin{array}{c} Elizabeth \\ \text{PHON} \quad \left[\begin{array}{c} np \\ \text{CAT} \quad \left[\begin{array}{c} Sigrid \\ \text{PHON} \quad lex-entry \end{array} \right] \end{array} \right] \end{array} \right]$$

Diese allerdings ergeben wenig Sinn — im Gegensatz zu den folgenden:

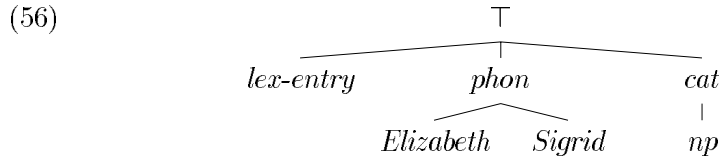
$$(54) \quad a. \quad \left[\begin{array}{c} lex-entry \\ \text{PHON} \quad Elizabeth \\ \text{CAT} \quad np \end{array} \right] \quad b. \quad \left[\begin{array}{c} lex-entry \\ \text{PHON} \quad Sigrid \\ \text{CAT} \quad np \end{array} \right]$$

Es ist also notwendig, sinnvolle von nicht sinnvollen Merkmalstrukturen durch ein formales Mittel unterscheiden zu können. Diesem Zweck dienen formale Bedingungen, die sogenannten *Constraints*: Nur solche Strukturen sind sinnvoll oder — um diesen inhaltlichen Begriff durch einen formalen zu ersetzen — *aufgelöst* (*resolved*), die alle Constraints des Systems erfüllen.

In unserem Fall wäre beispielsweise der folgende Constraint denkbar:

$$(55) \quad lex-entry \Rightarrow \left[\begin{array}{c} lex-entry \\ \text{PHON} \quad phon \\ \text{CAT} \quad cat \end{array} \right]$$

Ein Zeichen vom Typ *lex-entry* hat zwei Merkmale: PHON und CAT. Das Merkmal PHON kann Subtypen vom Typ *phon* als Werte annehmen; Das Merkmal CAT Subtypen vom Typ *cat*. Eine sinnvolle Typenhierarchie für diese Beschränkung wäre:



Aus diesem Beispiel läßt sich erahnen, wie der Modellbezug zwischen Gegenstandsbereich und Constraintsystem hergestellt wird: Die Modellierung eines Gegenstandsbereiches durch ein Constraintsystem besteht darin, Typen, Merkmale und Constraints so zu wählen, daß jedes Objekt des Gegenstandsbereiches durch eine aufgelöste Merkmalstruktur beschrieben wird und gleichzeitig jede aufgelöste Struktur einem Objekt des Gegenstandsbereiches entspricht.

Soll — wie in dieser Arbeit — ein Fragment einer natürlichen Sprache durch ein Constraintsystem beschrieben werden, muß also

1. jeder korrekte sprachliche Ausdruck durch eine aufgelöste Merkmalstruktur beschrieben werden,²⁶ und
2. jeder aufgelösten Merkmalstruktur ein korrekter sprachlicher Ausdruck entsprechen.²⁶

Die Rolle der Unifikation für Constraintsysteme läßt sich ebenfalls ableiten: Constraints, die in unserem System ebenfalls die Form von Merkmalstrukturen haben, formulieren *partielle* Anforderungen an Merkmalstrukturen. Diese partiellen Anforderungen oder Informationen werden mit dem Mittel der *Unifikation* kombiniert.

Typen als “Angriffspunkte” der Constraints

Anhand einiger einfacher Beispiele wollen wir uns die Funktionsweise von Constraints veranschaulichen:

²⁶ Wir abstrahieren hier von ambigen sprachlichen Ausdrücken, beispielsweise Homophonen, strukturell mehrdeutigen Sätzen etc., und setzen die Vereindeutigung sprachlicher Ausdrücke voraus, z.B. durch Indizierung, geeignete Klammerung etc.

Adäquatheit von Merkmalen — Constraints bestimmen, welche Merkmale in einer Klammerenebene adäquat sind: Die Unterstrukturen des Typs *sign* in der Struktur (1) beispielsweise enthalten immer die Merkmale PHON und CAT. Ein Constraint, der dieses sicherstellt, könnte z.B. folgendermaßen aussehen:

$$(57) \quad \textit{sign} \Rightarrow \begin{bmatrix} \textit{sign} \\ \text{PHON} \quad \top \\ \text{CAT} \quad \top \end{bmatrix}$$

Adäquatheit von Merkmalwerten — Ist ein Merkmal für einen Typ adäquat, so bestimmen Constraints, welche Werte dieses Merkmal annehmen darf. Der Constraint (57) beispielsweise läßt sich leicht in dieser Hinsicht erweitern:

$$(58) \quad \textit{sign} \Rightarrow \begin{bmatrix} \textit{sign} \\ \text{PHON} \quad \textit{list} \\ \text{CAT} \quad \textit{cat} \end{bmatrix}$$

(58) bestimmt, daß das Merkmal PHON einen Wert des Typs *list* und das Merkmal CAT einen Wert des Typs *cat* annehmen muß.

Abhängigkeiten unter Merkmalen — Zusammenhänge zwischen Merkmalwerten von Zeichen können ebenfalls durch Constraints sichergestellt werden:

$$(59) \quad \textit{phrase} \Rightarrow \begin{bmatrix} \textit{phrase} \\ \text{PHON} \quad \boxed{1} \circ \boxed{2} \\ \text{DTRS} \quad \left\langle \left[\text{PHON} \quad \boxed{1} \right], \left[\text{PHON} \quad \boxed{2} \right] \right\rangle \end{bmatrix}$$

(59) erzwingt, daß der PHON-Wert von Merkmalstrukturen des Typs *phrase* der Konkatenation der PHON-Werte seiner Töchter entspricht.²⁷

²⁷ Der Ausdruck “ $\boxed{1} \circ \boxed{2}$ ” beschreibt die Konkatenation der Listen “ $\boxed{1}$ ” und “ $\boxed{2}$ ” und wird ebenfalls durch Constraints realisiert (vgl. Abschnitt “Listen und Listen-Konkatenation”, S. 47).

Lexikalische Einträge als Constraints — Auch das Lexikon wird durch Constraints beschrieben. Lexikalische Einträge für

- (60) a. “Leonard” (np)
 b. “schläft” (vp)
 c. “trinkt” (vp)

lassen sich z.B. mit den folgenden Constraints beschreiben (∨ steht für die logische Disjunktion):

$$(61) \quad lex-entry \Rightarrow \left[\begin{array}{c} lex-entry \\ PHON \langle Leonard \rangle \\ CAT \quad np \end{array} \right] \vee \left[\begin{array}{c} lex-entry \\ PHON \langle schläft \rangle \\ CAT \quad vp \end{array} \right] \vee \left[\begin{array}{c} lex-entry \\ PHON \langle trinkt \rangle \\ CAT \quad vp \end{array} \right].$$

Phrasenstrukturregeln als Constraints — Phrasenstrukturregeln beschreiben, wie Wörter und Phrasen zu Phrasen und Sätzen kombiniert werden können. Die Regel

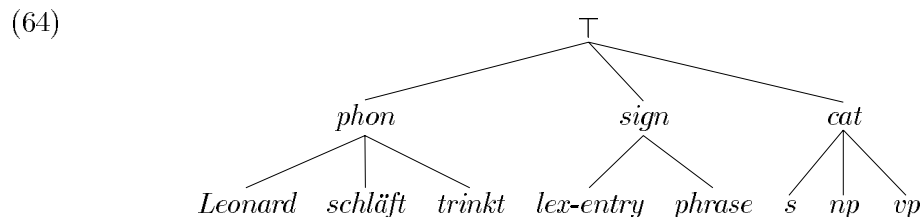
$$(62) \quad S \longrightarrow NP \ VP$$

beispielsweise kombiniert eine Nominalphrase mit einer Verbalphrase zu einem Satz. Als Constraint läßt sie sich folgendermaßen schreiben:

$$(63) \quad phrase \Rightarrow \left[\begin{array}{c} phrase \\ CAT \quad s \\ DTRS \left\langle \left[\begin{array}{c} sign \\ CAT \quad np \end{array} \right], \left[\begin{array}{c} sign \\ CAT \quad vp \end{array} \right] \right\rangle \end{array} \right].$$

Eine kleine Grammatik

Wenn wir zu den Constraints (58), (59), (61) und (63) die Typenhierarchie (64) annehmen, erhalten wir eine erste kleine Grammatik:



Wir wollen uns ansehen, wie das folgende Zeichen erweitert werden muß, damit es alle Constraints erfüllt:

$$(65) \quad \left[phrase \right]$$

Ein Constraint $\sigma \Rightarrow F$ muß auf alle Strukturen angewandt werden, die durch σ subsumiert werden. Da der Typ *phrase* ein Subtyp von *sign* ist, muß der Constraint (58) also auch auf (65) angewandt werden. Durch Unifikation der rechten Seite von (58) mit (65) erhalten wir:

$$(66) \quad \left[\begin{array}{ll} phrase \\ PHON & list \\ CAT & cat \end{array} \right]$$

Der Constraint (63) muß ebenfalls angewandt werden:

$$(67) \quad \left[\begin{array}{ll} phrase \\ PHON & list \\ CAT & s \\ DTRS & \left\langle \left[\begin{array}{ll} sign \\ CAT & np \end{array} \right], \left[\begin{array}{ll} sign \\ CAT & vp \end{array} \right] \right\rangle \end{array} \right]$$

Durch Anwendung von (58) auf die Töchter von (67) ergibt sich:

$$(68) \quad \left[\begin{array}{ll} phrase \\ PHON & list \\ CAT & s \\ DTRS & \left\langle \left[\begin{array}{ll} sign \\ PHON & list \\ CAT & np \end{array} \right], \left[\begin{array}{ll} sign \\ PHON & list \\ CAT & vp \end{array} \right] \right\rangle \end{array} \right]$$

Zeichen des typs *sign* subsumieren auch Zeichen des Typs *lex-entry* und *phrase*. Da die Töchter von (68) jedoch die Kategorie *np* bzw. *vp* haben müssen, läßt sich (63) nicht anwenden. (61) hingegen kann benutzt werden:

Die erste Tochter unifiziert mit

$$(69) \quad \left[\begin{array}{ll} lex-entry \\ PHON & \langle Leonard \rangle \\ CAT & np \end{array} \right]$$

die zweite Tochter wahlweise mit

$$(70) \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{schläft} \rangle \\ \text{CAT} \quad \textit{vp} \end{bmatrix} \quad \text{oder} \quad (71) \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{trinkt} \rangle \\ \text{CAT} \quad \textit{vp} \end{bmatrix}$$

Wir erhalten damit zwei mögliche Fortsetzungen:

$$(72) \quad \begin{array}{l} \text{a.} \\ \text{b.} \end{array} \begin{bmatrix} \textit{phrase} \\ \text{PHON} \quad \textit{list} \\ \text{CAT} \quad \textit{s} \\ \text{DTRS} \left\langle \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{Leonard} \rangle \\ \text{CAT} \quad \textit{np} \end{bmatrix}, \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{schläft} \rangle \\ \text{CAT} \quad \textit{vp} \end{bmatrix} \right\rangle \\ \text{DTRS} \left\langle \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{Leonard} \rangle \\ \text{CAT} \quad \textit{np} \end{bmatrix}, \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \langle \textit{trinkt} \rangle \\ \text{CAT} \quad \textit{vp} \end{bmatrix} \right\rangle \end{bmatrix}$$

Durch den Constraint (59) schließlich ergeben sich die folgenden Strukturen:

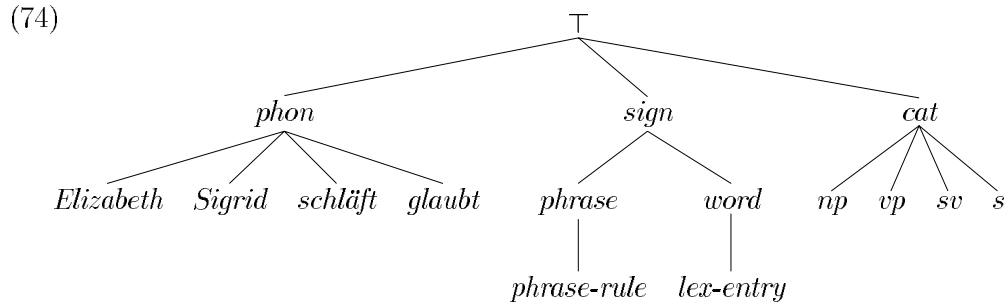
$$(73) \quad \begin{array}{l} \text{a.} \\ \text{b.} \end{array} \begin{bmatrix} \textit{phrase} \\ \text{PHON} \quad \boxed{1} \langle \textit{Leonard} \rangle \circ \boxed{2} \langle \textit{schläft} \rangle \\ \text{CAT} \quad \textit{s} \\ \text{DTRS} \left\langle \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \quad \boxed{1} \\ \text{CAT} \quad \textit{np} \end{bmatrix}, \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \quad \boxed{2} \\ \text{CAT} \quad \textit{vp} \end{bmatrix} \right\rangle \\ \text{DTRS} \left\langle \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \quad \boxed{1} \\ \text{CAT} \quad \textit{np} \end{bmatrix}, \begin{bmatrix} \textit{lex-entry} \\ \text{PHON} \quad \boxed{2} \\ \text{CAT} \quad \textit{vp} \end{bmatrix} \right\rangle \end{bmatrix}$$

(73a) und (73b) sind die einzigen Merkmalstrukturen, des Typs *phrase*, die alle Constraints erfüllen. Die Menge dieser beiden Strukturen ist damit die Sprache, die durch die Grammatik aus der Typenhierarchie (64) und den Constraints (58), (59), (61) und (63) hinsichtlich des Typs *phrase* beschrieben wird.

Constraintvererbung

Constraints zu einem Typ σ betreffen auch alle Subtypen von σ . Dieser Sachverhalt wird als *Constraintvererbung* bezeichnet.

Zur Veranschaulichung der Constraintvererbung erweitern wir das letzte Beispiel um einige Verben, Eigennamen und Phrasenstrukturregeln. Die folgende Typenhierarchie wird zugrunde gelegt:



Anstatt der graphischen Darstellung in (74) wird von nun an wahlweise die folgende textuelle Darstellung verwendet:

- (75)
- | | | |
|---------------|---------------|--|
| <i>phon</i> | $\Rightarrow$ | <i>Elizabeth</i> $\vee$ <i>Sigrid</i> $\vee$ <i>schläft</i> $\vee$ <i>glaubt</i> . |
| <i>cat</i> | $\Rightarrow$ | <i>np</i> $\vee$ <i>vp</i> $\vee$ <i>sv</i> $\vee$ <i>s</i> . |
| <i>sign</i> | $\Rightarrow$ | <i>phrase</i> $\vee$ <i>word</i> . |
| <i>phrase</i> | $\Rightarrow$ | <i>phrase-rule</i> . |
| <i>word</i> | $\Rightarrow$ | <i>lex-entry</i> . |

Jede Zeile stellt einen Typ zusammen mit seinen direkten Untertypen dar. Typen, die nicht auf der rechten Seite einer anderen Typendefinition vorkommen, sind direkte Subtypen von $\top$. Erscheint ein Typ auf keiner linken Seite einer Typendefinition, handelt es sich um einen atomaren Typ.

Constraints wurden bisher immer in der Form

$$(76) \quad \sigma \Rightarrow \begin{bmatrix} \sigma \\ \dots \end{bmatrix} \vee \begin{bmatrix} \sigma \\ \dots \end{bmatrix} \vee \dots$$

dargestellt. Da sich die linke Seite jedoch aus dem Typ der äußeren Strukturen auf der rechten Seite ergibt, können wir einfacher schreiben:

$$(77) \begin{bmatrix} \sigma \\ \dots \end{bmatrix} \begin{bmatrix} \sigma \\ \dots \end{bmatrix} \dots$$

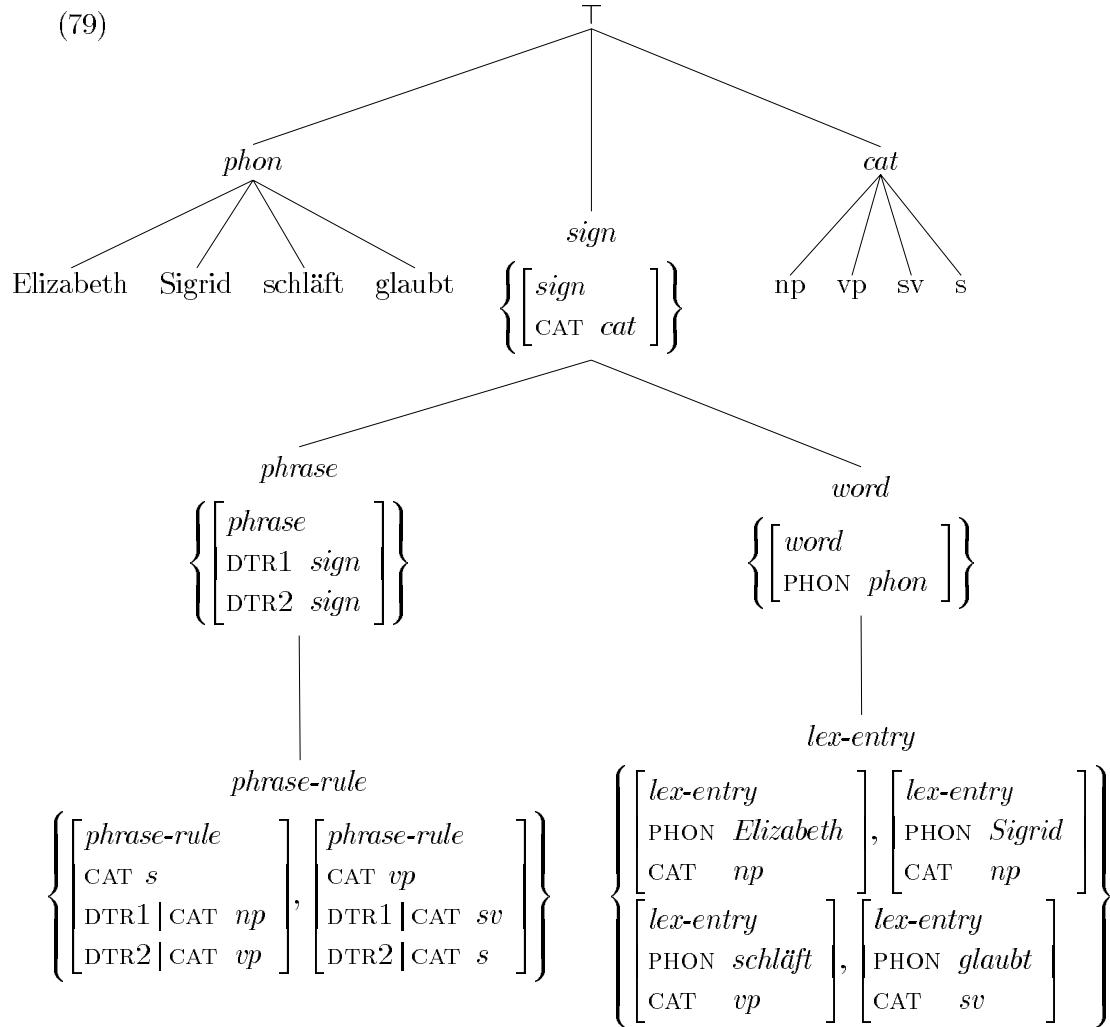
Mehrere Constraints des gleichen Typs sind als Alternativen zu verstehen.

Die Constraints für das neue Beispiel können nun folgendermaßen dargestellt werden:

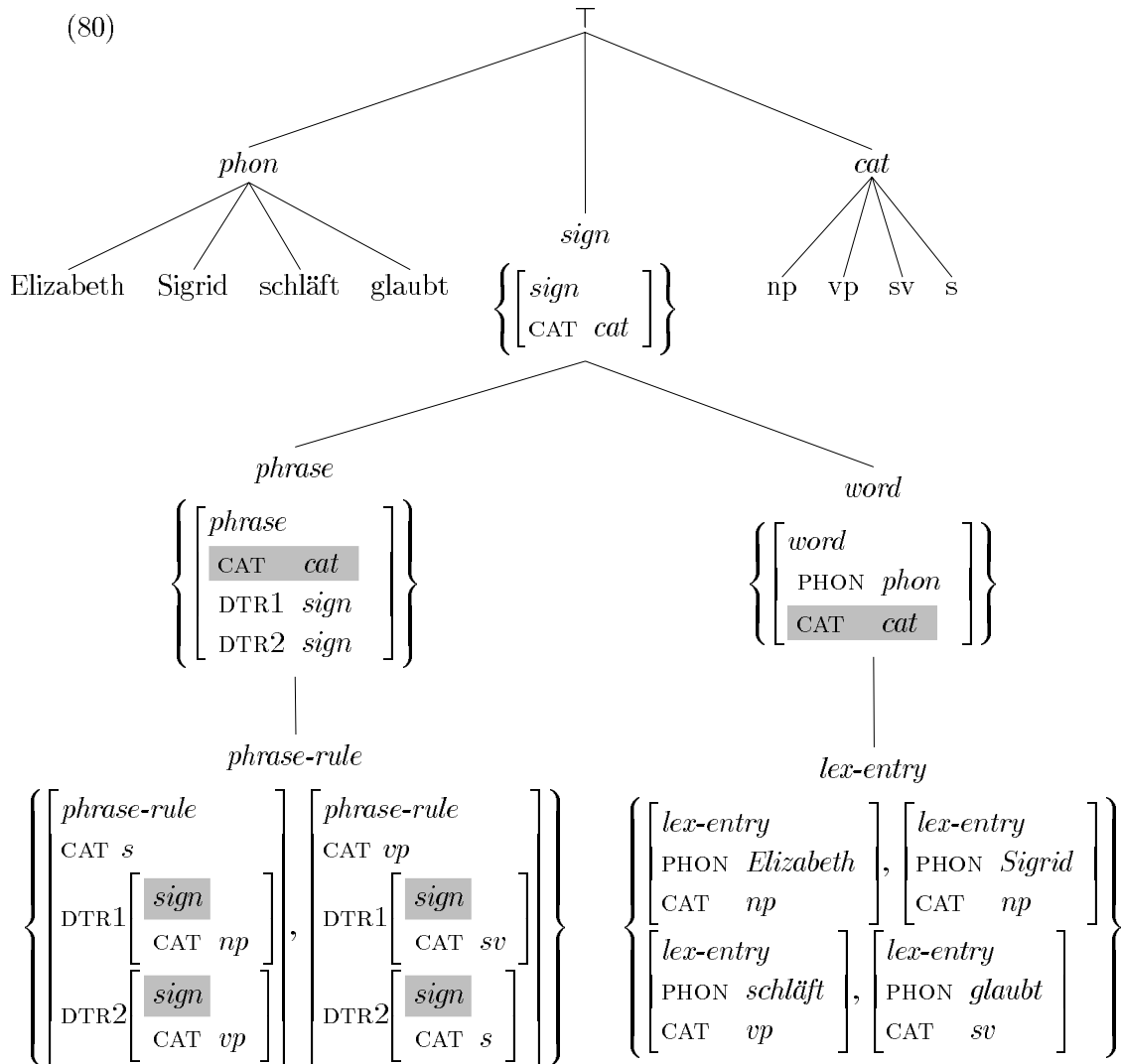
$$(78) \begin{bmatrix} sign \\ CAT \ cat \end{bmatrix} \\ \begin{bmatrix} phrase \\ DTR1 \ sign \\ DTR2 \ sign \end{bmatrix} \\ \begin{bmatrix} phrase-rule \\ CAT \ s \\ DTR1 \mid CAT \ np \\ DTR2 \mid CAT \ vp \end{bmatrix} \quad \begin{bmatrix} phrase-rule \\ CAT \ vp \\ DTR1 \mid CAT \ sv \\ DTR2 \mid CAT \ s \end{bmatrix} \\ \begin{bmatrix} word \\ PHON \ phon \end{bmatrix} \\ \begin{bmatrix} lex-entry \\ PHON \ Elizabeth \\ CAT \ np \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \ Sigrid \\ CAT \ np \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \ schläft \\ CAT \ vp \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \ glaubt \\ CAT \ sv \end{bmatrix}$$

Zur Veranschaulichung der Constraintvererbung werden Typenhierarchie und Constraints in einer gemeinsamen Hierarchie dargestellt. Jedem Typ wird die Menge seiner Constraints zugeordnet. Die Elemente dieser Menge verhalten

sich zueinander disjunktiv:



Die Constraint-Vererbung kann nun graphisch veranschaulicht werden. Die dabei entstehende Hierarchie unterscheidet sich in unserem Fall nur geringfügig von der ursprünglichen Darstellung. Bedingungen, die bei der Vererbung neu hinzukommenden, werden **grau unterlegt** dargestellt:



In Hierarchien mit disjunktiven Constraints für nicht-minimale Typen “multiplizieren” sich die Constraints, wie das folgende Beispiel einer “freien Wortstellung” veranschaulicht:

Die ursprüngliche Hierarchie:

$$(81) \quad \left\{ \left[\begin{array}{l} \textit{phrase} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{DTRS } \langle [\text{PHON } \boxed{1}], [\text{PHON } \boxed{2}] \rangle \end{array} \right], \left[\begin{array}{l} \textit{phrase} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{DTRS } \langle [\text{PHON } \boxed{2}], [\text{PHON } \boxed{1}] \rangle \end{array} \right] \right\}$$

$$\quad \quad \quad \textit{phrase-rule}$$

$$\left\{ \left[\begin{array}{l} \textit{phrase-rule} \\ \text{CAT } s \\ \text{DTRS } \langle [\text{CAT } np], [\text{CAT } vp] \rangle \end{array} \right], \left[\begin{array}{l} \textit{phrase-rule} \\ \text{CAT } vp \\ \text{DTRS } \langle [\text{CAT } sv], [\text{CAT } s] \rangle \end{array} \right] \right\}$$

Die gleiche Hierarchie mit vererbten Constraints:

$$(82) \quad \left\{ \left[\begin{array}{l} \textit{phrase} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{DTRS } \left\langle \left[\text{PHON } \boxed{1} \right], \left[\text{PHON } \boxed{2} \right] \right\rangle \end{array} \right], \left[\begin{array}{l} \textit{phrase} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{DTRS } \left\langle \left[\text{PHON } \boxed{2} \right], \left[\text{PHON } \boxed{1} \right] \right\rangle \end{array} \right] \right\}$$

$$\quad \quad \quad \textit{phrase-rule}$$

$$\left\{ \left[\begin{array}{l} \textit{phrase-rule} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{CAT } s \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{PHON } \boxed{1} \\ \text{CAT } np \end{array} \right], \left[\begin{array}{l} \text{PHON } \boxed{2} \\ \text{CAT } vp \end{array} \right] \right\rangle \end{array} \right], \left[\begin{array}{l} \textit{phrase-rule} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{CAT } s \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{PHON } \boxed{2} \\ \text{CAT } np \end{array} \right], \left[\begin{array}{l} \text{PHON } \boxed{1} \\ \text{CAT } vp \end{array} \right] \right\rangle \end{array} \right] \right\}$$

$$\left\{ \left[\begin{array}{l} \textit{phrase-rule} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{CAT } vp \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{PHON } \boxed{1} \\ \text{CAT } sv \end{array} \right], \left[\begin{array}{l} \text{PHON } \boxed{2} \\ \text{CAT } s \end{array} \right] \right\rangle \end{array} \right], \left[\begin{array}{l} \textit{phrase-rule} \\ \text{PHON } \boxed{1} \circ \boxed{2} \\ \text{CAT } vp \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{PHON } \boxed{2} \\ \text{CAT } sv \end{array} \right], \left[\begin{array}{l} \text{PHON } \boxed{1} \\ \text{CAT } s \end{array} \right] \right\rangle \end{array} \right] \right\}$$

Formale Beschreibung von Constraints und Constraintvererbung

Formal kann das Verhältnis von Typen und zugehörigen Constraints durch eine Funktion **Cons** beschrieben werden:

Definition 1.10 (Constraint-Funktion)

Die *Constraint-Funktion* **Cons** ist eine Funktion, die jedem Typ $\sigma \in \text{Type}$ eine Menge von Constraints $\text{Cons}(\sigma) \subseteq \text{MS}$ zuordnet:

$$\text{Cons} : \text{Type} \longrightarrow 2^{\text{MS}}$$

Die Elemente aus $\text{Cons}(\sigma)$ für ein $\sigma \in \text{Type}$ sind als alternative Constraints zu interpretieren. Die Menge $\text{Cons}(\sigma)$ ließe sich also auch als Disjunktion ihrer Elemente darstellen. Eine leere Constraintmenge wird immer erfüllt.

Um die Definition der Constraintvererbung zu vereinfachen, definieren wir die Funktion $\text{Cons}^+ : \text{Type} \longrightarrow 2^{\text{MS}}$ folgendermaßen:

$$\forall_{\sigma \in \text{Type}} : \text{Cons}^+(\sigma) = \begin{cases} \{\top\} & \text{falls } \text{Cons}(\sigma) = \emptyset, \\ \text{Cons}(\sigma) & \text{ansonsten} \end{cases}$$

Mit *Super* sei die Funktion $\text{Super} : \text{Type} \longrightarrow 2^{\text{Type}}$ bezeichnet, die einen Typ $\sigma \in \text{Type}$ auf die Menge seiner unmittelbaren Supertypen abbildet:

$$\forall_{\sigma \in \text{Type}} : \text{Super}(\sigma) = \{\tau \in \text{Type} \mid \tau \sqsupset \sigma \wedge \neg \exists_{v \in \text{Type}} : \tau \sqsupset v \wedge v \sqsupset \sigma\}$$

Definition 1.11 (Constraintvererbung²⁸)

$\text{Cons}^* : \text{Type} \longrightarrow 2^{\text{MS}}$ bildet jeden Typ auf die Menge seiner ererbten Constraints ab. Für jeden Typ $\sigma \in \text{Type}$ mit $\text{Super}(\sigma) = \{\tau_1, \dots, \tau_k\}$ ist $\text{Cons}^*(\sigma)$ definiert durch:

$$\text{Cons}^*(\sigma) = \bigcup \{ \bigsqcup_{0 \leq i \leq k} v_i \mid \begin{array}{l} v_0 \in \text{Cons}^+(\sigma), \\ v_i \in \text{Cons}^*(\tau_i) \quad \text{für } 1 \leq i \leq k \end{array} \}$$

Für $\top$ folgt aus dieser Definition wegen $\text{Super}(\top) = \emptyset$:

$$\text{Cons}^*(\top) = \bigcup \{ \bigsqcup_{0 \leq i \leq 0} v_i \mid v_0 \in \text{Cons}^+(\top) \} = \text{Cons}^+(\top)$$

²⁸ Vgl. Carpenter 1992, Def. 15.2 (Inherited Constraint), S. 229.

Aufgrund der hier benutzten Definition von Typenhierarchien und den daraus folgenden Konsequenzen für die Unifikation als einer Funktion auf *Mengen* von Merkmalstrukturen, sieht die hier vorgestellte Definition der Constraintvererbung auf den ersten Blick schwieriger aus, als die von Carpenter vorgeschlagene. Wird sie jedoch anhand von Beispielen durchgespielt, zeigt sich, daß sie unkompliziert ist und direkt in einen einfachen Algorithmus umgesetzt werden kann.

Definition von Constraintsystemen

Die Mengen **Type** und **Feat**, eine Typenhierarchie $\mathcal{H}$ über **Type** und eine Constraint-Funktion **Cons** ergeben zusammen ein Constraintsystem. Da Grammatiken in dieser Arbeit als Constraintsysteme formalisiert werden, besitzen wir damit alle Bausteine unseres formalen Modells.

Definition 1.12 (Constraintsystem²⁹)

Ein Constraintsystem ist ein 4-Tupel $\langle \text{Type}, \text{Feat}, \mathcal{H}, \text{Cons} \rangle$, mit:

1. **Type** ist eine endliche Menge von Typen;
2. **Feat** ist eine endliche Menge von Merkmalen;
3. $\mathcal{H} = \langle \text{Type}, \sqsubseteq, \top \rangle$ ist eine Typenhierarchie über **Type**;
4. **Cons** ist eine Constraint-Funktion.

Lösung eines Constraintsystems

Am Anfang dieses Abschnitts (S. 37) wurde anhand eines einfachen Beispiels gezeigt, wie Typenhierarchie und Constraints zusammenwirken, um eine Menge aufgelöster Merkmalstrukturen zu definieren: Von der Struktur *[phrase]* ausgehend wurden die Strukturen (73a) und (73b) abgeleitet.

Neben der Forderung, daß alle Constraints erfüllt werden müssen, gibt es eine weitere Anforderung an die Lösung eines Constraintsystems: Alle Typen einer Lösung müssen minimal sein. Diese Forderung läßt sich in der Syntax von Typendeklarationen wiederfinden:

$$(83) \quad \sigma \Rightarrow \tau_1 \vee \tau_2 \vee \dots \vee \tau_n.$$

läßt sich als Constraint des Types σ lesen:

“Wenn ein Zeichen vom Typ σ ist, muß es entweder vom Typ τ_1 , vom Typ τ_2 , ... oder vom Typ τ_n sein.”

Definition 1.13 (Lösung eines Constraintsystems³⁰)

Eine Merkmalstruktur $F \in \text{MS}^{\text{Type}, \text{Feat}}$ ist eine Lösung eines Constraintsystems $C = \langle \text{Type}, \text{Feat}, \mathcal{H}, \text{Cons} \rangle$, wenn:

1. Alle Typen in F minimal sind;

²⁹ vgl. Carpenter 1992, Def. 15.1 (Constraint System), S. 228.

2. $F@π \sqsubseteq σ$ für ein $σ \in \mathbf{Cons}^*(\theta(\delta(π, \bar{q})))$ für alle Pfade, für die $F@π$ definiert ist;
3. es gibt keine Lösung F' , die 1 und 2 erfüllt, mit $F' \sqsupseteq F$.

Die Menge aller Lösungen von C bezeichnen wir mit $L(C)$.

Die letzte Bedingung stellt sicher, daß eine Lösung keine Information enthält, die nicht durch die ersten beiden Bedingungen erzwungen wurde.

1.5 Resolution

Wir wissen nun, was Constraintsysteme sind und wie durch ein solches System eine Menge von Merkmalstrukturen definiert wird. Um die mathematische Spezifikation von Constraintsystemen und deren Lösungsmenge für die Praxis nutzbar zu machen, muß ein syntaktisches Verfahren gefunden werden, das zu einer Anfrage in Form einer unterspezifizierten Merkmalstruktur die Menge der Lösungen aufzählt. Ein solches Verfahren wird *Resolutionsverfahren* genannt.

Wie ein Resolutionsverfahren aussehen kann, läßt sich aus der Definition der Lösungsmenge leicht ersehen:

Von einer beliebigen unterspezifizierten Merkmalstruktur F ausgehend, wird rekursiv, einer fairen Strategie (beispielsweise Breitensuche) folgend, einer der folgenden zwei Schritte angewandt:

1. ein allgemeinerer Typ wird durch einen spezielleren ersetzt;
2. eine Teilstruktur von F des Typs $σ$ wird mit einem Element von $\mathbf{Cons}^*(σ)$ unifiziert;

Diese beiden Schritte werden solange wiederholt, bis sich eine Lösung ergibt.³¹

Formal läßt sich dieses Resolutionsverfahren folgendermaßen definieren:

Definition 1.14 ($π$ -Resolution³²)

Als $π$ -Resolution bezeichnen wir die kleinste Relation $\xRightarrow{π} \subseteq \mathbf{MS} \times \mathbf{MS}$, die die folgende Bedingung erfüllt:

³⁰ vgl. Carpenter 1992, Def. 15.3 (Resolved Feature Structure), S. 229.

³¹ Enthält die Ausgangsstruktur schon "überflüssige" Information, oder werden auf eine Substruktur vom Typ $σ$ mehrere Elemente von $\mathbf{Cons}^*(σ)$ nacheinander angewandt, kann eine Lösung mehr Information als nötig enthalten. Auch in diesem Fall wird sie jedoch durch eine Lösung subsumiert.

$$F \xRightarrow{\pi} F' \quad \text{falls} \quad \begin{array}{l} 1. \ \pi \in \mathbf{Path} \text{ und } \delta(\pi, \bar{q}) \text{ in } F \text{ definiert;} \\ 2. \ \sigma = \theta(\delta(\pi, \bar{q})); \\ 3. \ \phi \in \mathbf{Cons}^*(\sigma) \quad \text{oder} \quad \phi \sqsubset \sigma; \\ 4. \ F' \in F \sqcup [\pi \ \phi]. \end{array}$$

Sei π also ein beliebiger Pfad in F und σ der Typ der Unterstruktur, die sich in F am Ende des Pfades π befindet, so entsteht $F \sqcup [\pi \ \phi]$ entweder, indem σ durch einen direkten Untertyp ersetzt wird, oder indem $F @ \pi$ durch ein Element der Unifikation von $F @ \pi$ mit einem Element von $\mathbf{Cons}^*(\sigma)$ ersetzt wird.

Als *Resolution* bezeichnen wir die wiederholte Anwendung der π -Resolution bis eine vollspezifizierte Merkmalstruktur (Lösung) entsteht:

Definition 1.15 (Resolution³³)

Als Resolution bezeichnen wir die wiederholte Anwendung der π -Resolution auf eine Merkmalstruktur (Anfrage) nach einer fairen Strategie (z.B. Breitensuche), bis alle Typen der entstandenen Struktur minimal sind und mit einem Element ihres geerbten Constraints unifiziert wurden.

1.6 Beispiele zur Anwendung von Constraintsystemen

In den letzten Abschnitten wurden die wesentlichen Punkte des Systems besprochen, das in dieser Arbeit zur Modellierung der japanischen Verbmodifikation verwendet wird. Das bisher Gesagte soll nun anhand einiger Beispiele, in denen die später verwendeten Techniken vorgestellt werden, veranschaulicht werden:

Berechnungen von Funktionen — Die Realisierung von Berechnungen durch Constraintsysteme wird anhand der *Listenkonkatenation* dargestellt.

Semantik — Die Modellierung einer einfachen *Semantik* für natürliche Sprachen und ihr Zusammenspiel mit der Syntax wird an einem weiteren kleinen Beispiel vorgestellt.

1.6.1 Listen und Listen-Konkatenation

Constraintsysteme lassen sich zur Berechnung von Funktionen benutzen. Berechnungsvorschriften (“Programme”) für Constraintsysteme haben große

³² vgl. Carpenter 1992, Def. 15.5 (π -Resolution), S. 231.

³³ vgl. Carpenter 1992, Def. 15.8 (Resolution), S. 232.

Ähnlichkeit mit logischen Programmen, beispielsweise in *Pure Prolog*.³⁴ Um eine Vorstellung zu vermitteln, wie Berechnungsvorschriften für Constraintsysteme ausgedrückt werden können, wird in diesem Abschnitt die Konkatenation von Listen beschrieben.

Listen als rekursive Merkmalstrukturen

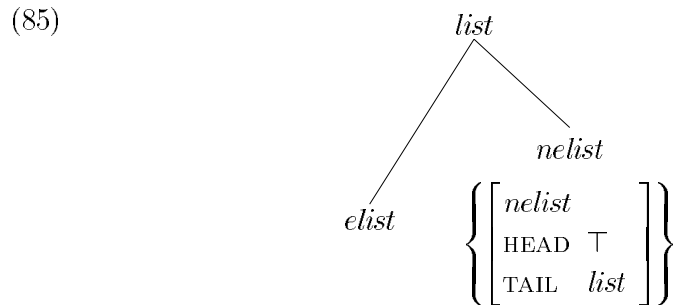
Listen sind rekursive Merkmalstrukturen³⁵. Sie genügen der folgenden Spezifikation:

- (84) **Typenhierarchie:**
 $list \Rightarrow elist \vee nelist.$

Constraints:

$$\left[\begin{array}{l} nelist \\ \text{HEAD } \top \\ \text{TAIL } list \end{array} \right]$$

Anschaulicher ist die graphische Darstellung:



Der leichten Lesbarkeit halber, wird diese explizite Form durch spitze Klammern abgekürzt:

³⁴ vgl. Sterling und Shapiro 1986 und Lloyd 1984.

³⁵ vgl. Abschnitt "Auflösung der Abkürzungen", S. 18.

(86) Abkürzung: abgekürzte Merkmalstruktur:

$$\begin{array}{ll}
 \langle \rangle & elist \\
 \langle a, b, c \rangle & \left[\begin{array}{l} nelist \\ \text{HEAD } a \\ \text{TAIL } \left[\begin{array}{l} nelist \\ \text{HEAD } b \\ \text{TAIL } \left[\begin{array}{l} nelist \\ \text{HEAD } c \\ \text{TAIL } elist \end{array} \right] \end{array} \right] \end{array} \right] \\
 \langle a, b \bullet list \rangle & \left[\begin{array}{l} nelist \\ \text{HEAD } a \\ \text{TAIL } \left[\begin{array}{l} nelist \\ \text{HEAD } b \\ \text{TAIL } list \end{array} \right] \end{array} \right]
 \end{array}$$

Rekursive Berechnung der Listen-Konkatenation

Rekursion läßt sich in Constraintsystemen durch rekursive Constraints darstellen. Bei der Konkatenation von Listen $\circ : Listen \times Listen \longrightarrow Listen$ lassen sich zwei Fälle unterscheiden:

1. das erste Argument ist die *leere Liste* (*Rekursionsverankerung*):
Die Konkatenation einer leeren Liste $\langle \rangle$ mit einer anderen Liste l ist l :

$$\langle \rangle \circ l = l$$

2. das erste Argument ist eine *nicht-leere Liste* (*Rekursionsschritt*):
Die Konkatenation einer nicht-leeren Liste $\langle e \bullet r \rangle$ aus einem ersten Element e und einer Restliste r mit einer Liste l entspricht einer Liste mit e als erstem Element und der Konkatenation von r mit l als Restliste:

$$\langle e \bullet r \rangle \circ l = \langle e \bullet (r \circ l) \rangle$$

Mit Hilfe der Merkmale LEFT, RIGHT und RESULT lassen sich diese beiden Fälle folgendermaßen beschreiben:

(87) a. *Rekursionsverankerung*:

$$\left[\begin{array}{ll} \text{append} & \\ \text{LEFT} & \langle \rangle \\ \text{RIGHT} & \boxed{1} \\ \text{RESULT} & \boxed{1} \end{array} \right]$$

b. *Rekursionsschritt:*

$$\left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \langle \boxed{1} \bullet \boxed{2} \rangle \\ \text{RIGHT} \quad \boxed{3} \\ \text{RESULT} \quad \langle \boxed{1} \bullet \boxed{4} \rangle \end{array} \right] \quad gdw. \quad \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \boxed{2} \\ \text{RIGHT} \quad \boxed{3} \\ \text{RESULT} \quad \boxed{4} \end{array} \right]$$

Auf eine Struktur

$$(88) \quad \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \text{Liste 1} \\ \text{RIGHT} \quad \text{Liste 2} \\ \text{RESULT} \quad ? \end{array} \right]$$

muß also im Falle einer leeren ersten Liste der Constraint 1 und im Falle einer nicht-leeren ersten Liste der Constraint 2 angewandt werden. Werden die beiden Bedingungen (87a) und (87b) als alternative Constraints zum Typ *append* aufgefaßt, wird genau dieses erreicht: [LEFT $\langle \rangle$] läßt sich nur mit einer Struktur mit leerer Liste als Wert von LEFT unifizieren, und [LEFT $\langle \boxed{1} \bullet \boxed{2} \rangle$] kann nur mit einer Struktur mit nicht-leerer Liste als erstem Argument unifiziert werden.

Die Bedingung im Rekursionsschritt, hier durch “*gdw.*” (“*genau dann wenn*”) ausgedrückt, läßt sich mit einem zusätzlichen Merkmal **CONDITION** mit Werten, die wiederum vom Typ *append* sind, formulieren. Insgesamt erhalten wir damit:

(89) **Constraints für den Typ *append*:**

$$\left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \langle \rangle \\ \text{RIGHT} \quad \boxed{1} \\ \text{RESULT} \quad \boxed{1} \end{array} \right] \quad \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \langle \boxed{1} \bullet \boxed{2} \rangle \\ \text{RIGHT} \quad \boxed{3} \\ \text{RESULT} \quad \langle \boxed{1} \bullet \boxed{4} \rangle \\ \text{CONDITION} \quad \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \boxed{2} \\ \text{RIGHT} \quad \boxed{3} \\ \text{RESULT} \quad \boxed{4} \end{array} \right] \end{array} \right]$$

Das Merkmal **CONDITION** erzwingt eine rekursive Berechnung der Konkatenation: Eine Struktur des Typs *append* mit nicht leerer Liste als Wert des Merkmals **LEFT** muß eines der beiden Constraints erfüllen. Da es mit

dem ersten Constraint wegen des nicht leeren ersten Elementes im Widerspruch steht, muß es dem zweiten Constraint genügen. Durch die Unifikation der Variablen $\boxed{2}$ mit der Restliste von LEFT erscheint diese als erstes Argument der Struktur des Merkmals CONDITION. Das zweite Argument dieser Struktur ist mit dem zweiten Argument der äußeren Struktur identisch. Da die CONDITION-Struktur wiederum vom Typ *append* ist, wiederholt sich der Berechnungsschritt mit der neuen, um ein Element kürzeren Argumentliste rekursiv, bis die Rekursionsverankerung angewandt werden kann. Über die koindizierten RESULT-Merkmale ergibt sich das Ergebnis der Listen-Konkatenation.

Im Falle der Ausgangsstruktur

$$(90) \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \langle a, b \rangle \\ \text{RIGHT} \quad \langle c, d \rangle \\ \text{RESULT} \quad ? \end{array} \right]$$

würde sich beispielsweise die folgende Ergebnisstruktur ergeben (der Übersichtlichkeit halber werden die Werte der Variablen wiederholt):

$$(91) \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \langle \boxed{1} a \bullet \boxed{2} \langle \boxed{3} b \bullet \boxed{4} \langle \rangle \rangle \rangle \\ \text{RIGHT} \quad \boxed{5} \langle c, d \rangle \\ \text{RESULT} \quad \langle \boxed{1} a \bullet \boxed{6} \langle \boxed{3} b \bullet \boxed{5} \langle c, d \rangle \rangle \rangle \\ \text{CONDITION} \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \boxed{2} \langle \boxed{3} b \bullet \boxed{4} \langle \rangle \rangle \\ \text{RIGHT} \quad \boxed{5} \langle c, d \rangle \\ \text{RESULT} \quad \boxed{6} \langle \boxed{3} b \bullet \boxed{5} \langle c, d \rangle \rangle \\ \text{CONDITION} \left[\begin{array}{l} \text{append} \\ \text{LEFT} \quad \boxed{4} \langle \rangle \\ \text{RIGHT} \quad \boxed{5} \langle c, d \rangle \\ \text{RESULT} \quad \boxed{5} \langle c, d \rangle \end{array} \right] \end{array} \right] \end{array} \right]$$

Das Merkmal CONDITION wird nur zur Berechnung der funktionalen Abhängigkeit von Merkmalen benötigt. Die Verwendung einer Funktion als Abkürzung ist deshalb übersichtlicher:

(92) Abkürzung: abgekürzte Merkmalstruktur:

$$\left[\begin{array}{l} \text{LIST1 } \boxed{1} \\ \text{LIST2 } \boxed{2} \\ \text{LIST } \boxed{1} \circ \boxed{2} \end{array} \right] \quad \left[\begin{array}{l} \text{LIST1 } \boxed{1} \\ \text{LIST2 } \boxed{2} \\ \text{LIST } \boxed{3} \\ \\ \text{CONDITION } \left[\begin{array}{l} \text{append} \\ \text{LEFT } \boxed{1} \\ \text{RIGHT } \boxed{2} \\ \text{RESULT } \boxed{3} \end{array} \right] \end{array} \right]$$

1.6.2 Syntax und Semantik

Der letzte Abschnitt dieses Kapitels ist dem Zusammenspiel von Syntax und Semantik gewidmet. Es soll beschrieben werden, wie diese beiden Komponenten der Grammatik mit ihren wechselseitigen Abhängigkeiten in einem Constraintsystem dargestellt werden können.

Aufbau der Semantik

Nicht die adäquate Modellierung semantischer Phänomene soll uns hier beschäftigen, sondern der Mechanismus der *Komposition* semantischer Beschreibungen entlang der syntaktischen Struktur. Aus diesem Grund wird eine schematisierte, kompositionale Semantik verwendet, die — soweit möglich — von allen semantischen Details abstrahiert, die zur Veranschaulichung des kompositionalen Prinzips nicht erforderlich sind.

Eigennamen werden als Atome dargestellt:

(93) *Semantik von Eigennamen:*
elizabeth, sigrid, leonard, sebastian

Adjektive und Verben durch Prädikate:

(94) *Semantik von Verben:*

$$\left[\begin{array}{l} \text{sleep} \\ \text{① cont} \end{array} \right] \quad \left[\begin{array}{l} \text{love} \\ \text{① cont} \\ \text{② cont} \end{array} \right] \quad \left[\begin{array}{l} \text{believe} \\ \text{① cont} \\ \text{② cont} \end{array} \right]$$

Die Semantik von Phrasen und Sätzen setzt sich kompositional aus der Semantik ihrer Bestandteile zusammen.

Der Eigenname *Leonard* und das Verb *schlafen* haben beispielsweise die folgende Semantik:

(95) a. Semantik von “Leonard”:
 leonard

b. Semantik von “schlafen”:

$$\begin{bmatrix} \text{sleep} \\ \textcircled{1} \text{ cont} \end{bmatrix}$$

Durch Komposition ergibt sich daraus die Semantik des Satzes *Leonard schläft*:

(96) Semantik von “Leonard schläft”:

$$\begin{bmatrix} \text{sleep} \\ \textcircled{1} \text{ leonard} \end{bmatrix}$$

Die Semantik transitiver Verben ergibt sich analog:

(97) Semantik von “Leonard liebt Sigrid”:

$$\begin{bmatrix} \text{love} \\ \textcircled{1} \text{ leonard} \\ \textcircled{2} \text{ sigrid} \end{bmatrix}$$

Das Verb *glauben* nimmt einen Satz als Argument. Die Semantik dieses Satzes erscheint als Wert des zweiten Argumentes in der semantischen Struktur des Matrixverbes:

(98) Semantik von “Sigrid glaubt, Elizabeth liebt Leonard”:

$$\begin{bmatrix} \text{believe} \\ \textcircled{1} \text{ sigrid} \\ \textcircled{2} \begin{bmatrix} \text{love} \\ \textcircled{1} \text{ elizabeth} \\ \textcircled{2} \text{ leonard} \end{bmatrix} \end{bmatrix}$$

Da die Semantik in Form von Merkmalstrukturen dargestellt werden kann, läßt sie sich als Wert des neu eingeführten Merkmals CONT in die lexikalischen Einträge und die Phrasenstrukturregeln integrieren.

Den lexikalischen Einträgen wird der Type zugeordnet, der ihre Semantik kodiert:

(99) a.
$$\begin{bmatrix} \text{lex-entry} \\ \text{PHON} \langle \text{Leonard} \rangle \\ \text{CAT} \text{ np} \\ \text{CONT} \text{ leonard} \end{bmatrix}$$

$$b. \left[\begin{array}{l} \textit{lex-entry} \\ \text{PHON } \langle \textit{schläft} \rangle \\ \text{CAT } \textit{vp} \\ \text{CONT } \textit{sleep} \end{array} \right]$$

In den Phrasenstrukturregeln wird angegeben, wie sich die Semantik der Tochterzeichen zu der Semantik des ganzen Zeichens zusammensetzen läßt:

Im Falle der Regel

$$(100) \quad S \longrightarrow NP \ VP$$

für den Satz *Leonard schläft* beispielsweise wird die Semantik des Verbes *schläft* auf die Semantik der Nominalphrase *Leonard* angewandt, und schließlich als Semantik des ganzen Zeichens übernommen:

$$(101) \left[\begin{array}{l} \textit{phrase-rule} \\ \text{CAT } \textit{s} \\ \text{CONT } \boxed{\text{I}} \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{CAT } \textit{np} \\ \text{CONT } \boxed{\text{2}} \end{array} \right], \left[\begin{array}{l} \text{CAT } \textit{vp} \\ \text{CONT } \boxed{\text{I}} [\textcircled{1} \boxed{\text{2}}] \end{array} \right] \right\rangle \end{array} \right]$$

Eine einfache Grammatik, die beispielsweise die Sätze

$$(102) \begin{array}{ll} a. & \textit{Elizabeth schläft} \\ & \left[\begin{array}{l} \textit{sleep} \\ \textcircled{1} \textit{elizabeth} \end{array} \right] \\ b. & \textit{Sebastian liebt Sigrid} \\ & \left[\begin{array}{l} \textit{love} \\ \textcircled{1} \textit{sebastian} \\ \textcircled{2} \textit{sigrid} \end{array} \right] \\ c. & \textit{Elizabeth glaubt, Sebastian liebt Sigrid} \\ & \left[\begin{array}{l} \textit{believe} \\ \textcircled{1} \textit{elizabeth} \\ \textcircled{2} \left[\begin{array}{l} \textit{love} \\ \textcircled{1} \textit{sebastian} \\ \textcircled{2} \textit{sigrid} \end{array} \right] \end{array} \right] \end{array}$$

beschreibt, läßt sich damit folgendermaßen als Constraintsystem modellieren:

(103) **Typenhierarchie:**

$phon \Rightarrow Elizabeth \vee Sigrid \vee Leonard \vee Sebastian \vee liebt \vee schläft \vee glaubt.$
 $cont \Rightarrow elizabeth \vee sigrid \vee leonard \vee sebastian \vee love \vee sleep \vee believe.$
 $cat \Rightarrow np \vee vp \vee tv \vee su \vee s.$
 $sign \Rightarrow phrase \vee word.$
 $phrase \Rightarrow phrase-rule.$
 $word \Rightarrow lex-entry.$

Constraints:

$$\begin{array}{l}
\begin{bmatrix} sleep \\ \textcircled{1} cont \end{bmatrix} \quad \begin{bmatrix} love \\ \textcircled{1} cont \\ \textcircled{2} cont \end{bmatrix} \quad \begin{bmatrix} believe \\ \textcircled{1} cont \\ \textcircled{2} cont \end{bmatrix} \\
\\
\begin{bmatrix} sign \\ PHON list \\ CAT cat \\ CONT cont \end{bmatrix} \\
\\
\begin{bmatrix} phrase \\ PHON \textcircled{1} \circ \textcircled{2} \\ DTRS \left\langle \begin{bmatrix} sign \\ PHON \textcircled{1} \end{bmatrix}, \begin{bmatrix} sign \\ PHON \textcircled{2} \end{bmatrix} \right\rangle \end{bmatrix} \\
\\
\begin{bmatrix} phrase-rule \\ CAT s \\ CONT \textcircled{1} \\ DTRS \left\langle \begin{bmatrix} CAT np \\ CONT \textcircled{2} \end{bmatrix}, \begin{bmatrix} CAT vp \\ CONT \textcircled{1} [\textcircled{1} \textcircled{2}] \end{bmatrix} \right\rangle \end{bmatrix} \quad \begin{bmatrix} phrase-rule \\ CAT vp \\ CONT \textcircled{1} \\ DTRS \left\langle \begin{bmatrix} CAT tv \\ CONT \textcircled{1} [\textcircled{2} \textcircled{2}] \end{bmatrix}, \begin{bmatrix} CAT np \\ CONT \textcircled{2} \end{bmatrix} \right\rangle \end{bmatrix} \\
\\
\begin{bmatrix} phrase-rule \\ CAT vp \\ CONT \textcircled{1} \\ DTRS \left\langle \begin{bmatrix} CAT su \\ CONT \textcircled{1} [\textcircled{2} \textcircled{2}] \end{bmatrix}, \begin{bmatrix} CAT s \\ CONT \textcircled{2} \end{bmatrix} \right\rangle \end{bmatrix} \\
\\
\begin{bmatrix} lex-entry \\ PHON \langle Elizabeth \rangle \\ CAT np \\ CONT elizabeth \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \langle Sigrid \rangle \\ CAT np \\ CONT sigrid \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \langle Leonard \rangle \\ CAT np \\ CONT leonard \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \langle Sebastian \rangle \\ CAT np \\ CONT sebastian \end{bmatrix} \\
\\
\begin{bmatrix} lex-entry \\ PHON \langle liebt \rangle \\ CAT tv \\ CONT love \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \langle schläft \rangle \\ CAT vp \\ CONT sleep \end{bmatrix} \quad \begin{bmatrix} lex-entry \\ PHON \langle glaubt \rangle \\ CAT su \\ CONT believe \end{bmatrix}
\end{array}$$

Gegeben eine Grammatik kann das Resolutionsverfahren für die folgenden Berechnungen benutzt werden:

Untersuchung der Korrektheit — Für ein gegebenes Zeichen wird ermittelt, ob es im Sinne der Grammatik korrekt ist;

Parsen — Zu einem Satz kann die Menge möglicher linguistischer Strukturen — und damit auch möglicher semantischer Interpretationen — berechnet werden;

Generierung — Zu einer gegebenen semantischen Struktur kann die Menge der möglichen, sprachlichen Realisierungen ermittelt werden;

Vervollständigung — Zu einem unterspezifizierten Zeichen können alle korrekten Zeichen aufgezählt werden, die durch dieses subsumiert werden;

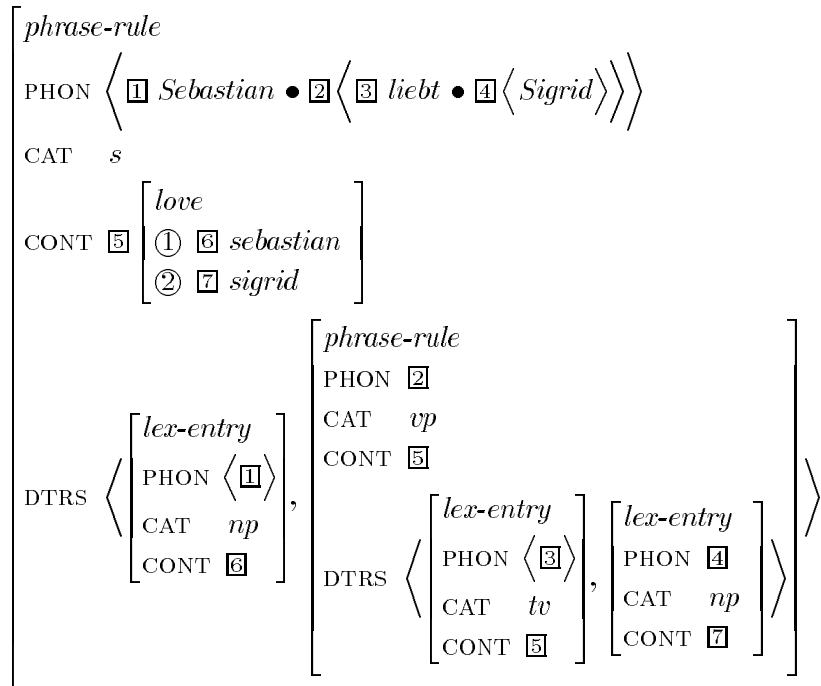
Aufzählung — Alle Zeichen, die durch die Grammatik beschrieben werden, können aufgezählt werden.

Für die drei Sätze (102a), (102b) und (102c) erhalten wir mit dieser Grammatik die folgenden Strukturen:

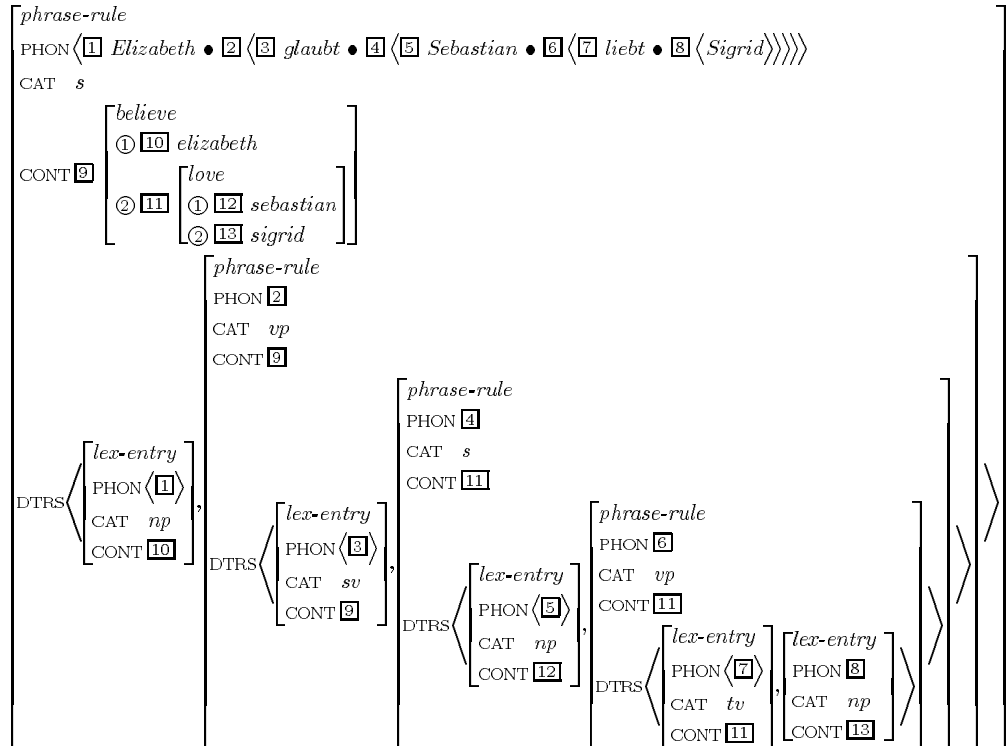
(104) a. *Elizabeth schläft*

$$\left[\begin{array}{l} \text{phrase-rule} \\ \text{PHON } \langle \boxed{1} \text{ Elizabeth } \bullet \boxed{2} \langle \text{schläft} \rangle \rangle \\ \text{CAT } s \\ \text{CONT } \boxed{3} \left[\begin{array}{l} \text{sleep} \\ \textcircled{1} \text{ } \boxed{4} \text{ elizabeth} \end{array} \right] \\ \text{DTRS } \left\langle \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON } \langle \boxed{1} \rangle \\ \text{CAT } np \\ \text{CONT } \boxed{4} \end{array} \right], \left[\begin{array}{l} \text{lex-entry} \\ \text{PHON } \boxed{2} \\ \text{CAT } vp \\ \text{CONT } \boxed{3} \end{array} \right] \right\rangle \end{array} \right]$$

b. *Sebastian liebt Sigrid*



c. *Elizabeth glaubt, Sebastian liebt Sigrid*



Vergleichen wir (104c) mit (1), und (104a) mit (2), so zeigt sich, daß wir am Anfang dieses Kapitels von den Zeichen ausgegangen sind, die wir nun durch

unsere Grammatik erklären können.³⁶

1.7 Defaults und Defaultvererbung

Monotone und nicht-monotone Logiken

Der Formalismus, der in den letzten Abschnitten vorgestellt wurde, kann als eine *monotone Logik* aufgefaßt werden. Monotone Logiken sind — einfach ausgedrückt — dadurch gekennzeichnet, daß Folgerungen, die aus einer Menge von Aussagen abgeleitet werden können, auch dann ihre Gültigkeit bewahren, wenn zu dieser Menge weitere Aussagen hinzugefügt werden. Diese Eigenschaft hat aus mathematischer und algorithmischer Sicht viele Vorteile; aus einer Perspektive, die sich an Eleganz und Ökonomie der Modelle orientiert, ist sie jedoch oft Ursache sperriger, unökonomischer Beschreibungen.

Logiken, in denen Folgerungen durch die Hinzunahme zusätzlicher Aussagen revidiert werden können, werden als *nicht-monotone Logiken* bezeichnet.

Defaultlogiken

Unter den nicht-monotonen Logiken sind aus Sicht der Linguistik insbesondere *Defaultlogiken* interessant: Sie gestatten die Formulierung von “Annahmen”, die solange gültig bleiben, wie sie nicht durch andere, höher bewertete Annahmen überschrieben werden:

- (105) a. *Alle Vögel können fliegen.*
 b. *Pinguine sind Vögel, die nicht fliegen können.*

(105a) beispielsweise behält seine Gültigkeit für alle Vögel, solange keine weitere Aussage ihr widerspricht. In dem Moment, in dem (105b) hinzugenommen wird, gilt (105a) nur noch eingeschränkt für Vögel, die keine Pinguine sind: (105b) überschreibt (105a) für die Untermenge der Pinguine.

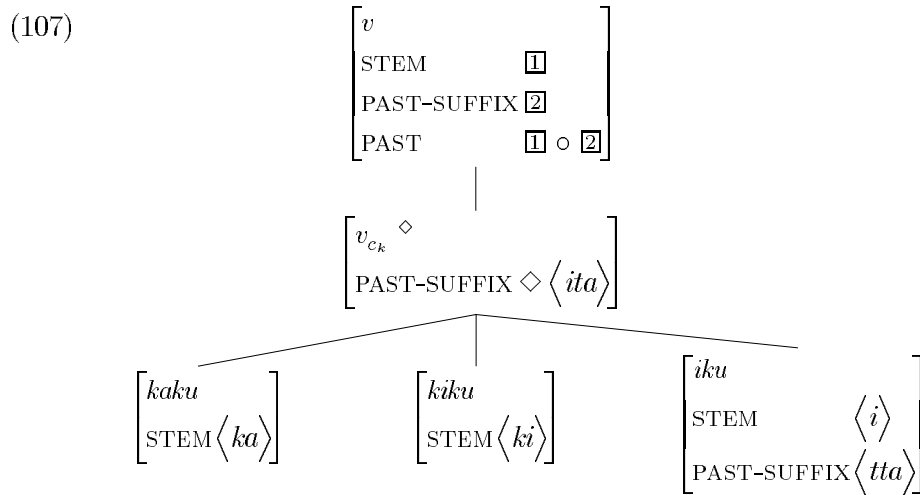
Defaulthierarchien

In der Linguistik gestatten Defaults eine elegante Modellierung von Ausnahmen: Das Perfekt japanischer Verben läßt sich beispielsweise durch Verknüpfung eines Stammes mit einem Perfekt-Suffix beschreiben. Das Perfekt-Suffix hat für konsonantische Verben die Form *ita*. Nur ein konsonantisches Verb weicht davon ab: das Verb 行く (*iku*; gehen, fahren), das das Perfekt mit dem Suffix *tta* bildet.

³⁶ Das Zeichen (2) wurde zur einfacheren Beschreibung vereinfacht dargestellt.

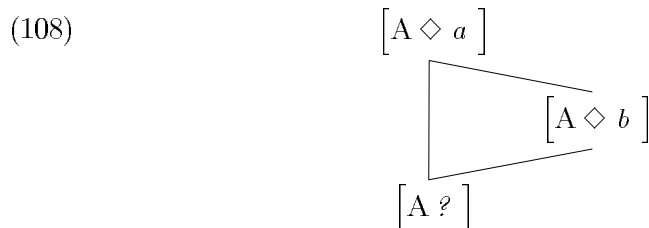
- (106) a. 書く $\rightarrow$ 書いた (*schreiben*)
 kaku *kaita*
- b. 聞く $\rightarrow$ 聞いた (*hören*)
 kiku *kiita*
- c. 行く $\rightarrow$ 行った (*gehen*)
 iku *itta*

Zur Darstellung solcher Ausnahmen verwenden wir Hierarchien mit einem abweichenden Vererbungmechanismus, die wir als *Defaulthierarchien* bezeichnen. Durch eine Defaulthierarchie läßt sich die Abweichung des Perfektes von *iku* folgendermaßen beschreiben:³⁷



Die Raute $\diamond$ dient der Hervorhebung von Defaults: Werte, die durch eine Raute markiert werden, lassen sich durch entsprechende Werte kleinerer Typen überschreiben. Damit Constraints, die Defaults enthalten, leichter erkannt werden können, markieren wir ihren äußeren Typ durch eine hochgestellte Raute.

Unterschiedliche Defaultwerte können kollidieren:



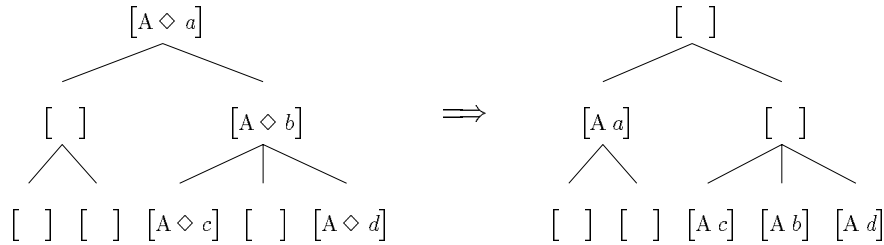
³⁷ Die hier dargestellte Modellierung des Perfekts von Verben dient nur der Veranschaulichung der Defaultvererbung; sie entspricht nicht dem Vorgehen in unserem Modell der japanischen Verbmorphologie.

Solche Fälle können in unterschiedlicher Weise aufgelöst werden. In dem hier vorgestellten Modell kommen Kollisionen nicht vor. Wir können uns jedoch vorstellen, daß aufeinandertreffende Defaultwerte genauso behandelt werden, wie normale Merkmalwerte: durch Unifikation.

Von Defaulthierarchien zu monotonen Hierarchien

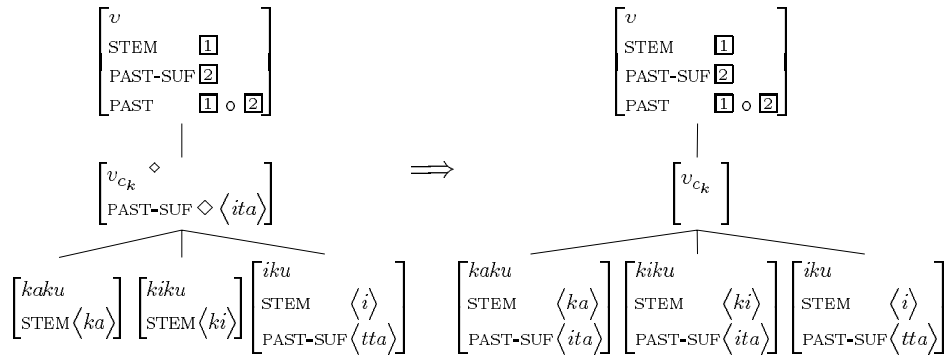
Defaulthierarchien lassen sich als eine Art “Metasprache” für monotone Hierarchien auffassen: Durch das “Herunterpropagieren” der Defaultwerte lassen sich alle Überschreibungen auflösen, so daß — wie das folgende Beispiel veranschaulicht — eine monotone Hierarchie entsteht:

(109) *Ableitung einer monotonen Hierarchie aus einer Defaulthierarchie:*



Für die Hierarchie (107), der Beschreibung des Perfekts k-konsonantischer Verben, würde sich in analoger Weise die folgende Hierarchie ergeben:

(110) *Ableitung der monotonen Hierarchie für die Perfektform:*



Defaulthierarchien, wie sie in dieser Arbeit verwendet werden, dienen also nur der verkürzten Notation monotoner Hierarchien. Für Beschreibungen, die mit Hilfe von Defaulthierarchien spezifiziert werden, läßt sich daher der Kalkül benutzen, der für monotone Beschreibungen im letzten Kapitel entwickelt wurde. Defaultlogiken, die sich nicht auf monotone Logiken reduzieren lassen, werden beispielsweise von Carpenter 1993 und Lascarides und Copestake 1999 behandelt.

2 Japanische Verbmorphologie

Das vorliegende zweite und das vierte Kapitel dieser Arbeit beschäftigen sich ausschließlich mit der Morphologie japanischer Verben: Im zweiten Kapitel wird die traditionelle japanische Verbmorphologie beschrieben und anschließend nach formalen Kriterien überarbeitet; Kapitel vier dient der Umsetzung des überarbeiteten Modells in eine Constraintbeschreibung und der Integration von Morphologie, Syntax und Semantik.

Die traditionelle japanische Verbmorphologie mag für die Vermittlung der Grammatik im Sprachenunterricht Vorzüge haben. Der Grund hierfür, ihr relativ einfacher und intuitiver Aufbau, ist jedoch auch dafür verantwortlich, daß sie den Kriterien der formalen Linguistik nur teilweise genügt.³⁸ Wie sich aus der traditionellen Verbmorphologie ein formales System ableiten läßt, ist daher, neben der Einführung in die Eigenschaften und die Morphologie japanischer Verben, einer der Schwerpunkte der folgenden Betrachtungen.

Einführung und Überarbeitung des morphologischen Systems geschehen in vier Schritten:

1. Einführung in die Charakteristika japanischer Verben und Adjektive;
2. Vorstellung der traditionellen Beschreibung der japanischen Verbmorphologie;
3. Ableitung und Beschreibung eines formalen verbmorphologischen Systems, das aus der traditionellen Verbmorphologie und der “Japanischen Morphosyntax” von Jens Rickmeyer (Rickmeyer 1995) hervorgeht;
4. Bemerkungen zu den Wechselwirkungen, die zwischen der Morphologie und anderen grammatischen Komponenten bestehen.

2.1 Verben und japanische Grammatik

In diesem Abschnitt sollen einige Charakteristika der japanischen Verbmorphologie vorgestellt werden, die für unser Thema wichtig sind: Zuerst wird die *Agglutination* erklärt, die als das grundlegende Prinzip der japanischen Morphologie bezeichnet werden kann. *Allomorphie* und *phonologische Verschleifungen*, zwei Phänomene, die die Agglutination der Morpheme begleiten und die Beschreibung der Morphologie schwierig machen, werden anschließend vorgestellt. Schließlich werden die Auswirkungen der morphologischen

³⁸ Für einen formalen, GPSG-basierten Ansatz, der trotz ihrer Mängel die traditionelle Verbmorphologie zugrunde legt, vergleiche Udo 1982.

Prozesse auf andere sprachliche Ebenen wie Syntax und Semantik beschreiben.

2.1.1 Agglutination

Morphologisch wird die japanische Sprache dem *agglutinierenden Sprachtyp* zugeordnet. Was mit dieser Klassifikation gemeint ist, läßt sich am besten anhand eines Beispiels darstellen:

(111) 私は先生にこの絵をほめられたかった。³⁹

watasi-ha señsei-ni kono e-wo home-rare-ta-ka.tt-a
 ich-TOP Lehrer-DAT dieses Bild-ACC loben-PASS-VOLU- ϕ -PAST
 Ich wollte vom Lehrer für dieses Bild gelobt werden.

Der Verbform *ほめられたかった* (*home-rare-ta-ka.tt-a*⁴⁰) liegt das Lexem *ほめる* (*home-ru*) mit der Bedeutung *loben* zugrunde. Durch Schrittweises Anfügen modifizierender Suffixe läßt sich die Bedeutung und das syntaktische Verhalten dieses Verbes derart modifizieren, daß sich schließlich die oben genannte Form ergibt, mit der komplexen Bedeutung “[jmd.] wollte [von jmd.] [für etwas] gelobt werden”.

Die morphologische Segmentierung dieser Verbform ergibt folgende Kette von Morphemen⁴¹:

(112) *home- rare- ta- ka.tt- a*
 {loben} + {Passiv} + {Voluntativ} + {Karu} + {Perfekt}
 loben -PASS -VOLU - ϕ -PAST

Jedes Präfix dieser Folge entspricht, wenn es durch ein Flexiv abgeschlossen wird, selbst wiederum einer möglichen Verbform⁴²:

³⁹ vgl. DBJG, S. 444, notes 2. (B) (7).

⁴⁰ Die Bindestriche bezeichnen Morphemgrenzen; die Punkte die Grenzen zwischen den Bauteilen (Morphemkern und Clitics) aus denen sich die Morphe zusammensetzen. Die Umschrift nimmt damit Aspekte des verbmorphologischen Systems vorweg, die erst an späterer Stelle erklärt werden.

⁴¹ In der dritten Zeile werden die in Kapitel 5 eingeführten Abkürzungen verwendet.

⁴² (113d) ist zwar im Sinne der morphologischen Regeln eine mögliche Verbform, wird jedoch so nicht verwendet:

Das Suffixverb *かる* (*ka.r-u*) dient nur der Überführung des adjektivischen *-ta.i* in ein Verb. Es ermöglicht damit das Anfügen des Perfekt-Suffixes *a*, das nicht direkt an eine adjektivische Form angeschlossen werden kann.

- (113) a. ほめる
 ~~home~~-ru
 loben-FLEX
 loben
- b. ほめられる
 ~~home~~-rare-ru
 loben-PASS-FLEX
 gelobt werden
- c. ほめられたい
 ~~home~~-rare-ta.i
 loben-PASS-VOLU-FLEX
 es gerne mögen, gelobt zu werden
- d. (ほめられたかる)⁴²
 ~~home~~-rare-ta-ka.r-u
 loben-PASS-VOLU- ϕ -FLEX
 es gerne mögen, gelobt zu werden
- e. ほめられたかった
 ~~home~~-rare-ta-ka.tt-a
 loben-PASS-VOLU- ϕ -PAST
 es gerne gemocht haben, gelobt zu werden

Diese Art der Modifikation von Lexemen durch schrittweises Affigieren von Morphemen wird als *agglutinierend* bezeichnet; die Sprachen, die sich dieser Form der morphologischen Variation bedienen, werden *agglutinierende Sprachen* genannt.

Allomorphie und phonologische Verschleifungen

Zwei Eigenschaften sind typisch für den agglutinierenden Sprachtyp⁴³:

1. Morpheme werden durch eine kleine Anzahl genetisch verwandter oder phonologisch ähnlicher Allomorphe realisiert.

⁴³ Ich beziehe mich auf den Morph-, Allomorph- und Morphembegriff von Rickmeyer 1995, S. 33, Abschnitt 12: “Durch eine sinnvolle Analyse eines Satzes wie vor allem durch Bildung von Paradigmen, gewinnen wir die kleinsten syntaktischen Elemente: die Morphe. Morphe haben im Japanischen segmentalen Charakter. [...]”

Zu bestimmten Morphen [...] gibt es genetisch verwandte oder phonologisch ähnliche Morphe, die semantisch einander äquivalent, aber zueinander komplementär distribuiert sind: die Allomorphe. Ein Morphem ist nun definiert als die Menge aller solcher Allomorphe bzw. für den Fall, daß es zu einem Morph keine Allomorphe gibt, als dieses Morph selbst.”

2. Morphe sind segmental; ihre Affigierung geschieht ohne größere Verschleifungen.

Dem Ideal dieses Sprachtyps würde eine Sprache entsprechen, in der jedes Morphem durch genau ein Morph realisiert wird, und die Verkettung der Morphe ohne jede Veränderung ihrer phonologischen Form geschieht.

Im Japanischen gibt es jedoch sowohl Allomorphie — d.h. mehrere Morphe pro Morphem — als auch phonologische Verschleifungen⁴⁴. Beide Prozesse führen zu einer engen Verbindung der morphologischen Elemente, so daß der Übergang zwischen den einzelnen Morphen oft fließend ist und die Morphemgrenze willkürlich festgesetzt werden muß.

Die Verschleifung der Morpheme ist charakteristisch für den flektierenden Sprachbau. Von manchen Linguisten wird das Japanische daher als agglutinierende Sprache mit flektierenden Zügen bezeichnet⁴⁵.

Wie wir im Abschnitt 2.2 *“Die Verbmorphologie der japanischen Schulgrammatik”* sehen werden, wird die Allomorphie in der traditionellen japanischen Schulgrammatik durch die Subklassifizierung der Verben in verschiedene Kategorien und ein System unterschiedlicher Stämme behandelt. Die Verschleifungen hingegen werden durch eine Menge phonologischer Regeln modelliert.

Syntaktische und semantische Modifikationen

Durch die morphologische Veränderung eines Wortes wird dessen semantisches und / oder syntaktisches Verhalten modifiziert:

- (114) a. 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
Tarō-NOM Zirō-ACC rufen-PAST
Tarō rief Zirō.

- b. 美智子が太郎に二郎を呼ばせた。⁴⁶

mitiko-ga tarou-ni zirou-wo yo.b.a-se.t-a
Mitiko-NOM Tarō-DAT Zirō-ACC rufen-CAUS-PAST
Mitiko ließ Tarō Zirō rufen.

- c. 太郎が美智子に二郎を呼ばせられた。⁴⁷

44 Die Grenze zwischen Verschleifung und Allomorphie ist dabei schwer zu ziehen: Wie wir später sehen werden, lassen sich die gleichen Phänomene sowohl durch Allomorphie als auch durch Verschleifungen modellieren.

45 Vergleiche beispielsweise Lewin 1990, S. 1, I Einleitung, Sprachstruktur.

46 LIJC, S. 34, (78)a.

47 LIJC, S. 34, (78)b.

tarou-ga mitiko-ni zirou-wo yo.b.a-se-rare.t-a
 Tarō-NOM Mitiko-DAT Zirō-ACC rufen-CAUS-PASS-PAST
 Tarō wurde von Mitiko veranlaßt, Zirō zu rufen.

Das Verb 呼ぶ (*yo.b-u*) mit der Bedeutung “rufen” in (114a) beispielsweise subkategorisiert für zwei Postpositionalphrasen: eine im Nominativ (das Subjekt) und eine im Akkusativ (das direkte Objekt).

Die Form 呼ばせた (*yo.b.a-se.t-a*) in (114b) hingegen, die durch Anfügen des Kausativ-Morphems an den Stamm von 呼ぶ (*yo.b-u*) gewonnen wird, hat nicht nur eine modifizierte Semantik, sondern verhält sich auch syntaktisch anders als ihre Ausgangsform: Aus dem zweiwertigen semantischen Prädikat *rufen*, der Bedeutung des Verbstammes, wird das komplexe Prädikat *zum Rufen veranlassen* der kausativischen Form. Der Akt des “Veranlassens” aber verlangt nach einem weiteren Objekt, dem “Veranlassenden”, so daß die entstehende Semantik drei Argumente fordert.

Die modifizierte Semantik spiegelt sich in der veränderten Syntax von 呼ばせた (*yo.b.a-se.t-a*): Das unterliegende Subjekt des Verbstammes erscheint als Dativ-Objekt, also als Postpositionalphrase mit der Postposition *ni*; der das Rufen “Veranlassende” hingegen wird zum Oberflächensubjekt des kausativischen Satzes.

Satz (114c) schließlich entsteht durch die Passivierung des kausativierten Verbes. Er hat zwar eine ähnliche Bedeutung wie sein aktivisches Äquivalent (114b), verhält sich jedoch syntaktisch anders: Das unterliegende Subjekt des Verbes, das im Kausativsatz zum Dativ-Objekt wurde, erscheint jetzt wieder im Nominativ. Dafür wird der durch das Kausativ eingeführte “Veranlassende” zum Dativ-Objekt.

2.1.2 Verben und Adjektive

Japanische Verben (動詞, *dousi*) und Adjektive (形容詞, *keiyousi*) unterscheiden sich hinsichtlich ihrer Morphologie und ihres syntaktischen Verhaltens von den entsprechenden Kategorien indoeuropäischer Sprachen. Insbesondere sind sich Verben und Adjektive im Japanischen morphologisch und syntaktisch so ähnlich, daß sie gemeinsam unter der Kategorie 用言 (*yougeñ*) zusammengefaßt werden. Wir verwenden den Begriff *Verb* in einem engeren und einem weiteren Sinne sowohl für die 動詞 (*dousi*) als auch für die 用言 (*yougeñ*). Die jeweils intendierte Bedeutung ist aus dem Kontext ersichtlich. Wo dieses nicht der Fall ist, wird explizit auf sie hingewiesen.

Eigenschaften von Verben (動詞, *dousi*) und Adjektiven (形容詞, *keiyousi*)

Im letzten Abschnitt wurde auf die Unterschiede hingewiesen, die zwischen japanischen Verben und Adjektiven einerseits und den entsprechenden Kategorien indoeuropäischer Sprachen andererseits bestehen. Einige der Unterschiede wollen wir uns nun näher ansehen.

1. Flexion

Im Japanischen sind nicht nur Verben, sondern auch Adjektive flektierbar: Ein großer Teil der Formen lassen sich von beiden Kategorien gleichermaßen bilden. Das Adjektiv 高い (*takai*; hoch, teuer) beispielsweise kann genauso ins Perfekt gesetzt werden, wie das Verb 食べる (*taberu*; essen):

(115) *Perfekt von Verben und Adjektiven:*

		<i>Präsens</i>		<i>Perfekt</i>
a.	Verb:	食べる <i>taberu</i>	→	食べた <i>tabeta</i>
b.	Adjektiv:	高い <i>takai</i>	→	高かった <i>takakatta</i>

In Anzahl und Art der Bildung bestehen jedoch Unterschiede zwischen Verben und Adjektiven.

Nominalverben und Nominaladjektive

Neben den Verben und Adjektiven kennt das Japanische auch *Nominalverben* und *Nominaladjektive* 形容動詞 (*keiyoudousi*). Beide haben ihren Ursprung im Chinesischen, einer isolierenden Sprache ohne jede Flexion, und können daher nicht flektiert werden. Formal gesehen ähneln sie damit den ebenfalls unflektierbaren Nomen, bzw., dank ihrer verbalen Bedeutung, nominalisierten Verben bzw. Adjektiven:

(116) a.	<i>Nominalverb:</i>	b.	<i>Nominaladjektiv:</i>
	勉強		きれい
	<i>beñkyou</i>		<i>kirei</i>
	“Studieren”		“Schönheit”

Um die Nominalverben und -adjektive in der japanische Syntax verwenden zu können, muß auch von ihnen das Paradigma verbaler bzw. adjektivischer Formen gebildet werden. Diese Aufgabe wird im Falle der Nominalverben von dem Hilfsverb する (*suru*) übernommen; für die Nominaladjektive dient dem selben Zweck das Hilfsverb だ (*da*):

(117)

	Verb	Adjektiv	Nominalverb	Nominaladj.
	食べる <i>taberu</i> (essen)	高い <i>takai</i> (hoch, teuer)	勉強 <i>beñkyou</i> (“Studieren”)	きれい <i>kirei</i> (“Schönheit”)
Präsens	食べる <i>taberu</i>	高い <i>takai</i>	勉強 する <i>beñkyou suru</i>	きれい だ <i>kirei da</i>
Perfekt	食べた <i>tabeta</i>	高かった <i>takakatta</i>	勉強 した <i>beñkyou sita</i>	きれい だった <i>kirei datta</i>
Honorativ	食べます <i>tabemasu</i>	高 かい です <i>takai desu</i>	勉強 します <i>beñkyou simasu</i>	きれい です <i>kirei desu</i>
...	...	...	...	...

Die folgenden Beispiele illustrieren die Verwendung der Nominalverben und -adjektive in der Syntax:

- (118) a. i. *Nominalverb mit prädikativer Kopula:* ii. *Nominalverb mit adnominaler Kopula:*
- 学生は勉強する。 勉強する学生
- gakusei-ha beñkyou* *beñkyou suru gakusei*
Student-TOP Studieren Studieren COP-FLEX Student
suru
COP-FLEX
Studenten studieren.
- b. i. *Nominaladjektiv mit prädikativer Kopula:* ii. *Nominaladjektiv mit adnominaler Kopula:*
- 花はきれいだ。 きれいな花
- hana-ha kirei* *kirei-na hana*
Blume-TOP schön Blume-COP Blume
da
COP-FLEX
Die Blume ist schön.

Nominalverben und -adjektive sind nicht Thema dieser Arbeit. Sie werden höchstens als synthetische Formen betrachtet: 勉強する (*beñkyou suru*; studieren) beispielsweise kann als synthetisches Verb ins Lexikon eingetragen werden, das sich in seiner Formbildung genau wie する (*suru*; tun, machen) verhält.

2. Syntakisches Verhalten

Japanische Adjektive können nicht nur adnominal bzw. als Komplement eines kopulativen Verbes verwendet werden, sondern treten auch selber als

Prädikat auf:

- (119) a. *prädikatives Adjektiv:* b. *adnominales Adjektiv:*

花は赤い。

hana-ha aka-i
Blume-TOP rot-FLEX
Die Blume ist rot.

赤い花

aka-i hana
rot-FLEX Blume
die / eine rote Blume

Auch adnominale Adjektive können dabei konjugiert werden: Die Form 赤かった (*aka-ka.tt-a*; rot gewesen sein) beispielsweise entspricht dem Perfekt des Adjektives 赤い (*aka.i*; rot), und kann ebenfalls sowohl prädikativ als auch adnominal benutzt werden:

- (120) a. *prädikatives Adjektiv* b. *adnominales Adjektiv*
im Perfekt: *im Perfekt:*

花は赤かった。

hana-ha aka-ka.tt-a
Blume-TOP rot-PAST
Die Blume war rot.

赤かった花

aka-ka.tt-a hana
rot-PAST Blume
die / eine Blume, die rot war

Ebenso wie Adjektive prädikativ verwendet werden können, lassen sich Verben adnominal benutzen. Genau wie Adjektive können sie im Japanischen einem Nomen direkt vorangestellt werden:

- (121) a. *prädikatives Verb:*

ジョンはステーキを食べた。⁴⁸

Jon-ha sutēki-wo tabe.t-a
John-TOP Steak-ACC essen-PAST
John hat ein Steak gegessen.

- b. *adnominales Verb:*

ジョンが食べたステーキ⁴⁹

Jon-ga tabe.t-a sutēki
John-NOM essen-PAST Steak
Das Steak, das John gegessen hat

Im Japanischen sind sich Adjektive und Verben also auch syntaktisch weit ähnlicher, als in indoeuropäischen Sprachen. Wie oben bereits angedeutet, verwenden wir den Begriff *Verb* daher oft für Verben und Adjektive zusammen.

⁴⁸ vgl. DBJG, S. 377, notes 1. (1)a..

⁴⁹ vgl. DBJG, S. 376, KS(A).

3. Überführung von Verben in Adjektive und umgekehrt

Einige Suffixe überführen Verben in Adjektive bzw. Adjektive in Verben:

(122) *Das Negationssuffix -na-i formt Verben in Adjektive um:*

食べる (V)	→	食べない (Adj)
<i>tabe-ru</i>		<i>tabe-na-i</i>
essen-FLEX		essen-NEG-FLEX
essen		nicht essen

(123) *Das (semantisch leere) Suffix -ka.r-u formt Adjektive in Verben um:*

高い (Adj)	→	高かる (V)
<i>taka-i</i>		<i>taka-ka.r-u</i>
hoch-FLEX		hoch- ϕ -FLEX
hoch		hoch

Mit Hilfe des Suffixes *-ka.r-u* wird beispielsweise das Perfekt der Adjektive gebildet:

(124) *Bildung des Perfekts von Adjektiven:*

高い (Adj)	→	高かった (V)
<i>taka-i</i>		<i>taka-ka.tt-a</i>
hoch-FLEX		hoch- ϕ -PAST
hoch		hoch gewesen sein

Valenz

Lexikalische Verb- und Adjektiv-Stämme unterscheiden sich hinsichtlich ihrer Valenz ähnlich wie ihre Entsprechungen in indoeuropäischen Sprachen: Adjektive subkategorisieren meist nur für ein Objekt, das sie semantisch modifizieren (z.B. *das große Haus*), während Verben ein oder mehrere Objekte fordern.

Fakultativität der grammatischen Elemente

Am Ende dieses Abschnittes soll auf ein weiteres wesentliches Merkmal agglutinierender Sprachen hingewiesen werden: den *fakultativen Charakter* der Modifikationen und die damit einhergehenden Folgen für die “Logik” der Sprache.

In flektierenden Sprachen wie dem Deutschen und Englischen können Verbformen durch eine Menge “morphologischer Baupläne” beschrieben werden. Jedes Merkmal, das durch die Morphologie bestimmt wird, läßt sich an einem obligatorischen Morphem ablesen, für dessen Realisierung notfalls ein *Null-Morph* (ϕ -Morph) angenommen wird. In agglutinierenden Sprachen hingegen

ist das Anfügen von Morphemen fakultativ: Ein entsprechendes Morphem wird nur dann angehängt, wenn die mit ihm verbundene Information notwendig ist, oder gesellschaftliche Normen ihren Ausdruck erforderlich machen.

Dieser “fakultative Charakter” der sprachlichen Elemente — und damit auch der in ihnen ausgedrückten Informationen — ist eine Tendenz, die sich durch alle Ebenen der japanischen Grammatik zieht: Ausgedrückt wird nur, was für das Gelingen der sprachlichen Kommunikation wirklich erforderlich ist. Alles, was für das intendierte Verständnis nicht wirklich nötig ist bzw. was aus dem Kontext oder der Situation erschlossen werden kann, wird nicht explizit erwähnt.⁵⁰

Das Japanische kennt dafür allerdings grammatische Kategorien, die in indoeuropäischen Sprachen nicht existieren: Dazu gehört die Höflichkeitssprache und ein kompliziertes System satzabschließender Partikeln, durch die beispielsweise die emotionale Einstellung des Sprechers gegenüber dem Gesagten ausgedrückt werden kann.

2.2 Die Verbmorphologie der japanischen Schulgrammatik

Im letzten Abschnitt wurde das Japanische hinsichtlich seiner Verbmorphologie als agglutinierende Sprache mit flektierenden Zügen charakterisiert. Als Kennzeichen der Annäherung an den flektierenden Sprachbau wurde die Verschleifung der Morphe angeführt, aus denen sich die Verbformen zusammensetzen. Folge dieser Verschleifung ist die Unmöglichkeit, Morphemgrenzen eindeutig zu bestimmen: Die Verbformen liefern keine objektiven Anhaltspunkte für ihre Zergliederung in Morphe.

Eine systematische Beschreibung der Verbformen setzt jedoch die Segmentierung in Morphe voraus. Da sich keine objektiven Maßstäbe für diesen Zweck angeben lassen, müssen andere Kriterien für die Untergliederung der Verbformen gefunden werden. Diese anderen Kriterien sind in den meisten Fällen ästhetischer Natur und betreffen die *Eleganz* und die *Ökonomie* der Beschreibung.

Ästhetische Kriterien sind subjektiv. Sie hängen von der Einschätzung des Linguisten bzw. der intendierten Verwendung der Sprachbeschreibung ab.

⁵⁰ Bei der Übersetzung in indoeuropäische Sprachen resultieren aus dieser Eigenschaft des Japanischen maschinell kaum zu lösende Probleme: Informationen, die im Japanischen nicht ausgedrückt werden (Definitheit, Numerus, Argumente des Verbs etc.), im Deutschen oder Englischen jedoch obligatorisch sind, müssen während der Übersetzung aus Kontext und Weltwissen erschlossen werden (vergleiche beispielsweise Siegel 1997, Siegel 1996a und Siegel 1995).

Es ist daher nicht verwunderlich, daß unterschiedliche, miteinander konkurrierende Beschreibungen der japanischen Morphologie existieren.

In diesem Abschnitt soll das phonologisch-morphologische System der traditionellen japanischen Schulgrammatik genauer dargestellt werden. Sie bedient sich — wie auch der Großteil der konkurrierenden Beschreibungen — eines Systems von Kategorien und Stammformen, um die Allomorphie der an der Verbmodifikation beteiligten Elemente zu modellieren. Verschleifungen, die mit diesem Mittel nur schwer darstellbar sind, werden durch phonologische Regeln abgebildet.

Unter den Begriffen *japanische Schulgrammatik* bzw. *traditionelle japanische Grammatik* fassen wir Sprachbeschreibungen zusammen, die sich an dem in der Edo-Zeit erarbeiteten System von sechs Verbstämmen bzw. Flexionsformen (活用形, *katuyoukei*) orientieren:

(125) *Die Verbstämme der klassischen Schriftsprache:*

未然形	<i>mizenkei</i>	Indefinitform
連用形	<i>renyoukei</i>	Konjunkionalform
終止形	<i>syuusikei</i>	Finalform
連体形	<i>rentaikei</i>	Attributivform
已然形	<i>izenkei</i>	Konditionalform
命令形	<i>meireikei</i>	Imperativform

Diese Stämme oder Flexionsformen wurden ursprünglich zur Beschreibung der klassischen Schriftsprache entwickelt (vgl. Lewin 1990, Paragraph 118, S.105). Zur Beschreibung des Neujapanischen finden daher verschiedene leicht modifizierte Systeme Anwendung. In unserer Beschreibung der japanischen Schulgrammatik beziehen wir uns weitgehend auf die Variante von McClain 1981. Es werden jedoch auch Aspekte anderer Beschreibungen miteinbezogen, beispielsweise Lewin 1990.

In Neujapanischen fallen 終止形 (*syuusikei*) und 連体形 (*rentaikei*) zusammen. McClain 1981 unterscheidet daher nicht zwischen diesen beiden Stämmen. Weiterhin wird der 已然形 (*izenkei*) in 假定形 (*kateikei*) umbenannt, und das Futur durch einen eigenen Stamm, den 推量形 (*suiryoukei*), beschrieben:

(126) *Die Verbstämme nach McClain 1981:*

未然形	<i>mizenkei</i>	Indefinitform
連用形	<i>renyoukei</i>	Konjunkionalform
終止形 bzw. 連体形	<i>syuusikei</i> bzw. <i>rentaikei</i>	Finalform bzw. Attributivform
假定形	<i>kateikei</i>	Konditionalform
命令形	<i>meireikei</i>	Imperativform
推量形	<i>suiryoukei</i>	Futurform

2.2.1 Flexionsklassen

Die japanische Schulgrammatik⁵¹ unterscheidet zwischen *vokalischen* (一段活用, *itidañ katuyou*), *konsonantischen* (五段活用, *godañ katuyou*) und *unregelmäßigen Verben* (不規則活用, *fukisoku katuyou*)⁵². Vokalische Verben unterteilen sich weiter in *iru*- und *eru*-Verben; konsonantische Verben können hinsichtlich des letzten Konsonanten ihres Stammes — dem sogenannten *Stammkonsonanten* — weiter unterteilt werden.⁵³ Unregelmäßige Verben bilden die Klassen der *suru*- (サ行変格活用, *sagyou henkaku katuyou*), *ku-ru*- (カ行変格活用, *kagyou henkaku katuyou*), *nasaru*- und *gozaru*-Verben (不規則五段活用, *fukisoku godañ katuyou*). Zusammengefaßt ergibt sich das folgende Klassifikationsschema:

(127) *Verbkategorien*⁵⁴:

1. vokalische Verben (一段活用, *itidañ katuyou*)
 - a. *iru*-Verben (上一段活用, *kamiitidañ katuyou*)

z.B.: 見る *miru* sehen
 いる *iru* sein
 ...
 - b. *eru*-Verben (下一段活用, *simoitidañ katuyou*)

z.B.: 食べる *taberu* essen
 出る *deru* verlassen
 ...
2. konsonantische Verben (五段活用, *godañ katuyou*)

z.B.: ある *aru* sein
 書く *kaku* schreiben
 読む *yomu* lesen
 言う *iu* sagen
 ...
3. unregelmäßige Verben
 - a. *suru*-Verben (サ行変格活用, *sagyou henkaku katuyou*)

する *suru* tun, machen

51 Die folgende Systematik bezieht sich auf McClain 1981 und weicht in einigen Punkten von anderen Beschreibungen ab.

52 Die Verben der 不規則活用 (*fukisoku katuyou*) werden meist nicht, wie hier geschehen, in einer Kategorie zusammengefasst, sondern durch サ行変格活用 (*sagyou henkaku katuyou*) und カ行変格活用 (*kagyou henkaku katuyou*) getrennt behandelt. Die Verben der 不規則五段活用 (*fukisoku godañ katuyou*) hingegen werden der 五段活用 (*godañ katuyou*) zugeordnet, und mit einer Beschreibung ihrer Unregelmäßigkeiten versehen. Die Formen des Honoratives *-masu* werden getrennt behandelt.

53 Das Verb *kak-u* (書く, *schreiben*) beispielsweise hat den Stammkonsonanten *k*. Folgende Stammkonsonanten kommen vor: *k*, *g*, *s*, *t*, *n*, *b*, *m*, *r* und *w*.

- b. *kuru*-Verben (力行変格活用, *kagyou henkaku katuyou*)
 来る *kuru* kommen
- c. *nasaru*-Verben (不規則五段活用, *fukisoku godaŋ katuyou*)
 なさる *nasaru* tun, machen (Honorativ von *suru*)
 下さる *kudasaru* geben (Honorativ von *kureru*)
 おっしゃる *ossyaru* sagen (Honorativ von *iu*)
 いらっしゃる *irassyaruru* sein, gehen, kommen
 (Honorativ von *iru*, *iku*, *kuru*)
- d. *gozaru*-Verben (不規則五段活用, *fukisoku godaŋ katuyou*)
 ごさる *gozaru* sein (höflich für *aru*)

2.2.2 Verbstämme

In der japanischen Schulgrammatik werden sechs Verbstämme unterschieden, die von jedem Verb gebildet werden können. Die Funktion der Verbstämme ist heterogen: Ein Teil beschreibt morphologisch abgeschlossene Formen, die direkt in der Syntax benutzt werden können; ein anderer Teil dient dem Anschluß sogenannter *Hilfsverben* (助動詞, *zyodousi*).⁵⁵

Die Hilfsverben der traditionellen japanischen Grammatik entsprechen den Suffixmorphemen, die wir am Anfang dieses Kapitels (S. 62) als charakteristisches Mittel der morphologischen Modifikation agglutinierender Sprachen bezeichnet haben.

Die folgende Übersicht stellt die japanischen Verbstämme mit einigen ihrer jeweiligen Funktionen vor:

(128) Verbstämme und Funktionen:

⁵⁴ Dieses Schema ist weitgehend identisch zu dem von McClain 1981.

⁵⁵ Lewin 1990 übersetzt 助動詞 (*zyodousi*) durch *Verbalsuffix*.

Stamm	syntaktische Funktion	Basis für ⁵⁶
mizeñkei	ϕ	<i>Negation</i> <i>negiertes Partizip</i> <i>Passiv</i> <i>Kausativ</i> <i>Kausativ-Passiv</i> ...
reñyoukei	<i>Nominalisierung</i> <i>Koordination</i>	<i>Verbkomposita</i> <i>Honorativ</i> <i>Voluntativ</i> ...
syuusikei reñtaikei	<i>Satzschlußform</i> <i>Adnominalform</i>	ϕ ϕ
kateikei	ϕ	<i>Konditional</i>
meireikei	<i>Imperativ</i>	ϕ
suiryokei	<i>Futur</i>	ϕ

2.2.3 Hilfsverben (助動詞, *zyodousi*)

Wie im letzten Abschnitt gesagt, dienen die Verbstämme teilweise der Bildung von Formen, die direkt in der Syntax verwendet werden können, und teilweise als Basis der sogenannten *Hilfsverben* (助動詞, *zyodousi*). Diese Hilfsverben werden wiederum einer Kategorie zugeordnet und können ihrerseits Stammformen bilden und damit als Basis weiterer Suffixe bzw. Hilfsverben fungieren. Die Hilfsverben der traditionellen Grammatik dienen damit der Modellierung des charakteristischen Mittels der japanischen Verbmodifikation: der Agglutination.

Die Bildung der Stammformen und ihre Funktion in der Syntax bzw. als Basis für die Agglutination von Hilfsverben soll beispielhaft anhand einiger Tabellen dargestellt werden.⁵⁷ Diese teilen sich in fünf Spalten:

1. *Name des Stammes.*

⁵⁶ Um auf die späteren Kapitel vorzubereiten, werden die japanischen Hilfsverben in dieser Übersicht nicht mit ihren traditionellen Namen, sondern mit den Bezeichnungen aus Rickmeyer 1995 benannt: Diese werden auch in der hier entwickelten morphologischen Beschreibung verwendet.

⁵⁷ Die Tabellen sind in leicht veränderter Form McClain 1981 entnommen.

2. *Stammkern*⁵⁸ — Alle Stämme haben einen gemeinsamen ersten Teil. Diesen bezeichnen wir als *Kern* des jeweiligen Stammes.
3. *Stammformativ*⁵⁸ — Der Stamm ergibt sich aus der Konkatination von *Stammkern* und *Stammformativ*. Da der Kern bei allen Stämmen gleich ist, ist das *Stammformativ* das eigentliche Charakteristikum des jeweiligen Stammes.
4. *Suffix*⁵⁸ — In einigen Fällen entsteht aus Kern und Stammformativ bereits eine morphologisch abgeschlossene Form; die meisten Formen werden allerdings mittels eines Suffixes gebildet, das oft wiederum flektiert werden kann.
5. *Name Verbform* — Die Form aus Stamm und Suffix wird durch den hier angegebenen Namen bezeichnet.

Die erste Tabelle beschreibt die Stammformen eines vokalischen Verbes.

(129) *Konjugation des vokalischen Verbes* 見る (*miru*; sehen):

Stamm			abgeleitete Formen	
Name des Stammes	Kern ⁵⁸	Stammformativ ⁵⁸	Suffix ⁵⁸	Name der abgeleiteten Form
mizeñkei	<i>mi</i>		-nai -zu -rareru -saseru, -sasu -saserareru ...	Negation negiertes Partizip Passiv Kausativ Kausativ-Passiv ...
reñyoukei	<i>mi</i>		-masu -tai ...	Honorativ Voluntativ ...
syuusikei rentaikei	<i>mi</i>	<i>ru</i>		Satzschlußform Adnominalform
kateikei	<i>mi</i>	<i>re</i>	-ba	Konditional
meireikei	<i>mi</i>	<i>ro</i> <i>yo</i>		Imperativ
suiryokei	<i>mi</i>	<i>you</i>		Futur / Volitional

Die folgende Tabelle zeigt die Stammformen der konsonantischen Verben anhand des Verbes 書く (*kaku*; schreiben):

⁵⁸ Lewin 1990 bezeichnet den *Stammkern* als *Stammsilbe*, das *Stammformativ* als *Endungssilbe* und die *Suffixe* als *silbisches Endungsmorphem*, *Verbalsuffix* etc.

(130) Konjugation des konsonantischen Verbes 書く (*kaku*; schreiben):

Stamm			abgeleitete Formen	
Name des Stammes	Kern ⁵⁸	Stamm-formativ ⁵⁸	Suffix ⁵⁸	Name der abgeleiteten Form
mizenkei	<i>ka</i>	<i>ka</i>	-nai -zu -reru -seru, -su -serareru ...	Negation negiertes Partizip Passiv Kausativ Kausativ-Passiv ...
reñyoukei	<i>ka</i>	<i>ki</i>	-masu -tai ...	Honorativ Voluntativ ...
syuusikei reñtaikei	<i>ka</i>	<i>ku</i>		Satzschlußform Adnominalform
kateikei	<i>ka</i>	<i>ke</i>	-ba	Konditional
meireikei	<i>ka</i>	<i>ke</i>		Imperativ
suiryokei	<i>ka</i>	<i>kou</i>		Futur / Volitional

Als Beispiel für die Variationen bei der Bildung der Stammformen unregelmäßiger Verben wird das Verb する (*suru*; tun, machen) ausgewählt:

(131) Konjugation des unregelmäßigen Verbes する (*suru*; tun, machen):

Stamm			abgeleitete Formen	
Name des Stammes	Kern ⁵⁸	Stamm-formativ ⁵⁸	Suffix ⁵⁸	Name der abgeleiteten Form
mizenkei		<i>si</i> <i>se</i> <i>sa</i>	-nai -zu -reru -seru, -su -serareru ...	Negation negiertes Partizip Passiv Kausativ Kausativ-Passiv ...
reñyoukei		<i>si</i>	-masu -tai ...	Honorativ Voluntativ ...
syuusikei reñtaikei	<i>su</i>	<i>ru</i>		Satzschlußform Adnominalform
kateikei	<i>su</i>	<i>re</i>	-ba	Konditional
meireikei	<i>se</i> <i>si</i>	<i>yo</i> <i>ro</i>		Imperativ
suiryokei	<i>si</i>	<i>you</i>		Futur / Volitional

Wie die Tabelle zeigt, bestehen Mizenkei und Reñyoukei von する (*suru*) aus nur einer Silbe und lassen sich daher nicht in Kern und Stammformativ trennen; Mizenkei und Meireikei können aufgrund der Unregelmäßigkeiten von する (*suru*) bei der Bildung der Formen nicht eindeutig bestimmt werden.

2.2.4 Phonologische Transformationen

Nicht alle Verbformen lassen sich als Verbindung einer der traditionellen Verbstämme mit einem Hilfsverb beschreiben. Bei der Bildung des Perfektes oder des Partizips beispielsweise kommt es zu einer so engen Verbindung zwischen Stammform und Hilfsverb, daß unter Voraussetzung des traditionellen Beschreibungssystems die Konkatenation zur Erklärung der Formen nicht hinreichend ist:

(132) *Phonologische Transformationen bei der Bildung der onbin-form:*

a. *vokalische Verben:*

<i>Verb</i>	<i>Perfekt</i>	<i>Partizip</i>
<i>iru</i>	<i>ita</i>	<i>ite</i>
<i>miru</i>	<i>mita</i>	<i>mite</i>
<i>taberu</i>	<i>tabeta</i>	<i>tabete</i>

b. *konsonantische Verben:*

<i>Verb</i>	<i>Perfekt</i>	<i>Partizip</i>
<i>yomu</i>	<i>yoñda</i>	<i>yoñde</i>
<i>yobu</i>	<i>yoñda</i>	<i>yoñde</i>
<i>sinu</i>	<i>siñda</i>	<i>siñde</i>
<i>iu</i>	<i>itta</i>	<i>itte</i>
<i>matu</i>	<i>matta</i>	<i>matte</i>
<i>wakaru</i>	<i>wakatta</i>	<i>wakatte</i>
<i>kaku</i>	<i>kaita</i>	<i>kaite</i>
<i>isogu</i>	<i>isoida</i>	<i>isoide</i>
<i>dasu</i>	<i>dasita</i>	<i>dasite</i>

c. *Ausname:*

<i>Verb</i>	<i>Perfekt</i>	<i>Partizip</i>
<i>iku</i>	<i>itta</i>	<i>itte</i>

d. *unregelmäßige Verben:*

<i>Verb</i>	<i>Perfekt</i>	<i>Partizip</i>
<i>kuru</i>	<i>kita</i>	<i>kite</i>
<i>suru</i>	<i>sita</i>	<i>site</i>

Diese Formen werden deshalb als Resultat zweier Prozesse erklärt:

1. Konkatenation von Stamm und Suffix;
2. Verschleifung von Stamm und Suffix durch eine phonologische Regel:

- (133) Modellierung der onbin-Form in der traditionellen Grammatik,
dargestellt am Beispiel des Perfektes von 書く (*kaku*; schreiben):

$$kaki + ta \xrightarrow{\text{Konkatenation}} kaki-ta \xrightarrow{\text{Transformation}} kaita$$

2.2.5 Ausnahmen

Die Systematik der traditionellen Verbklassen beruht auf einem Kompromiß zwischen dem Streben nach Einfachheit und Genauigkeit der Beschreibung: Einige Verben bilden einzelne Formen relativ zu diesem Schema unregelmäßig und lassen sich daher nicht eindeutig einer der oben angegebenen Klassen zuordnen. Diese abweichenden Formen werden als *Ausnahmen* bei den betreffenden Verben bzw. Formen explizit aufgeführt. 行く (*iku*; gehen) z.B. verhält sich nur bei der Bildung der Onbin-Formen — wie dem Partizip oder Perfekt — anders als die anderen Verben seiner Gruppe:

(134)	<i>Verb</i>	<i>Partizip</i>
<i>regelmäßig:</i>	書く <i>kaku</i>	書いて <i>kaite</i>
<i>unregelmäßig:</i>	行く <i>iku</i>	行つて <i>itte</i>

Wird auf diese Unregelmäßigkeit hingewiesen — beispielsweise als Ausnahme bei der Beschreibung des Partizips⁵⁹ — kann es trotzdem wie ein regelmäßiges Verb behandelt werden. Durch diese Art der Behandlung von Unregelmäßigkeiten gelingt es, die Verbklassifikation und die Regeln zur Ableitung der Formen systematisch und übersichtlich zu halten.

2.3 Die überarbeitete Verbmorphologie

Im letzten Abschnitt wurde das morphologische System der traditionellen japanischen Schulgrammatik vorgestellt. Durch die Kombination von systematischer Beschreibung, Ausnahmeregelungen und phonologischen Regeln gelingt es, dieses System für den Benutzer einfach und übersichtlich zu halten. Was jedoch aus der menschlichen Perspektive als “einfach” und “übersichtlich” erscheint, erweist sich aus der Sicht der formalen Linguistik als Problem: Ausnahmen und mehrstufige Beschreibungen sind schwerer zu formalisieren als Komplexe, für den Menschen unübersichtliche, Systematiken.

Es liegt daher nahe, das traditionelle System hinsichtlich dieser “formalen” Kriterien zu überarbeiten.

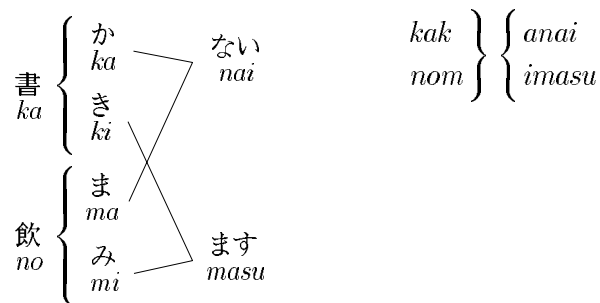
⁵⁹ Das unregelmäßige Partizip von 行く (*iku*) kann natürlich auch als Ausnahme bezüglich der phonologischen Verschleifung bei der Bildung der onbin-Form interpretiert werden.

2.3.1 Kana-Schreibweise vs. phonologische Umschrift

Das traditionelle System verwendet die japanische Kana-Schreibweise bei der Beschreibung der Verbformen. Ein Kana-Zeichen stellt mehrere Phoneme dar. Die Verbformen können daher nicht immer an ihren durch systematischen Vergleich ermittelten Morphemgrenzen segmentiert werden. Daraus resultiert — verglichen mit einem System, das eine Umschrift verwendet — eine unnötig große Anzahl von Allomorphen:

(135) *Vergleich der Ableitung von Verbformen:*

a. *Kana-Kanji-Mischschrift:* b. *phonologische Umschrift:*



Um zu einer einfacheren Beschreibung zu kommen, gehen wir den Weg, der in (135b) vorgeschlagen wird, und legen eine geeignete Umschrift⁶⁰ zugrunde.

2.3.2 Eliminierung der phonologischen Transformationen

Durch die Verwendung der phonologischen Transformationen gelingt es der traditionellen Grammatik, das System der Verbstämme einfach zu halten. Ein solches zweistufiges System aus Komposition und Transformation ist jedoch aus formaler Sicht schwieriger zu handhaben als ein komplexes System von Verbstämmen und morphologischen Bauteilen.⁶¹ In unserer Beschreibung werden daher auch die Formen, die traditionell auf phonologische Verschleifungen zurückgeführt werden, durch die Konkatenation von Verbstämmen und Endungen erklärt.

⁶⁰ Wir verwenden eine leicht modifizierte Form der Nipponsiki-Umschrift (vgl. Bollmann 1996), die eine eindeutige Abbildung von Kana-Kanji-Mischschrift auf die Umschrift erlaubt. Außerdem werden Teile der Wörter, die in japanischen Texten als Kanji geschrieben werden, mit den entsprechenden Kanji annotiert. Damit sind alle Formen der Verschriftlichung des Japanischen darstellbar und können ineinander überführt werden.

⁶¹ Unter dem Namen *Two-Level Morphology* entwickelte der finnische Computerlinguist Kimmo Koskeniemi einen sehr effizienten und inzwischen weit verbreiteten Formalismus zur Beschreibung von morphophonologischen Transformationen (Koskeniemi 1983). Im Gegensatz zu der klassischen generativen Phonologie (Chomsky und Halle 1968) werden in Koskeniemis Formalismus die Transformationsregeln nicht sequentiell, sondern simultan angewandt, und können daher durch einen regulären Automaten berechnet werden.

Die Form 書いた (*kaita*; geschrieben haben) beispielsweise wird in der traditionellen Grammatik als Produkt der Konkatenation von reñyoukei (連用形) und Perfekt-Suffix *ta* erklärt, an die sich eine morphologische Verschleifung (*onbin-Verschleifung*) von *ki-t* zu *it* anschließt:

$$(136) \text{ } kaki + ta \xrightarrow{\text{Konkatenation}} kaki-ta \xrightarrow{\text{Transformation}} kaita$$

In dem hier entwickelten System hingegen wird auch diese Form unter Voraussetzung nur einer Ebene alleine durch Konkatenation erklärt:

$$(137) \text{ } ka + it + a \xrightarrow{\text{Konkatenation}} kaita$$

Nach traditioneller Auffassung entsteht diese Form durch morphologische *und* phonologische Operationen. Dieser Ansicht zufolge handelt es sich bei den Formprimitiven *ka*, *it* und *a* also um *morphophonologische* Elemente als Teile einer *morphophonologischen* Beschreibung. Da die Grenze zwischen Morphologie und Phonologie aber fließend ist und viele Phänomene einfach aus pragmatischen Überlegungen der einen oder anderen Ebene zugeordnet werden, sprechen wir meistens einfach von *morphologischen* Elementen und einer *morphologischen* Beschreibung.

2.3.3 Eliminierung der Ausnahmen zugunsten einer komplizierteren Systematik

Wie oben erwähnt, motiviert sich die Verwendung von Ausnahmen aus dem Bestreben, die Systematik aus Verbklassen und Stämmen einfach zu halten. Aus formaler Perspektive empfiehlt sich wiederum ein entgegengesetztes Vorgehen: Eine komplexe Systematik läßt sich einfacher modellieren als eine große Zahl von Ausnahmen.

Motiviert wurde Koskeniemi durch die Probleme herkömmlicher Systeme bei der Beschreibung der finnischen Verbmorphologie (vgl. Koskeniemi 1985b, Koskeniemi und Church 1988 und Koskeniemi 1985a). Da die finnische Verbmorphologie große Ähnlichkeiten mit der japanischen hat — auch das Finnische ist eine agglutinierende Sprache — könnte eine Two-Level-Beschreibung der japanischen Verbmorphologie eine interessante Variante zu dem hier vorgeschlagenen Modell darstellen.

Interessanter erscheint mir allerdings die Verwendung dieses Formalismus zur Beschreibung der phonologischen Komponente. Selbst die Onbin-Verschleifungen, die in der traditionellen japanischen Grammatik als phonologische Prozesse interpretiert werden, sind an spezielle morphologische Mechanismen gebunden, und treten außerhalb dieser nicht auf. Auch bei der Integration einer Two-Level-Komponente in ein sprachverarbeitendes System erscheint mir daher die hier vorgeschlagene „morphologische“ Modellierung adäquater.

In unserem Modell werden die meisten Ausnahmen der traditionellen Grammatik durch eigene Verbklassen und spezielle Regeln modelliert. Nur ein relativ kleiner Teil der unregelmäßigen Formen der japanischen Schulgrammatik wird direkt ins Lexikon eingetragen.

2.3.4 Von der traditionellen Grammatik zum überarbeiteten System

Das in dieser Arbeit entwickelte System ist durch die folgenden Punkte charakterisiert:

1. Verwendung einer phonologischen Umschrift;
2. Eine einzige morphophonologische Beschreibungsebene;
3. Konkatenation bzw. Segmentierung als einzige Grundoperation auf morphophonologischen Formen;
4. Segmentierung der Verbformen in Formprimitiven auf morphophonologischer Ebene;
5. Ein fein gegliedertes, hierarchisches System von Verbkategorien;
6. Ein fein gegliedertes, hierarchisches System von Verbstämmen.

Bei der Entwicklung einer Beschreibung, die diese Kriterien⁶² erfüllt, können wir neben der eben beschriebenen traditionellen Grammatik auf die “Japanische Morphosyntax” von Jens Rickmeyer (Rickmeyer 1983, Rickmeyer 1984, Rickmeyer 1994, Rickmeyer 1995) zurückgreifen. Rickmeyer verwendet, wie wir später in diesem Abschnitt sehen werden, ebenfalls eine phonologische Umschrift und reduziert die Zahl der Verbstämme auf zwei. Zur Ableitung der Onbin-Formen werden in seinem System allerdings ähnliche phonologische Regeln verwendet wie in der traditionellen japanischen Grammatik.

Durch die in der vorliegenden Arbeit vorgeschlagene Eliminierung der traditionellen phonologischen Regeln und die Annahme von fünf auseinander ableitbaren, hierarchisch gegliederten Stämmen entsteht daher ein System, das sich bezüglich Anzahl, Art und Ableitung der Stämme und bezüglich der Form der Suffixmorpheme sowohl von der traditionellen Grammatik als auch von Rickmeyers Beschreibung grundlegend unterscheidet. Die Unterteilung und Nomenklatur der Morpheme hingegen stimmt mit der von Rickmeyer 1995 vorgeschlagenen überein.

⁶² Die Idee, morphologische Beschreibungen alleine auf die Konkatenation fein klassifizierter Formprimitiven aufzubauen, ist auch in der Computerlinguistik nicht neu: MORPHIX (Finkler und Neumann 1986, Finkler und Neumann 1988) beispielsweise, eine klassifikationsbasierte Beschreibung der Morphologie des Deutschen, beruht auf einer ähnlichen Grundentscheidung. Allerdings werden in MORPHIX morphologische Regeln und Strategien nicht deklarativ beschrieben, sondern direkt in Form von Algorithmen kodiert.

2.3.4.1 Verbkategorien

Die Ausarbeitung einer formalen grammatischen Beschreibung ist ein aufwendiger Prozeß: Segmentierung und Klassifizierung bedingen sich gegenseitig, und jede Veränderung einer Komponente zieht Veränderungen der anderen nach sich — ein sich selbst anstoßender, zirkulärer Prozeß. Der Übergang von dem traditionellen System zu dem neu entwickelten Modell kann daher nur ansatzweise und in Form einiger repräsentativer Schritte beschrieben werden.

Zunächst soll das System der traditionellen Verbkategorien formalisiert werden. Dazu führen wir eine formale Nomenklatur für die japanischen bzw. deutschen Namen der Verbklassen ein:

(138) *Eine formale Nomenklatur für die traditionellen Verbkategorien:*

1. v_v *vokalische Verben* (一段活用, *itidañ katuyou*)
(die Unterscheidung in *iru-* und *eru-*Verben ist für die folgenden Betrachtungen nicht interessant.)
2. v_c *konsonantische Verben* (五段活用, *godañ katuyou*)
Für die Subkategorisierung entsprechend des Stammkonsonanten werden die folgenden Symbole verwendet:

<i>Symbol</i>	<i>Stammvokal</i>	<i>Beispiel</i>
v_{c_k}	<i>k</i>	書く <i>ka.k-u</i> schreiben
v_{c_g}	<i>g</i>	泳ぐ <i>oyo.g-u</i> schwimmen
v_{c_s}	<i>s</i>	話す <i>hana.s-u</i> sprechen
v_{c_t}	<i>t</i>	待つ <i>ma.t-u</i> warten
v_{c_n}	<i>n</i>	死ぬ <i>si.n-u</i> sterben
v_{c_b}	<i>b</i>	呼ぶ <i>yo.b-u</i> rufen
v_{c_m}	<i>m</i>	読む <i>yo.m-u</i> lesen
v_{c_r}	<i>r</i>	分かる <i>waka.r-u</i> verstehen
v_{c_w}	<i>w</i>	笑う <i>war.a-u</i> lachen

3. v_i *unregelmäßige Verben*
 - a. $v_{i_{suru}}$ *suru-Verben* (サ行変格活用, *sagyou henkaku katuyou*)
 - b. $v_{i_{kuru}}$ *kuru-Verben* (カ行変格活用, *kagyou henkaku katuyou*)
 - c. $v_{i_{nasaru}}$ *nasaru-Verben* (不規則五段活用, *fukisoku godañ katuyou*)
 - d. $v_{i_{gozaru}}$ *gozaru-Verben* (不規則五段活用, *fukisoku godañ katuyou*)

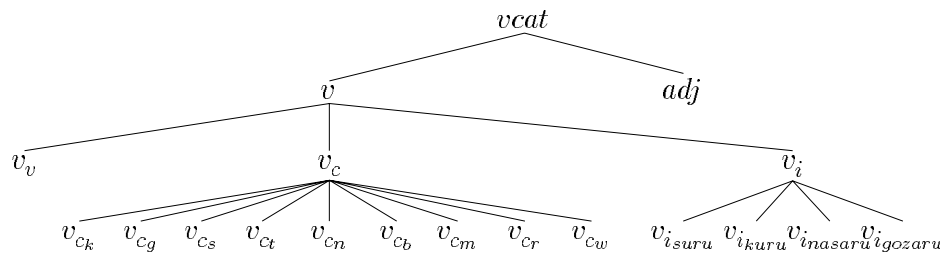
Verben allgemein kürzen wir mit *v* ab, Adjektive mit *adj*. Für Verben und Adjektive zusammen wird *vcat* verwandt:

(139) *Abkürzungen für Verben und Adjektive:*

1. *vcat* Verben und Adjektive (用言, *youngeñ*)
 - a. *adj* Adjektive (形容詞, *keiyousi*)
 - b. *v* Verben (動詞, *dousi*)

Mit Hilfe der formalen Nomenklatur können die Verbklassen der japanischen Schulgrammatik als Klassenhierarchie dargestellt werden:

(140) *Flexionsklassen der Verben in der japanischen Schulgrammatik:*



Von der klassischen Beschreibung zu einer Defaultbeschreibung

In der traditionellen Beschreibung werden unregelmäßige Verben entweder den Verbklassen zugeordnet, denen sie am ähnlichsten sind, oder — unter Aufgabe der Ähnlichkeiten mit anderen Klassen — durch eine eigene Klasse modelliert. Im ersten Fall werden Unterschiede im Verhalten von Verben verwischt, im zweiten Fall werden Ähnlichkeiten zwischen den Verbklassen ignoriert, so daß gleiches Verhalten an verschiedenen Stellen erklärt werden muß und Redundanzen entstehen.

Eine elegantere und weniger redundante Darstellung ergibt sich, wenn die Ähnlichkeiten der Verben bezüglich ihrer Flexion durch hierarchische Klassen explizit gemacht und bei der Beschreibung ihres Verhaltens ausgenutzt werden:

1. Alle Verben werden durch eine Kategorie beschrieben;
2. Die Kategorien werden nach der Ähnlichkeit bei der Formbildung und der Anzahl der durch sie beschreibenden Verben in Ober- und Unterklassen hierarchisch geordnet.

Einer Kategorienhierarchie, die nach diesen Kriterien entwickelt wird, liegt implizit ein Denken in Defaults zugrunde: Das Verb 行< (*i.k-u*; gehen, fahren) beispielsweise verhält sich weitgehend wie ein *k*-konsonantisches Verb. Nur bei der Bildung der onbin-Formen zeigt es sich abweichend:⁶³

⁶³ vgl. Abschnitt 1.7 *Defaults und Defaultvererbung*, S. 58.

- (141) 1. Die Menge der *k-konsonantische Verben* bilden das Partizip durch Suffigierung von *ite* (Default)
2. das Verb 行く (*iku*; gehen) bildet das Partizip durch Suffigierung von *tte*

Den eben angegebenen Kriterien zufolge müßte es als Instanz eines Untertyps $v_{c_{iku}}$ der *k-konsonantischen Verben* modelliert werden:

$$(142) \quad \begin{array}{c} v_{c_k} \\ | \\ v_{c_{iku}} \end{array}$$

Diese Art der Kategorienbildung geht davon aus, daß Untertypen alle Regeln ihrer Supertypen erben — es sei denn, abweichende Regeln (“Ausnahmeregeln”) überschreiben diese:

(143) *Defaultregeln zur Bildung des Partizips:*

- a. Verben der Kategorie v_{c_k} Stamm + *ite* (Default)
- b. Verben der Kategorie $v_{c_{iku}}$ Stamm + *tte* (überschreibt Default)

Diese Art, Ähnlichkeiten und Abweichungen durch hierarchische Klassen und Regelüberschreibungen zu modellieren, entspricht den Mechanismen von Defaultlogiken. Offensichtlich sind Defaultbeschreibungen also intuitiv — sie lassen sich leicht verstehen und ableiten — und führen durch die Vermeidung von Redundanzen zu eleganteren Modellen.

In einem ersten Schritt arbeiten wir deshalb das traditionelle System in eine Defaultbeschreibung um. Folgende Änderungen der Kategorienhierarchie bieten sich dafür an:

Überarbeitung der Kategorienhierarchie:

Zuerst führen wir neue Verbklassen ein. Bestimmte Formen, die bisher als Ausnahmen behandelt wurden, lassen sich dadurch systematisch beschreiben:

- Das Verb 行く (*iku*; gehen, fahren) verhält sich, von der Bildung der onbin-Formen abgesehen, wie ein v_{c_k} -Verb. Zu seiner Beschreibung führen wir die Kategorie $v_{c_{iku}}$ als Unterkategorie von v_{c_k} ein;
- Das Honorativsuffix *masu* wird als neue Unterkategorie der unregelmäßigen Verben dargestellt;

- Hilfsverb ある (*aru*; sein) und Kopulaverb だ (*da*; sein) ähneln weitgehend den v_{cr} -Verben, und werden durch entsprechende Untertypen repräsentiert.

Eine Verbkategorie fällt weg:

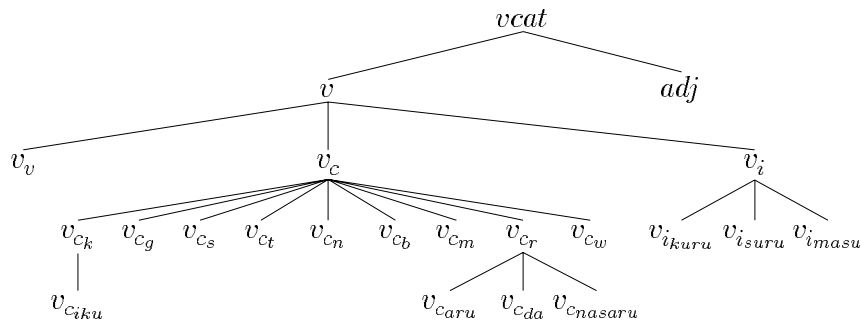
- Das Verb ごさる (*gozaru*; sein; formale Variante von ある) kommt im modernen Japanischen nur noch in seiner honorativen Form ごぜいます (*gozaimasu*) vor. Es läßt sich daher als v_{imasu} -Verb beschreiben, und wir können auf die Kategorie $v_{igozaru}$ verzichten.

Eine Verbkategorie wird neu eingeordnet:

- Die Kategorie der nasaru-Verben ($v_{inasaru}$) wird “umgehängt”: Die Flexion der $v_{inasaru}$ -Verben entspricht weitgehend der der r-konsonantischen Verben. Wir können $v_{inasaru}$ daher adäquater als Untertyp von v_{cr} behandeln, und nennen $v_{inasaru}$ entsprechend in $v_{cnasaru}$ um.

Zusammenfassend ergibt sich die folgende modifizierte Hierarchie:

(144) *Flexionsklassen der Verben in einer defaultbasierten Beschreibung:*



Defaultlogiken ermöglichen eine elegantere Form der linguistischen Beschreibung als monotone Logiken.⁶⁴ Sie sind jedoch theoretisch und algorithmisch schwerer zu handhaben und für die Verwendung in computerlinguistischen Systemen daher oft zu langsam. In dieser Arbeit wird deshalb eine monotone Beschreibung angestrebt.

⁶⁴ Vgl. auch meine Studienarbeit (Bollmann 1997), in der ein defaultbasierter Ansatz für ein Modell der japanischen Verbmorphologie vorgestellt wird.

Von der Defaultbeschreibung zu einer monotonen Beschreibung

Um aus der Hierarchie (144) eine Hierarchie abzuleiten, die sich für eine monotone Beschreibung eignet, erweitern wir diese derart, daß jede Verbkategorie durch ein minimales Kategoriensymbol repräsentiert wird. Die Idee, die diesem Vorgehen zugrunde liegt, läßt sich wiederum am besten anhand eines Beispiels darstellen.

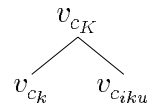
Das Verhalten des Verbes 行く (*iku*; gehen, fahren), das wir als Modell zur Darstellung der Defaultlogik benutzt haben, beschreiben wir nunmehr monoton:

(145) *Monotone Beschreibung der Bildung des Partizips von 行く (*iku*):*

1. Die Menge der *k*-konsonantischen Verben ohne das Verb 行く (*iku*; gehen) bilden das Partizip durch Suffigierung von *ite*
2. das Verb 行く (*iku*; gehen) bildet das Partizip durch Suffigierung von *tte*

Formal können diese Regeln entweder mit der Hilfe von logischen Operatoren oder Mengenoperatoren formuliert werden ($v_{c_k} \wedge \neg v_{c_{iku}}$), oder durch Einführung einer neuen Verbkategorie für die *k*-konsonantischen Verben ohne 行く. Der letztgenannte Ansatz führt zu folgenden Regeln:

(146) a. *Hierarchie der Verbkategorien:*



b. *Regeln zur Bildung des Partizips:*

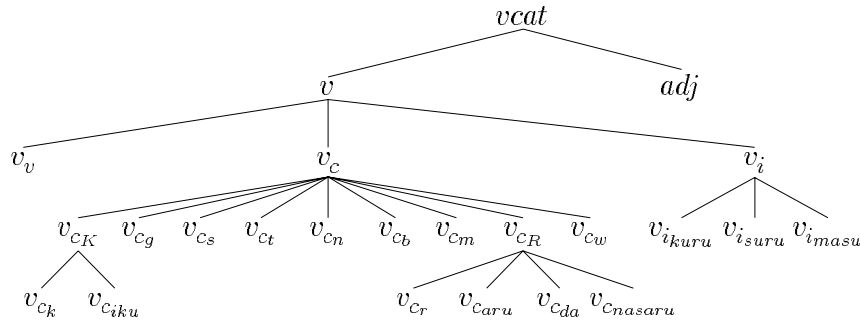
- i. Verben der Kategorie v_{c_k} Stamm + *ite*
- ii. Verben der Kategorie $v_{c_{iku}}$ Stamm + *tte*

Regeln, die sowohl für v_{c_k} -Verben als auch für $v_{c_{iku}}$ -Verben gelten, lassen sich als Regeln zum Typ v_{c_K} formulieren. Regeln hingegen, die nur für v_{c_k} - bzw. $v_{c_{iku}}$ -Verben gelten, werden relativ zu dem entsprechenden Untertyp formuliert.

Modifizieren wir die Kategorienhierarchie in der gleichen Weise auch bezüglich

der r-konsonantischen Verben (v_{cr}), ergibt sich das folgende endgültige Kategorienschema:

(147) *Flexionsklassen der Verben:*



Abkürzungen

Ein wesentlicher Vorteil von Kategorienhierarchien besteht in der Möglichkeit, Ähnlichkeiten im Verhalten durch hierarchische Klassifizierung darzustellen: Gilt eine Regel für alle konsonantischen Verben, muß sie nur einmal an der entsprechenden Kategorie ausgedrückt werden, um kenntlich zu machen, daß sie auch für alle Unterkategorien gültig ist.

Nicht jede Ähnlichkeit im morphologischen Verhalten läßt sich jedoch durch eine allgemeinste Kategorie unseres Schemas ausdrücken. Beispielsweise existiert keine Kategorie, die es gestattet, nur v_{cn} -, v_{cb} - und v_{cm} -Verben gemeinsam anzusprechen. Um dennoch Klassen dieser Art bilden zu können ohne das Kategorienschema unnötig komplex gestalten zu müssen, erlauben wir die Verwendung von Disjunktionen bei der Formulierung morphologischer Regeln:

$$(148) \quad v_{cn} \vee v_{cb} \vee v_{cm}.$$

Noch übersichtlicher wird die Beschreibung von morphologischen Regeln durch die Verwendung von Abkürzungen:

(149) Abkürzung: abgekürzte Form:

$$v_{cn,b,m} = v_{cn} \vee v_{cb} \vee v_{cm}$$

$$v_{v,i} = v_v \vee v_i$$

...

Neben der Disjunktion erlauben wir die Mengensubtraktion. Wegen der Endlichkeit der Kategorien läßt sich jede Disjunktion auch als Mengensubtraktion

formulieren und umgekehrt. Die Verwendung der Mengensubtraktion dient also nur der Übersichtlichkeit der Darstellung:

$$(150) \quad v \setminus v_i = (v_v \vee v_c \vee v_i) \setminus v_i = v_v \vee v_c$$

Auch die Mengensubtraktion kann abgekürzt werden:

$$(151) \quad \begin{array}{ll} \text{Abkürzung:} & \text{abgekürzte Form:} \\ v \setminus v_i & = v \setminus v_i \end{array}$$

Ein komplizierteres Beispiel, das den Unterschied in der Übersichtlichkeit der verschiedenen Notationen deutlicher macht:

$$(152) \quad \begin{aligned} v \setminus_{cda} &= v \setminus v_{cda} \\ &= (v_v \vee v_c \vee v_i) \setminus v_{cda} \\ &= v_v \vee v_c \setminus v_{cda} \vee v_i \\ &= v_v \vee (v_{c_K} \vee v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \vee v_{c_R} \vee v_{c_w}) \setminus v_{cda} \vee v_i \\ &= v_v \vee v_{c_K} \vee v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \vee v_{c_R} \setminus v_{cda} \vee v_{c_w} \vee v_i \\ &= v_v \vee v_{c_K} \vee v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \vee (v_{c_r} \vee v_{c_{aru}} \vee v_{cda} \vee v_{c_{nasaru}}) \setminus v_{cda} \vee v_{c_w} \vee v_i \\ &= v_v \vee v_{c_K} \vee v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \vee v_{c_r} \vee v_{c_{aru}} \vee v_{c_{nasaru}} \vee v_{c_w} \vee v_i \end{aligned}$$

Die folgenden Beispiele zeigen einige der am meisten verwendeten Abkürzungen:

(153) Abkürzung: abgekürzte Form:

$$\begin{aligned} v_{c_{n,b,m}} &= v_{c_n} \vee v_{c_b} \vee v_{c_m} \\ v_{c \setminus_{K,R,w}} &= v_c \setminus v_{c_{K,R,w}} \\ &= (v_{c_K} \vee v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \vee v_{c_R} \vee v_{c_w}) \setminus (v_{c_K} \vee v_{c_R} \vee v_{c_w}) \\ &= v_{c_g} \vee v_{c_s} \vee v_{c_t} \vee v_{c_n} \vee v_{c_b} \vee v_{c_m} \\ v_{c \setminus_{da,imasu}} &= v_{c \setminus_{da}} \vee v_{imasu} \\ &= \dots \end{aligned}$$

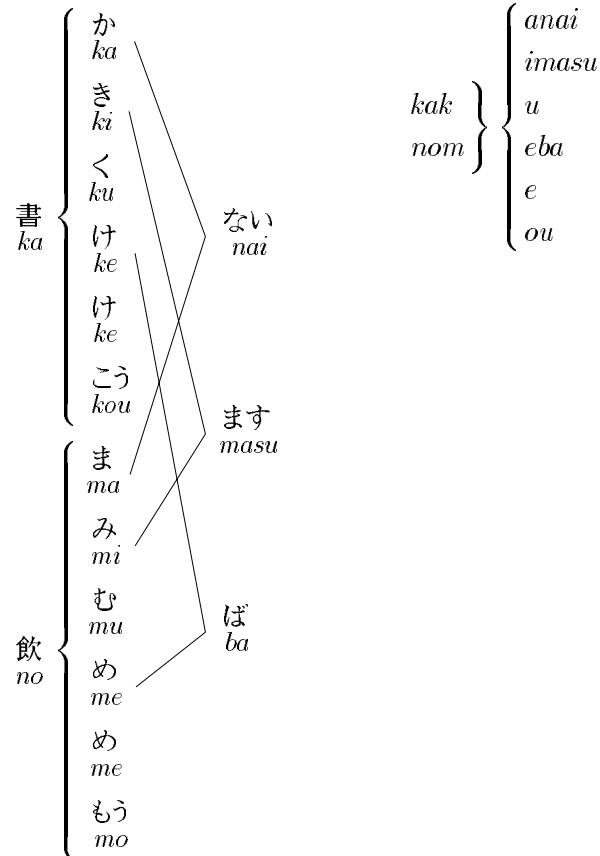
2.3.4.2 Verbstämme

Die japanische Schulgrammatik benutzt ein System von sechs Verbstämmen. Die Form dieser Stämme erklärt sich weitgehend aus der Verwendung der Kana-Schreibweise. Didaktisch motivierte Systeme, die eine phonologische Umschrift verwenden, kommen teilweise mit nur einem Stamm aus⁶⁵:

⁶⁵ vgl. beispielsweise Ishii et al. 1989.

(154) *Vergleich der Ableitung von Verbformen:*

- a. *japanische Schulgrammatik:*
(Kana-Kanji-Mischschrift)
- b. *didaktisches System ohne Verbstämme.*⁶⁵
(phonologische Umschrift)



Dieser Vergleich bezieht sich nur auf die regelmäßigen konsonantischen Verben. Werden andere Verbklassen mit berücksichtigt, wird sichtbar, daß sich die Komplexität, die bei der Beschreibung der morphologischen Variation unvermeidbar ist, bei einem System ohne Verbstämme nur verlagert: Die Verwendung von Verbstämmen gleicht die Unterschiede zwischen den verschiedenen Verbklassen teilweise aus. Ein Suffix kann daher in einem System *mit* Verbstämmen unter Umständen bei allen Verbklassen durch das gleiche Morph realisiert werden. In einem System *ohne* Verbstämme hingegen verlagert sich die Allomorphie auf die Suffixe: Unter Umständen muß in einem solchen System ein Suffix bei jeder Verbklasse durch ein anderes Allomorph realisiert werden.

Die Optimierung des Beschreibungsaufwandes unter Voraussetzung einer phonologischen Umschrift führt zu einer mittleren Zahl von Stämmen und Suffix-

Allomorphen. Rickmeyer 1995 beispielsweise reduziert die Anzahl der Stämme auf zwei:

(155) *Das System von Rickmeyer 1995:*

$$\begin{array}{l} \left. \begin{array}{l} kak \\ nom \end{array} \right\} \left\{ \begin{array}{l} anai \\ u \\ eba \\ e \\ ou \end{array} \right. \\ \\ \left. \begin{array}{l} kaki \\ nomi \end{array} \right\} \left\{ \begin{array}{l} masu \end{array} \right.$$

Er übernimmt den 連用形 (*reñyoukei*) der japanischen Schulgrammatik als sogenannte *Verbbasis* und nimmt als zweiten Stamm den allen traditionellen Stammformen gemeinsamen ersten Teil, den er als *Verbstamm* bezeichnet.

Die Motivation, den 連用形 (*reñyoukei*) zu übernehmen, besteht sicherlich in der Häufigkeit der Formen, die mit diesem Stamm gebildet werden. Wie oben erklärt, läßt sich durch die Verwendung unterschiedlicher Stämme die Anzahl der Suffixallomorphe verringern.

Die gleiche Motivation führt uns zu der Übernahme eines weiteren (leicht modifizierten) traditionellen Stammes: des 未然形 (*mizenkei*). Alle anderen Stämme lassen sich einfacher durch die Allomorphie der entsprechenden Suffixe erklären. Wir nennen den 連用形 (*reñyoukei*), da er meist auf *i* endet, *i-Stamm* (*i-stem*), und den 未然形 (*mizenkei*), der im Allgemeinen mit einem *a* abschließt, *a-Stamm* (*a-stem*); Rickmeyers *Verbstamm* bezeichnen wir — da er bei den konsonantischen Verben auf dem jeweiligen Stammkonsonanten auslautet — als *konsonantischen Stamm* (*cons-stem*).

Verbklasse und *Verbstamm* charakterisieren die morphologischen Formprimitiven unseres Modells: Der Lexikoneintrag eines Verbes ist gekennzeichnet durch die Verbklasse, der das Verb angehört, und die Stammform, in der es sich befindet. Suffixe hingegen lassen sich durch Verbklasse und Stamm der Formen, an die sie sich anschließen können, und Verbklasse und Stamm der Formen, die durch ihre Affigierung gebildet werden, charakterisieren. Die Kombination $\langle \text{Verbklasse}, \text{Stamm} \rangle$ — bzw. die Kreuzklassifikation morphologischer Primitiven entsprechend dieser zwei Dimensionen — ist so zentral für unser Modell, daß wir für sie ein eigenes System von Abkürzungen definieren:

(156) *Abkürzungen für Verbkategorie und Stammform:*

Abkürzung	Verbkategorie	Stamm
v^{cons}	v	<i>cons-stem</i>
adj^v	<i>adj</i>	<i>v-stem</i>
$v_{c_k}^{onbin}$ <i>usw.</i>	v_{c_k}	<i>onbin-stem</i>

Onbin-Form

Sowohl die japanische Schulgrammatik als auch Rickmeyer verwenden *phonologische Regeln* zur Modellierung der onbin-Verschleifungen:

(157) *Ableitung des Perfekt unter Verwendung phonologischer Regeln:*

- a. *tabe + ta → tabe-ta → tabeta*
- b. *kaki + ta → kaki-ta → kaita*
- c. *nomi + ta → nomi-ta → noñda*

Um auch diese Formen ohne Transformationen erklären zu können, verwenden wir einen zusätzlichen Stamm und nennen ihn *onbin-Stamm* (*onbin-stem*):

(158) *Ableitung des Perfekt unter Verwendung des onbin-Stammes:*

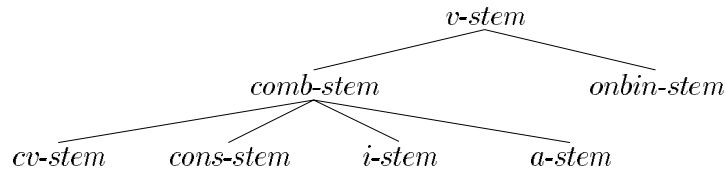
- a. *tabet + a → tabeta*
- b. *kait + a → kaita*
- c. *nond + a → noñda*

Die Verbstämme unterscheiden sich in ihrer Endung und teilen einen gemeinsamen ersten Teil. Diesen gemeinsamen Teil übernehmen wir als weiteren Stamm, und nennen ihn, da er bei vokalischen Verben nicht existiert, *konsonantischen Verbstamm*, oder kürzer: *cv-Stamm* (*cv-stem*).

Bei vokalischen Verben sind cv-Stamm, konsonantischer Stamm, i-Stamm und a-Stamm identisch; nur der onbin-Stamm unterscheidet sich von ihnen. Wir fassen die ersten vier Stämme der vokalischen Verben daher unter einem Stamm zusammen, und bezeichnen ihn als *Kombinationsstamm* (*comb-stem*). Alle Stämme zusammen subsumieren wir unter der Kategorie *Verbstamm* (v-

stem). Damit ergibt sich die folgende Hierarchie von Verbstämmen:

(159) *Stammformen der Verben:*



Die Formen aus Schema (154), (155), (157) und (158) leiten sich in unserem System folgendermaßen ab:

(160) a. *Stämme konsonantischer Verben:*

cons-stem:	$\left. \begin{array}{l} kak \\ nom \end{array} \right\} \left\{ \begin{array}{l} u \\ eba \\ e \\ ou \end{array} \right.$
a-stem:	$\left. \begin{array}{l} kaka \\ noma \end{array} \right\} \left\{ nai \right.$
i-stem:	$\left. \begin{array}{l} kaki \\ nomi \end{array} \right\} \left\{ masu \right.$
onbin-stem:	$\left. \begin{array}{l} kait \\ noñd \end{array} \right\} \left\{ a \right.$

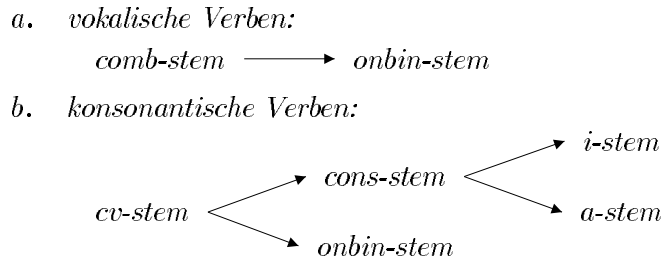
b. *Stämme vokalischer Verben:*

comb-stem =	$\left\{ \begin{array}{l} \text{cons-stem:} \\ \text{a-stem:} \\ \text{i-stem:} \\ \text{onbin-stem:} \end{array} \right. \left. \begin{array}{l} tabe \\ tabe \\ tabe \\ tabet \end{array} \right\} \left\{ \begin{array}{l} \left\{ \begin{array}{l} ru \\ reba \\ ro \\ you \end{array} \right\} \\ nai \\ masu \\ a \end{array} \right.$
--------------------	--

Ableitung der Verbstämme

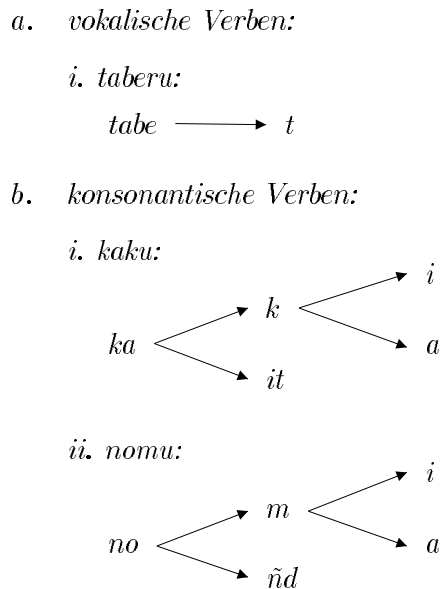
Die in unserem System vorgeschlagenen Stämme lassen sich durch Anhängen bestimmter Endungen systematisch auseinander erklären:

(161) *Ableitung der Verbstämme auseinander:*



Die Stämme aus (160) beispielsweise lassen sich folgendermaßen bilden:

(162) *Ableitung der Verbstämme auseinander:*



Wir bezeichnen die Endungen, die der Ableitung der Stämme dienen, als *Clitics*.

Eigenschaften und Form der Clitics

Clitics lassen sich durch zwei Eigenschaften charakterisieren:

1. Verbkasse und Stamm der Formen, an die sie sich anschließen können;
2. Verbkasse und Stamm der Form, die durch ihre Affigierung gebildet wird.

Werden für Stamm und Verbkategorie die Abkürzungen aus (156) (S. 91) verwendet, können die Clitics aus (162) beispielsweise folgendermaßen beschrieben werden:

(163) *Clitics aus Beispiel (162):*

a. *cons-stem:*

- | | | | | |
|-----|-----|----------------|-----------------------------------|--------------------------|
| i. | MAJ | $v_{c_K}^{cv}$ | $\longrightarrow v_{c_K}^{cons};$ | PHON $\langle k \rangle$ |
| ii. | MAJ | $v_{c_m}^{cv}$ | $\longrightarrow v_{c_m}^{cons};$ | PHON $\langle m \rangle$ |

b. *i-stem:*

- | | | | | |
|----|-----|------------------------------------|--|--------------------------|
| i. | MAJ | $v_{c_{\backslash nasaru}}^{cons}$ | $\longrightarrow v_{c_{\backslash nasaru}}^i;$ | PHON $\langle i \rangle$ |
|----|-----|------------------------------------|--|--------------------------|

c. *a-stem:*

- | | | | | |
|----|-----|-------------------------------|---|--------------------------|
| i. | MAJ | $v_{c_{\backslash w}}^{cons}$ | $\longrightarrow v_{c_{\backslash w}}^a;$ | PHON $\langle a \rangle$ |
|----|-----|-------------------------------|---|--------------------------|

d. *onbin-stem:*

- | | | | | |
|------|-----|----------------------|--|-----------------------------------|
| i. | MAJ | v_v^{cv} | $\longrightarrow v_v^{onbin};$ | PHON $\langle t \rangle$ |
| ii. | MAJ | $v_{c_k,i}^{cv}$ | $\longrightarrow v_{c_k,i}^{onbin};$ | PHON $\langle it \rangle$ |
| iii. | MAJ | $v_{c_{n,b,m}}^{cv}$ | $\longrightarrow v_{c_{n,b,m}}^{onbin};$ | PHON $\langle \tilde{n}d \rangle$ |

Es sei darauf hingewiesen, daß sich durch das Suffigieren eines Clitics nur der Stamm, nicht aber die Kategorie ändert.

Eine vollständige Beschreibung der Clitics findet sich in Kapitel 5.1.2, S. 198 ff.

2.3.4.3 Ausnahmen

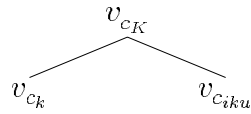
Das überarbeitete System von Verbkategorien und Verbstämmen erlaubt, einen großen Teil der Verformen, die in der traditionellen Beschreibung durch Ausnahmeregelungen behandelt werden, als regelmäßige Instanzen neuer Klassen zu beschreiben. Formen jedoch, die nur bei einem Verb vorkommen, werden auch in diesem System einfacher durch Sonderregeln beschrieben.

Dies geschieht durch die folgenden zwei Schritte:

1. Die Ableitung der Form entsprechend der gültigen "regelmäßigen" Vorschriften wird blockiert — beispielsweise indem die entsprechende morphologische Regel auf die übrigen Verben begrenzt wird;
2. Die abweichende Verform wird als zusätzlicher synthetischer Eintrag im Lexikon vermerkt.

Der unregelmäßige onbin-Stamm von 行く (*i.k-u*; gehen, fahren) beispielsweise läßt sich folgendermaßen darstellen⁶⁶:

1. Die Regel zur Ableitung des onbin-Stammes k-konsonantischer Verben wird auf k-konsonantische Verben ohne 行く (*i.k-u*) eingeschränkt. Dieses könnte durch die Abkürzung $v_{c_k \setminus iku}$ geschehen, oder durch die Verfeinerung der Kategorienhierarchie: In unserem System wurde der zusätzliche Typ v_{c_K} eingeführt, der sich in v_{c_k} und $v_{c_{iku}}$ unterteilt:



Die Kategorie v_{c_k} bezieht sich also ohnehin auf die Menge der k-konsonantischen Verben ohne $v_{c_{iku}}$.

2. Die Form *itt* wird ins Lexikon aufgenommen:

$$\left[\begin{array}{l} \text{PHON } \langle itt \rangle \\ \dots | \text{MAJ } v_{c_{iku}}^{\text{onbin}} \\ \dots \end{array} \right]$$

2.4 Morphologie und Grammatik

Die Aufteilung der Linguistik in Unterdisziplinen, die sich mit den verschiedenen Beschreibungsebenen der Sprache befassen, suggeriert eine weitgehende Unabhängigkeit der sprachlichen Komponenten: Die Phonologie beschäftigt sich mit den Lauten einer Sprache, die Morphologie mit ihren Wörtern, die Syntax mit den Sätzen, und die Semantik schließlich mit der Bedeutung. Der vergebliche Versuch jedoch, eine der sprachlichen Ebenen ohne Rückgriff auf die anderen zu beschreiben, macht die Fragwürdigkeit einer strengen Trennung der Beschreibungsebenen deutlich.

Wörter lassen sich beispielsweise nicht ohne Berücksichtigung von Phonologie und Semantik morphologisch in Morphe bzw. Morpheme segmentieren: Es ist eine wesentliche, in vielen Theorien sogar notwendige Eigenschaft der Morpheme, Bedeutung zu tragen. Auf der anderen Seite hängen Gestalt

⁶⁶ Verben der Kategorie $v_{c_{t,w,R}}$ formen den onbin-Stamm — ebenso wie 行く (*iku*) — mit *tt*. 行く (*iku*) wurde deshalb nicht, wie hier dargestellt, durch einen unregelmäßigen Lexikoneintrag modelliert, sondern durch Erweiterung dieser onbin-Regel auf $v_{c_{t,w,R,iku}}$.

und Anzahl der Morphe wesentlich von dem vorausgesetzten phonologischen Modell ab.

Die Trennung von Syntax und Morphologie ist ebenso schwierig: Ob eine sprachliche Form syntaktisch oder morphologisch erklärt wird, ist oft einfach Folge einer mehr oder weniger willkürlichen Entscheidung. Im Falle des japanischen Kausatives beispielsweise sprechen ähnlich viele und schwerwiegende Argumente für eine syntaktische wie für eine morphologische Modellierung. Entscheidet man sich für eine syntaktische Modellierung, sind Auswirkungen auf die Morphologie zu berücksichtigen; bei einer morphologischen Ableitung des Kausatives ist es umgekehrt.

In dem hier vorgestellten Modell der japanischen Verbmorphologie werden traditionell als phonologisch interpretierte Prozesse weitgehend von der Morphologie mitbehandelt. Durch die Ausrichtung auf die orthographische Form und eine Verfeinerung der morphologischen Kategorien und Grundelemente können Morphologie und Phonologie auf die Konkatenation von morphophonologischen Grundbausteinen reduziert werden. Die Onbin-Verschleifung bei der Bildung des Perfektes beispielsweise, die sonst als Resultat phonologischer Regeln betrachtet wird, behandeln wir durch eine größere Anzahl von Morphen, die sich wiederum aus der Verkettung morphophonologischer Formprimitiven ergeben. Strenggenommen handelt es sich also nicht um eine Verbmorphologie, sondern um eine “Verbmorphophonologie”.

Die Wechselwirkungen zwischen Semantik und Morphologie sind ebenfalls komplex. Der Semantik wurde deshalb ein eigenes Unterkapitel (3.3 “Semantik”) gewidmet. Es wurde versucht, semantische Darstellung und Probleme auf das Minimum zu reduzieren, das zur Beschreibung der Abhängigkeiten zwischen Morphologie und Phonologie notwendig ist.

Am schwierigsten sind die Wechselwirkungen zwischen Verbmorphologie und Syntax zu beschreiben. Ein Teil der morphologischen Prozesse — wie der oben schon erwähnte Kausativ — motivieren sich geradezu aus der Syntax. Um diese Formen zu beschreiben, muß daher ein bestimmtes syntaktisches Modell zugrundegelegt werden.

Der Vorstellung eines einfachen syntaktisch-semantischen Modells dient Kapitel 3. Dort werden die grundlegenden Schemata und Prinzipien vorgestellt und die Geometrie der Zeichen erklärt, die ihnen entspricht. Syntax und Semantik orientieren sich dabei an der HPSG (Pollard und Sag 1987 und Pollard und Sag 1994), an der JPSG (Gunji 1987 und Gunji 1991) und an der japanischen Syntax des Verbmobilprojektes (Siegel 1996b⁶⁷, Wahlster 1997,

⁶⁷ Wahlster 1997 und Wahlster 1993 bieten eine Übersicht über das Verbmobilprojekt. Siegel 1996b beschreibt eine ältere Version der japanischen Grammatik des Verbmobilprojektes, die in dem Formalismus TrUG (Block und Schachtel 1992 und Caspari und Schmid

Wahlster 1993). Genau wie bei der Entwicklung der Semantik wurde versucht, die Beschreibung möglichst einfach zu halten und auf die Phänomene zu begrenzen, die für die Modellierung der Verbmorphologie notwendig sind.

1994) implementiert wurde. Einige Aspekte der aktuellen Version werden in Siegel 1999, Siegel 1998 und Siegel 1997 beschrieben.

3 Syntax und Semantik als Rahmen der Morphologie

3.1 Einleitung

Form und *Bedeutung* sind nach Saussure im sprachlichen Zeichen so untrennbar miteinander verbunden wie die zwei Seiten eines Blattes Papier.⁶⁸ Diese Aussage bezieht sich auf die Stärke der Assoziation zwischen Form und Bedeutung sprachlicher Zeichen, nicht auf die Art ihrer Zuordnung. Atomare sprachliche Zeichen sind nach Saussure arbiträr: die Verbindung von Bedeutung und Form entsteht durch Konvention. Anders ist es bei komplexen Zeichen. Hier wird seit Frege 1879 von der *Kompositionalität der Bedeutung* gesprochen: Die Bedeutung komplexer Zeichen entsteht durch die Komposition der Bedeutung ihrer Bestandteile. Die Struktur der intendierten Bedeutung motiviert die Struktur der Form; die Struktur der Form bedingt die Struktur der Bedeutung.

Die Wechselwirkungen zwischen Form und Bedeutung können die Ebenen der Sprache überschreiten: Die morphologische Form eines Verbes bestimmt seine Bedeutung. Die Verbbedeutung erfordert eine bestimmte Anzahl semantischer Argumente. Semantische Argumente werden syntaktisch realisiert. Morphologische Form und syntaktische Form beeinflussen sich also gegenseitig.

In Kapitel 2 „*Japanische Verbmorphologie*“ wurde der Kausativ-Passiv als Beispiel für die Wechselwirkungen von Morphologie, Syntax und Semantik anhand des folgenden Beispiels vorgestellt:⁶⁹

(164) a. 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
Tarō-NOM Zirō-ACC rufen-PAST
Tarō rief Zirō.

b. 美智子が太郎に二郎を呼ばせた。⁷⁰

68 «*La langue est encore comparable à une feuille de papier : la pensée est le recto et le son le verso ; on ne peut découper le recto sans découper en même temps le verso ; de même dans la langue, on ne saurait isoler ni le son de la pensée, ni la pensée du son ; on n’y arriverait que par une abstraction dont le résultat serait de faire de la psychologie pure ou de la phonologie pure.*» (de Saussure 1916, S. 157 der Ausgabe von 1972.)

69 Die Umschrift nimmt Aspekte des verbmorphologischen Systems aus dem nächsten Kapitel vorweg: Bindestriche bezeichnen Morphemgrenzen und Punkte die Grenzen zwischen den Bauteilen (Morphemkern und Clitics), aus denen sich die Morphe zusammensetzen.

70 LIJC, S. 34, (78)a.

mitiko-ga tarou-ni zirou-wo yo.b.a-se.t-a
 Mitiko-NOM Tarō-DAT Zirō-ACC rufen-CAUS-PAST
 Mitiko ließ Tarō Zirō rufen.

c. 太郎が美智子に二郎を呼ばせられた。⁷¹

tarou-ga mitiko-ni zirou-wo yo.b.a-se-rare.t-a
 Tarō-NOM Mitiko-DAT Zirō-ACC rufen-CAUS-PASS-PAST
 Tarō wurde von Mitiko veranlaßt, Zirō zu rufen.

Wir fassen die in Kapitel 2 dargestellten, über die Morphologie des Verbes vermittelten Abhängigkeiten zwischen den sprachlichen Ebenen noch einmal tabellarisch zusammen:

(165) Morphologie	Syntax (Valenz)	Semantik ⁷²
<i>yo.ñd-a</i>	$\left\{ \begin{array}{l} \text{PP}[\text{subj}, \text{ga}] \\ \text{PP}[\text{obj}, \text{wo}] \end{array} \right\}$	$\left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{call} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right] \end{array} \right]$
<i>yo.b.a-se.t-a</i>	$\left\{ \begin{array}{l} \text{PP}[\text{subj}, \text{ga}] \\ \text{PP}[\text{obj}, \text{wo}] \\ \text{PP}[\text{obj2}, \text{ni}] \end{array} \right\}$	$\left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ I cont}(\text{obj2}) \\ \textcircled{3} \left[\begin{array}{c} \text{call} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right] \end{array} \right] \end{array} \right]$
<i>yo.b.a-se-rare.t-a</i>	$\left\{ \begin{array}{l} \text{PP}[\text{subj}, \text{ga}] \\ \text{PP}[\text{obj}, \text{wo}] \\ \text{PP}[\text{obj2}, \text{ni}] \end{array} \right\}$	$\left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ cont}(\text{obj2}) \\ \textcircled{2} \text{ I cont}(\text{subj}) \\ \textcircled{3} \left[\begin{array}{c} \text{call} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right] \end{array} \right] \end{array} \right]$

⁷¹ LIJC, S. 34, (78)b.

⁷² **cont(subj)**, **cont(obj)** und **cont(obj2)** sind Abkürzungen, die in dieser Reihenfolge die Semantik vom Subjekt, vom direkten Objekt und vom indirekten Objekt bezeichnen. Formal lassen sie sich folgendermaßen definieren:

Andere Veränderungen von Syntax und Semantik durch verbmorphologische Mittel sind denkbar: Modifikationen können Satzfinalstellung, Adnominalstellung oder Adverbalstellung bedingen; sie können den Subkategorisierungsrahmen beeinflussen; sie können die Modifikation durch Adverben obligatorisch machen etc.

Veränderungen des semantischen bzw. syntaktischen Verhaltens sprachlicher Zeichen gehen mit Veränderungen der semantischen bzw. syntaktischen Merkmale einher. Merkmale jedoch haben keine Bedeutung an sich, sondern erhalten ihre Bedeutung aus der linguistischen Theorie, auf die sie sich beziehen. Ein morphologisches Modell muß daher auf Modellen der Syntax und Semantik aufbauen.

Dieses Kapitel dient der Darstellung des syntaktisch-semantischen Rahmens, auf dem die vorgestellte Morphologie beruht. Sowohl Semantik als auch Syntax werden dabei auf ihre wesentlichen, zur Erklärung der Morphologie notwendigen Prinzipien reduziert. Durch diese Beschränkung auf das Prinzipielle lassen sich beide Komponenten jedoch leicht an komplexere Beschreibungen anpassen.

Dominanzschemata, Prinzipien und Zeichen

Chomsky hebt die *Rekursivität* als wesentliches Merkmal der menschlichen Sprache hervor⁷³: Aus einem beschränkten Vorrat sprachlicher Grundbausteine läßt sich dank einer ebenfalls beschränkten Menge rekursiver Kombinationsschemata eine unendliche Vielfalt sprachlicher Äußerungen ableiten. Sprachliche Grundbausteine, wie auch die durch Kombination aus ihnen gewonnenen Elemente, werden in der HPSG als *Zeichen* dargestellt; die Kombinationsschemata werden durch Dominanzschemata und Prinzipien abgebildet.

(166)	Abkürzung:	abgekürzte AVM:
	$[... X \text{ cont}(\text{subj})]$	$\left[\begin{array}{l} ... X \text{ } \boxed{\text{SUBJ}} \\ ... \text{SUBCAT} \{ \text{PP}[\text{subj}]:\boxed{\text{SUBJ}} \dots \} \end{array} \right]$
	$[... X \text{ cont}(\text{obj})]$	$\left[\begin{array}{l} ... X \text{ } \boxed{\text{OBJ}} \\ ... \text{SUBCAT} \{ \text{PP}[\text{obj}]:\boxed{\text{OBJ}} \dots \} \end{array} \right]$
	$[... X \text{ cont}(\text{obj2})]$	$\left[\begin{array}{l} ... X \text{ } \boxed{\text{OBJ2}} \\ ... \text{SUBCAT} \{ \text{PP}[\text{obj2}]:\boxed{\text{OBJ2}} \dots \} \end{array} \right]$

$\boxed{\text{SUBJ}}$ beschreibt innerhalb von Subkategorisierungsmengen die Semantik eines Zeichens:

(167)	Abkürzung:	abgekürzte AVM:
	$\boxed{\text{SUBJ}}$	$[\text{SYNSEM} \text{LOCAL} \text{CONT} \text{ } \boxed{\text{SUBJ}}]$

73 vgl. Chomsky 1957.

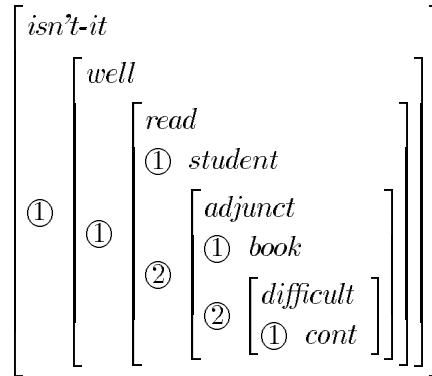
Beispielsatz

Die Beschreibung formaler Theorien ist ohne einen Anwendungskontext abstrakt und schwer verständlich. Um die linguistischen Motivationen sichtbar zu machen, liegt es daher nahe, die einzelnen Aspekte der Theorie durch ein konkretes Beispiel zu verdeutlichen. Der folgende Satz — hier zusammen mit seiner erst an späterer Stelle erklärten semantischen Struktur dargestellt — soll diesem Zwecke dienen:

(168) a. 学生は難しい本を良く読むね。

gakusei-ha muzukasii hoñ-wo yoku yomu ne
 Student-TOP schwer buch-ACC gut lesen nicht wahr
 Der Student liest das schwierige Buch gut, nicht wahr?

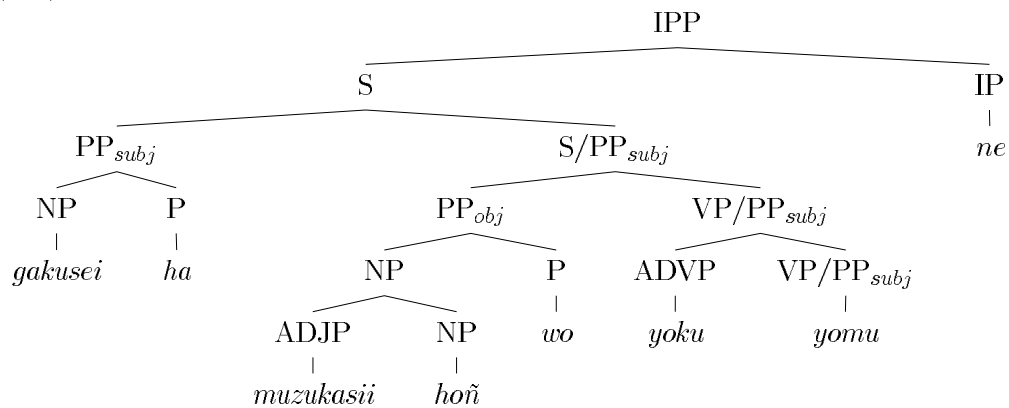
b. Semantische Darstellung des Satzes:



Anhand dieses Satzes, dessen Ableitung in unserem Modell schrittweise durchgeführt wird, werden Dominanzschemata, Prinzipien und sprachliche Zeichen eingeführt. Bevor wir uns jedoch diesen zuwenden, soll die Struktur des Satzes informell und von unserem Modell unabhängig dargestellt werden. Gleichzeitig wird damit ein Überblick über das folgende Kapitel ermöglicht.

Satz (168) läßt sich durch den folgenden Syntaxbaum beschreiben:

(169) a. Struktur des Beispielsatzes:



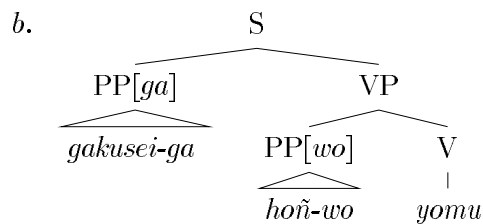
b. *Liste der verwendeten Abkürzungen:*

N	<i>Nomen</i>
P	<i>Postposition</i>
V	<i>Verb</i>
ADJ	<i>Adjektiv</i>
ADV	<i>Adverb</i>
IP	<i>Interjektionspartikel</i>
XP	<i>X-Phrase</i>
S	<i>Satz</i>
X/Y	<i>X mit nicht realisiertem Y (Gap)</i>
<i>subj</i>	<i>Subjekt</i>
<i>obj</i>	<i>Objekt</i>

Zentrum des Satzes ist das Verb 読む (*yomu*; lesen). Es fordert zwei Komplemente: ein *Subjekt* — der *Lesende* — und ein *direktes Objekt* — das *Gelesene*. Im Standardfall wird das Subjekt durch eine Postpositionalphrase (PP) mit der Postposition (P) が (*ga*) realisiert und das Objekt durch eine Postpositionalphrase mit der Postposition を (*wo*):

(170) a. 学生が本を読む。

gakusei ga hon wo yomu
 Student-NOM buch-ACC lesen
 Der Student liest das Buch.

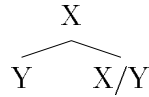


In Satz (169) hingegen wurde die Subjekt-PP topikalisiert — d.h. aus dem Satz heraus an dessen Anfang verschoben. Topikalisierte PPs werden durch die Postposition は (*ha*) gekennzeichnet.⁷⁴ Im Syntaxbaum drückt sich die Topikalisierung durch Kategorien mit einem Schrägstrich aus: X/Y steht für eine Kategorie X, in der die Kategorie Y nicht realisiert (sondern *extrahiert*) wurde. VP/PP_{subj} beispielsweise bezeichnet eine Verbalphrase⁷⁵ mit extrahierter PP_{subj}. Ein Teilbaum der Form

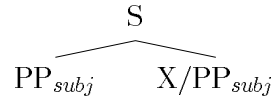
⁷⁴ は (*ha*) ersetzt hier die Postposition が (*ga*).

⁷⁵ Wir unterscheiden an dieser Stelle nicht zwischen den verschiedenen Projektionen von Verben: Eine VP bezeichnet hier einfach eine Projektion eines Verbes, die bezüglich des Subjekts noch ungesättigt ist.

(171) a. *Sättigung eines extrahierten Komplementes:*



b. *Beispiel:*



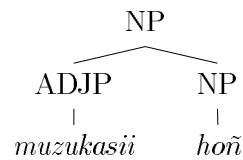
beschreibt die nachträgliche Sättigung dieser Kategorie — hier als Topic des Satzes.

読む (*yomu*; lesen) und 本 (*hon*; Buch) werden durch Attribute näher bestimmt: 読む (*yomu*) durch das Adverb 良く (*yoku*; gut), 本 (*hon*) durch das Adjektiv 難しい (*muzukasii*; schwierig):

(172) a. i. 難しい本

muzukasii hon
schwer buch
schwieriges Buch

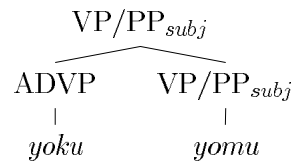
ii.



b. i. 良く読

yoku yomu
gut lesen
liest gut

ii.



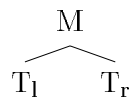
Phrasen aus Attribut und Bezugselement haben die gleiche Kategorie wie das Bezugselement ohne Attribut.

Schließlich wird der Satz in seiner Gesamtheit dem Interjektionspartikel ね (*ne*; nicht wahr) zugeordnet.

Dominanzschemata

Jeder Teilbaum von (169) ist binär:

(173)



Die binären Teilbäume lassen sich durch drei Typen syntaktischer Beziehungen erklären:

1. T_l ist *Komplement* von T_r ;
2. T_l ist *Adjunkt* zu T_r ;

3. T_I ist *Füller* von T_r ;

In allen Fällen ist die rechte Tochter Head der Phrase: Japanisch ist eine binär verzweigende, Head-finale Sprache⁷⁶.

Die drei Arten syntaktischer Beziehungen werden durch sogenannte *Dominanzschemata* modelliert. Entsprechend der drei unterschiedlichen Dominanzschemata gliedert sich dieses Kapitel in drei Unterabschnitte: “*Komplemente*”, “*Adjunkte*” und “*Nichtlokale Abhängigkeiten*”. Der letzte Unterabschnitt ist den *Relativsätzen* und *Attributsätzen* des Japanischen gewidmet.

Prinzipien

Dominanzschemata sind für die Bildung der lokalen Bäume verantwortlich. Die Gestalt der Zeichen, die diese Bäume repräsentieren, wird aber nur zu einem Teil von ihnen bestimmt.

Linguistische Zeichen modellieren das Verhalten bzw. die Eigenschaften der beschriebenen sprachlichen Elemente durch Merkmale. Die Merkmalwerte lexikalischer Zeichen werden im Lexikon spezifiziert. Werden zwei Zeichen durch einen lokalen Baum zu einem Zeichen zusammengefaßt, errechnen sich die Merkmale des Gesamtzeichens aus den Merkmalen der einzelnen Tochterzeichen. Die dazu nötigen “Rechenvorschriften” werden durch *Prinzipien* beschrieben.

Dominanzschemata lassen sich nicht ohne ihre Wechselwirkungen mit den einzelnen Prinzipien beschreiben. Die beteiligten Prinzipien werden deshalb an der Stelle eingeführt, an der sie zur Erklärung der Schemata benötigt werden.

Zeichenarchitektur

Die Definition der Dominanzschemata und Prinzipien erfordert eine einheitliche Architektur der Zeichen. Diese sollte die linguistischen Informationen so

⁷⁶ Diese Aussage ist ein gutes Beispiel für die wechselseitige Beeinflussung von Theorie und Objekt: Japanisch läßt sich auch als nicht binäre, nur in Teilen Head-finale Sprache beschreiben. Auch wenn Head-Finalität und Binarität der Verzweigungen von den japanischen Sprachdaten in vielen Fällen nahegelegt werden, lassen sich eine Reihe von Phänomenen auch anders interpretieren. Die Konsequenz, mit der diese beiden Eigenschaften in unserer Beschreibung auf alle Konstruktionen angewandt wird, entstammt also an manchen Stellen auch einer ästhetischen oder systematischen Motivation.

In diesem Sinne sind Binarität und Head-Finalität also *Axiome* unserer Beschreibung, die dem Japanischen — zumindestens teilweise — von der Theorie “vorgeschrieben” werden. Es wäre daher besser zu sagen: Das Japanische läßt sich aus der Perspektive der Beschreibungsökonomie und Eleganz am besten als Head-finale, binär verzweigende Sprache beschreiben.

organisieren, daß die Formulierung der Schemata und Prinzipien möglichst einfach geschehen kann. Da sich die Merkmalstruktur also offensichtlich aus den Prinzipien und Schemata motiviert, werden Merkmale und ihre Wertebereiche an der Stelle eingeführt, an der sie zum ersten Mal benötigt werden. Aspekte der Merkmalsarchitektur, die sich aus dem Kontext erschließen lassen, bleiben unerklärt.

Relativsätze im Japanischen

Wie oben festgestellt, ist das Japanische eine Head-finale Sprache — bzw. wird hier als eine solche beschrieben: Der Head einer Phrase steht immer an ihrem Ende. Aus dieser Eigenschaft lassen sich zwei weitere ableiten:

1. Der Head eines Satzes ist das Verb. Das Verb steht folglich am Ende des Satzes. Eine Ausnahme bilden Sätze mit Interjektions- oder Fragepartikeln: in diesen Fällen ist das Partikel Head und steht am Ende des Satzes.
2. Der Head von Phrasen aus Attribut und Bezugswort ist das Bezugswort: Das *Modifizierende* steht also immer vor dem *Modifizierten*.

Aus Punkt 2. läßt sich eine weitere Besonderheit des Japanischen erschließen: Relativsätze sind Attribute; sie stehen damit nicht wie im Deutschen oder Englischen *hinter* ihrem Bezugswort, sondern *vor* diesem.

In Kapitel 2 wurde auf die Ähnlichkeit zwischen Verben und Adjektiven eingegangen. Diese Ähnlichkeit und die Möglichkeit, Verben in Adjektive und umgekehrt Adjektive in Verben zu überführen⁷⁷, legen nahe, beide Kategorien auch in ihrem syntaktischen Verhalten ähnlich zu modellieren: werden also adnominale Verben mit ihren Ergänzungen als Relativsatz modelliert, müssen auch adnominale Adjektive als Relativsatz dargestellt werden.

Die morphologische Ableitung der Adnominalform von Verben und Adjektiven ist Teil der Verbmorphologie. Voraussetzung ihrer Behandlung ist jedoch ein formales Modell der Relativ- und Attributsätze im Japanischen. Den japanischen Relativ- und Attributsätzen ist deshalb ebenfalls ein Abschnitt dieses Kapitels gewidmet.

Der zweite Teil dieses Kapitels widmet sich der Semantik. Von den semantischen Problemen interessieren uns jedoch nur solche, die die *Interaktion* von Morphologie und Semantik — oder anders ausgedrückt: die *Kompositionalität* semantischer Strukturen — betreffen. Wir verwenden daher einen semantischen Formalismus, der die Mechanismen der Kompositionalität hervorhebt und von allen anderen Problemen abstrahiert.

⁷⁷ vgl. Paragraph “Überführung von Verben in Adjektive und umgekehrt” 69.

3.2 Syntax

Alle Sätze des Japanischen lassen sich auf binär verzweigende lokale Bäume zurückführen. Den lokalen Bäumen liegt ein rekursives Bildungsschema zugrunde: Die Töchter oder Teilbäume, die sich in einem lokalen Baum zu einer neuen Einheit verbinden, entsprechen entweder lexikalischen Einträgen oder rekursiv nach dem gleichen Schema gebildeten komplexen Einheiten. Lokale Bäume stehen damit für das Prinzip syntaktischer Konstruktion. Diese Konstruktion ist nicht willkürlich, sondern erfolgt nach bestimmten, genau festgelegten Konstruktionsvorschriften — den *Schemata*. In der Tradition der Generativen Transformationsgrammatik werden diese Schemata als *$\bar{X}$ -Schemata* (gelesen: “X-Bar-Schemata”) bezeichnet; in der HPSG heißen sie *Dominanzschemata*⁷⁸.

Dominanzschemata und Merkmalslogiken

Im ersten Kapitel haben wir Typenhierarchien und Constraints als Beschreibungsmittel der HPSG kennengelernt. Syntaxbäume, lokale Bäume, lexikalische Einträge, Dominanzschemata etc. — alle linguistischen Begrifflichkeiten müssen auf diese beiden Mittel abgebildet werden.

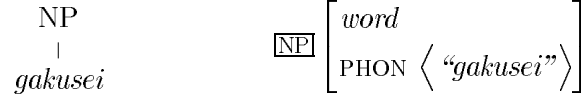
Im Syntaxbaum (169) gibt es zwei Typen lokaler Bäume:

- (174) a. *unäre lokale Bäume*
 bzw. *Knoten mit einer terminalen Tochter:*
- | | |
|--|--|
| $\begin{array}{c} M \\ \\ T \end{array}$ | $\begin{array}{c} NP \\ \\ \textit{gakusei} \end{array}$ |
|--|--|
- b. *binäre lokale Bäume*
 bzw. *Knoten mit zwei nicht-terminalen Töchtern:*
- | | |
|--|---|
| $\begin{array}{cc} & M & \\ & / \backslash & \\ T_l & & T_r \end{array}$ | $\begin{array}{ccc} & PP & \\ & / \backslash & \\ NP & & P \end{array}$ |
|--|---|

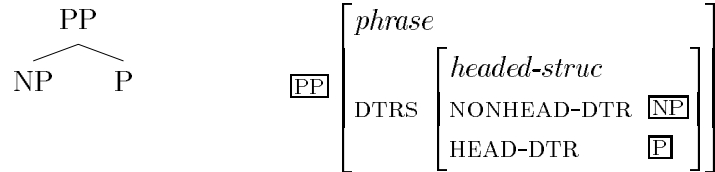
Beide Typen lokaler Bäume lassen sich auch als AVM beschreiben — wir vernachlässigen dabei die Kodierung der Kategorien, und stellen diese als Tags dar:

⁷⁸ Syntaktische Konstruktion ist vielfältig. In den Anfängen der generativen Grammatik gab es deshalb eine große Anzahl sehr spezieller Schemata. Im Laufe der Zeit stellte sich jedoch heraus, daß sich diese unter Berücksichtigung der lexikalischen Information der involvierten Töchter auf wenige allgemeine Schemata zurückführen lassen. Tendenz der modernen formalen Grammatik ist es deshalb, von wenigen allgemeinen Schemata auszugehen, deren Verhalten weitgehend durch die Eigenschaften der Töchter — insbesondere der Head-Tochter — bestimmt wird.

(175) a. *Knoten mit einer terminalen Tochter:*



b. *Knoten mit zwei nicht-terminalen Töchtern:*

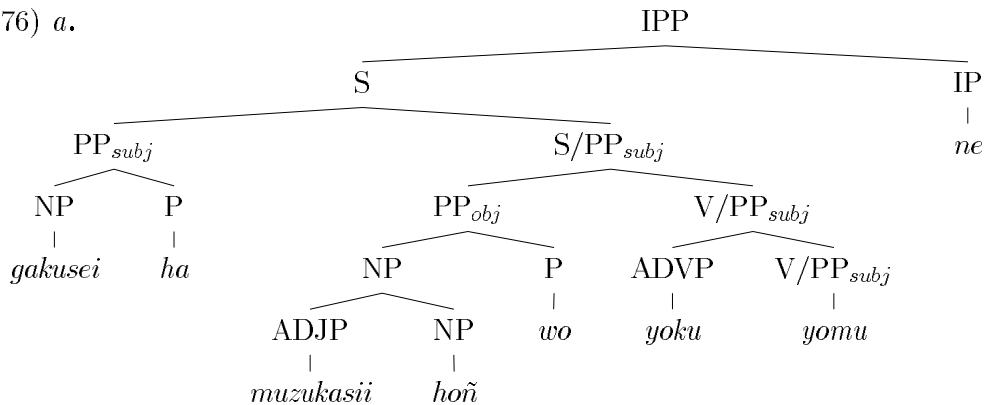


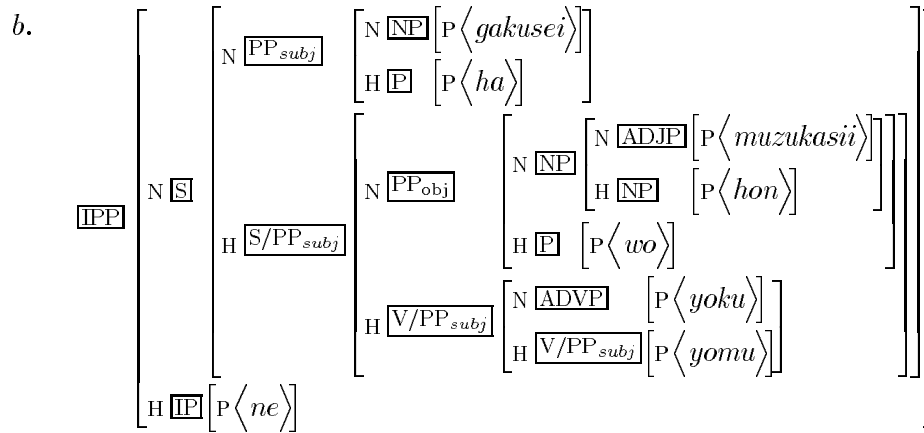
Zeichen, wie das aus (175a), entsprechen Wörtern und sind — soweit sie nicht durch eine morphologische Komponente beschrieben werden — Bestandteile des Lexikons. Sie werden deshalb durch den Typ *word* gekennzeichnet und als *lexikalische Zeichen* bezeichnet. Zeichen wie (175b) hingegen stehen nicht im Lexikon, sondern werden aus lexikalischen Zeichen durch syntaktische Kombination abgeleitet. Wir bezeichnen sie von nun an als *phrasale Zeichen* und ordnen ihnen den Typ *phrase* zu.

Durch rekursive Einbettung läßt sich der gesamte Baum (169) — hier als (176a) noch einmal dargestellt — in eine AVM übersetzen.

Die Struktur $\left[\text{DTRS} \left[\begin{array}{l} \text{NONHEAD-DTR} \boxed{\text{N}} \\ \text{HEAD-DTR} \quad \boxed{\text{H}} \end{array} \right] \right]$ wird dabei durch $\left[\begin{array}{l} \text{N} \boxed{\text{N}} \\ \text{H} \boxed{\text{H}} \end{array} \right]$ abgekürzt und $\left[\text{PHON} \langle \boxed{\text{P}} \rangle \right]$ durch $\left[\text{P} \langle \boxed{\text{P}} \rangle \right]$:

(176) a.

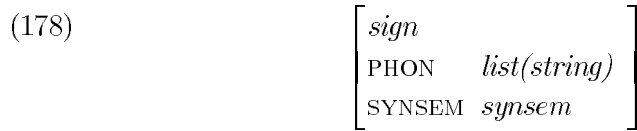




Zeichen Die zwei Typen der Zeichen — *word* für lexikalische und *phrase* für phrasale Zeichen — lassen sich unter dem allgemeineren Typ *sign* zusammenfassen:



Die Gestalt der Zeichen kann durch Constraints zu diesen Typen beschrieben werden. Die folgenden Constraints nehmen einige Eigenschaften von Zeichen vorweg, die erst an späterer Stelle motiviert werden:



Constraints vererben sich — wie in Kapitel 1 beschrieben — entlang der Typenhierarchie von Supertypen auf ihre Subtypen, hier also von *sign* auf *word* und *phrase*. Allen Zeichen — lexikalischen wie phrasalen — sind also zwei Eigenschaften gemeinsam:⁷⁹

1. Sie haben eine *phonologische Form* — dargestellt durch das Merkmal PHON, das als Wert Listen phonologischer Formen annimmt;
2. Sie haben *syntaktisch-semantische Eigenschaften* — dargestellt durch das Merkmal SYNSEM.

Vorerst interessiert uns nur das Merkmal PHON bei lexikalischen Zeichen. Alle anderen Merkmale und Merkmalwerte werden an späterer Stelle beschrieben.

⁷⁹ In der Standardversion der HPSG (Pollard und Sag 1994) gibt es zusätzlich noch das Merkmal QSTORE. Dank der vereinfachten Semantik kann in dieser Darstellung jedoch auf QSTORE verzichtet werden.

Phrasale Zeichen

Phrasale Zeichen haben zusätzlich zu den von *sign* ererbten Merkmalen PHON und SYNSEM das Merkmal DTRS.

$$(179) \quad \left[\begin{array}{l} \textit{phrase} \\ \text{DTRS } \textit{headed-structure} \end{array} \right]$$

Dieses wird durch folgenden Constraint näher bestimmt:⁸⁰

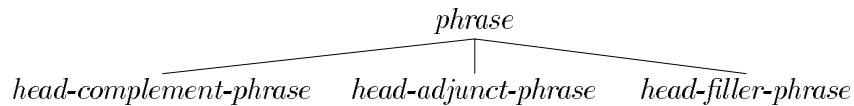
$$(180) \quad \left[\begin{array}{ll} \textit{headed-structure} & \\ \text{HEAD-DTR} & \textit{sign} \\ \text{NONHEAD-DTR} & \textit{sign} \end{array} \right]$$

Alle phrasalen Zeichen haben also zwei Töchter: Eine *Head-Tochter* — dargestellt durch HEAD-DTR — und eine *Non-Head-Tochter* — dargestellt durch NON-HEAD-DTR.

Dominanzschemata

Die Dominanzschemata für Komplemente, Attribute und Füller stellen weitere Bedingungen an phrasale Zeichen. Sie lassen sich daher durch Subtypen zu *phrase* mit entsprechenden Constraints realisieren:

(181) *Typenhierarchie:*



Typenconstraints:

$$\left[\begin{array}{l} \textit{head-complement-phrase} \\ \dots \end{array} \right] \quad \left[\begin{array}{l} \textit{head-adjunct-phrase} \\ \dots \end{array} \right] \quad \left[\begin{array}{l} \textit{head-filler-phrase} \\ \dots \end{array} \right]$$

Durch das Resolutionskalkül wird jedes Vorkommen von *phrase* durch einen seiner minimalen Subtypen ersetzt — d.h. durch *head-complement-phrase*, *head-adjunct-phrase* oder *head-filler-phrase*. Jede Substruktur einer Lösung muß durch ein Element des ererbten Constraints zu seinem Typ subsumiert werden. Das aber bedeutet, daß jeder lokale Teilbaum eines Syntaxbaumes ein Dominanzschema erfüllen muß.

⁸⁰ In unserem Modell kommen wir mit dem Typ *headed-structure* für DTRS aus, und könnten daher auf diese Merkmalsebene verzichten. In der Standardversion hingegen gibt es *headed-structure*, *coordinate-structure*, *head-complement-structure*, *head-marker-structure*, *head-adjunct-structure* und *head-filler-structure*.

Subkategorisierung

Wir haben die syntaktische Konstruktion auf drei Dominanzschemata zurückgeführt: Das *Head-Komplement-Schema*, das *Head-Adjunkt-Schema* und das *Head-Füller-Schema*. Diese Schemata wurden als syntaktische Kombinationsmittel vorgestellt, die durch Eigenschaften ihrer Töchter gesteuert werden. Es stellt sich damit die Frage, wie diese Eigenschaften beschaffen sind.

Die Antwort, die von den meisten modernen Syntax-Theorien vorgeschlagen wird, läßt sich durch den Begriff *Subkategorisierung* charakterisieren. Dieser Begriff erhält in unserem syntaktischen Modell eine einfache, technische Bedeutung. Unsere bisherige Beschreibung der Dominanzschemata betrifft nur deren binären, rekursiven Charakter. Würden keine weiteren Beschränkungen zugrunde gelegt werden, könnte jedes Wort bzw. jede Wortgruppe mit jedem anderen Wort bzw. jeder anderen Wortgruppe kombiniert werden. Nur bestimmte Kombinationen sind jedoch sinnvoll. Diese hängen von Eigenschaften der beteiligten zwei Elementen ab und können daher am besten auch an diesen beschrieben werden.

Vereinfacht ausgedrückt kann gesagt werden, daß eine Tochter die andere Tochter eines lokalen Baumes als mögliche Schwester *legitimiert* bzw. die Töchter sich gegenseitig *legitimieren*. Diese Legitimierung bezeichnen wir von nun an als *Subkategorisierung* und sagen: *A subkategorisiert für B*, wenn *A B* als (direkte oder mittelbare) Schwester legitimiert.

Damit benutzen wir den Begriff *Subkategorisierung* in einem weiteren Sinne als gemeinhin üblich:⁸¹ In unserem Modell *subkategorisieren* u.a.

- Verben, Postpositionen, Interjektionspartikel etc. für ihre Komplemente;
- Adverben und Adnomen für ihre Bezugswörter;
- Verben für Adverben (Verben und Adverben subkategorisieren sich also gegenseitig);
- Phrasen mit einer "Lücke" für ihre Füller (beispielsweise bei Topikalisierung);

Die Idee, die dem Begriff *Subkategorisierung* zugrunde liegt, ist seit Tesnière 1959 — ursprünglich unter der Bezeichnung *Valenz* — zentral für die moderne Linguistik. Entsprechend vielfältig ist die Nomenklatur, die in diesem

⁸¹ In den meisten Fällen wird von Subkategorisierung nur im Zusammenhang mit Komplementen gesprochen. Von Adjunkten etc. wird hingegen meist gesagt, daß sie ihr Bezugswort etc. *wählen*.

Zusammenhang verwendet wird: Wir benutzen in dieser Arbeit — um Unklarheiten von vornherein vorzubeugen — die folgenden Begriffe synonym:

1. *Subkategorisierung* und *Valenz* — den Begriff *Kasus* bringen wir mit den Postpositionen des Japanischen in Verbindung, und benutzen daher im Zusammenhang mit Verben weiterhin den Ausdruck *Kasusrahmen*;
2. *Komplement*, *Ergänzung* und *Objekt*;
3. *Adjunkt* und *Angabe* — adnominale Adjunkte bzw. Angaben bezeichnen wir auch als *Attribute*.

3.2.1 Komplemente

Der Begriff *Subkategorisierung* soll zunächst anhand des Dominanzschematas für Komplemente dargestellt werden. Wir gehen an dieser Stelle nicht auf Motivation und Problematik der Abgrenzung zwischen *Komplement*, *Adjunkt*, *Füller* ein, sondern beschreiben diese Begriffe anhand von Beispielen und deren Umsetzung in unserem Modell.⁸² Dazu greifen wir auf eine vereinfachte Version des Beispielsatzes (168) zurück:

(182) 学生が本を読む。

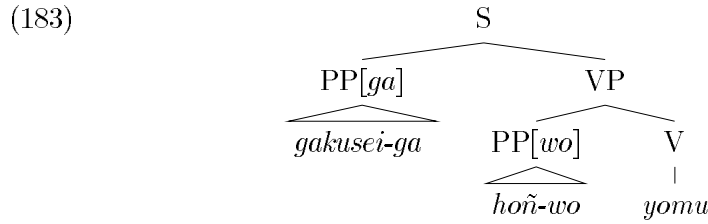
gakusei ga hoñ wo yomu
 Student-NOM buch-ACC lesen
 Der Student liest das Buch.

Das Zentrum des Satzes wird durch das Verb 読む (*yomu*; lesen) gebildet. Dieses fordert zwei Objekte: Ein Subjekt — der Lesende — und ein direktes Objekt — das Gelesene. Das Subjekt wird im Japanischen durch eine Postpositionalphrase (PP) mit der Postposition が (*ga*) (abgekürzt als PP[*ga*]) realisiert; das Objekt durch eine PP mit der Postposition を (*wo*) (abgekürzt als PP[*wo*]). Unter Annahme der binären Verzweigung ergibt sich die folgende Darstellung:⁸³

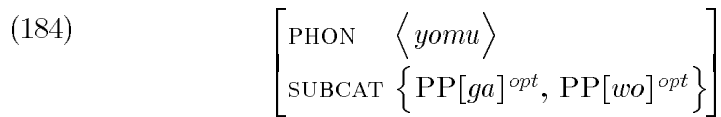
82 Für den Versuch, die Begriffe *Ergänzung* und *Angabe* im Zusammenhang mit dem Begriff *Satzglied* zu definieren, vergleiche Engel 1994.

83 S und VP stehen respektive für *Satz* und *Verbalphrase*. Diese Abkürzungen werden in Analogie zu Theorien wie der Standardversion der HPSG (Pollard und Sag 1987, Pollard und Sag 1994) verwendet: S steht für einen vollständigen Satz; VP für einen Satz, bei dem das Subjekt noch nicht gesättigt wurde. Strenggenommen haben sie keine formale Entsprechung in unserem Modell: Phrasen werden — dem fakultativen Charakter der meisten japanischen Komplemente entsprechend — anders als in diesen, vornehmlich an indoeuropäischen Sprachen orientierten Theorien definiert.

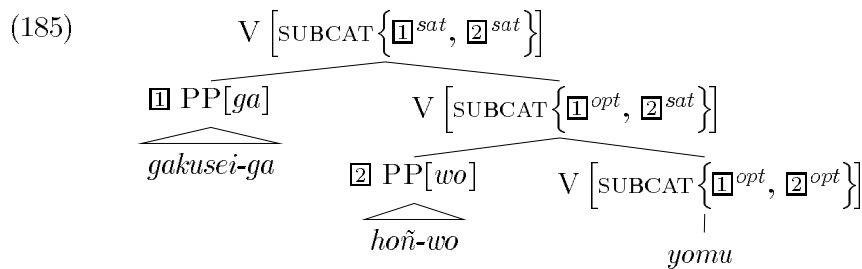
Zur Verwendung des Ausdruckes *Phrase* in dieser Arbeit vergleiche (206), S. 126.



Die Eigenschaft von 読む (*yomu*; lesen), für eine PP[*ga*] und eine PP[*wo*] zu subkategorisieren, wird mit Hilfe des Merkmals SUBCAT formal folgendermaßen ausgedrückt:



Mit der AVM-Darstellung des Verbes *yomu* läßt sich (183) präzisieren:



Dieser Syntaxbaum veranschaulicht den Mechanismus, durch den die Subkategorisierung in Unifikationsgrammatiken wie der hier vorgestellten modelliert wird: Durch das *Head-Komplement-Schema* wird die Non-Head-Tochter (die linke Tochter) mit einem Element aus der SUBCAT-Menge unifiziert. Als Non-Head-Tochter sind daher nur Zeichen zulässig, die mit einem SUBCAT-Element der Head-Tochter unifizierbar sind.

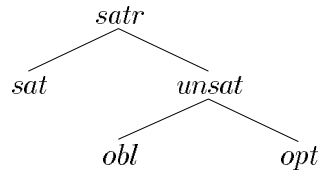
In der Standard-HPSG werden gesättigte Elemente aus der SUBCAT-Menge gelöscht. Diese Art der Kodierung erlaubt es jedoch nur, zwei Zustände darzustellen, und erschwert damit die Modellierung fakultativer Elemente. In Anlehnung an die Verbmobilgrammatik von Melanie Siegel (Siegel 1996b) stellen wir den Sättigungsgrad eines Zeichens deshalb durch ein eigenes Merkmal dar, dessen Wert als Index an der rechten oberen Ecke der SUBCAT-Elemente notiert wird:⁸⁴ der Typ *sat* markiert ein gesättigtes Element; *unsat* und seine Untertypen hingegen ein ungesättigtes Komplement.

⁸⁴ Die Darstellung der Sättigung von Komplementen durch einen Index dient also nur der besseren Lesbarkeit. Formal werden Index und subkategorisiertes Objekt in einer Struktur des Types *subcat-element* (*sc-elem*) zusammengefaßt:

Für ungesättigte Elemente kann damit genau spezifiziert werden, ob sie *obligatorisch* (*obl*) sind oder *optional* (*opt*): nur Elemente mit dem Index *obl* müssen realisiert werden; solche mit *opt* hingegen sind fakultativ.

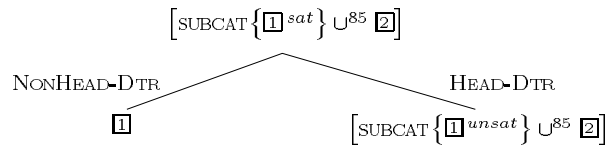
Als Typenhierarchie lassen sich diese Indexe folgendermaßen darstellen:

- (188) *Typenhierarchie zur Kennzeichnung der Sättigung von Komplementen in SUBCAT-Mengen:*



Das Head-Komplement-Schema kann auf ein Element der SUBCAT-Menge nur dann angewandt werden, wenn dessen Index durch *unsat* subsumiert wird, es also als noch nicht realisiert gekennzeichnet ist. Nach der Realisierung wird das entsprechende Element durch *sat* als realisiert markiert. Damit kann das Head-Komplement-Schema folgenderweise dargestellt werden:

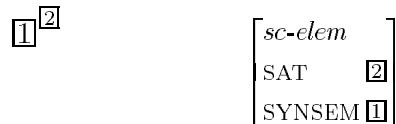
- (189) *Das Head-Komplement-Schema (vorläufige Version):*



- (186)
$$\left[\begin{array}{l} sc\text{-}elem \\ SAT \quad satr \\ SYNSEM \quad synsem \end{array} \right]$$

Die Indexschreibweise ist eine Abkürzung für diese Struktur:

- (187) Abkürzung: abgekürzte AVM:



85 $SUBCAT \{ [1]^{(un)sat} \} \cup [2]$ symbolisiert, daß ein beliebiges Element der SUBCAT-Menge ausgewählt werden kann. $\cup$ läßt sich — wie die anderen Abkürzungen auch — durch

Ein Element aus der SUBCAT-Menge wird mit der Non-Head-Tochter unifiziert. Der SUBCAT-Wert der Head-Tochter wird von der Mutter mit einer Variation übernommen: Der Index des realisierten Komplementes wird durch *sat* als gesättigt gekennzeichnet.

Die Modellierung der Komplemente durch Mengen und Indizes wurde von Melanie Siegel in ihrem Modell der japanischen Syntax (vgl. Siegel 1996b) vorgeschlagen.⁸⁶ In der Standardversion der HPSG nimmt SUBCAT die Liste der ungesättigten Komplemente als Wert. Gesättigte Komplemente werden aus der Liste gestrichen. Es können also nur zwei Zustände dargestellt werden:

(190) Abkürzung:

$$\left[\begin{array}{c} \dots | \text{SUBCAT } \{ \boxed{1}^{sat} \} \cup \boxed{2} \\ \text{DTRS} \left[\begin{array}{c} \text{H-DTR} \left[\dots | \text{SUBCAT } \{ \boxed{1}^{unsat} \} \cup \boxed{2} \right] \\ \text{NH-DTR} \left[\text{SYNSEM } \boxed{1} \right] \end{array} \right] \end{array} \right]$$

abgekürzte AVM:

$$\left[\begin{array}{c} \dots | \text{SUBCAT } \boxed{1} \\ \text{DTRS} \left[\begin{array}{c} \text{H-DTR} \left[\dots | \text{SUBCAT } \boxed{2} \right] \\ \text{NH-DTR} \left[\text{SYNSEM } \boxed{3} \right] \end{array} \right] \\ \text{COND} \left[\begin{array}{c} \text{substitute} \\ \text{LIST1 } \boxed{2} \\ \text{ELEM1} \left[\begin{array}{cc} \text{SAT} & unsat \\ \text{SYNSEM} & \boxed{3} \end{array} \right] \\ \text{LIST2 } \boxed{1} \\ \text{ELEM2} \left[\begin{array}{cc} \text{SAT} & sat \\ \text{SYNSEM} & \boxed{3} \end{array} \right] \end{array} \right] \end{array} \right]$$

(191) Constraints für den Typ *substitute*:

$$\left[\begin{array}{c} \text{substitute} \\ \text{LIST1 } \langle \boxed{1} \bullet \boxed{3} \rangle \\ \text{ELEM1 } \boxed{1} \\ \text{LIST2 } \langle \boxed{2} \bullet \boxed{3} \rangle \\ \text{ELEM2 } \boxed{2} \end{array} \right] \quad \left[\begin{array}{c} \text{substitute} \\ \text{LIST1 } \langle \boxed{1} \bullet \boxed{2} \rangle \\ \text{ELEM1 } \boxed{3} \\ \text{LIST2 } \langle \boxed{1} \bullet \boxed{4} \rangle \\ \text{ELEM2 } \boxed{5} \\ \text{CONDITION} \left[\begin{array}{c} \text{substitute} \\ \text{LIST1 } \boxed{2} \\ \text{ELEM1 } \boxed{3} \\ \text{LIST2 } \boxed{4} \\ \text{ELEM2 } \boxed{5} \end{array} \right] \end{array} \right]$$

⁸⁶ Siegel modelliert die SUBCAT-Mengen allerdings nicht — wie hier vorgeschlagen — durch rekursive Merkmalstrukturen (*Mengen*), sondern nichtrekursiv mittels der Merkmale SUBJ, OBJ, OBJ2, ...:

(192) SUBCAT-Menge:

$$\{ [subj, \dots], [obj, \dots], [obj2, \dots], \dots \}$$

SUBCAT-Array:

$$\left[\begin{array}{c} \text{SUBJ } [\dots] \\ \text{OBJ } [\dots] \\ \text{OBJ2 } [\dots] \\ \dots \end{array} \right]$$

den: *ungesättigt* (Element der SUBCAT-Liste) und *gesättigt* (aus der SUBCAT-Liste gestrichen). Diese Art der Modellierung eignet sich somit nicht für Sprachen wie das Japanische, in denen Komplemente meist nicht obligatorisch sind. Für Sprachen wie das Deutsche oder Englische ist sie hingegen ausreichend.⁸⁷

Wortstellung

Strenggenommen dienen Syntaxbäume nur der Darstellung der Abhängigkeitsverhältnisse zwischen Verb und Komplementen. Implizit haben wir jedoch die Reihenfolge der terminalen Töchter im Syntaxbaum mit der Oberflächenreihenfolge der Wörter im Satz identifiziert.⁸⁸ Noch deutlicher wird diese Annahme durch die Modellierung der Subkategorisierung mittels einer Menge: Durch das Head-Komplement-Schema kann ein beliebiges Element der SUBCAT-Menge ausgewählt und realisiert werden. Die dadurch ermöglichten Strukturen entsprechen genau den Stellungsmöglichkeiten der Konstituenten in einem Satz:

(193) *Der Student liest das Buch.*

a. 学生が本を読む。

gakusei ga hoñ wo yomu
Student-NOM buch-ACC lesen

b. 本を学生が読む。

hoñ wo gakusei ga yomu
buch-ACC Student-NOM lesen

Japanisch ist eine Head-finale Sprache: Der Head einer Phrase steht immer am Ende. Von dieser Einschränkung abgesehen ist die Wortstellung

87 Auch in Grammatiken des Deutschen oder Englischen wird von *fakultativen Komplementen* gesprochen (vgl. Helbig und Buscha 1998 oder Engel 1994). In den meisten Fällen dienen diese jedoch der vereinfachten Beschreibung von Verben, die in leicht veränderter Bedeutung mit verschiedenen Kasusrahmen stehen können: Wird der Begriffes *Fakultativität* in diesem Sinne verwendet, läßt sich *essen* beispielsweise als Verb mit fakultativem Akkusativobjekt beschreiben.

Um dieser Interpretation der Fakultativität vorzubeugen, wird das “Weglassen” von Komplementen im Japanischen von vielen Autoren, durch Begriffe wie *Ellipse*, *Nullanapher* oder *Nullpronomen* beschrieben (vgl. beispielsweise Siegel 1997).

88 Für einen Ansatz, der die Reihenfolge der terminalen Töchter von der Reihenfolge ihres Auftretens trennt, vgl. Gunji 1999.

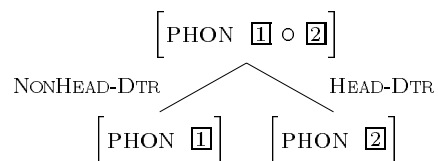
Das von Gunji vorgeschlagene “morphologische Prinzip” erlaubt eine interessante Interaktion von Syntax und Semantik — und damit eine ganz andere Art der “Trennung” (oder besser Parallelität) dieser Ebenen als hier vorgeschlagen. (Zu dem hier vorgeschlagenen Ansatz vergleiche Gunji 1999 mit Manning et al. 1999.)

im Japanischen frei: Die Komplemente des Heads können also in beliebiger Reihenfolge stehen.⁸⁹ Genau diese Wortstellung ergibt sich, wenn die Head Tochter in lokalen Bäumen immer rechts von der Nicht-Head-Tochter angeordnet wird.

PHON-Merkmal Prinzip

In der AVM-Darstellung wird die phonologische Form eines Zeichens im Merkmal PHON kodiert. Um die gleiche Reihenfolge der Wörter zu erhalten, wie bei der Darstellung als Baum, müssen die PHON-Werte der Töchter in lokalen Bäumen in der richtigen Reihenfolge konkateniert und dem PHON-Wert der Mutter zugewiesen werden. Dieses einfache Stellungsprinzip läßt sich bei der derzeitigen Formulierung der syntaktischen Schemata durch einen einzigen Constraint zum Typ *phrase* realisieren: das sogenannte PHON-Merkmal-Prinzip.

(194) *Das PHON-Merkmal Prinzip.*⁹⁰



Schemata und lexikalische Information

Lexikalischer Head H einer syntaktischen Struktur P heißt das Zeichen, das sich ergibt, wenn die Linie der Head-Töchter von P bis zu einem lexikalischen Zeichen verfolgt wird. Alle Konstruktionen, die durch rekursive Anwendung von Schemata aus dem lexikalischen Head H entstehen, heißen *Projektionen* von H.

Das Verhalten des Head-Komplement-Schemas beruht weitgehend auf Informationen, die durch die Töchter geliefert werden. Insbesondere der SUBCAT-Wert der Head-Tochter spielt dabei eine wesentliche Rolle. Dieser Wert geht — modifiziert durch andere Anwendungen des Head-Komplement-Schemas — auf den SUBCAT-Wert des lexikalischen Heads zurück.

⁸⁹ Es existiert eine Standardreihenfolge der Konstituenten im Japanischen. Eine von dieser abweichende Konstituentenreihenfolge kann mit einer Verschiebung der Bedeutung, einer Betonung bestimmter Elemente etc. verbunden sein. Von diesen Aspekten der Wortstellung wird hier jedoch abstrahiert.

⁹⁰ Zur Realisierung der Konkatenation $\circ$ vergleiche Abschnitt “Listen und Listen-Konkatenation”, S. 47.

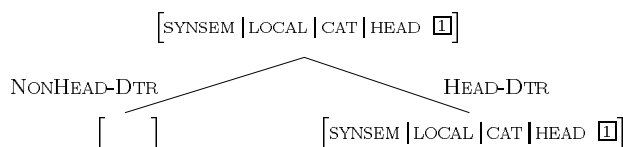
Das HEAD-Merkmal-Prinzip

Der SUBCAT-Wert komplexer Zeichen wird vom Head-Komplement-Schema aus dem entsprechenden Wert seiner Head-Tochter berechnet. Eine Reihe anderer Merkmale werden direkt von der Head-Tochter übernommen und daher als HEAD-Merkmale bezeichnet. In unserem Fall gehören dazu zwei Merkmale: MAJ, das Merkmal zur Darstellung der lexikalischen Kategorie, das für Bezeichnungen wie *Verb*, *Verbalphrase*, *Nomen*, *Nominalphrase* verantwortlich ist, und MOD, das im Abschnitt über Adjunkte näher erklärt wird:

$$\begin{bmatrix} \text{head} \\ \text{MAJ} \quad \text{maj} \\ \text{MOD} \quad \text{mod-synsem} \end{bmatrix}$$

Das Prinzip, das für den Transport der Head-Merkmale zur Mutter verantwortlich ist, heißt entsprechend HEAD-Merkmal-Prinzip und läßt sich als weiterer Constraint des Types *phrase* formalisieren. Da alle HEAD-Merkmale unter dem Merkmal HEAD zusammengefaßt sind, hat es eine einfache Form:

(195) *Das HEAD-Merkmal-Prinzip:*



Die letzten Abschnitte veranschaulichen die zentrale Bedeutung lexikalischer Informationen für das vorgestellte syntaktische Modell:⁹¹ Durch die Darstellung der grammatischen Informationen im Lexikon können die Schemata allgemein gehalten werden. Der Unterschied im syntaktischen Verhalten von Verben beispielsweise wird durch das Zusammenspiel von Head-Komplement-Schema, HEAD-Merkmal-Prinzip und lexikalischen Informationen modelliert — in diesem Fall insbesondere durch den SUBCAT-Wert des entsprechenden lexikalischen Eintrags:⁹²

91 Die *Lexikalisierung* der Informationen ist eine häufig zu beobachtende Tendenz moderner Theorien, vgl. beispielsweise Chomsky 1981, Pollard und Sag 1994 oder Gunji 1987.

92 Der SUBCAT-Wert eines Verbes definiert dessen *Valenz*. Eine Liste der gebräuchlichsten japanischen Verben mit Angabe ihrer Valenz und der Abhängigkeiten zwischen Valenz und Bedeutung findet sich in Rickmeyer 1977. Die Verwendung der Verben im Zusammenhang mit den verschiedenen möglichen Kasusrahmen wird in diesem Lexikon anhand von Beispielen und einer semantischen Prädikatorensprache veranschaulicht.

(196) Verb:	Valenz:
寝る (schlafen): <i>neru</i>	$\left[\text{SUBCAT } \{ \text{PP}[ga]^{opt} \} \right]$
食べる (essen): <i>taberu</i>	$\left[\text{SUBCAT } \{ \text{PP}[ga]^{opt}, \text{PP}[wo]^{opt} \} \right]$
書く (schreiben): <i>kaku</i>	$\left[\text{SUBCAT } \{ \text{PP}[ga]^{opt}, \text{PP}[ni]^{opt}, \text{PP}[wo]^{opt} \} \right]$

Am Ende dieses Abschnitts werden wir noch einige andere Anwendungen des Head-Komplement-Schema sehen, die seine Universalität im Zusammenhang mit lexikalischen Informationen veranschaulichen.

Kategorien

Der Begriff *Kategorie* wird in dieser Arbeit in zwei verschiedenen Zusammenhängen verwendet:

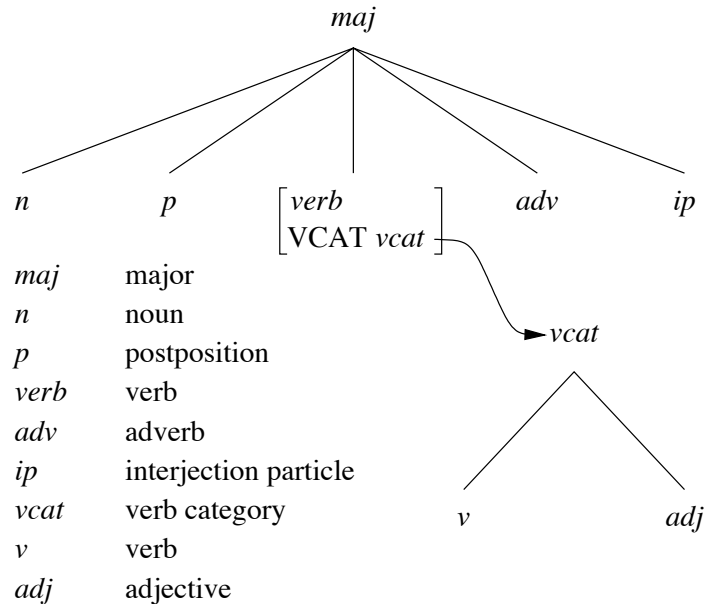
1. als *lexikalische Kategorie* — beispielsweise *v* für Verb, *n* für Nomen usw. Lexikalische Kategorien werden durch das Merkmal MAJ bezeichnet.
2. als *syntaktische Kategorie* — wie wir später sehen werden, entsprechen die syntaktischen Kategorien einer Kombination des MAJ-Wertes mit anderen Merkmalwerten.⁹³

Die Unterteilung der Wörter entsprechend der lexikalischen Kategorie ist für die Charakterisierung der Elemente der SUBCAT-Menge wesentlich. In unserem Modell kommen wir mit fünf lexikalischen Kategorien aus (Verben i.w.S. unterteilen sich weiter in Verben i.e.S. und Adjektive):

⁹³ Die *syntaktische Kategorie* wird in HPSG durch das Merkmal CAT modelliert: Unter CAT werden die Merkmale zusammengefaßt, die die syntaktische Kategorie definieren — beispielsweise das MAJ- und das SUBCAT-Merkmal.

In der neueren Version der HPSG (Pollard und Sag 1994) wird MAJ nicht als eigenes Merkmal, sondern als Subtyp von *head*, dem Wertetyp von HEAD, modelliert.

(197)



Wir unterscheiden:

(198) *Kategorien:*

1. *n* — Nomen:

z.B.:	雨	<i>ame</i>	Regen
	映画	<i>eiga</i>	Kino
	手紙	<i>tegami</i>	Brief
2. *verb* — Verben i.w.S.:
 - a. *v* — Verben i.e.S.:

z.B.:	買う	<i>kau</i>	kaufen
	泳ぐ	<i>oyogu</i>	schwimmen
	読む	<i>yomu</i>	lesen
 - b. *adj* — Adjektive:

z.B.:	広い	<i>hiro</i>	weit, breit
	高い	<i>takai</i>	hoch, teuer
	安い	<i>yasui</i>	billig
3. *adv* — Adverben:

z.B.:	早く	<i>hayaku</i>	schnell
	すぐ	<i>sugu</i>	sofort
	良く	<i>yoku</i>	gut
4. *p* — Postpositionen:

z.B.:	が	<i>ga</i>	〈Subjekt〉
	を	<i>wo</i>	〈direktes Objekt〉
	に	<i>ni</i>	〈indirektes Objekt〉

5. *ip* — Interjektionspartikeln:

z.B.:	ね	<i>ne</i>	“[...], nicht war?”
	か	<i>ka</i>	⟨Frage⟩
	わ	<i>wa</i>	⟨Frage (Frauensprache)⟩

Postpositionen bilden mit Nominalphrasen zusammen die schon vorgestellten *Postpositionalphrasen*.

Interjektionspartikel sind ein Charakteristikum der japanischen Sprache. Sie bezeichnen u.a. die Einstellung des Sprechenden gegenüber dem im Satz ausgesagten Sachverhalt, kennzeichnen Fragen⁹⁴ etc.

Die Kategorie *adv* wird nur für lexikalische Adverben benutzt. Sie spielt in der Syntax keine Rolle: Adverben sind Adjunkte, und Adjunkte wählen ihre Bezugswörter. Wesentlich für Adverben ist also nicht ihre Kategorie, sondern das Vermögen, Bezugswörter zu kennzeichnen. Wie dieses geschieht, werden wir im Abschnitt 3.2.2 “Adjunkte” beschreiben.

Kasus

Der *Kasus* wird im Japanischen durch Postpositionen ausgedrückt. Wir haben Kasus und Postpositionen schon als kennzeichnendes Merkmal von Postpositionalphrasen kennengelernt:

(199)	Funktion:	Kasus / Postposition:
	<i>Subjekt</i> :	PP[<i>ga</i>]
	<i>direktes Objekt</i> :	PP[<i>wo</i>]
	<i>indirektes Objekt</i> :	PP[<i>ni</i>]
	. . .	

Der Kasus entspricht — abgesehen von den Postpositionen *ha* und *mo* — der Oberflächenform der Postpositionen. Er wird durch das Merkmal CASE beschrieben.

Die Architektur der Zeichen

Merkmale lassen sich in Untergruppen unterteilen: Es gibt syntaktische Merkmale, semantische Merkmale etc. Prinzipien und Schemata beziehen sich oft auf solche Untergruppen. Merkmale werden daher in rekursiven Strukturen hierarchisch gegliedert.

In der Gliederung der Merkmale spiegelt sich das gesamte sprachliche Modell: Prinzipien und Schemata sind auf die Struktur der Zeichen angewiesen; die

⁹⁴ Wir fassen auch Fragepartikel wie か (*ka*) und わ (*wa*) als Interjektionspartikel auf.

Struktur der Zeichen hingegen läßt sich ohne die Prinzipien nicht verstehen. Um einen Aspekt der Theorie zu erklären, müssen deshalb oft andere, erst später thematisierte Aspekte vorläufig unerklärt hingenommen werden.

Bisher wurden erst die Merkmale PHON, SUBCAT, MAJ und CASE eingeführt. Trotzdem sollen diese Merkmale — alle übrigen Prinzipien und Schemata vorwegnehmend — in den Zeichen aus Beispiel (183) mit vollständigen Pfaden dargestellt werden. Die beteiligten Merkmale werden dabei nur mit einigen Worten vorgestellt:⁹⁵

(200) Zeichen für das Verb 読む (*yomu*; lesen) und seine Komplemente:

- a.
$$\left[\begin{array}{l} \text{phrase} \\ \text{PHON } \langle \text{gakusei } ga \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \text{MAJ } p \\ \text{CASE } ga \end{array} \right] \\ \text{SUBCAT } \{ \text{NP}^{sat} \} \end{array} \right] \end{array} \right]$$
- b.
$$\left[\begin{array}{l} \text{phrase} \\ \text{PHON } \langle \text{hoñ } wo \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \text{MAJ } p \\ \text{CASE } wo \end{array} \right] \\ \text{SUBCAT } \{ \text{NP}^{sat} \} \end{array} \right] \end{array} \right]$$
- c.
$$\left[\begin{array}{l} \text{word} \\ \text{PHON } \langle \text{yomu} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{l} \text{HEAD} \mid \text{MAJ} \left[\begin{array}{l} \text{verb} \\ \text{VCAT } v \end{array} \right] \\ \text{SUBCAT } \{ \text{PP}[ga]:\text{I}^{opt}, \text{PP}[wo]:\text{II}^{opt} \} \end{array} \right] \end{array} \right]$$

Es folgt eine kurze Beschreibung der vorkommenden Merkmale:⁹⁶

⁹⁵ Die vorgeschlagene Architektur syntaktischer Zeichen entspricht einer Kombination von Aspekten der HPSG (Pollard und Sag 1987, Pollard und Sag 1994), JPSG (Gunji 1987, Gunji 1991) und der japanischen Grammatik des Verbmobilprojektes (vgl. Siegel 1996b).

⁹⁶ Zum Teil wurden die Merkmale schon vorgestellt, zum Teil werden sie in späteren Abschnitten ausführlicher erklärt. Merkmale, die nicht Teil eines Pfades zu einem schon eingeführten Merkmal (PHON, SUBCAT, MAJ und CASE) sind, werden an späterer Stelle behandelt.

PHON — beschreibt die phonologische Realisierung des beschriebenen Zeichens in Form einer Liste.

SYNSEM — faßt die syntaktischen und semantischen Merkmale zusammen.

LOCAL — dient zur Abgrenzung von NONLOCAL. NONLOCAL wird im Zusammenhang mit nicht-lokalen Konstituenten benutzt.⁹⁷ Unter LOCAL werden alle anderen Eigenschaften der Zeichen zusammengefaßt.

CAT — Das Merkmal MAJ zur Beschreibung der lexikalischen Kategorie wurde schon eingeführt. In einem weiteren Sinne lassen sich jedoch auch andere Eigenschaften zur Kategorie eines Zeichen zählen: beispielsweise sein Sättigungsgrad, der durch das Merkmal SUBCAT dargestellt wird,⁹⁸ oder das Merkmal MOD, das beschreibt, ob sich das Zeichen als Adjunkt verwenden läßt.

HEAD — Der Bezeichnung *Head* sind wir schon bei der Unterscheidung der Töchter komplexer Zeichen in *Head-Tochter* und *Non-Head-Tochter* begegnet. Die Head-Merkmale motivieren sich aus der Beobachtung, daß bestimmte wesentliche Eigenschaften komplexer Zeichen mit denen ihrer Head-Tochter übereinstimmen. Diese Eigenschaften werden unter dem Merkmal HEAD gruppiert und durch das *Head-Merkmal-Prinzip*⁹⁹ von der Head-Tochter übernommen.

MAJ — Das Merkmal MAJ wurde im letzten Abschnitt eingeführt. Die Werte von MAJ entsprechen den traditionellen Wortarten. Sie sind Head-Merkmale und werden daher bei komplexen Zeichen von der Head-Tochter übernommen.

CASE — Der KASUS wird im Japanischen durch Postpositionen realisiert. Dieses Merkmal existiert daher nur bei ihren Projektionen (P und PP).

Das Head-Komplement-Schema kann nun in seiner endgültigen Form beschrieben werden:

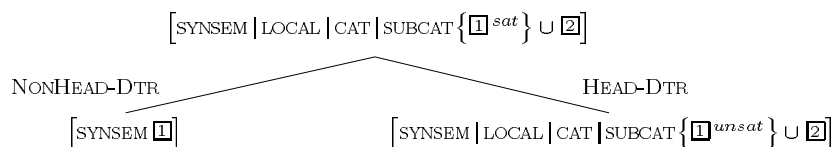
(201) *Das Head-Komplement-Schema:*

⁹⁷ In Abschnitt 3.2.3 wird die Modellierung nichtlokaler Konstituenten am Beispiel des Topic-Partikels *ha* (*ha*) dargestellt.

⁹⁸ In der Standardversion der HPSG wird das *Kategoriensymbol* VP beispielsweise für Zeichen mit [MAJ *v*] verwendet, wenn alle Objekte, abgesehen vom Subjekt, gesättigt sind. S hingegen entsteht aus VP durch Sättigung des Subjektes.

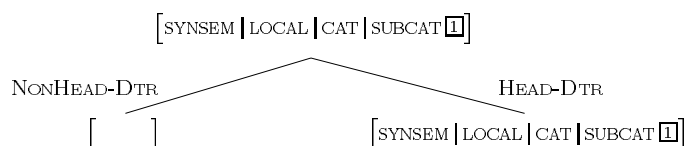
Da im Japanischen die Objekte im Allgemeinen fakultativ sind, lassen sich beide Symbole nur begrenzt auf das Japanische übertragen.

⁹⁹ Vgl. den entsprechenden Abschnitt auf S. 118.



Head-Adjunkt-Schema und Head-Füller-Schema — beide werden in späteren Kapiteln vorgestellt — lassen den SUBCAT-Wert unberührt. In lokalen Bäumen, die durch diese Schemata legitimiert werden, muß sein Wert daher durch den folgenden Constraint an die Mutter weitergegeben werden:

(202) *Constraint zur Berechnung des SUBCAT-Wertes
für Head-Adjunkt- und Head-Füller-Schema:*



(201) unterscheidet sich nur in zwei Punkten von der ersten Version in (189):

1. Das Subkategorisierungsmerkmal SUBCAT wird durch seinen gesamten Pfad dargestellt;
2. Nicht die gesamte Head-Tochter wird mit dem Element aus der SUBCAT-Menge unifiziert, sondern nur ihr SYNSEM-Wert.

Alle Merkmale außerhalb von SYNSEM sind für die Subkategorisierung nicht interessant. In unserem Modell gibt es allerdings nur ein einziges Merkmal, das nicht unter SUBCAT fällt: Das Merkmal PHON. SYNSEM wird nur verwendet, um nicht von der Standardversion der HPSG abzuweichen.

Abkürzungen

In manchen Kontexten führt die Verwendung expliziter Zeichen zu Unübersichtlichkeiten: Eine SUBCAT-Menge beispielsweise wird durch die Angabe der vollständigen Zeichen unlesbar, ein Text unleserlich. Wir definieren daher Abkürzungen, um Zeichen einfacher zu charakterisieren:

1. Abkürzungen der Kategorie

Anstatt ein Nomen umständlich durch

$$(203) \quad \left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } n \right]$$

zu beschreiben, schreiben wir das Kategoriensymbol n groß: N. Indem wir für die anderen Kategorien entsprechend verfahren, erhalten wir folgende Kurzformen:

(204)	Abkürzung:	abgekürzte AVM:
	N	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } n \right]$
	P	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } p \right]$
	Vcat	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } verb \right]$
	ADV	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } adv \right]$
	IP	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } ip \right]$

Verben im engeren Sinne und Adjektive kürzen wir folgendermaßen ab:

(205)	Abkürzung:	abgekürzte AVM:
	V	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } \left[\begin{array}{l} verb \\ \text{VCAT } v \end{array} \right] \right]$
	ADJ	$\left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MAJ } \left[\begin{array}{l} verb \\ \text{VCAT } adj \end{array} \right] \right]$

2. Die Sättigung von Zeichen

Subkategorisiert wird in den meisten Fällen für gesättigte Zeichen — also Zeichen, die zumindest hinsichtlich ihrer obligatorischen Komplemente gesättigt sind. Um diese zu kennzeichnen, hängen wir an das Kategoriensymbol den Buchstaben P an. Den Erklärungen im Abschnitt “Adjunkte” vorgehend, wird dabei auch das zweite Subkategorisierungsmerkmal ADJS mit dargestellt:

(206) Abkürzung: abgekürzte AVM:

$$XP \quad X \sqcup \left[\begin{array}{c} \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{cc} \textit{cat} & \\ \text{SUBCAT} & \textit{saturated}^{100} \\ \text{ADJS} & \textit{saturated} \end{array} \right] \end{array} \right]$$

Wenn ein *Satz* gemeint ist — also über ein saturiertes Verb gesprochen wird, das als Satz verstanden werden soll und nicht als Basis weiterer Anwendungen des Head-Komplement-Schemas — verwenden wir jedoch anstelle von VP das Symbol S; VP hingegen wird für verbale Zeichen benutzt, die als *Verbalphrase* verstanden werden sollen. Da Subjekte fast immer fakultativ sind, lassen sich die meisten VPs — formal gesehen — auch als S beschreiben; Ein S hingegen kann nur dann als VP beschrieben werden, wenn sein Subjekt nicht gesättigt wurde. S und VP enthalten also oft eine “Interpretation” einer Wortfolge als Satz oder als Verbalphrase.

3. Andere Merkmalwerte

Bestimmte Merkmalwerte zur Charakterisierung von Zeichen werden ohne Angabe von Pfad und Merkmal den Abkürzungen in eckigen Klammern hinterhergestellt. Das *ga* in PP[*ga*] beispielsweise steht für

100 *saturated* ist eine Abkürzung:

(207) Abkürzung: abgekürzte Merkmalstruktur:

$$\left[\text{LIST } \textit{saturated} \right] \quad \left[\begin{array}{c} \textit{versuch} \\ \text{LIST } \boxed{1} \\ \text{CONDITION } \left[\begin{array}{c} \textit{saturated} \\ \text{LIST } \boxed{1} \end{array} \right] \end{array} \right]$$

Sie läßt sich durch rekursive Typenconstraints folgendermaßen definieren:

(208) **Constraints für den Typ *saturated*:**

$$\left[\begin{array}{c} \textit{saturated} \\ \text{LIST } \langle \rangle \end{array} \right] \quad \left[\begin{array}{c} \textit{saturated} \\ \text{LIST } \left\langle \left[\begin{array}{c} \textit{sc-elem} \\ \text{SAT } \textit{sat} \vee \textit{opt} \end{array} \right] \bullet \boxed{1} \right\rangle \\ \text{CONDITION } \left[\begin{array}{c} \textit{saturated} \\ \text{LIST } \boxed{1} \end{array} \right] \end{array} \right]$$

$\vee$ wird hier als Abkürzung für zwei alternative AVMs verwendet, die sich nur bezüglich der durch $\vee$ verbundenen Werte unterscheiden. $\vee$ läßt sich also “ausmultiplizieren”.

$$[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{CASE } ga],$$

so daß sich insgesamt die folgende Entsprechung ergibt:

(209) Abkürzung: abgekürzte AVM:

$$PP[ga] \quad \left[\begin{array}{c} \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \end{array} \left[\begin{array}{c} cat \\ \text{HEAD} \left[\begin{array}{c} p\text{-head} \\ \text{MAJ } p \\ \text{CASE } ga \end{array} \right] \\ \text{SUBCAT } \textbf{saturated} \\ \text{ADJS } \textbf{saturated} \end{array} \right] \right]$$

Kontextabhängigkeit der Abkürzungen

Die Abkürzungen sind Kontextabhängig: In Subkategorisierungskontexten — beispielsweise zur Darstellung eines SUBCAT-Elementes — stehen sie für den Wert des Merkmals SYNSEM; außerhalb derartiger Kontexte bezeichnen sie das gesamte Zeichen.

Meistens ergibt sich das Gemeinte aus dem Kontext; ist dieses nicht der Fall, wird explizit auf die intendierte Bedeutung hingewiesen.

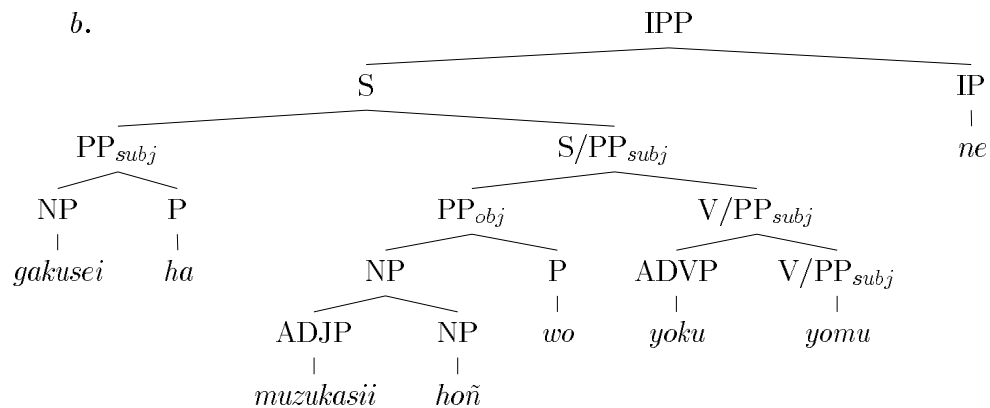
Beispiele für das Head-Komplement-Schema

Bisher haben wir das Head-Komplement-Schema im Zusammenhang mit Verben und ihren Komplementen kennengelernt. Die Eigenschaft, für Komplemente zu subkategorisieren, ist jedoch nicht auf Verben beschränkt; vielmehr läßt sich — zumindestens in unserem Modell — der größte Teil der syntaktischen Konstruktionen auf dieses Schema zurückführen.

Kehren wir zu unserm Beispieltext (168) zurück — hier als (210) noch einmal dargestellt:

(210) a. 学生は難しい本を良く読むね。

gakusei-ha muzukasii hoñ-wo yoku yomu ne
 Student-TOP schwer buch-ACC gut lesen nicht wahr
 Der Student liest das schwierige Buch gut, nicht wahr?



In diesem Satz lassen sich drei weitere Beziehungen durch das Head-Komplement-Schema erklären:

(211)	NONHEAD — Komplement —	HEAD — subkategorisiert für das Komplement —
a.	学生 <i>gakusei</i>	は <i>ha</i>
b.	難しい 本 <i>muzukasii hon</i>	を <i>wo</i>
c.	学生 は 難しい 本 を よく 読む <i>gakusei ha muzukasii hon wo yoku yomu</i>	ね <i>ne</i>

Postpositionen subkategorisieren für Nominalphrasen:

Die Postpositionalphrasen (211a) und (211b) entstehen aus der Komposition einer Nominalphrase und einer Postposition:

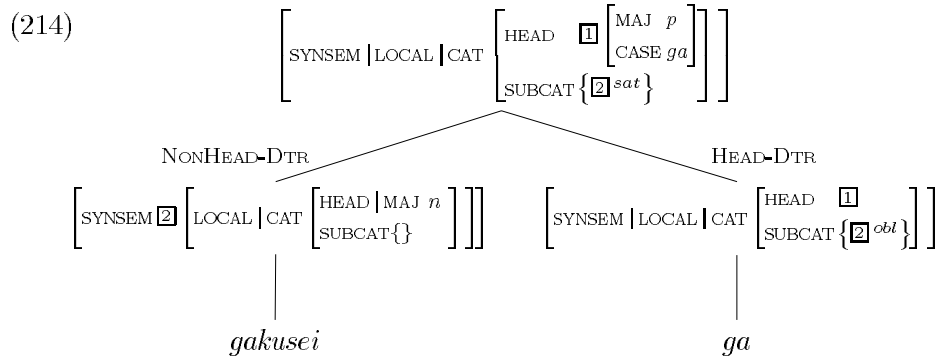
(212) a.	i. 学生は <i>gakusei ha</i> Student TOP der Student	ii. 難しい本を <i>muzukasii-hon wo</i> schwierige-Buch ACC das schwierige Buch
b. i.	<pre> graph TD PP1 --> NP1[NP] PP1 --> P1[P] NP1 --> gakusei P1 --> ha </pre>	<pre> graph TD PP2 --> NP2[NP] PP2 --> P2[P] NP2 --> muzukasii_hon[muzukasii-hon] P2 --> wo </pre>

Postpositionen werden vor allem durch ihren Kasus charakterisiert. Der Kasus wird in den meisten Fällen durch die Oberflächenform der Postposition

bezeichnet — also durch $\text{CASE} \in \{ga, wo, he, ni, \dots\}$. Die Postposition ist der Head der Postpositionalphrase und subkategorisiert für die Nominalphrase:

$$(213) \quad \left[\begin{array}{l} \text{PHON } \langle ga \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \text{MAJ } p \\ \text{CASE } ga \end{array} \right] \\ \text{SUBCAT } \{ \text{NP:} \boxed{1}^{obl} \} \end{array} \right] \end{array} \right]$$

CASE und MAJ sind Head-Merkmale. Ihre Werte vererben sich daher über das Head-Merkmal-Prinzip von der Postposition auf die Postpositionalphrase. Als Merkmalstruktur läßt sich (212*b-i*) damit folgendermaßen darstellen:



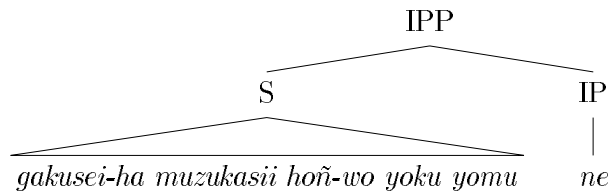
Interjektionspartikel subkategorisieren für Sätze:

In unserer Syntax werden Interjektionspartikel als Head einer entsprechenden Phrase — der Interjektionspartikelphrase IPP — aufgefaßt:

(215) a. 学生は難しい本を良く読むね。

gakusei-ha muzukasii hoñ-wo yoku yomu ne
 Der Student liest das schwierige Buch gut nicht wahr?
 Der Student liest das schwierige Buch gut, nicht wahr?

b.



Interjektionspartikel werden daher als Zeichen modelliert, die für Sätze — d.h. saturierte verbale Elemente — subkategorisieren:

$$(216) \quad \left[\begin{array}{c} \text{PHON } \langle ne \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } ip \\ \text{SUBCAT } \{S\} \end{array} \right] \end{array} \right]$$

3.2.2 Adjunkte

Komplemente vs. Adjunkte

Syntaktische Kategorien wie *Nomen*, *Verb*, *Nominalphrase*, *Verbalphrase*, *Satz* etc. werden durch den Wert des Merkmals CAT beschrieben¹⁰¹ — und insbesondere durch die Kombination der Merkmale MAJ und SUBCAT. Das Head-Komplement-Schema modifiziert den SUBCAT-Wert — und damit die syntaktische Kategorie. In lokalen Bäumen, die durch dieses Schema legitimiert werden, sind deshalb die syntaktischen Kategorien von Mutter und Head-Tochter verschieden. Wenn man unter diesem Aspekt den Syntaxbaum unseres Beispielsatzes (210) noch einmal betrachtet, fallen zwei Verzweigungen auf, in denen sich die Kategorien von Mutter und Head-Tochter nicht unterscheiden:

- (217) a. i. 難しい本
 muzukasii hoñ
 schwer buch
 schwieriges Buch
- ii.
$$\begin{array}{c} \text{NP} \\ \swarrow \quad \searrow \\ \text{ADJP} \quad \text{NP} \\ | \quad | \\ \text{muzukasii} \quad \text{hoñ} \end{array}$$
- b. i. 良く読
 yoku yomu
 gut lesen
 liest gut
- ii.
$$\begin{array}{c} \text{VP/PP}_{\text{subj}} \\ \swarrow \quad \searrow \\ \text{ADVP} \quad \text{VP/PP}_{\text{subj}} \\ | \quad | \\ \text{yoku} \quad \text{yomu} \end{array}$$

Die Nicht-Head-Töchter dieser lokalen Bäume werden als *Adjunkte* bezeichnet und das Schema, durch welches sie syntaktisch eingebunden werden, als Head-Adjunkt-Schema.¹⁰²

101 vgl. die Abgrenzung der Begriffe *lexikalische Kategorie* — dargestellt durch das Merkmal MAJ — und *syntaktische Kategorie* — dargestellt durch CAT (vgl. Paragraph “Kategorien”, S. 119).

102 Modifizierende Ausdrücke stehen im Japanischen vor dem modifizierten Element: Die modifizierte Form hat die gleiche Kategorie wie die nicht modifizierte Form. Da die Kategorie durch Head-Merkmale beschrieben wird, muß das modifizierte Element also Head des lokalen Baumes sein.

Syntaktisch verhalten sich Adjunkt und Bezugswort zusammen genauso wie das nicht modifizierte Bezugswort für sich allein. Da das syntaktische Verhalten durch die syntaktische Kategorie bestimmt wird, muß diese bei Mutter und Head-Tochter gleich sein.

Traditionell wird die Möglichkeit eines Wortes oder einer Wortgruppe, als Adjunkt verwendet werden zu können, durch die lexikalische Kategorie kodiert: Die Bezeichnung *Adverb* läßt sich durch “*Adjunkt zum Verb*” umschreiben und *Adjektiv* als “*Adjunkt (Attribut) zum Nomen*”. Da die lexikalische Kategorie durch den MAJ-Wert des Zeichens dargestellt wird, entsprechen dieser Art der grammatischen Kodierung Regeln wie die folgenden:

$$(218) \quad \begin{array}{ll} a. & VP \longrightarrow V \text{ NP Adv.} \\ b. & N \longrightarrow \text{Adj N.} \end{array}$$

bzw. als SUBCAT-Menge ausgedrückt:

$$(219) \quad \begin{array}{ll} a. & \left[\begin{array}{l} \dots | \text{MAJ} \quad v \\ \dots | \text{SUBCAT} \{ \text{NP}, \text{Adv} \} \end{array} \right] \\ b. & \left[\begin{array}{l} \dots | \text{MAJ} \quad n \\ \dots | \text{SUBCAT} \{ \text{Adj} \} \end{array} \right] \end{array}$$

Beiden Arten der Modellierung ist gemeinsam, daß durch eine Regel oder das Bezugswort *für* das Adjunkt subkategorisiert wird, und nicht das Adjunkt für sein Bezugswort subkategorisiert.

In dem hier vorgestellten Modell und seinen Vorbildern HPSG und JPSG wird deshalb die lexikalische Kategorie von der Aufgabe entlastet, die Verwendung einer Wortgruppe als Adjunkt zu kennzeichnen. Die Eigenschaft, die den adverbialen oder adnominalen Gebrauch von Elementen legitimiert, wird durch ein eigenes Merkmal (MOD) beschrieben. Im Japanischen wird dieses Merkmal durch die Flexionsform des lexikalischen Heads einer Konstruktion bestimmt: Verben und Adjektive beispielsweise können Satzfinal-, Adnominal- oder Adverbialformen bilden und entsprechend als Zentrum eines Hauptsatzes oder als Adjunkt verwendet werden.

Anders als beim Head-Komplement-Schema ist es also nicht die Head-Tochter, die die Nicht-Head-Tochter legitimiert, sondern umgekehrt die Nicht-Head-Tochter, die die Head-Tochter legitimiert. Das modifizierende Element selbst ist also verantwortlich für seine Verwendung als Adjunkt und wählt sein Bezugselement — oder anders ausgedrückt: Adjunkte *subkategorisieren* für ihre Bezugselemente.

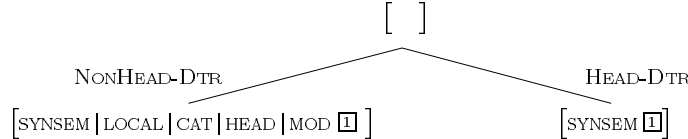
Modellierung von Adjunkten

Die Eigenschaft der Adjunkte, für Bezugselemente zu subkategorisieren, wird durch das Merkmal MOD beschrieben. Die beiden Adjunkte 難しい (*muzukasii*; schwierig) und 良く (*yoku*; gut) aus (210) beispielsweise lassen sich damit vereinfacht folgendermaßen beschreiben:

- (220) a.
$$\left[\begin{array}{l} \text{PHON } \langle \text{muzukasii} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MOD NP} \end{array} \right]$$
- b.
$$\left[\begin{array}{l} \text{PHON } \langle \text{yoku} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \mid \text{MOD Vcat} \end{array} \right]$$

Ähnlich wie die Elemente der SUBCAT-Menge, die die möglichen Elemente des Heads beschreiben, beschreibt der MOD-Wert die möglichen Bezugswörter des Adjunktes. Genau wie bei Komplementen wird diese Spezifikation durch die Unifikation mit dem SYNSEM-Wert des Bezugswortes überprüft.

(221) *Attribute wählen ihre Bezugswörter:*



Zeichen, die nicht als Adjunkt benutzt werden können, werden durch [MOD *none*] gekennzeichnet. MOD kann also entweder Strukturen des Typs *synsem* oder *none* als Wert annehmen. Um dieses formal darzustellen, wird der Typ *mod-synsem* eingeführt, und MOD als Wertspezifikation zugewiesen:

- (222) a.
$$\left[\begin{array}{l} \text{head} \\ \dots \\ \text{MOD } \text{mod-synsem} \end{array} \right]$$
- b.
$$\begin{array}{c} \text{mod-synsem} \\ \swarrow \quad \searrow \\ \text{none} \quad \text{synsem} \end{array}$$

Adjunkte und Satzstellung

Adverbien haben im Japanischen eine relativ freie Wortstellung. So kann das Adverb 良く (*yoku*; gut) in unserem Beispielsatz an folgenden Stellen stehen:¹⁰³

¹⁰³ Die Stellung des Adverbis kann seinen Skopus beeinflussen oder mit einer Verschiebung der Bedeutung verbunden sein. Von diesen Phänomenen wird hier jedoch abstrahiert.

(223) mögliche Stellungen des Adverbs 良く (*yoku*; gut) im Satz
Der Student liest das Buch gut:

a. 良く学生が本を読む。

<i>yoku</i>	<i>gakusei-ga</i>	<i>hoñ-wo</i>	<i>yomu</i>
gut	Student-NOM	buch-ACC	lesen

b. 学生が良く本を読む。

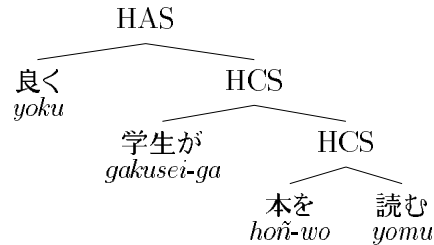
<i>gakusei-ga</i>	<i>yoku</i>	<i>hoñ-wo</i>	<i>yomu</i>
Student-NOM	gut	buch-ACC	lesen

c. 学生が本を良く読む。

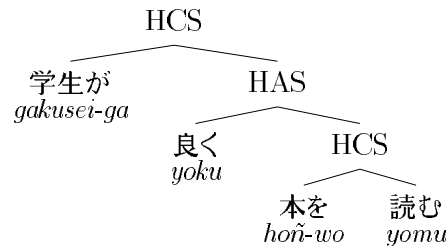
<i>gakusei-ga</i>	<i>hoñ-wo</i>	<i>yoku</i>	<i>yomu</i>
Student-NOM	buch-ACC	gut	lesen

Diese Variation in der Wortstellung kann dadurch erfaßt werden, daß Ad-
 verben nur für die Aspekte ihrer Bezugszeichen subkategorisieren, die nicht
 durch das Head-Komplement-Schema beeinflußt werden: beispielsweise für
 den MAJ-, nicht aber den SUBCAT-Wert von Zeichen. Da das Head-Adjunkt-
 Schema die Eigenschaften unverändert läßt, die das Head-Komplement-Sche-
 ma steuern, kann es an beliebiger Stelle mit dem Head-Komplement-Schema
 interferieren:

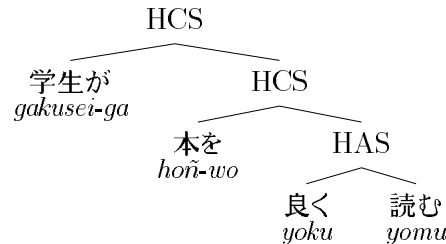
(224) a.



b.



c.



Head-Adjunkt-Schema und Semantik

Um die endgültige Version des Head-Adjunkt-Schemas formulieren zu können, ist ein Vorgriff auf den Abschnitt zur Semantik nötig: Die Semantik von Adjunkten wird durch Einbettung der modifizierten semantischen Struktur in die semantische Struktur des Adjunktes modelliert. Die Bedeutungen des Adverbs 良く (*yoku*; gut) beispielsweise wird folgendermaßen dargestellt:¹⁰⁴

- (228) a. 良く
 yoku
 gut
- b. $\left[\begin{array}{l} well \\ \textcircled{1} cont \end{array} \right]$

Bei der Modifikation komplexer Verben durch Adverbien kann es zu mehreren Lesarten kommen:

- (229) 健が黙って直美を座らせた。¹⁰⁵

ken-ga damatte naomi-wo suwa.r.a-se.t-a
 Ken-NOM schweigend Naomi-ACC setzen-CAUS-PAST
 i. Ken brachte Naomi schweigend dazu, sich zu setzen.
 ii. Ken brachte Naomi dazu, sich schweigend zu setzen.

Im Falle von (229*b-i*) bezieht sich das Adverb auf den Aspekt des “Veranlassens”, im Falle von (229*b-ii*) hingegen auf den Akt des “sich Setzens”.

104 Die Semantik von Adjektiven müßte in Analogie zu der Semantik von Adverbien die folgende Struktur haben:

- (225) i. おいしいステーキ
 oisii sutēki
 gut schmecken Steak
 gut schmeckende Steaks
- ii. $\left[\begin{array}{l} delicious \\ \textcircled{1} steak \end{array} \right]$

Adjektive können jedoch auch zur Bildung von Nebensätzen wie

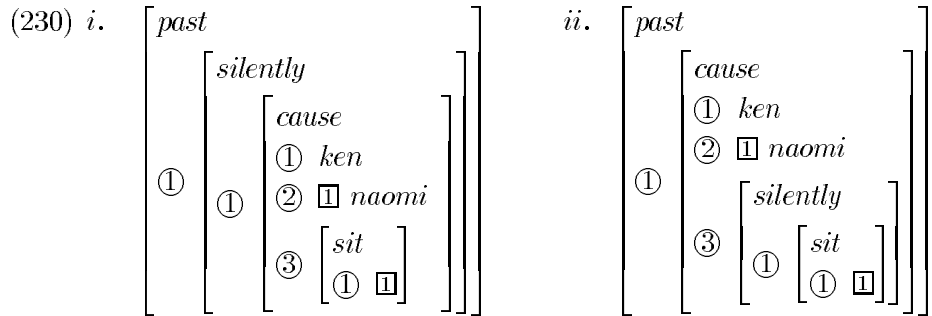
- (226) (ステーキが) おいしいレストラン (vgl. DBJG, S. 376, KS(B))
 (*sutēki-ga oisii resutoran*)
 (Steak-NOM) gut schmecken Restaurant
 Ein Restaurant, in dem (Steaks) / [das Essen] gut schmecken/t

benutzt werden (vgl. 3.3.3). Da beide Arten von Adjunkten auf morphosyntaktischer Ebene nicht unterschieden werden können, werden sie gleichbehandelt:

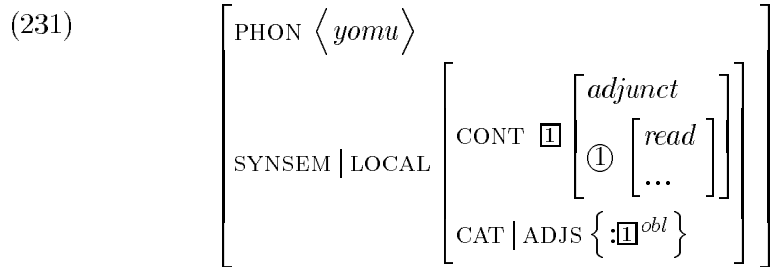
- (227) a. $\left[\begin{array}{l} adjunct \\ \textcircled{1} steak \\ \textcircled{2} \left[\begin{array}{l} delicious \\ \textcircled{1} cont \end{array} \right] \end{array} \right]$
- b. $\left[\begin{array}{l} adjunct \\ \textcircled{1} restaurant \\ \textcircled{2} \left[\begin{array}{l} delicious \\ \textcircled{1} cont / steak \end{array} \right] \end{array} \right]$

105 vgl. Manning et al. 1999, Beispiel (24)b.

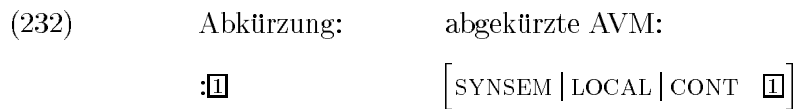
Wird das Perfekt in der Semantik durch $[past]$ und *Ken* und *Naomi* durch *ken* und *naomi* dargestellt, lassen sich diese beiden Lesarten folgendermaßen veranschaulichen:



Um diese semantische Struktur kompositional aus der Semantik von *suwaraseta* und *damatte* errechnen zu können, muß die Semantik des Adverbs an einer bestimmten Stelle in die Semantik des Verbes “eingebaut” werden. Dieses “Einbauen” wird in unserem Modell dadurch modelliert, daß die Bezugswörter mittels des Merkmals ADJS semantisch für ihre Adjunkte subkategorisieren:

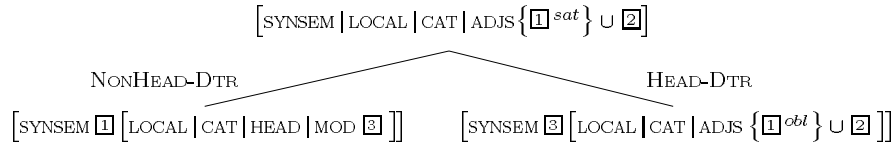


$: \boxed{1}$ beschreibt dabei innerhalb von Subkategorisierungsmengen die Semantik eines Zeichens:

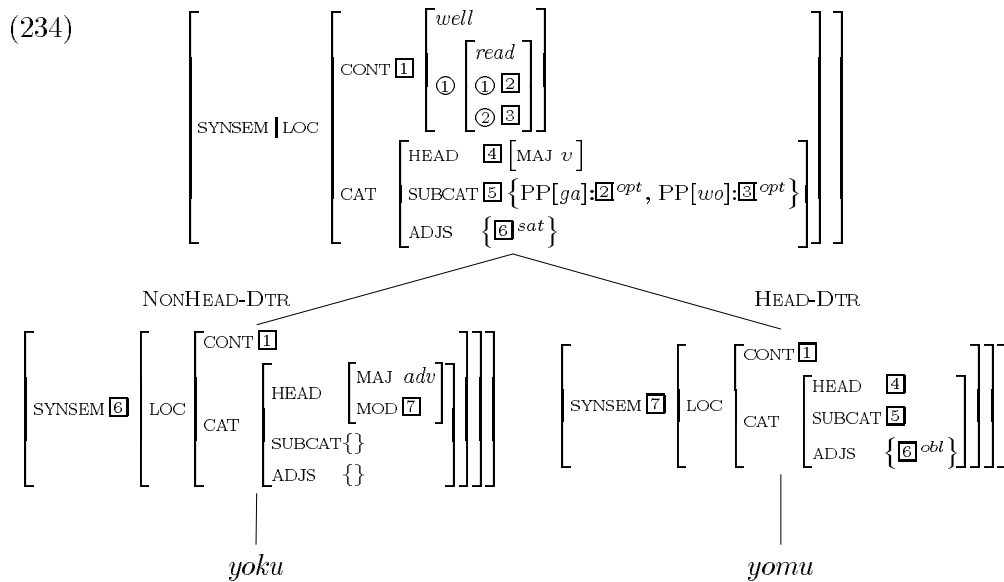


Durch ein entsprechend modifiziertes Head-Adjunkt-Schema läßt sich die Semantik des Adjunktes an entsprechender Stelle in die Semantik des Verbes einbauen:

(233) *Head-Adjunkt-Schema, endgültige Version:*¹⁰⁶



Für 良く読む (*yoku yomu*; gut lesen) ergibt sich beispielsweise die folgende Struktur:



Die beiden Lesarten von (229) können durch zwei entsprechende Lexikoneinträge für 座らせた (*suwaraseta*; zum Setzen veranlaßt haben) erklärt werden, die sich nur bezüglich der Stelle unterscheiden, die für die Semantik des Adjunktes vorgesehen ist:

¹⁰⁶ Strenggenommen müssen Adjunkte bezüglich ihrer obligatorischen Komplemente gesättigt sein. Diese Forderung würde jedoch eine genauere Modellierung von Relativsätzen erfordern als hier vorgenommen und wird deshalb fortgelassen. Vgl. auch Fußnote 122, S. 147.

- (235) a.
$$\left[\begin{array}{c} \text{PHON } \langle \text{suvaraseru} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \\ \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \\ \textcircled{3} \textcircled{3} \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \left[\begin{array}{c} \text{sit} \\ \textcircled{1} \textcircled{2} \end{array} \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } v \\ \text{SUBCAT} \{ \text{PP}[\text{ga}]:\textcircled{1}^{opt}, \text{PP}[\text{wo}]:\textcircled{2}^{opt} \} \\ \text{ADJS} \{ :\textcircled{3}^{obl} \} \end{array} \right] \end{array} \right] \end{array} \right]$$
- b.
$$\left[\begin{array}{c} \text{PHON } \langle \text{suvaraseru} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \\ \left[\begin{array}{c} \text{CONT} \textcircled{3} \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \\ \textcircled{3} \left[\begin{array}{c} \text{sit} \\ \textcircled{1} \textcircled{2} \end{array} \right] \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } v \\ \text{SUBCAT} \{ \text{PP}[\text{ga}]:\textcircled{1}^{opt}, \text{PP}[\text{wo}]:\textcircled{2}^{opt} \} \\ \text{ADJS} \{ :\textcircled{3}^{obl} \} \end{array} \right] \end{array} \right] \end{array} \right]$$

Lexikalische Regel zur Adjunkteinführung

Adjunkte sind fakultativ. Diese Fakultativität kann aber nicht — wie im Falle der Komplemente — durch den Index *opt* der entsprechenden Elemente der ADJS-Menge dargestellt werden: Die Semantik der Adjunkte wird mit einer “Leerstelle” in der Semantik des Bezugswortes unifiziert; Zeichen, die für Adjunkte subkategorisieren, haben also eine andere semantische Struktur als Zeichen, die nicht für Adjunkte subkategorisieren. Die Fakultativität der Adjunkte muß daher durch verschiedene Zeichenvarianten für das entsprechende Wort modelliert werden.

Nomen beispielsweise können durch ein oder mehrere Adjektive modifiziert werden. Für das Nomen 本 (*hon*; Buch) beispielsweise müssen daher die folgenden Lexikoneinträge existieren:¹⁰⁷

¹⁰⁷ Die Zeichen wurden auf die hier wesentlichen Merkmale vereinfacht.

- (236) a.
$$\left[\begin{array}{c} \text{PHON } \langle ho\tilde{n} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{c} \text{CONT } book \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } n \\ \text{ADJS } \{ \} \end{array} \right] \end{array} \right] \end{array} \right]$$
- b.
$$\left[\begin{array}{c} \text{PHON } \langle ho\tilde{n} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{c} \text{CONT } \boxed{1} \left[\begin{array}{c} adjunct \\ \textcircled{1} book \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } n \\ \text{ADJS } \{ : \boxed{1}^{obl} \} \end{array} \right] \end{array} \right] \end{array} \right]$$
- c.
$$\left[\begin{array}{c} \text{PHON } \langle ho\tilde{n} \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{c} \text{CONT } \boxed{1} \left[\begin{array}{c} adjunct \\ \textcircled{1} \boxed{2} \left[\begin{array}{c} adjunct \\ \textcircled{1} book \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \mid \text{MAJ } n \\ \text{ADJS } \{ : \boxed{1}^{obl}, : \boxed{2}^{obl} \} \end{array} \right] \end{array} \right] \end{array} \right]$$
- d. usw.

Zusammen mit dem Adjektiv 難しい (*muzukasii*; schwierig)

- (237)
$$\left[\begin{array}{c} \text{PHON } \langle muzukasii \rangle \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} adjunct \\ \textcircled{1} cont \\ \textcircled{2} \left[\begin{array}{c} difficult \\ \textcircled{1} cont \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ } adj \\ \text{MOD } NP \end{array} \right] \\ \text{SUBCAT } \{ \} \\ \text{ADJS } \{ \} \end{array} \right] \end{array} \right] \end{array} \right]$$

ergibt sich beispielsweise für (217a) die folgende Semantik:

- (238) a. 難しい本
 muzukasii hoñ
 schwer buch
 schwieriges Buch
- b.
$$\left[\begin{array}{l} adjunct \\ \textcircled{1} book \\ \textcircled{2} \left[\begin{array}{l} difficult \\ \textcircled{1} cont \end{array} \right] \end{array} \right]$$

Zeichen, die für Adjunkte subkategorisieren, unterscheiden sich nur bezüglich der Werte von CONT und ADJS von den korrespondierenden Standardeinträgen für Verben und Nomen. In den meisten Fällen können sie daher systematisch mit sogenannten *lexikalischen Regeln* aus einem einzigen Lexikoneintrag abgeleitet werden. Regel (239)¹⁰⁸ beispielsweise gestattet es, aus einem Lexikoneintrag für ein Nomen, der keine Adjunkte vorsieht, die entsprechende Variante für ein Adjunkt — oder, bei mehrfacher Anwendung der Regel, mehrere Adjunkte — abzuleiten. Dargestellt werden nur die Merkmale, die durch die lexikalische Regel modifiziert werden; alle anderen Merkmale werden unverändert übernommen:

(239) *Lexikalische Regel zur Adjunkteinführung*:¹⁰⁹

$$\left[\begin{array}{l} \text{S} | \text{L} \\ \text{CAT} \left[\begin{array}{l} \text{CONT} \textcircled{1} \\ \text{HEAD} | \text{MAJ } vcat \vee n \\ \text{ADJS} \textcircled{2} \end{array} \right] \end{array} \right] \Rightarrow \left[\begin{array}{l} \text{S} | \text{L} \\ \text{CAT} | \text{ADJS} \left[\begin{array}{l} \text{CONT} \textcircled{3} \left[\begin{array}{l} adjunct \\ \textcircled{1} \textcircled{1} \end{array} \right] \\ \textcircled{2} \cup \{ [] : \textcircled{3}^{obl} \} \end{array} \right] \end{array} \right]$$

Die Ableitung entsprechender Formen für komplexe Verben wie 座らせる (*suwaraseru*; jemanden dazu veranlassen, sich zu setzen), bei denen die Semantik des Adjunktes in die komplexe Semantik “eingebaut” wird, ist komplizierter und wird erst in Abschnitt 4.3 “Lexikalisch-morphologische Regeln” des nächsten Kapitels behandelt. Bis dahin können wir uns die entsprechenden Zeichen als alternative Lexikoneinträge vorstellen.

108 Der Ansatz, bei der Adjunktion im Japanischen den Head für seine Adverben subkategorisieren zu lassen, geht auf Iida et al. 1994, Manning et al. 1999 zurück. Er wird dort im Rahmen einer stark vereinfachten Morphologie vorgestellt, die ausschließlich auf lexikalischen Regeln basiert.

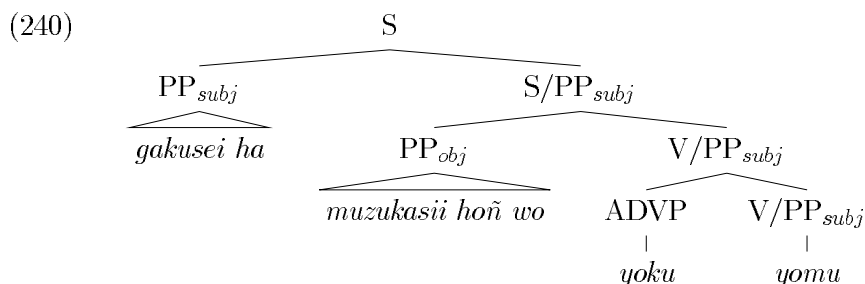
Iida et al. 1994 bezieht sich bei der Angabe seiner lexikalischen Regel (*adverbial type raising*) auf ähnliche Vorschläge in Abeillé und Godard 1994 (Französisch), Bouma und van Noord 1994 (Holländisch) und Kim und Sag 1995 (Englisch).

109 SYNSEM und LOCAL werden durch ihre Anfangsbuchstaben S und L abgekürzt.

Für lexikalische Regeln wird die Schreibweise aus Pollard und Sag 1994 verwendet. In unserem Constraint-System lassen sich lexikalische Regeln ähnlich wie syntaktische Schemata realisieren; anstatt zweier Töchter fordert eine lexikalische Regel nur eine Tochter. Die Merkmalwerte der Tochter werden — mit Ausnahme der Werte, die durch lexikalische Regeln modifiziert werden — von der Mutter übernommen.

3.2.3 Nichtlokale Abhängigkeiten

Mit dem Head-Komplement-Schema und dem Head-Adjunkt-Schema können wir unseren Beispielsatz (168), S. 102 nahezu vollständig ableiten. Nur die Phrase 学生は (*gakusei-ha*; “was den Studenten betrifft”) und ihr Zusammenhang mit den Kategorien, die im Syntaxbaum (169), S. 102 durch einen Slash markiert sind (S/PP_{subj} und V/PP_{subj}), kann noch nicht erklärt werden.



Die Phrase *gakusei ga* wird als *Topic* des Satzes bezeichnet. Das Topic steht am Satzanfang und wird durch die Postposition は (*ha*) — den sogenannten *Topicpartikel* — gekennzeichnet. Das Topic hat entweder einen thematischen Zusammenhang mit dem Satz und läßt sich syntaktisch wie eine Adverb behandeln

(241) 象は鼻が長い。 ¹¹⁰

zou-ha hana-ga nagai
 Elephant-TOP Nase-NOM lang
 Elephanten haben lange Rüssel.
 (Was die Elephanten betrifft: Die Rüssel sind lang.)

oder es entspricht einem Komplement des Verbes, das aus dem Verbalkomplex heraus an den Anfang des Satzes verschoben wurde. Letzteres ist auch in unserem Beispielsatz (168) der Fall, der hier als (242) noch einmal dargestellt wird:

(242) 学生は難しい本を良く読むね。

gakusei-ha muzukasii hoñ-wo yoku yomu ne
 Student-TOP schwer buch-ACC gut lesen nicht wahr
 Der Student liest das schwierige Buch gut, nicht wahr?

¹¹⁰ Vgl. den gleichnamigen Aufsatz von Mikami Akira (Mikami 1960) und DBJG, -wa -ga -は -が, example (d), S. 525.

Der erste Fall wird durch ein Zeichen für は (*ha*) modelliert, das mit [MOD *Vcat*] markiert ist und damit als Adjunkt fungieren kann; der zweite Fall, in dem は (*ha*) eine topikalisierte Komplement-PP bezeichnet, wird in Analogie zu der Behandlung nichtlokaler Konstituenten in der HPSG mit Hilfe des Merkmals SLASH modelliert. In beiden Fällen ersetzt das Topicpartikel は (*ha*) entweder die ursprüngliche Postposition oder steht nach dieser:

(243) *Der Topicpartikel は (*ha*) und seine Verwendung mit Postpositionen:*

a. は (*ha*) ersetzt die Postposition:

が	→	は
<i>ga</i>		<i>ha</i>
を	→	は
<i>wo</i>		<i>ha</i>

b. は (*ha*) ersetzt die Postposition oder steht nach dieser:

に	→	(に) は
<i>ni</i>		<i>ni ha</i>
へ	→	(へ) は
<i>he</i>		<i>he ha</i>

c. は (*ha*) steht nach der Postposition:

に	→	に は
<i>ni</i>		<i>ni ha</i>
で	→	で は
<i>de</i>		<i>de ha</i>
と	→	と は
<i>to</i>		<i>to ha</i>
から	→	から は
<i>kara</i>		<i>kara ha</i>
まで	→	まで は
<i>made</i>		<i>made ha</i>
より	→	より は
<i>yor</i>		<i>yor ha</i>

Topikalisierung

Die Topikalisierung ist ein syntaktischer Prozeß und interagiert — im Gegensatz zu Head-Komplement- und Head-Adjunkt-Schema — nicht mit morphologischen Prozessen. Aus diesem Grunde sollen die Mechanismen, die ihr zugrunde liegen, nur kurz skizziert werden.

Topikalisierung kommt durch Interaktion dreier Mechanismen zustande:

1. *Lexikalische Regel zur Extraktion von Komplementen:*

Ein Komplement wird durch eine lexikalische Regel (oder ein spezielles ϕ -Komplement) aus der SUBCAT-Menge extrahiert und der Wertemenge des Merkmals SLASH hinzugefügt.

2. *SLASH-Merkmal-Prinzip:*

Die Werte der SLASH-Merkmale werden durch ein eigenes Prinzip, das *SLASH-Merkmal-Prinzip*, im Syntaxbaum nach oben transportiert.

3. *Head-Füller-Schema:*

Nachdem die Komplemente und Adjunkte des Verbes realisiert wurden, werden die SLASH-Elemente durch ein eigenes Dominanzschema, das *Head-Füller-Schema*, realisiert.

Das SLASH-Merkmal ist das einzige nicht-lokale Merkmal, das in unserer Syntax benutzt wird — und damit das einzige Merkmal, das nicht unter LOCAL steht, sondern unter NONLOCAL:

$$(244) \quad \left[\begin{array}{l} \text{nonlocal} \\ \text{SYNSEM} \mid \text{NONLOCAL} \mid \text{SLASH}^{111} \text{set}(\text{local}) \end{array} \right]$$

Bei lexikalischen Elementen wird es durch die leere Menge spezifiziert. Nur durch die Anwendung der lexikalischen Regel zur Extraktion von Komplementen kann es Elemente aufnehmen:

(246) *Komplement-Extraktionsregel:*¹¹²

$$\left[\begin{array}{l} \text{S} \left[\begin{array}{l} \text{LOCAL} \mid \text{SUBCAT} \{ \text{1}^{unsat} \} \cup \text{2} \\ \text{NONLOCAL} \mid \text{SLASH} \text{3} \end{array} \right] \end{array} \right] \Rightarrow \left[\begin{array}{l} \text{S} \left[\begin{array}{l} \text{LOCAL} \mid \text{SUBCAT} \{ \text{1} \mid \text{LOCAL} \text{4}^{sat} \} \cup \text{2} \\ \text{NONLOCAL} \mid \text{SLASH} \{ \text{4}^{obl} \} \cup \text{3} \end{array} \right] \end{array} \right]$$

Die Komplement-Extraktions-Regel markiert ein Element der SUBCAT-Menge als gesättigt (*sat*) und übernimmt dessen LOCAL-Wert als neues Element in die SLASH-Menge. Im Falle unseres Beispielsatzes, in dem das Subjekt topikalisiert wird, hat die Anwendung der Komplement-Extraktions-Regel folgende Auswirkungen:

(247) *Extraktion des Subjektes von 読む (yomu; lesen):*

$$\left[\begin{array}{l} \text{PHON} \langle \text{yomu} \rangle \\ \text{S} \left[\begin{array}{l} \text{LOCAL} \mid \text{SUBCAT} \left\{ \text{PP}[\text{ga}]^{opt} \right\} \\ \text{PP}[\text{wo}]^{opt} \end{array} \right] \\ \text{NONLOCAL} \mid \text{SLASH} \{ \} \end{array} \right] \Rightarrow \left[\begin{array}{l} \text{PHON} \langle \text{yomu} \rangle \\ \text{S} \left[\begin{array}{l} \text{LOCAL} \mid \text{SUBCAT} \left\{ \text{PP}[\text{ga}, \text{LOCAL} \text{1}^{sat}] \right\} \\ \text{PP}[\text{wo}]^{opt} \end{array} \right] \\ \text{NONLOCAL} \mid \text{SLASH} \{ \text{1}^{obl} \} \end{array} \right]$$

¹¹¹ In Pollard und Sag 1994 findet sich das Merkmal SLASH unter folgendem Pfad:

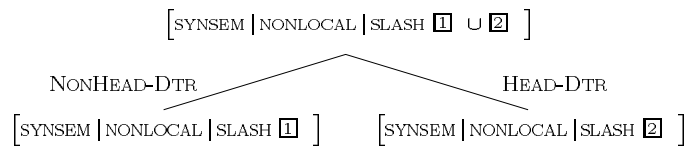
$$(245) \quad \left[\text{SYNSEM} \mid \text{NONLOCAL} \mid \text{INHERITED} \mid \text{SLASH} \text{set}(\text{local}) \right]$$

Hier wird jedoch von einer vereinfachten Version der Behandlung nicht-lokaler Phänomene ausgegangen, so daß auf das Merkmal INHERITED verzichtet werden kann.

¹¹² Vgl. Pollard und Sag 1994: Chapter 9.5. Alle durch die lexikalische Regel nicht modifizierten Merkmale werden unverändert übernommen.

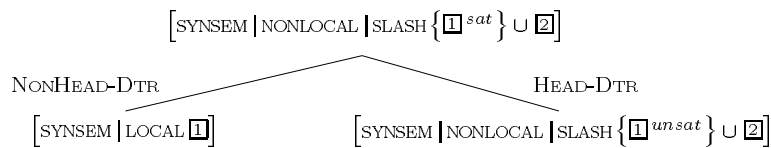
In lokalen Bäumen errechnet sich der SLASH-Wert der Mutter aus der Vereinigung der SLASH-Werte der Töchter. Der Constraint, der hierfür verantwortlich ist, wird als SLASH-Merkmal-Constraint bezeichnet und läßt sich für das Head-Komplement-Schema und das Head-Adjunkt-Schema folgendermaßen darstellen:¹¹³

(248) *Das SLASH-Merkmal-Prinzip für HKS und HAS:*



Nachdem die obligatorischen Elemente der SUBCAT- und ADJS-Menge gesättigt wurden, können die Elemente der SLASH-Menge mit dem Head-Füller-Schema realisiert werden:¹¹⁴

(249) *Das HEAD-Füller-Schema (vereinfacht):*



Im Falle unseres Beispielsatzes würde das Zusammenspiel von Extraktion, Transport und Realisation des topikalisierten Elementes vereinfacht folgendermaßen aussehen:

¹¹³ In Pollard und Sag 1994 wurden alle Prinzipien so formuliert, daß sie für *alle* Dominanzschemata gelten. Da die Prinzipien dadurch schwerer zu lesen sind, wird hier eine vereinfachte Version benutzt, die nur für Head-Komplement-Schema und Head-Adjunkt-Schema gültig ist. Das Head-Füller-Schema hingegen wird so beschrieben, daß es die bei seiner Anwendung gültige Version des SLASH-Merkmal-Prinzips miteinschließt.

¹¹⁴ Die Anwendung des Head-Komplement-Schemas nach Anwendung des Head-Füller-Schemas zur Realisierung noch nicht gesättigter, optionaler SUBCAT-Elemente kann durch das Ersetzen der Indizes *opt* durch einen anderen Wert unterbunden werden. Dieser Mechanismus zur Verhinderung weiterer Anwendungen des Head-Komplement-Schemas wurde — genau wie das Überprüfen des Sättigungsgrades der Head-Tochter — in (249) weggelassen.

Das Steak, das Herr Tanaka gegessen hat, war teuer.

Die folgenden Eigenschaften charakterisieren den Relativsatz im Japanischen¹¹⁷:

1. Das Verb steht in einer Attributivform.¹¹⁸
2. Das Bezugswort des Relativsatzes repräsentiert ein Komplement des Nebensatzes. Dieses kann folglich nicht ein zweites Mal im Nebensatz realisiert werden. Die Postposition, die normalerweise die Beziehung zwischen Objekt und Verb markiert, fällt weg.
3. Relativsätze enthalten keine topikalisierten PPs; das Partikel は (*ha*) kommt in ihnen also nicht vor. Das Subjekt des Relativsatzes wird entweder durch が (*ga*) oder durch の (*no*) realisiert.

Attributsätze

Von den Relativsätzen auf morphosyntaktischer Ebene in den meisten Fällen nicht unterscheidbar sind die *Attributsätze*: Der einzige Unterschied zwischen Relativsätzen und Attributsätzen besteht darin, daß das Bezugswort von Attributsätzen nicht Komplement des Nebensatzes ist.

- (253) a. i. [ステーキが] おいしいレストラン¹¹⁹
 [sutēki-ga] oisii resutoran
 Steak-NOM gut schmecken Restaurant
 Ein Restaurant, in dem Steaks gut schmecken
- ii. おいしいステーキ
 oisii sutēki
 gut schmecken Steak
 gut schmeckende Steaks
- b. i. [ジョンが] [ステーキを] 食べたレストラン
 [jon-ga] [sutēki-wo] tabeta resutoran
 John-NOM Steak-ACC aß Restaurant
 Das Restaurant, in dem John Steak aß
- ii. [ジョンが] 食べたステーキ¹²⁰
 [jon-ga] tabeta sutēki
 John-NOM aß Steak
 Das Steak, das John aß

117 vgl. DBJG, S. 377, Notes 1.

118 Im modernen Japanisch fällt die Attributivform mit der satzfinalen Form zusammen. Das klassische Japanisch hingegen unterscheidet 終止形 (*syuusikei*; Finalform) und 連体形 (*reñtaikei*; Attributivform).

119 DBJG, S. 376, KS(B).

120 DBJG, S. 378, Notes (6).

Die unter *i.* aufgeführten Nebensätze sind Attributsätze; die unter *ii.* aufgeführten hingegen Relativsätze. Der in [eckige Klammern] gesetzte Teil der Nebensätze kann — wie die meisten Komplemente japanischer Verben — weggelassen werden. Syntaktisch sind Relativ- und Attributsätze dann nicht mehr zu unterscheiden. Nur durch die Berücksichtigung von Semantik und Weltwissen kann erschlossen werden, daß ein Restaurant an sich weder lecker ist, noch gegessen werden kann, daß es sich bei diesem Wort folglich nicht um ein Komplement des Nebensatzverbes handeln kann — und daß die Nebensätze daher nicht Relativsätze sondern Attributsätze sein müssen.

In Fällen, die auf morphosyntaktischer Ebene nicht aufgelöst werden können, bieten sich zwei mögliche Strategien an:

1. *Alternative Analysen:*

Alle möglichen Analysen werden ermittelt und als alternative Lösungen vorgeschlagen;

2. *Unterspezifikation:*¹²¹

Es wird eine Form der Darstellung gewählt, die ebenso mehrdeutig ist wie die analysierte Form selber.

In beiden Fällen läßt sich durch Integration von Heuristiken oder Weltwissen die Mehrdeutigkeit an einer späteren Stelle eventuell auflösen. In unserer Syntax wählen wir den zweiten Weg: Relativsätze und Attributsätze werden semantisch beide auf die gleiche Weise als (evt. komplexe) Adjunkte behandelt.

Relativsätze als Adjunkte

Im Abschnitt 3.2.2 “Adjunkte” wurde beschrieben, daß für Adjunkte ein einziges Merkmal notwendig und hinreichend ist: das Merkmal MOD. Mittels dieses Merkmals subkategorisieren Adjunkte für ihre Bezugswörter und legitimieren die Anwendung des Head-Adjunkt-Schemas. Adnominale Verben können aus satzfinalen Verben also einfach durch das Ersetzen von [MOD *none*] durch [MOD N] abgeleitet werden.

¹²¹ Für semantische Formalismen, die eine elegantere Form der Unterspezifikation erlauben, vgl. Pinkal 1995, Pinkal 1999 und Copestake et al. 1997. In diesen Aufsätzen wird unter anderem gezeigt, wie sich semantische Wertebereiche besser eingrenzen lassen und wie auch Skopusmehrdeutigkeiten in einer Struktur dargestellt werden können.

Der im Verbmobilprojekt verwendete semantische Formalismus wird in Bos et al. 1994 beschrieben.

Verben und Adjektive

In Kapitel 2 wurde auf die Ähnlichkeiten von Verben und Adjektiven in ihrem syntaktischen Verhalten hingewiesen. Diese Ähnlichkeiten beider Kategorien und die Möglichkeit der Überführung von Verben in Adjektive und umgekehrt legt ihre Gleichbehandlung nahe: Beide werden im Lexikon daher als Worte beschrieben, die für Objekte subkategorisieren und als Hauptverb mit diesen zusammen einen Satz bilden können. Die Adnominal- bzw. Relativsatzform von Adjektiven und Verben wird daher bei beiden Kategorien ebenfalls auf die gleiche Weise durch die folgende lexikalische Regel abgeleitet:

- (254) *Lexikalische Regel zur Ableitung der
Adnominalform von Verben und Adjektiven:*

$$\left[\begin{array}{c} \text{S} | \text{L} | \text{CAT} | \text{HEAD} \left[\begin{array}{c} \text{MAJ } vcat \\ \text{MOD } none \end{array} \right] \end{array} \right] \Rightarrow \left[\text{S} | \text{L} | \text{CAT} | \text{HEAD} | \text{MOD } N \right]$$

Eine zweite Modifikation der verbalen Zeichen ist notwendig: Die Semantik der Verben muß so modifiziert werden, daß sie zur Modifikation der Semantik von Nomina benutzt werden kann. Auf diese Modifikation werden wir im Kapitel 3.3 “Semantik” näher eingehen.

Regel (254) könnte durch andere Modifikationen ergänzt werden: Bei der Ableitung adnominaler Verben zur Bildung von Relativsätzen beispielsweise kann ein Komplement aus der SUBCAT-Menge als gesättigt gekennzeichnet und mit dem Wert von MOD unifiziert werden. Dadurch würde das Bezugswort semantisch mit dem Komplement identifiziert werden. Wie jedoch gesagt wurde, werden in der hier vorgestellten Semantik Relativsätze und Attributsätze gleichbehandelt, und die Beziehung vom Bezugswort zum Nebensatz offengelassen.¹²²

Relativsätze und Attributsätze mit oder ohne vorangestellte Komplemente werden also durch Zeichen beschrieben, die sich syntaktisch und semantisch genauso verhalten, wie die in Abschnitt 3.2.2, S. 130, beschriebenen Adjunkte. Als Alternative zur lexikalischen Regel (254) wird im Kapitel zur Morphologie die Behandlung durch spezielle Morpheme (Adnominal-Flexive) vorgeschlagen. Beide Modellierungen sind jedoch letztlich gleichwertig.

3.3 Semantik

Thema dieser Arbeit ist die Morphologie von Verben und ihre Interaktion mit der Semantik. Nicht Probleme der Semantik an sich sollen deshalb in die-

¹²² Da bei Relativsätzen das relativierte Komplement nicht als gesättigt gekennzeichnet wird, kann im Head-Adjunkt-Schema nicht gefordert werden, daß Relativsätze hinsichtlich ihrer obligatorischen Komplemente gesättigt sind. Vgl. Fußnote 106, S. 136.

sem Abschnitt beschrieben werden, sondern die semantischen Aspekte dieser *Interaktion*. Von *Interaktion* zwischen Morphologie und Semantik wird dann gesprochen, wenn morphologische Modifikationen von semantischen Modifikationen begleitet werden.

In formalen Theorien der Syntax und Semantik wird meistens von zwei Grundannahmen ausgegangen:

1. Den Grundelementen der Theorie, d.h. den Einträgen des morphologischen Lexikons, sind formale Repräsentationen ihrer Semantik zugeordnet;
2. Die Verknüpfung syntaktischer Grundelemente wird von der Verknüpfung ihrer semantischen Repräsentationen begleitet.

Punkt 2 wird seit Frege 1879 als *Kompositionsprinzip der Bedeutung* bezeichnet. Die Parallelität der *syntaktischen* und *semantischen Komposition* ist gemeint, wenn in dieser Arbeit von *Interaktion* zwischen Semantik und Morphologie gesprochen wird.

Wie alle anderen Komponenten unifiktionsbasierter Theorien werden auch die semantischen Elemente durch getypte Merkmalstrukturen dargestellt. Die semantische Komposition wird entsprechend durch Unifikation modelliert.

Da uns die Interaktion von Semantik und Morphologie interessiert und nicht die Semantik an sich, empfiehlt es sich, von der inneren Struktur der semantischen Merkmalstrukturen so weit es möglich ist zu abstrahieren und damit den Schwerpunkt auf die Komposition zu legen. Durch diese Abstraktion bleibt die Semantik mit ihren Wechselwirkungen mit der Morphologie übersichtlich und allgemein und kann im Bedarfsfall leicht durch eine andere semantische Theorie ersetzt werden.¹²³

Die Forderung nach Abstraktion wird durch folgende Prinzipien erfüllt:

1. Morphologischen bzw. syntaktischen Grundelementen wird möglichst eindeutig ein semantisches Grundelement zugeordnet.
2. Die morphologische bzw. syntaktische Valenz wird soweit möglich eindeutig in die semantische Valenz übersetzt. Es wird nicht zwischen Konstanten, semantischen Funktionen, Prädikaten, Quantoren, logischen Verknüpfungen etc. unterschieden.

¹²³ Innerhalb der HPSG wird ein situationssemantischer Ansatz verfolgt. Vgl. beispielsweise Barwise und Perry 1983, Cooper et al. 1990, Barwise et al. 1991 und Aczel et al. 1993 oder die entsprechenden Kapitel in Pollard und Sag 1987 und Pollard und Sag 1994.

3. Um die Beziehung zwischen Semantik und morphologischen bzw. syntaktischen Elementen hervorzuheben und dennoch auf den qualitativen Unterschied zwischen beiden Beschreibungen hinzuweisen, werden die morphologischen bzw. syntaktischen Grundelemente semantisch durch eine Übersetzung oder Umschreibung ins Englische dargestellt. Die semantischen Argumente werden nicht durch Merkmale oder semantische Rollen interpretiert, sondern einfach durchnummeriert:

(255) *syntaktische Grundelemente und ihre semantischen Entsprechungen:*

PHON	CONT	Übersetzung
<i>gakusei</i>	<i>student</i>	Student
<i>muzukasii</i>	$\left[\begin{array}{l} \textit{difficult} \\ \textcircled{1} \textit{ cont} \end{array} \right]$	schwer
<i>yomu</i>	$\left[\begin{array}{l} \textit{read} \\ \textcircled{1} \textit{ cont} \\ \textcircled{2} \textit{ cont} \end{array} \right]$	lesen
...		

Diese Prinzipien führen zu einer naiven Semantik, die einer isomorphen, “notationellen Variante” der morphologischen Struktur nahekommt. Gerade durch diese Einfachheit entspricht sie jedoch den hier gestellten Anforderungen.

Im Folgenden werden anhand von Beispielen die semantischen Strukturen und Prinzipien vorgestellt.

3.3.1 Lexikalische Einträge

Die Semantik lexikalischer Einträge wird im Lexikon spezifiziert. Wie in den folgenden Abschnitten gezeigt, ergibt sich für jede morphologische oder syntaktische Kategorie eine typische Form der semantischen Repräsentation:

Nomen

Nomen werden durch minimale semantische Typen dargestellt:

(256) <i>Nomen:</i>	<i>Übersetzung:</i>	<i>CONT:</i>
学生 <i>gakusei</i>	<i>Student:</i>	$\left[\dots \text{CONT } \textit{student} \right]$
電話 <i>denwa</i>	<i>Telephon:</i>	$\left[\dots \text{CONT } \textit{telephone} \right]$
本 <i>hon</i>	<i>Buch:</i>	$\left[\dots \text{CONT } \textit{book} \right]$

Verben

Verben werden durch semantische Strukturen modelliert, deren Typ die Bedeutung angibt und deren Argumente die Valenz spiegeln:

(257) *Nomen:* *Übersetzung:* *CONT:*

寝る <i>neru</i>	<i>schlafen:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{sleep} \\ \textcircled{1} \textcircled{1} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\textit{ga}]:\textcircled{1}^{opt} \right\} \end{array} \right]$
食べる <i>taberu</i>	<i>essen:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{eat} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \begin{array}{l} \text{PP}[\textit{ga}]:\textcircled{1}^{opt} \\ \text{PP}[\textit{wo}]:\textcircled{2}^{opt} \end{array} \right\} \end{array} \right]$
書く <i>kaku</i>	<i>schreiben:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{write} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \\ \textcircled{3} \textcircled{3} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \begin{array}{l} \text{PP}[\textit{ga}]:\textcircled{1}^{opt} \\ \text{PP}[\textit{ni}]:\textcircled{2}^{opt} \\ \text{PP}[\textit{wo}]:\textcircled{3}^{opt} \end{array} \right\} \end{array} \right]$

Wie (257) zeigt, wird die Verbindung zwischen Syntax und Semantik — d.h. zwischen den syntaktischen Komplementen und den semantischen Argumenten — durch Unifikation der entsprechenden Argumente der Verbsemantik mit den CONT-Werten der Komplemente schon im Lexikoneintrag hergestellt.

Im Abschnitt 3.2.2 “*Adjunkte*” wurde auf Verben mit komplexer Semantik hingewiesen:

(258)	$\left[\begin{array}{l} \text{PHON} \left\langle \textit{suwaraseru} \right\rangle \\ \text{SYNSEM} \text{LOCAL} \left[\begin{array}{l} \text{CONT} \left[\begin{array}{l} \textit{cause} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \\ \textcircled{3} \left[\begin{array}{l} \textit{sit} \\ \textcircled{1} \textcircled{2} \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{l} \text{HEAD} \text{MAJ } v \\ \text{SUBCAT} \left\{ \text{PP}[\textit{ga}]:\textcircled{1}^{opt}, \text{PP}[\textit{ni} \vee \textit{wo}]:\textcircled{2}^{opt} \right\} \\ \text{ADJS} \quad \{ \} \end{array} \right] \end{array} \right] \end{array} \right]$
-------	---

Im Kapitel 4 “*Struktur der morphologischen Beschreibung*” werden wir sehen, daß sich die Semantik solcher Verben kompositional aus der Semantik der Morpheme ergibt, aus denen sich diese Verben zusammensetzen.

Adjektive

Auf die Ähnlichkeit zwischen Adjektiven und Verben und die Möglichkeit, beide als Hauptverb oder adnominal zu verwenden, wurde schon hingewiesen. Adjektive und Verben werden auch semantisch ähnlich behandelt:

(259) *Nomen:* *Übersetzung:* *CONT:*

難しい <i>muzukasii</i>	<i>schwierig:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{difficult} \\ \textcircled{1} \text{ } \square \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\textit{ga}]:\square^{\textit{opt}} \right\} \end{array} \right]$
高い <i>takai</i>	<i>hoch:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{high} \\ \textcircled{1} \text{ } \square \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\textit{ga}]:\square^{\textit{opt}} \right\} \end{array} \right]$
面白い <i>omoshiroi</i>	<i>interessant:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{interesting} \\ \textcircled{1} \text{ } \square \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\textit{ga}]:\square^{\textit{opt}} \right\} \end{array} \right]$

Im Gegensatz zu Verben subkategorisieren Adjektive jedoch in den meisten Fällen nur für ein Komplement.

Adverben

Die semantische Modellierung von Adverben wurde schon im Abschnitt 3.2.2 “*Adjunkte*” beschrieben: Sie werden als Strukturen mit einem Argument modelliert, das durch das Head-Adjunkt-Schema in die Semantik ihrer Bezugswörter eingebaut wird.

(260) *Nomen:* *Übersetzung:* *CONT:*

良く <i>yoku</i>	<i>gut:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{well} \\ \textcircled{1} \textit{cont} \end{array} \right] \\ \dots \text{SUBCAT} \{ \} \end{array} \right]$
黙って <i>damatte</i>	<i>schweigend:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{silently} \\ \textcircled{1} \textit{cont} \end{array} \right] \\ \dots \text{SUBCAT} \{ \} \end{array} \right]$
毎日 <i>mainiti</i>	<i>täglich:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \textit{every-day} \\ \textcircled{1} \textit{cont} \end{array} \right] \\ \dots \text{SUBCAT} \{ \} \end{array} \right]$

Postpositionen

Postpositionen haben die Aufgabe, die syntaktisch-semantische Relation zwischen dem Head und seinen Komplementen herzustellen. Sie haben daher meist keine wirkliche lexikalische Bedeutung. In unserem Modell übernehmen sie daher in den meisten Fällen einfach die Bedeutung der Nominalphrasen, für die sie subkategorisieren. In den wenigen Fällen, in denen sie in unserem Modell Bedeutung tragen (も (*mo*; auch) und は (*ha*; <Topic>), werden sie semantisch ähnlich wie Adjektive modelliert:

(261) *Nomen:* *Übersetzung:* *CONT:*

が <i>ga</i>	<Subjekt>:	$\begin{bmatrix} \dots \text{CONT} & \mathbb{I} \\ \dots \text{SUBCAT} & \{ \text{N:} \mathbb{I}^{opt} \} \end{bmatrix}$
は <i>ha</i>	<Topic>:	$\begin{bmatrix} \dots \text{CONT} & \begin{bmatrix} \text{topic} \\ \textcircled{1} \mathbb{I} \end{bmatrix} \\ \dots \text{SUBCAT} & \{ \text{N:} \mathbb{I}^{opt} \} \end{bmatrix}$
も <i>mo</i>	<i>auch:</i>	$\begin{bmatrix} \dots \text{CONT} & \begin{bmatrix} \text{also} \\ \textcircled{1} \mathbb{I} \end{bmatrix} \\ \dots \text{SUBCAT} & \{ \text{N:} \mathbb{I}^{opt} \} \end{bmatrix}$

Interjektionspartikel

Interjektionspartikel subkategorisieren für Sätze und modifizieren deren Bedeutung. Semantisch verhalten sie sich daher genauso wie Adjektive — und werden entsprechend modelliert:

(262) *Nomen:* *Übersetzung:* *CONT:*

か <i>ka</i>	<Frage>:	$\begin{bmatrix} \dots \text{CONT} & \begin{bmatrix} \text{question} \\ \textcircled{1} \mathbb{I} \end{bmatrix} \\ \dots \text{SUBCAT} & \{ \text{S:} \mathbb{I}^{opt} \} \end{bmatrix}$
ね <i>ne</i>	<i>nicht wahr?:</i>	$\begin{bmatrix} \dots \text{CONT} & \begin{bmatrix} \text{isn-t-it} \\ \textcircled{1} \mathbb{I} \end{bmatrix} \\ \dots \text{SUBCAT} & \{ \text{S:} \mathbb{I}^{opt} \} \end{bmatrix}$
な <i>na</i>	<Verbot>:	$\begin{bmatrix} \dots \text{CONT} & \begin{bmatrix} \text{prohibition} \\ \textcircled{1} \mathbb{I} \end{bmatrix} \\ \dots \text{SUBCAT} & \{ \text{S:} \mathbb{I}^{opt} \} \end{bmatrix}$

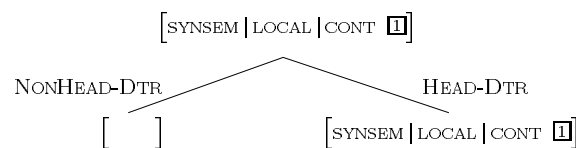
3.3.2 Dominanzschemata

Nachdem die semantischen Grundelemente, wie sie sich bei den lexikalischen Einträgen finden, vorgestellt wurden, sollen nun die Mechanismen der semantischen Komposition dargestellt werden:

Das CONTENT-Merkmal-Prinzip

Allen Dominanzschemata gemeinsam ist die Identität der Semantik von Mutter und Head-Tochter. Diese wird durch das *CONTENT-Merkmal-Prinzip* — ein Constraint phrasaler Zeichen — hergestellt:¹²⁴

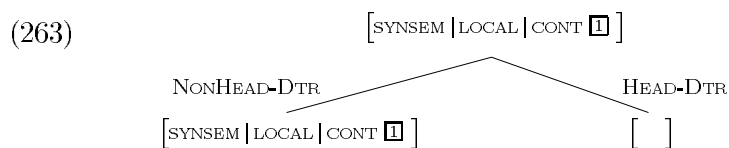
(264) *Das CONTENT-Merkmal-Prinzip:*



Head-Komplement-Schema

Das Head-Komplement-Schema wird durch die SUBCAT-Menge der Head-Tochter gesteuert. Wie im letzten Abschnitt beschrieben, erfolgt auch die Verknüpfung zwischen der Semantik der HEAD-Tochter und der Semantik ihrer Komplemente im Lexikon:

¹²⁴ In der Standardversion der HPSG wird die Semantik beim Head-Adjunkt-Schema vom Adjunkt übernommen:



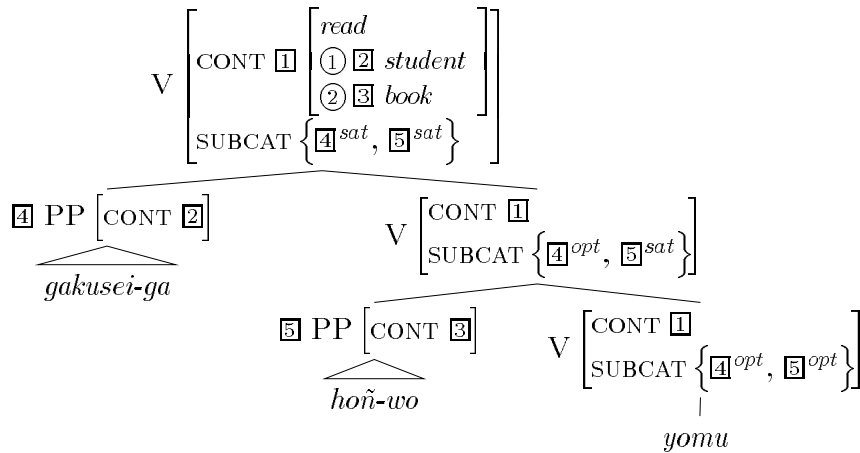
Dieser Mechanismus ist einfacher, weil auf lexikalische Regeln zur Adjunkteinführung verzichtet werden kann. Er läßt sich jedoch nur verwenden, wenn davon ausgegangen wird, daß die Semantik der Bezugsworte in die Semantik des Adjunktes eingebunden wird. Für Verben mit komplexer Bedeutung, wie hier vorgestellt, ist dieser Ansatz daher nicht verwendbar.

(265) *Nomen: Übersetzung: CONT:*

寝る <i>neru</i>	<i>schlafen:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \text{sleep} \\ \textcircled{1} \textcircled{1} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\text{ga}]:\textcircled{1}^{\text{opt}} \right\} \end{array} \right]$
難しい <i>muzukasii</i>	<i>schwierig:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \text{difficult} \\ \textcircled{1} \textcircled{1} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{PP}[\text{ga}]:\textcircled{1}^{\text{opt}} \right\} \end{array} \right]$
か <i>ka</i>	<i><Frage>:</i>	$\left[\begin{array}{l} \dots \text{CONT} \left[\begin{array}{l} \text{question} \\ \textcircled{1} \textcircled{1} \end{array} \right] \\ \dots \text{SUBCAT} \left\{ \text{S}:\textcircled{1}^{\text{opt}} \right\} \end{array} \right]$

Das Head-Komplement-Schema unifiziert die Elemente der SUBCAT-Menge mit den SYNSEM-Werten der Zeichen, durch die diese realisiert werden. Durch diese Unifikation kommt gleichzeitig die Komposition der Semantik zustande. Durch das CONTENT-Merkmal-Prinzip schließlich wird die semantische Struktur der Head-Tochter zur Mutter transportiert:

(266) *Semantik und Head-Komplement-Schema (schematische Darstellung):*



Head-Adjunkt-Schema

Der Mechanismus, der für die Komposition der Semantik bei Adjunkten verantwortlich ist, wurde schon im Abschnitt 3.2.2 “Adjunkte”, S. 130 vorgestellt: Durch eine lexikalische Regel wird in die Semantik des Bezugswortes eine “Leerstelle” für die Semantik des Adjunktes eingeführt. Gleichzeitig wird diese an das CONT-Merkmal des entsprechenden, ebenfalls durch die lexikalische Regel neu eingeführten ADJS-Elementes gebunden. Durch das Head-Adjunkt-Schema wird das ADJS-Element mit dem Zeichen unifiziert,

das das Adjunkt repräsentiert. Das CONTENT-Merkmal-Prinzip schließlich transportiert die Semantik von der Head Tochter (dem Bezugswort) zur Mutter:

- (267) a. *Anwendung der lexikalischen Regel zur Adjunkteinführung auf den lexikalischen Eintrag von 本 (hoñ; Buch) (schematische Darstellung):*

$$\left[\begin{array}{c} \text{CONT } book \\ \text{ADJS } \{ \} \end{array} \right] \Rightarrow \left[\begin{array}{c} \text{CONT } \boxed{1} \left[\begin{array}{c} adjunct \\ \textcircled{1} book \end{array} \right] \\ \text{ADJS } \{ : \boxed{1}^{obl} \} \end{array} \right]$$

- b. CONT^{125} und MOD-Wert des lexikalischen Eintrags des Adjektivs 難しい (muzukasii; schwierig) (schematische Darstellung):

$$\left[\begin{array}{c} \text{CONT } \left[\begin{array}{c} adjunct \\ \textcircled{1} cont \\ \textcircled{2} \left[\begin{array}{c} difficult \\ \textcircled{1} cont \end{array} \right] \end{array} \right] \\ \text{MOD } NP \end{array} \right]$$

- c. *Semantik und Head-Adjunkt-Schema (schematische Darstellung):*

$$\begin{array}{c} NP \left[\begin{array}{c} \text{CONT } \boxed{1} \left[\begin{array}{c} adjunct \\ \textcircled{1} book \\ \textcircled{2} \left[\begin{array}{c} difficult \\ \textcircled{1} cont \end{array} \right] \end{array} \right] \\ \text{ADJS } \{ \boxed{2}^{sat} \} \end{array} \right] \\ \swarrow \quad \searrow \\ \boxed{2} \text{ ADV } \left[\begin{array}{c} \text{CONT } \boxed{1} \\ \text{MOD } \boxed{4} \end{array} \right] \quad \boxed{4} \text{ NP } \left[\begin{array}{c} \text{CONT } \boxed{1} \\ \text{ADJS } \{ \boxed{2}^{obl} \} \end{array} \right] \\ | \qquad \qquad \qquad | \\ muzukasii \qquad \qquad \qquad hoñ \end{array}$$

125 Auf syntaktischer Ebene kann nicht zwischen Relativsatz und Attributsatz unterschieden werden. Der semantische Bezug zwischen adnominalem Adjektiv und Bezugsnomen wird daher offengelassen. Würde das syntaktische Bezugswort immer die semantische Argumentstelle von Adjektiven belegen, könnte 難しい (muzukasii; schwierig) einfacher durch

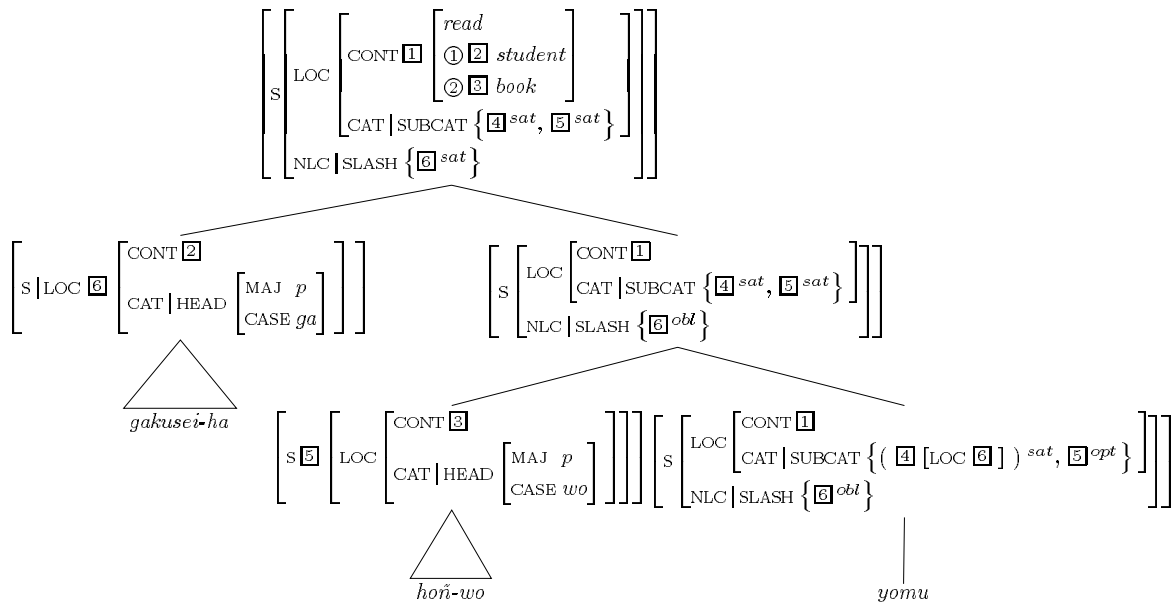
$$(268) \left[\begin{array}{c} difficult \\ \textcircled{1} cont \end{array} \right]$$

modelliert werden. Als Ergebnis der semantischen Komposition durch das Head-Adjunkt-Schema würde sich dann für 難しい 本 (muzukasii hoñ; schwieriges Buch) folgende Semantik ergeben:

Head-Füller-Schema

Beim Head-Füller-Schema ist der Mechanismus, durch den die Semantik des topikalisierten Elementes mit dem entsprechenden Argument der Verbsemantik unifiziert wird, komplizierter: Durch die Komplement-Extraktionsregel wird der LOCAL-Wert — und damit auch das CONTENT-Merkmal — des extrahierten Elementes der SUBCAT-Menge mit dem entsprechenden Element der SLASH-Menge unifiziert. Das SLASH-Merkmal-Prinzip transportiert dieses Element bis zu dem lokalen Baum, an dem es durch das Head-Füller-Schema realisiert wird. Der LOCAL-Wert des extrahierten Elementes wird dabei mit dem LOCAL-Wert seiner Realisierung unifiziert. Gleichzeitig kommt es dadurch zur Unifikation vom CONTENT-Wert des Topics und dem entsprechenden Merkmalwert der Verbsemantik:

(270) *Beispiel für die Anwendung des Head-Füller-Schemas (schematische Darstellung):*



3.3.3 Relativsätze

Im Abschnitt 3.2.4 "Relativsätze und Attributsätze", S. 144, wurde die Ähnlichkeit zwischen Relativsätzen und Attributsätzen anhand der folgenden Beispiele dargestellt:

(269) $\left[\begin{array}{l} \text{difficult} \\ \text{① book} \end{array} \right]$

Vgl. Abschnitt 3.2.4 "Relativsätze und Attributsätze", S. 144, und Abschnitt 3.3.3 "Relativsätze", S. 156.

- (271) a. i. [ステーキが] おいしいレストラン¹²⁶
 [sutēki-ga] oisii resutoran
 Steak-NOM gut schmecken Restaurant
 Ein Restaurant, in dem Steaks gut schmecken
 ii. おいしいステーキ
 oisii sutēki
 gut schmecken Steak
 gut schmeckende Steaks
- b. i. [ジョンが] [ステーキを] 食べたレストラン
 [jon-ga] [sutēki-wo] tabeta resutoran
 John-NOM Steak-ACC aß Restaurant
 Das Restaurant, in dem John Steak aß
 ii. [ジョンが] 食べたステーキ¹²⁷
 [jon-ga] tabeta sutēki
 John-NOM aß Steak
 Das Steak, das John aß

Da aus syntaktischer Perspektive allein nicht zu entscheiden ist, ob es sich um einen Relativsatz oder einen Attributsatz handelt, kann auch die Art der semantischen Relation zwischen Nebensatz und Bezugswort aus dieser Perspektive nicht entschieden werden. Im Falle eines Relativsatzes müßte die Semantik des Bezugswortes mit einem Argument der Verbsemantik des Nebensatzes koindiziert werden; im Falle eines Attributsatzes hingegen müßte eine andere, ebenfalls formal kaum zu erschließende semantische Relation R zwischen Bezugswort und Nebensatz angenommen werden:

- (272) a. i. $\left\{ \boxed{1} \text{ restaurant}, \boxed{2} \left[\begin{array}{c} \text{delicate} \\ \textcircled{1} \text{ steak} \end{array} \right], R? (\boxed{1}, \boxed{2}) \right\}$
 ii. $\left\{ \boxed{1} \text{ steak}, \left[\begin{array}{c} \text{delicate} \\ \textcircled{1} \boxed{1} \end{array} \right] \right\}$
- b. i. $\left\{ \boxed{1} \text{ restaurant}, \boxed{2} \left[\begin{array}{c} \text{eat} \\ \textcircled{1} \text{ john} \\ \textcircled{2} \text{ steak} \end{array} \right], R? (\boxed{1}, \boxed{2}) \right\}$
 ii. $\left\{ \boxed{1} \text{ steak}, \left[\begin{array}{c} \text{eat} \\ \textcircled{1} \text{ John} \\ \textcircled{1} \boxed{1} \end{array} \right] \right\}$

126 DBJG, S. 376, KS(B).

127 DBJG, S. 378, Notes (6).

Diese auf syntaktischer Ebene nicht auflösbare Mehrdeutigkeit kann auf verschiedene Weise behandelt werden:

1. Durch verschiedene alternative Strukturen;
2. Durch Unterspezifikation — entweder durch Offenlassen der semantischen Beziehung oder durch Angabe eines Wertebereiches für die beteiligten Merkmale.

Da die Auflösung semantischer Mehrdeutigkeiten nicht Thema dieser Arbeit ist, werden die semantischen Relationen von Relativ- bzw. Attributivsätzen in dem hier vorgestellten Modell offengelassen. Im übernächsten Kapitel wird bei den entsprechenden Flexiveinträgen allerdings auch die Modellierung durch alternative Analysen dargestellt.

Adjektive, Relativsätze, Attributsätze etc. werden semantisch einheitlich durch folgende Struktur beschrieben:

$$(273) \left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ <Semantik des Bezugswortes>} \\ \textcircled{2} \text{ <Semantik des Adjunktes>} \end{array} \right]$$

Die gesamte Merkmalstruktur erscheint an der Stelle der semantischen Darstellung, an der sonst die Semantik des Bezugswortes ohne Adjunkt stehen würde. Für die Sätze 本を読む (*hoñ-wo yomu*; ein Buch lesen) und 難しい本を良く読む (*muzukasii hoñ-wo yoku yomu*; ein schwieriges Buch gut lesen) erhalten wir damit folgende semantische Darstellungen:

$$(274) \begin{array}{ll} a. & i. \quad \text{本を読む} \\ & \text{hoñ-wo} \quad \text{yomu} \\ & \text{Buch-ACC} \quad \text{lesen} \\ & \text{ein Buch lesen} \\ & ii. \quad \left[\begin{array}{l} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ book} \end{array} \right] \\ b. & i. \quad \text{難しい本を良く読む} \\ & \text{muzukasii} \quad \text{hoñ-wo} \quad \text{yoku} \quad \text{yomu} \\ & \text{schwierig} \quad \text{Buch-ACC} \quad \text{gut} \quad \text{lesen} \\ & \text{ein schwieriges Buch gut lesen} \\ & ii. \quad \left[\begin{array}{l} \text{well} \\ \textcircled{1} \left[\begin{array}{l} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ book} \\ \textcircled{1} \left[\begin{array}{l} \text{difficult} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

Dieses Beispiel macht auch den Unterschied zwischen der Semantik von Adverben und der Semantik von Relativsätzen, Attributsätzen und Adjektiven deutlich: Bei Adverben wie 良く (*yoku*; gut) ist die semantische Relation zwischen Adjunkt und Bezugswort eindeutig und kann daher schon im Lexikon spezifiziert werden; bei Relativsätzen, Attributsätzen und Adjektiven wie 難しい (*muzukasii*; schwierig) hingegen bleibt sie unbestimmt und führt damit zu unterspezifizierten semantischen Strukturen.

Die Semantik unseres Beispielsatzes

Sämtliche syntaktischen und semantischen Mechanismen, die an der Ableitung des Beispielsatzes (168) beteiligt sind, wurden damit vorgestellt:

(275) a. 学生は難しい本を良く読むね。

gakusei-ha muzukasii hon-wo yoku yomu ne
 Student-TOP schwer buch-ACC gut lesen nicht wahr
 Der Student liest das schwierige Buch gut, nicht wahr?

b.
$$\left[\begin{array}{c} \text{isn't-it} \\ \left[\begin{array}{c} \text{well} \\ \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ student} \\ \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ book} \\ \textcircled{2} \left[\begin{array}{c} \text{difficult} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

Die Gegenüberstellung von Oberflächenform und semantischer Struktur, die wir diesem Kapitel vorangestellt haben, ist damit durch unser syntaktisches und semantisches Modell erklärt.

4 Struktur der morphologischen Beschreibung

4.1 Einleitung

Die Ableitung einer systematischen Beschreibung der japanischen Verbmorphologie und die Darstellung eines einfachen Modells der japanischen Syntax und Semantik waren Thema der letzten zwei Kapitel. Diese beiden Komponenten sollen nun zusammengefügt werden: Die Beschreibung der Verbmorphologie wird formalisiert und in das syntaktisch-semantische Modell integriert.

Die Mechanismen, auf denen das im Folgenden dargestellte morphologische Modell basiert, lassen sich in vier einfachen Punkten zusammenfassen:

1. *Morphologische Klassifizierung* — Morphologische Elemente (Stämme und Suffixe) werden entsprechend zweier Hierarchien kreuzklassifiziert:
 - (a) Hierarchie der *Verbklassen*;
 - (b) Hierarchie der *Stammformen*.

Suffixe subkategorisieren für Stammform und Verbkategorie ihrer Stämme.

2. *Morphologisches Lexikon* — Die klassifizierten morphologischen Formprimitiven (atomare morphologische Elemente) werden durch eine lexikalische Defaultvererbungshierarchie beschrieben.
3. *Agglutinationsschema* — Durch das Agglutinationsschema können Suffixe an atomare (lexikalische) und komplexe Stämme suffigiert werden.
4. *Stammerkmaldefault* — Suffixe modifizieren die syntaktischen und semantischen Eigenschaften ihrer Stämme; alle nicht durch das Suffix betroffenen Eigenschaften werden vom Stamm übernommen.

Der Schwerpunkt dieses Kapitels besteht in der Ausführung dieser Punkte. Die Reduktion der morphologischen Erscheinungen auf diese grundlegenden vier Mechanismen beruht auf den übrigen Kapiteln: Die Grundlagen der morphologischen Komponente wurden schon dargestellt; ihre Anwendung auf ein Fragment des Japanischen wird im nächsten Kapitel vorgeführt. Die *Integration* der schon beschriebenen Komponenten ist deshalb ebenfalls Aufgabe dieses Kapitels.

Im ersten Abschnitt wird das Verbmorphologische System, das im Kapitel 2 „*Japanische Verbmorphologie*“ aus der traditionellen Beschreibung abgeleitet wurde, als Constraintsystem formalisiert. Dieses Modell wird anschließend um die Mechanismen der syntaktischen und semantischen Modifikation erweitert, durch die die morphologischen Prozesse begleitet werden.

4.2 Agglutination

Genau wie die Mechanismen der Syntax und Semantik sind auch die Mechanismen der Morphologie abstrakt und ohne konkretes Anschauungsmaterial schwer nachvollziehbar. Wie in Kapitel 3 “*Syntax und Semantik als Rahmen der Morphologie*” sollen sie deshalb anhand von Beispielen vorgestellt werden. Die Ableitung der folgenden zwei Formen bildet daher das Gerüst des folgenden Kapitels:

- (276) a. 行った

i.tt-a
gehen-PAST
gegangen sein

- b. 行かせた

i.k.a-se.t-a
gehen-CAUS-PAST
zum Gehen veranlaßt haben.

Sie können beispielsweise in der folgenden Weise gebraucht werden:

- (277) a. むすめは大学へ行った。

musume-ha daigaku-he i.tt-a
Tochter-TOP Universität-DIR gehen-PAST
Die Tochter ging zur Universität.

- b. 鈴木さんはむすめを/に大学へ行かせた。¹²⁸

suzuki-sa~n-ha musume-wo/-ni daigaku-he i.k.a-se.t-a
Herr Suzuki-TOP Tochter-ACC/-DAT Universität-DIR gehen-CAUS-PAST
Herr Suzuki ließ seine Tochter zur Universität gehen.

Beide sind Formen des Verbes 行く (*i.k-u*; gehen, fahren), das — zumindest in der hier vorgestellten Gebrauchsweise — für zwei Komplemente subkategorisiert: eine Subjekt-PP, die den *Gehenden* bezeichnet, und eine Objekt-PP, durch die der *Zielort* des Gehens beschrieben wird. (276a) entspricht dem Perfekt von 行く, (276b) entspricht dem Perfekt des Kausativs von 行く. (276a) wird semantisch nur bezüglich der Zeit modifiziert und durch folgende Struktur dargestellt:

- (278)
$$\left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{go-to} \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right] \end{array} \right]$$

¹²⁸ DBJG, S. 387, KS (A).

Bei der Ableitung der Kausativ-Perfekt-Form (276*b*) hingegen verändert sich sowohl die Semantik als auch der Kasusrahmen von 行く:

$$(279) \quad \left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \boxed{\text{cont}}(\text{obj}) \\ \textcircled{3} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \boxed{\text{cont}} \\ \textcircled{2} \text{ cont}(\text{obj2}) \end{array} \right] \end{array} \right] \end{array} \right]$$

Der Kausativ von Verben wird im Deutschen und Englischen durch syntaktische Einbettung gebildet:

$$(280) \quad \begin{array}{ll} \text{Zugrunde liegendes Verb:} & \text{Kausativ:} \\ A \text{ geht nach } B & C \text{ veranlasst } A, \text{ nach } B \text{ zu gehen} \end{array}$$

Der Kausativ führt ein neues Objekt ein: den *Veranlassenden*. Semantisch wird der Kausativ durch Einbettung in eine Struktur des Typs *cause* modelliert:

$$(281) \quad \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) = \text{cont}(\text{subj}(\text{stem}))^{129} \\ \textcircled{3} \text{ cont}(\text{stem}) \end{array} \right]$$

Das neue Objekt — der *Veranlassende* — wird im Japanischen durch eine Verschiebung des Kasusrahmens des kausativierten Verbes modelliert:

$$(283) \quad \text{Transformation des SUBCAT-Wertes bei der Kausativbildung:}$$

$$\left\{ \text{PP}[\text{subj}]:\boxed{\text{cont}}, \dots \right\} \longrightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]^{\text{opt}}, \text{PP}[\text{obj}, \text{ni} \vee \text{wo}]:\boxed{\text{cont}}, \dots \right\}$$

Das unterliegende Subjekt des kausativierten Verbes erscheint an der Oberfläche als PP[*wo*] oder PP[*ni*]; der Veranlassende wird zum syntaktischen Subjekt des kausativischen Verbs.

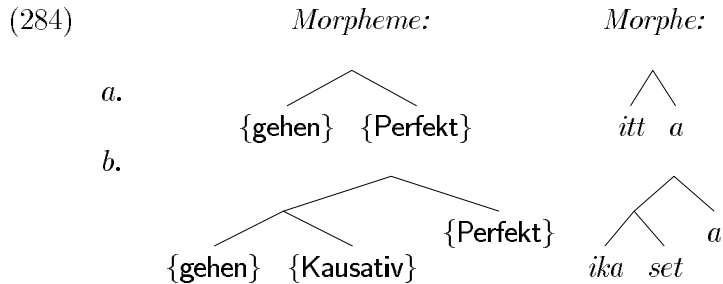
¹²⁹ Die Abkürzung **cont(subj(stem))** ist folgendermaßen definiert:

$$(282) \quad \begin{array}{ll} \text{Abkürzung:} & \text{abgekürzte AVM:} \\ \dots | \text{X cont}(\text{subj}(\text{stem})) & \left[\begin{array}{c} \dots | \text{X } \boxed{\text{cont}} \\ \dots | \text{STEM } [\dots | \text{SUBCAT } \{ \text{PP}[\text{subj}]:\boxed{\text{cont}}, \dots \}] \end{array} \right] \end{array}$$

mit $\text{X}:\boxed{\text{cont}}$ für $\boxed{\text{cont}} = \text{cont}(\text{X})$ in Subkategorisierungskontexten.

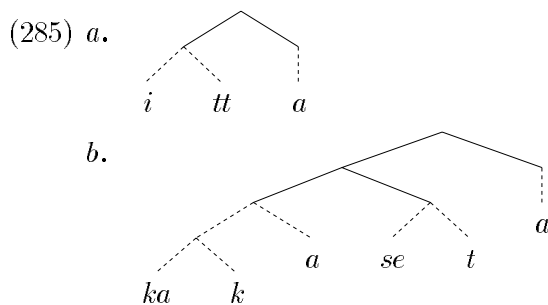
Analyse der Verbformen in Morpheme

Morphologisch haben die Beispiele (276) die folgende Struktur:

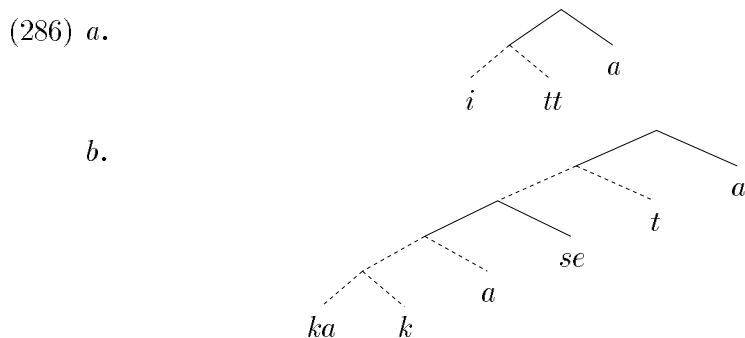


Analyse der Morphe in Kern und Clitics

In Kapitel 2 haben wir gesehen, daß sich die Allomorphe eines Morphems systematisch aus einem Kern, der allen Allomorphen gemeinsamen ist, und einer Menge von Clitics ableiten lassen. Für unser Beispiel (276) ergeben sich damit folgende Strukturen:



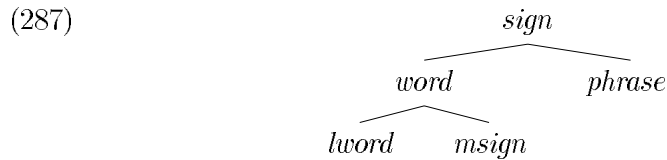
Die Struktur der Formen aus Morphemen als auch die darin eingebetteten Strukturen der Morphe aus Kern und Clitics sind linksrekursiv. Die zweistufige Ableitung (1. Ableitung der Morphe; 2. Ableitung der Verbform) kann daher durch eine einfachere, einstufige ersetzt werden:



Die Mechanismen, die zur Formalisierung dieser Art von Ableitung benötigt werden, sind minimal: Wir unterscheiden zwischen *atomaren* und *komplexen Strukturen* und kommen mit einem einzigen Schema — dem *Agglutinations-schema* — aus, um alle Formen zu erklären.

Morphologische Zeichen

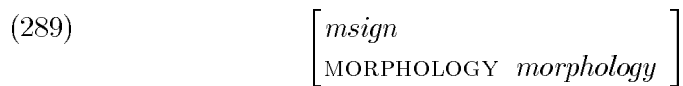
Syntaktische Zeichen werden in Wörter (*word*) und Phrasen (*phrase*) unterschieden. Formen, die morphologisch abgeleitet werden, sind ebenfalls Wörter. Wir führen deshalb den neuen Typ *lword* für Vollform-Lexikoneinträge ein, der gemeinsam mit den morphologischen Zeichen (*msign*) von *word* subsumiert wird:



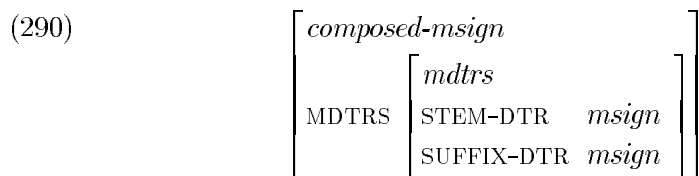
Morphologische Zeichen sind atomar oder komplex:



Zusätzlich zu den Merkmalen PHON und SYNSEM, die allen Zeichen gemeinsam sind, haben morphologische Zeichen das Merkmal MORPHOLOGY, unter dem ihre morphologischen Eigenschaften zusammengefaßt werden:



Komplexe morphologische Zeichen haben — genauso wie komplexe syntaktische Zeichen — ein Merkmal zur Darstellung ihrer Töchter (MDTRS). Agglutination wird linksrekursiv modelliert, als Anfügung eines Suffixes (SUFFIX-DTR) an einen atomaren oder komplexen Stamm (STEM-DTR):



Unter Verwendung der Akkürzung

$$(291) \quad \begin{array}{ll} \text{Abkürzung:} & \text{abgekürzte AVM:} \\ \left[\begin{array}{l} \text{ST } \boxed{1} \\ \text{SU } \boxed{2} \end{array} \right] & \left[\text{MDTRS } \left[\begin{array}{l} \text{STEM-DTR } \boxed{1} \\ \text{SUFFIX-DTR } \boxed{2} \end{array} \right] \right] \end{array}$$

können die morphologischen Strukturen (286a) und (286b) der Beispielformen (276a) und (276b) damit folgendermaßen dargestellt werden:

$$(292) \quad a. \quad \left[\begin{array}{l} \text{ST} \left[\begin{array}{l} \text{ST} \left[\text{P } \langle i \rangle \right] \\ \text{SU} \left[\text{P } \langle tt \rangle \right] \end{array} \right] \\ \text{SU} \left[\text{P } \langle a \rangle \right] \end{array} \right]$$

$$b. \quad \left[\begin{array}{l} \text{ST} \left[\begin{array}{l} \text{ST} \left[\begin{array}{l} \text{ST} \left[\begin{array}{l} \text{ST} \left[\text{P } \langle ka \rangle \right] \\ \text{SU} \left[\text{P } \langle k \rangle \right] \end{array} \right] \\ \text{SU} \left[\text{P } \langle a \rangle \right] \end{array} \right] \\ \text{SU} \left[\text{P } \langle se \rangle \right] \end{array} \right] \\ \text{SU} \left[\text{P } \langle t \rangle \right] \\ \text{SU} \left[\text{P } \langle a \rangle \right] \end{array} \right]$$

Morphologische Subkategorisierung

Damit ein Suffix an einen Stamm angeschlossen werden kann, muß sich dieser in einer bestimmten *morphologischen Form* befinden. Wie wir in der Einführung in die japanische Verbmorphologie gesehen haben, wird die *Form* durch die Kombination von Stammform (*v-stem*) und Verbkategorie (*vcat*) ausgedrückt. Die Bestandteile der Form *itta* können in der Darstellungsweise, die schon in Kapitel 2 (vgl. (163), S. 95) benutzt wurde, folgendermaßen charakterisiert werden:

(293) *Die morphologischen Bestandteile der Form itta:*

$$\begin{array}{llll} a. & \text{MAJ } v_{iku}^{cv}; & & \text{PHON } \langle i \rangle \\ b. & \text{MAJ } v_{t,w,R,iku}^{cv} \longrightarrow v_{t,w,R,iku}^{onbin}; & & \text{PHON } \langle tt \rangle \\ c. & \text{MAJ } v^{onbin} \longrightarrow v; & & \text{PHON } \langle a \rangle \end{array}$$

Die “Abbildung”, die durch die Suffixe realisiert wird, kann durch *morphologische Subkategorisierung* ausgedrückt werden:

(294) *Morphologische Subkategorisierung:*

	Suffix subkategorisiert für Formen der Kategorie:		Kategorie des Verbsuffixes / der Gesamtform:		phonologische Realisierung:
MAJ	$v_{ct,w,R,iku}^{cv}$	→	$v_{ct,w,R,iku}^{onbin}$;		PHON $\langle tt \rangle$

Zur Unterscheidung der morphologischen Kategorien wird das syntaktische Merkmal MAJ benutzt, das allerdings innerhalb der morphologischen Komponente differenziertere Werte annimmt:

(295) $[\text{SYNSEM} | \text{LOCAL} | \text{CAT} | \text{HEAD} | \text{MAJ } maj]$

In der Syntax wurden Verben und Adjektive unter der Kategorie *verb* zusammengefaßt. Verben und Adjektive konnten durch das Merkmal VCAT unterschieden werden:

(296) a. i. Verben: $\begin{bmatrix} verb \\ \text{VCAT } v \end{bmatrix}$ ii. Adjektive: $\begin{bmatrix} verb \\ \text{VCAT } adj \end{bmatrix}$

b.

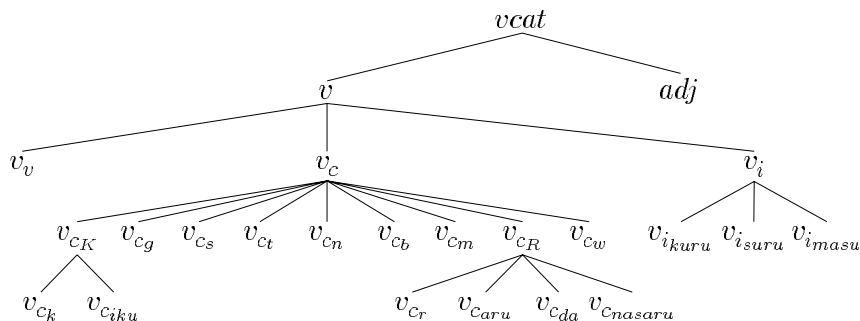
```

      vcat
     /  \
    v    adj

```

Um die morphologischen Kategorien in dieses Schema integrieren zu können, wird die *vcat*-Hierarchie so erweitert, daß sich die in Kapitel 2 “*Japanische Verbmorphologie*” entwickelte Hierarchie (147) (S. 88) ergibt. Diese wird hier als (297) noch einmal dargestellt:

(297) *Flexionsklassen der Verben:*



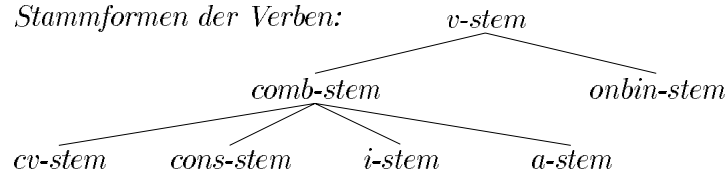
Ein zweites Merkmal — VSTEM — wird zur Unterscheidung der Stammformen verwendet. Es kann die Werte annehmen, die als mögliche Stammfor-

men im Kapitel 2 “*Japanische Verbmorphologie*” zusammengefaßt wurden (vgl. (159), S. 93):

(298) a. Merkmale zur Darstellung der Stammformen:

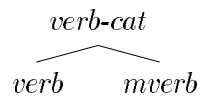
$$\begin{bmatrix} \text{VCAT} & \text{vcat} \\ \text{VSTEM} & \text{v-stem} \end{bmatrix}$$

b. Stammformen der Verben:



Um die Verbkategorien, wie sie in der Syntax verwendet werden, von den morphologischen Kategorien unterscheiden zu können, wird der Typ *verb-cat* eingeführt, der *verb* — für die Syntax — und *mverb* — für die Morphologie — subsumiert. VCAT ist für beide Subtypen adäquat, VSTEM nur für die Darstellung in der Morphologie:

(299) a.



b. i. $\begin{bmatrix} \text{verb-cat} \\ \text{VCAT} & \text{vcat} \end{bmatrix}$

ii. $\begin{bmatrix} \text{mverb} \\ \text{VSTEM} & \text{v-stem} \end{bmatrix}$

Der besseren Lesbarkeit halber werden auch in morphologischen Zeichen die Abkürzungen aus Kapitel 2 “*Japanische Verbmorphologie*” verwendet:

(300)	Abkürzung:	abgekürzte AVM:
	$\boxed{1}\boxed{2}$	$\begin{bmatrix} \text{verb-stem} \\ \text{VCAT} \quad \boxed{1} \\ \text{VSTEM} \quad \boxed{2}\text{-stem} \end{bmatrix}$
	v^{cons}	$\begin{bmatrix} \text{verb-stem} \\ \text{VCAT} \quad v \\ \text{VSTEM} \quad \text{cons-stem} \end{bmatrix}$
	adj^v	$\begin{bmatrix} \text{verb-stem} \\ \text{VCAT} \quad adj \\ \text{VSTEM} \quad v\text{-stem} \end{bmatrix}$
	$v_{c_k}^{onbin}$	$\begin{bmatrix} \text{verb-stem} \\ \text{VCAT} \quad v_{c_k} \\ \text{VSTEM} \quad \text{onbin-stem} \end{bmatrix}$
	$v_{c_n, b, m}^{onbin}$	$\begin{bmatrix} \text{verb-stem} \\ \text{VCAT} \quad v_{c_n} \vee v_{c_b} \vee v_{c_m} \\ \text{VSTEM} \quad \text{onbin-stem} \end{bmatrix}$
	<i>usw.</i>	

Morphologische Kategorie und morphologische Subkategorisierung

Als *lexikalische Kategorie* wird der Wert des Merkmals MAJ — also beispielsweise n , p , v , $v_{c_k}^{onbin}$ etc. — bezeichnet. In Analogie zur *syntaktischen Kategorie*, die durch das Merkmal CAT bestimmt wird, bezeichnen wir das Merkmal MCAT als *morphologische Kategorie*.

Wie alle Eigenschaften von Zeichen werden auch morphologische Kategorie und morphologische Subkategorisierung durch Merkmale ausgedrückt: die morphologische Kategorie durch MCAT; die morphologische Subkategorisierung durch STEM.

(301) *Morphologische Kategorie MCAT*
und *morphologische Subkategorisierung STEM*:

$$\left[\text{MORPHOLOGY} \mid \text{MCAT} \begin{bmatrix} \text{MHEAD} \mid \text{MAJ} \quad maj \\ \text{STEM} \quad stem \end{bmatrix} \right]$$

Das Merkmal zur Darstellung der lexikalischen Kategorie MAJ hat den gleichen Wertebereich wie das MHEAD-Merkmal MAJ und wird mit diesem, wie später gezeigt, koindiziert. STEM, durch das die morphologische Subkategorisierung ausgedrückt wird, hat den Wertetyp *stem*:

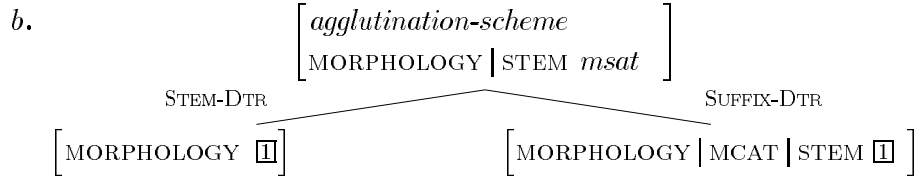
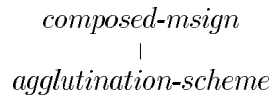


Ist ein morphologisches Zeichen gesättigt, wird dieses durch *msat* (für *morphological satisfied*) markiert; ungesättigte Suffixe subkategorisieren für den MORPHOLOGY-Wert ihres Stammes.

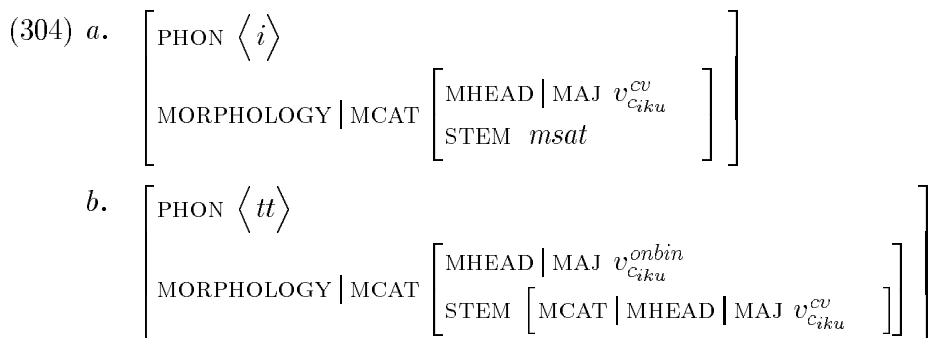
4.2.1 Das Agglutinationsschema

Durch das Agglutinationsschema kann ein Suffix an einen lexikalischen bzw. schon abgeleiteten komplexen Stamm angeschlossen werden. Das Agglutinationsschema ist zur Zeit das einzige morphologische Schema. Es wird als Constraint zum Typ *agglutination-scheme*, dem einzigen Subtyp von *composed-msign*, formuliert:

(303) a. *Das Agglutinationsschema:*



Die Form 行った (*i.tt-a*) aus (276a) beispielsweise setzt sich aus den folgenden Zeichen zusammen (es werden nur die Merkmale PHON, MAJ und STEM dargestellt):



$$c. \left[\begin{array}{l} \text{PHON } \langle a \rangle \\ \text{MORPHOLOGY | MCAT } \left[\begin{array}{l} \text{MHEAD | MAJ } v \\ \text{STEM } \left[\text{MCAT | MHEAD | MAJ } v^{onbin} \right] \end{array} \right] \end{array} \right]$$

Durch Anwendung des Agglutinationsschemas ergibt sich für 行った (*i.tt-a*) folgende Ableitung:

$$(305) \quad \begin{array}{c} \left[\text{MORPHOLOGY | MCAT } \left[\begin{array}{l} \text{MHEAD } \boxed{1} \text{ [MAJ } v \text{]} \\ \text{STEM } msat \end{array} \right] \right] \\ \swarrow \quad \searrow \\ \left[\text{MORPHOLOGY } \boxed{2} \left[\text{MCAT } \left[\begin{array}{l} \text{MHEAD } \boxed{3} \text{ [MAJ } v_{iku}^{onbin} \text{]} \\ \text{STEM } msat \end{array} \right] \right] \right] \quad \left[\begin{array}{l} \text{PHON } \langle a \rangle \\ \text{MORPHOLOGY | MCAT } \left[\begin{array}{l} \text{MHEAD } \boxed{1} \\ \text{STEM } \boxed{2} \end{array} \right] \end{array} \right] \\ \swarrow \quad \searrow \\ \left[\begin{array}{l} \text{PHON } \langle i \rangle \\ \text{MORPHOLOGY } \boxed{4} \left[\text{MCAT } \left[\begin{array}{l} \text{MHEAD | MAJ } v_{iku}^{cu} \\ \text{STEM } msat \end{array} \right] \right] \end{array} \right] \quad \left[\begin{array}{l} \text{PHON } \langle tt \rangle \\ \text{MORPHOLOGY | MCAT } \left[\begin{array}{l} \text{MHEAD } \boxed{3} \\ \text{STEM } \boxed{4} \end{array} \right] \end{array} \right] \end{array}$$

In den lokalen Bäumen dieser Ableitung sind die Merkmale MHEAD von Mutter und Suffixtochter koindiziert. Wie es zu dieser Koindizierung kommt, wird im nächsten Abschnitt erklärt.

4.2.2 Das morphologische Head-Merkmal-Prinzip

Wie bei komplexen syntaktischen Zeichen errechnen sich die Eigenschaften komplexer morphologischer Zeichen aus den Merkmalen der Töchter. Genau wie in der Syntax werden Merkmale, die sich aus dem morphologischen Head — also der Suffixtochter — ergeben, unter einem Head-Merkmal — in diesem Fall MHEAD für *morphologischer Head* — zusammengefaßt. Das morphologische Head-Merkmal-Prinzip ähnelt ebenfalls seiner syntaktischen Entsprechung:

(306) *Das morphologische Head-Merkmal-Prinzip:*

$$\begin{array}{c} \left[\text{MORPHOLOGY | MCAT | MHEAD } \boxed{1} \right] \\ \swarrow \quad \searrow \\ \text{STEM-DTR} \quad \text{SUFFIX-DTR} \\ \left[\quad \right] \quad \left[\text{MORPHOLOGY | MCAT | MHEAD } \boxed{1} \right] \end{array}$$

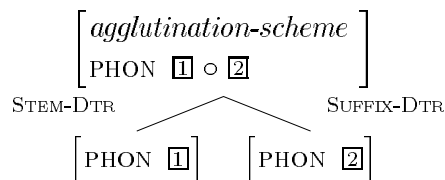
In der hier entwickelten Morphologie gibt es nur drei Merkmale, die durch das Suffix bestimmt und durch das morphologische Head-Merkmal-Prinzip betroffen werden:

1. MAJ — die Kategorie der Verbform.
2. MOD — Dieses Merkmal ist mit dem entsprechenden syntaktischen Merkmal koindiziert: Der MOD-Wert von flektierten Verben wird durch das Flexiv (das letzte Suffix) bestimmt (vgl. 4.2.4, S. 172).
3. LEVEL — Nur eine bestimmte Menge von Suffixen können den morphologischen Prozess abschließen: Das Resultat ihrer Agglutinierung bildet damit eine Verbform, die in der Syntax verwendet werden kann. Suffixe mit dieser Eigenschaft (Flexive) werden von den übrigen Suffixen durch das LEVEL-Merkmal unterschieden (vgl. Abschnitt 4.4 “*Das LEVEL-Merkmal*”, S. 185).

4.2.3 Das morphologische PHON-Merkmal-Prinzip

Der phonologische Wert komplexer morphologischer Zeichen ergibt sich auf die gleiche Weise wie bei komplexen syntaktischen Zeichen: durch Konkatination der entsprechenden Werte der Tochterzeichen.

(307) *Das morphologische PHON-Merkmal-Prinzip:*



4.2.4 Das MORPHOLOGY-Merkmal-Prinzip (MMP)

Linguistische Prinzipien betreffen in den meisten Fällen nicht einzelne Merkmale, sondern Gruppen von Merkmalen. Die Unterteilung der Merkmale in Gruppen spiegelt sich in der Geometrie der Zeichen: Unter dem Merkmal HEAD beispielsweise finden sich die Merkmale, die von dem HEAD-Merkmal-Prinzip betroffen werden; unter NONLOCAL solche, die für nicht-lokale Abhängigkeiten verantwortlich sind etc. Da sich diese Gruppierungen auf “natürliche Weise” aus dem linguistischen Verhalten ergeben, werden sie als *natürliche Klassen* bezeichnet.

Die Merkmale, die für das syntaktisch-semantische Verhalten der Zeichen verantwortlich sind, werden unter dem Pfad SYNSEM zusammengefaßt. Die Geometrie der SYNSEM-Werte spiegelt daher die natürlichen Klassen hinsichtlich Syntax und Semantik.

Einige SYNSEM-Merkmale — beispielsweise MAJ und MOD — spielen auch für die Morphologie eine wichtige Rolle. Zudem können Suffixe im Japanischen alle syntaktisch-semantischen Eigenschaften ihrer Stämme modifizieren bzw. überschreiben. MAJ, MOD und SYNSEM als Gesamtes sind also auch für die morphologischen Prozesse relevant.

Syntaktische und semantische Prinzipien erfordern eine andere Gruppierung der Merkmale als morphologische Prinzipien. Die natürlichen Klassen hinsichtlich der Morphologie sind daher anders als die natürlichen Klassen hinsichtlich Syntax und Semantik. Durch das MORPHOLOGY-Merkmal-Prinzip werden die morphologisch relevanten Merkmale in natürliche Klassen entsprechend morphologischer Gesichtspunkte neu geordnet:

(308) *Das MORPHOLOGY-Merkmal-Prinzip:*

$$\left[\begin{array}{c} \text{MORPHOLOGY} \\ \text{SYNSEM } \boxed{3} \end{array} \left[\begin{array}{c} \text{MCAT} \mid \text{MHEAD} \left[\begin{array}{c} \text{MAJ } \boxed{1} \\ \text{MOD } \boxed{2} \end{array} \right] \\ \text{SYNSEM } \boxed{3} \left[\text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \left[\begin{array}{c} \text{MAJ } \boxed{1} \\ \text{MOD } \boxed{2} \end{array} \right] \right] \end{array} \right] \right]$$

Die morphologischen Merkmale MAJ, MOD und SYNSEM unter MORPHOLOGY sind mit den entsprechenden syntaktisch-semantischen Merkmalen unter SYNSEM koindiziert. Bis auf PHON können daher alle Merkmale auch über MORPHOLOGY erreicht werden. In den morphologischen Betrachtungen können wir uns also künftig auf die Darstellung von PHON und MORPHOLOGY beschränken.

4.2.5 Das STEM-Merkmal-Default

In komplexen morphologischen Zeichen wird

- PHON durch das morphologische PHON-Merkmal-Prinzip,
- STEM durch das Agglutinationsschema und
- MHEAD durch das morphologische Head-Merkmal-Prinzip

berechnet. Alle anderen Merkmale werden entweder durch das Suffix überschrieben oder — unverändert oder in modifizierter Form — vom Stamm übernommen. Für diese Übernahme bzw. Modifikation der Stammerkmale ist das STEM-Merkmal-Default verantwortlich.

Defaultmäßig werden alle Stammerkmale übernommen; nur wenn in der Spezifikation der Suffixe explizit ein anderer Wert angegeben wird, werden die Werte des Stammes überschrieben bzw. modifiziert. Die Merkmale des Stammes können daher ebenfalls als “Defaults” angesehen werden; Abweichungen von diesen “Defaults” müssen durch Suffixe spezifiziert werden.

In diesem Defaultverhalten spiegelt sich ein wichtiger Unterschied zwischen Flexion und Agglutination, auf den schon in Kapitel 2 “*Japanische Verbmorphologie*” hingewiesen wurde: Während sich die Verben in flektierenden Sprachen nach festen Bauplänen aus einzelnen Morphemen zusammensetzen, werden die Verbformen agglutinierender Sprachen relativ frei durch Suffigierung von Morphemen gebildet.¹³⁰ In flektierenden Sprachen kann damit jedes morphologisch beeinflussbare Merkmal durch ein obligatorisches Morphem festgelegt werden; im Fall einer agglutinierenden Morphologie hingegen werden Merkmale nur dann ausgedrückt, wenn sie von einem sonst angenommenen Standard abweichen. Das Prinzip, nur solche grammatischen oder inhaltlichen Sachverhalte explizit auszudrücken, die nicht einem angenommenen Standard entsprechen oder aus dem Kontext erschlossen werden können, findet sich im Japanischen nicht nur in der Morphologie, sondern auch in anderen Bereichen, etwa der Syntax: Objekte des Verbes müssen in indoeuropäischen Sprachen ausgedrückt werden, notfalls durch semantisch leere Pronomina; im Japanischen hingegen werden Objekte weggelassen, wenn sie sich aus dem Zusammenhang ergeben.

Modellierung des Stammerkmaldefaults durch Defaultunifikation

Das *Default*-Verhalten japanischer Verben läßt sich auch auf formaler Seite am besten durch Defaults repräsentieren:¹³¹ Lexikalische Verbstämme werden durch grammatische *Defaults* gekennzeichnet, die durch Suffixe über-

130 Daß in den flektierenden Sprachen feste Baupläne mit obligatorischen Morphemen für Numerus, Kasus etc. angenommen werden, wird insbesondere durch das sogenannte ϕ -Morphem deutlich: Ist die Realisierung einer morphologischen Kategorie an der Oberfläche nicht sichtbar, wird von einem Morph der Länge 0 — einem sogenannten ϕ -Morph — ausgegangen.

131 Wie in Kapitel 1 erklärt, werden aus Defaulthierarchien vor der Anwendung in der Analyse bzw. Generierung monotone Hierarchien berechnet. Defaults sind also prinzipiell für unser Modell nicht erforderlich. Bei seiner Entwicklung wurde sogar auf sie verzichtet: Für das Perfekt-Suffix beispielsweise wurde nicht nur angegeben, was sich bezüglich der Stammerkmale verändert, sondern ebenso, was durch das Perfektsuffix *nicht* verändert wird. Ein solches Modell ist jedoch nicht nur redundant und aufwendiger zu erstellen, sondern auch schwerer zu verstehen. Es ist daher sinnvoll, das Modell um einen Mechanismus

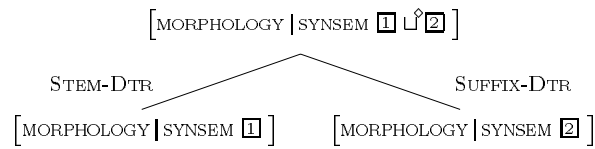
geschrieben werden können. Für Verbstämme wird beispielsweise als Default die Verwendung als Prädikat eines Hauptsatzes spezifiziert. Nur wenn ein Suffix — beispielsweise das adnominale Präsens-Flexiv — dieses Default überschreibt, weicht eine Verbform von dieser Defaultannahme ab.

Umgangssprachlich läßt sich das STEM-Merkmal-Default folgendermaßen beschreiben:

- (309) *Von den Merkmalen abgesehen, die durch das Suffix überschrieben werden, ist der SYNSEM-Wert eines komplexen morphologischen Zeichens gleich dem SYNSEM-Wert seines Stammes.*

In einer intuitiven, halbformalen Schreibweise läßt sich das Stammerkmal-Default folgendermaßen veranschaulichen:

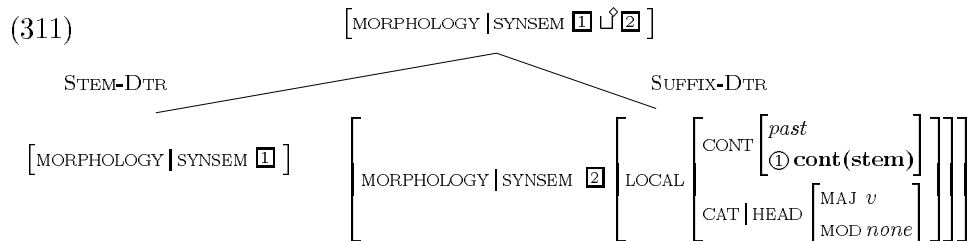
- (310) *Das STEM-Merkmal-Default, erste Näherung:*



Das Symbol $\sqcup^{\hat{}}$ steht dabei für eine Operation, die als *Defaultunifikation* bezeichnet werden kann: Ergebnis dieser Operation ist eine Struktur, die — abgesehen von den Werten, die in $\boxed{2}$ spezifiziert werden — mit $\boxed{1}$ übereinstimmt; die Werte, die in $\boxed{2}$ angegeben werden, überschreiben also die entsprechenden Werte in $\boxed{1}$.

Beispiel

Die Berechnung der Merkmalwerte der Perfektform von 行< (*i.k-u*; gehen, fahren) aus Beispiel (276a) läßt sich mit diesem Schema folgendermaßen darstellen:

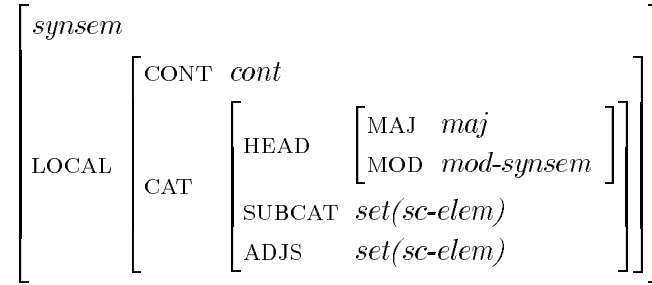


zu erweitern, der es gestattet, die zu übernehmenden Merkmale zu erschließen: die Defaulthierarchien. Das Ergebnis ist die hier dargestellte Defaultbeschreibung der japanischen Verbmorphologie.

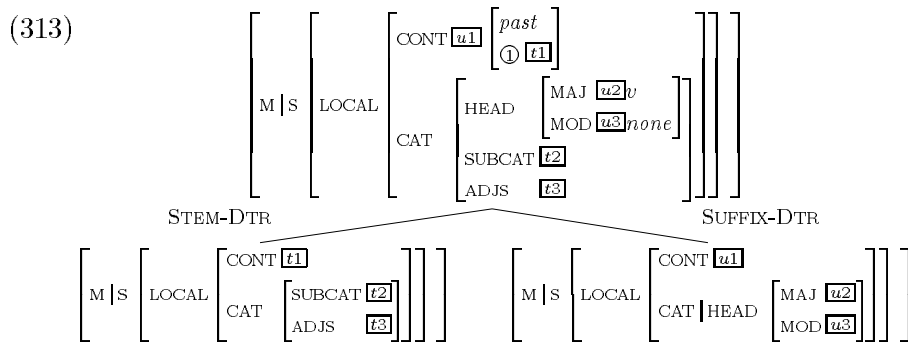
cont(stem) ist eine Abkürzung, mit der der Wert des Merkmals CONT des Stammes bezeichnet wird.

Zur Darstellung der Ergebnisse von (311) setzen wir die folgende Grundstruktur der SYNSEM-Werte voraus:

(312) *synsem-Grundstruktur der morphologischen Zeichen:*



Das Ergebnis der Auflösung der Defaultunifikation in (311) sieht damit folgendermaßen aus:¹³²



Modellierung des Stammerkmaldefaults durch Defaulthierarchien

Der Constraint (310) veranschaulicht die Funktion des Stammerkmaldefaults. Da es aber als Constraint zum Agglutinationsschema formuliert ist, dessen Töchter erst zum Zeitpunkt der Resolution instantiiert werden, erfordert es ein nichtmonotones Resolutionsverfahren. Es kann also in dieser Form nicht

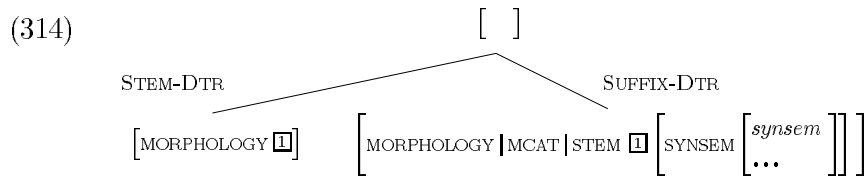
¹³² In den folgenden Strukturen werden einige Merkmale aus Platzgründen durch ihre Anfangsbuchstaben abgekürzt: M $\hat{=}$ MORPHOLOGY, S $\hat{=}$ SYNSEM.

Tags mit dem Präfix *t* (z.B. $\boxed{t1}$, $\boxed{t2}$, ...) bezeichnen Werte, die vom Stamm übernommen werden; vom Suffix hingegen werden die Werte übernommen, deren Tags mit *u* beginnen (z.B. $\boxed{u1}$, $\boxed{u2}$, ...).

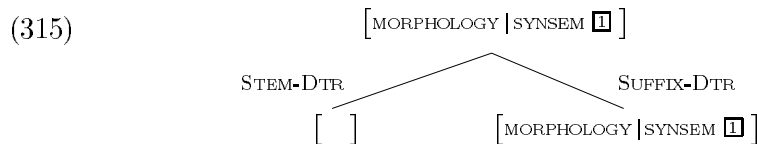
durch Defaulthierarchien formalisiert werden, wie sie in Abschnitt 1.7 eingeführt wurden: Defaults in Defaulthierarchien lassen sich nur auf Merkmalwerte anwenden, deren Ursprung schon aus der Constrainthierarchie erschlossen werden kann. In den folgenden Abschnitten wird daher das Stammerkmaldefault in entsprechender Weise umformuliert.

1. Verlegung des Defaults vom Schema auf das Suffix

Der Constraint (310) wird so umformuliert, daß die SYNSEM-Werte von Stamm- und Mutterzeichen im Suffixzeichen zugänglich sind. Der SYNSEM-Wert des Stammes ist dank der morphologischen Subkategorisierung schon vom Suffix aus erreichbar:

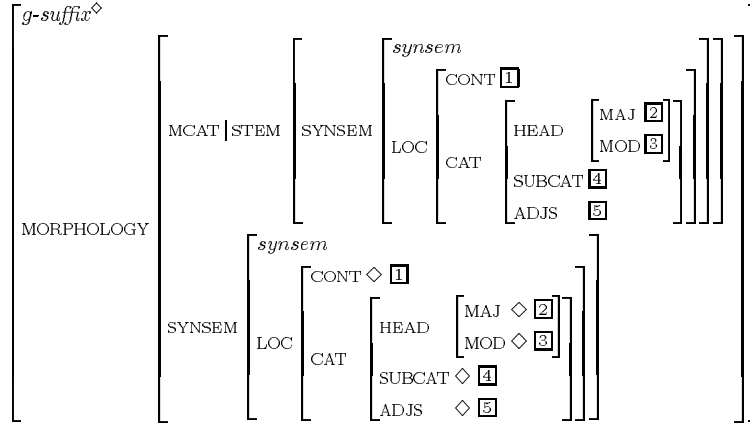


Durch einen neuen Constraint zum Typ *agglutination-scheme* wird der SYNSEM-Wert von Mutter und Suffix koindiziert:

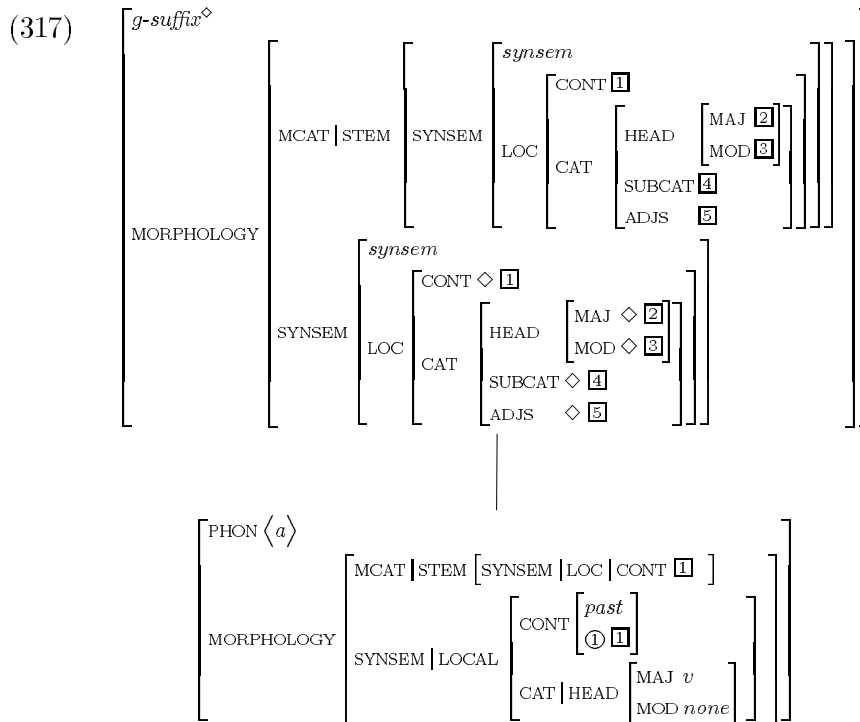


Im Suffixzeichen sind damit beide SYNSEM-Merkmale — das der Mutter und das des Stammes — verfügbar. Wird die Merkmalstruktur (312) als Grundstruktur der SYNSEM-Werte angenommen, läßt sich das Stammerkmaldefault damit folgendermaßen formulieren:

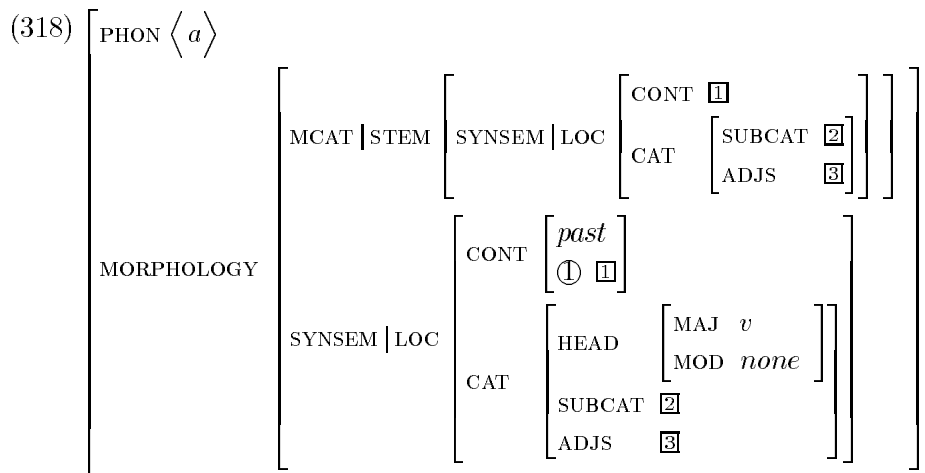
(316) *Das STEM-Merkmal-Default als Defaulttypenconstraint:*



Der Typ *g-suffix*, als dessen Default-Constraint das Stammerkmaldefault formuliert wird, subsumiert alle morphologischen Suffixe, wie auf Seite 184 gezeigt wird. Das Perfekt-Suffix, das zur Veranschaulichung des Stammerkmaldefaults in (311) und (313) herangezogen wurde, wird damit ebenfalls durch den Default-Constraint (316) subsumiert:



Durch das Berechnen der monotonen Constrainthierarchie aus der Default-Constrainthierarchie ergibt sich damit aus der Defaultdarstellung des Perfektsuffixes das folgende monotone Zeichen:



4.3 Lexikalisch-morphologische Regeln

In Abschnitt 3.2.2, S. 137 wurden lexikalische Regeln als Mittel eingeführt, systematische Beziehungen zwischen lexikalischen Elementen abzubilden: Mit ihrer Hilfe lassen sich beispielsweise aus Lexikoneinträgen von Nomen, die keine Adjunkte vorsehen, systematisch Zeichen ableiten, die durch ein oder mehrere Adjunkte modifiziert werden können. Ein weiteres Beispiel betraf die Modellierung nichtlokaler Abhängigkeiten: Durch eine lexikalische Regel können Objekte von der SUBCAT-Menge in die SLASH-Menge überführt und damit als topikalisiertes Objekt realisiert werden (vgl. (246), S. 142).

Es wurden jedoch auch Zeichen vorgestellt, deren Ableitung nicht — oder nur mit erheblichem Aufwand — durch lexikalische Regeln modelliert werden kann, und die daher als alternative Lexikoneinträge modelliert wurden. Zu diesen nicht ableitbaren Zeichen gehören Verben mit komplexer Semantik, deren innere, eingebettete Struktur durch Adverbien modifiziert werden kann. Die Verbform 行かせた (*i.k.a-se.t-a*; zum Gehen veranlassen) ist ein Beispiel für ein solches Verb mit komplexer Semantik:

(319) 健が黙って直美を大学へ行かせた。¹³³

- | | | | | |
|---------------|----------------|-----------------|-------------------|---------------------|
| <i>keñ-ga</i> | <i>damatte</i> | <i>naomi-wo</i> | <i>daigaku-he</i> | <i>i.k.a-se.t-a</i> |
| Ken-NOM | schweigend | Naomi-ACC | Universität-DIR | gehen-CAUS-PAST |
- a.* Ken brachte Naomi schweigend dazu, zur Schule zu gehen.
b. Ken brachte Naomi dazu, schweigend zur Schule zu gehen.

Setzt man für 行かせた (*i.k.a-se.t-a*; zum Gehen veranlassen)

¹³³ vgl. Manning et al. 1999, Beispiel (24)b.

$$(320) \quad \left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \boxed{1} \text{ cont(obj)} \\ \textcircled{3} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \boxed{1} \\ \textcircled{2} \text{ cont(obj2)} \end{array} \right] \end{array} \right] \end{array} \right]$$

als Semantik voraus, ergeben sich die folgenden Lesarten für (319):

- (321) i. *Ken brachte Naomi schweigend dazu, zur Schule zu gehen.* ii. *Ken brachte Naomi dazu, schweigend zur Schule zu gehen.*

$$\left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{silently} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ ken} \\ \textcircled{2} \boxed{1} \text{ naomi} \\ \textcircled{3} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \boxed{1} \\ \textcircled{2} \text{ school} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \quad \left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ ken} \\ \textcircled{2} \boxed{1} \text{ naomi} \\ \textcircled{3} \left[\begin{array}{c} \text{silently} \\ \textcircled{1} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \boxed{1} \\ \textcircled{2} \text{ school} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

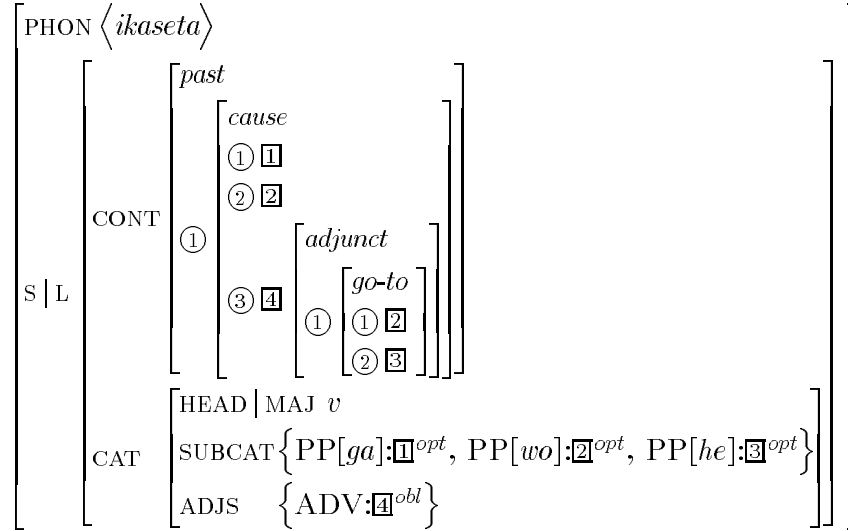
Die Lesarten werden aus folgenden alternativen Zeichen für 行かせた (*i.k.a-se.t-a*; zum Gehen veranlassen) abgeleitet:¹³⁴

- (322) a. 行かせた (*ikaseta*) mit weitem Adverbscopus:

$$\left[\begin{array}{c} \text{PHON } \langle \text{ikaseta} \rangle \\ \text{S | L} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \text{past} \\ \textcircled{1} \boxed{4} \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \boxed{1} \\ \textcircled{2} \boxed{2} \\ \textcircled{3} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \boxed{2} \\ \textcircled{2} \boxed{3} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{HEAD | MAJ } v \\ \text{SUBCAT } \{ \text{PP}[\text{ga}]:\boxed{1}^{\text{opt}}, \text{PP}[\text{wo}]:\boxed{2}^{\text{opt}}, \text{PP}[\text{he}]:\boxed{3}^{\text{opt}} \} \\ \text{ADJS } \{ \text{ADV}:\boxed{4}^{\text{obl}} \} \end{array} \right] \end{array} \right] \end{array} \right]$$

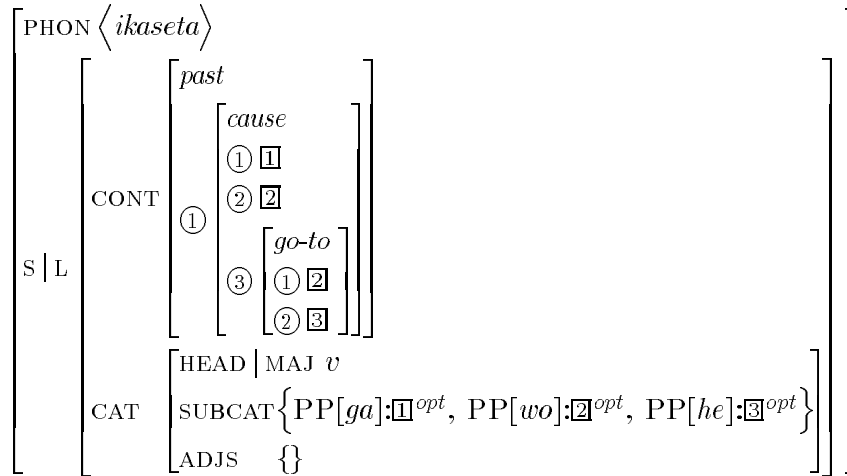
¹³⁴ Dargestellt werden nur die Merkmale, die für die folgenden Überlegungen relevant sind.

b. 行かせた (*ikaseta*) mit engem Adverbscopus:



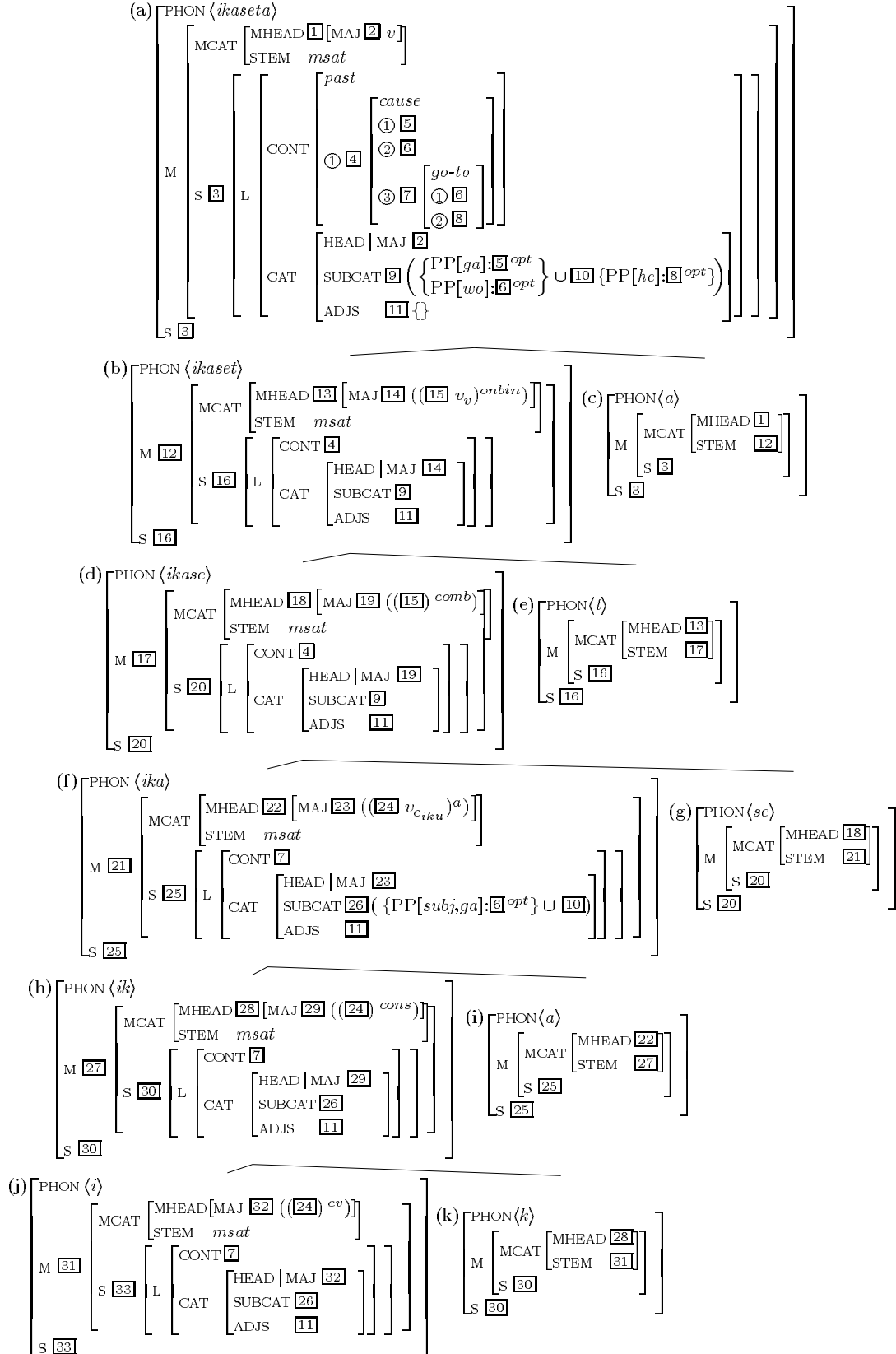
Beiden liegt das folgende Zeichen zugrunde, das das gleiche Verb ohne Adjunkte beschreibt.

(323) 行かせた (*ikaseta*) ohne Adverben:



Dieses leitet sich morphologisch folgendermaßen ab:

(324) *Ableitung der Form 行かせた (i.k.a-se.t-a; zum Laufen veranlassen).*¹³⁵

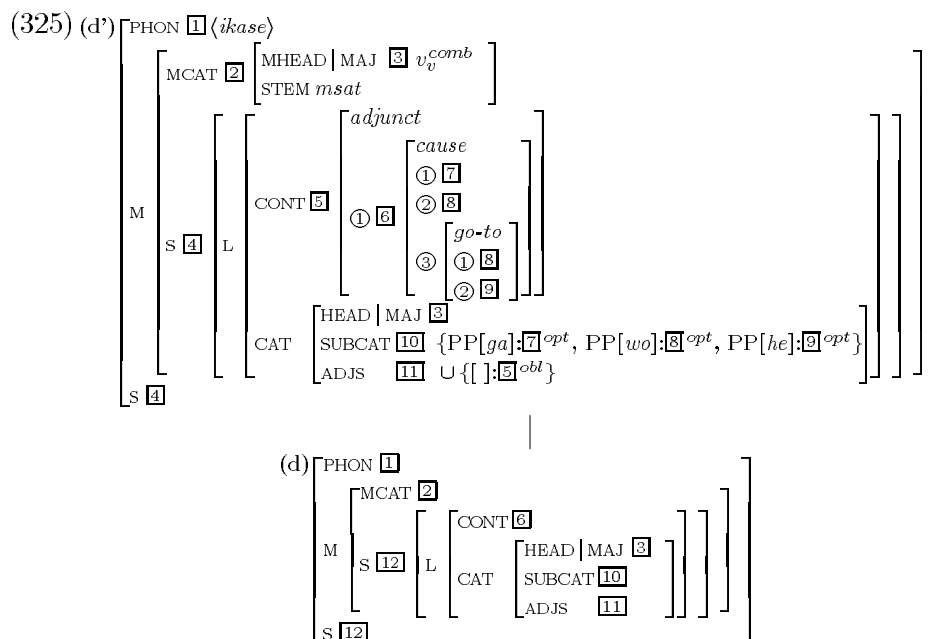


¹³⁵ Es werden nur die Merkmale dargestellt, die für die folgenden Betrachtungen wichtig sind.

(324) zeigt, wie sich die semantische Struktur von 行かせた (*i.k.a-se.t-a*) aus der Semantik des lexikalischen Stammes von 行く (*iku*) und den semantischen Modifikationen der Suffixe ergibt. Aus dem Vergleich dieser Ableitung mit den Lexikoneinträgen (322a) und (322b) wird ersichtlich, daß sich diese durch den Einschub eines Ableitungsschrittes erzeugen lassen. (322a) beispielsweise ergibt sich durch den Einschub eines neuen Ableitungsschrittes in (324), der die Struktur (d) ähnlich modifiziert wie die schon vorgestellte lexikalische Regel zur Adjunkteinführung (vgl. (239), S. 139):

1. Der CONT-Wert wird in eine Struktur des Typs *adjunct* eingebettet;
2. ADJS wird um ein neues Element erweitert, dessen CONT-Wert mit der neuen semantischen Struktur koindiziert ist.

Wir erhalten damit den folgenden zusätzlichen Ableitungsschritt:



(322*b*) läßt sich durch Einschub eines entsprechenden Ableitungsschrittes an der Stelle (j) der Ableitung (324) erzeugen.

In beiden Fällen ähnelt der Ableitungsschritt, durch den das Adjunkt eingeführt wird, der lexikalischen Regel zur Adjunkteinführung (239), S. 139, die wir in Kapitel 3 kennengelernt haben. Im Unterschied zu dieser setzt die hier vorgestellte Ableitung jedoch nicht auf dem Lexikon auf, sondern greift in die morphologische Ableitung ein. Wegen ihres Status zwischen morphologischen und lexikalischen Regeln bezeichnen wir sie als *lexikalisch-morphologische Regel*.

4.3.1 Lexikalisch-morphologische Regel zur Adjunkteinführung

Die lexikalisch-morphologische Regel zur Adjunkteinführung darf nur an bestimmten Stellen in die morphologische Ableitung eingreifen. Die Zeichen, die durch sie modifiziert werden dürfen, werden deshalb durch

$$(326) \quad \left[\text{MORPHOLOGY} \mid \text{MCAT} \mid \text{MHEAD} \mid \text{ADJUNCT} \mid + \right]$$

markiert.

Das Merkmal ADJUNCT wird bei atomaren morphologischen Zeichen durch den Lexikoneintrag und bei komplexen morphologischen Zeichen durch das letzte Suffix bestimmt. ADJUNCT läßt sich daher als morphologisches Head-Merkmal beschreiben.

Zur Beschreibung lexikalisch-morphologischer Regeln wird die gleiche Notation verwendet wie zur Beschreibung lexikalischer Regeln: Dargestellt werden nur die modifizierten Merkmale; alle anderen Merkmalwerte werden übernommen.^{136, 137}

$$(327) \quad \left[\begin{array}{c} \text{M} \mid \left[\begin{array}{c} \text{MCAT} \mid \text{MHEAD} \mid \text{ADJUNCT} \mid +^{138} \\ \text{S} \mid \text{L} \left[\begin{array}{c} \text{CONT} \mid \text{①} \\ \text{CAT} \mid \text{ADJS} \mid \text{②} \end{array} \right] \end{array} \right] \end{array} \right] \Rightarrow \left[\begin{array}{c} \text{M} \mid \text{S} \mid \text{L} \left[\begin{array}{c} \text{CONT} \mid \text{③} \left[\begin{array}{c} \text{adjunct} \\ \text{④} \mid \text{⑤} \end{array} \right] \\ \text{CAT} \mid \text{ADJS} \mid \text{②} \cup \{ \text{⑥} \mid \text{③}^{obl} \} \end{array} \right] \end{array} \right]$$

4.4 Klassifizierung morphologischer Formen

Morphologische Formen (*m-form*) unterteilen sich — wie im Kapitel 5 dargestellt — in Verbstämme (*v-stem*) und Suffixe im weiteren Sinne (*g-suffix*); Suffixe im weiteren Sinne wiederum werden in Suffixe im engeren Sinne (*m-suffix*) und Clitics (*clitic*) untergliedert. Suffixe im engeren Sinne können in Suffixe (*suffix*) und Flexive (*flexive*) unterschieden werden:

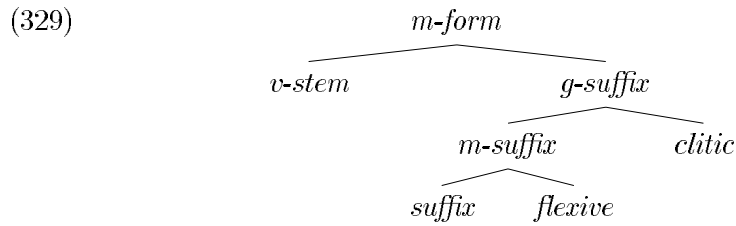
¹³⁶ Für die Modellierung von lexikalischen und lexikalisch-morphologischen Regeln bietet sich daher ebenfalls ein Default an. Dieses könnte ähnlich wie das STEM-Merkmal-Default aussehen.

¹³⁷ Lexikalisch-morphologische Regeln lassen sich — genau wie lexikalische Regeln — als “Schema” mit nur einer Tochter modellieren. Die dazu nötigen Typen und Constraints werden hier jedoch nicht vorgestellt.

¹³⁸ Das ADJUNCT-Merkmal wird nicht verändert, so daß beliebig viele Adjunkte eingeführt werden können. In den meisten Fällen reichen jedoch zwei Adjunkte.

Um endlose Rekursionen zu vermeiden, kann für ADJUNCT eine natürliche Zahl als Wert spezifiziert werden, die durch das Schema heruntergezählt wird:

$$(328) \quad \left[\begin{array}{c} \text{M} \mid \text{MCAT} \mid \text{MHEAD} \mid \text{ADJUNCT} \mid \text{⑥} > 0 \\ \dots \end{array} \right] \Rightarrow \left[\begin{array}{c} \text{M} \mid \text{MCAT} \mid \text{MHEAD} \mid \text{ADJUNCT} \mid \text{⑥} - 1 \\ \dots \end{array} \right]$$

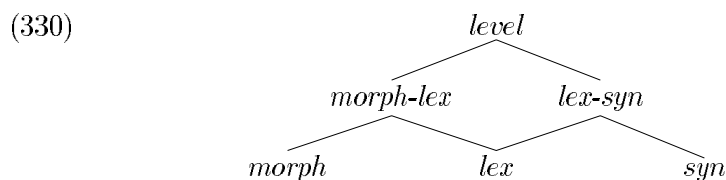


Clitics verändern nur die phonologische Repräsentation und die Stammform der Verben. *Suffixe* können die Syntax und / oder Semantik der Stämme modifizieren. Formen, an die ein Suffix angeschlossen wurde, erfordern jedoch weitere morphologische Modifikationen, bevor sie als morphologisch abgeschlossen in der Syntax verwendet werden können. *Flexive*, die ebenfalls die Syntax und Semantik ihrer Stämme modifizieren können, schließen den morphologischen Bildungsprozeß ab. Formen, die durch Flexive beendet werden, sind damit die einzigen, die in die Syntax übernommen werden.

Das LEVEL-Merkmal

Um Formen, die morphologisch abgeschlossen sind, von nicht abgeschlossenen Formen unterscheiden zu können, wird das Merkmal **LEVEL** eingeführt. **LEVEL** ersetzt das Boolesche Merkmal **LEX**, das in der Standard-HPSG lexikalische von phrasalen Zeichen unterscheidet.

LEVEL kann folgende Werte annehmen:



morph — für “Morphologie” — bezeichnet Formen, die noch nicht morphologisch abgeschlossen sind. *lex* — für “Lexikon” — bezeichnet Formen, die morphologisch abgeschlossen sind. Auf syntaktischer Ebene sind Elemente des Vollformlexikons und morphologisch abgeleitete Zeichen nicht unterscheidbar. Mit *lex* markierte Zeichen entsprechen damit atomaren syntaktischen Zeichen. *syn* bezeichnet komplexe syntaktische Zeichen bzw. Phrasen.

Um zu gewährleisten, daß in syntaktischen Schemata keine Zeichen verwendet werden können, die morphologisch noch nicht abgeschlossen sind, wird für das **LEVEL**-Merkmal der Töchter der Wert *lex-syn* spezifiziert. Die Elemente des Vollformlexikons werden mit [**LEVEL** *lex*] gekennzeichnet und die Mutterzeichen der Schemata durch [**LEVEL** *syn*].

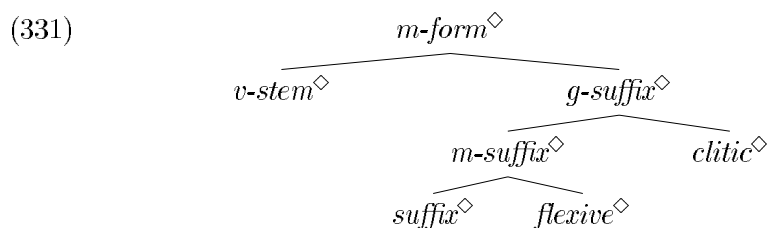
Der LEVEL-Wert komplexer morphologischer Zeichen wird durch das letzte Suffix bestimmt. LEVEL kann damit als morphologisches Head-Merkmal beschrieben werden. Da LEVEL für alle Flexive den Wert *lex* annimmt, und bei allen anderen atomaren morphologischen Zeichen durch *morph* beschrieben wird, kann sein Wert aus der Typenhierarchie atomarer morphologischer Zeichen ererbt werden.

5 Ein Fragment der japanischen Verbmorphologie

In den bisherigen Kapiteln wurden die morphologischen Mechanismen, die der vorliegenden Beschreibung zugrunde liegen, und ihr syntaktisch-semantischer Rahmen vorgestellt. Auf diesen Grundlagen aufbauend sollen die bisher nur beispielhaft eingeführten Verbformen in diesem Kapitel systematisch zu einem repräsentativen Fragment der japanischen Verbmorphologie ausgearbeitet werden.

Im ersten Abschnitt (5.1, S. 188) wird die Organisation des morphologischen Lexikons dargestellt. Da die Form der Einträge einem bestimmten Schema folgt, wird dieses in Unterabschnitt 5.1.1 “*Lexikalische Einträge*”, S. 188 beschrieben.

Im letzten Kapitel wurde die Hierarchie der morphologischen Zeichen vorgestellt:



Jeder Lexikoneintrag ist durch einen minimalen Typ dieser Hierarchie gekennzeichnet. Die Lexikoneinträge eines Typs haben viele Gemeinsamkeiten. Für jeden Typ lexikalischer Zeichen lassen sich daher die “typischen Eigenschaften” in Form eines Defaultzeichens angeben. In den einzelnen lexikalischen Einträgen müssen dann nur noch die *Abweichungen* von dem jeweiligen Prototyp — d.h. die *idiosynkratischen Merkmale* — dargestellt werden. Der Definition und Erläuterung der Defaultzeichen dient der zweite Abschnitt (5.1.2, S. 194) dieses Unterkapitels.

Nach diesen Bemerkungen zur Organisation des morphologischen Lexikons folgt das Lexikon selber (5.2, S. 205). Zuerst wird ein exemplarisches Verb- und Adjektivlexikon (5.2.1, S. 205) vorgestellt. Diesem folgt die Darstellung der Clitics (5.1.2, S. 198). Anschließend werden die Flexive (5.2.3, S. 226) und die Suffixverben und -adjektive (5.2.4, S. 248) aufgelistet.

5.1 Organisation des morphologischen Lexikons

Das Bestreben, Redundanzen zu vermeiden, könnte als der eigentliche Motor der Theoriebildung bezeichnet werden: Aus Listen korrekter Sätze oder Formen, der redundantesten Form der linguistischen “Beschreibung”, entstehen durch “Herauskürzen” der Redundanzen hierarchische Beschreibungssysteme, in denen jeder sprachliche Sachverhalt nur einmal und so allgemein wie möglich, aufgeführt wird.

Die gleiche Motivation liegt auch der Struktur des morphologischen Lexikons zugrunde. Durch die Organisation der morphologisch-lexikalischen Eigenschaften in einer Defaultvererbungshierarchie, müssen diese nicht für jeden Lexikoneintrag einzeln aufgelistet werden: Lexikalischen Einträge werden nach ihren Gemeinsamkeiten in Gruppen geordnet, denen ein lexikalischer Typ zugewiesen wird. Die “typischen”¹³⁹ Eigenschaften einer Gruppe von Zeichen können nun dem jeweiligen Typ zugeordnet werden. Von diesem können sie dann — soweit sie nicht durch abweichende Werte überschrieben werden — geerbt werden. Das gleiche Ordnungsprinzip kann auf die gewonnenen Typen rekursiv angewandt werden, so daß eine Hierarchie von Defaulttypen entsteht. Ein weiterer Teil der Eigenschaften lexikalischer Einträge wird durch morphologische Prinzipien, die wiederum Typen zugordnet sind, instantiiert.

Explizit angegeben werden müssen nur die *idiosynkratischen* Eigenschaften: Merkmalwerte also, die von den Defaultwerten ihrer lexikalischen Klasse abweichen bzw. von diesen nicht festgelegt werden.

Das Zusammenwirken von lexikalischer Defaultvererbungshierarchie, morphologischen Prinzipien und idiosynkratischen Eigenschaften soll im Folgenden anhand des Perfekt-Flexivs *a* beispielhaft dargestellt werden. Anschließend wird das vollständige System lexikalischer Typen eingeführt.

5.1.1 Lexikalische Einträge

Im morphologischen Lexikon werden nur die *idiosynkratischen Eigenschaften* der morphologischen Einträge angegeben. Das Perfekt-Flexives *a* beispielsweise ist durch vier Eigenschaften gekennzeichnet:

139 Die Kennzeichnung “typische Eigenschaft” unterscheidet *Defaultvererbungshierarchien* von *monotonen* Vererbungshierarchien: *monotone* Eigenschaften einer Gruppe von Zeichen müssen von *allen* ihren Elementen geteilt werden. *Typische* bzw. *Default*-Eigenschaften hingegen gelten nur für den *Großteil* der Elemente und können für einzelne Einträge mit dem Mittel der *Merkmalsüberschreibung* abweichend festgelegt werden.

“Typische” bzw. *Default*-Eigenschaften werden im Folgenden einfach als *Defaults* bezeichnet.

- (332) a. $flexiv^{\diamond}$
 b. PHON $\langle a \rangle$
 c. MAJ $v^{onbin} \longrightarrow v$
 d. CONT $\left[\begin{array}{l} past \\ \textcircled{1} \text{ cont}(\text{stem}) \end{array} \right]$

Statt vollständiger Pfade wird in der Regel nur das letzte Merkmal angegeben. Da sich Lexikoneinträge immer auf einen Defaulttyp beziehen (in diesem Falle den Defaulttyp $flexiv^{\diamond}$ aus (332a)), kann der entsprechende Pfad aus der Defaultvererbungshierarchie erschlossen werden.

Die Eigenschaften unseres Beispiels haben folgende Bedeutung:

Eigenschaft (332a): $flexiv^{\diamond}$

Morphologische Lexikoneinträge beziehen sich immer auf einen Defaulttyp. Dieser wird in der ersten Zeile des Lexikoneitragers festgelegt. In unserem Beispiel handelt es sich um den Defaulttyp $flexiv^{\diamond}$, der — setzt man die Defaultvererbung voraus — folgende Gestalt hat:

(333) Das Defaultzeichen für Flexive ($flexiv^{\diamond}$):

$$\left[\begin{array}{l} flexiv^{\diamond} \\ MORPHOLOGY | MCAT \left[\begin{array}{l} MHEAD \left[\begin{array}{ll} MOD & \diamond none \\ ADJUNCT & \diamond - \end{array} \right] \\ STEM \left[MCAT | STEM msat \right] \end{array} \right] \end{array} \right]$$

Eigenschaft (332b): PHON $\langle a \rangle$

Die zweite Eigenschaft des Präsens-Flexivs betrifft seine phonologische Realisierung — in diesem Falle $\langle a \rangle$. Es ergibt sich folgendes Zeichen:

$$(334) \quad \left[\text{PHON } \langle a \rangle \right]$$

Eigenschaft (332c): MAJ $v^{onbin} \longrightarrow v$

Suffixe subkategorisieren für Stämme bestimmter Kategorien und überführen die Gesamtform in eine neue Kategorie. Der Pfeil “ $\longrightarrow$ ” symbolisiert diese Transition.¹⁴⁰

Die Schreibweise “ $v^{onbin} \longrightarrow v$ ” bedeutet damit, daß

¹⁴⁰ Auf beiden Seiten des Pfeiles werden Abkürzungen verwendet: in diesem Fall v^{onbin} und v . Für ihre Definition vgl. (300), S. 168.

1. das Suffix für einen Stamm der Kategorie v^{onbin} subkategorisiert, und
2. das Suffix — und damit auch jede durch Suffigierung mit diesem gebildete komplexe Form — der Kategorie v angehört.

In AVM-Schreibweise entspricht $v^{onbin} \longrightarrow v$ damit:

$$(335) \quad \left[\text{MORPHOLOGY} \mid \text{MCAT} \left[\begin{array}{l} \text{MHEAD} \mid \text{MAJ } v \\ \text{STEM} \left[\text{MCAT} \mid \text{MHEAD} \mid \text{MAJ } v^{onbin} \right] \end{array} \right] \right]$$

Viele Suffixe werden in Abhängigkeit der Verbklasse phonologisch unterschiedlich realisiert. Das Präsens-Flexiv beispielsweise kann durch ru , u , uru oder i dargestellt werden. Im Minimum werden also vier Zeichen¹⁴¹ benötigt, die sich nur bezüglich PHON und MAJ unterscheiden. Um die übrigen Eigenschaften nicht ebenfalls wiederholen zu müssen, werden die PHON- und MAJ-Werte einander zugeordnet in disjunktiven Listen dargestellt:

$$(336) \quad \begin{array}{llll} \text{b. } i. & \text{MAJ } v_{v,ikuru}^{cons} & \longrightarrow v; & \text{PHON } \langle ru \rangle \\ & ii. & \text{MAJ } v_{c\backslash da,imasu}^{cons} & \longrightarrow v; & \text{PHON } \langle u \rangle \\ & iii. & \text{MAJ } v_{i,suru}^{cons} & \longrightarrow v; & \text{PHON } \langle uru \rangle \\ & iv. & \text{MAJ } adj^v & \longrightarrow adj; & \text{PHON } \langle i \rangle \end{array}$$

$$\text{Eigenschaft (332d): CONT } \left[\begin{array}{l} past \\ \textcircled{1} \text{ cont(stem)} \end{array} \right]$$

Ergänzt man den Pfad des CONT-Merkmals, ergibt sich:

$$(337) \quad \left[\text{MORPHOLOGY} \mid \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT } \left[\begin{array}{l} past \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \right]$$

Die Funktion **cont(stem)** bezieht sich auf das Merkmal CONT des Subkategorisierten Stammes¹⁴². Mit dem Mittel der Koindizierung ausgedrückt entspricht (337) der AVM:

¹⁴¹ Die Abkürzungen zur Angabe des MAJ-Wertes stehen meist für komplexe Disjunktionen von Verbkategorien. Da Disjunktionen innerhalb von Zeichen in unserem System zu mehreren Zeichen “ausmultipliziert” werden, ist die Anzahl der Präsens-Zeichen in Wirklichkeit sehr viel größer.

¹⁴² **cont(stem)** ist also wieder nur eine Abkürzung für eine AVM mit Koindizierung:

$$(340) \left[\begin{array}{c} \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \mid \text{STEM} \left[\text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \boxed{\mathbb{I}} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \quad \boxed{\mathbb{I}} \end{array} \right] \end{array} \right] \end{array} \right]$$

Das Defaultzeichen für Flexive $\textit{flexiv}^\diamond$, auf das der Eintrag des Perfekt-Flexives referiert, schlägt die Identität der CONT-Merkmale von Stamm und Flexiv vor. (332*d*) überschreibt also den Defaultwert für Flexive.

Zusammenfassung der lexikalischen Angaben:

Die Angaben aus dem Lexikoneintrag ((332), S. 189) für PHON-, MAJ- und CONT-Merkmal ergeben das folgende Zeichen:

$$(341) \left[\begin{array}{c} \text{PHON} \langle a \rangle \\ \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \mid \text{MAJ } v \\ \text{STEM} \left[\begin{array}{c} \text{MCAT} \mid \text{MHEAD} \mid \text{MAJ } v^{\textit{onbin}} \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \boxed{\mathbb{I}} \end{array} \right] \end{array} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \quad \boxed{\mathbb{I}} \end{array} \right] \end{array} \right] \end{array} \right]$$

Zusammen mit den Defaultwerten des Referenzzeichens für Flexive ($\textit{flexiv}^\diamond$) ergibt sich die folgende AVM:

(338)	Abkürzung:	abgekürzte AVM:
	$[\dots \mid \text{X cont}(\textit{stem})]$	$\left[\begin{array}{c} \dots \mid \text{X } \boxed{\mathbb{I}} \\ \dots \mid \text{STEM } [\dots \mid \text{CONT } \boxed{\mathbb{I}}] \end{array} \right]$

Sie setzt sich aus den folgenden zwei Abkürzungen zusammen:

(339)	Abkürzung:	abgekürzte AVM:
	$[\dots \mid \text{X stem}]$	$\left[\begin{array}{c} \dots \mid \text{X } \boxed{\mathbb{I}} \\ \dots \mid \text{STEM } \boxed{\mathbb{I}} \end{array} \right]$
	$[\dots \mid \text{X cont}(\textit{Y})]$	$\left[\begin{array}{c} \dots \mid \text{X } \boxed{\mathbb{I}} \\ \dots \mid \text{Y } [\dots \mid \text{CONT } \boxed{\mathbb{I}}] \end{array} \right]$

$$(342) \left[\begin{array}{c} \text{PHON} \langle a \rangle \\ \\ \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \left[\begin{array}{cc} \text{MAJ} & v \\ \text{MOD} & none \\ \text{ADJUNCT} & - \end{array} \right] \\ \text{STEM} \left[\begin{array}{c} \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \mid \text{MAJ } v^{onbin} \\ \text{STEM } msat \end{array} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT } \boxed{1} \end{array} \right] \end{array} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \left[\begin{array}{c} past \\ \textcircled{1} \quad \boxed{1} \end{array} \right] \end{array} \right] \end{array} \right]$$

Anwendung des MORPHOLOGY-Merkmal-Prinzips:

Alle morphologischen Zeichen (*msign*) sind Instanzen des MORPHOLOGY-Merkmal-Prinzips (MMP)¹⁴³, das hier als (343) nochmal aufgeführt wird:

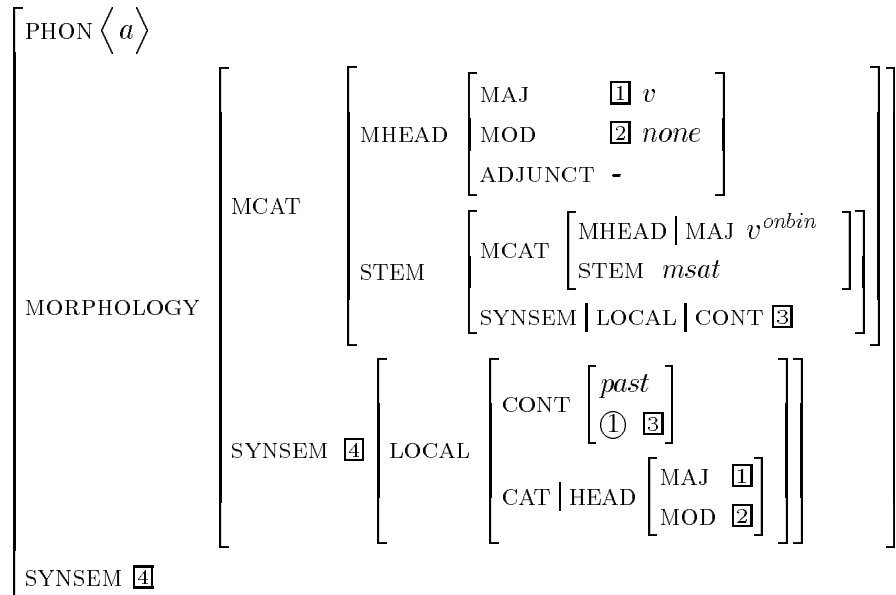
(343) *Das MORPHOLOGY-Merkmal-Prinzip:*

$$\left[\begin{array}{c} \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \mid \text{MHEAD} \left[\begin{array}{cc} \text{MAJ } \boxed{1} \\ \text{MOD } \boxed{2} \end{array} \right] \\ \text{SYNSEM } \boxed{3} \left[\text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \left[\begin{array}{cc} \text{MAJ } \boxed{1} \\ \text{MOD } \boxed{2} \end{array} \right] \right] \end{array} \right] \\ \text{SYNSEM } \boxed{3} \end{array} \right]$$

Durch Anwenden des MMPs auf die bisher abgeleitete Form des Perfekt-Flexivs (341) ergibt sich die folgende AVM:

¹⁴³ vgl. 4.2.4, S. 172.

(344) Nach Anwendung des MORPHOLOGY-Merkmal-Prinzips:



Anwendung des STEM-Merkmal-Defaults:

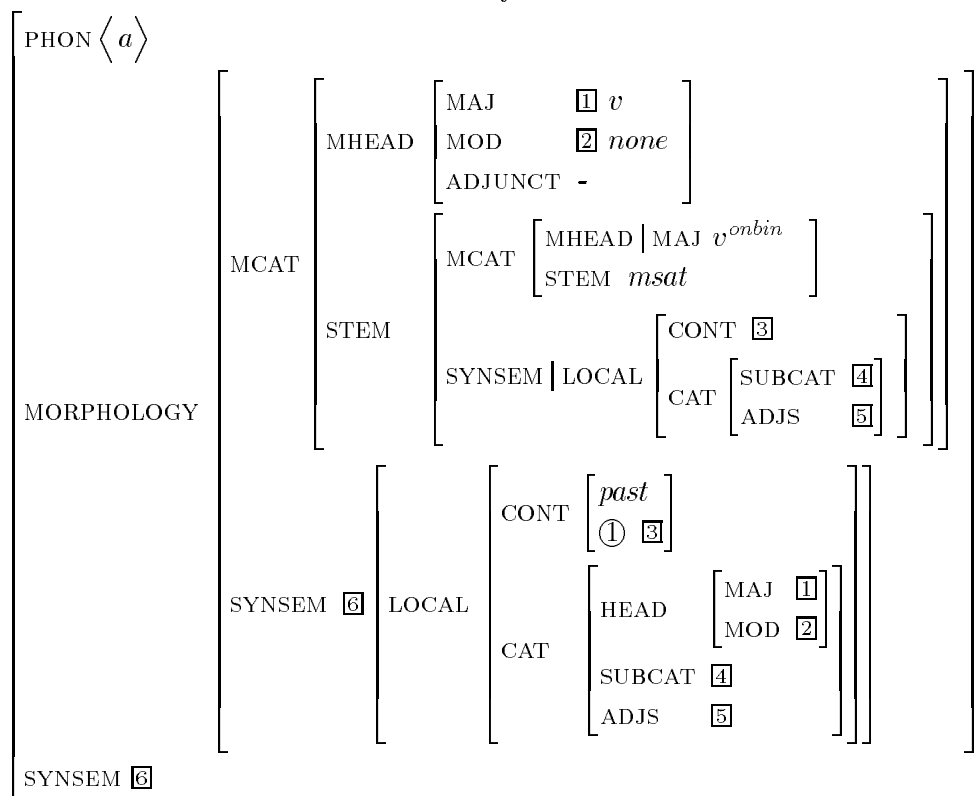
Schließlich werden alle Suffixe (*g-suffix*) durch das STEM-Merkmal-Default (4.2.5, S. 173) subsumiert. Diesem zufolge werden alle SYNSEM-Merkmale, die nicht durch den Lexikoneintrag und das MORPHOLOGY-Merkmal-Prinzip bestimmt werden, vom Stamm übernommen.

Im Falle des Perfekt-Flexivs sind die *cat*-Merkmale SUBCAT und ADJS noch unbestimmt. Sie werden daher mit den entsprechenden Merkmalen des Stammes koindiziert.

Lexikoneintrag, Defaultzeichen, MORPHOLOGY-Merkmal-Prinzip und STEM-Merkmal-Default ergeben das vollständige Zeichen. Für das Perfekt-Flexiv

erhalten wir damit:

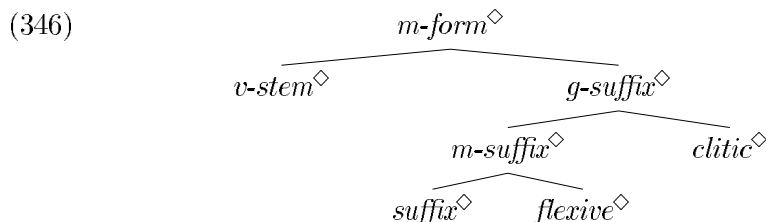
(345) *Nach Anwenden des STEM-Merkmal-Defaults:*



5.1.2 Lexikalische Defaults

5.1.2.1 Lexikalische Defaultvererbungshierarchie

Alle lexikalischen Zeichen beziehen sich auf ein Defaultzeichen. Die Defaultzeichen ergeben sich aus der Defaulthierarchie lexikalischer Typen und entsprechenden Defaultconstraints, wie in Abschnitt 1.7 (S. 58) beschrieben. Die Defaulthierarchie lexikalischer Typen ist mit der *m-form*-Hierarchie identisch und wird an dieser Stelle als (346) noch einmal aufgeführt:



Defaultconstraints und Defaultzeichen werden im nächsten Abschnitt vorgestellt.

5.1.2.2 Morphologische Lexikoneinträge ($m\text{-form}^\diamond$)

Der allgemeinste Typ für morphologisch-lexikalische Zeichen ist $m\text{-form}$. Seine Untertypen — Verbstämme und Suffixe im engeren Sinne — unterscheiden sich so stark, daß für $m\text{-form}^\diamond$ nur eine Defaultspezifikation sinnvoll ist. Diese hat die Form $[\text{MOD} \diamond \phi]$. Das Symbol ϕ dient nur der besseren Lesbarkeit der TFS-Spezifikation und kann durch einen beliebigen minimalen Wert ersetzt werden, der für das betreffende Merkmal zulässig ist: ϕ wird immer dann verwendet, wenn ein Merkmal für eine Merkmalstruktur zwar definiert ist, der entsprechende Merkmalwert in der aktuellen Struktur jedoch keine Bedeutung hat. Ein nicht-minimaler Wert an dieser Stelle würde zu Lösungen führen, die sich nur bezüglich dieses Merkmals unterscheiden und damit redundant sind.

Das Merkmal MOD ist ein MHEAD-Merkmal, das nur für die Syntax von Bedeutung ist. Es ist jedoch für alle morphologischen Formen (Strukturen des Typs $m\text{-form}$) definiert. Da nur die MHEAD-Merkmale von Flexiven in der Syntax sichtbar sind, hat es bei allen anderen morphologischen Zeichen keine Funktion und sollte deshalb durch einen beliebigen minimalen Wert spezifiziert werden. Diese Belegung mit einem beliebigen minimalen Wert wird durch ϕ signalisiert.

Für den Typ $m\text{-form}^\diamond$ erhalten wir damit:

(347) *Constraint für $m\text{-form}^\diamond$:*

$$\left[m\text{-form}^\diamond \mid \text{MORPHOLOGY} \mid \text{MCAT} \mid \text{MHEAD} \mid \text{MOD} \diamond \phi \right]$$

Um den Defaultstatus der Defaultwerte optisch hervorzuheben, kennzeichnen wir diese durch ein vorgestelltes $\diamond$.

5.1.2.3 Verbstämme ($v\text{-stem}^\diamond$)

Verbstämme können durch lexikalisch-morphologische Regeln zur Einführung von Adjunkten (Adverben) modifiziert werden. Die Anwendbarkeit dieser Regel wird durch die Spezifikation $[\text{ADJUNCT} +]$ ausgedrückt. Diese Spezifikation wird daher als Default am Typ $v\text{-stem}^\diamond$ angetragen.

Das zweite Default betrifft das Merkmal für morphologische Subkategorisierung STEM: Verbstämme bilden den Ausgangspunkt der agglutinierenden Verbmodifikation. Sie subkategorisieren daher selber nicht für Stämme, und es gilt $[\text{STEM } msat]$.

Das Subkategorisierungsmerkmal für Adjunkte ADJS bei komplexen morphologischen Formen wird prinzipiell vom Stamm geerbt. Es entspricht bei fast

allen lexikalischen Einträgen der leeren Menge und wird nur durch die oben genannte lexikalisch-morphologische Regel zur Adjunkteinführung verändert. Nur bei komplexen, unregelmäßigen Formen wie *nai*, der Negation von *suru*, kann es auch bei lexikalischen Einträgen Elemente enthalten.

Für $v\text{-stem}^\diamond$ ergibt sich damit das folgende Defaultzeichen:

(348) *Constraint für $v\text{-stem}^\diamond$:*

$$\left[\begin{array}{c} v\text{-stem}^\diamond \\ \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \mid \text{ADJUNCT} \mid \diamond + \\ \text{STEM} \mid \diamond msat \end{array} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{ADJS} \mid \diamond \{ \} \end{array} \right] \end{array} \right]$$

Das ererbte Defaultzeichen hat folgende Form:

(349) *Ererbter Constraint für $v\text{-stem}^\diamond$:*

$$\left[\begin{array}{c} v\text{-stem}^\diamond \\ \text{MORPHOLOGY} \left[\begin{array}{c} \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \left[\begin{array}{c} \text{MOD} \mid \diamond \phi \\ \text{ADJUNCT} \mid \diamond + \end{array} \right] \\ \text{STEM} \mid \diamond msat \end{array} \right] \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{ADJS} \mid \diamond \{ \} \end{array} \right] \end{array} \right]$$

Aus dem ererbten Defaultzeichen, den idiosynkratischen Merkmalen des Lexikoneintrags, dem MORPHOLOGY-Merkmal-Prinzip und dem STEM-Merkmal-Default ergibt sich das vollständige Zeichen.

Dem Lexikoneintrag für das Verb 飲む (*no.m-u*; trinken) beispielsweise

(350) a. $v\text{-stem}^\diamond$

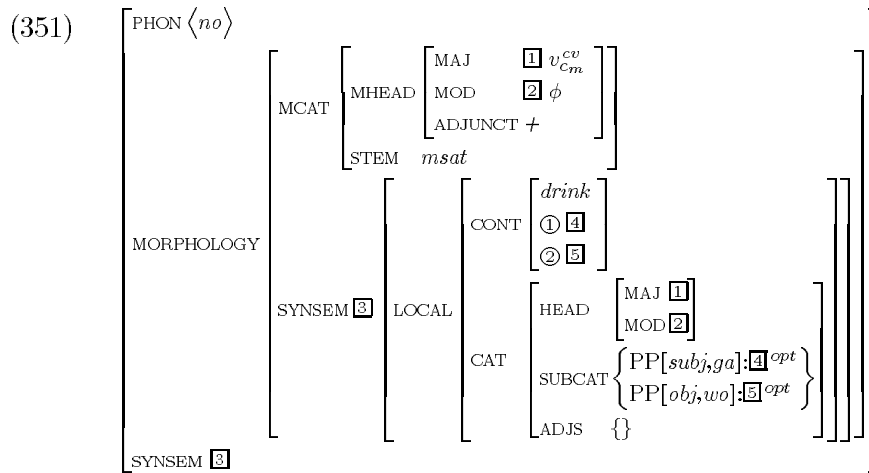
b. PHON $\langle no \rangle$

c. MAJ v_{cm}^{cv}

d. SUBCAT $\{ \text{PP}[\text{subj}, \text{ga}]^{opt}, \text{PP}[\text{obj}, \text{wo}]^{opt} \}$

e. CONT $\left[\begin{array}{c} \text{drink} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$

entspricht das folgende ausformulierte Zeichen:



5.1.2.4 Suffixe im weiteren Sinne ($g\text{-suffix}^\diamond$)

Suffixverben, Flexive und Clitics — alle drei durch $g\text{-suffix}^\diamond$ subsumiert — haben zwei weitere gemeinsame Defaulteigenschaften:

1. Im Allgemeinen können sie nicht durch die lexikalisch-morphologische Regel zur Adjunkteinführung modifiziert werden. Sie werden deshalb durch das Default $[\text{ADJUNCT } \diamond -]$ spezifiziert. Einzelne Suffixe, wie beispielsweise das Kausativ-Suffix, weichen allerdings von dieser Eigenschaft ab und überschreiben das Default.
2. Suffixverben, Flexive und Clitics können sich prinzipiell nur an Formen anschließen, die selber schon gesättigt sind. Da diese Eigenschaft ohne Ausnahme für alle $g\text{-suffix}^\diamond$ -Zeichen gilt, wird sie nicht als Default, sondern als monotone Eigenschaft beschrieben.

(352) *Constraint für $g\text{-suffix}^\diamond$:*

$$\left[\begin{array}{c} g\text{-suffix}^\diamond \\ \\ \text{MORPHOLOGY} \mid \text{MCAT} \end{array} \left[\begin{array}{c} \text{MHEAD} \mid \text{ADJUNCT } \diamond - \\ \text{STEM} \left[\text{MCAT} \mid \text{STEM } msat \right] \end{array} \right] \right]$$

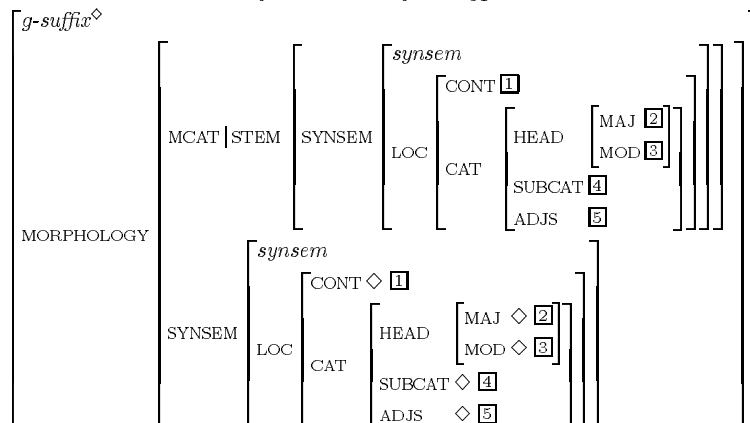
Der ererbte Constraint hat folgende Form:

(353) *ererbter Constraint für $g\text{-suffix}^\diamond$ (ohne STEM-Merkmal-Default):*

$$\left[\begin{array}{c} g\text{-suffix}^\diamond \\ \\ \text{MORPHOLOGY} \mid \text{MCAT} \end{array} \left[\begin{array}{c} \text{MHEAD} \left[\begin{array}{c} \text{MOD} \quad \diamond \, \phi \\ \text{ADJUNCT} \quad \diamond \, - \end{array} \right] \\ \text{STEM} \left[\text{MCAT} \mid \text{STEM } msat \right] \end{array} \right] \right]$$

Auch das STEM-Merkmal-Default, hier als (354) wiederholt, ist als Default zum Typ $g\text{-suffix}^\diamond$ definiert. Der Übersichtlichkeit halber wird es bei der Angabe der Defaultzeichen jedoch weggelassen.

(354) *Das STEM-Merkmal-Default als Defaulttypenconstraint:*



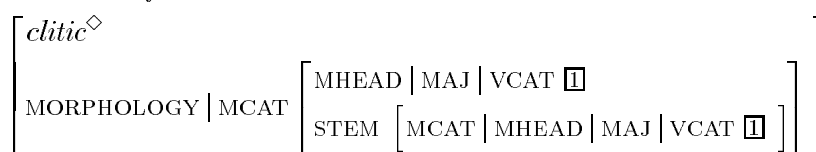
5.1.2.5 Clitics ($\text{clitic}^\diamond$)

Clitics dienen der Ableitung der Verbstämme. Sie modifizieren nur zwei Merkmale der Formen, an die sie sich anschließen:

1. Die phonologische Gestalt — also das Merkmal PHON;
2. Den Verbstamm und somit das Merkmal VSTEM.

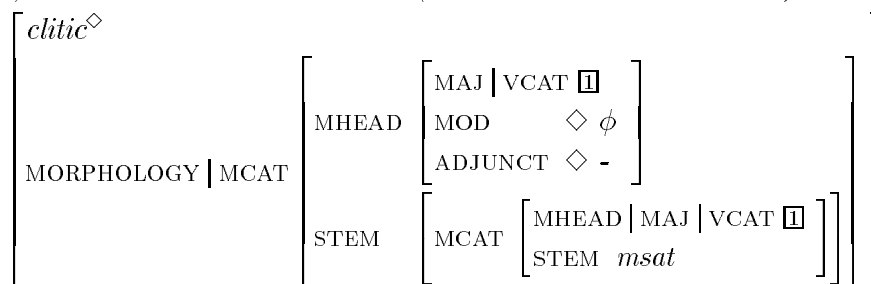
Die Kategorie der Verben, dargestellt durch das Merkmal VCAT, bleibt unverändert. Das Defaultzeichen hat daher die folgende Form:

(355) *Constraint für $\text{clitic}^\diamond$:*



Das ererbte Defaultzeichen für Clitics sieht folgendermaßen aus:

(356) *ererbter Constraint für $\text{clitic}^\diamond$ (ohne STEM-Merkmal-Default):*



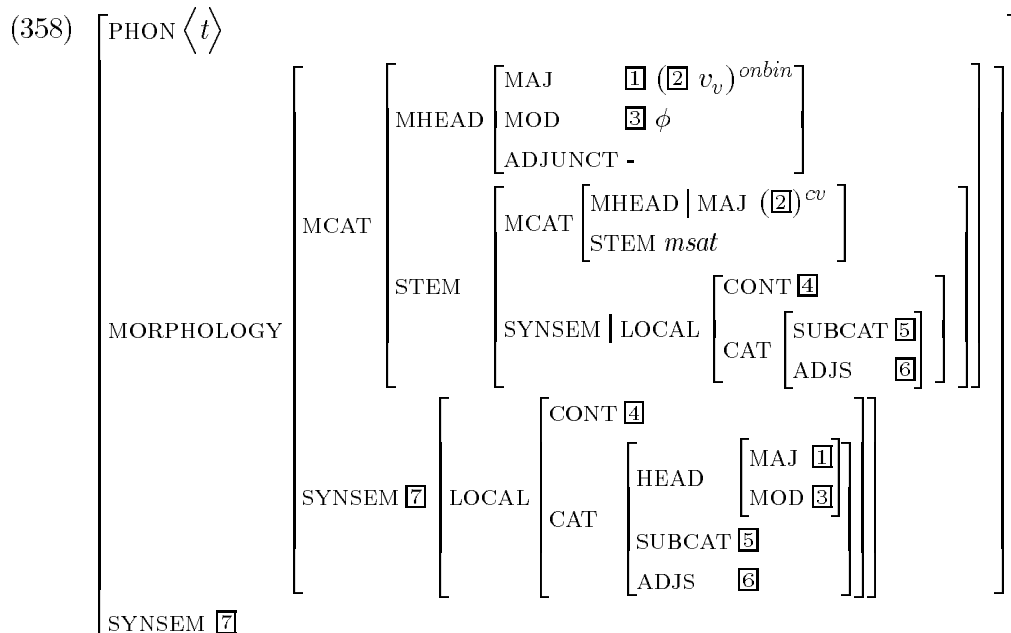
Das Clitic *t* zur Bildung des onbin-Stammes der vokalischen Verben soll den Zusammenhang zwischen lexikalischer Spezifikation und vollständigem Zeichen bei Clitics veranschaulichen:

*Lexikalische Spezifikation:*¹⁴⁴

(357) *clitic*[◇]

PHON $\langle t \rangle$
MAJ $v_v^{cv} \rightarrow v_v^{onbin}$

Das abgeleitete Zeichen:



Die Form der Lexikoneinträge ist für alle Clitics gleich und umfasst die beiden Merkmale PHON und MAJ. In Abschnitt 5.2.2 “Clitics”, S. 222, der vollständigen Aufzählung der in dieser Arbeit verwendeten Clitics, werden sie daher einfach durch Tabellen beschrieben.

Die Tabelle

(359)

<i>Phon</i>	<i>Maj</i>
<i>k</i>	$v_{c_K}^{cv} \rightarrow v_{c_K}^{cons}$
<i>s</i>	$v_{c_s}^{cv} \rightarrow v_{c_s}^{cons}$
	...

z.B. beschreibt ausschnittsweise die Bildung des konsonantischen Stammes der Verben.

¹⁴⁴ Der besseren Lesbarkeit halber wird die Kategorie, die durch Clitics nicht verändert wird, auf beiden Seiten des Pfeils wiederholt.

5.1.2.6 Suffixe im engeren Sinne (*m-suffix*[◇])

Im Gegensatz zu den Clitics, die nur der Ableitung der Stammformen dienen, beschreiben Suffixe des Typs *m-suffix*[◇] (das *m* steht für *morph*) den lexikalischen “Kern”¹⁴⁵. Alle typischen Eigenschaften von m-Suffixen werden schon vom Default zum Typ *g-suffix*[◇] beschrieben. Das Defaultzeichen für m-Suffixe hat daher die folgende, einfache Gestalt:

$$(360) \text{ Constraint für } m\text{-suffix}^\diamond: \left[m\text{-suffix}^\diamond \right]$$

Der ererbte Constraint entspricht damit dem ererbten Constraint des Typs *g-suffix*[◇]:

(361) *ererbter Constraint für m-suffix*[◇] (ohne STEM-Merkmal-Default):

$$\left[\begin{array}{c} m\text{-suffix}^\diamond \\ \text{MORPHOLOGY} \mid \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \left[\begin{array}{cc} \text{MOD} & \diamond \phi \\ \text{ADJUNCT} & \diamond - \end{array} \right] \\ \text{STEM} \left[\text{MCAT} \mid \text{STEM } msat \right] \end{array} \right] \end{array} \right]$$

5.1.2.7 Flexive (*flexive*[◇])

MOD ist als Merkmal, dessen Wert nur für die Syntax interessant ist, alleine für Flexive sinnvoll: Als MHEAD-Merkmal wird sein Wert bei morphologisch abgeleiteten syntaktischen Zeichen durch das abschließende Flexiv bestimmt. Bei allen anderen morphologischen Zeichen spielt sein Wert keine Rolle.

Als Standardform der Verben wird [MOD *none*] gewählt, so daß sich das folgende Defaultzeichen für Flexive ergibt:

$$(362) \text{ Constraint für flexiv}^\diamond: \left[\begin{array}{c} flexiv^\diamond \\ \text{MORPHOLOGY} \mid \text{MCAT} \mid \text{MHEAD} \mid \text{MOD} \diamond none \end{array} \right]$$

Flexive, die zur Bildung der japanischen Variante der Relativsätze dienen, wie beispielsweise das *i* in 高い本 (*taka.i hon*; ein teures Buch) überschreiben dieses Default.

Das ererbte Defaultzeichen für Flexive hat folgende Gestalt:

¹⁴⁵ Morphe lassen sich also durch den regulären Ausdruck *Kern* ◦ *Clitic** beschreiben.

(363) *ererbter Constraint für Flexive flexiv[◇] (ohne STEM-Merkmal-Default):*

$$\left[\begin{array}{c} flexiv^{\diamond} \\ \\ MORPHOLOGY \mid MCAT \left[\begin{array}{c} MHEAD \left[\begin{array}{cc} MOD & \diamond none \\ ADJUNCT & \diamond - \end{array} \right] \\ STEM \left[MCAT \mid STEM msat \right] \end{array} \right] \end{array} \right]$$

Als Beispiel für ein ausformuliertes Flexiv wird die Satzschlußform des Perfekt-Flexives dargestellt, deren Ableitung schon in Abschnitt 5.1.1 (S. 188) ausführlich beschrieben wurde. Der Lexikoneintrag dieser Form sieht folgendermaßen aus:

- (364) a. *flexiv[◇]*
 b. PHON $\langle a \rangle$
 c. MAJ $v^{onbin} \rightarrow v$
 d. CONT $\left[\begin{array}{c} past \\ \textcircled{1} \text{ cont(stem)} \end{array} \right]$

Das vollständige Zeichen, das diesem Lexikoneintrag entspricht, hat die Form:

$$(365) \left[\begin{array}{c} PHON \langle a \rangle \\ \\ MORPHOLOGY \left[\begin{array}{c} MCAT \left[\begin{array}{c} MHEAD \left[\begin{array}{cc} MAJ \textcircled{1} v \\ MOD \textcircled{2} none \\ ADJUNCT - \end{array} \right] \\ MCAT \left[\begin{array}{c} MHEAD \mid MAJ v^{onbin} \\ STEM msat \end{array} \right] \\ STEM \left[\begin{array}{c} SYNSEM \mid LOCAL \left[\begin{array}{c} CONT \textcircled{3} \\ CAT \left[\begin{array}{cc} SUBCAT \textcircled{4} \\ ADJS \textcircled{5} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \\ \\ SYNSEM \textcircled{6} \left[\begin{array}{c} LOCAL \left[\begin{array}{c} CONT \left[\begin{array}{c} past \\ \textcircled{1} \textcircled{3} \end{array} \right] \\ CAT \left[\begin{array}{c} HEAD \left[\begin{array}{cc} MAJ \textcircled{1} \\ MOD \textcircled{2} \end{array} \right] \\ SUBCAT \textcircled{4} \\ ADJS \textcircled{5} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

5.1.2.8 Suffixverben und -adjektive ($\text{suffix}^\diamond$)

Das Defaultzeichen für Suffixverben ist mit dem seines Supertypen identisch und kann daher in der folgenden Weise spezifiziert werden:

(366) *Constraint für $\text{suffix}^\diamond$:*

$$\left[\text{suffix}^\diamond \right]$$

Das ererbte Defaultzeichen ist damit — von seinem Defaulttyp abgesehen — mit dem ererbten Defaultzeichen von $m\text{-suffix}^\diamond$ identisch:

(367) *ererbter Constraint für $\text{suffix}^\diamond$ (ohne STEM-Merkmal-Default):*

$$\left[\begin{array}{c} \text{suffix}^\diamond \\ \text{MORPHOLOGY} \mid \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \left[\begin{array}{cc} \text{MOD} & \diamond \phi \\ \text{ADJUNCT} & \diamond - \end{array} \right] \\ \text{STEM} \left[\text{MCAT} \mid \text{STEM} \text{ } msat \right] \end{array} \right] \end{array} \right]$$

Das Suffix zur Bildung der Honorativform *-mas.-u* soll zur Veranschaulichung des Zusammenhangs von Lexikoneintrag und vollständigem Zeichen bei den Suffixverben dienen:

Der Lexikoneintrag sieht folgendermaßen aus:

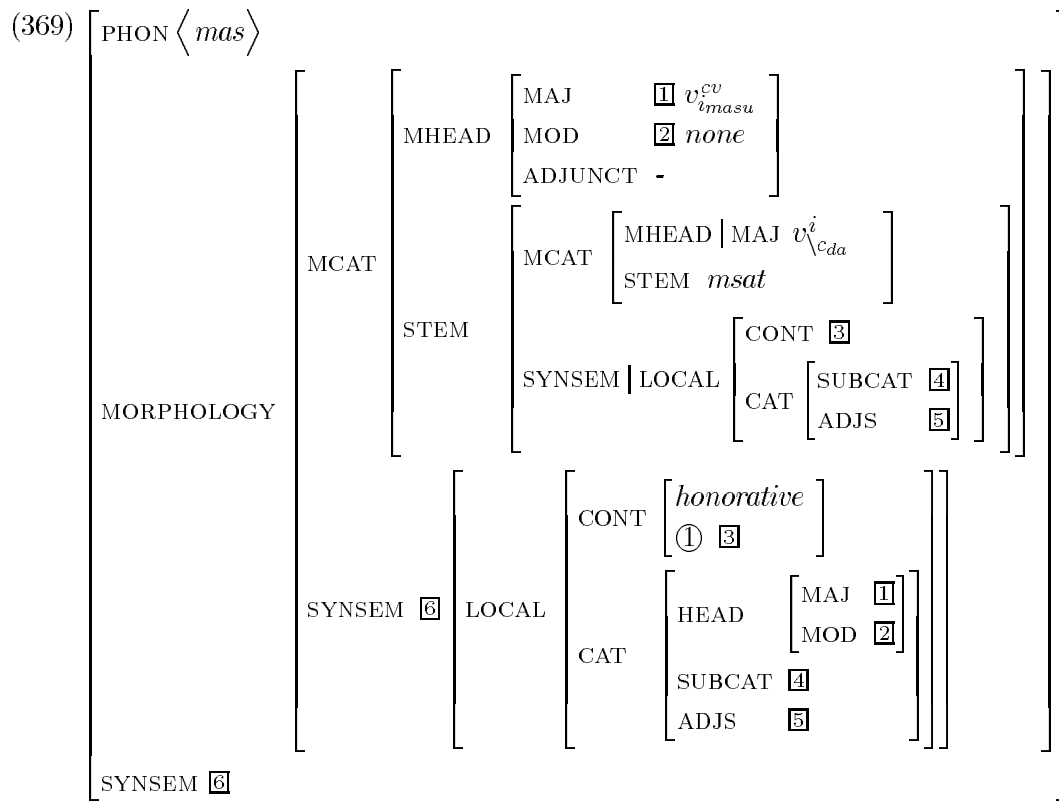
Morphologie:

(368) a. $\text{suffix}^\diamond$

b. CONT $\left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \text{ cont}(\text{stem}) \end{array} \right]$

c. MAJ $v_{cda}^i \longrightarrow v_{imasu}^{cv}$; PHON $\langle mas \rangle$

Ihm entspricht das folgende vollspezifizierte Zeichen:



5.2 Das morphologische Lexikon

Nachdem die Organisation des morphologischen Lexikons und die Ableitung morphologischer Zeichen dargestellt wurde, folgen nun die Lexikoneinträge selber. Zuerst wird ein fragmentarisches Lexikon von Verb- und Adjektivstämmen beschrieben; diesem folgt die Auflistung der Clitics zur Ableitung der Stammformen; eine repräsentative Auflistung der japanischen Flexive, Suffixverben und Suffixadjektive schließt sich daran an.

5.2.1 Verbstämme

Für jede Verbkategorie wird der Lexikoneintrag mindestens eines Verbes als Beispiel dargestellt.¹⁴⁶ Das Lexikon kann damit leicht durch neue Einträge erweitert werden. Die Beschreibung der Einträge folgt einem einheitlichen Schema: Zuerst wird das Verb im Kontext eines Beispielsatzes vorgestellt. Anschließend folgt die formale Beschreibung seiner idiosynkratischen Eigenschaften, d.h. sein Lexikoneintrag. Beispielsatz und Lexikoneintrag werden durch Erklärungen weiterer Besonderheiten ergänzt.

5.2.1.1 v_v : いる – *i-ru* – sein

Beispiel:

(370) a. 花子はあそこにあります。

<i>Hanako-ha</i>	<i>asoko-ni</i>	<i>i-ru</i>	/	<i>i-mas.-u</i>
Hanako-TOP	dort drüben-LOC	sein-FLEX	/	sein-HON-FLEX
Hanaok ist dort drüben.				

b. 彼女は寝ている / います。

<i>kanozjo-ha</i>	<i>ne.t-e</i>	<i>i-ru</i>	/	<i>i-mas.-u</i>
Sie-TOP	schlafen-PART	KONT-FLEX	/	KONT-HON-FLEX
Sie schläft gerade.				

Morphologie:

¹⁴⁶ Die Abkürzungen, die bei der Quellenangaben der Beispielsätze Verwendung finden, werden im Literaturverzeichnis erklärt.

- (371) a. $v\text{-stem}^\diamond$
 b. PHON $\langle i \rangle$
 c. MAJ v_v^{comb}
 d. i. SUBCAT $\left\{ \begin{array}{l} \text{PP}[\text{subj}, \text{ga}]^{opt} \\ \text{PP}[\text{obj}, \text{ni}]^{opt} \end{array} \right\};$ CONT $\left[\begin{array}{l} be \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$
 ii. SUBCAT $\{ \};$ CONT *continuation*

5.2.1.2 v_v : 見る – *mi-ru* – sehen

Beispiel:

(372) きのお面白くない映画を見ました。

kinou omosiro-ku na-i eiga-wo mi-mas.it-a
 gestern interessant-ADV NEG-FLEX Film-ACC sehen-HON-PAST
 Gestern habe ich einen uninteressanten Film gesehen.

Morphologie:

- (373) a. $v\text{-stem}^\diamond$
 b. PHON $\langle mi \rangle$
 c. MAJ v_v^{comb}
 d. SUBCAT $\left\{ \text{PP}[\text{subj}, \text{ga}]^{opt}, \text{PP}[\text{obj}, \text{wo}]^{opt} \right\}$
 e. CONT $\left[\begin{array}{l} see \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$

5.2.1.3 v_v : 寝る – *ne-ru* – schlafen

Beispiel:

(374) 彼女は寝ている / います。

kanozyo-ha ne.t-e i-ru / i-mas.-u
 Sie-TOP schlafen-PART KONT-FLEX / KONT-HON-FLEX
 Sie schläft gerade.

Morphologie:

- (375) a. $v\text{-stem}^\diamond$
 b. PHON $\langle ne \rangle$
 c. MAJ v_v^{comb}
 d. SUBCAT $\{PP[subj,ga]^{opt}\}$
 e. CONT $\begin{bmatrix} sleep \\ \textcircled{1} \text{ cont(subj)} \end{bmatrix}$

5.2.1.4 v_v : 食べる – *tabe-ru* – essen

Beispiel:

- (376) 僕は今ピザを/が食べたい / 食べたいです¹⁴⁷

boku-ha ima pizza-wo/-ga tabe-ta-i / tabe-ta-i
 ich-TOP jetzt Pizza-ACC/-NOM essen-VOLU-FLEX / essen-VOLU-FLEX
des.-u
 HON-FLEX
 Ich möchte jetzt Pizza essen.

Morphologie:

- (377) a. $v\text{-stem}^\diamond$
 b. PHON $\langle tabe \rangle$
 c. MAJ v_v^{comb}
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj,wo]^{opt}\}$
 e. CONT $\begin{bmatrix} eat \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{bmatrix}$

5.2.1.5 v_{c_k} : 書く – *ka.k-u* – schreiben

Beispiel:

- (378) 直美が手紙を毎日書く約束をした。¹⁴⁸

naomi-ga tegami-wo mainiti ka.k-u yakusoku-wo
 Naomi-NOM Brief-ACC täglich schreiben-FLEX Versprechen-ACC
s.it-a
 machen-PAST
 Naomi versprach, jeden Tag einen Brief zu schreiben.

¹⁴⁷ DBJG, S. 442, KS(B).

¹⁴⁸ JPSG, S. 53, (3.63).

Morphologie:

- (379) a. *v-stem*[◇]
 b. PHON $\langle ka \rangle$
 c. MAJ $v_{c_k}^{cv}$
 d. SUBCAT $\left\{ \text{PP}[\text{subj}, \text{ga}]^{opt}, \text{PP}[\text{obj}, \text{ni}]^{opt} \text{PP}[\text{obj2}, \text{wo}]^{opt} \right\}$
 e. CONT $\left[\begin{array}{l} \text{drink} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj2}) \\ \textcircled{3} \text{ cont}(\text{obj}) \end{array} \right]$

5.2.1.6 $v_{c_{iku}}$: 行く – *i.k-u* – gehen, fahren

Beispiel:

- (380) 映画に行きましょう。 ¹⁴⁹

eiga-ni i.k.i-mas.-you
 Film-DIR gehen-HON-VOLI
 Laßt und ins Kino gehen!

Morphologie:

- (381) a. *v-stem*[◇]
 b. PHON $\langle i \rangle$
 c. MAJ $v_{c_{iku}}^{cv}$
 d. SUBCAT $\left\{ \text{PP}[\text{subj}, \text{ga}]^{opt}, \text{PP}[\text{obj}, \text{ni} \vee \text{he}]^{opt} \right\}$
 e. CONT $\left[\begin{array}{l} \text{go-to} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$

5.2.1.7 v_{c_g} : 泳ぐ – *oyo.g-u* – schwimmen

Beispiel:

- (382) 彼は海で泳ぐ。

kare-ha umi-de oyo.g-u
 er-TOP Meer-LOC schwimmen-FLEX
 Er schwimmt im Meer.

149 DBJG, S. 240, KS(B).

Morphologie:

- (383) a. *v-stem*[◇]
 b. PHON $\langle oyo \rangle$
 c. MAJ v_{cg}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}\}$
 e. CONT $\begin{bmatrix} swim \\ \textcircled{1} \text{ cont(subj)} \end{bmatrix}$

5.2.1.8 v_{cs} : 話す – *hana.s-u* – sprechen

Beispiel:

- (384) 私は英語を話します。 ¹⁵⁰

watasi-ha eigo-wo hana.s.i-mas.-u
 ich-TOP Englisch-ACC sprechen-HON-FLEX
 Ich spreche Englisch.

Morphologie:

- (385) a. *v-stem*[◇]
 b. PHON $\langle hana \rangle$
 c. MAJ v_{cs}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj2,ni]^{opt}, PP[obj,wo]^{opt}\}$
 e. CONT $\begin{bmatrix} speak \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj2)} \\ \textcircled{3} \text{ cont(obj)} \end{bmatrix}$

5.2.1.9 v_{ct} : 待つ – *ma.t-u* – warten

Beispiel:

- (386) 電車の到着を待つ。

densya-no toutyaku-wo ma.t-u
 Zug-gen Ankunft-ACC warten-FLEX
 [Er] wartet auf die Ankunft des Zuges.

150 DBJG, S. 371, notes (1)a.

Morphologie:

- (387) a. *v-stem*[◇]
 b. PHON $\langle ma \rangle$
 c. MAJ v_{ct}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj,wo]^{opt}\}$
 e. CONT $\begin{bmatrix} wait-for \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{bmatrix}$

5.2.1.10 v_{cn} : 死ぬ – *si.n-u* – sterben

Beispiel:

- (388) 人間は死ぬべきだ。

niñgen-ha si.n-u beki da
 Mensch-TOP sterben-FLEX Notwendigkeit COP-FLEX
 Der Mensch ist sterblich.

Morphologie:

- (389) a. *v-stem*[◇]
 b. PHON $\langle si \rangle$
 c. MAJ v_{cn}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}\}$
 e. CONT $\begin{bmatrix} die \\ \textcircled{1} \text{ cont(subj)} \end{bmatrix}$

5.2.1.11 v_{cb} : 呼ぶ – *yo.b-u* – rufen

Beispiel:

- (390) 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
 Tarō-NOM Zirō-ACC rufen-PAST
 Tarō rief Zirō.

Morphologie:

- (391) a. $v\text{-stem}^\diamond$
 b. PHON $\langle yo \rangle$
 c. MAJ $v_{c_b}^{cv}$
 d. SUBCAT $\{ \text{PP}[\text{subj}, ga]^{opt}, \text{PP}[\text{obj}, wo]^{opt} \}$
 e. CONT $\left[\begin{array}{l} call \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$

5.2.1.12 v_{c_m} : 飲む – *no.m-u* – trinken

Beispiel:

- (392) 健はコーヒーを飲む。

keñ-ha kouhii-wo no.m-u
 Ken-TOP Kaffee-ACC trinken-FLEX
 Ken trinkt Kaffee.

Morphologie:

- (393) a. $v\text{-stem}^\diamond$
 b. PHON $\langle no \rangle$
 c. MAJ $v_{c_m}^{cv}$
 d. SUBCAT $\{ \text{PP}[\text{subj}, ga]^{opt}, \text{PP}[\text{obj}, wo]^{opt} \}$
 e. CONT $\left[\begin{array}{l} drink \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{array} \right]$

5.2.1.13 v_{c_r} : 降る – *fu.r-u* – $\langle \text{Regen} \rangle$ fällt

Beispiel:

- (394) 雨がすごくふっている。 ¹⁵¹

ame-ga sugo-ku hu.tt-e i-ru
 Regen-NOM furchtbar stark-ADV fallen-PART KONT-FLEX
 Es regnet in Strömen.

Morphologie:

151 JM, S. 223, 32-2.1.

- (395) a. $v\text{-stem}^\diamond$
 b. PHON $\langle fu \rangle$
 c. MAJ v_{cr}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}\}$
 e. CONT $\begin{bmatrix} fall \\ \textcircled{1} \text{ cont}(\text{subj}) \end{bmatrix}$

5.2.1.14 v_{caru} : ある – *a.r-u* – sein

Beispiel:

- (396) よかった。なくなったゆびわがこんな所にあったわ。 ¹⁵²

yo-ka.tt-a. na-ku na.tt-a yubiwa-ga koñna
 gut-φ-PAST nicht sein-ADV werden-PAST Fingerring-NOM solch ein
 tokoro-ni a.tt-a
 Platz-LOC sein-PAST
 Welch ein Glück! Der verlorengegangene Ring hat an einem solchen Platz
 gesteckt.

Morphologie:

- (397) a. $v\text{-stem}^\diamond$
 b. PHON $\langle a \rangle$
 c. MAJ v_{caru}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj,ni]^{opt}\}$
 e. CONT $\begin{bmatrix} is-in \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{bmatrix}$

Die Negation von ある (*aru*) ist unregelmäßig und erfordert einen eigenen Lexikoneintrag:

Morphologie:

¹⁵² vgl. JM, S. 210, 30-3.1.1. c) 6.

- (398) a. $v\text{-stem}^\diamond$
 b. PHON $\langle na \rangle$
 c. MAJ adj^v
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj,ni]^{opt}\}$
 e. CONT $\left[\begin{array}{c} not \\ \textcircled{1} \left[\begin{array}{c} is-in \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right] \end{array} \right]$

5.2.1.15 v_{cda} : だ – *da* – sein

Beispiel:

- (399) 長い旅でした。¹⁵³

naga-i tabi des.it-a
 lang-FLEX Reise ist-PAST
 Es war eine lange Reise.

da flektiert weitgehend wie ein v_{cr} -Verb, bildet jedoch einzelne Formen unregelmäßig. Nur ein Teil der theoretisch ableitbaren Formen wird tatsächlich verwendet. Für die Formen, die morphologisch vom Stamm *da* ableitbar sind, dient der folgende Lexikoneintrag:

Morphologie:

- (400) a. $v\text{-stem}^\diamond$
 b. PHON $\langle da \rangle$
 c. MAJ v_{cda}^{cv}
 d. SUBCAT $\{PP[subj,ga]^{opt}, N[obj]^{obl}\}$
 e. CONT $\left[\begin{array}{c} is \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right]$

Das Präsens von *da* wird unregelmäßig gebildet und entspricht nicht **daru*, wie es bei regelmäßigem Verhalten erwartet werden könnte, sondern *da*. Es wird durch einen eigenen Eintrag ins Lexikon realisiert¹⁵⁴:

¹⁵³ JM, S. 206, 30-2.1. a) 2.

¹⁵⁴ Lexikalisch-morphologische Regeln können auf unregelmäßige Verbformen wie *da* und unregelmäßige komplexe Verbstämme wie *desu* (Honorativ von *da*) nicht angewandt wer-

$$(401) \left[\begin{array}{c} \text{PHON } \langle da \rangle \\ \\ \text{SYNSEM} \mid \text{LOCAL} \end{array} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} is \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ} \mid \text{VCAT } v \\ \text{MOD } none \end{array} \right] \\ \text{SUBCAT} \left\{ \text{PP}[subj,ga]:\textcircled{1}^{opt}, \text{N}[obj]:\textcircled{2}^{obl} \right\} \\ \text{ADJS} \quad \{ \} \end{array} \right] \end{array} \right] \right]$$

de, das Partizip von *da*, ist ebenfalls unregelmäßig. Es entspricht dem folgenden Lexikoneintrag¹⁵⁴:

$$(402) \left[\begin{array}{c} \text{PHON } \langle de \rangle \\ \\ \text{SYNSEM} \mid \text{LOCAL} \end{array} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} adjunct \\ \textcircled{1} cont^{155} \\ \\ \textcircled{2} \left[\begin{array}{c} participle \\ \textcircled{1} \left[\begin{array}{c} is \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \end{array} \right] \end{array} \right] \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ} \mid \text{VCAT } v \\ \text{MOD } Vcat \end{array} \right] \\ \text{SUBCAT} \left\{ \text{PP}[subj,ga]:\textcircled{1}^{opt}, \text{N}[obj]:\textcircled{2}^{obl} \right\} \\ \text{ADJS} \quad \{ \} \end{array} \right] \end{array} \right] \right]$$

Der Honorativ wird nicht mit dem Suffixverb *-ma.s-u* gebildet, sondern mit der unregelmäßigen Basis *des-u*¹⁵⁴:

Morphologie:

den.

Formen, die normalerweise durch dieses Mittel definiert werden (wie beispielsweise die Formen, die für Adverben subkategorisieren) müssen daher direkt ins Lexikon eingetragen werden. Die entsprechenden Einträge werden jedoch an dieser Stelle der Übersichtlichkeit halber nicht aufgeführt.

¹⁵⁵ Das Argument $\textcircled{1}$, das Bezugswort des Adjunktes, wird durch das Head-Adjunkt-Schema instantiiert.

- (403) a. $v\text{-stem}^\diamond$
 b. PHON $\langle des \rangle$
 c. MAJ v_{imasu}^{cv}
 d. SUBCAT $\{ PP[subj,ga]^{opt}, N[obj]^{obl} \}$
 e. CONT $\left[\begin{array}{c} \text{honorable} \\ \textcircled{1} \left[\begin{array}{c} is \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right] \end{array} \right]$

desu dient auch zur Bildung des Honorativs von Adjektiven; in diesem Fall nimmt es einen Satz als Argument (vgl. den Eintrag zum Honorativ-Suffix, S. 248).

5.2.1.16 $v_{cnasaru}$: くださる – *kudasa.r-u* – bitte

Beispiel:

- (404) パンを薄く切ってください。 ¹⁵⁶

pañ-wo usu-ku ki.tt-e kudasai
 Brot-ACC dünn-ADV schneiden-PART bitte
 Schneide das Brot bitte dünn.

Morphologie:

- (405) a. $v\text{-stem}^\diamond$
 b. PHON $\langle kudasa \rangle$
 c. MAJ $v_{cnasaru}^{cv}$
 d. SUBCAT $\{ \}$
 e. CONT $[please\text{-}do\text{-}for\text{-}me]$

Der Imperativ wird unregelmäßig gebildet:^{157, 158}

¹⁵⁶ JM, S. 226, 32-2.3.2.

¹⁵⁷ Die Varianten dieses Zeichens, die für Adverben subkategorisieren bzw. adnominal stehen, wurden weggelassen (vgl. Fußnote 154 auf Seite 213).

¹⁵⁸ *kudasai* wird in der Bedeutung, die wir hier betrachten, immer zusammen mit einem adverbialen Satz verwendet (vgl. z.B. (404)). Die Merkmale CONT und ADJS des entsprechenden Eintrages seien deshalb als Beispiel angeführt:

$$(407) \left[\begin{array}{c} \text{PHON } \langle kudasai \rangle \\ \\ \text{SYNSEM} \mid \text{LOCAL} \end{array} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \textit{imperative} \\ \textcircled{1} \textit{ please-do-for-me} \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ} \mid \text{VCAT} \quad v \\ \text{MOD} \quad \textit{none} \end{array} \right] \\ \text{SUBCAT} \{ \} \\ \text{ADJS} \{ \} \end{array} \right] \end{array} \right] \right]$$

5.2.1.17 v_{cw} : 買う – *ka.-u* – kaufen

Beispiel:

(408) 太郎は自転車を買いました。

tarou-ha jitensya-wo ka..i-mas.it-a
 Tarō-TOP Fahrrad-ACC kaufen-HON-PAST
 Tarō kaufte ein Fahrrad.

Morphologie:

$$(409) \begin{array}{ll} a. & v\text{-stem}^\diamond \\ b. & \text{PHON } \langle ka \rangle \\ c. & \text{MAJ } v_{cw}^{cv} \\ d. & \text{SUBCAT } \{ \text{PP}[\textit{subj,ga}]^{opt}, \text{PP}[\textit{obj,wo}]^{opt} \} \\ e. & \text{CONT } \left[\begin{array}{c} \textit{buy} \\ \textcircled{1} \text{ cont}(\textit{subj}) \\ \textcircled{3} \text{ cont}(\textit{obj}) \end{array} \right] \end{array}$$

5.2.1.18 v_{ikuru} : 来る – *ku-ru* – kommen

Beispiel:

(410) 健は空港まで来た。

keñ-ha kuukou-made kit-a
 Ken-TOP Flughafen-DIR kommen-PAST
 Ken kam zum Flughafen.

$$(406) \left[\begin{array}{c} \dots \mid \text{CONT} \left[\begin{array}{c} \textit{imperative} \\ \textcircled{1} \textcircled{2} \left[\begin{array}{c} \textit{adjunct} \\ \textcircled{1} \textit{ please-do-for-me} \end{array} \right] \end{array} \right] \\ \\ \dots \mid \text{ADJS} \left\{ [] : \textcircled{2}^{obl} \right\} \end{array} \right]$$

Morphologie:

- (411) a. *v-stem*[◇]
 b. SUBCAT $\{ \text{PP}[\text{subj}, \text{ga}]^{\text{opt}}, \text{PP}[\text{obj}, \text{ni} \vee \text{he} \vee \text{made}]^{\text{opt}} \}$
 c. CONT $\begin{bmatrix} \text{come} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{bmatrix}$
 d. i. MAJ $v_{i\text{kuru}}^{\text{cons}}$; PHON $\langle \text{ku} \rangle$
 ii. MAJ $v_{i\text{kuru}}^i$; PHON $\langle \text{ki} \rangle$
 iii. MAJ $v_{i\text{kuru}}^a$; PHON $\langle \text{ko} \rangle$
 iv. MAJ $v_{i\text{kuru}}^{\text{onbin}}$; PHON $\langle \text{kit} \rangle$

Der Imperativ von *kuru* wird unregelmäßig gebildet:¹⁵⁹

$$(412) \left[\begin{array}{c} \text{PHON } \langle \text{koi} \rangle \\ \\ \text{SYNSEM} \mid \text{LOCAL} \end{array} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \text{imperative} \\ \textcircled{1} \left[\begin{array}{c} \text{come} \\ \textcircled{1} \quad \boxed{1} \\ \textcircled{2} \quad \boxed{2} \end{array} \right] \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ} \mid \text{VCAT } v \\ \text{MOD } \textit{none} \end{array} \right] \\ \text{SUBCAT} \left\{ \text{PP}[\text{subj}, \text{ga}]:\boxed{1}^{\text{opt}}, \right. \\ \left. \text{PP}[\text{obj}, \text{ni} \vee \text{he} \vee \text{made}]:\boxed{2}^{\text{opt}} \right\} \\ \text{ADJS } \{ \} \end{array} \right] \end{array} \right] \right]$$

Ebenfalls unregelmäßig ist der Volitional (Futur)¹⁵⁹:

¹⁵⁹ Die Varianten dieses Zeichens, die für Adverben subkategorisieren, bzw. adnominal stehen, wurden weggelassen (vgl. Fußnote 154, S. 213 und Fußnote 157, S. 215)

$$(413) \left[\begin{array}{c} \text{PHON } \langle kyou \rangle \\ \\ \text{SYNSEM | LOCAL} \end{array} \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \text{future} \\ \textcircled{1} \left[\begin{array}{c} \text{come} \\ \textcircled{1} \quad \boxed{1} \\ \textcircled{2} \quad \boxed{2} \end{array} \right] \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{MAJ | VCAT } v \\ \text{MOD } none \end{array} \right] \\ \text{SUBCAT} \left\{ \text{PP}[\text{subj}, \text{ga}]:\boxed{1}^{opt}, \right. \\ \left. \text{PP}[\text{obj}, \text{ni} \vee \text{he} \vee \text{made}]:\boxed{2}^{opt} \right\} \\ \text{ADJS} \quad \{ \} \end{array} \right] \end{array} \right] \right]$$

Eine der zwei möglichen Potentialformen¹⁵⁹ wird durch einen unregelmäßigen vokalischen Verbstamm realisiert:

Morphologie:

$$(414) \begin{array}{ll} a. & v\text{-stem}^\diamond \\ b. & \text{PHON } \langle kore \rangle \\ c. & \text{MAJ } v_v^{comb} \\ d. & \text{SUBCAT } \left\{ \text{PP}[\text{subj}, \text{ga}]:\boxed{1}^{opt}, \right. \\ & \left. \text{PP}[\text{obj}, \text{ni} \vee \text{he} \vee \text{made}]:\boxed{2}^{opt} \right\} \\ e. & \text{CONT} \left[\begin{array}{c} \text{can} \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \left[\begin{array}{c} \text{come} \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right] \end{array} \right] \end{array}$$

5.2.1.19 $v_{i_{suru}}$: する – *s.-uru* – tun, machen

Beispiel:

(415) あの人は英語を勉強する。

ano hito-ha eigo-wo beñkyou s.-uru
 jener Mensch-TOP Englisch-ACC Studium tun-FLEX
 Er studiert Englisch.

Morphologie:

- (416) a. *v-stem*[◇]
 b. PHON $\langle s \rangle$
 c. MAJ $v_{i_{suru}}^{cv}$
 d. SUBCAT $\{PP[subj,ga]^{opt}, PP[obj,wo]^{opt}\}$
 e. CONT $\begin{bmatrix} do \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) \end{bmatrix}$

Beide Imperativformen von する (*suru*) sind unregelmäßig:¹⁶⁰

- (417) $\left[\begin{array}{c} \text{PHON } \langle siro \rangle \vee \langle seyo \rangle \\ \\ \text{SYNSEM} \mid \text{LOCAL} \\ \\ \text{CONT} \begin{bmatrix} imperative \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \begin{bmatrix} do \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textcircled{2} \end{bmatrix} \end{bmatrix} \\ \\ \text{CAT} \begin{bmatrix} \text{HEAD} \begin{bmatrix} \text{MAJ} \mid \text{VCAT } v \\ \text{MOD } none \end{bmatrix} \\ \text{SUBCAT} \left\{ PP[subj,ga]:\textcircled{1}^{opt}, PP[obj,ni \vee he \vee made]:\textcircled{2}^{opt} \right\} \\ \text{ADJS } \{ \} \end{bmatrix} \end{array} \right]$

Das Verb できる (*dekiru*) realisiert die Potentialform von する (*suru*). Die Änderung der Kasusreaktion beim Potential ist in diesem Fall obligatorisch: Das direkte Objekt von できる (*dekiru*) wird immer durch が (*ga*) markiert:¹⁶¹

Beispiel:

- (418) 私はチェスが / *を出来る。 ¹⁶²

*watasi-ha tyesu-ga / *-wo deki-ru*
 ich-TOP Schach-NOM / *-ACC können-FLEX
 Ich kann Schach spielen.

¹⁶⁰ Varianten dieses Zeichens, die für Adverben subkategorisieren bzw. adnominal stehen, werden aus Platzgründen weggelassen.

¹⁶¹ vgl. DBJG, S. 371, notes 3.

¹⁶² vgl. DBJG, S. 372, notes 3. (3).

Morphologie:

- (419) a. *v-stem*[◇]
 b. PHON $\langle deki \rangle$
 c. MAJ v_v^{comb}
 d. SUBCAT $\left\{ \begin{array}{l} PP[subj,ga]:\boxed{1}^{opt}, \\ PP[obj,ga]:\boxed{2}^{opt} \end{array} \right\}$
 e. CONT $\left[\begin{array}{l} can \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \left[\begin{array}{l} come \\ \textcircled{1} \text{ cont(subj)} \\ \textcircled{2} \text{ cont(obj)} \end{array} \right] \end{array} \right]$

Andere Valenzrahmen für *suru* sind möglich. Eine der Hauptfunktionen von *suru* ist es beispielsweise, mit einem Nomen zusammen ein komplexes Verb zu bilden. Das Nomen erscheint dabei entweder als Akkusativobjekt (PP[*wo*]), oder es wird *suru* vorangestellt. In beiden Fällen bestimmt das Nomen die Valenz von *suru*:

Beispiele:

- (420) a. SUBCAT $\{ PP[ga]^{opt}, PP[ni]^{opt}, PP[wo]^{opt} \}$

ジョンは先生に質問をした。 ¹⁶³

joñ-ha señsei-ni situmõñ-wo s.it-a
 John-TOP Lehrer-DAT Frage-ACC machen-PAST
 John stellte dem Lehrer eine Frage.

- b. SUBCAT $\{ N^{obl}, PP[ga]^{opt}, PP[wo]^{opt} \}$

日本語を勉強しなければならない。 ¹⁶⁴

nihoñgo-wo beñkyou s.-ina-kereba na.r.a-na-i
 Japanisch-ACC Studium machen-NEG-COND werden-NEG-FLEX
 Du mußt Japanisch lernen.

- c. SUBCAT $\{ N^{obl}, PP[ga]^{opt}, PP[to]^{opt} \}$

あんな男とは結婚するな! ¹⁶⁵

anna otoko-to-ha kekkoñ s.-uru na
 jene Art von Mann-KOM-TOP Heirat machen-FLEXNEG IMP
 Heirate keinen solchen Mann!

163 DBJG, S. 366, notes 3. (2)a.

164 vgl. JM, S. 97, 20-3.1.2. c).

165 DBJG, S. 266, ex. (c).

5.2.1.20 $v_{i_{masu}}$: -ます – **-mas.-u** – **Honorativ***Beispiele:*(421) a. 日本へ来ましてから一年たちました。¹⁶⁶

nihoñ-he ki-mas.it-ekara itineñ ta.t.i-mas.it-a
 Japan-DIR kommen-HON-PARTSeit ein Jahr vergehen-HON-PAST
 Es ist ein Jahr vergangen, seitdem ich nach Japan gekommen bin.

b. 田中さんが食べたステーキは高かったです。¹⁶⁷

tanaka-san-ga tabe.t-a sutēki-ha taka-ka.tt-a des-u
 Herr Tanaka-NOM essen-PAST Steak-TOP teuer-φ-PAST HON-FLEX
 Das Steak, das Herr Tanaka gegessen hat, war teuer.

Diese Kategorie findet nur bei der Bildung des Honorativs Verwendung. Vergleiche den Abschnitt zum Honorativsuffix **-mas.-u** auf Seite 248 und den Eintrag von *desu* ((403), S. 214), der unregelmäßigen Honorativform von *da*.

5.2.1.21 *adj*: 高い – **taka-i** – hoch, teuer*Beispiel:*(422) この本はおそろしく高い。¹⁶⁸

kono hoñ-ha osorosi-ku taka-i
 dieses Buch-TOP furchtbar-ADV teuer-FLEX
 Dieses Buch ist furchtbar teuer.

Morphologie:(423) a. $v\text{-stem}^\diamond$ b. PHON $\langle taka \rangle$ c. MAJ adj^v d. SUBCAT $\{PP[ga]^{opt}\}$ e. CONT $\left[\begin{array}{l} high \\ \textcircled{1} \text{ cont(subj)} \end{array} \right]$

¹⁶⁶ JM, S. 97, 20-3.2.1.1.

¹⁶⁷ DBJG, S. 376, KS(A).

¹⁶⁸ JM, S. 205, 30-2.1. a) 1.

5.2.2 Clitics

Clitics dienen, wie in Abschnitt “*Verbstämme*”, S. 89 beschrieben, der Ableitung der Stammformen der Verben, d.h. der systematischen Beschreibung der Allomorphie der Verb- oder Suffixmorpheme. Die Kategorie der Stämme wird von Clitics nicht verändert.¹⁶⁹

Clitics werden durch Tabellen beschrieben. Zur Erinnerung wird die Hierarchie der Stammformen, auf der die morphologische Beschreibung beruht,

169 Clitics übernehmen die Verbkategorie des Stammes und überschreiben nur die Stammform. Diese Eigenschaft folgt schon aus dem Defaultzeichen für Clitics (vgl. (356), S. 198):

$$(424) \quad \left[\begin{array}{c} \text{clitic}^\diamond \\ \text{MORPHOLOGY} \mid \text{MCAT} \left[\begin{array}{c} \text{MHEAD} \mid \text{MAJ} \mid \text{VCAT} \quad \boxed{} \\ \text{STEM} \left[\text{MCAT} \mid \text{MHEAD} \mid \text{MAJ} \mid \text{VCAT} \quad \boxed{} \right] \end{array} \right] \end{array} \right]$$

Obwohl die Schreibweise $v_{c_k}^{cv} \longrightarrow v_{c_k}^{onbin}$ damit redundant ist, da sie die Verbkategorie zwei Mal aufführt, wird sie, der Anschaulichkeit und Einheitlichkeit halber, auch für Clitics benutzt.

Die Kategorien in den Disjunktionen auf der linken und rechten Seite des Pfeils sind “der Reihenfolge entsprechend” aufzulösen. Die Abbildung

$$a_1 \vee a_2 \vee \dots \vee a_n \quad \longrightarrow \quad b_1 \vee b_2 \vee \dots \vee b_n$$

beispielsweise ist als Abkürzung von

$$\begin{array}{ccc} a_1 & \longrightarrow & b_1 \\ a_2 & \longrightarrow & b_2 \\ & \dots & \\ a_n & \longrightarrow & b_n \end{array}$$

zu verstehen.

Das Clitic *i* beispielsweise dient zur Ableitung des i-Stammes der konsonantischen Verben mit Ausnahme von *nasaru* und der unregelmäßigen Verben mit Ausnahme von *kuru*. Es wird durch folgende Abbildung beschrieben:

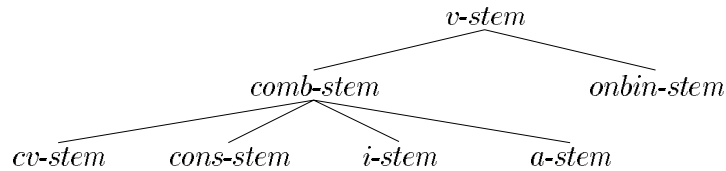
$$v_{c \setminus \text{nasaru}, i \text{ suru}, \text{masu}}^{\text{cons}} \quad \longrightarrow \quad v_{c \setminus \text{nasaru}, i \text{ suru}, \text{masu}}^i$$

Diese Darstellung steht damit für die folgenden Abbildungen (die sich allerdings noch weiter auflösen lassen):

$$\begin{array}{ccc} v_{c \setminus \text{nasaru}}^{\text{cons}} & \longrightarrow & v_{c \setminus \text{nasaru}}^i \\ v_{i \text{ suru}}^{\text{cons}} & \longrightarrow & v_{i \text{ suru}}^i \\ v_{i \text{ masu}}^{\text{cons}} & \longrightarrow & v_{i \text{ masu}}^i \end{array}$$

hier noch einmal dargestellt:

(425) *Stammformen der Verben:*



5.2.2.1 Konsonantischer Stamm (*cons-stem*)

Der konsonantische Stamm (*cons-stem*) entsteht bei konsonantischen Verben aus der Konkatenation von traditionellem Verbstamm (hier als *cv-stem* modelliert) und Stammvokal. Bei vokalischen Verben ist er, genauso wie *cv-*, *i-* und *a-Stamm*, mit dem Kombinationsstamm (*comb-stem*) identisch. Der konsonantische Stamm entspricht damit dem gemeinsamen Präfix der traditionellen Verbstämme, wenn diese in der hier benutzten Umschrift geschrieben werden.¹⁷⁰

Der konsonantische Stamm ist von allen Verbstämmen unseres Systems der produktivste: *i-* und *a-Stamm*, *Präsens*, *Futur*, *Imperativ*, *Konditional* und *Potential*¹⁷¹ werden mit seiner Hilfe gebildet.

(426)

<i>Phon</i>	<i>Maj</i>		
<i>k</i>	$v_{c_K}^{cv}$	$\longrightarrow$	$v_{c_K}^{cons}$
<i>s</i>	$v_{c_s}^{cv}$	$\longrightarrow$	$v_{c_s}^{cons}$
<i>t</i>	$v_{c_t}^{cv}$	$\longrightarrow$	$v_{c_t}^{cons}$
<i>n</i>	$v_{c_n}^{cv}$	$\longrightarrow$	$v_{c_n}^{cons}$
<i>m</i>	$v_{c_m}^{cv}$	$\longrightarrow$	$v_{c_m}^{cons}$
<i>r</i>	$v_{c_R}^{cv}$	$\longrightarrow$	$v_{c_R}^{cons}$
ϕ	$v_{c_w,i}^{cv}$	$\longrightarrow$	$v_{c_w,i}^{cons}$
<i>g</i>	$v_{c_g}^{cv}$	$\longrightarrow$	$v_{c_g}^{cons}$
<i>b</i>	$v_{c_b}^{cv}$	$\longrightarrow$	$v_{c_b}^{cons}$

¹⁷⁰ In der traditionell benutzten Kana-Schreibweise entspricht der traditionelle Verbstamm (*cv-stem*) dem gemeinsamen Präfix.

¹⁷¹ Ausgehend vom *a-stem* bilden die vokalischen Verben und das Verb 来る (*kuru*) eine zweite Potentialform.

5.2.2.2 i-Stamm (*i-stem*)

Der i-Stamm (*i-stem*) — traditionell als 連用形 (*reñyoukei*) bezeichnet — dient als Basis zur Ableitung von *Honorativ* und *Voluntativ*. Abgesehen von den $v_{c_{nasaru}}$ -Verben wird er bei allen konsonantischen Verben durch Suffigierung des Vokales *i* an den konsonantischen Stamm gebildet. Bei den vokalischen Verben ist er mit dem Kombinationsstamm (*comb-stem*) identisch.

(427)

<i>Phon</i>	<i>Maj</i>
<i>i</i> $v_{c_{nasaru}, i_{suru}, masu}^{cons}$	$\longrightarrow$ $v_{c_{nasaru}, i_{suru}, masu}^i$
<i>i</i> $v_{c_{nasaru}}^{cv}$	$\longrightarrow$ $v_{c_{nasaru}}^i$

5.2.2.3 a-Stamm (*a-stem*)

Der a-Stamm (*a-stem*) dient der Ableitung der *Negation* und des *negierten Partizips*. Auch *Kausativ*, *Passiv* und *Potential* werden mit ihm gebildet. In der traditionellen Beschreibung der Verbmorphologie hat er den 未然形 (*mizenkei*) als Entsprechung.

(428)

<i>Phon</i>	<i>Maj</i>
<i>a</i> $v_{c_w}^{cons}$	$\longrightarrow$ $v_{c_w}^a$
<i>wa</i> $v_{c_w}^{cons}$	$\longrightarrow$ $v_{c_w}^a$
<i>o</i> $v_{i_{kuru}}^{cv}$	$\longrightarrow$ $v_{i_{kuru}}^a$
ϕ $v_{i_{suru}}^{cv}$	$\longrightarrow$ $v_{i_{suru}}^a$
<i>e</i> $v_{i_{masu}}^{cv}$	$\longrightarrow$ $v_{i_{masu}}^a$

5.2.2.4 Onbin-Stamm (*onbin-stem*)

Die Onbin-Form (音便, *onbin*) ist sprachgeschichtlich aus der Assimilation des Vokales *i*, auf dem der 連用形 (*reñyoukei*) auslautet, mit dem Anlautkonsonanten *t* bestimmter sich anschließender Suffixe entstanden. In den meisten Grammatiken wird er deshalb auch heute noch als “Verschleifung” aufgefasst und durch “phonologische” Transformationsregeln modelliert. Da jedoch aus sprachhistorischer Perspektive vermutlich auch alle anderen Stämme aus Assimilationserscheinungen entstanden sind und die Onbin-Verschleifung im modernen Japanisch obligatorisch ist, liegt es nahe, sie mit den gleichen Mitteln zu behandeln wie die anderen morphologischen Variationen: durch eine eigene Stammform, hier als *Onbin-Stamm* (*onbin-stem*) bezeichnet.

Die Annahme eines Onbin-Stammes ist — zusammen mit den Konsequenzen, die sich daraus ergeben — der eigentliche Unterschied zwischen dem hier vorgeschlagenen System und anderen Modellen der japanischen Verbmorphologie.

Der Onbin-Stamm dient in unserem Modell u.a. der Ableitung von *Perfekt*, *Partizip* und *Perfekt-Konditional*.

(429)

<i>Phon</i>		<i>Maj</i>	
<i>t</i>	v_v^{cv}	$\longrightarrow$	v_v^{onbin}
<i>it</i>	$v_{c_k,i}^{cv}$	$\longrightarrow$	$v_{c_k,i}^{onbin}$
<i>id</i>	$v_{c_g}^{cv}$	$\longrightarrow$	$v_{c_g}^{onbin}$
<i>sit</i>	$v_{c_s}^{cv}$	$\longrightarrow$	$v_{c_s}^{onbin}$
<i>tt</i>	$v_{c_{t,w,R,iku}}^{cv}$	$\longrightarrow$	$v_{c_{t,w,R,iku}}^{onbin}$
<i>ñd</i>	$v_{c_{n,b,m}}^{cv}$	$\longrightarrow$	$v_{c_{n,b,m}}^{onbin}$

5.2.3 Flexive

Dieser Abschnitt dient der Beschreibung der *Flexive* des vorgestellten morphologischen Fragmentes. Als *Flexive* werden Suffixe bezeichnet, die den morphologischen Ableitungsprozess abschließen: Ein Verb, das auf einem Flexiv endet, entspricht einem atomaren syntaktischen Zeichen und kann nicht durch die Suffigierung weiterer Morpheme modifiziert werden.

5.2.3.1 Präsens (-FLEX)

In Kapitel 2 wurde die *Ökonomie* als ein Prinzip der japanischen Sprache vorgestellt: Im Allgemeinen werden nur solche Informationen explizit ausgedrückt, die sich nicht aus der Situation und den sprachlich-kommunikativen Grundannahmen erschließen lassen. Das Präsens kann als “neutrale” Zeit angesehen werden und wird daher, im Gegensatz zum Futur oder Perfekt, im Japanischen nicht explizit ausgedrückt. Das in diesem Abschnitt eingeführte Suffix fügt der semantischen Darstellung des Verbes dementsprechend keine zeitliche Information hinzu. Es erfüllt damit nur die formale Aufgabe eines Flexives: den morphologischen Bildungsprozess abzuschließen und eine Form zu bilden, die in der Syntax verwendet werden kann.¹⁷²

Verben, die mit dem Präsens-Flexiv abschließen, können satzabschließend zur Bildung von Hauptsätzen verwendet werden¹⁷³ oder adnominal zur Bildung von Nebensätzen. Zuerst wird die Hauptsatz-bildende Form beschrieben, die wir auch einfach als *Satzschluß-* oder *prädikative* Form bezeichnen; anschließend folgt die Darstellung der *Adnominalform*.

Satzschlußform:

Beispiele:

(430) a. コーヒーを飲む / 飲みます。¹⁷⁴

kouhii-wo no.m-u / no.m.i-mas.-u
Kaffee-ACC trinken-FLEX / trinken-HON-FLEX
Ich trinke Kaffee.

b. この本はおそろしく高い。¹⁷⁵

172 Diese Überlegung zugrunde legend, wäre es sinnvoller, nicht von einem *Präsens-Flexiv* zu sprechen, sondern einfach von einem (*sematisch neutralen*) Flexiv. Wenn dieses hier nicht geschieht, hat es den Grund, daß wir der von Rickmeyer 1995 vorgeschlagenen Einteilung der Flexive und Suffixe folgen und dort der Begriff “Präsens-Flexiv” (vgl. JM, s. 77, 20-2.1.1. -Ru ‘Präsens’) verwendet wird.

173 Nach einem satzabschließenden Verb können nur noch Interjektionspartikel stehen.

174 JM, S. 78, 20-2.1.1. b) 1.

175 JM, S. 205, 30-2.1. a) 1.

kono hoñ-ha osorosi-ku taka-i
 dieses Buch-TOP furchtbar-ADV teuer-FLEX
 Dieses Buch ist furchtbar teuer.

c. いま考えている / います。¹⁷⁶

ima kañgae.t-e i-ru / i-mas.-u
 jetzt denken-PART KONT-FLEX / KONT-HON-FLEX
 Ich denke gerade darüber nach.

Morphologie:

(431) a. *flexiv*[◇]

- | | | | | | |
|-------|-----|------------------------------------|------------------------|------|-----------------------|
| b. i. | MAJ | v_{v,i_kuru}^{cons} | $\longrightarrow v;$ | PHON | $\langle ru \rangle$ |
| ii. | MAJ | $v_{c\backslash da,i masu}^{cons}$ | $\longrightarrow v;$ | PHON | $\langle u \rangle$ |
| iii. | MAJ | $v_{i suru}^{cons}$ | $\longrightarrow v;$ | PHON | $\langle uru \rangle$ |
| iv. | MAJ | adj^v | $\longrightarrow adj;$ | PHON | $\langle i \rangle$ |

Adnominalform:

Die Adnominalform gestattet es, Verb- oder Adjektivkonstruktionen einem Nomen attributiv zuzuordnen:

Beispiele:

(432) a. 英語ができる人¹⁷⁷

eigo-ga deki-ru hito
 Englisch-NOM können-FLEX Mensch
 jemand, der Englisch kann

b. 長い旅でした。¹⁷⁸

naga-i tabi des.it-a
 lang-FLEX Reise ist-PAST
 Es war eine lange Reise.

Verben bzw. Adjektive in dieser Form modifizieren also ein Nomen und haben die Semantik von Adjunkten:

¹⁷⁶ JM, S. 78, 20-2.1.1. b) 1.

¹⁷⁷ JM, S. 78, 20-2.1.1. b) 2.

¹⁷⁸ JM, S. 206, 30-2.1. a) 2.

Morphologie:

- (433) *c.* MOD N
- d.* CONT $\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont}^{179} \\ \textcircled{2} \text{ cont(stem)}^{180} \end{array} \right]$

In unserem Modell werden Adjektive und Verben syntaktisch und morphologisch sehr ähnlich behandelt: Adjektive werden als hauptsatzbildende Prädikate, die für ein Subjekt subkategorisieren, im Lexikon aufgenommen:

- (435) *das Adjektiv 高い (takai; hoch) in prädikativer Verwendung:*

建物が高い。

tatemono-ga taka-i
Gebäude-NOM hoch-FLEX
Das Gebäude ist hoch.

Die adnominale Verwendung von Adjektiven hingegen wird, genau wie die adnominale Verwendung von Verben, als Relativsatz modelliert:

- (436) *das Adjektiv 高い (takai; hoch) in adnominaler Verwendung:*

高い建物

taka-i tatemono
hoch-FLEX Gebäude
ein hohes Gebäude

Es wird die semantische Struktur

- (437) CONT $\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont(stem)} \end{array} \right]$

179 Das Merkmal $\textcircled{1}$ wird durch das Head-Adjunkt-Schema instantiiert und hat als Wert die durch den Nebensatz modifizierte semantische Struktur. Das nominale Bezugswort muß, um durch ein Adjunkt modifiziert werden zu können, vorher durch eine lexikalische bzw. lexikalisch-morphologische Regel transformiert werden (vgl. 3.2.2 “Adjunkte”, S. 130).

180 **cont(stem)** ist folgendermaßen definiert:

- (434) Abkürzung: abgekürzte AVM:
- $$\left[\dots | X \text{ cont(stem)} \right] \left[\begin{array}{l} \dots | X \boxed{} \\ \dots |_{\text{STEM}} \left[\dots |_{\text{CONT}} \boxed{} \right] \end{array} \right]$$

verwendet, in der das erste Argument die Semantik des Bezugsnomens (Postcedent) aufnimmt und das zweite Argument die Semantik des Relativsatzes.

Die semantische Beziehung zwischen dem Bezugswort und den Argumenten des Relativsatzes kann durch Koindizierung ausgedrückt werden. Verschiedene Beziehungen sind möglich:

1. Das Bezugswort entspricht dem Subjekt des Relativsatzes:

(438) 日本語を教えている先生¹⁸¹

nihongo-wo osie.t-e i-ru señsei
 Japanisch-ACC unterrichten-PART KONT-FLEX Lehrer
 Der Lehrer, der Japanisch unterrichtet

Diese Beziehung kann folgendermaßen ausgedrückt werden:

(439) *e. i.* CONT [① **cont(subj)**]¹⁸²

2. Das Bezugswort entspricht dem direkten Objekt des Relativsatzes:

(441) ジョンが¹⁸³ 食べたステーキ¹⁸⁴

Jon-ga tabe.t-a sutēki
 John-NOM essen-PAST Steak
 Das Steak, das John gegessen hat

Diesem Fall entspricht die folgende Spezifikation:

¹⁸¹ vgl. DBJG, S. 377, ex. (a).

¹⁸² Die Abkürzung **cont(subj)** ist folgendermaßen definiert:

(440) Abkürzung: abgekürzte AVM:

$$[... | X \text{ cont(subj)}] \left[\begin{array}{l} ... | X \text{ ①} \\ ... | \text{SUBCAT} \{ \text{PP}[\text{subj}]: \text{①} \} \end{array} \right]$$

mit $X:\text{①}$ für $\text{①} = \text{cont}(X)$ in Subkategorisierungskontexten.

¹⁸³ Das Partikel *ga* kann in bestimmten Fällen auch durch *no* ersetzt werden:

(442) ジョンの食べたステーキ

Jon-no tabe.t-a sutēki
 John-GEN essen-PAST Steak
 Das Steak, das John gegessen hat

¹⁸⁴ DBJG, S. 378, notes 3. (6).

(443) *e. ii.* CONT $\left[\textcircled{1} \text{ cont}(\mathbf{obj})^{185} \right]$

3. Das Bezugswort entspricht dem indirekten Objekt des Relativsatzes (beispielsweise einer PP[*ni*]):

(445) 道子が行く学校¹⁸⁶

Mitiko-ga i.k-u gakkou
 Mitiko-NOM gehen-FLEX Schule
 Die Schule, in die Mitiko geht

Die folgende AVM drückt diesen Bezug aus:

(446) *e. i.* CONT $\left[\textcircled{1} \text{ cont}(\mathbf{obj2})^{187} \right]$

4. Das Bezugswort entspricht keinem der Argumente des Nebensatzes — es handelt sich also streng genommen nicht um eine Relativsatzkonstruktion:

(448) 魚がこげるにおい¹⁸⁸

sakana-ga koge-ru nioi
 Fisch-NOM anbrennen-FLEX Gestank
 Der Gestank von anbrennendem Fisch

Diesem Fall entsprechen die Flexive, wie sie durch (431) zusammen mit (433) ausgedrückt werden.

Die Auflösung der semantischen Beziehung zwischen Nebensatz und Bezugswort setzt in den meisten Fällen das “Verständnis” des Satzes und der beschriebenen Situation voraus. Am sinnvollsten ist es daher, die Beziehung

185 Die Abkürzung **cont(obj)** ist folgendermaßen definiert:

$$(444) \quad \begin{array}{ll} \text{Abkürzung:} & \text{abgekürzte AVM:} \\ \left[\dots | X \text{ cont}(\mathbf{obj}) \right] & \left[\dots | X \boxed{} \right. \\ & \left. \dots | \text{SUBCAT} \{ \text{PP}[\mathbf{obj}]: \boxed{}, \dots \} \right] \end{array}$$

186 vgl. DBJG, S. 377, ex. (e).

187 Die Abkürzung **cont(obj2)** ist folgendermaßen definiert:

$$(447) \quad \begin{array}{ll} \text{Abkürzung:} & \text{abgekürzte AVM:} \\ \left[\dots | X \text{ cont}(\mathbf{obj2}) \right] & \left[\dots | X \boxed{} \right. \\ & \left. \dots | \text{SUBCAT} \{ \text{PP}[\mathbf{obj2}]: \boxed{}, \dots \} \right] \end{array}$$

188 DBJG, S. 379, notes 4. (8).

zwischen Nebensatz und Postcedent unspezifiziert zu lassen und nur den Bereich der möglichen Bezüge anzugeben. Eine in dieser Weise *unterspezifizierte* semantische Darstellung setzt allerdings eine komplexere semantische Beschreibungssprache voraus als hier verwendet.¹⁸⁹

5.2.3.2 Perfekt (-PAST)

Das *Perfekt-Flexiv*¹⁹⁰ dient dem Ausdruck der Vergangenheit bzw. Vorzeitigkeit bezüglich des Äußerungszeitpunktes oder dem in der Matrixphrase dargestellten Sachverhalt:

Satzschlußform:

Beispiele:

- (449) a. あの人の全集は残らず読んだ。¹⁹¹

ano hito-no zeñsyuu-ha nokor-azu
jener Mensch-GEN gesamte Werke-TOP übrigbleiben-NEG PART
yo.ñd-a
lesen-PAST
Ich habe seine gesamten Werke ausnahmslos gelesen.

- b. 彼は先月ここへ来ました。¹⁹²

kare-ha señgetu koko-he ki.mas.it-a
er-TOP letzter Monat hier-DIR kommen-HON-PAST
Er ist letzten Monat hierher gekommen.

- c. 次の日の試験は難しかったなあ。¹⁹³

tugi-no hi-no sikeñ-ha muzukasi.ka.tt-a naa
nächster-GEN Tag-GEN Prüfung-TOP schwer-φ-PAST INTERJEKTION
Die Prüfung am nächsten Tag war aber schwer!

Morphologie:

- (450) a. *flexiv*[◇]

b. CONT $\left[\begin{array}{l} \textit{past} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right]$

c. MAJ $v^{onbin} \longrightarrow v$; PHON $\langle a \rangle$

¹⁸⁹ Vgl. beispielsweise Pinkal 1995, Pinkal 1999 oder Copestake et al. 1997.

¹⁹⁰ JM, S. 84, 20-2.2.2. -Ta 'Perfekt'.

¹⁹¹ JM, S. 84, 20-2.2.2. a) 1.

¹⁹² JM, S. 85, 20-2.2.2. b) 1.

¹⁹³ vgl. JM, S. 210, 30-3.1.1. c) 6.

Die Perfektform von Adjektiven wird unter Zuhilfenahme des semantisch leeren Suffixverbes *-Ka.r-u*¹⁹⁴ gebildet:

- (451) 高い $\longrightarrow$ 高かる $\longrightarrow$ 高かった
taka-i *taka-ka.r-u* *taka-ka.tt-a*
hoch-FLEX hoch- ϕ -FLEX hoch- ϕ -PAST

Adnominalform:

Das Perfekt-Flexiv kann genauso wie das Präsens-Flexiv adnominal gebraucht werden:

Beispiele:

- (452) a. きのう見た映画はつまらなかった。¹⁹⁵
kinou mi.ta eiga-ha tumarana.ka.tt-a
gestern sehen-PAST Film-TOP langweilig- ϕ -PAST
Der Film, den ich gestern gesehen hatte, war langweilig
- b. 田中さんが食べたステーキは高かった / 高かったです。¹⁹⁶
tanaka-sa $\tilde{n}$ -ga tabe.t-a sutēki-ha taka-ka.tt-a / taka-ka.tt-a
Herr Tanaka-NOM essen-PAST Steak-TOP teuer- ϕ -PAST / teuer- ϕ -PAST
des-u
HON-FLEX
Das Steak, das Herr Tanaka gegessen hat, war teuer.

Morphologie:

- (453) d. MOD N
- e. CONT $\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \left[\begin{array}{l} \text{past} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right]$

Für die Beziehungen zwischen dem Relativsatz und dem Postcedent gibt es die gleichen Möglichkeiten wie beim Präsens-Flexiv:

- (454) f. i. CONT $\left[\textcircled{1} \text{ cont(subj)} \right]$
ii. CONT $\left[\textcircled{1} \text{ cont(obj)} \right]$
iii. CONT $\left[\textcircled{1} \text{ cont(obj2)} \right]$

¹⁹⁴ vgl. 5.2.4

¹⁹⁵ JM, S. 84, 20-2.2.2. a) 2.

¹⁹⁶ DBJG, S. 376, KS(A).

5.2.3.3 Futur / Volitional (-VOLI)

Vermutungen und Möglichkeiten (vgl. (455a)), Absichten (vgl. (455b), (455c), (455d)) und Aufforderungen (vgl. (455e)) können durch das *Futur*¹⁹⁷ ausgedrückt werden:

Satzschlußform:

Beispiele:

- (455) a. あの人が日本の着物を着たら、きっと美しかろう。¹⁹⁸

ano hito-ga nihon-no kimono-wo ki.t-ara, kitto
 jener Mensch-NOM Japan-GEN Kimono-ACC tragen-PCOND, bestimmt
utukusi-ka.r-ou
 schön-φ-VOLI

Sie wird bestimmt schön aussehen, wenn sie einen japanischen Kimono trägt.

- b. 真面目に働こうと努めた。¹⁹⁹

mazime-ni hatara.k-ou to tutome.t-a
 ernsthaft-ADV arbeiten-VOLI CIT bemühen-PAST

Er bemühte sich, ernsthaft zu arbeiten.

- c. 私は日本歴史を読もうと思う / 思います。²⁰⁰

watasi-ha nihon rekisi-wo yo.m-ou to omo.-u /
 ich-TOP japanische Geschichte-ACC lesen-VOLI CIT denken-FLEX /
omo..i-mas.-u
 denken-HON-FLEX

Ich glaube, ich werde (Bücher) über die japanische Geschichte lesen.

- d. 私が彼に話しましょう。²⁰¹

watasi-ga kare-ni hana.s.i-mas.-you
 ich-NOM er-DAT sprechen-HON-VOLI

Ich werde mit ihm sprechen.

- e. 映画に行きましょう。²⁰²

eiga-ni i.k.i-mas.-you
 Film-DIR gehen-HON-VOLI

Laßt und ins Kino gehen!

¹⁹⁷ JM, S. 79, 20-2.1.3. -Yoo 'Futur'.

¹⁹⁸ JM, S. 209, 30-3.1.1. c) 2.

¹⁹⁹ JM, S. 80, 20-2.1.3. c) 2.

²⁰⁰ DBJG, S. 569, KS(A).

²⁰¹ DBJG, S. 240, KS(A).

²⁰² DBJG, S. 240, KS(B).

Morphologie:

(456) a. *flexiv*[◇]

b. CONT $\left[\begin{array}{l} \text{future} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right]$

c. i. MAJ $v_{v,i\text{masu}}^{\text{cons}} \longrightarrow v$; PHON $\langle \text{you} \rangle$

ii. MAJ $v_c^{\text{cons}} \longrightarrow v$; PHON $\langle \text{ou} \rangle$

iii. MAJ $v_{i\text{suru}}^{\text{cons}} \longrightarrow v$; PHON $\langle \text{iyou} \rangle$

Da für Vermutungen oft die Verbindung eines Verbes im Präsens oder Perfekt mit dem Futur des Partikelverbes *da* verwendet wird, führen viele Grammatiken diese Verbindung als eigene Form auf²⁰³.

(457) 今夜は冷えこむだろう。 ²⁰⁴

koñya-ha hieko.m-u da.r-ou
 Heute Nacht-TOP abkühlen-FLEX sein-VOLI
 Heute Nacht wird es abkühlen.

Adnominalform:

Die Adnominalform des Futurs wird nur selten benutzt. Sie bildet sich auf die gleiche Weise wie die Adnominalform des Präsens- bzw. Perfekt-Flexives:

Beispiel:

(458) そんなこと、あろうはずがない。 ²⁰⁵

sonna koto, arou hazu-ga na.i
 solch eine Sache, sein-VOLI Möglichkeit-NOM sein-NEG-FLEX
 So etwas kann einfach nicht vorkommen.

Morphologie:

(459) f. MOD N

g. CONT $\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \left[\begin{array}{l} \text{future} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right]$

²⁰³ vgl. die Verbtafeln in B, die den Einträgen von Ishii et al. 1989 nachempfunden sind.

²⁰⁴ JM, S. 105, 20-5.1.2.-1.

²⁰⁵ JM, S. 80, 20-2.1.3. b) s..

Für die Beziehungen zwischen Relativsatz und Postcedent gibt es auch in diesem Fall drei alternative Möglichkeiten:

- (460) *h.* *i.* CONT $\left[\textcircled{1} \text{ cont}(\text{subj}) \right]$
ii. CONT $\left[\textcircled{1} \text{ cont}(\text{obj}) \right]$
iii. CONT $\left[\textcircled{1} \text{ cont}(\text{obj2}) \right]$

5.2.3.4 Imperativ (-IMP)

Mit dem *Imperativ* werden Aufforderungen, Befehle oder Bitten formuliert. Er wird durch zwei unterschiedliche Morpheme realisiert:

1. $-E^{206}$ bei Verben der Klasse v_c , v_{imasu} und v_{ikuru} ;
2. $-ro^{207}$ nur bei Verben der Klasse v_v und v_{isuru} .

Beispiele:

- (461) *a.* 座れ! ²⁰⁸
suwa.r-e
 setzen-IMP
 Setz dich!
- b.* お金を貯めろ! ²⁰⁹
o-kane-wo tame-ro
 HON-kane-ACC sparen-IMP
 Spar dein Geld!

Morphologie:

- (462) *a.* *flexiv*[◇]
- b.* CONT $\left[\begin{array}{l} \text{imperative} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{stem}) \end{array} \right]$
- c.* *i.* MAJ $v_{c,imasu}^{cons} \longrightarrow v$; PHON $\langle e \rangle$
ii. MAJ $v_{imasu}^{cons} \longrightarrow v$; PHON $\langle i \rangle$
iii. MAJ $v_v^{cons} \longrightarrow v$; PHON $\langle ro \rangle$

206 JM, S. 80, 20-2.1.4.-1. $-E$ 'Imperativ'.

207 JM, S. 86, 20-2.2.6. $-ro$ 'Imperativ'.

208 CJVG, S. 15.

209 CJVG, S. 15.

Manche Verben bilden den Imperativ abweichend von diesen Regeln. Solche unregelmäßigen Formen werden durch einen eigenen Eintrag ins Lexikon behandelt:

Ausnahmen:

	Verb		Imperativ	
(463) a.	くれる <i>kure-ru</i>	(mir) geben	くれ <i>kure</i>	<i>gib (mir)!</i>
b.	来る <i>ku-ru</i>	kommen	来い <i>koi</i>	<i>komm!</i>
c.	する <i>s-uru</i>	tun, machen	せよ, しろ <i>seyo siro</i>	<i>tu, mach!</i>

Einige Verben haben mehrere Imperativformen:

(464) a.	$v_{i\text{masu}}$:	くださいませ <i>kudasa.i-mas-e</i>	<i>geben sie bitte!</i>
		いらっしゃいまし <i>irassy.a.i-mas-i</i>	<i>willkommen! (Frauensprache)</i>
b.	$v_{c\text{nasaru}}$:	なさい ²¹⁰ <i>nasa-i</i>	<i>...bitte!</i>
		なされ ²¹¹ <i>nasa.r-e</i>	<i>...bitte!</i>

(464a) wird durch mehrere alternative Regeln modelliert (vgl. (462c-i) und (462c-ii)); die unregelmäßige Form なさい (*nasai*) — die erste Variante des Imperativs von なさる (*nasaru*) — wird durch einen eigenen Lexikoneintrag modelliert. Die Form なされ (*nasa.r-e*) hingegen entspricht der regelmäßigen Imperativform konsonantischer Verben.

Bei manchen Verben wird der Imperativ nicht benutzt.

Prohibitiv bzw. negierter Imperativ

Verbote, die negative Entsprechung des Imperativs, werden durch das Interjektionspartikel *na*²¹² ausgedrückt und sind somit nicht Thema der Morphologie:²¹³

²¹⁰ JM, S. 81, 20-2.1.4.-2.

²¹¹ CJVG, S. 188.

²¹² JM, S. 120, 21-2.2.5.-1.

²¹³ Zur Modellierung von Interjektionspartikeln vergleiche S. 129.

Beispiele:

(465) a. たばこを吸うな! ²¹⁴

tabako-wo su.-u na
Zigarette-ACC rauchen-FLEX NEGIMP
Nicht rauchen!

b. あんな男とは結婚するな! ²¹⁵

anna otoko-to-ha kekkoñ s.-uru na
jene Art von Mann-KOM-TOP Heirat machen-FLEX NEGIMP
Heirate keinen solchen Mann!

5.2.3.5 Partizip (-PART)

Durch die Bildung der *Partizipialform*²¹⁶ kann eine Verbal- oder Adjektivphrase einem anderen Verb oder Adjektiv adverbial²¹⁷ zugeordnet werden. Die semantische Beziehung zwischen Neben- und Matrixsatz bleibt dabei weitgehend unbestimmt und hängt vom syntaktischen und semantischen Kontext und der beschriebenen Situation ab. Unter anderem sind folgende Interpretationen möglich:

1. Zeitliche Abfolge von Geschehnissen, vgl. (466a);
2. Konjunktion zweier Aussagen, vgl. (466b) und (466c);
3. Kausale Beziehung zwischen Neben- und Matrixsatz, vgl. (466d);
4. Kontinuativ²¹⁸: als Attributsatz zum Hilfsverb *いる* (*iru*), vgl. (467);

usw.

In unserer semantischen Darstellung wird die Relation *participle* verwendet.

Beispiele:

(466) zur Bildung eines adverbial zugeordneten Nebensatzes:

²¹⁴ DBJG, S. 266, KS.

²¹⁵ DBJG, S. 266, ex. (c).

²¹⁶ JM, S. 84, 20-2.2.1. -*Te* 'Partizip' und S. 208, 30-2.4. -*kute* 'Partizip'.

²¹⁷ Warum Konstruktionen aus einem Verb in Partizipialform (V-*te*) und einem sich anschließenden Satz besser als adverbiale Zuordnung und nicht als Koordination analysiert werden, läßt sich in Iida et al. 1994 bzw. Manning et al. 1999 nachlesen.

²¹⁸ den Kontinuativ betreffend vergleiche auch

JM, S. 139, 2.3.1. vV+*Te*(=p) i.ru / or.u / irassyar.u und
DBJG, S. 155, iru² *いる*.

- a. ジムは日本へ行って勉強した / 勉強しました。 ²¹⁹

jimu-ha nihoñ-he itt-e beñkyous.it-a /
 Jim-TOP Japan-DIR gehen-PART Studium machen-PAST /
beñkyous..i-mas.it-a
 Studium machen-HON-PAST
 Jim fuhr nach Japan und studierte (dort).

- b. ここのステーキは安くておいしい / おいしいです。 ²²⁰

koko-no sutēki-ha yasu-kute oisi-i /
 hier-GEN Steak-TOP billig-PART wohlschmeckend-FLEX /
oisi-i des.-u
 wohlschmeckend-FLEX HON-FLEX
 Hier sind Steaks billig und schmecken gut.

- c. 私の父は先生で高校で英語を教えている / います。 ²²¹

watasi-no titi-ha señsei de kōkō-de eigo-wo
 Ich-GEN Vater-TOP Lehrer ist-PART Oberschule-LOC Englisch-ACC
osie.t-e i-ru / i-mas.-u
 unterrichten-PART KONT-FLEX / KONT-HON-FLEX
 Mein Vater ist Lehrer und unterrichtet Englisch an einer Oberschule.

- d. このアパートは広くていい / いいです。 ²²²

kono apāto-ha hiro-kute i-i / i-i des.-u
 Dieses Apartment-TOP groß-PART gut-FLEX / gut-FLEX HON-FLEX
 Dieses Apartment ist groß und gut.

(467) zur Bildung des Kontinuativs:

- a. 佐々木さんは酒を飲んでいる / います。 ²²³

sasaki-sañ-ha sake-wo no.ñd-e i-ru /
 Herr Sasaki-TOP Sake-ACC trinken-PART KONT-FLEX /
i-mas.-u
 KONT-HON-FLEX
 Herr Sasaki trinkt (gerade) Sake.

- b. 和江は新聞を読んでいる。 ²²⁴

kazue-ha sinbuñ-wo yo.ñd-e i-ru
 Kazue-TOP Zeitung-ACC lesen-PART KONT-FLEX
 Kazue liest (gerade) Zeitung.

219 DBJG, S. 464, KS(1).

220 DBJG, S. 464, KS(2).

221 DBJG, S. 464, KS(4).

222 vgl. DBJG, S. 464, KS(3).

223 DBJG, S. 155, KS.

224 DBJG, S. 155, ex. (a).

- c. このりんごはくさっている。 ²²⁵

kono riŋgo-ha kusa.tt-e i-ru
 dieser Apfel-TOP verfaulen-PART KONT-FLEX
 Dieser Apfel ist verfault.

- d. 私は鈴木さんを知っています。 ²²⁶

watasi-ha suzuki-saŋ-wo si.tt-e i-mas.-u
 ich-TOP Fräulein Suzuki-ACC kennen-PART KONT-HON-FLEX
 Ich kenne Fräulein Suzuki.

- e. 木田さんはみんなから尊敬されている。 ²²⁷

kida-saŋ-ha miŋna-kara soŋkeis.-are.t-e i-ru
 Herr Kida-TOP alle-REL Achtung tun-PASS-PART KONT-FLEX
 Herr Kida wird von allen geachtet.

Morphologie:

- (468) a. *flexiv*[◇]

- b. MOD Vcat

- c. CONT
$$\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont}^{228} \\ \textcircled{2} \left[\begin{array}{l} \text{participle} \\ \textcircled{1} \text{ cont}(\text{stem}) \end{array} \right] \end{array} \right]$$

- d. i. MAJ $v^{onbin} \longrightarrow v$; PHON $\langle e \rangle$
 ii. MAJ $adj^{onbin} \longrightarrow adj$; PHON $\langle kute \rangle$

Partizip + *kara*

Das Partizip kann auch zusammen mit dem Partikel *kara* (seit) adverbial benutzt werden. Wie die Übersetzung *seit* schon andeutet, wird dadurch eine zeitliche Abfolge (*Vorzeitigkeit*) ausgedrückt: das Ereignis, das durch den Matrixsatz beschrieben wird, findet nach dem des Nebensatzes statt.²²⁹ Als semantische Relation wird *since* verwendet.

Da sich die zeitliche Bedeutung²³⁰ von *kara* aus der Verwendung mit dem

²²⁵ DBJG, S. 155, ex. (b).

²²⁶ DBJG, S. 156, ex. (d).

²²⁷ DBJG, S. 367, notes 4. (5) b.

²²⁸ Das Argument ①, das Bezugswort des Adjunktes, wird durch das Head-Adjunkt-Schema instantiiert.

²²⁹ vgl. JM, S. 110, 20-5.5.2.1.-2.

²³⁰ Mit dem Präsens und dem Perfekt bezeichnet *kara* eine kausale Beziehung: vgl. JM, S. 110, 20-5.5.2.1.-1.

Partizip ergibt, wurde die Verbindung von Partizip + *kara* nicht syntaktisch modelliert, sondern als synthetische Form ins Lexikon aufgenommen.

Beispiel:

(469) 日本へ来ましてから一年たちました。²³¹

Nihoñ-he ki-mas.it-e kara iti neñ ta.t.i-mas.it-a.
 Japan-DIR kommen-HON-PART seit ein Jahr vergehen-HON-PAST
 Es ist ein Jahr vergangen, seitdem ich nach Japan gekommen bin.

Morphologie:

(470) a. *flexiv*[◇]

b. MOD Vcat

c. CONT $\left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ cont}^{232} \\ \textcircled{2} \left[\begin{array}{c} \text{since} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right]$

d. i. MAJ $v^{onbin} \rightarrow v$; PHON $\langle ekara \rangle$
 ii. MAJ $adj^{onbin} \rightarrow adj$; PHON $\langle kutekara \rangle$

5.2.3.6 Negiertes Partizip (-NEGPART)

Das *negierte Partizip*²³³ steht wie das Partizip adverbial und kann in den meisten Fällen durch “ohne zu tun” bzw. “ohne getan zu haben” übersetzt werden. (“tun” steht hier stellvertretend für die Bedeutung des jeweiligen Verbes.)

Beispiel:

(471) あの人は怒りもせず笑っていた。²³⁴

ano hito-ha okori-mo s-ezu wara.tt-e it-a
 jener Mensch-TOP sogar Ärger-ACC tun-NEGPART lachen-PART KONT-FLEX
 Er lachte, ohne sich darüber gar aufzuregen.

²³¹ JM, S. 97, 20-3.2.1.1.

²³² Das Argument ①, das Bezugswort des Adjunktes, wird durch das Head-Adjunkt-Schema instantiiert.

²³³ JM, S. 81, 20-2.1.5. -*Azu* ‘Negiertes Partizip’.

²³⁴ JM, S. 82, 20-2.1.5. b) 1.

Morphologie:

(472) a. *flexiv*[◇]

b. MOD Vcat

c. CONT $\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont}^{235} \\ \textcircled{2} \left[\begin{array}{l} \text{participle} \\ \textcircled{1} \left[\begin{array}{l} \text{not} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right] \end{array} \right]$

d. i. MAJ $v_{v,c,i_{kuru}}^a \longrightarrow v$; PHON $\langle zu \rangle^{236}$
 ii. MAJ $v_{i_{suru}}^{cv} \longrightarrow v$; PHON $\langle ezu \rangle$

Negiertes Partizip + *ni*

Das Partizip wird oft mit dem Partikel *ni* verwendet. Partikel, die sich mit Verbalphrasen zu einem adverbialen Komplex verbinden, werden durch unsere Grammatik nicht erfasst. Endung und Partikel werden deshalb zur Zeit als synthetische Form im Lexikon modelliert.

Beispiele:

(474) a. ナンシーはきのう朝ご飯を食べずに学校へ行った。²³⁷

235 Argument ①, durch das die Semantik des Bezugswortes bezeichnet wird, wird durch das Head-Adjunkt-Schema instantiiert.

236 Rickmeyer 1995 analysiert das Partizip von 来る (*kuru*; kommen) in *k + ozu*. Diese Segmentierung wird hier nicht verwendet, da sie gegen unsere Regel verstößt, Kanji als untrennbare Einheiten zu behandeln:

(473) Form	Stamm	Endung
Präsens	来 <i>ku</i>	る <i>ru</i>
negiertes Präsens	来 <i>ko</i>	ない <i>nai</i>
Honorativ Präsens	来 <i>ki</i>	ます <i>masu</i>

Das Kanji 来 verändert in Abhängigkeit der Verbform seine Aussprache. Um sowohl die Kana-Kanji-Schreibung als auch die Umschrift darstellen zu können, werden Kanji und Aussprache einander im Lexikon zugeordnet. Dieses ist aber nur dann ohne übermäßigen Aufwand möglich, wenn die Kanji — und damit auch ihre Umschrift — als untrennbar behandelt werden.

237 DBJG, S. 272, related expressions I. [1] a.

nañsii-ha kinou asagohañ-wo tabe-zuni gakkou-he
 Nancy-TOP gestern Frühstück-ACC essen-NEGPART Schule-DIR
itt-a
 gehen-PAST

Nancy ging gestern zur Schule, ohne Frühstück zu essen.

- b. 中田さんは大阪に行かずに京都に行った。 ²³⁸

nakada-saň-ha oosaka-ni i.k.a-zuni kyouto-ni itt-a
 Herr Nakada-TOP Ōsaka-DIR fahren-NEGPART Kyōto-DIR fahren-PAST
 Herr Nakada fuhr nach Kyōto, ohne nach Ōsaka zu fahren.

- c. 辞書を使わずに読んでください。 ²³⁹

jisyo-wo tuka..wa-zuni yo.ñd-e kudasai
 Wörterbuch-ACC benutzen-NEGPART lesen-PART bitte
 Bitte lesen sie es, ohne ein Wörterbuch zu benutzen.

Morphologie:

- (475) a. *flexiv*[◇]

- b. MOD Vcat

- c. CONT
$$\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont }^{240} \\ \textcircled{2} \left[\begin{array}{l} \text{participle} \\ \textcircled{1} \left[\begin{array}{l} \text{not} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right] \end{array} \right]$$

- d. i. MAJ $v_{v,C,i_{kuru}}^a \rightarrow v$; PHON $\langle zuni \rangle$
 ii. MAJ $v_{i_{suru}}^{CV} \rightarrow v$; PHON $\langle ezuni \rangle$

5.2.3.7 Konditional (-COND)

Das *Konditional*²⁴¹ wird adverbial einem anderen Verb oder Adjektiv zugeordnet und dient zur Kennzeichnung einer Bedingung, unter der die Aussage des Matrixsatzes wahr wird. Es teilt sich diese Funktion unter anderem mit dem Perfekt-Konditional (vgl. Abschnitt 5.2.3). Für den Bedeutungsunterschied bzw. die genauen Bedingungen zur Auswahl der passenden Konditionalform

²³⁸ DBJG, S. 272, related expressions I. [1] b.

²³⁹ DBJG, S. 272, related expressions I. [1] c.

²⁴⁰ Das Argument ②, das Bezugswort des Adjunktes, wird durch das Head-Adjunkt-Schema instantiiert.

²⁴¹ JM, S. 79, 20-2.1.2. *-Reba* ‘Konditional’ und S. 208, 30-2.5. *-kereba* ‘Konditional’

vergleiche die entsprechenden Einträge in Makino und Tsutsui 1986²⁴² und Rickmeyer 1995²⁴³

Beispiele:

- (476) a. 雨が降ればぬれます。 ²⁴⁴

ame-ga hu.r-eba nure-mas.-u
 Regen-NOM fallen-COND naß werden-HON-FLEX
 Wenn es regnet, wird man naß.

- b. これは松本先生に聞けば分かります。 ²⁴⁵

kore-ha matumoto-señsei-ni ki.k-eba waka.r.i-mas.-u
 dieses-TOP Prof. Matumoto-DAT fragen-COND verstehen-HON-FLEX
 Wenn Sie Prof. Matumoto fragen, werden sie es verstehen.

- c. 安ければ買います。 ²⁴⁶

yasu-kereba ka..i-mas.-u
 billig-COND kaufen-HON-FLEX
 Wenn es billig ist, werde ich es kaufen.

- d. もっと安ければ買いました。 ²⁴⁷

motto yasu-kereba ka..i-mas.it-a
 mehr billig-COND kaufen-HON-PAST
 Ich hätte es gekauft, wenn es billiger gewesen wäre.

Morphologie:

- (477) a. *flexiv*[◇]

- b. MOD Vcat

- c. CONT
$$\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ cont}^{248} \\ \textcircled{2} \left[\begin{array}{l} \text{conditional} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right]$$

- d. i. MAJ $v_{v,i_kuru}^{cons} \rightarrow v$; PHON $\langle reba \rangle$
 ii. MAJ $v_c^{cons} \rightarrow v$; PHON $\langle eba \rangle$
 iii. MAJ $v_{i_suru,masu}^{cons} \rightarrow v$; PHON $\langle ureba \rangle$
 iv. MAJ $adj^{cons} \rightarrow v$; PHON $\langle kereba \rangle$

242 S. 81 ff, *ba* ば (Konditional); S. 452 ff, $\sim tara$ たら (Konditional-Perfekt).

243 s. Fußnote 241 und Abschnitt 5.2.3

244 JM, S. 79, 20-2.1.2. c).

245 DBJG, S. 82, ex. (a).

246 DBJG, S. 82, ex. (c).

247 DBJG, S. 83, notes 5. (3).

5.2.3.8 Perfekt-Konditional (-PCOND)

Das *Perfekt-Konditional*²⁴⁹ steht adverbial und bezeichnet eine Bedingung, deren Abschluß das Geschehen zur Folge hat, das in der Matrixphrase beschrieben wird. Es kann auch zur Beschreibung einer zeitlichen Abfolge benutzt werden.

Für die Abgrenzung zum Konditional vergleiche Abschnitt 5.2.3.

Beispiele:

- (478) a. 山田さんが来たら私は帰る / 帰ります。²⁵⁰

yamada-sa^ñ-ga kit-ara watasi-ha kae.r-u
Herr Yamada-NOM kommen-PCOND ich-TOP nach Hause fahren-FLEX
/ kae.r.i-mas.-u
/ nach Hause fahren-HON-FLEX
Wenn Herr Yamada kommt, fahre ich nach Hause.

- b. 先生に聞いたらすぐ分かった。²⁵¹

se^ñsei-ni ki.it-ara sugu waka.tt-a
Lehrer-DAT fragen-PCOND sofort verstehen-PAST
Als ich meinen Lehrer fragte, verstand ich es sofort.

- c. おもしろかったら読みます。おもしろくなかったら読みません。²⁵²

omosiro-ka.tt-ara yo.m.i-mas.-u. omosiro-ku na-ka.tt-ara
interessant- ϕ -PCOND lesen-HON-FLEX interessant-ADV NEG- ϕ -PCOND
yo.m.i-mas.e-^ñ
lesen-HON-NEG
Wenn es interessant ist, lese ich es. Ist es jedoch nicht interessant, lese ich es nicht.

Morphologie:

248 Das Argument ①, enthält die Semantik des Bezugswortes und wird durch das Head-Adjunkt-Schema instantiiert.

249 JMS. 85, 20-2.2.3. -Tara, -Taraba 'Perfekt-Konditional'

250 DBJG, S. 452, KS.

251 DBJG, S. 452, ex. (a).

252 vgl. DBJG, S. 453, ex. (c).

- (479) a. *flexiv*[◇]
 b. MOD Vcat
 c. CONT $\left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ cont}^{253} \\ \textcircled{2} \left[\begin{array}{c} \text{perfect-conditional} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right] \end{array} \right]$
 d. i. MAJ $v^{onbin} \rightarrow v$; PHON $\langle ara \rangle$
 ii. MAJ $v^{onbin} \rightarrow v$; PHON $\langle araba \rangle$

5.2.3.9 Adverbial (-ADV)

Die *Adverbialform*²⁵⁴ dient — wie der Name schon sagt — der adverbialen Zuordnung: Sie gestattet es, eine Verbal- oder Adjektivalphrase zur Modifikation einer anderen Verbal- oder Adjektivalphrase zu verwenden. Die Funktion des Adverbials ist dabei vielfältig:

1. *Idiomatisierte Verbindungen*, z.B. *adverbial* + なる (*naru*), vgl. (480a);
2. Bildung der *Negation* von Adjektiven bzw. adjektivischen Verbformen vgl. (480b);
3. *Adverbale Verwendung* von Adjektiven, vgl. (480c) - (480f);
4. *Koordination* von Adjektiven bzw. Verben in adjektivischen Formen, vgl. (480g);

usw.

Die Bedeutung kann sehr unterschiedlich sein und muß aus dem semantisch-syntaktischen Kontext und der beschriebenen Situation erschlossen werden. In idiomatisierten Verbindungen, wie in 1 und 2, ist sie jedoch weitgehend festgelegt.

Beispiele:

- (480) a. 彼女は赤くなる。 ²⁵⁵

²⁵³ Das Argument ①, das Bezugswort des Adjunktes, wird durch das Head-Adjunkt-Schema instantiiert.

²⁵⁴ JM, S. 207, 30-2.2. *-ku* 'Adverbial'.

²⁵⁵ vgl. JM, S. 207, 30-2.2 a) 1.

- kanozyo-ha aka-ku na.r-u*
sie-TOP rot-ADV werden-FLEX
Sie wird rot.
- b. その本はおもしろくない。 ²⁵⁶
sono hoñ-ha omosiro-ku na-i
dieses Buch-TOP interessant-ADV NEG-FLEX
Dieses Buch ist uninteressant.
- c. 少し甘く点数をつけました。 ²⁵⁷
sukosi ama-ku teñsuu-wo take-mas.it-a
ein wenig großzügig-ADV Punkte-ACC anbringen-HON-PAST
Ich habe nicht so streng zensiert.
- d. 雨がすごくふっている。 ²⁵⁸
ame-ga sugo-ku hu.tt-e i-ru
Regen-NOM furchtbar stark-ADV fallen-PART KONT-FLEX
Es regnet in Strömen.
- e. パンを薄く切ってください。 ²⁵⁹
pañ-wo usu-ku ki.tt-e kudasai
Brot-ACC dünn-ADV schneiden-PART bitte
Schneide das Brot bitte dünn.
- f. ヘやは恐ろしくあつかった。 ²⁶⁰
heya-ha osorosi-ku atu-ka.tt-a
Zimmer-TOP furchtbar-ADV heiß- ϕ -PAST
Das Zimmer war furchtbar heiß.
- g. 夏は涼しく、冬はあたたかい所に住みたい。 ²⁶¹
natu-ha suzusi-ku, huyu-ha atataka-i tokoro-ni
Sommer-TOP kühl-ADV, Winter-TOP warm-FLEX Ort-LOC
su.m.i-ta-i
wohnen-VOLU-FLEX
Ich möchte dort wohnen, wo es im Sommer kühl und im Winter warm ist.

Morphologie:

-
- 256 vgl. JM, S. 207, 30-2.2 a) 1.
257 JM, S. 223, 32-2.1.
258 JM, S. 223, 32-2.1.
259 JM, S. 226, 32-2.3.2.
260 JM, S. 233, 33-2.
261 JM, S. 233, 33-2.

- (481) a. *flexiv*[◇]
 b. MOD Vcat
 c. CONT $\left[\begin{array}{l} \textit{adjunct} \\ \textcircled{1} \textit{cont}^{262} \\ \textcircled{2} \left[\begin{array}{l} \textit{adverbial} \\ \textcircled{1} \textbf{cont}(\textbf{stem}) \end{array} \right] \end{array} \right]$
 d. MAJ $\textit{adj}^v \longrightarrow \textit{adj};$ PHON $\langle ku \rangle$

262 Das Argument ①, die Semantik des Bezugswortes, wird durch das Head-Adjunkt-Schema instantiiert.

5.2.4 Suffixverben und -adjektive (*suffix*[◇])

Suffixverben und -adjektive sind das eigentliche Charakteristikum der japanischen Verbmorphologie. Sie stehen in vielfältiger Kombination und Anzahl zwischen initialem Verbstamm und abschließendem Flexiv und sind damit verantwortlich für den agglutinierenden Zug der japanischen Verbmorphologie. In diesem Kapitel werden die Suffixverben und -adjektive unseres Fragmentes aufgelistet und beschrieben.

5.2.4.1 Honorativ (-HON)

Das Japanische verfügt über zahlreiche lexikalische, morphologische und syntaktische Mittel, um Wertschätzungen auszudrücken und die relative soziale Stellung des Gesprächspartners bzw. der Person, über die gesprochen wird, zu würdigen. Kategorien zum Ausdruck von Unterschieden in der gesellschaftlichen Stellung, Achtung und Höflichkeit etc. spielen in der Morphologie eine ähnliche Rolle wie Numerus und Person im Deutschen.

Der *Honorativ*²⁶³ steht für eine der Möglichkeiten, Achtung vor dem Gesprächspartner auszudrücken. Er ist Standard der formalen, höflichen Sprechweise, wie sie beispielsweise im Gespräch mit Unbekannten in den meisten Fällen üblich ist:

Beispiele:

- (482) a. 日本へ来ましてから一年たちました。 ²⁶⁴

nihon-he ki-mas.it-ekara itineñ ta.t.i-mas.it-a
Japan-DIR kommen-HON-PART Seit ein Jahr vergehen-HON-PAST

Es ist ein Jahr vergangen, seitdem ich nach Japan gekommen bin.

- b. この漢字の書き方が分かりません。 ²⁶⁵

kono kañzi-no kakikata-ga waka.r.i-mas.e-ñ
dieses Kanji-GEN Schreibweise-NOM verstehen-HON-NEG

Ich weiß nicht, wie man dieses Kanji schreibt.

- c. 映画に行きましょう。 ²⁶⁶

eiga-ni i.k.i-mas.-you
Film-DIR gehen-HON-VOLI

Gehen wir ins Kino!

- d. この薬を飲めばよくなります。 ²⁶⁷

263 JM, S. 97, 20-3.2.1.1. -mas.u 'Honorativ'.

264 JM, S. 97, 20-3.2.1.1.

265 DBJG, S. 183, ex. (a).

266 DBJG, S. 240, KS(B).

kono kusuri-wo no.m-eba yo-ku na.r.i-mas.-u
 diese Medizin-ACC trinken-COND gut-ADV werden-HON-FLEX
 Wenn Du diese Medizin nimmst, wird es Dir besser gehen.

Der Honorativ wird durch das Verbsuffix *-mas.-u* gekennzeichnet und semantisch durch die Einbettung in eine Struktur des Typs *honorative* dargestellt.

Morphologie:

(483) a. *suffix*[◇]

b. CONT $\left[\begin{array}{l} \textit{honorative} \\ \textcircled{1} \textbf{cont(stem)} \end{array} \right]$

c. MAJ $v_{cda}^i \longrightarrow v_{imasu}^{cv}; \quad \text{PHON } \langle mas \rangle$

Der unregelmäßige Honorativ des Verbes *da* wird durch einen eigenen Lexikoneintrag modelliert.

Der Honorativ von Adjektiven wird syntaktisch durch das Nachstellen des Partikelverbes です (*desu*) — der honorativen Variante des Verbes だ (*da*) — gebildet:

Beispiele:

(484) a. この映画は面白い / 面白いです。 ²⁶⁸

kono eiga-ha omosiro-i / omosiro-i des.-u
 dieser Film-TOP interessant-FLEX / interessant-FLEX HON-FLEX
 Dieser Film ist interessant.

Dieses Nachstellen von です (*desu*) wird in unserem Modell durch die Einbettung des adjektivischen Satzes beschrieben:

Morphologie:

(485) a. *v-stem*[◇]

b. PHON $\langle des \rangle$

c. MAJ v_{imasu}^{cv}

d. SUBCAT $\{S[obl]^{obl}\}$

e. CONT $\left[\begin{array}{l} \textit{honorative} \\ \textcircled{1} \textbf{cont(obj)} \end{array} \right]$

267 DBJG, S. 81, KS.

268 JM, S. 214, 30-5.1.2.

5.2.4.2 Negation (-NEG)

Im Gegensatz zum Deutschen oder Englischen, wo die *Negation* von Verben meist auf syntaktischer Ebene ausgedrückt wird, bedient sich das Japanische zu ihrer Bildung morphologischer Mittel; die Negation wird durch das Suffixadjektiv *-na.i* bezeichnet²⁶⁹:

Beispiele:

(486) a. 本を見ないでください。 ²⁷⁰

hoñ-wo mi-na-i de kudasai
Buch-ACC sehen-NEG-FLEX PART bitte
Sehen Sie bitte nicht ins Buch.

b. 日本語を勉強しなければならない。 ²⁷¹

nihoñgo-wo beñkyou s.-ina-kereba na.r.a-na-i
Japanisch-ACC Studium machen-NEG-COND werden-NEG-FLEX
Du mußt Japanisch lernen.

c. 仕事が進まない。 ²⁷²

sigoto-ga susu.m.a-na-i
Arbeit-NOM fortschreiten-NEG-FLEX
Die Arbeit schreitet nicht fort.

Der Lexikoneintrag sieht folgendermaßen aus:

Morphologie:

(487) a. *suffix*[◇]

b. CONT $\left[\begin{array}{l} \text{not} \\ \text{① cont(stem)} \end{array} \right]$

c. i. MAJ $v_{v,c \setminus aru, i_kuru}^a \longrightarrow adj^v$; PHON $\langle na \rangle$
ii. MAJ $v_{i_suru}^a \longrightarrow adj^v$; PHON $\langle ina \rangle$

ない (*na.i*), die Negation von ある (*a.r-u*), wird durch einen eigenen Lexikoneintrag realisiert.²⁷³

269 JM, S. 96, 20-3.1.2. Suffixadjektive: V+a: *-Ana.i* 'Negation'.

270 JM, S. 97, 20-3.1.2. b).

271 vgl. JM, S. 97, 20-3.1.2. c).

272 IJDW, vgl. Eintrag für 進む (*susumu*; fortschreiten).

273 s. (398), S. 212.

Die Negation von Adjektiven

Adjektive werden durch adverbale Zuordnung zu *ない* (*na.i*), der Negation von *ある* (*a.r-u*), negiert.²⁷⁴

Beispiele:

- (490) a. おもしろかったら読めます。おもしろくなかったら読みません。²⁷⁶

omosiro-ka.tt-ara yo.m.i-mas.-u. omosiro-ku na-ka.tt-ara
 interessant- ϕ -PCOND lesen-HON-FLEX interessant-ADV NEG- ϕ -PCOND
yo.m.i-mas.e-ñ
 lesen-HON-NEG

Wenn es interessant ist, lese ich es. Ist es jedoch nicht interessant, lese ich es nicht.

- b. きのう面白くない映画を見ました。

kinou omosiro-ku na-i eiga-wo mi-mas.it-a
 gestern interessant-ADV NEG-FLEX Film-ACC sehen-HON-PAST
 Gestern habe ich einen uninteressanten Film gesehen.

Zur Bildung der Adverbialform vgl. Abschnitt “Adverbial”, S. 245.

Die Negation von Verben der Kategorie *v_{i masu}*

Die Honorativform *-masu* wird nach einem alten Flexionsparadigma gebildet²⁷⁷. Es handelt sich dabei nicht — wie im Falle von *Ana.i* — um ein Suffixadjektiv, sondern um ein Flexiv:

274 Die Negation von Adjektiven läßt sich auch als synthetische Form darstellen:

- (488) c. iii. MAJ $adj^a \rightarrow adj^v$; PHON $\langle kuna \rangle$

Ein solcher Eintrag würde jedoch zu zwei möglichen Analysen führen: 1. als synthetische Form und 2. als adverbial stehende Phrase zum Hilfsverb *ない* (*nai*; Negation).

Ein weiteres Argument gegen eine synthetische Form ist der mögliche Einschub anderer Elemente zwischen Adjektiv und *ない* (*nai*):

- (489) 痛くもかゆくもない。²⁷⁵
ita-ku-mo kayu-ku-mo na.i
 Es kratzt und juckt mich nicht.

275 vgl. JM, S. 233, 33-2.

276 vgl. DBJG, S. 453, ex. (c).

277 JM, S. 96, 20-3.1.2. *-An.u* ‘Negation’.

Beispiele:

- (491) a. おもしろかったら読みます。おもしろくなかったら読みません。 ²⁷⁸

omosiro-ka.tt-ara yo.m.i-mas.-u. omosiro-ku na-ka.tt-ara
 interessant- ϕ -PCOND lesen-HON-FLEX interessant-ADV NEG- ϕ -PCOND
yo.m.i-mas.e-ñ
 lesen-HON-NEG

Wenn es interessant ist, lese ich es. Ist es jedoch nicht interessant, lese ich es nicht.

Der Lexikoneintrag sieht folgendermaßen aus:

Morphologie:

- (492) a. *flexiv*[◇]

b. CONT $\left[\begin{array}{l} \text{not} \\ \textcircled{1} \text{ cont(stem)} \end{array} \right]$

c. MAJ $v_{imasu}^a \longrightarrow v$; PHON $\langle \hat{n} \rangle$

5.2.4.3 Voluntativ (-VOLU)

Durch den *Voluntativ*²⁷⁹ kann der Wunsch nach einer eigenen Tat, der Tat eines anderen oder nach dem Eintreffen eines Zustandes ausgedrückt werden:

Beispiele:

- (493) a. 私は日本へ行きたい / 行きたいです。 ²⁸⁰

watasi-ha nihon-he i.k.i-ta-i / i.k.i-ta-i
 ich-TOP Japan-DIR fahren-VOLU-FLEX / fahren-VOLU-FLEX
des.-u
 HON-FLEX

Ich möchte nach Japan fahren.

- b. 私は先生にこの絵をほめられたい。 ²⁸¹

watasi-ha señsei-ni kono e-wo home-rare-ta-i
 ich-TOP Lehrer-DAT dieses Bild-ACC loben-PASS-VOLU-FLEX
 Ich möchte vom Lehrer für dieses Bild gelobt werden.

- c. 和男はとても行きたかった。 ²⁸²

²⁷⁸ vgl. DBJG, S. 453, ex. (c).

²⁷⁹ JM, S. 98, 20-3.2.2.2. *-ta.i* ‘Voluntativ’.

²⁸⁰ DBJG, S. 442, KS(A).

²⁸¹ DBJG, S. 444, notes 2. (B) (7).

²⁸² DBJG, S. 443, notes 1. (1).

kazuo-ha totemo i.k.i-ta-ka.tt-a
 Kazuo-TOP unbedingt gehen-VOLU- ϕ -PAST
 Kazuo wollte unbedingt gehen.

d. 和男は早く行きたかった。

kazuo-ha hayaku i.k.i-ta-ka.tt-a
 Kazuo-TOP schnell gehen-VOLU- ϕ -PAST
 Kazuo wollte schnell gehen.

Adverben zu einem voluntativen Verb können sich sowohl auf das zugrunde liegende Verb als auch auf den Akt des *Wünschens* beziehen. *-ta.i* wird daher mit [ADJUNCT +] markiert, so daß die lexikalisch-morphologische Regel zur Adjunkteinführung auch auf die semantische Struktur des Voluntativs angewendet werden kann.

Semantisch wird der Voluntativ durch eine Struktur des Typs *want* beschrieben, deren erstes Argument den Wünschenden bezeichnet:

Morphologie:

(494) a. *suffix*[◇]

b. CONT $\left[\begin{array}{l} \textit{want} \\ \textcircled{1} \textbf{cont(subj)} \\ \textcircled{2} \textbf{cont(stem)} \end{array} \right]$

c. ADJUNCT +

d. MAJ $v^i \rightarrow adj^v$; PHON $\langle ta \rangle$

Der Voluntativ kann mit einer Änderung der Kasusreaktion einhergehen:²⁸³

Beispiele:

(495) a. 僕は今ピザを/が食べたい / 食べたいです²⁸⁴

boku-ha ima piza-wo/-ga tabe-ta-i / tabe-ta-i
 ich-TOP jetzt Pizza-ACC/-NOM essen-VOLU-FLEX / essen-VOLU-FLEX
des.-u
 HON-FLEX
 Ich möchte jetzt Pizza essen.

Wir beschreiben diese Änderung der Valenz durch die folgende fakultative Modifikation des SUBCAT-Wertes:

²⁸³ vgl. JM, S. 98, 20-3.2.2.2. *-ta.i* 'Voluntativ'.

²⁸⁴ DBJG, S. 442, KS(B).

Morphologie:

$$(496) \text{ e. SUBCAT } \left\{ \text{PP}[\text{obj}, \text{wo}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\} \longrightarrow \left\{ \text{PP}[\text{obj}, \text{ga}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\}$$

5.2.4.4 Kausativ (-CAUS)

Situationen, in denen eine Person eine andere veranlaßt, etwas zu tun, gehören zum Repertoire menschlicher Grunderfahrungen. In der Linguistik werden Konstruktionen mit entsprechendem Inhalt als *Kausativ* bezeichnet. Während im Deutschen der Kausativ durch syntaktische Einbettung ausgedrückt wird, kann er im Japanischen morphologisch durch die Suffigierung eines Kausativmorphems gebildet werden.

Beispiele:

(497) a. i. 健は空港まで来た。

keñ-ha kuukou-made kit-a
Ken-TOP Flughafen-DIR kommen-PAST
Ken kam zum Flughafen.

ii. 直美は健を空港まで来させた。 ²⁸⁵

naomi-ha keñ-wo kuukou-made ko-sase.t-a
Naomi-TOP Ken-ACC Flughafen-DIR kommen-CAUS-PAST
Naomi ließ Ken zum Flughafen kommen.

b. i. 直美は富雄を訪ねた。

naomi-ha tomio-wo tazune.t-a
Naomi-TOP Tomio-ACC besuchen-PAST
Naomi besuchte Tomio.

ii. 健は直美に富雄を訪ねさせた。 ²⁸⁶

keñ-ha naomi-ni tomio-wo tazune-sase.t-a
Ken-TOP Naomi-DAT Tomio-ACC besuchen-CAUS-PAST
Ken veranlasste Naomi, Tomio zu besuchen.

c. i. 直美が手紙を毎日書く約束をした。 ²⁸⁷

naomi-ga tegami-wo mainiti ka.k-u yakusoku-wo
Naomi-NOM Brief-ACC täglich schreiben-FLEX Versprechen-ACC
s.it-a
machen-PAST

Naomi versprach, jeden Tag einen Brief zu schreiben.

ii. 健が直美に手紙を毎日書く約束をさせた。 ²⁸⁸

285 JPSG, S. 51, (3.58)a.

286 JPSG, S. 52, (3.58)b.

287 JPSG, S. 53, (3.63).

288 JPSG, S. 53, (3.62).

keñ-ga naomi-ni tegami-wo mainiti ka.k-u
 Ken-NOM Naomi-DAT Brief-ACC täglich schreiben-FLEX
yakusoku-wo s.-ase.t-a
 Versprechen-ACC machen-CAUS-PAST
 Ken ließ Naomi versprechen, jeden Tag einen Brief zu schreiben.

Morphologie:

(498) a. *suffix*[◇]

b. SUBCAT $\left\{ \text{PP}[\text{subj}]:\boxed{1}^{\boxed{2}}, \dots \right\}$
 $\longrightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]^{\text{opt}}, \text{PP}[\text{obj}, \text{ni} \vee \text{wo}]:\boxed{1}^{\boxed{2}}, \dots \right\}$

c. CONT $\left[\begin{array}{l} \text{cause} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{obj}) = \text{cont}(\text{subj}(\text{stem}))^{289} \\ \textcircled{3} \text{ cont}(\text{stem}) \end{array} \right]$

d. ADJUNCT +

e. i. MAJ $v_{v, i_{kuru}}^a \longrightarrow v_v^{comb}; \quad \text{PHON} \langle sase \rangle$
 ii. MAJ $v_c^a \longrightarrow v_v^{comb}; \quad \text{PHON} \langle se \rangle$
 iii. MAJ $v_{i_{suru}}^a \longrightarrow v_v^{comb}; \quad \text{PHON} \langle ase \rangle$

Für einige Verben läßt sich eine Variante des Kausatives nach den folgenden Regeln bilden:

Morphologie:

(500) e. i. MAJ $v_{v, i_{kuru}}^a \longrightarrow v_{c_s}^{cv}; \quad \text{PHON} \langle sa \rangle$
 ii. MAJ $v_c^a \longrightarrow v_{c_s}^{cv}; \quad \text{PHON} \langle \phi \rangle$

Die morphologischen, syntaktischen und semantischen Mechanismen, die bei der Bildung des Kausatives auftreten, lassen sich am besten anhand eines Beispiels erklären:

289 Die Abkürzung **cont(subj(stem))** ist folgendermaßen definiert:

(499) Abkürzung: abgekürzte AVM:

$$\left[\dots | \text{X cont}(\text{subj}(\text{stem})) \right] \left[\dots | \text{X } \boxed{1} \left[\dots | \text{STEM } \left[\dots | \text{SUBCAT } \left\{ \text{PP}[\text{subj}]:\boxed{1}, \dots \right\} \right] \right] \right]$$

mit $\text{X}:\boxed{1}$ für $\boxed{1} = \text{cont}(\text{X})$ in Subkategorisierungskontexten.

(501) 健は手紙を書く。

keñ-ha tegami-wo ka.k-u
 Ken-TOP Brief-ACC schreiben-FLEX
 Ken schreibt einen Brief.

Das Verb 書く (*kaku*; schreiben) subkategorisiert in der hier vorgestellten Bedeutung für zwei Postpositionalphrasen: die Subjektphrase PP[*ga*] (健, *Ken*), durch die der *Schreibende* ausgedrückt wird, und die Objektphrase PP[*wo*] (手紙, *tegami*; Brief), die das Objekt des Schreibens bezeichnet.

Aus dem Verb 書く (*kaku*; schreiben) läßt sich durch Einfügen eines Kausativmorphems das kausativische Verb 書かせる (*kakaseru*) mit der Bedeutung *zum Schreiben veranlassen* ableiten:

(502) 直美は健に手紙を書かせる。

naomi-ha keñ-ni tegami-wo ka.k.a-se-ru
 Naomi-TOP Ken-DIR Brief-ACC schreiben-CAUS-FLEX
 Naomi läßt Ken einen Brief schreiben.

(502) ist die kausativische Variante des Satzes (501). Seine Bedeutung läßt sich im Deutschen nur durch syntaktische Einbettung umschreiben: “*Naomi läßt Ken einen Brief schreiben*” oder “*Naomi veranlaßt Ken, einen Brief zu schreiben*”.

Betrachtet man die ursprüngliche Verbbedeutung und das *Veranlassen* als semantische Primitiven, läßt sich der Kausativ durch Einbettung der Verbsemantik in eine Kausativstruktur darstellen. Wenn wir für den Satz (501) die folgende Semantik annehmen:

(503)
$$\left[\begin{array}{l} \textit{write} \\ \textcircled{1} \textit{ ken} \\ \textcircled{2} \textit{ letter} \end{array} \right]$$

entspricht der kausativen Variante (502) damit folgende Semantik:

(504)
$$\left[\begin{array}{l} \textit{cause} \\ \textcircled{1} \textit{ naomi} \\ \textcircled{2} \boxed{\textit{ ken}} \\ \textcircled{3} \left[\begin{array}{l} \textit{write} \\ \textcircled{1} \boxed{} \\ \textcircled{2} \textit{ letter} \end{array} \right] \end{array} \right]$$

Neben die Objekte des zugrunde liegenden Verbes tritt ein weiteres Objekt zur Beschreibung des *Veranlassenden*. Dieses erscheint in der Oberflächenstruktur als Nominativ (Partikel *ga*) und erzwingt eine Verschiebung des Oberflächenkasus des ursprünglichen Subjekts: In unserem Beispiel erscheint das unterliegende Subjekt des Verbs in der Oberflächenstruktur als indirektes Objekt, gekennzeichnet durch das Partikel *ni*. Zusätzlich zu der Veränderung der morphologischen *Form* (498*e*) und der *Semantik* (498*c*) wird also die *Valenz* (498*b*) des Verbs modifiziert.

(498*b*) beschreibt die Ableitung der neuen Subkategorisierungsmenge: Das Verbsubjekt wird durch ein neues Objekt — eine PP[*obj, ni*] oder eine PP[*obj, wo*] — ersetzt. Der Wert des CONT-Merkmals des ursprünglichen Subjekts wird mit dem CONT-Wert der neuen PP koindiziert (vgl. (498*c*)). Die Wahl der Postposition kann zu Bedeutungsunterschieden führen und wird durch Constraints eingegrenzt; beispielsweise darf *wo* nur dann verwendet werden, wenn kein anderes Objekt diese Postposition erfordert. Diese Constraints werden hier jedoch nicht thematisiert.²⁹⁰

Die Spezifikation (498*d*) ermöglicht, wie in Abschnitt 4.3 “*Lexikalisch-morphologische Regeln*”, S. 179 beschrieben, die Einführung von Adverbien. Adverbien können also sowohl die Bedeutung modifizieren, die durch den Verbstamm ausgedrückt wird, als auch das kausativische “Veranlassen”:

- (505) a. 健が一人で直美に本を読ませた。²⁹¹

keñ-ga hitori-de naomi-ni hoñ-wo yo.m.a-se.t-a
 Ken-NOM alleine-ADV Naomi-DAT Buch-ACC lesen-CAUS-PAST
 Ken brachte Naomi alleine dazu, das Buch zu lesen. /
 Ken brachte Naomi dazu, das Buch alleine zu lesen.

- b. 健が黙って直美を座らせた。²⁹²

keñ-ga damatte naomi-wo suwa.r.a-se.t-a
 Ken-NOM schweigend Naomi-ACC setzen-CAUS-PAST
 Ken brachte Naomi schweigend dazu, sich zu setzen. /
 Ken brachte Naomi dazu, sich schweigend zu setzen.

- c. 健が自分のペンで直美に作文を書かせた。²⁹³

keñ-ga zibuñ-no peñ-de naomi-ni sakubun-wo
 Ken-NOM selbst-GEN Füller-ADV Naomi-DAT Aufsatz-ACC
ka.k.a-se.t-a
 schreiben-CAUS-PAST
 Ken brachte Naomi mit seinem Füller dazu, einen Aufsatz zu

²⁹⁰ vgl. beispielsweise Makino und Tsutsui 1986, S. 387, *saseru* させる

²⁹¹ LIJC, S. 15, (24)a.

²⁹² LIJC, S. 15, (24)b.

²⁹³ LIJC, S. 15, (24)c.

schreiben. / Ken brachte Naomi dazu, mit ihrem Füller einen Aufsatz zu schreiben.

Die *Kausativ-Konstruktion*²⁹⁴ gab und gibt Anlaß zu Kontroversen innerhalb der Linguistik: Während das semantische Verhalten und Phänomene wie die Bindung von Anaphern und Pronomina, der Skopus des Honorativs etc. für eine Modellierung durch Satzeinbettung auf syntaktischer Ebene sprechen, legt das morphologische und syntaktische Verhalten — Kasuszuweisung, Wortstellung etc. — eine Behandlung auf morphologischer Ebene nahe. Eine Übersicht der Argumente beider Seiten läßt sich in Gunji 1996b nachlesen: Manning et al. 1996 faßt die Argumente für eine Behandlung im Rahmen der Morphologie zusammen, während Gunji 1996a mit ähnlichen Argumenten für eine syntaktische Modellierung des Kausativs plädiert.²⁹⁵

Beide Ansätze haben eine strategische Gemeinsamkeit: die “morphologische” bzw. “syntaktische” Darstellung wird derart erweitert, daß auch auf die relevanten Informationen der jeweils anderen Beschreibungsebene zugegriffen werden kann. So sind die für die Bindungstheorie erforderlichen syntaktischen Informationen zur Konstituentenstruktur in dem morphologischen Ansatz als “hierarchische Argumentstruktur” auch auf morphologischer Ebene zugänglich; das “morphologische Prinzip” des syntaktischen Ansatzes hingegen gestattet, die Klammerung der Morpheme durch die Konstituentenstruktur von der linearen Anordnung der Morpheme zu trennen, so daß Morpheme, die in der Konstituentenstruktur getrennt erscheinen, morphologisch gesehen eine Einheit bilden können.

Beide Ansätze weichen also die Grenze zwischen Morphologie und Syntax auf und ähneln sich — von einem abstrakteren Standpunkt aus gesehen — damit eher, als daß sie sich unterscheiden.

Probleme wie die Anaphernbindung oder die Modellierung der Mechanismen zum Ausdruck der Höflichkeit im Japanischen liegen außerhalb unserer Fragestellung. Damit können wir den Schwierigkeiten, die bei einer morphologischen Behandlung des Kausativs auftreten, weitgehend ausweichen. In unserem Modell wird deshalb — der morphologischen Ausrichtung entsprechend — der morphologische Ansatz verfolgt. Letztendlich sind jedoch beide Ansätze gleichwertig und gestatten — jeder auf seine Art — die Lösung der beschriebenen Probleme.

²⁹⁴ JM, S. 94, 20-3.1.1.3. ‘Kausativ’.

²⁹⁵ Beide Aufsätze wurden an verschiedener Stelle (und Manning et al. 1996 auch in verschiedenen Versionen) veröffentlicht: Manning et al. 1996 erschien auch als Iida et al. 1994 und Manning et al. 1999; Gunji 1996a wurde auch als Gunji 1999 veröffentlicht.

5.2.4.5 Passiv (-PASS)

Im Japanischen gibt es zwei verschiedene Arten des Passivs²⁹⁶: den *direkten* und den *indirekten Passiv*. Der *direkte Passiv* ähnelt dem Passiv des Deutschen oder Englischen: Das direkte oder indirekte Objekt des unterliegenden Verbs erscheint im Passivsatz als Subjekt, während das unterliegende Subjekt zum indirekten Objekt des passivischen Satzes wird:

Beispiele:

- (506) a. i. 花子は一郎をだました / だましました。²⁹⁷
hanako-ha itirou-wo dama.sit-a / dama.s.i-mas.it-a
 Hanako-TOP Ichirō-ACC täuschen-PAST / täuschen-HON-PAST
 Hanako täuschte Ichirō.
- ii. 一郎は花子にだまされた / だまされました。²⁹⁸
itirou-ha hanako-ni dama.s.a-re.t-a /
 Ichirō-TOP Hanako-DAT täuschen-PASS-PAST /
dama.s.a-re-mas.it-a
 täuschen-PASS-HON-PAST
 Ichirō wurde von Hanako getäuscht.
- b. i. ジョンは先生に質問をした。²⁹⁹
joñ-ha señsei-ni situmoñ-wo s.it-a
 John-TOP Lehrer-DAT Frage-ACC machen-PAST
 John stellte dem Lehrer eine Frage.
- ii. 先生はジョンに質問をされた。³⁰⁰
señsei-ha joñ-ni situmoñ-wo s.-are.t-a
 Lehrer-TOP John-DAT Frage-ACC machen-PASS-PAST
 Dem Lehrer wurde von John eine Frage gestellt.

Der *indirekte Passiv* wird auch als “*Betroffenheitspassiv*” bezeichnet: Dem Subkategorisierungsrahmen wird ein weiteres Objekt hinzugefügt, das in den meisten Fällen ein Lebewesen bezeichnet, welches durch die Verbhandlung in einer oft unangenehmen Weise “*betroffen*” wird:

Beispiele:

- (507) i. 弟はケーキを食べた / 食べました。³⁰¹

296 JM, S. 92, 20-3.1.1.1. -*Rare.ru* ‘Passiv’.

297 vgl. DBJG, S. 366, notes 1. (1).

298 DBJG, S. 364, KS(A).

299 DBJG, S. 366, notes 3. (2)a.

300 DBJG, S. 366, notes 3. (2)b.

301 DBJG, S. 364, KS(C).

otouto-ha kēki-wo tabe.t-a / tabe-mas.it-a
 jüngerer Bruder-TOP Kuchen-ACC essen-PAST / essen-HON-PAST
 Mein jüngerer Bruder aß den Kuchen.

ii. 私は弟にケーキを食べられた / 食べられました。³⁰²

watasi-ha otouto-ni kēki-wo tabe-rare.t-a /
 ich-TOP jüngerer Bruder-DAT Kuchen-ACC essen-PASS-PAST /
tabe-rare-mas.it-a
 essen-PASS-HON-PAST
 Mein jüngerer Bruder aß den Kuchen (worüber ich mich nicht freute).
 / “Mein jüngerer Bruder aß mir den Kuchen weg.”

In diesem Fall wird das unterliegende Subjekt des Verbes zu einer PP[*ni*], während das Oberflächensubjekt das neu eingeführte Objekt — den Betroffenen der Handlung — beschreibt.

Direkter Passiv

In (506) wird das unterliegende Subjekt des Verbs mit Hilfe der Postposition に (*ni*) ausgedrückt. Unter bestimmten Umständen können jedoch auch die Postpositionen によって (*ni yotte*) und から (*kara*) Verwendung finden.³⁰³

Beispiele:

(508) a. i. 学生は私に日本の大学のことを聞いた。

gakusei-ha watasi-ni nihon-no daigaku-no
 Student-TOP ich-DAT Japan-GEN Universität-GEN
koto-wo ki.it-a
 Angelegenheit-ACC fragen-PAST
 Die Studenten befragten mich über japanische Universitäten.

ii. 私は学生から日本の大学のことを聞かれた。³⁰⁴

watasi-ha gakusei-kara nihon-no daigaku-no
 ich-TOP Student-REL Japan-GEN Universität-GEN
koto-wo ki.k.a-re.t-a
 Angelegenheit-ACC fragen-PASS-PAST
 Ich wurde von den Studenten über japanische Universitäten
 befragt.

b. この絵はピカソによってかかれた。³⁰⁵

kono e-ha pikaso-niyotte ka.k.a-re.t-a
 dieses Bild-TOP Picasso-REL malen-PASS-PAST

³⁰² DBJG, S. 364, KS(C).

³⁰³ Für die Bedingungen ihrer Anwendung vergleiche DBJG, S. 366, *rareru* られる, Notes 4.

³⁰⁴ DBJG, S. 367, notes 4. (5)a.

³⁰⁵ DBJG, S. 366, notes 4. (3)a.

Dieses Bild wurde von Picasso gemalt.

- c. 電話はベルによって発明された。³⁰⁶

deñwa-ha beru-niyotte hatumeis.-are.t-a
Telephon-TOP Bell-REL erfinden-PASS-PAST
Das Telephon wurde von Bell erfunden.

- d. 木田さんはみんなから尊敬されている。³⁰⁷

kida sañ-ha miñna-kara soñkeis.-are.t-e i-ru
Herr Kida-TOP alle-REL respektieren-PASS-PART KONT-FLEX
Herr Kida wird von allen respektiert.

Das Oberflächensubjekt eines passivischen Satzes entsteht in den meisten Fällen aus einem unterliegenden direkten Objekt, also einer PP[*wo*]. In selteneren Fällen kann es jedoch auch aus einem indirekten Objekt (PP[*ni*]) entstehen:

Beispiele:

- (509) a. i. 先生が生徒に成績が悪いと言ったこと。³⁰⁸

señsei-ga seito-ni seiseki-ga waru-i to
Lehrer-NOM Schüler-DAT Ergebnis-NOM schlecht-FLEX CIT
i.tt-a koto
sagen-PAST Sachverhalt

[...] daß der Lehrer dem Schüler gesagt hat, er habe schlechte Zensuren.

- ii. 生徒が先生に成績が悪いと言われたこと。³⁰⁹

seito-ga señsei-ni seiseki-ga waru-i to
Schüler-NOM Lehrer-DAT Ergebnis-NOM schlecht-FLEX CIT
i..wa-re.t-a koto
sagen-PASS-PAST Sachverhalt

[...] daß dem Schüler vom Lehrer gesagt worden ist, er habe schlechte Zensuren.

- b. i. 裁判所がスサンに無罪を言い渡した。³¹⁰

saibansyo-ga susañ-ni muzai-wo iiwata.sit-a
Gericht-NOM Susan-DAT Unschuld-ACC urteilen-PAST
Das Gericht sprach Susan frei.

- ii. スサンが裁判所に無罪を言い渡された。³¹¹

306 DBJG, S. 366, notes 4. (3)b.

307 DBJG, S. 367, notes 4. (5)b.

308 JM, S. 93, 3.1.1.1. b) 1.4.

309 JM, S. 93, 3.1.1.1. b) 1.4.

310 JPSG, S. 64, (3.84).

311 JPSG, S. 63, (3.83)a.

susān-ga saibansyo-ni muzai-wo iiwata.s.a-re.t-a
 Susan-NOM Gericht-DAT Unschuld-ACC urteilen-PASS-PAST
 Susan wurde vom Gericht freigesprochen.

Die Modifikation des Subkategorisierungsrahmens kann damit folgendermaßen zusammengefaßt werden:

$$(510) \text{ SUBCAT } \left\{ \text{PP}[\text{subj}]:\boxed{2}, \text{PP}[\text{obj} \vee \text{obj}2, \text{wo} \vee \text{ni}]:\boxed{4}, \dots \right\} \\ \longrightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]:\boxed{3}, \text{PP}[\text{obj}, \text{ni} \vee \text{niyotte} \vee \text{kara}]:\boxed{1}, \dots \right\}$$

Wie im Deutschen unterscheiden sich Aktiv- und Passivsatz kaum in ihrer Bedeutung. In der hier dargestellten Modellierung des direkten Passivs entspricht die Semantik daher der Semantik des entsprechenden aktivischen Satzes. Modifiziert werden folglich nur die Merkmale PHON, MAJ und SUBCAT:

Morphologie:

$$(511) \text{ a. } \text{suffix}^\diamond \\ \text{b. SUBCAT } \left\{ \text{PP}[\text{subj}]:\boxed{2}, \text{PP}[\text{obj} \vee \text{obj}2, \text{wo} \vee \text{ni}]:\boxed{4}, \dots \right\} \\ \longrightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]:\boxed{3}, \text{PP}[\text{obj}, \text{ni} \vee \text{niyotte} \vee \text{kara}]:\boxed{1}, \dots \right\} \\ \text{c. i. MAJ } v_{v, i_{kuru}}^a \longrightarrow v_v^{comb}; \quad \text{PHON } \langle rare \rangle \\ \text{ii. MAJ } v_c^a \longrightarrow v_v^{comb}; \quad \text{PHON } \langle re \rangle \\ \text{iii. MAJ } v_{i_{suru}}^a \longrightarrow v_v^{comb}; \quad \text{PHON } \langle are \rangle$$

Indirekter Passiv

Beispiele:

(512) i. フレッドは夜おそくアパートに来た / 来ました。

fureddo-ha yoru osoku apāto-ni kit-a /
 Fred-TOP Nacht spät Apartment-DIR kommen-PAST /
ki-mas.it-a
 kommen-HON-PAST
 Fred kam spät in der Nacht ins Apartment.

ii. ジェーンはフレッドに夜おそくアパートに来られた / 来られました。 ³¹²

312 DBJG, S. 364, KS(B).

jēñ-ha fureddo-ni yoru osoku apāto-ni ko-rare.t-a
 Jane-TOP Fred-DAT Nacht spät Apartment-DIR kommen-PASS-PAST
 / *ko-rare-mas.it-a*
 / kommen-PASS-HON-PAST
 Fred kam spät in der Nacht in Janes Apartment (worüber sich Jane nicht freute).

Die Verbformen des indirekten Passivs unterscheiden sich morphologisch nicht von denen des direkten Passivs. Im Gegensatz zu letzterem vergrößert sich jedoch beim indirekten Passiv die Menge der Objekte um eine PP, die den von der Handlung Betroffenen bezeichnet:

$$(513) \quad a. \quad \text{SUBCAT} \quad \left\{ \text{PP}[\text{subj}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\} \\
\rightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]^{\text{opt}}, \text{PP}[\text{obj}, \text{ni}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\} \\
\text{CONT} \quad \left[\begin{array}{l} \text{be-affected-by} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{stem}) \end{array} \right]$$

Die semantische Modifikation wird durch Einbettung der Verbsemantik in eine Struktur des Types *be-affected-by* beschrieben. Das erste Argument dieser Struktur wird durch die Semantik des zusätzlichen Objektes instantiiert:

$$(514) \quad \text{CONT} \quad \left[\begin{array}{l} \text{be-affected-by} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{stem}) \end{array} \right]$$

Damit ergibt sich für das indirekte Passiv der folgende Lexikoneintrag:

Morphologie:

$$(515) \quad a. \quad \text{suffix}^{\diamond} \\
b. \quad \text{SUBCAT} \quad \left\{ \text{PP}[\text{subj}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\} \\
\rightarrow \left\{ \text{PP}[\text{subj}, \text{ga}]^{\text{opt}}, \text{PP}[\text{obj}, \text{ni}]:\boxed{\text{I}}^{\boxed{2}}, \dots \right\} \\
c. \quad \text{CONT} \quad \left[\begin{array}{l} \text{be-affected-by} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{stem}) \end{array} \right] \\
d. \quad i. \quad \text{MAJ} \quad v_{v, i_{kuru}}^a \rightarrow v_v^{\text{comb}}; \quad \text{PHON} \quad \langle \text{rare} \rangle \\
ii. \quad \text{MAJ} \quad v_c^a \rightarrow v_v^{\text{comb}}; \quad \text{PHON} \quad \langle \text{re} \rangle \\
iii. \quad \text{MAJ} \quad v_{i_{suru}}^a \rightarrow v_v^{\text{comb}}; \quad \text{PHON} \quad \langle \text{are} \rangle$$

5.2.4.6 Kausativ-Passiv (-CAUS-PASS)

Kausativ und Passiv werden — in dieser Reihenfolge — oft zusammen verwendet. Zur Ableitung und Verwendung dieser Kombination vergleiche das folgende Beispiel:

Beispiele:

- (516) a. 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
Tarō-NOM Zirō-ACC rufen-PAST
Tarō rief Zirō.

- b. 美智子が太郎に二郎を呼ばせた。³¹³

mitiko-ga tarou-ni zirou-wo yo.b.a-se.t-a
Mitiko-NOM Tarō-DAT Zirō-ACC rufen-CAUS-PAST
Mitiko ließ Tarō Zirō rufen.

- c. 太郎が美智子に二郎を呼ばせられた。³¹⁴

tarou-ga mitiko-ni zirou-wo yo.b.a-se-rare.t-a
Tarō-NOM Mitiko-DAT Zirō-ACC rufen-CAUS-PASS-PAST
Tarō wurde von Mitiko veranlaßt, Zirō zu rufen.

Die Modifikationen des Subkategorisierungsrahmens durch Kausativ und Passiv heben sich teilweise auf: Das unterliegende Subjekt des Verbes erscheint auch als Subjekt an der syntaktischen Oberfläche; zusätzlich zu den Objekten des Verbes erscheint ein Dativobjekt, das durch den Kausativ eingeführt wird und den Veranlassenden der Handlung beschreibt.

5.2.4.7 Potential (-POT)

Durch das *Potential*^{B15} wird die *Möglichkeit* oder *Fähigkeit* bezeichnet, etwas zu tun:

Beispiele:

- (517) a. i. 私は英語を話します。³¹⁶

watasi-ha eigo-wo hana.s.i-mas.-u
ich-TOP Englisch-ACC sprechen-HON-FLEX
Ich spreche Englisch.

³¹³ LIJC, S. 34, (78)a.

³¹⁴ LIJC, S. 34, (78)b.

³¹⁵ JM, S. 93, 20-3.1.1.2. -*Re.ru* 'Potential'.

³¹⁶ DBJG, S. 371, notes (1)a.

- ii. 私は英語が/を話せます。 ³¹⁷
watasi-ha eigo-ga/-wo hana.s-e-mas.-u
 ich-TOP Englisch-NOM/-ACC sprechen-POT-HON-FLEX
 Ich kann Englisch sprechen.
- b. i. 彼が漢字を書くこと ³¹⁸
kare-ga kañzi-wo ka.k-u koto
 er-NOM chinesische Zeichen-ACC schreiben-FLEX Sachverhalt
 [...] daß er chinesische Zeichen schreibt.
- ii. 彼が/に漢字が書けること ³¹⁹
kare-ga/-ni kañzi-ga ka.k-e-ru
 er-NOM/-DAT chinesische Zeichen-NOM schreiben-POT-FLEX
koto
 Sachverhalt
 [...] daß er chinesische Zeichen schreiben kann.
- c. 私は日本語が読める / 読めます。 ³²⁰
watasi-ha nihongo-ga yo.m-e-ru / yo.m-e-mas.-u
 ich-TOP Japanisch-NOM lesen-POT-FLEX / lesen-POT-HON-FLEX
 Ich kann Japanisch lesen.
- d. この水は飲めない / 飲めません。 ³²¹
kono mizu-ha no.m-e-na-i / no.m-e-mas.e-ñ
 dieses Wasser-TOP trinken-POT-NEG-FLEX / trinken-POT-HON-NEG
 Dieses Wasser ist nicht trinkbar.

Der Potential kann fakultativ mit einer Änderung der Kasusreaktion einhergehen (517*a*). Die Variante ohne Änderung der Valenz wird durch den folgenden Lexikoneintrag beschrieben:³²²

Morphologie (ohne Änderung der Kasusreaktion):

³¹⁷ DBJG, S. 371, notes (1)b.

³¹⁸ JM, S. 94, 20-3.1.1.2. b) 1.

³¹⁹ JM, S. 94, 20-3.1.1.2. b) 1.

³²⁰ DBJG, S. 370, KS(A).

³²¹ DBJG, S. 370, KS(B).

³²² Wie der folgende Eintrag zeigt, bilden einige der vokalischen Verben mit dem Suffix (518*d-iii*) *re* eine zweite Potentialform.

(518) a. *suffix*[◇]

b. CONT $\left[\begin{array}{l} \text{can} \\ \textcircled{1} \text{ cont}(\text{subj}) \\ \textcircled{2} \text{ cont}(\text{stem}) \end{array} \right]$

c. ADJUNCT +

d. i. MAJ $v_c^{cons} \longrightarrow v_v^{comb};$ PHON $\langle e \rangle$
 ii. MAJ $v_{v,i_{kuru}}^a \longrightarrow v_v^{comb};$ PHON $\langle rare \rangle$
 iii. MAJ $v_v^{cons} \longrightarrow v_v^{comb};$ PHON $\langle re \rangle$

Adverben, die mit einem Verb im Potential stehen, können sich sowohl auf die lexikalische Bedeutung des Verbes beziehen als auch auf den Aspekt des “Könnens”, der durch das Potentialsuffix eingeführt wird:

(519) a. *Peter kann gut Briefe schreiben:*

$$\left[\begin{array}{l} \text{gut} \\ \left[\begin{array}{l} \text{kann} \\ \textcircled{1} \text{ Ⅰ Peter} \\ \left[\begin{array}{l} \text{schreiben} \\ \textcircled{1} \text{ Ⅰ} \\ \textcircled{2} \text{ Briefe} \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

b. *Peter kann Briefe gut schreiben:*

$$\left[\begin{array}{l} \text{kann} \\ \textcircled{1} \text{ Ⅰ Peter} \\ \left[\begin{array}{l} \text{gut} \\ \left[\begin{array}{l} \text{schreiben} \\ \textcircled{1} \text{ Ⅰ} \\ \textcircled{2} \text{ Briefe} \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

Die Potentialsuffixe werden deshalb mit [ADJUNCT +] gekennzeichnet.³²³

Bei einer Änderung der Kasusreaktion wird das unterliegende Subjekt an der Oberfläche mit der Postposition *ha* (*ha*) oder *ni* (*ni*) realisiert. Dem Topicpartikel *ha* (*ha*) liegt in der SUBCAT-Menge ein *ga* (*ga*) zugrunde:

Morphologie:

323 Vgl. den Abschnitt 4.3 “Lexikalisch-morphologische Regeln”, S. 179.

$$(520) \text{ SUBCAT } \left\{ \text{PP}[\text{subj}, \text{ga}]:\text{I}^{\text{[2]}}, \text{PP}[\text{obj}, \text{wo}]:\text{III}^{\text{[4]}}, \dots \right\} \\ \longrightarrow \left\{ \text{PP}[\text{subj}, \text{ga} \vee \text{ni}]:\text{I}^{\text{[2]}}, \text{PP}[\text{obj}, \text{ga}]:\text{III}^{\text{[4]}}, \dots \right\}$$

来る (*kuru*; kommen) bildet neben der regelmäßigen Potentialform 来られる (*korareru*) einen zweiten, unregelmäßigen Potential; der Potential von する (*suru*; tun, machen) ist ebenfalls unregelmäßig³²⁴:

(521)	<i>Verb</i>	<i>Kausativ</i>
	来る	→ 来られる, 来れる
	<i>kuru</i>	<i>korareru koreru</i>
	する	→ 出来る
	<i>suru</i>	<i>dekiru</i>

5.2.4.8 Hilfssuffix ‘-Karu’ (- ϕ)

Ein großer Teil der Suffixe, wie das Volitional (vgl. (522a)), das Perfekt-Flexiv (vgl. (522b) und (522c)), das Perfekt-Konditional (vgl. (522d)) etc. können nicht an adjektivische, sondern nur an verbale Formen angeschlossen werden. In diesem Fall dient das semantisch neutrale Suffixverb *karu*³²⁵ als Mittler zwischen adjektivischem Stamm und Suffix, indem es den Stamm in eine verbale Form überführt:

Beispiele:

(522) a. あの人が日本の着物を着たら、きっと美しかろう。³²⁶

ano hito-ga nihoñ-no kimono-wo ki.t-ara kitto
 jener Mensch-NOM Japan-GEN Kimono-ACC tragen-PCOND bestimmt
utukusi-ka.r-ou
 schön- ϕ -VOLI

Sie wird bestimmt schön aussehen, wenn sie einen japanischen Kimono trägt.

b. 次の日の試験は難しかったなあ。³²⁷

tugi-no hi-no sikeñ-ha muzukasi-ka.tt-a naa
 nächster-GEN Tag-GEN Prüfung-TOP schwer- ϕ -PAST INTERJEKTION

Die Prüfung am nächsten Tag war aber schwer!

324 Vergleiche auch die betreffenden Lexikoneinträge: für 来る (*kuru*): (414), S. 216; und für する (*suru*): (419), S. 218.

325 JM, S. 209, 30-3.1.1. -*kar.u*.

326 JM, S. 209, 30-3.1.1. c) 2.

327 vgl. JM, S. 210, 30-3.1.1. c) 6.

- c. よかった。なくなったゆびわがこんな所にあったわ。³²⁸

yo-ka.tt-a. na-ku na.tt-a yubiwa-ga koñna
 gut- ϕ -PAST nicht sein-ADV werden-PAST Fingerring-NOM solch ein
tokoro-ni a.tt-a
 Platz-LOC sein-PAST

Welch ein Glück! Der verlorengegangene Ring hat an einem solchen Platz gesteckt.

- d. およろしかったら、ちょっと私のうちへいらっしゃいませんか。³²⁹

oyorosi-ka.tt-ara tyotto watakusi-no uti-he irassya.i-mas.e-ñ
 HON-gut- ϕ -PCOND ein wenig ich-GEN Haus-DIR kommen-HON-NEG
ka
 QUESTION

Könnten Sie nicht, wenn es Ihnen genehm sein sollte, für einen Moment zu uns nach Hause kommen?

Das Suffixverb *-Karu* modifiziert demzufolge nur PHON- und MAJ-Wert des Stammes:

Morphologie:

(523) a. *suffix*[◇]

b. MAJ *adj*^v $\longrightarrow$ *v*_{c_r}^{cv}; PHON $\langle ka \rangle$

328 vgl. JM, S. 210, 30-3.1.1. c) 6.

329 JM, S. 210, 30-3.1.1. c) 7.

6 Entwicklung und Implementierung des vorgestellten Modells

6.1 TFS als experimentelles System

Constraintbeschreibungen — d.h. Typenhierarchien und Typenconstraints — lassen sich als Programm auffassen: Ähnlich wie im Paradigma der *logischen Programmierung* logische Beschreibungen eines Gegenstandsbereiches durch Hornformeln als Programm interpretiert werden, lassen sich Typenhierarchien und Typenconstraints direkt zur Berechnung der Strukturen verwenden, die durch sie beschrieben werden. Der “Mechanismus”, der diesen Berechnungen zugrunde liegt, beruht auf einem mathematischen Kalkül und wurde im Abschnitt 1.5 “*Resolution*” beschrieben.

Um Constraintbeschreibung und Kalkül auf einem Computer ausführen zu können, müssen diese in eine Form gebracht werden, die durch einen Rechner manipuliert werden kann. Eine Möglichkeit ist es, den Kalkül direkt in ein Programm umzusetzen, das die Constraintbeschreibung interpretiert. Ein Beispiel für eine solche informatische Umsetzung ist TFS (*The Typed Feature Structure Representation Formalism*)³³⁰, das System, das für die Umsetzung der vorgestellten Grammatik verwendet wurde.

Der vorgestellte Kalkül läßt sich also als Spezifikation der “Programmiersprache TFS”, die in den letzten Kapiteln angegebene Grammatik als “TFS-Programm” verstehen. Die Implementierung war demnach schon Thema der vorangehenden Betrachtungen, so daß sich die Frage nach der Motivation des vorliegenden Kapitels stellt. Die Antwort ergibt sich aus der Komplexität der vorgestellten Grammatik: TFS folgt bei seinen Berechnungen einer sehr allgemeinen Strategie und kann Abhängigkeiten zwischen den Merkmalen, die sich aus den Constraints und der Art der Anfragen ergeben, nicht berücksichtigen. Daraus resultiert in den meisten Fällen ein sehr zeit- und speicheraufwendiges Backtracking. Obwohl TFS bei unbegrenzter Zeit und unbegrenztem Speicher die meisten Constraintbeschreibungen berechnen kann,³³¹ lassen sich in der Praxis daher nur kleine, experimentelle Grammatiken wirklich ausführen.

TFS ist damit als experimentelles System zu verstehen: Es macht die Rolle von Constraintsystemen als Formalismus zur eleganten, deklarativen Spezifikation von Grammatiken deutlich, indem es den Zusammenhang zwischen Spezifikation und formaler Semantik aufzeigt. Beschreibungen kleiner

³³⁰ vgl. Emele und Zajac 1990a, Emele und Zajac 1990b, Emele 1991, Emele 1993, Emele 1994 und Zajac 1992.

³³¹ TFS benutzt keine faire Strategie bei der Traversierung der Ableitungsräume, sondern Tiefensuche. Bestimmte Formen der Rekursion können daher zu Endlosableitungen führen.

Sprachausschnitte können direkt überprüft werden, ohne durch zusätzliches strategisches Wissen verfälscht zu werden. Für die praktische Anwendung von Grammatiken bei der Analyse und Generierung von Sprache hingegen empfiehlt sich die Verwendung spezieller, für diese Zwecke optimierter Systeme wie ALE³³² oder PAGE³³³.

6.2 Strategien für die Implementierung komplexer Grammatiken

Die Grammatik des Japanischen, die in dieser Arbeit entwickelt wurde, ist zu komplex, um direkt durch TFS interpretiert werden zu können. Es wurden deshalb verschiedene Strategien angewandt, um den Aufwand der Berechnungen zu verringern. Alle Strategien beziehen sich auf die Implementierung des Ableitungskalküls in TFS. Diese beruht auf Backtracking: Der Zustandsraum, der durch eine Anfrage und den Kalkül aufgespannt wird, wird durch Tiefensuche traversiert. Enthält die Anfrage mehrere Teilprobleme, kommt es daher zu einer Multiplikation ihres Zeitaufwands. Ein Satz mit einem Adjektiv, einem Adverb und einem Verb beispielsweise enthält drei morphologisch komplexe Formen. Auch wenn jede dieser Formen einfach ist, ist der Aufwand bei der Analyse des Gesamtsatzes erheblich.

Um den Rechenaufwand gering zu halten, bieten sich zwei Ansätze an:

1. Vereinfachung der Anfrage,
2. Vereinfachung der Constraintbeschreibung.

Bei der Entwicklung von Grammatiken empfiehlt sich die Zerlegung in Subsysteme, die getrennt oder inkrementell entwickelt werden.

6.2.1 Vereinfachung der Anfrage

Für die Berechnung einfacher Probleme ist die Rechenkapazität von TFS ausreichend. Einzelne Aspekte unserer Grammatik lassen sich damit durch TFS alleine berechnen. Es können verschiedene Strategien angewandt werden, um Anfragen einfach zu halten:

³³² vgl. Carpenter und Penn 1999

³³³ Die Module, aus denen sich das PAGE-System zusammensetzt, werden in Uszkoreit et al. 1994, Kiefer und Fettig 1995, Kiefer und Scherf 1994, Krieger und Schäfer 1994a, Krieger und Schäfer 1994b, Krieger und Schäfer 1994c und Backofen und Weyers 1994 dokumentiert.

1. *Beschränkung auf ein Problem pro Anfrage*

Morphologische Formen, syntaktische Schemata etc. lassen sich anhand von Anfragen untersuchen, die neben dem Aspekt von Interesse keine weiteren aufwendigen Probleme enthalten.

2. *Zerlegung komplexer Probleme in Teilprobleme*

Diese Strategie ist ein Spezialfall der zuletzt genannten: Komplexe Anfragen können in mehrere Anfragen zerlegt werden. Beispielsweise lassen sich Satzgefüge in mehreren Schritten analysieren oder die Berechnung morphologischer Strukturen von der Berechnung syntaktisch-semanticischer trennen. In beiden Fällen ist zur Integration der Ergebnisse ein weiterer Schritt notwendig, der als “*Vorberechnen von Teilstrukturen*” an späterer Stelle erklärt wird.

3. *Hinzunahme zusätzlicher Informationen*

Die Anfragen an TFS können durch zusätzliche Informationen angereichert werden. Da vorgegebene Informationen nicht mehr durch Backtracking berechnet werden müssen, läßt sich mit diesem Mittel der Suchraum verkleinern und damit der Rechenaufwand unter Umständen drastisch verringern. Sollen beispielsweise komplexe Sätze untersucht werden, können Teile ihrer morphologischen oder syntaktischen Struktur vorgegeben werden.

6.2.2 Vereinfachung der Constraintbeschreibung

Der Aufwand von Berechnungen durch TFS resultiert nicht nur aus der Struktur der Anfrage, sondern auch aus der Anzahl alternativer Typen und Typenconstraints, die für eine Berechnung zu berücksichtigen sind, d.h. der Komplexität der Grammatik. Zur Vereinfachung der Grammatik relativ zu einer Anfrage lassen sich verschiedene Ansätze verfolgen:

1. *Vorberechnung endlicher Probleme*

- (a) *Auflösung von Rekursionen durch Betrachtung endlicher Fälle*

Rekursionen können als mächtiges und elegantes Abstraktionsmittel bei der Theoriebildung bezeichnet werden. In bestimmten Fällen kann jedoch eine Abbildung rekursiver auf endliche Strukturen die Berechnung vereinfachen. Subkategorisierungsmengen beispielsweise haben selten mehr als fünf Elemente und können daher auch als Merkmalstrukturen mit den Merkmalen ELEMENT1, ELEMENT2, ..., ELEMENT5 dargestellt werden. Rekursive Constraints, die sich auf derart reduzierte Strukturen beziehen, lassen

sich vorberechnen, so daß sich die Berechnungen zur Laufzeit entsprechend vereinfachen.³³⁴

(b) *Vorbereitung von Teilstrukturen*

Die Heuristik “*Speichern statt Berechnen*” aus der Algorithmik kann auch auf linguistische Probleme angewendet werden: Teile von Constraintbeschreibungen lassen sich in manchen Fällen durch vorberechnete Constraints ersetzen. Bei komplexen Anfragen beispielsweise kann die Berechnung in mehreren Phasen geschehen. So kann die Grammatik in einem ersten Schritt dazu verwendet werden, ein Vollformlexikon der in der Anfrage vorkommenden Wörter zu generieren, das anschließend von der Syntax zur Analyse der Satzstruktur verwendet wird. Ebenso lassen sich natürlich einzelne Phrasen oder Teilsätze getrennt analysieren oder generieren.

2. *Verwendung von Ausschnitten der Grammatik*

Das Verfahren des “*dynamischen Ladens*” aus der Informatik läßt sich ebenfalls auf TFS anwenden: Werden nur die lexikalischen Zeichen, Schemata und Prinzipien geladen, die für eine Anfrage notwendig sind, kann u.U. eine erhebliche Vereinfachung der Berechnung erzielt werden.

6.2.3 Entwicklung von Subsystemen

Bei der Entwicklung einer Constraintbeschreibung lassen sich weitere Strategien anwenden. In den meisten Fällen können Teilaspekte der Beschreibung getrennt oder aber inkrementell entwickelt werden: Syntax und Morphologie beispielsweise lassen sich weitgehend unabhängig voneinander modellieren und untersuchen; die Semantik kann inkrementell auf der Syntax aufgebaut werden. Bei der Integration dieser Subsysteme bzw. der inkrementellen Weiterentwicklung können Strategien wie die Vorbereitung von Strukturen, die Hinzunahme von Informationen etc. verfolgt werden.

Um die Anwendung dieser Strategien bei der Entwicklung der Grammatik und ihrer Verifikation anhand eines Korpus von Beispielsätzen automatisieren zu können, wurde TFS in ein “Meta-System” integriert. In unserem Fall wurde dazu die Programmiersprache *Expect*³³⁵ verwendet: Komplexe Probleme können damit durch *Expect*-Programme in Unterprobleme zerlegt

³³⁴ Diese Strategie wird auch in der Verbmobilgrammatik von Melanie Siegel angewandt.

³³⁵ TFS wird als Lisp-Image ausgeliefert; die Lisp-Quellen sind nicht zugänglich. Da TFS keine Programmierschnittstelle hat, wurde in *Expect* ein Wrapper geschrieben, der mit TFS über die Benutzershell kommuniziert. Auf dieser Weise wurde es möglich, TFS wie eine normale Funktion von *Expect*-Programmen aus aufzurufen.

werden; die Teilprobleme können entweder an TFS weitergegeben, oder direkt von Expect berechnet werden; die Integration der Teilergebnisse wird — gesteuert wiederum von *Expect*-Programmen — von TFS vorgenommen.

6.3 Entwicklung der japanischen Grammatik

Die Grammatik des Japanischen, die in den vorangehenden Kapiteln vorgestellt wurde, kann nur für einfache Anfragen direkt durch TFS interpretiert werden. Schon ein Großteil der Beispielsätze und morphologischen Formen aus Anhang A sind für die direkte Berechnung durch TFS zu aufwendig. Bei der Entwicklung der japanischen Grammatik wurden deshalb schon zu einem frühen Zeitpunkt die im letzten Abschnitt vorgestellten Strategien angewandt. Der Darstellung der verschiedenen Stadien dieser Entwicklung im Zusammenhang mit den angewandten Strategien dient die folgende Auflistung:

6.3.1 Verbmorphologie

Verbformen und Morphe

Ausgehend von den Verbformen des *“The Complete Japanese Verb Guide”* (CJVG, Ishii et al. 1989)³³⁶, der traditionellen Grammatik (vgl. Lewin 1990) und der *“Japanische[n] Morphosyntax”* (Rickmeyer 1995) wurde ein Lexikon von Morphen erstellt. Die Beschränkungen bezüglich der Kombination und Realisierung der Morphe wurde durch morphologische Subkategorisierung ausgedrückt; die möglichen Kombinationen der Morphe bei der Bildung komplexer Formen entsprechend dieser Beschränkungen wurden durch ein morphologisches Schema modelliert.

Zur Berechnung der so entstandenen Morphologie war die Kapazität von TFS ausreichend; Strategien mußten zu diesem Zeitpunkt noch nicht angewandt werden.

³³⁶ Der *“The Complete Japanese Verb Guide”* (CJVG, Ishii et al. 1989) erfüllt alle Voraussetzungen, um als Korpus für eine relativ vollständige Beschreibung der japanischen Verbformen zu dienen: Für die häufigsten Verben des modernen Japanisch listet er alle wichtigen Grundformen auf. (Die Tabellen aus Anhang B wurden den Tabellen des CJVG nachempfunden und vermitteln damit eine Vorstellung von seinen Verbeinträgen.) Da die Auswahl der Verben und Verbformen repräsentativ für das moderne Japanisch ist, lassen sich nahezu alle Verbformen aus seinen Tabellen ableiten: Wird ein Verb nicht aufgelistet, läßt es sich anhand eines nach gleichen Regeln flektierenden erklären; wird eine Verbform nicht in der Tabelle dargestellt, läßt sie sich aus der Kombination ihrer Grundformen ableiten.

Morphe und morphologische Formprimitiven

Die Morphe wurden in morphologische Formprimitiven (Kern und Clitics) zerlegt und ihre Struktur systematisch beschrieben. Um das morphologische Schema auch zur Ableitung der Morphe verwenden zu können, wurde die Hierarchie der Stämme erweitert.

Bei der Ableitung der morphologischen Beschreibung diente eine Menge repräsentativer Verben und Verbformen — beide wurden aus dem *“The Complete Japanese Verb Guide”* entnommen — zur Verifikation der entwickelten Beschreibung: Alle relevanten Verbformen wurden mit dem entwickelten Modell generiert und mit den Formen aus dem CJVG verglichen. Anhand der ermittelten Fehler wurde anschließend die Beschreibung korrigiert. Dieser Vorgang wurde so lange wiederholt, bis alle Formen korrekt generiert wurden. Auch wenn es sich bei den so überprüften Formen nur um einfache Grundformen handelt, folgt aus ihrer korrekten Ableitung die Richtigkeit des gesamten morphologischen Modells: Komplexe Verbformen leiten sich aus der rekursiven Kombination der Grundformen ab. Durch “strukturelle Induktion” kann damit aus der Korrektheit der Grundformen auf die Korrektheit der komplexen Formen geschlossen werden.

Bei der Ableitung des morphologischen Modells wurde *Expect* nur zur Automatisierung der Generierung und zur Überprüfung der Verbformen verwendet: Die zugrunde gelegten Verbformen sind relativ einfach strukturiert — sie bestehen aus maximal drei Morphemem — und lassen sich damit noch durch TFS berechnen.

Die Analyse und Generierung komplexerer Formen ist jedoch zu zeit- und speicheraufwendig für TFS. Für ihre Berechnung wurde daher erstmals eine der im letzten Abschnitt beschriebenen Strategien angewandt: die Formen wurden inkrementell berechnet. Ein Expect-Programm benutzte die morphologische Beschreibung, um Formen aus zwei morphologischen Primitiven abzuleiten. Die abgeleiteten Merkmalstrukturen wurden in ein temporäres Lexikon geschrieben und zur Generierung von Formen der Länge drei verwendet usw. Dieser Vorgang wurde so lange wiederholt, bis die Verbformen vollständig abgeleitet waren. Während die morphologische Beschreibung von TFS zur Generierung, Analyse und Vervollständigung von Zeichen verwendet werden kann, ist die Expect-basierte Ableitung richtungsabhängig: sie läßt sich nur zur Analyse von Formen verwenden.

6.3.2 Syntax

Die syntaktischen und semantischen Modifikationen, die mit der agglutinierenden Verbmorphologie einhergehen, lassen sich am besten anhand eines Korpus untersuchen. Aus diesem Grund wurden zu jeder morphologischen

Form aus verschiedenen Grammatiken und Wörterbüchern³³⁷ repräsentative Beispielsätze gesucht.³³⁸ Um die Beispielsätze modellieren und untersuchen zu können, wurde ein syntaktisch-semantisches Modell benötigt. Die Entwicklung des syntaktischen Teils dieses Modells war Ziel der zweiten Arbeitsphase.

Bei der Entwicklung der Syntax wurde zuerst von stark vereinfachten Beispielsätzen und einem ebenso grundlegenden Vollformlexikon ausgegangen. Die so entstandene Grammatik wurde anhand komplizierterer Beispielsätze nach und nach erweitert, bis alle syntaktischen Phänomene, die in den Beispielsätzen vorkommen, behandelt werden konnten.

Für die Berechnung einfacher Sätze war die Kapazität von TFS ausreichend; bei komplexen Sätzen hingegen mußte die syntaktische Struktur zumindest teilweise vorgegeben werden.

6.3.3 Semantik

In der nächsten Entwicklungsphase wurde die Syntax um eine einfache kompositionale Semantik erweitert: Die Lexikoneinträge wurden durch semantische Beschreibungen ergänzt und die Schemata und Prinzipien so erweitert, daß sie diese kombinieren konnten. Bei komplexen Sätzen mußte wie bei der Ableitung der Syntax die syntaktische Struktur teilweise vorgegeben werden.

6.3.4 Integration von Morphologie, Syntax und Semantik

Erweiterung der morphologischen Zeichen um syntaktische und semantische Merkmale

Um das syntaktisch-semantische Verhalten der Verbformen zu modellieren, wurden die morphologischen Lexikoneinträge durch entsprechende Merkmale angereichert: Die Verbstämme wurden mit Merkmalen versehen, die ihr Defaultverhalten modellieren; die Suffixe mit Spezifikationen der Modifikationen dieser Merkmale, die durch sie bewirkt werden. Das morphologische Schema und die morphologischen Prinzipien wurden entsprechend modifiziert. Das Ziel dieser Entwicklungsphase war es, aus den Einträgen der morphologischen Primitiven das Vollformlexikon ableiten zu können, das für die Syntax und Semantik verwendet wurde.

Wie bei der Entwicklung der Morphologie wurden die Formen unter Zuhilfenahme von *Expect* inkrementell berechnet.

³³⁷ Eine Liste der verwendeten Grammatiken und Wörterbüchern findet sich im Anhang A auf Seite 287.

³³⁸ vgl. Anhang A.

Integration des morphologischen Modells

Morphologie, Syntax und Semantik konnten nun in einer gemeinsamen Beschreibung integriert werden. Mit der gewonnenen Grammatik und TFS ließen sich einfache Sätze mit einfachen morphologischen Formen direkt generieren und analysieren; komplexe Sätze und Sätze mit komplexen morphologischen Formen wurden wiederum durch ein hybrides System aus *Expect* und TFS analysiert. Es wurde folgendermaßen vorgegangen:

1. Die atomaren syntaktischen Zeichen wurden aus dem Satz extrahiert und mittels der morphologischen Komponente inkrementell berechnet. Die Ergebnisse wurden in ein "Vollformlexikon" für den entsprechenden Satz gespeichert.
2. Unter Zuhilfenahme des generierten Lexikons wurde der Satz syntaktisch und semantisch analysiert. Bei komplexen Sätzen mußte die syntaktische Struktur vorgegeben werden.

Damit das Expect-Programm die syntaktischen und morphologischen Strukturen erkennen konnte, mußten diese in den Anfragen vorgegeben werden.

Motivation dieses Vorgehens war die Entwicklung und Verifikation des entwickelten sprachlichen Modells, nicht die Anwendung der Grammatik zur Analyse oder Generierung. Auch wenn das hybride System aus TFS und *Expect* damit besser als experimentelle "Entwicklungsumgebung" für Grammatiken beschrieben wird, soll im nächsten Abschnitt seine Eignung als "computerlinguistisches System" bewertet werden.

6.4 Kritik des benutzten Verfahrens

Um das dargestellte Vorgehen zu bewerten, lassen sich zwei grundlegende Dimensionen heranziehen:

1. *Deklarative Sicht:*
Die deklarative Eleganz der entstandenen linguistischen Beschreibung im Zusammenhang mit dem durch sie beschriebenen Sprachausschnitt.
2. *Anwendungsperspektive:*
Die Eignung des entwickelten Gesamtsystems zur Anwendung für linguistische Berechnungen, d.h. das konkrete Laufzeitverhalten bei der Analyse (Parsen) bzw. Synthese (Generierung) von sprachlichen Formen unter Berücksichtigung von Zeit- und Speicheraufwand.

Schon das Thema der Arbeit betont den zuerstgenannten, deklarativen Aspekt: Die japanische Verbmorphologie soll im Zusammenhang mit der Semantik — und damit auch der Syntax — computerlinguistisch *beschrieben* werden. Eine Betonung der Deklarativität als Ziel einer ersten Arbeitsphase ist jedoch auch bei anwendungsorientiertem Vorgehen vorteilhaft: Durch das Erstellen einer deklarativen Beschreibung und die anschließende Überarbeitung nach prozeduralen Gesichtspunkten ergeben sich im Allgemeinen effizientere, robustere und leichter erweiterbare Systeme als bei der direkten Ausrichtung auf Anwendbarkeit.

Im Folgenden sollen kurz die Einschränkungen hinsichtlich der direkten Anwendbarkeit dargestellt werden, die sich aus der Verwendung eines deklarativen Systems wie TFS ergeben. Anschließend jedoch soll dargelegt werden, warum dieses Vorgehen, verglichen mit der Verwendung anwendungsorientierter Systeme, trotzdem Vorteile hat.

TFS und die japanische Grammatik

Typenhierarchien und Constraints erlauben eine elegante Spezifikation von linguistischen Sachverhalten. Diese Eleganz resultiert aus der deklarativen Form der Darstellung und ihrer mathematisch definierten Semantik. Die „algorithmisierte Form“ dieser mathematischen Semantik — der in Kapitel 1 beschriebene Kalkül — wird von TFS unter Anwendung der Tiefensuche direkt zur Ableitung von Lösungen angewandt. TFS erfüllt damit — im Gegensatz zu anwendungsorientierten Systemen — alle Anforderungen, die aus theoretischer Sicht an deklarative Systeme gestellt werden.

Aus eben diesen Qualitäten ergeben sich jedoch Nachteile bei der Anwendung von TFS zur Analyse oder Generierung von Sprache: Berechnungen sind aufgrund der allgemeinen Ableitungsstrategie extrem zeit- und speicheraufwendig, so daß TFS nur für relativ kleine Constraintbeschreibungen benutzt werden kann.

In dem hier beschriebenen Ansatz wurde dieses Problem mittels eines hybriden Systems aus TFS und *Expect* umgangen. Die daraus resultierende Aufteilung der computerlinguistischen Berechnungen auf TFS und *Expect* und die Vereinfachung der Berechnungen durch die im letzten Abschnitt vorgestellten Strategien führten jedoch zu erheblichen Einschränkungen des Gesamtsystems:

1. Bei komplexen Sätzen oder morphologischen Formen kann das erstellte System nur zur Analyse benutzt werden.
2. Bei komplexen Eingaben muß die morphologisch-syntaktische Struktur vorgegeben werden. Das Erkennen der syntaktischen Struktur — sonst

eine der wichtigsten Aufgaben der syntaktischen Analyse — ist zu komplex. Die Aufgabe von TFS beschränkt sich also darauf, zu überprüfen, ob die vorgegebene Struktur mit der Grammatik verträglich ist.

3. Durch die Zerlegung in Teilsysteme, das Vorgeben von Informationen etc. kommen wesentliche Aspekte der Grammatik nicht zur Anwendung. Das Verhalten der grammatischen Spezifikation kann ohne diese äußeren Einschränkungen für komplexe Sätze nicht experimentell überprüft werden.

Für die konkrete Anwendung ist das System aus TFS, *Expect* und der japanischen Grammatik damit nicht geeignet: Weder zur Generierung noch zur Analyse läßt es sich direkt verwenden. Anders verhält es sich mit der entwickelten Grammatik: Als deklarative Beschreibung ist sie von TFS unabhängig und kann auf jedem anderen computerlinguistischen System ohne allzu großen Aufwand implementiert werden.

Anwendungsbezogene Systeme

Um trotz der syntaktischen und morphologischen Kombinatorik zu einem akzeptablen Laufzeitverhalten zu gelangen, werden anwendungsorientierte computerlinguistische Systeme, wie z.B. PAGE oder ALE, für die Generierung oder Analyse von Sätzen optimiert. Diese Optimierung besteht in der festen Integration bestimmter Berechnungsstrategien in das betreffende System. Diese Strategien jedoch führen genau wie die im letzten Abschnitt vorgestellten zu einer Abweichung des Systems von der deklarativen Semantik der linguistischen Beschreibungssprache. Im Gegensatz zu der Vorgehensweise, der bei der Entwicklung der dargestellten Grammatik gefolgt wurde, kann diese Abweichung nicht vom Autor der Grammatik kontrolliert werden. Selbst bei der Betrachtung eines kleinen Ausschnittes eines sprachlichen Modells erfolgt die Ableitung nicht entsprechend dessen deklarativer Semantik, die damit nicht experimentell überprüft werden kann.

Damit wird der Vorteil des in dieser Arbeit verfolgten Ansatzes deutlich: Durch die Möglichkeit, über die Verwendung von Strategien bewußt zu entscheiden, kann die Semantik von Grammatiken besser überprüft werden, als bei der Verwendung eines Systems mit fest integrierten Strategien. Bei bewußtem Einsatz der Strategien ist es daher leichter, zu einer korrekten Spezifikation der linguistischen Sachverhalte zu kommen.

Zusammenfassung

Die Implementierung unseres Systems dient also nicht der Anwendung, sondern der Ableitung der vorgestellten Grammatik und sollte deshalb auch

nicht an anwendungsbezogenen Kriterien gemessen werden. Die Anwendung der Grammatik — und damit verbunden die Optimierung hinsichtlich der Generierung oder Analyse von Sprache — läßt sich besser im Rahmen einer gesonderten Betrachtung diskutieren und bildet damit ein interessantes Thema für die Fortführung dieser Arbeit.

7 Ergebnisse und Ausblick

In den vorangehenden Kapiteln wurde ein repräsentatives Fragment der japanischen Verbmorphologie unter Berücksichtigung von Syntax und Semantik formal beschrieben. Um einen Überblick über die Ergebnisse des vorgestellten Ansatzes zu erhalten und auf interessante mögliche Weiterführungen dieser Arbeit hinzuweisen, sollen die wichtigsten Grundlagen, Charakteristika und auch Beschränkungen des vorgelegten Modells nun noch einmal zusammenfassend dargestellt werden.

Ergebnisse

Der linguistische Rahmen des entwickelten Modells wurde durch die Orientierung an der traditionellen japanischen Verbmorphologie und an modernen linguistischen Theorien wie der HPSG und der JPSG festgelegt. Hieraus ergibt sich auch der formale Rahmen, da HPSG und JPSG auf einer speziellen Klasse mathematischer Formalismen basieren, den sogenannten *getypten Merkmalslogiken*. Diese sind auch ein geeignetes Mittel zur Formalisierung der traditionellen japanischen Verbmorphologie.

Merkmalslogiken begünstigen — auch wenn sie Turing-mächtig sind — eine bestimmte Art der linguistischen Modellbildung: Sie sind Resultat eines langen Annäherungsprozesses linguistischer und mathematischer Ansätze und stellen einen gelungenen Kompromiß zwischen den Anforderungen und Möglichkeiten beider Seiten dar.

Die vorgestellte Verbmorphologie ist damit Resultat der Abwägung folgender drei Anforderungen:

1. die inhärenten Prinzipien der japanischen Grammatik und Morphologie in möglichst eleganter und ökonomischer Weise herauszuarbeiten;
2. die linguistischen Ansätze, wie sie sich in der Tradition von HPSG und JPSG bewährt haben, auf die Beschreibung der japanischen Grammatik und Verbmorphologie fortzusetzen;
3. die Beschreibung so zu gestalten, daß die formalen Mittel von Merkmalslogiken und die Prinzipien der japanischen Sprache, soweit dieses möglich ist, zur Deckung gebracht werden.

Aus diesen Forderungen ergeben sich die folgenden linguistischen und formalen Grundprinzipien unseres Modells:

Rekursive Komposition als einzige Grundoperation

Alle morphologischen und syntaktischen Formen werden unter Voraussetzung nur einer Grundoperation modelliert, der rekursiven Komposition von hierarchisch klassifizierten morphophonologischen Formprimitiven. Diese Komposition ist durch zwei Eigenschaften gekennzeichnet:

1. Sie wird gesteuert durch das komplexe Zusammenspiel von *Schemata* und *Prinzipien*, die auf Merkmalstrukturen operieren. Die morphologische Kombinatorik wird dabei insbesondere durch die feinstrukturierte Kreuzklassifizierung hinsichtlich Verbkategorie (*vcat*) und Verbstamm (*v-stem*) eingegrenzt.
2. Der Komposition entspricht auf der Seite der sprachlichen Oberfläche die *Konkatenation*: Alle Formen — phonologische, morphologische und syntaktische — beruhen auf der Konkatenation von entsprechenden morphophonologischen bzw. morphosyntaktischen Formprimitiven.

Im Gegensatz zu anderen Beschreibungen, wie der traditionellen Schulgrammatik (Lewin 1990), der Morphosyntax von Jens Rickmeyer (Rickmeyer 1995) oder einer ebenfalls denkbaren Two-Level-Modellierung der japanischen Morphologie, gibt es in der hier vorgestellten Beschreibung also nur eine operationale Ebene: die Komposition von Zeichen bzw. die Konkatenation von Formprimitiven. Auf Permutationen der Elemente (*Transformationen*) und phonologische Transformationen (*phonologische Regeln*) kann damit verzichtet werden. Diese Reduktion hat mehrere Vorteile: Sie führt zu einem hierarchisch gegliederten, übersichtlichen und gut verständlichen Modell und kann leicht und effizient in computerlinguistischen Systemen implementiert werden.

Gleichbehandlung von Verben und Adjektiven

Verben und Adjektive verhalten sich im Japanischen morphologisch und syntaktisch sehr ähnlich und werden dementsprechend in dem hier vorgestellten Modell auch ähnlich modelliert: Beide Wortklassen werden durch eine gemeinsame Kategorie subsumiert und im Lexikon bezüglich ihrer Komplemente mittels des Subkategorisierungsmerkmals SUBCAT charakterisiert. Durch Agglutination entsprechender Suffixe entstehen adnominale (relativsatz- bzw. attributsatzbildende), adverbale oder satzfinale (hauptsatzbildende) Formen.

Komplexe Verben

Die Hilfsverben (助動詞, *zyodousi*) der japanischen Schulgrammatik werden, wie in Manning et al. 1999 vorgeschlagen, auf morphologischer Ebene

als Verbsuffixe modelliert und nicht, wie in Gunji 1999, durch syntaktische Einbettung. Die Kausativ-Passiv-Form eines Verbes beispielsweise erscheint damit in der Syntax als ein Wort und nicht als komplexe Konstruktion mehrerer ineinander eingebetteter Verbalphrasen.

Zur Diskussion der Vor- und Nachteile einer solchen Interpretation der verbalen Formen vergleiche Manning et al. 1999 und Gunji 1999. Dort kann auch nachgelesen werden, wie sich durch einen solchen Ansatz bedingte, anscheinende Widersprüche und Probleme — beispielsweise bei der Anaphernresolution — auflösen lassen.

Agglutination vs. Flexion

Verbformen in flektierenden Sprachen lassen sich am besten durch eine Menge nichtrekursiver morphologischer Schemata beschreiben. Alle durch die morphologische Form bestimmten Merkmale werden in diesen durch obligatorische Morpheme instantiiert. Die Werte der Merkmale, die durch die Verbformen beeinflusst werden, können also immer an den entsprechenden Morphen abgelesen werden und lassen sich daher monoton beschreiben.

Anders ist es bei agglutinierenden Sprachen: In dem vorliegenden Modell des Japanischen beispielsweise werden die Merkmale morphologischer Formen initial durch eine Menge von Grundannahmen (Defaults) beschrieben. Nur wenn ein Wert erfordert ist, der von diesen Grundannahmen abweicht, wird dieser durch ein zusätzliches Morphem ausgedrückt und der Defaultwert damit überschrieben. Das morphologische Bildungsschema ist also rekursiv.

Die Grundannahmen werden in unserem Modell durch Defaulthierarchien modelliert: Lexikalische Einträge werden durch Defaultzeichen beschrieben; die Defaultwerte, die in den Defaultzeichen spezifiziert werden, können durch Suffixe überschrieben werden.

Wechselwirkungen zwischen Morphologie, Syntax und Semantik

Der Mechanismus, durch den die Auswirkungen morphologischer Modifikationen auf die Syntax und Semantik von Verben modelliert wird, wurde im letzten Absatz beschrieben: Die lexikalischen Verbeinträge werden mit einer Defaultbeschreibung ihrer syntaktischen und semantischen Eigenschaften versehen, die durch Suffixe, wie den Kausativ, das Passiv oder den Voluntativ, modifiziert werden können. Die Beschreibungsebenen werden also nicht als voneinander unabhängige Module aufgefaßt, sondern als interagierende Teile eines komplexen Systems.

Um die Mechanismen dieser Wechselwirkungen deutlich hervortreten zu lassen, wurde die syntaktische und semantische Darstellung auf die Aspekte reduziert, die zur Darstellung der Interaktionen zwischen Syntax, Semantik

und Morphologie wesentlich sind. Das Gesamtsystem bleibt dadurch transparent und verständlich und kann leicht an detailliertere Beschreibungen angepaßt werden.

Ausblick

Wie in allen computerlinguistischen Modellen kann auch in dem hier entwickelten nur ein kleiner Ausschnitt der tatsächlichen sprachlichen Komplexität erfaßt werden. Dabei werden auf linguistischer und mathematisch-informatischer Seite zahlreiche Punkte und Problemfelder berührt, deren weitere Untersuchung eine interessante Fortführung dieser Arbeit bilden könnte. Einige dieser Punkte sollen hier kurz angeführt werden:

Ausweitung der morphologischen Beschreibung

Das vorliegende Modell beschreibt die Formen der Verben und Adjektive des modernen Japanisch. Geht man von einer Einteilung der Morphologie in Flexion und Wortbildung aus, wobei letztere weiter in Derivation und Komposition unterteilt wird, deckt das vorgelegte Modell Flexion und Teilbereiche der Derivation von Verben ab. Die Ausweitung des Modells auf die verbleibenden Bereiche der Verbmorphologie wäre damit ein interessantes Thema einer weiteren Untersuchung. Die Beschreibung der Morphologie der anderen Wortarten, die — zumindest teilweise — mit ähnlichen Mitteln behandelt werden kann, könnte sich daran anschließen.

Integration der Morphologie in eine repräsentative Beschreibung der Syntax und Semantik des Japanischen

Syntax und Semantik werden in unserem Modell nur zur Veranschaulichung ihrer Interaktion mit der Morphologie behandelt. Die Ausweitung dieses Fragmentes bzw. die Integration der vorgestellten Morphologie in eine repräsentative Beschreibung dieser sprachlichen Ebenen bietet sich daher als Fortsetzung dieser Arbeit an.

Eine solche Beschreibung könnte an Korpora des Japanischen — beispielsweise Sammlungen von Zeitungstexten, Transkriptionen des gesprochenen Japanisch etc. — evaluiert und verfeinert werden. Dieses Vorgehen wäre insbesondere deshalb interessant, weil das Japanische zur Bildung von Verschleifungen und Verkürzungen neigt, die unter anderem auch die Verbformen betreffen.

Heuristiken und korpusbasierte Methoden

Das vorliegende Modell des Japanischen läßt nur eine binäre Bewertung von Analysen zu: Sätze sind entweder ableitbar — d.h. im Sinne der Grammatik

korrekt — oder nicht ableitbar — und damit relativ zur Grammatik falsch. Sätze, deren Strukturen nur leicht von den Regeln der Grammatik abweichen, können nicht analysiert werden; Analysen von Sätzen hingegen, die von der Grammatik akzeptiert werden, sind oft in extremer Weise mehrdeutig. Die Bewertung alternativer Ergebnisse kann nur nach externen Kriterien geschehen.

Die Abschwächung der starren grammatischen Regeln verbunden mit der Integration von Heuristiken, die manuell spezifiziert oder anhand von Korpora erlernt werden, könnte diese Schwäche des vorgelegten Modells ausgleichen. Eine einfache Möglichkeit, Heuristiken auszudrücken, könnte beispielsweise darin bestehen, die lexikalischen Einträge mit Gewichtungen zu versehen, die durch die morphologischen und syntaktischen Schemata nach bestimmten Formeln zu einem Gesamtgewicht verrechnet werden.³³⁹

Optimierung der informatischen Verfahren

Merkmalslogiken, wie sie in unserem Modell zugrunde gelegt werden, sind ein mächtiges Mittel zur Modellierung komplexer Sprachen. Diese Mächtigkeit erweist sich jedoch bei der Anwendung in der Analyse und Generierung konkreter Äußerungen als Nachteil: Ihre Berechnung ist schon bei einfachen Modellen extrem zeit- und speicheraufwendig.

Eine genaue Untersuchung der spezifizierten Grammatik im Zusammenhang mit der erwünschten Berechnung ermöglicht jedoch oft wesentliche Optimierungen: Das Parsen von Formen beispielsweise kann durch spezielle Strategien, wie sie beispielsweise von Chart-Parsern verwendet werden, wesentlich optimiert werden. Entsprechend verhält es sich bei der Generierung.

Weitere Optimierungen lassen sich aus Eigenschaften der Grammatik ableiten: Die Verbmorphologie, wie sie hier beschrieben wurde, ist beispielsweise linksrekursiv. Zur ihrer Analyse reicht damit ein regulärer Automat aus. Eine Trennung von Morphologie und Syntax, ähnlich der Teilung von lexikalischer und struktureller Analyse bei Compilern in der Informatik, würde zu erheblich besserem Laufzeitverhalten führen.

Diskrete Modelle vs. "fuzzy"-Beschreibungen

Alle bisher vorgeschlagenen Fortführungen des beschriebenen sprachlichen Modells lassen einen wesentlichen Schwachpunkt unangetastet: Dieser resultiert aus der Verwendung eines Formalismus, der auf diskreten Mitteln beruht, wie sie für die klassische symbolverarbeitende KI (Künstliche Intelligenz) typisch sind.

³³⁹ vgl. Mitsuishi et al. 1998, Tsujii 1998 und Torisawa et al. 1998.

Versuche, Formalismen zu entwickeln, die — wie beispielsweise die Fuzzy-Logik, statistische Modelle³⁴⁰ oder Neuronale Netze — auf kontinuierlichen Werten aufbauen, haben in einzelnen Bereichen zu großen Verbesserungen geführt. Letztendlich scheitern die bisherigen Ansätze aber an Schwierigkeiten, Rekursionen zu modellieren. Rekursive Einbettungen sind für kontinuierliche Systeme schwer zu erkennen und zu repräsentieren; Rekursionen in Strukturen — beispielsweise der Semantik — können nicht dargestellt und daher auch nicht manipuliert werden.

Gerade die Morphologie könnte hier ein interessantes Experimentierfeld bieten: Mit ihren relativ einfachen Strukturen und Regeln bietet sie die idealen Voraussetzungen für die Erforschung und Verbesserung dieser Schwachpunkte kontinuierlicher Modelle. Der Versuch, die hier dargestellten Strukturen und Mechanismen durch Fuzzy-Logik-basierte, probabilistische bzw. neuronale Modelle abzubilden, wäre sicherlich eine der interessantesten — und schwierigsten — Möglichkeiten einer Fortführung dieser Arbeit.

³⁴⁰ vgl. Manning und Schütze 1999

A Beispielsätze zum Fragment der japanischen Verbmorphologie

Ausgangspunkt des vorgestellten morphosyntaktischen Modells war der *“The Complete Japanese Verb Guide”* (Ishii et al. 1989): Alle Formen³⁴¹, die in dieser Verbtabelle aufgeführt werden, sollten durch ein computerlinguistisches System beschrieben werden.³⁴²

Um Syntax und Semantik der Verben zu integrieren, wurden zu den Morphemen, durch die diese Verbformen beschrieben werden, aus verschiedenen Aufsätzen, Grammatiken, Wörterbüchern etc. Beispielsätze herausgesucht, anhand derer die Beschreibung entwickelt wurde. Diese Beispielsätze werden auf den folgenden Seiten zusammen mit ihrer abgeleiteten semantischen Beschreibung aufgelistet.³⁴³

Die Sätze entstammen — soweit sie nicht in Absprache mit japanischen Muttersprachlern extra gebildet wurden — den folgenden Quellen:

- | | |
|------------|---|
| (524) CJVG | <i>The Complete Japanese Verb Guide</i> , Ishii et al. 1989. |
| DBJG | <i>A Dictionary of Basic Japanese Grammar</i> ,
Makino und Tsutsui 1986. |
| IJDW | <i>IKUBUNDO Japanisch-Deutsches Wörterbuch</i> ,
Tomiya et al. 1996. |
| JM | <i>Japanische Morphosyntax</i> , Rickmeyer 1995. |
| JPSG | <i>Japanese Phrase Structure Grammar</i> , Gunji 1987. |
| LIJC | <i>The Lexical Integrity of Japanese Causatives</i> ,
Manning et al. 1999. |

Abkürzungen

Die erzeugten semantischen Strukturen sind teilweise komplex und schwer zu verstehen. In manchen Fällen werden deshalb Teilstrukturen durch leicht-

341 Eine Ausnahme bilden die Verben der Höflichkeitssprache (敬語, *keigo*), die in vielen Fällen besser als Varianten der Verben durch eigene Lexikoneinträge modelliert werden.

342 Verbtabelle können dank ihrer statischen Natur nur die häufigsten “Grundformen” veranschaulichen. Die große Anzahl möglicher Kombinationen, die aus diesen Grundformen in agglutinierenden Sprachen abgeleitet werden können, müssen vom Benutzer selbst erschlossen werden. In der hier entwickelten Beschreibung hingegen können — dank der rekursiven morphologischen Regeln — auch alle Kombinationen dieser Grundformen abgeleitet werden.

343 Syntax und Morphologie alleine führen zu einer großen Menge möglicher Analysen. Die Analysen, die der vom Sprecher intendierten Bedeutung entsprechen, können nur durch das “Verstehen” der Äußerung und die Integration von Kontext und Weltwissen herausgefiltert werden. Aus Platzgründen werden in diesem Anhang jedoch nur diese intendierten semantischen Strukturen dargestellt.

ter lesbare Abkürzungen substituiert. Wo dieses die Lesbarkeit der Analysen erleichtert, wurde weiterhin die Reihenfolge der semantischen Argumente verändert.

Für Beispielsätze, die sowohl in neutraler als auch in honorativer Form aufgeführt werden, wird nur die Semantik der letzteren Variante dargestellt. Die Semantik der neutralen Äquivalente ergibt sich durch Weglassen der Honorativ-Semantik:

$$(525) \quad \begin{bmatrix} \text{honorative} \\ \textcircled{1} \quad \boxed{1} \end{bmatrix} \longrightarrow \boxed{1}$$

Die folgenden Abkürzungen wurden verwendet:

Bei der Negation von Adjektiven:

$$(526) \quad \begin{bmatrix} \text{adjunct} \\ \textcircled{1} \quad \text{not} \\ \textcircled{2} \quad \begin{bmatrix} \text{adverbial} \\ \textcircled{1} \quad \boxed{1} \end{bmatrix} \end{bmatrix} \longrightarrow \begin{bmatrix} \text{not} \\ \textcircled{1} \quad \boxed{1} \end{bmatrix}$$

Beim Partizip:

$$(527) \quad \begin{bmatrix} \text{adjunct} \\ \textcircled{1} \quad \boxed{1} \\ \textcircled{2} \quad \begin{bmatrix} \text{participle} \\ \textcircled{1} \quad \boxed{2} \end{bmatrix} \end{bmatrix} \longrightarrow \begin{bmatrix} \text{participle} \\ \textcircled{1} \quad \boxed{2} \\ \textcircled{2} \quad \boxed{1} \end{bmatrix}$$

Beim Kontinuativ:

$$(528) \quad \begin{bmatrix} \text{adjunct} \\ \textcircled{1} \quad \text{continuation} \\ \textcircled{2} \quad \begin{bmatrix} \text{participle} \\ \textcircled{1} \quad \boxed{1} \end{bmatrix} \end{bmatrix} \longrightarrow \begin{bmatrix} \text{continuation} \\ \textcircled{1} \quad \boxed{1} \end{bmatrix}$$

Beim Partizip + kara:

$$(529) \quad \begin{bmatrix} \text{adjunct} \\ \textcircled{1} \quad \boxed{1} \\ \textcircled{2} \quad \begin{bmatrix} \text{since} \\ \textcircled{1} \quad \boxed{2} \end{bmatrix} \end{bmatrix} \longrightarrow \begin{bmatrix} \text{since} \\ \textcircled{1} \quad \boxed{2} \\ \textcircled{2} \quad \boxed{1} \end{bmatrix}$$

Beim negierten Partizip:

$$(530) \quad \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \left[\begin{array}{c} \text{participle} \\ \textcircled{1} \left[\begin{array}{c} \text{not} \\ \textcircled{1} \text{ II} \end{array} \end{array} \right] \end{array} \right] \end{array} \right] \longrightarrow \left[\begin{array}{c} \text{without-doing} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ I} \end{array} \right]$$

Beim Konditional:

$$(531) \quad \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \left[\begin{array}{c} \text{conditional} \\ \textcircled{1} \text{ II} \end{array} \right] \end{array} \right] \longrightarrow \left[\begin{array}{c} \text{if} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ I} \end{array} \right]$$

Beim Perfekt-Konditional:

$$(532) \quad \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \left[\begin{array}{c} \text{perfect-conditional} \\ \textcircled{1} \text{ II} \end{array} \right] \end{array} \right] \longrightarrow \left[\begin{array}{c} \text{if} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ I} \end{array} \right]$$

Beim Adverbial:

$$(533) \quad \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \left[\begin{array}{c} \text{adverbial} \\ \textcircled{1} \text{ II} \end{array} \right] \end{array} \right] \longrightarrow \left[\begin{array}{c} \text{adverbial} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ I} \end{array} \right]$$

Beim Verb くたさい (kudasai; bitte):

$$(534) \quad \left[\begin{array}{c} \text{imperative} \\ \textcircled{1} \left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ please-do-for-me} \\ \textcircled{2} \left[\begin{array}{c} \text{participle} \\ \textcircled{1} \text{ I} \end{array} \right] \end{array} \right] \end{array} \right] \longrightarrow \left[\begin{array}{c} \text{please-do-for-me} \\ \textcircled{1} \text{ I} \end{array} \right]$$

A.1 Flexive

A.1.1 Präsens

Satzschlußform

- (535) a. いま考えている。 ³⁴⁴

ima kañgae.t-e i-ru
jetzt denken-PART KONT-FLEX
Ich denke gerade darüber nach.

- b.
$$\left[\begin{array}{c} \textit{continuation} \\ \textcircled{1} \left[\begin{array}{c} \textit{now} \\ \textcircled{1} \left[\begin{array}{c} \textit{consider} \\ \textcircled{1} \textit{cont} \\ \textcircled{2} \textit{cont} \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (536) a. コーヒーを飲みます。 ³⁴⁵

kouhii-wo no.m.i-mas.-u
Kaffee-ACC trinken-HON-FLEX
Ich trinke Kaffee.

- b.
$$\left[\begin{array}{c} \textit{honorative} \\ \textcircled{1} \left[\begin{array}{c} \textit{drink} \\ \textcircled{1} \textit{cont} \\ \textcircled{2} \textit{coffee} \end{array} \right] \end{array} \right]$$

- (537) a. この本はおそろしく高い。 ³⁴⁶

kono hon-ha osorosi-ku taka-i
dieses Buch-TOP furchtbar-ADV teuer-FLEX
Dieses Buch ist furchtbar teuer.

- b.
$$\left[\begin{array}{c} \textit{adverbial} \\ \textcircled{1} \left[\begin{array}{c} \textit{terrible} \\ \textcircled{1} \textit{cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \textit{expensive} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \left[\begin{array}{c} \textit{this} \\ \textcircled{1} \textit{book} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

344 JM, S. 78, 20-2.1.1. b) 1.

345 JM, S. 78, 20-2.1.1. b) 1.

346 JM, S. 205, 30-2.1. a) 1.

Adnominalform

(538) a. 英語ができる人³⁴⁷

eigo-ga deki-ru hito
 Englisch-NOM können-FLEX Mensch
 jemand, der Englisch kann

b.
$$\left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ II } \textit{person} \\ \left[\begin{array}{c} \textit{can} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \left[\begin{array}{c} \textit{do} \\ \textcircled{2} \textcircled{1} \text{ II} \\ \textcircled{2} \textit{english} \end{array} \right] \end{array} \right] \end{array} \right]$$

(539) a. 長い旅でした。³⁴⁸

naga-i tabi des.it-a
 lang-FLEX Reise ist-PAST
 Es war eine lange Reise.

b.
$$\left[\begin{array}{c} \textit{past} \\ \left[\begin{array}{c} \textit{honorative} \\ \left[\begin{array}{c} \textit{is} \\ \textcircled{1} \textit{cont} \\ \left[\begin{array}{c} \textit{adjunct} \\ \textcircled{1} \text{ II } \textit{travel} \\ \textcircled{2} \left[\begin{array}{c} \textit{long} \\ \textcircled{1} \text{ II} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

A.1.2 Perfekt

Satzschlußform

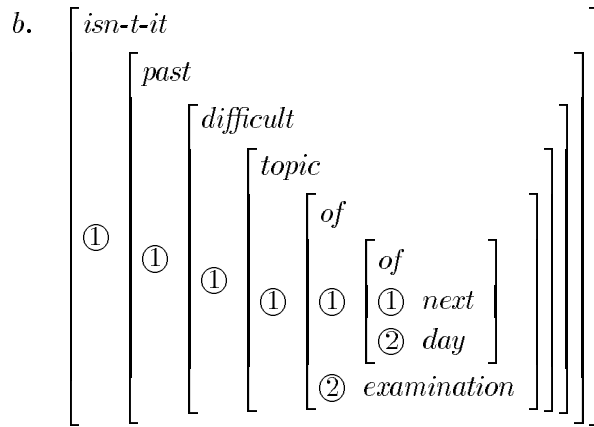
(540) a. あの人の全集は残らず読んだ。³⁴⁹

ano hito-no zeñsyuu-ha nokor-azu
 jener Mensch-GEN gesamte Werke-TOP übrigbleiben-NEG PART
yo.ñd-a
 lesen-PAST
 Ich habe seine gesamten Werke ausnahmslos gelesen.

347 JM, S. 78, 20-2.1.1. b) 2.

348 JM, S. 206, 30-2.1. a) 2.

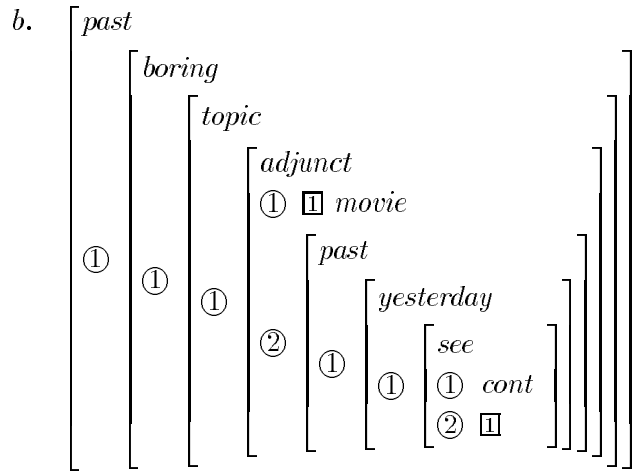
349 JM, S. 84, 20-2.2.2. a) 1.



Adnominalform

(543) a. きのう見た映画はつまらなかった。 ³⁵²

kinou mi.ta eiga-ha tumarana.ka.tt-a
 gestern sehen-PAST Film-TOP langweilig- ϕ -PAST
 Der Film, den ich gestern gesehen hatte, war langweilig



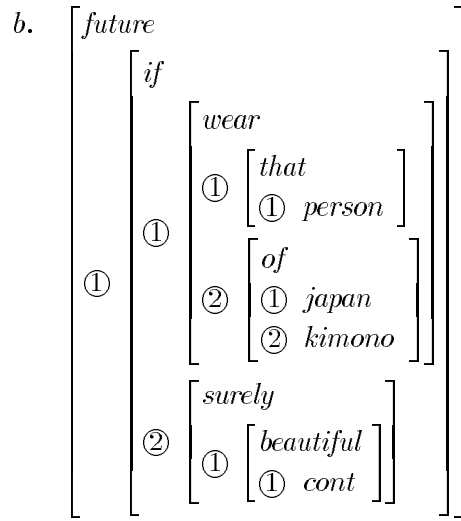
(544) a. 田中さんが食べたステーキは高かった / 高かったです。 ³⁵³

tanaka-san-ga tabe.t-a sutēki-ha taka-ka.tt-a / taka-ka.tt-a
 Herr Tanaka-NOM essen-PAST Steak-TOP teuer- ϕ -PAST / teuer- ϕ -PAST
des-u
 HON-FLEX

Das Steak, das Herr Tanaka gegessen hat, war teuer.

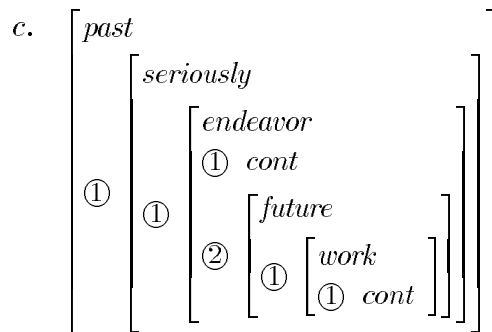
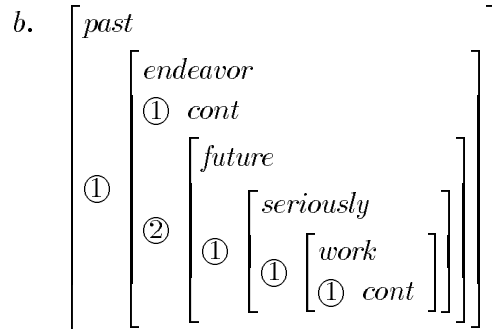
³⁵² JM, S. 84, 20-2.2.2. a) 2.

³⁵³ DBJG, S. 376, KS(A).



(547) a. 真面目に働こうと努めた。³⁵⁶

mazime-ni hatara.k-ou to tutome.t-a
 ernsthaft-ADV arbeiten-VOLI CIT bemühen-PAST
 Er bemühte sich, ernsthaft zu arbeiten.



(548) a. 私は日本歴史を読もうと思う / 思います。³⁵⁷

³⁵⁶ JM, S. 80, 20-2.1.3. c) 2.

³⁵⁷ DBJG, S. 569, KS(A).

watasi-ha nihoñ rekisi-wo yo.m-ou to omo.-u /
 ich-TOP japanische Geschichte-ACC lesen-VOLI CIT denken-FLEX /
omo.i-mas.-u
 denken-HON-FLEX

Ich glaube, ich werde (Bücher) über die japanische Geschichte lesen.

- b.
$$\left[\begin{array}{c} \text{honorable} \\ \left[\begin{array}{c} \text{think} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} i \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{future} \\ \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{cont} \\ \textcircled{2} \text{japanese-history} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (549) a. 私が彼に話しましょう。 ³⁵⁸

watasi-ga kare-ni hana.s.i-mas.-you
 ich-NOM er-DAT sprechen-HON-VOLI
 Ich werde mit ihm sprechen.

- b.
$$\left[\begin{array}{c} \text{future} \\ \left[\begin{array}{c} \text{honorable} \\ \textcircled{1} \left[\begin{array}{c} \text{speak} \\ \textcircled{1} i \\ \textcircled{2} he \end{array} \right] \end{array} \right] \end{array} \right]$$

- (550) a. 映画に行きましょう。 ³⁵⁹

eiga-ni i.k.i-mas.-you
 Film-DIR gehen-HON-VOLI
 Laßt und ins Kino gehen!

- b.
$$\left[\begin{array}{c} \text{future} \\ \left[\begin{array}{c} \text{honorable} \\ \textcircled{1} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \text{cont} \\ \textcircled{2} \text{movie} \end{array} \right] \end{array} \right] \end{array} \right]$$

358 DBJG, S. 240, KS(A).

359 DBJG, S. 240, KS(B).

A.1.4 Imperativ

- (551) a. 座れ!
- ³⁶⁰

suwa.r-e
 setzen-IMP
 Setz dich!

- b.
$$\left[\begin{array}{l} \textit{do-imperative} \\ \textcircled{1} \text{ I } \textit{cont} \\ \textcircled{2} \left[\begin{array}{l} \textit{sit} \\ \textcircled{1} \text{ I} \end{array} \right] \end{array} \right]$$

- (552) a. お金を貯めろ!
- ³⁶¹

o-kane-wo tame-ro
 HON-kane-ACC sparen-IMP
 Spar dein Geld!

- b.
$$\left[\begin{array}{l} \textit{do-imperative} \\ \textcircled{1} \text{ I } \textit{cont} \\ \textcircled{2} \left[\begin{array}{l} \textit{save} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \textit{money} \end{array} \right] \end{array} \right]$$

- (553) a. たばこを吸うな!
- ³⁶²

tabako-wo su.-u na
 Zigarette-ACC rauchen-FLEX NEGIMP
 Nicht rauchen!

- b.
$$\left[\begin{array}{l} \textit{do-not-prohibitive} \\ \textcircled{1} \text{ I } \textit{cont} \\ \textcircled{2} \left[\begin{array}{l} \textit{smoke} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \textit{cigarettes} \end{array} \right] \end{array} \right]$$

- (554) a. あんな男とは結婚するな!
- ³⁶³

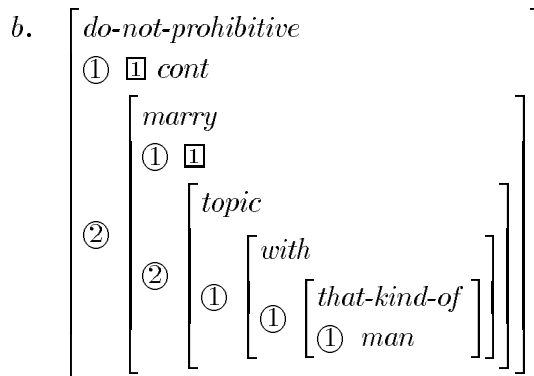
anna otoko-to-ha kekkoñ s.-uru na
 jene Art von Mann-KOM-TOP Heirat machen-FLEX NEGIMP
 Heirate keinen solchen Mann!

360 CJVG, S. 15.

361 CJVG, S. 15.

362 DBJG, S. 266, KS.

363 DBJG, S. 266, ex. (c).

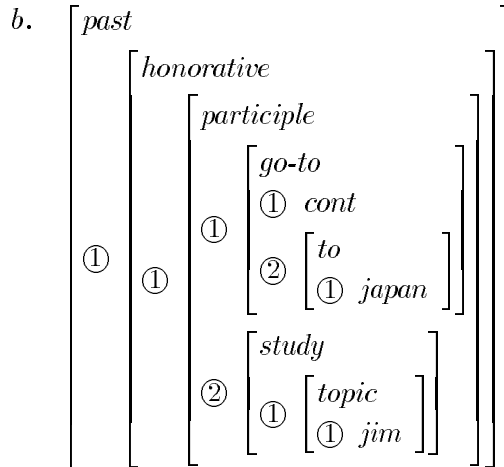


A.1.5 Partizip

Adverbialform

(555) a. ジムは日本へ行って勉強した / 勉強しました。³⁶⁴

jimu-ha nihon-he itt-e beñkyous.it-a /
 Jim-TOP Japan-DIR gehen-PART Studium machen-PAST /
beñkyous..i-mas.it-a
 Studium machen-HON-PAST
 Jim fuhr nach Japan und studierte (dort).

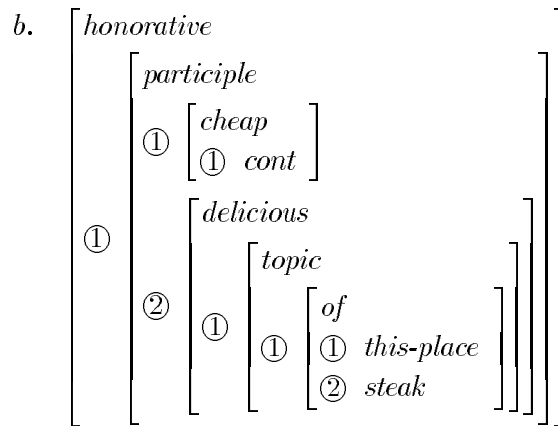


(556) a. このステーキは安くておいしい / おいしいです。³⁶⁵

koko-no sutēki-ha yasu-kute oisi-i /
 hier-GEN Steak-TOP billig-PART wohlschmeckend-FLEX /
oisi-i des.-u
 wohlschmeckend-FLEX HON-FLEX
 Hier sind Steaks billig und schmecken gut.

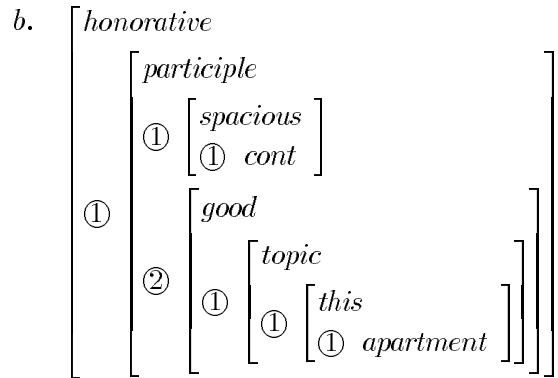
364 DBJG, S. 464, KS(1).

365 DBJG, S. 464, KS(2).



(557) a. このアパートは広くていい / いいです。 ³⁶⁶

kono apāto-ha hiro-kute i-i / i-i des.-u
 Dieses Apartment-TOP groß-PART gut-FLEX / gut-FLEX HON-FLEX
 Dieses Apartment ist groß und gut.

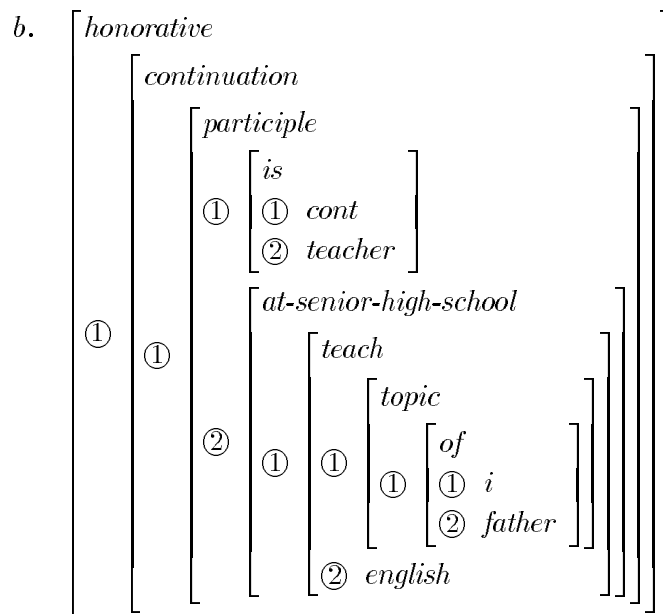


(558) a. 私の父は先生で高校で英語を教えている / います。 ³⁶⁷

watasi-no titi-ha señsei de kōkō-de eigo-wo
 Ich-GEN Vater-TOP Lehrer ist-PART Oberschule-LOC Englisch-ACC
osie.t-e i-ru / i-mas.-u
 unterrichten-PART KONT-FLEX / KONT-HON-FLEX
 Mein Vater ist Lehrer und unterrichtet Englisch an einer Oberschule.

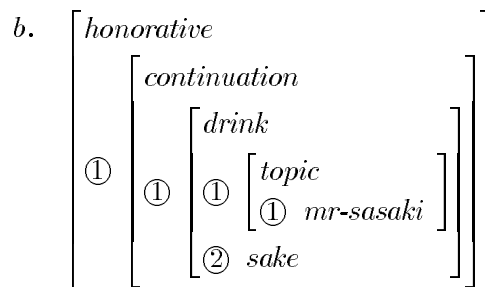
³⁶⁶ vgl. DBJG, S. 464, KS(3).

³⁶⁷ DBJG, S. 464, KS(4).



(559) a. 佐々木さんは酒を飲んでいる / います。 ³⁶⁸

sasaki-saŋ-ha sake-wo no.ñd-e i-ru /
 Herr Sasaki-TOP Sake-ACC trinken-PART KONT-FLEX /
i-mas.-u
 KONT-HON-FLEX
 Herr Sasaki trinkt (gerade) Sake.



(560) a. 和江は新聞を読んでいる。 ³⁶⁹

kazue-ha sinbuñ-wo yo.ñd-e i-ru
 Kazue-TOP Zeitung-ACC lesen-PART KONT-FLEX
 Kazue liest (gerade) Zeitung.

368 DBJG, S. 155, KS.

369 DBJG, S. 155, ex. (a).

$$b. \left[\begin{array}{c} continuation \\ \textcircled{1} \left[\begin{array}{c} read \\ \textcircled{1} \left[\begin{array}{c} topic \\ \textcircled{1} kazue \end{array} \right] \\ \textcircled{2} newspaper \end{array} \right] \end{array} \right]$$

(561) a. このりんごはくさっている。 ³⁷⁰

kono riŋgo-ha kusa.tt-e i-ru
 dieser Apfel-TOP verfaulen-PART KONT-FLEX
 Dieser Apfel ist verfault.

$$b. \left[\begin{array}{c} continuation \\ \textcircled{1} \left[\begin{array}{c} rot \\ \textcircled{1} \left[\begin{array}{c} topic \\ \textcircled{1} \left[\begin{array}{c} this \\ \textcircled{1} apple \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(562) a. 私は鈴木さんを知っています。 ³⁷¹

watasi-ha suzuki-saŋ-wo si.tt-e i-mas.-u
 ich-TOP Fräulein Suzuki-ACC kennen-PART KONT-HON-FLEX
 Ich kenne Fräulein Suzuki.

$$b. \left[\begin{array}{c} honorative \\ \textcircled{1} \left[\begin{array}{c} continuation \\ \textcircled{1} \left[\begin{array}{c} know \\ \textcircled{1} \left[\begin{array}{c} topic \\ \textcircled{1} i \end{array} \right] \\ \textcircled{2} miss-suzuki \end{array} \right] \end{array} \right] \end{array} \right]$$

(563) a. 木田さんはみんなから尊敬されている。 ³⁷²

kida-saŋ-ha miŋna-kara soŋkeis.-are.t-e i-ru
 Herr Kida-TOP alle-REL Achtung tun-PASS-PART KONT-FLEX
 Herr Kida wird von allen geachtet.

370 DBJG, S. 155, ex. (b).

371 DBJG, S. 156, ex. (d).

372 DBJG, S. 367, notes 4. (5) b.

- b.
$$\left[\begin{array}{c} \text{continuation} \\ \textcircled{1} \left[\begin{array}{c} \text{respect} \\ \textcircled{1} \left[\begin{array}{c} \text{from} \\ \textcircled{1} \text{ everybody} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ mr-kida} \end{array} \right] \end{array} \right] \end{array} \right]$$

A.1.6 Negiertes Partizip

- (564) a. あの人は怒りもせず笑っていた。 ³⁷³

ano hito-ha okori-mo s-ezu wara.tt-e
 jener Mensch-TOP sogar Ärger-ACC tun-NEG PART lachen-PART
i.t-a
 KONT-FLEX
 Er lachte, ohne sich darüber gar aufzuregen.

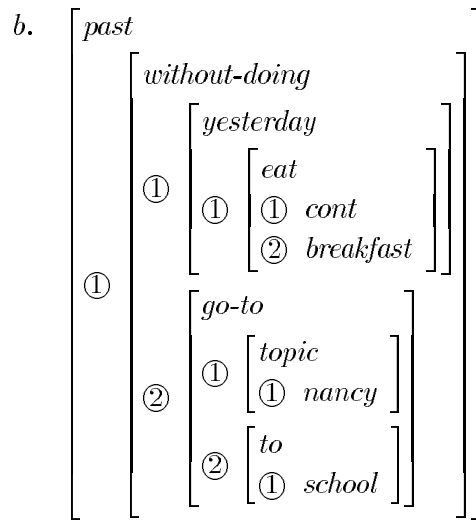
- b.
$$\left[\begin{array}{c} \text{past} \\ \textcircled{1} \left[\begin{array}{c} \text{continuation} \\ \textcircled{1} \left[\begin{array}{c} \text{without-doing} \\ \textcircled{1} \left[\begin{array}{c} \text{do} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \left[\begin{array}{c} \text{even} \\ \textcircled{1} \text{ anger} \end{array} \right] \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{laugh} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \left[\begin{array}{c} \text{that} \\ \textcircled{1} \text{ person} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (565) a. ナンシーはきのう朝ご飯を食べずに学校へ行った。 ³⁷⁴

nañsü-ha kinou asagohan-wo tabe-zuni gakkou-he
 Nancy-TOP gestern Frühstück-ACC essen-NEG PART Schule-DIR
itt-a
 gehen-PAST
 Nancy ging gestern zur Schule, ohne Frühstück zu essen.

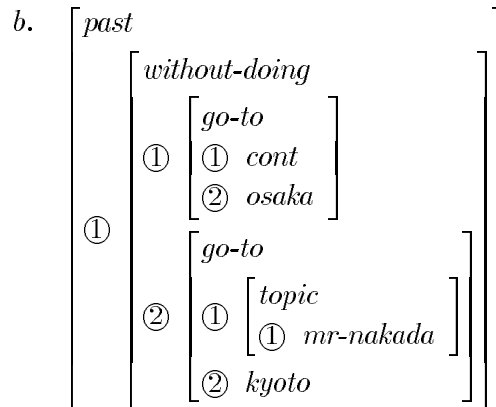
373 JM, S. 82, 20-2.1.5. b) 1.

374 DBJG, S. 272, related expressions I. [1] a.



(566) a. 中田さんは大阪に行かずに京都に行った。 ³⁷⁵

nakada-san-ha oosaka-ni i.k.a-zuni kyouto-ni itt-a
 Herr Nakada-TOP Ōsaka-DIR fahren-NEGPART Kyōto-DIR fahren-PAST
 Herr Nakada fuhr nach Kyōto, ohne nach Ōsaka zu fahren.



(567) a. 辞書を使わずに読んでください。 ³⁷⁶

jisyo-wo tuka..wa-zuni yo.ñd-e kudasai
 Wörterbuch-ACC benutzen-NEGPART lesen-PART bitte
 Bitte lesen sie es, ohne ein Wörterbuch zu benutzen.

³⁷⁵ DBJG, S. 272, related expressions I. [1] b.

³⁷⁶ DBJG, S. 272, related expressions I. [1] c.

- b.
$$\left[\begin{array}{c} \text{please-do-for-me} \\ \left[\begin{array}{c} \text{without-doing} \\ \left[\begin{array}{c} \text{use} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ dictionary} \end{array} \right] \\ \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

A.1.7 Konditional

- (568) a. 雨が降ればぬれます。 ³⁷⁷

ame-ga hu.r-eba nure-mas.-u
 Regen-NOM fallen-COND naß werden-HON-FLEX
 Wenn es regnet, wird man naß.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \left[\begin{array}{c} \text{fall} \\ \textcircled{1} \text{ rain} \end{array} \right] \\ \left[\begin{array}{c} \text{get-wet} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (569) a. これは松本先生に聞けば分かります。 ³⁷⁸

kore-ha matumoto-señsei-ni ki.k-eba waka.r.i-mas.-u
 dieses-TOP Prof. Matumoto-DAT fragen-COND verstehen-HON-FLEX
 Wenn Sie Prof. Matumoto fragen, werden sie es verstehen.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \left[\begin{array}{c} \text{ask} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ prof-matsumoto} \\ \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ this} \end{array} \right] \end{array} \right] \\ \left[\begin{array}{c} \text{be-understood} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

377 JM, S. 79, 20-2.1.2. c).

378 DBJG, S. 82, ex. (a).

- (570) a. 安ければ買います。
- ³⁷⁹

yasu-kereba ka..i-mas.-u
 billig-COND kaufen-HON-FLEX
 Wenn es billig ist, werde ich es kaufen.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{cheap} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{buy} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (571) a. もっと安ければ買いました。
- ³⁸⁰

motto yasu-kereba ka..i-mas.it-a
 mehr billig-COND kaufen-HON-PAST
 Ich hätte es gekauft, wenn es billiger gewesen wäre.

- b.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{more} \\ \textcircled{1} \left[\begin{array}{c} \text{cheap} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{buy} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

A.1.8 Perfekt-Konditional

- (572) a. 山田さんが来たら私は帰る / 帰ります。
- ³⁸¹

yamada-san-ga kit-ara watasi-ha kae.r-u
 Herr Yamada-NOM kommen-PCOND ich-TOP nach Hause fahren-FLEX
 / *kae.r.i-mas.-u*
 / nach Hause fahren-HON-FLEX
 Wenn Herr Yamada kommt, fahre ich nach Hause.

379 DBJG, S. 82, ex. (c).

380 DBJG, S. 83, notes 5. (3).

381 DBJG, S. 452, KS.

- b.
$$\left[\begin{array}{c} \text{honorable} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{come} \\ \textcircled{1} \text{ mr-yamada} \\ \textcircled{2} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{return} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} i \end{array} \right] \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (573) a. 先生に聞いたらすぐ分かった。 ³⁸²

señsei-ni ki.it-ara sugu waka.tt-a
 Lehrer-DAT fragen-PCOND sofort verstehen-PAST
 Als ich meinen Lehrer fragte, verstand ich es sofort.

- b.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{ask} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ teacher} \\ \textcircled{3} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{immediately} \\ \textcircled{1} \left[\begin{array}{c} \text{be-understood} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (574) a. おもしろかったら読みます。おもしろくなかったら読みません。 ³⁸³

omosiro-ka.tt-ara yo.m.i-mas.-u. omosiro-ku na-ka.tt-ara
 interessant-φ-PCOND lesen-HON-FLEX interessant-ADV NEG-φ-PCOND
yo.m.i-mas.e-ñ
 lesen-HON-NEG

Wenn es interessant ist, lese ich es. Ist es jedoch nicht interessant, lese ich es nicht.

382 DBJG, S. 452, ex. (a).

383 vgl. DBJG, S. 453, ex. (c).

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{interesting} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$
- c.
$$\left[\begin{array}{c} \text{not} \\ \left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{not} \\ \textcircled{1} \left[\begin{array}{c} \text{interesting} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

A.1.9 Adverbial

- (575) a. 彼女は赤くなる。 ³⁸⁴

kanozyo-ha aka-ku na.r-u
 sie-TOP rot-ADV werden-FLEX
 Sie wird rot.

- b.
$$\left[\begin{array}{c} \text{adverbial} \\ \textcircled{1} \left[\begin{array}{c} \text{red} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{become} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ she} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (576) a. その本はおもしろくない。 ³⁸⁵

sono hoñ-ha omosiro-ku na-i
 dieses Buch-TOP interessant-ADV NEG-FLEX
 Dieses Buch ist uninteressant.

³⁸⁴ vgl. JM, S. 207, 30-2.2 a) 1.

³⁸⁵ vgl. JM, S. 207, 30-2.2 a) 1.

- b.
$$\left[\begin{array}{c} not \\ \textcircled{1} \left[\begin{array}{c} interesting \\ \textcircled{1} \left[\begin{array}{c} topic \\ \textcircled{1} \left[\begin{array}{c} this \\ \textcircled{1} book \end{array} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(577) a. 少し甘く点数をつけました。 ³⁸⁶

sukosi ama-ku teñsuu-wo tuke-mas.it-a
 ein wenig großzügig-ADV Punkte-ACC anbringen-HON-PAST
 Ich habe nicht so streng zensiert.

- b.
$$\left[\begin{array}{c} past \\ \textcircled{1} \left[\begin{array}{c} honorative \\ \textcircled{1} \left[\begin{array}{c} adverbial \\ \textcircled{1} \left[\begin{array}{c} a-little \\ \textcircled{1} \left[\begin{array}{c} generous \\ \textcircled{1} cont \end{array} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(578) a. 雨がすごくふっている。 ³⁸⁷

ame-ga sugo-ku hu.tt-e i-ru
 Regen-NOM furchtbar stark-ADV fallen-PART KONT-FLEX
 Es regnet in Strömen.

- b.
$$\left[\begin{array}{c} continuation \\ \textcircled{1} \left[\begin{array}{c} adverbial \\ \textcircled{1} \left[\begin{array}{c} terrible \\ \textcircled{1} cont \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

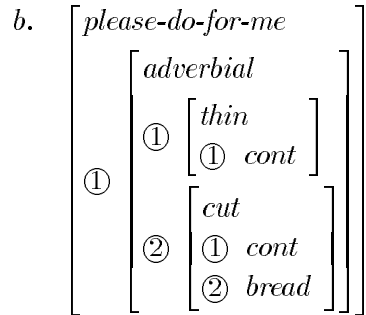
(579) a. パンを薄く切ってください。 ³⁸⁸

pañ-wo usu-ku ki.tt-e kudasai
 Brot-ACC dünn-ADV schneiden-PART bitte
 Schneide das Brot bitte dünn.

386 JM, S. 223, 32-2.1.

387 JM, S. 223, 32-2.1.

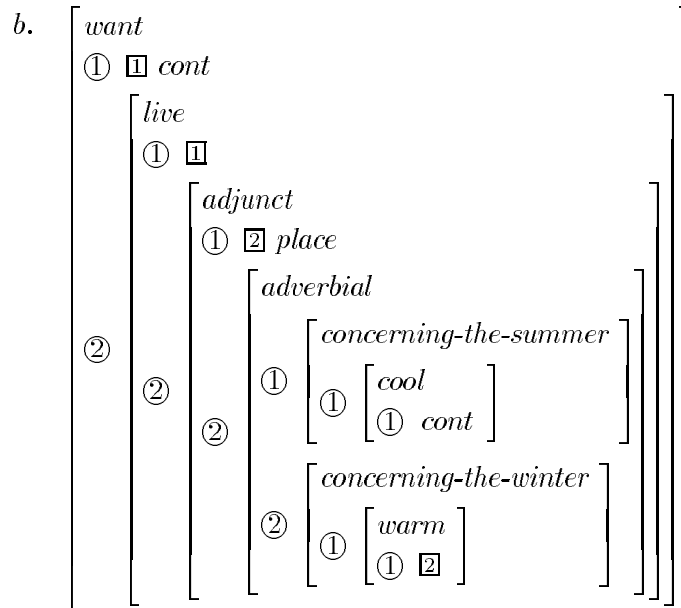
388 JM, S. 226, 32-2.3.2.



(580) a. 夏は涼しく、冬はあたたかい所に住みたい。 ³⁸⁹

natu-ha suzusi-ku, huyu-ha atataka-i tokoro-ni
 Sommer-TOP kühl-ADV, Winter-TOP warm-FLEX Ort-LOC
su.m.i-ta-i
 wohnen-VOLU-FLEX

Ich möchte dort wohnen, wo es im Sommer kühl und im Winter warm ist.



(581) a. ヘやは恐ろしくあつかった。 ³⁹⁰

heya-ha osorosi-ku atu-ka.tt-a
 Zimmer-TOP furchtbar-ADV heiß-φ-PAST

Das Zimmer war furchtbar heiß.

³⁸⁹ JM, S. 233, 33-2.

³⁹⁰ JM, S. 233, 33-2.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{adverbial} \\ \textcircled{1} \left[\begin{array}{c} \textit{terrible} \\ \textcircled{1} \textit{cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \textit{hot} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{room} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- b.
$$\left[\begin{array}{c} \text{future} \\ \textcircled{1} \left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ movie} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (585) a. この薬を飲めばよくなります。 ³⁹⁴

kono kusuri-wo no.m-eba yo-ku na.r.i-mas.-u
 diese Medizin-ACC trinken-COND gut-ADV werden-HON-FLEX
 Wenn Du diese Medizin nimmst, wird es Dir besser gehen.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{drink} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \left[\begin{array}{c} \text{this} \\ \textcircled{1} \text{ medicine} \end{array} \right] \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{adverbial} \\ \textcircled{1} \left[\begin{array}{c} \text{good} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{become} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (586) a. この映画は面白い / 面白いです。 ³⁹⁵

kono eiga-ha omosiro-i / omosiro-i des.-u
 dieser Film-TOP interessant-FLEX / interessant-FLEX HON-FLEX
 Dieser Film ist interessant.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \left[\begin{array}{c} \text{interesting} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \left[\begin{array}{c} \text{this} \\ \textcircled{1} \text{ movie} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

³⁹⁴ DBJG, S. 81, KS.

³⁹⁵ JM, S. 214, 30-5.1.2.

A.2.2 Negation

- (587) a. 本を見ないでください。
- ³⁹⁶

hon-wo mi-na-i de kudasai
 Buch-ACC sehen-NEG-FLEX PART bitte
 Sehen Sie bitte nicht ins Buch.

- b.
$$\left[\begin{array}{c} \textit{please-do-for-me} \\ \left[\begin{array}{c} \textit{not} \\ \left[\begin{array}{c} \textcircled{1} \left[\begin{array}{c} \textit{see} \\ \textcircled{1} \textit{cont} \\ \textcircled{2} \textit{book} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (588) a. 日本語を勉強しなければならない。
- ³⁹⁷

nihoŋgo-wo beŋkyou s.-ina-kereba na.r.a-na-i
 Japanisch-ACC Studium machen-NEG-COND werden-NEG-FLEX
 Du mußt Japanisch lernen.

- b.
$$\left[\begin{array}{c} \textit{not} \\ \left[\begin{array}{c} \textit{if} \\ \left[\begin{array}{c} \textit{not} \\ \left[\begin{array}{c} \textit{do} \\ \textcircled{2} \textit{cont} \\ \textcircled{1} \textit{study} \\ \textcircled{3} \textit{japanese} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \left[\begin{array}{c} \textcircled{2} \left[\begin{array}{c} \textit{become} \\ \textcircled{1} \textit{cont} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (589) a. 仕事が進まない。
- ³⁹⁸

sigoto-ga susu.m.a-na-i
 Arbeit-NOM fortschreiten-NEG-FLEX
 Die Arbeit schreitet nicht fort.

- b.
$$\left[\begin{array}{c} \textit{not} \\ \left[\begin{array}{c} \textit{progress} \\ \textcircled{1} \textit{work} \end{array} \right] \end{array} \right]$$

396 JM, S. 97, 20-3.1.2. b).

397 vgl. JM, S. 97, 20-3.1.2. c).

398 IJDW, vgl. Eintrag für 進む (*susumu*; fortschreiten).

- (590) a. おもしろかったら読みます。おもしろくなかったら読みません。³⁹⁹

omosiro-ka.tt-ara yo.m.i-mas.-u. omosiro-ku na-ka.tt-ara
 interessant- ϕ -PCOND lesen-HON-FLEX interessant-ADV NEG- ϕ -PCOND
yo.m.i-mas.e-ñ
 lesen-HON-NEG

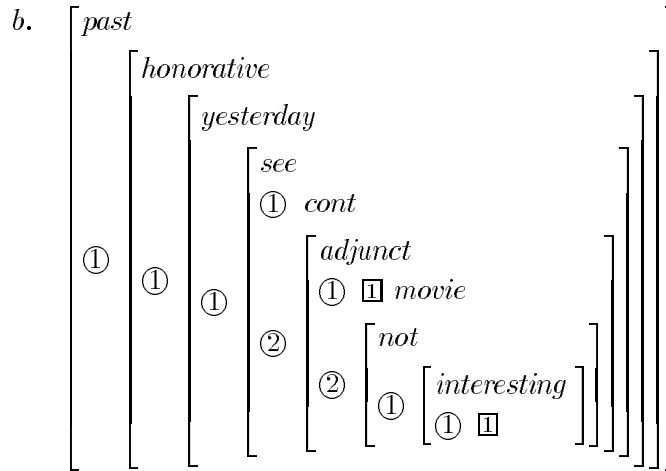
Wenn es interessant ist, lese ich es. Ist es jedoch nicht interessant, lese ich es nicht.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{interesting} \\ \textcircled{1} \text{ cont} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right]$$
- $$\left[\begin{array}{c} \text{not} \\ \textcircled{1} \left[\begin{array}{c} \text{honorative} \\ \textcircled{1} \left[\begin{array}{c} \text{if} \\ \textcircled{1} \left[\begin{array}{c} \text{not} \\ \textcircled{1} \left[\begin{array}{c} \text{interesting} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{read} \\ \textcircled{1} \text{ cont} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (591) a. きのう面白くない映画を見ました。

kinou omosiro-ku na-i eiga-wo mi-mas.it-a
 gestern interessant-ADV NEG-FLEX Film-ACC sehen-HON-PAST
 Gestern habe ich einen uninteressanten Film gesehen.

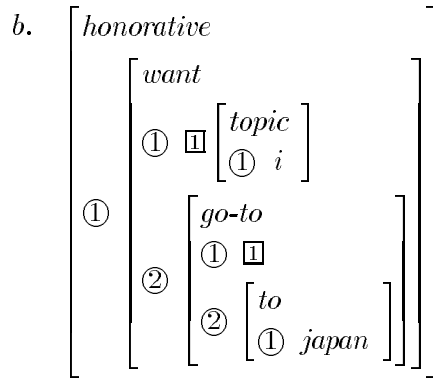
³⁹⁹ vgl. DBJG, S. 453, ex. (c).



A.2.3 Voluntativ

(592) a. 私は日本へ行きたい / 行きたいです。⁴⁰⁰

watasi-ha nihon-he i.k.i-ta-i / i.k.i-ta-i
 ich-TOP Japan-DIR fahren-VOLU-FLEX / fahren-VOLU-FLEX
des.-u
 HON-FLEX
 Ich möchte nach Japan fahren.

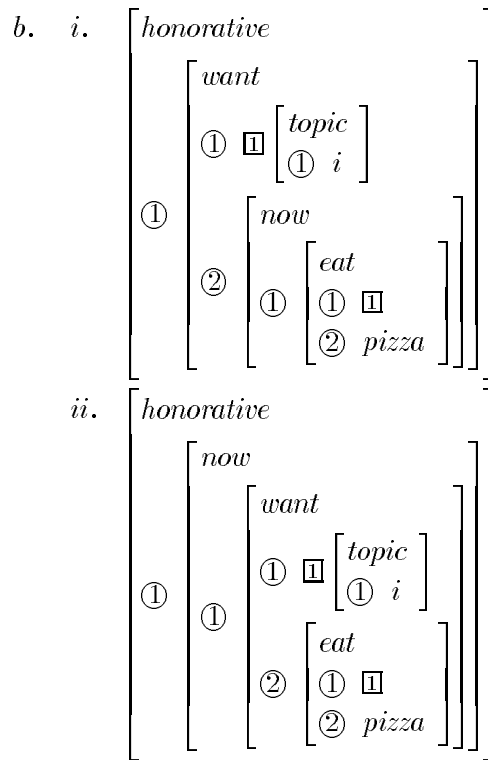


(593) a. 僕は今ピザを/が食べたい / 食べたいです⁴⁰¹

boku-ha ima piza-wo/-ga tabe-ta-i / tabe-ta-i
 ich-TOP jetzt Pizza-ACC/-NOM essen-VOLU-FLEX / essen-VOLU-FLEX
des.-u
 HON-FLEX
 Ich möchte jetzt Pizza essen.

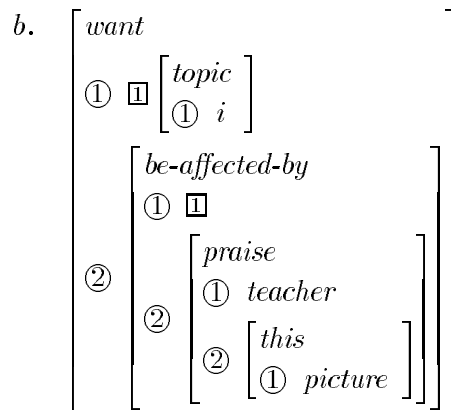
⁴⁰⁰ DBJG, S. 442, KS(A).

⁴⁰¹ DBJG, S. 442, KS(B).



(594) a. 私は先生にこの絵をほめられたい。⁴⁰²

watasi-ha seŋsei-ni kono e-wo home-rare-ta-i
 ich-TOP Lehrer-DAT dieses Bild-ACC loben-PASS-VOLU-FLEX
 Ich möchte vom Lehrer für dieses Bild gelobt werden.



(595) a. 和男はとても行きたかった。⁴⁰³

kazuo-ha totemo i.k.i-ta-ka.tt-a
 Kazuo-TOP unbedingt gehen-VOLU-φ-PAST

402 DBJG, S. 444, notes 2. (B) (7).

403 DBJG, S. 443, notes 1. (1).

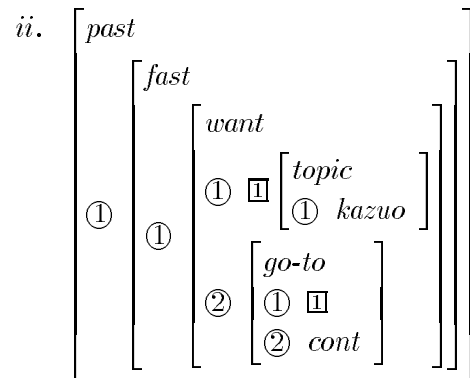
Kazuo wollte unbedingt gehen.

- b. i.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{want} \\ \textcircled{1} \text{ II} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ kazuo} \end{array} \right] \\ \textcircled{1} \left[\begin{array}{c} \text{very} \\ \textcircled{2} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$
- ii.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{very} \\ \textcircled{1} \left[\begin{array}{c} \text{want} \\ \textcircled{1} \text{ II} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ kazuo} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(596) a. 和男は早く行きたかった。

kazuo-ha hayaku i.k.i-ta-ka.tt-a
 Kazuo-TOP schnell gehen-VOLU- ϕ -PAST
 Kazuo wollte schnell gehen.

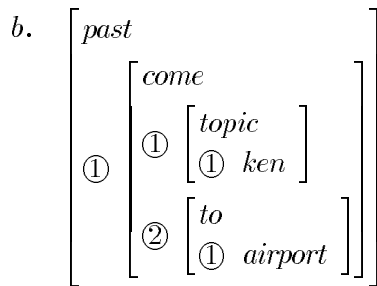
- b. i.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{want} \\ \textcircled{1} \text{ II} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ kazuo} \end{array} \right] \\ \textcircled{1} \left[\begin{array}{c} \text{fast} \\ \textcircled{2} \left[\begin{array}{c} \text{go-to} \\ \textcircled{1} \text{ II} \\ \textcircled{2} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$



A.2.4 Kausativ

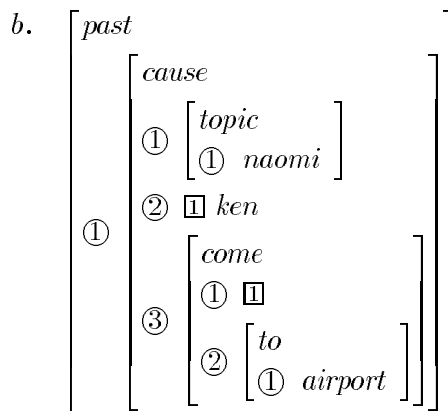
(597) a. 健は空港まで来た。

keñ-ha kuukou-made kit-a
 Ken-TOP Flughafen-DIR kommen-PAST
 Ken kam zum Flughafen.



(598) a. 直美は健を空港まで来させた。⁴⁰⁴

naomi-ha keñ-wo kuukou-made ko-sase.t-a
 Naomi-TOP Ken-ACC Flughafen-DIR kommen-CAUS-PAST
 Naomi ließ Ken zum Flughafen kommen.



⁴⁰⁴ JPSG, S. 51, (3.58)a.

- (599) a. 直美は富雄を訪ねた。

naomi-ha tomio-wo tazune.t-a
 Naomi-TOP Tomio-ACC besuchen-PAST
 Naomi besuchte Tomio.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{visit} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{naomi} \end{array} \right] \\ \textcircled{2} \textit{tomio} \end{array} \right] \end{array} \right]$$

- (600) a. 健は直美に富雄を訪ねさせた。
- ⁴⁰⁵

keñ-ha naomi-ni tomio-wo tazune-sase.t-a
 Ken-TOP Naomi-DAT Tomio-ACC besuchen-CAUS-PAST
 Ken veranlasste Namoi, Tomio zu besuchen.

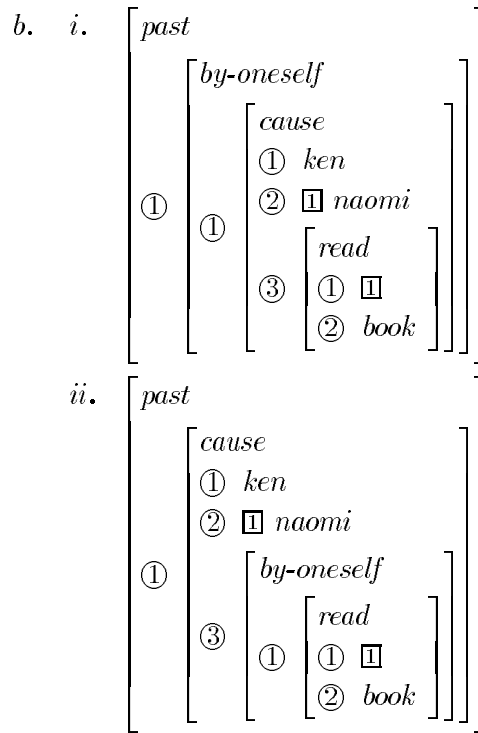
$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{cause} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{ken} \end{array} \right] \\ \textcircled{2} \boxed{\textit{naomi}} \\ \textcircled{3} \left[\begin{array}{c} \textit{visit} \\ \textcircled{1} \boxed{} \\ \textcircled{2} \textit{tomio} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (601) a. 直美が手紙を毎日書く約束をした。
- ⁴⁰⁶

naomi-ga tegami-wo mainiti ka.k-u yakusoku-wo
 Naomi-NOM Brief-ACC täglich schreiben-FLEX Versprechen-ACC
s.it-a
 machen-PAST
 Naomi versprach, jeden Tag einen Brief zu schreiben.

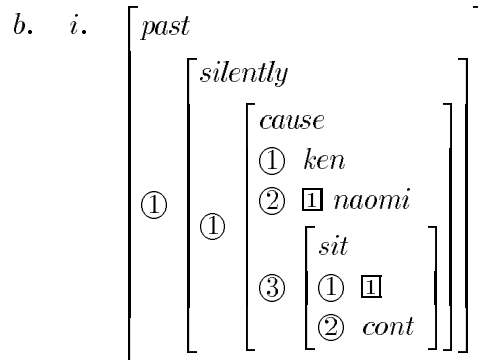
⁴⁰⁵ JPSG, S. 52, (3.58)b.

⁴⁰⁶ JPSG, S. 53, (3.63).

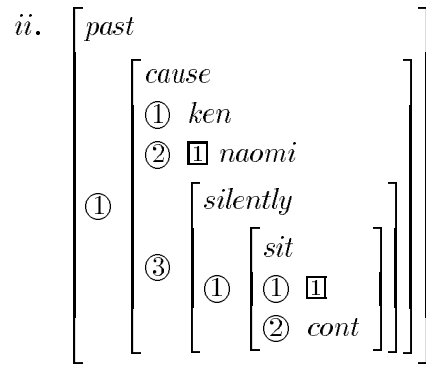


(604) a. 健が黙って直美を座らせた。⁴⁰⁹

ken-ga damatte naomi-wo suwa.r.a-se.t-a
 Ken-NOM schweigend Naomi-ACC setzen-CAUS-PAST
 Ken brachte Naomi schweigend dazu, sich zu setzen. /
 Ken brachte Naomi dazu, sich schweigend zu setzen.



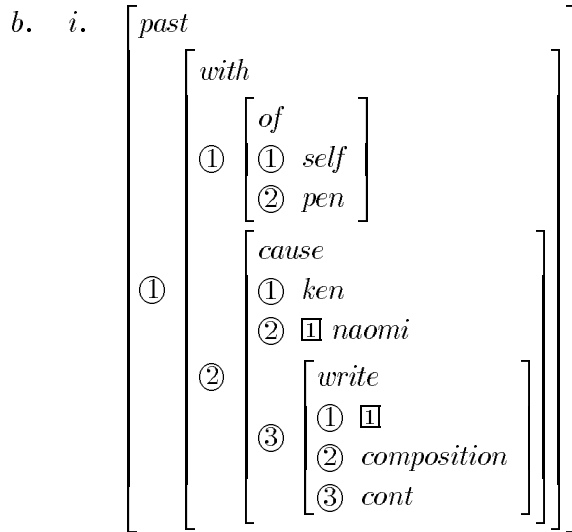
⁴⁰⁹ LIJC, S. 15, (24)b.



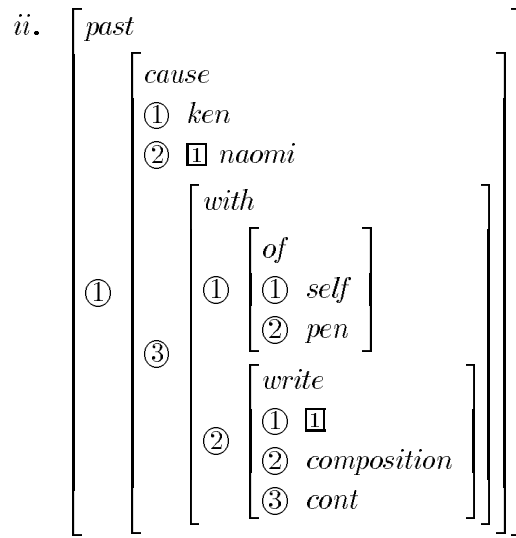
(605) a. 健が自分のペンで直美に作文を書かせた。 ⁴¹⁰

keñ-ga zibuñ-no peñ-de naomi-ni sakubun-wo
 Ken-NOM selbst-GEN Füller-ADV Naomi-DAT Aufsatz-ACC
ka.k.a-se.t-a
 schreiben-CAUS-PAST

Ken brachte Naomi mit seinem Füller dazu, einen Aufsatz zu schreiben. / Ken brachte Naomi dazu, mit ihrem Füller einen Aufsatz zu schreiben.

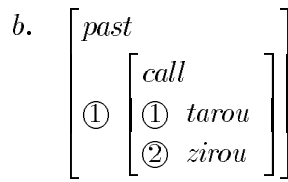


410 LIJC, S. 15, (24)c.



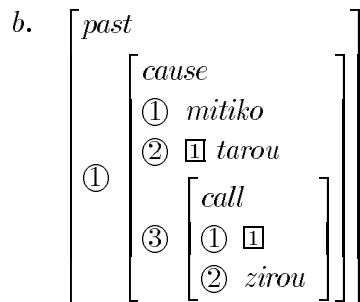
(606) a. 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
 Tarō-NOM Zirō-ACC rufen-PAST
 Tarō rief Zirō.



(607) a. 美智子が太郎に二郎を呼ばせた。⁴¹¹

mitiko-ga tarou-ni zirou-wo yo.b.a-se.t-a
 Mitiko-NOM Tarō-DAT Zirō-ACC rufen-CAUS-PAST
 Mitiko ließ Tarō Zirō rufen.



(608) a. 太郎が美智子に二郎を呼ばせられた。⁴¹²

411 LIJC, S. 34, (78)a.

412 LIJC, S. 34, (78)b.

tarou-ga mitiko-ni zirou-wo yo.b.a-se-rare.t-a
 Tarō-NOM Mitiko-DAT Zirō-ACC rufen-CAUS-PASS-PAST
 Tarō wurde von Mitiko veranlaßt, Zirō zu rufen.

- b.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{cause} \\ \textcircled{1} \text{ mitiko} \\ \textcircled{2} \textcircled{\text{I}} \text{ tarou} \\ \textcircled{3} \left[\begin{array}{c} \text{call} \\ \textcircled{1} \textcircled{\text{I}} \\ \textcircled{2} \text{ zirou} \end{array} \right] \end{array} \right] \end{array} \right]$$

A.2.5 Passiv

direkter Passiv

- (609) a. 花子は一郎をだました / だましました。⁴¹³

hanako-ha itirou-wo dama.sit-a / dama.s.i-mas.it-a
 Hanako-TOP Ichirō-ACC täuschen-PAST / täuschen-HON-PAST
 Hanako täuschte Ichirō.

- b.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{deceive} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ hanako} \end{array} \right] \\ \textcircled{2} \text{ itirou} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (610) a. 一郎は花子にだまされた / だまされました。⁴¹⁴

itirou-ha hanako-ni dama.s.a-re.t-a / dama.s.a-re-mas.it-a
 Ichirō-TOP Hanako-DAT täuschen-PASS-PAST / täuschen-PASS-HON-PAST
 Ichirō wurde von Hanako getäuscht.

- b.
$$\left[\begin{array}{c} \text{past} \\ \left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{deceive} \\ \textcircled{1} \text{ hanako} \\ \textcircled{2} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \text{ itirou} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

413 vgl. DBJG, S. 366, notes 1. (1).

414 DBJG, S. 364, KS(A).

- (611) a. ジョンは先生に質問をした。
- ⁴¹⁵

joñ-ha señsei-ni situmõñ-wo s.it-a
 John-TOP Lehrer-DAT Frage-ACC machen-PAST
 John stellte dem Lehrer eine Frage.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{do} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{john} \end{array} \end{array} \right] \\ \textcircled{2} \textit{question} \\ \textcircled{3} \textit{teacher} \end{array} \right] \end{array} \right]$$

- (612) a. 先生はジョンに質問をされた。
- ⁴¹⁶

señsei-ha joñ-ni situmõñ-wo s.-are.t-a
 Lehrer-TOP John-DAT Frage-ACC machen-PASS-PAST
 Dem Lehrer wurde von John eine Frage gestellt.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{do} \\ \textcircled{1} \textit{john} \\ \textcircled{2} \textit{question} \\ \textcircled{3} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{teacher} \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

- (613) a. この絵はピカソによってかかれた。
- ⁴¹⁷

kono e-ha pikaso-niyotte ka.k.a-re.t-a
 dieses Bild-TOP Picasso-REL malen-PASS-PAST
 Dieses Bild wurde von Picasso gemalt.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{paint} \\ \textcircled{1} \textit{picasso} \\ \textcircled{2} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \left[\begin{array}{c} \textit{this} \\ \textcircled{1} \textit{picture} \end{array} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (614) a. 電話はベルによって発明された。
- ⁴¹⁸

415 DBJG, S. 366, notes 3. (2)a.

416 DBJG, S. 366, notes 3. (2)b.

417 DBJG, S. 366, notes 4. (3)a.

deñwa-ha beru-niyotte hatumeis.-are.t-a
 Telefon-TOP Bell-REL erfinden-PASS-PAST
 Das Telefon wurde von Bell erfunden.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{invent} \\ \textcircled{1} \textit{bell} \\ \textcircled{2} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{telephone} \end{array} \end{array} \right] \end{array} \right] \end{array} \right]$$

(615) a. 学生は私に日本の大学のことを聞いた。

gakusei-ha watasi-ni nihon-no daigaku-no koto-wo
 Student-TOP ich-DAT Japan-GEN Universität-GEN Angelegenheit-ACC
ki.it-a
 fragen-PAST
 Die Studenten befragten mich über japanische Universitäten.

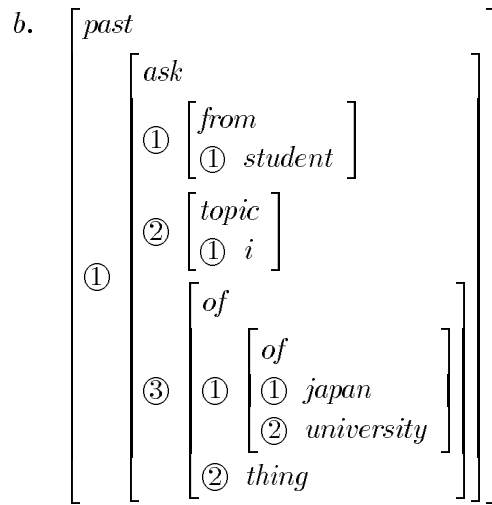
$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{ask} \\ \textcircled{1} \left[\begin{array}{c} \textit{topic} \\ \textcircled{1} \textit{student} \end{array} \right] \\ \textcircled{2} \textit{i} \\ \textcircled{3} \left[\begin{array}{c} \textit{of} \\ \textcircled{1} \left[\begin{array}{c} \textit{of} \\ \textcircled{1} \textit{japan} \\ \textcircled{2} \textit{university} \end{array} \right] \\ \textcircled{2} \textit{thing} \end{array} \right] \end{array} \right] \end{array} \right]$$

(616) a. 私は学生から日本の大学のことを聞かれた。⁴¹⁹

watasi-ha gakusei-kara nihon-no daigaku-no koto-wo
 ich-TOP Student-REL Japan-GEN Universität-GEN Angelegenheit-ACC
ki.k.a-re.t-a
 fragen-PASS-PAST
 Ich wurde von den Studenten über japanische Universitäten befragt.

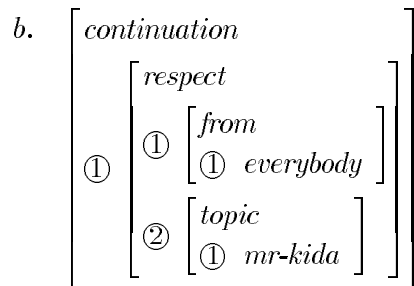
418 DBJG, S. 366, notes 4. (3)b.

419 DBJG, S. 367, notes 4. (5)a.



(617) a. 木田さんはみんなから尊敬されている。⁴²⁰

kida sañ-ha miñna-kara soñkeis.-are.t-e i-ru
 Herr Kida-TOP alle-REL respektieren-PASS-PART KONT-FLEX
 Herr Kida wird von allen respektiert.

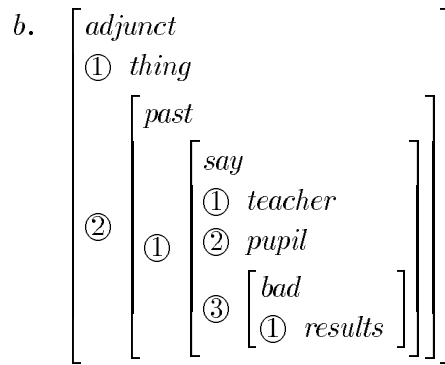


(618) a. 先生が生徒に成績が悪いと言ったこと。⁴²¹

señsei-ga seito-ni seiseki-ga waru-i to i.tt-a
 Lehrer-NOM Schüler-DAT Ergebnis-NOM schlecht-FLEX CIT sagen-PAST
koto
 Sachverhalt
 [...] daß der Lehrer dem Schüler gesagt hat, er habe schlechte
 Zensuren.

⁴²⁰ DBJG, S. 367, notes 4. (5)b.

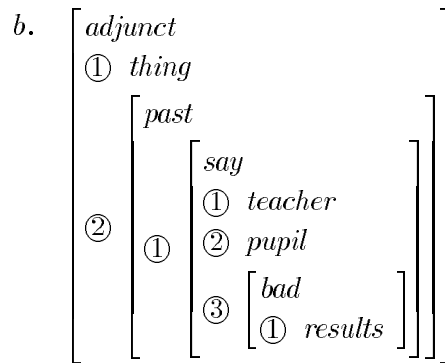
⁴²¹ JM, S. 93, 3.1.1.1. b) 1.4.



(619) a. 生徒が先生に成績が悪いと言われたこと。⁴²²

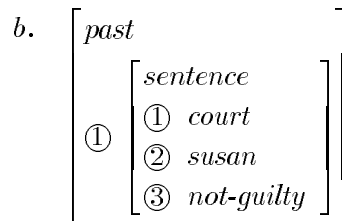
seito-ga señsei-ni seiseki-ga waru-i to
 Schüler-NOM Lehrer-DAT Ergebnis-NOM schlecht-FLEX CIT
i..wa-re.t-a koto
 sagen-PASS-PAST Sachverhalt

[...] daß dem Schüler vom Lehrer gesagt worden ist, er habe schlechte Zensuren.



(620) a. 裁判所がスサンに無罪を言い渡した。⁴²³

saibansyo-ga susan-ni muzai-wo iwata.sit-a
 Gericht-NOM Susan-DAT Unschuld-ACC urteilen-PAST
 Das Gericht sprach Susan frei.



422 JM, S. 93, 3.1.1.1. b) 1.4.

423 JPSCG, S. 64, (3.84).

- (621) a. スサンが裁判所に無罪を言い渡された。
- ⁴²⁴

susān-ga saibansyo-ni muzai-wo iiwata.s.a-re.t-a
 Susan-NOM Gericht-DAT Unschuld-ACC urteilen-PASS-PAST
 Susan wurde vom Gericht freigesprochen.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{sentence} \\ \textcircled{1} \textit{court} \\ \textcircled{2} \textit{susan} \\ \textcircled{3} \textit{not-guilty} \end{array} \right] \end{array} \right]$$

- (622) a. 太郎が二郎を呼んだ。

tarou-ga zirou-wo yo.ñd-a
 Tarō-NOM Zirō-ACC rufen-PAST
 Tarō rief Zirō.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{call} \\ \textcircled{1} \textit{tarou} \\ \textcircled{2} \textit{zirou} \end{array} \right] \end{array} \right]$$

- (623) a. 美智子が太郎に二郎を呼ばせた。
- ⁴²⁵

mitiko-ga tarou-ni zirou-wo yo.b.a-se.t-a
 Mitiko-NOM Tarō-DAT Zirō-ACC rufen-CAUS-PAST
 Mitiko ließ Tarō Zirō rufen.

$$b. \left[\begin{array}{c} \textit{past} \\ \textcircled{1} \left[\begin{array}{c} \textit{cause} \\ \textcircled{1} \textit{mitiko} \\ \textcircled{2} \textcircled{1} \textit{tarou} \\ \textcircled{3} \left[\begin{array}{c} \textit{call} \\ \textcircled{1} \textcircled{1} \\ \textcircled{2} \textit{zirou} \end{array} \right] \end{array} \right] \end{array} \right]$$

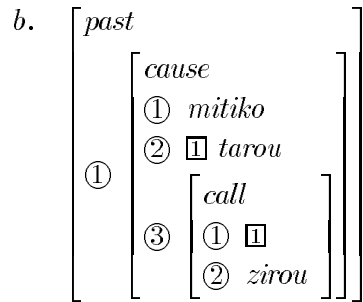
- (624) a. 太郎が美智子に二郎を呼ばせられた。
- ⁴²⁶

tarou-ga mitiko-ni zirou-wo yo.b.a-se-rare.t-a
 Tarō-NOM Mitiko-DAT Zirō-ACC rufen-CAUS-PASS-PAST
 Tarō wurde von Mitiko veranlaßt, Zirō zu rufen.

424 JPSG, S. 63, (3.83)a.

425 LIJC, S. 34, (78)a.

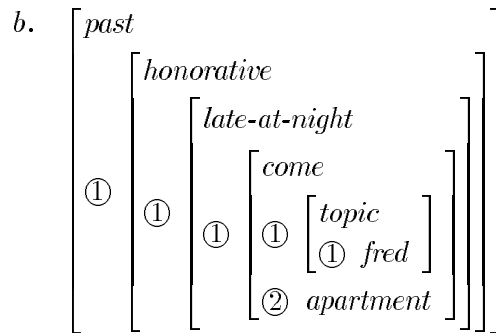
426 LIJC, S. 34, (78)b.



indirekter Passiv

(625) a. フレッドは夜おそくアパートに来た / 来ました。

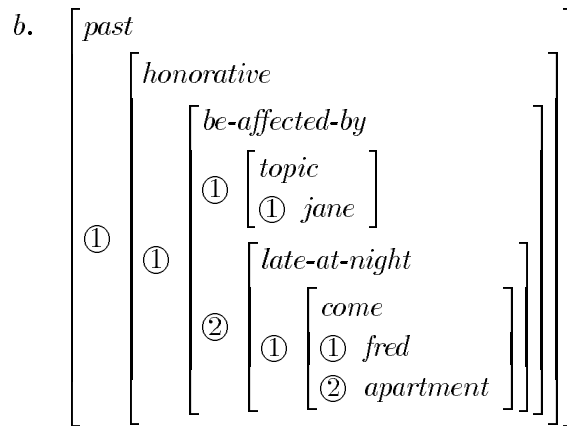
fureddo-ha yoru osoku apāto-ni kit-a /
 Fred-TOP Nacht spät Apartment-DIR kommen-PAST /
ki-mas.it-a
 kommen-HON-PAST
 Fred kam spät in der Nacht ins Apartment.



(626) a. ジェーンはフレッドに夜おそくアパートに来られた / 来られました。⁴²⁷

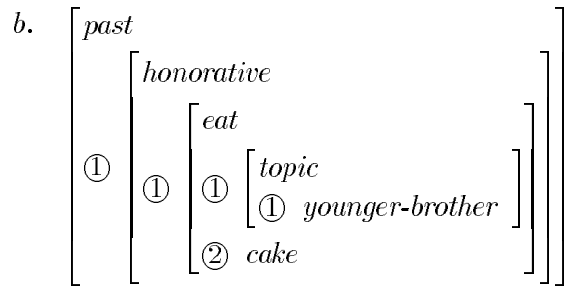
jēn-ha fureddo-ni yoru osoku apāto-ni ko-rare.t-a
 Jane-TOP Fred-DAT Nacht spät Apartment-DIR kommen-PASS-PAST
 / *ko-rare-mas.it-a*
 / kommen-PASS-HON-PAST
 Fred kam spät in der Nacht in Janes Apartment (worüber sich Jane
 nicht freute).

427 DBJG, S. 364, KS(B).



(627) a. 弟はケーキを食べた / 食べました。⁴²⁸

otouto-ha kēki-wo tabe.t-a / tabe-mas.it-a
 jüngerer Bruder-TOP Kuchen-ACC essen-PAST / essen-HON-PAST
 Mein jüngerer Bruder aß den Kuchen.

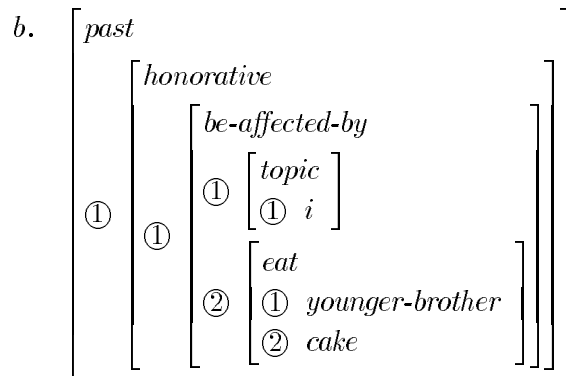


(628) a. 私は弟にケーキを食べられた / 食べられました。⁴²⁹

watasi-ha otouto-ni kēki-wo tabe-rare.t-a /
 ich-TOP jüngerer Bruder-DAT Kuchen-ACC essen-PASS-PAST /
tabe-rare-mas.it-a
 essen-PASS-HON-PAST
 Mein jüngerer Bruder aß den Kuchen. (worüber ich mich nicht freute.)
 / “Mein jüngerer Bruder aß mir den Kuchen weg.”

⁴²⁸ DBJG, S. 364, KS(C).

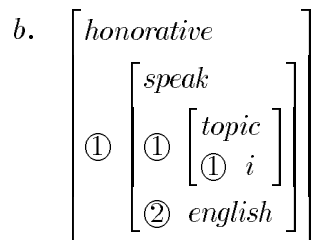
⁴²⁹ DBJG, S. 364, KS(C).



A.2.6 Potential

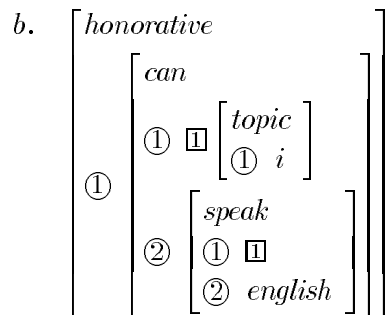
(629) a. 私は英語を話します。 ⁴³⁰

watasi-ha eigo-wo hana.s.i-mas.-u
 ich-TOP Englisch-ACC sprechen-HON-FLEX
 Ich spreche Englisch.



(630) a. 私は英語が/を話せます。 ⁴³¹

watasi-ha eigo-ga/-wo hana.s-e-mas.-u
 ich-TOP Englisch-NOM/-ACC sprechen-POT-HON-FLEX
 Ich kann Englisch sprechen.



(631) a. 私は日本語が読める / 読めます。 ⁴³²

430 DBJG, S. 371, notes (1)a.

431 DBJG, S. 371, notes (1)b.

watasi-ha nihoŋgo-ga yo.m-e-ru / yo.m-e-mas.-u
 ich-TOP Japanisch-NOM lesen-POT-FLEX / lesen-POT-HON-FLEX
 Ich kann Japanisch lesen.

- b.
$$\left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{can} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} i \end{array} \right] \\ \textcircled{2} \left[\begin{array}{c} \text{read} \\ \textcircled{1} \left[\begin{array}{c} \text{read} \\ \textcircled{2} \text{ japanese} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(632) a. この水は飲めない / 飲めません。⁴³³

kono mizu-ha no.m-e-na-i / no.m-e-mas.e-ñ
 dieses Wasser-TOP trinken-POT-NEG-FLEX / trinken-POT-HON-NEG
 Dieses Wasser ist nicht trinkbar.

- b.
$$\left[\begin{array}{c} \text{not} \\ \left[\begin{array}{c} \text{honorative} \\ \left[\begin{array}{c} \text{can} \\ \textcircled{1} \left[\begin{array}{c} \text{cont} \\ \textcircled{1} \left[\begin{array}{c} \text{drink} \\ \textcircled{1} \left[\begin{array}{c} \text{topic} \\ \textcircled{1} \left[\begin{array}{c} \text{this} \\ \textcircled{1} \text{ water} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

(633) a. 彼が漢字を書くこと⁴³⁴

kare-ga kañzi-wo ka.k-u koto
 er-NOM chinesische Zeichen-ACC schreiben-FLEX Sachverhalt
 [...] daß er chinesische Zeichen schreibt.

- b.
$$\left[\begin{array}{c} \text{adjunct} \\ \textcircled{1} \text{ thing} \\ \left[\begin{array}{c} \text{write} \\ \textcircled{1} \text{ he} \\ \textcircled{2} \text{ kanji} \end{array} \right] \end{array} \right]$$

432 DBJG, S. 370, KS(A).

433 DBJG, S. 370, KS(B).

434 JM, S. 94, 20-3.1.1.2. b) 1.

- (634) a. 彼が^s/に漢字が書けること⁴³⁵

kare-ga/-ni kañzi-ga ka.k-e-ru koto
 er-NOM/-DAT chinesische Zeichen-NOM schreiben-POT-FLEX Sachverhalt
 [...] daß er chinesische Zeichen schreiben kann.

- b.
$$\left[\begin{array}{l} \text{adjunct} \\ \textcircled{1} \text{ thing} \\ \left[\begin{array}{l} \text{can} \\ \textcircled{1} \text{ I he} \\ \textcircled{2} \left[\begin{array}{l} \text{write} \\ \textcircled{1} \text{ I} \\ \textcircled{2} \text{ kanji} \end{array} \right] \end{array} \right] \end{array} \right]$$

A.2.7 Hilfssuffix ‘Karu’

- (635) a. あの人が日本の着物を着たら、きっと美しかろう。 ⁴³⁶

ano hito-ga nihon-no kimono-wo ki.t-ara kitto
 jener Mensch-NOM Japan-GEN Kimono-ACC tragen-PCOND bestimmt
utukusi-ka.r-ou
 schön-φ-VOLI

Sie wird bestimmt schön aussehen, wenn sie einen japanischen Kimono trägt.

- b.
$$\left[\begin{array}{l} \text{future} \\ \left[\begin{array}{l} \text{if} \\ \textcircled{1} \left[\begin{array}{l} \text{wear} \\ \textcircled{1} \left[\begin{array}{l} \text{that} \\ \textcircled{1} \text{ person} \end{array} \right] \\ \textcircled{2} \left[\begin{array}{l} \text{of} \\ \textcircled{1} \text{ japan} \\ \textcircled{2} \text{ kimono} \end{array} \right] \end{array} \right] \\ \textcircled{2} \left[\begin{array}{l} \text{surely} \\ \textcircled{1} \left[\begin{array}{l} \text{beautiful} \\ \textcircled{1} \text{ cont} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

- (636) a. 次の日の試験は難しかったなあ。 ⁴³⁷

435 JM, S. 94, 20-3.1.1.2. b) 1.

436 JM, S. 209, 30-3.1.1. c) 2.

437 vgl. JM, S. 210, 30-3.1.1. c) 6.

B Verbtafeln

Ausgangspunkt der vorgestellten Beschreibung waren die Verbtafeln des *“The Complete Japanese Verb Guide”* (Ishii et al. 1989). Die in diesem Buch aufgelisteten Verbformen sollten durch ein System morphologischer Regeln beschrieben werden. Um die aufgestellten Regeln zu überprüfen, wurde für jede Verbklasse mindestens ein repräsentatives Verb ausgewählt. Die durch das System generierten Formen wurden mit denen des *“The Complete Japanese Verb Guide”* durch ein Programm verglichen. Die dabei ermittelten Fehler wurden analysiert und korrigiert. Dieser Prozess wurde so lange wiederholt, bis alle Formen korrekt gebildet wurden.

Da die Auswahl der Verben repräsentativ für die japanische Verbmorphologie ist, kann gesagt werden, daß (von den schon in Anhang A erwähnten Formen der Höflichkeitssprache abgesehen) *alle* Formen des *“The Complete Japanese Verb Guide”* — und damit auch alle aus ihnen abgeleiteten komplexen Formen — durch das vorgestellte System korrekt beschrieben werden.

Um einen Überblick über die bewerteten Verbformen zu geben, werden die generierten Verbtafeln der ausgewählten Verben in diesem Anhang dargestellt.

Zur Ableitung der Verbtafeln

Da im *“The Complete Japanese Verb Guide”* (CJVG) auch synthetische Formen vorkommen, wurde für diese die dort verwendete Nomenklatur übernommen. Alle anderen Bezeichnungen stimmen mit den in Kapitel 5 *“Ein Fragment der japanischen Verbmorphologie”* verwendeten überein. Die Ableitung der Formen ergibt sich aus der folgenden Tabelle:⁴⁴¹

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	-FLEX	-NEG-FLEX
	FORMELL	-HON-FLEX	-HON-NEG
PERFEKT	NEUTRAL	-PAST	-NEG- ϕ -PAST
	FORMELL	-HON-PAST	-HON-NEG DA-HON-PAST
PARTIZIP		-PART	-NEG-PART
IMPERATIV		-IMP	-FLEX NEGIMP
VOLITIONAL	NEUTRAL	-VOLI	
	FORMELL	-HON-VOLI	
PRÄSUMPTIV	NEUTRAL	-FLEX DA-VOLI	-NEG-FLEX DA-VOLI
	FORMELL	-FLEX DA-HON-VOLI	-NEG-FLEX DA-HON-VOLI
KONDITIONAL	NEUTRAL	-COND	-NEG-COND
		-PCOND	-NEG- ϕ -PCOND
	FORMELL	-HON-PCOND	-HON-NEG DA-HON-PCOND
POTENTIAL		-POT-FLEX	
		-POT-FLEX	
PASSIV		-PASS-FLEX	
KAUSATIV		-CAUS-FLEX	
KAUSATIV-PASSIV		-CAUS-PASS-FLEX	
		-CAUS-PASS-FLEX	

⁴⁴¹ In [eckigen Klammern] gesetzte Formen werden im CJVG nicht angegeben; Formen in (runden Klammern) werden nur eingeschränkt verwendet und stehen auch im CJVG in runden Klammern. Adjektive und das Verb *だ* (*da*; sein) werden im CJVG nicht aufgelistet.

いる – *iru* – sein, <Kontinuative>

v_v

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	いる <i>iru</i>	いない <i>inai</i>
	FORMELL	います <i>imasu</i>	いません <i>imaseñ</i>
PERFEKT	NEUTRAL	いた <i>ita</i>	いなかった <i>inakatta</i>
	FORMELL	いました <i>imasita</i>	いませんでした <i>imaseñ desita</i>
PARTIZIP		いて <i>ite</i>	いなくて <i>inakute</i>
IMPERATIV		いろ <i>iro</i>	いるな <i>iru na</i>
VOLITIONAL	NEUTRAL	いよう <i>iyou</i>	
	FORMELL	いましょう <i>imasyou</i>	
PRÄSUMPTIV	NEUTRAL	いるだろう <i>iru darou</i>	いないだろう <i>inai darou</i>
	FORMELL	いるでしょう <i>iru desyou</i>	いないでしょう <i>inai desyou</i>
KONDITIONAL	NEUTRAL	いれば <i>ireba</i>	いなければ <i>inakereba</i>
		いたら <i>itara</i>	いなかったら <i>inakattara</i>
	FORMELL	いましたら <i>imasitara</i>	いませんでしたら <i>imaseñ desitara</i>
POTENTIAL		いれる <i>ireru</i>	
		いられる <i>irareru</i>	
PASSIV		いられる <i>irareru</i>	
KAUSATIV		いさせる <i>isaseru</i>	
KAUSATIV-PASSIV		いさせられる <i>isaserareru</i>	
		[いざされる] <i>isasareru</i>	

見る – *miru* – sehen v_v

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	見る <i>miru</i>	見ない <i>minai</i>
	FORMELL	見ます <i>mimasu</i>	見ません <i>mimaseñ</i>
PERFEKT	NEUTRAL	見た <i>mita</i>	見なかった <i>minakatta</i>
	FORMELL	見ました <i>mimasita</i>	見ませんでした <i>mimaseñ desita</i>
PARTIZIP		見て <i>mite</i>	見なくて <i>minakute</i>
IMPERATIV		見ろ <i>miro</i>	見るな <i>miru na</i>
VOLITIONAL	NEUTRAL	見よう <i>miyou</i>	
	FORMELL	見ましょう <i>mimasyou</i>	
PRÄSUMPTIV	NEUTRAL	見るだろう <i>miru darou</i>	見ないだろう <i>minai darou</i>
	FORMELL	見るでしょう <i>miru desyou</i>	見ないでしょう <i>minai desyou</i>
KONDITIONAL	NEUTRAL	見れば <i>mireba</i>	見なければ <i>minakereba</i>
		見たら <i>mitara</i>	見なかったら <i>minakattara</i>
	FORMELL	見ましたら <i>mimasitara</i>	見ませんでしたら <i>mimaseñ desitara</i>
POTENTIAL		見れる <i>mireru</i>	
		見られる <i>mirareru</i>	
PASSIV		見られる <i>mirareru</i>	
KAUSATIV		見させる <i>misaseru</i>	
KAUSATIV-PASSIV		見させられる <i>misaserareru</i>	
		〔見さされる〕 <i>misasareru</i>	

寝る – *neru* – schlafen v_v

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	寝る <i>neru</i>	寝ない <i>nenai</i>
	FORMELL	寝ます <i>nemasu</i>	寝ません <i>nemaseñ</i>
PERFEKT	NEUTRAL	寝た <i>neta</i>	寝なかった <i>nenakatta</i>
	FORMELL	寝ました <i>nemasita</i>	寝ませんでした <i>nemaseñ desita</i>
PARTIZIP		寝て <i>nete</i>	寝なくて <i>nenakute</i>
IMPERATIV		寝ろ <i>nero</i>	寝るな <i>neru na</i>
VOLITIONAL	NEUTRAL	寝よう <i>neyou</i>	
	FORMELL	寝ましょう <i>nemasyou</i>	
PRÄSUMPTIV	NEUTRAL	寝るだろう <i>neru darou</i>	寝ないだろう <i>nenai darou</i>
	FORMELL	寝るでしょう <i>neru desyou</i>	寝ないでしょう <i>nenai desyou</i>
KONDITIONAL	NEUTRAL	寝れば <i>nereba</i>	寝なければ <i>nenakereba</i>
		寝たら <i>netara</i>	寝なかったら <i>nenakattara</i>
	FORMELL	寝ましたら <i>nemasitara</i>	寝ませんでしたら <i>nemaseñ desitara</i>
POTENTIAL		[寝れる] <i>nereru</i> 寝られる <i>nerareru</i>	
PASSIV		寝られる <i>nerareru</i>	
KAUSATIV		寝させる <i>nesaseru</i>	
KAUSATIV-PASSIV		寝させられる <i>nesaserareru</i> [寝さされる] <i>nesasareru</i>	

食べる – *taberu* – essen

v_v

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	食べる <i>taberu</i>	食べない <i>tabenai</i>
	FORMELL	食べます <i>tabemasu</i>	食べません <i>tabemaseñ</i>
PERFEKT	NEUTRAL	食べた <i>tabeta</i>	食べなかった <i>tabenakatta</i>
	FORMELL	食べました <i>tabemasita</i>	食べませんでした <i>tabemaseñ desita</i>
PARTIZIP		食べて <i>tabete</i>	食べなくて <i>tabenakute</i>
IMPERATIV		食べろ <i>tabero</i>	食べるな <i>taberu na</i>
VOLITIONAL	NEUTRAL	食べよう <i>tabeyou</i>	
	FORMELL	食べましょう <i>tabemasyou</i>	
PRÄSUMPTIV	NEUTRAL	食べるだろう <i>taberu darou</i>	食べないだろう <i>tabenai darou</i>
	FORMELL	食べるでしょう <i>taberu desyou</i>	食べないでしょう <i>tabenai desyou</i>
KONDITIONAL	NEUTRAL	食べれば <i>tabereba</i>	食べなければ <i>tabenakereba</i>
		食べたら <i>tabetara</i>	食べなかったら <i>tabenakattara</i>
	FORMELL	食べましたら <i>tabemasitara</i>	食べませんでしたら <i>tabemaseñ desitara</i>
POTENTIAL		食べれる <i>tabereru</i>	
		食べられる <i>taberareru</i>	
PASSIV		食べられる <i>taberareru</i>	
KAUSATIV		食べさせる <i>tabesaseru</i>	
KAUSATIV-PASSIV		食べさせられる <i>tabesaserareru</i>	
		〔食べさされる〕 <i>tabesasareru</i>	

書く – *kaku* – schreiben v_{c_k}

		AFFIRMATIV	NEGIIERT
PRÄSENS	NEUTRAL	書く <i>kaku</i>	書かない <i>kakanai</i>
	FORMELL	書きます <i>kakimasu</i>	書きません <i>kakimaseñ</i>
PERFEKT	NEUTRAL	書いた <i>kaita</i>	書かなかった <i>kakanakatta</i>
	FORMELL	書きました <i>kakimasita</i>	書きませんでした <i>kakimaseñ desita</i>
PARTIZIP		書いて <i>kaite</i>	書かなくて <i>kakanakute</i>
IMPERATIV		書け <i>kake</i>	書くな <i>kaku na</i>
VOLITIONAL	NEUTRAL	書こう <i>kakou</i>	
	FORMELL	書きましょう <i>kakimasyou</i>	
PRÄSUMPTIV	NEUTRAL	書くだろう <i>kaku darou</i>	書かないだろう <i>kakanai darou</i>
	FORMELL	書くでしょう <i>kaku desyou</i>	書かないでしょう <i>kakanai desyou</i>
KONDITIONAL	NEUTRAL	書けば <i>kakeba</i>	書かなければ <i>kakanakereba</i>
		書いたら <i>kaitara</i>	書かなかったら <i>kakanakattara</i>
	FORMELL	書きましたら <i>kakimasitara</i>	書きませんでしたら <i>kakimaseñ desitara</i>
POTENTIAL		書ける <i>kakeru</i>	
PASSIV		書かれる <i>kakareru</i>	
KAUSATIV		書かせる <i>kakaseru</i>	
KAUSATIV-PASSIV		書かせられる <i>kakaserareru</i>	
		書かされる <i>kakasareru</i>	

行く – *iku* – gehen, fahren v_{ciku}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	行く <i>iku</i>	行かない <i>ikanai</i>
	FORMELL	行きます <i>ikimasu</i>	行きません <i>ikimaseñ</i>
PERFEKT	NEUTRAL	行った <i>itta</i>	行かなかった <i>ikanakatta</i>
	FORMELL	行きました <i>ikimasita</i>	行きませんでした <i>ikimaseñ desita</i>
PARTIZIP		行って <i>itte</i>	行かなくて <i>ikanakute</i>
IMPERATIV		行け <i>ike</i>	行くな <i>iku na</i>
VOLITIONAL	NEUTRAL	行こう <i>ikou</i>	
	FORMELL	行きましょう <i>ikimasyou</i>	
PRÄSUMPTIV	NEUTRAL	行くだろう <i>iku darou</i>	行かないだろう <i>ikanai darou</i>
	FORMELL	行くでしょう <i>iku desyou</i>	行かないでしょう <i>ikanai desyou</i>
KONDITIONAL	NEUTRAL	行けば <i>ikeba</i>	行かなければ <i>ikanakereba</i>
		行ったら <i>ittara</i>	行かなかったら <i>ikanakattara</i>
	FORMELL	行きましたら <i>ikimasitara</i>	行きませんでしたら <i>ikimaseñ desitara</i>
POTENTIAL		行ける <i>ikeru</i>	
PASSIV		行かれる <i>ikareru</i>	
KAUSATIV		行かせる <i>ikaseru</i>	
KAUSATIV-PASSIV		行かせられる <i>ikaserareru</i>	
		行かされる <i>ikasareru</i>	

泳ぐ – *oyogu* – schwimmen v_{cg}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	泳ぐ <i>oyogu</i>	泳がない <i>oyoganai</i>
	FORMELL	泳ぎます <i>oyogimasu</i>	泳ぎません <i>oyogimaseñ</i>
PERFEKT	NEUTRAL	泳いだ <i>oyoida</i>	泳がなかった <i>oyoganakatta</i>
	FORMELL	泳ぎました <i>oyogimasita</i>	泳ぎませんでした <i>oyogimaseñ desita</i>
PARTIZIP		泳いで <i>oyoide</i>	泳がなくて <i>oyoganakute</i>
IMPERATIV		泳げ <i>oyoge</i>	泳ぐな <i>oyogu na</i>
VOLITIONAL	NEUTRAL	泳ごう <i>oyogou</i>	
	FORMELL	泳ぎましょう <i>oyogimasyou</i>	
PRÄSUMPTIV	NEUTRAL	泳ぐだろう <i>oyogu darou</i>	泳がないだろう <i>oyoganai darou</i>
	FORMELL	泳ぐでしょう <i>oyogu desyou</i>	泳がないでしょう <i>oyoganai desyou</i>
KONDITIONAL	NEUTRAL	泳げば <i>oyogeba</i>	泳がなければ <i>oyoganakereba</i>
		泳いだら <i>oyoidara</i>	泳がなかったら <i>oyoganakattara</i>
	FORMELL	泳ぎましたら <i>oyogimasitara</i>	泳ぎませんでしたら <i>oyogimaseñ desitara</i>
POTENTIAL		泳げる <i>oyogeru</i>	
PASSIV		泳がれる <i>oyogareru</i>	
KAUSATIV		泳がせる <i>oyogaseru</i>	
KAUSATIV-PASSIV		泳がせられる <i>oyogaserareru</i>	
		泳がされる <i>oyogasareru</i>	

話す – *hanasu* – sprechen v_{cs}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	話す <i>hanasu</i>	話さない <i>hanasanai</i>
	FORMELL	話します <i>hanasimasu</i>	話しません <i>hanasimaseñ</i>
PERFEKT	NEUTRAL	話した <i>hanasita</i>	話さなかった <i>hanasanakatta</i>
	FORMELL	話しました <i>hanasimasita</i>	話しませんでした <i>hanasimaseñ desita</i>
PARTIZIP		話して <i>hanasite</i>	話さなくて <i>hanasanakute</i>
IMPERATIV		話せ <i>hanase</i>	話すな <i>hanasu na</i>
VOLITIONAL	NEUTRAL	話そう <i>hanasou</i>	
	FORMELL	話しましょう <i>hanasimasyou</i>	
PRÄSUMPTIV	NEUTRAL	話すだろう <i>hanasu darou</i>	話さないだろう <i>hanasanai darou</i>
	FORMELL	話すでしょう <i>hanasu desyou</i>	話さないでしょう <i>hanasanai desyou</i>
KONDITIONAL	NEUTRAL	話せば <i>hanaseba</i>	話さなければ <i>hanasanakereba</i>
		話したら <i>hanasitara</i>	話さなかったら <i>hanasanakattara</i>
	FORMELL	話しましたら <i>hanasimasitara</i>	話しませんでしたら <i>hanasimaseñ desitara</i>
POTENTIAL		話せる <i>hanaseru</i>	
PASSIV		話される <i>hanasareru</i>	
KAUSATIV		話させる <i>hanasaseru</i>	
KAUSATIV-PASSIV		話させられる <i>hanasaserareru</i>	

待つ – *matu* – warten v_{ct}

		AFFIRMATIV	NEGIIERT
PRÄSENS	NEUTRAL	待つ <i>matu</i>	待たない <i>matanai</i>
	FORMELL	待ちます <i>matimasu</i>	待ちません <i>matimaseñ</i>
PERFEKT	NEUTRAL	待った <i>matta</i>	待たなかった <i>matanakatta</i>
	FORMELL	待ちました <i>matimasita</i>	待ちませんでした <i>matimaseñ desita</i>
PARTIZIP		待って <i>matte</i>	待たなくて <i>matanakute</i>
IMPERATIV		待て <i>mate</i>	待つな <i>matu na</i>
VOLITIONAL	NEUTRAL	待とう <i>matou</i>	
	FORMELL	待ちましょう <i>matimasyou</i>	
PRÄSUMPTIV	NEUTRAL	待つだろう <i>matu darou</i>	待たないだろう <i>matanai darou</i>
	FORMELL	待つでしょう <i>matu desyou</i>	待たないでしょう <i>matanai desyou</i>
KONDITIONAL	NEUTRAL	待てば <i>mateba</i>	待たなければ <i>matanakereba</i>
		待ったら <i>mattara</i>	待たなかったら <i>matanakattara</i>
	FORMELL	待ちましたら <i>matimasitara</i>	待ちませんでしたら <i>matimaseñ desitara</i>
POTENTIAL		待てる <i>materu</i>	
PASSIV		待たれる <i>matareru</i>	
KAUSATIV		待たせる <i>mataseru</i>	
KAUSATIV-PASSIV		待たせられる <i>mataserareru</i>	
		待たされる <i>matasareru</i>	

死ぬ – *sinu* – sterben v_{cn}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	死ぬ <i>sinu</i>	死なない <i>sinanai</i>
	FORMELL	死にます <i>sinimasu</i>	死にません <i>sinimaseñ</i>
PERFEKT	NEUTRAL	死んだ <i>siñda</i>	死ななかった <i>sinanakatta</i>
	FORMELL	死にました <i>sinimasita</i>	死にませんでした <i>sinimaseñ desita</i>
PARTIZIP		死んで <i>siñde</i>	死ななくて <i>sinanakute</i>
IMPERATIV		死ね <i>sine</i>	死ぬな <i>sinu na</i>
VOLITIONAL	NEUTRAL	死のう <i>sinou</i>	
	FORMELL	死にましよう <i>sinimasyou</i>	
PRÄSUMPTIV	NEUTRAL	死ぬだろう <i>sinu darou</i>	死なないだろう <i>sinanai darou</i>
	FORMELL	死ぬでしょう <i>sinu desyou</i>	死なないでしょう <i>sinanai desyou</i>
KONDITIONAL	NEUTRAL	死ねば <i>sineba</i>	死ななければ <i>sinanakereba</i>
		死んだら <i>siñdara</i>	死ななかったら <i>sinanakattara</i>
	FORMELL	死にましたら <i>sinimasitara</i>	死にませんでしたら <i>sinimaseñ desitara</i>
POTENTIAL		死ねる <i>sineru</i>	
PASSIV		死なれる <i>sinareru</i>	
KAUSATIV		死なせる <i>sinaseru</i>	
KAUSATIV-PASSIV		死なせられる <i>sinaserareru</i>	
		死なされる <i>sinasareru</i>	

呼ぶ – *yobu* – rufen

v_{cb}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	呼ぶ <i>yobu</i>	呼ばない <i>yobanai</i>
	FORMELL	呼びます <i>yobimasu</i>	呼びません <i>yobimaseñ</i>
PERFEKT	NEUTRAL	呼んだ <i>yoñda</i>	呼ばなかった <i>yobanakatta</i>
	FORMELL	呼びました <i>yobimasita</i>	呼びませんでした <i>yobimaseñ desita</i>
PARTIZIP		呼んで <i>yoñde</i>	呼ばなくて <i>yobanakute</i>
IMPERATIV		呼べ <i>yobe</i>	呼ぶな <i>yobu na</i>
VOLITIONAL	NEUTRAL	呼ぼう <i>yobou</i>	
	FORMELL	呼びましょう <i>yobimasyou</i>	
PRÄSUMPTIV	NEUTRAL	呼ぶだろう <i>yobu darou</i>	呼ばないだろう <i>yobanai darou</i>
	FORMELL	呼ぶでしょう <i>yobu desyou</i>	呼ばないでしょう <i>yobanai desyou</i>
KONDITIONAL	NEUTRAL	呼べば <i>yobeba</i>	呼ばなければ <i>yobanakereba</i>
		呼んだら <i>yoñdara</i>	呼ばなかったら <i>yobanakattara</i>
	FORMELL	呼びましたら <i>yobimasitara</i>	呼びませんでしたら <i>yobimaseñ desitara</i>
POTENTIAL		呼べる <i>yoberu</i>	
PASSIV		呼ばれる <i>yobareru</i>	
KAUSATIV		呼ばせる <i>yobaseru</i>	
KAUSATIV-PASSIV		呼ばせられる <i>yobaserareru</i>	
		呼ばされる <i>yobasareru</i>	

飲む – *nomu* – trinken v_{cm}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	飲む <i>nomu</i>	飲まない <i>nomanai</i>
	FORMELL	飲みます <i>nomimasu</i>	飲みません <i>nomimaseñ</i>
PERFEKT	NEUTRAL	飲んだ <i>noñda</i>	飲まなかった <i>nomanakatta</i>
	FORMELL	飲みました <i>nomimasita</i>	飲みませんでした <i>nomimaseñ desita</i>
PARTIZIP		飲んで <i>noñde</i>	飲まなくて <i>nomanakute</i>
IMPERATIV		飲め <i>nome</i>	飲むな <i>nomu na</i>
VOLITIONAL	NEUTRAL	飲もう <i>nomou</i>	
	FORMELL	飲みましょう <i>nomimasyou</i>	
PRÄSUMPTIV	NEUTRAL	飲むだろう <i>nomu darou</i>	飲まないだろう <i>nomanai darou</i>
	FORMELL	飲むでしょう <i>nomu desyou</i>	飲まないでしょう <i>nomanai desyou</i>
KONDITIONAL	NEUTRAL	飲めば <i>nomeba</i>	飲まなければ <i>nomanakereba</i>
		飲んだら <i>noñdara</i>	飲まなかったら <i>nomanakattara</i>
	FORMELL	飲みましたら <i>nomimasitara</i>	飲みませんでしたら <i>nomimaseñ desitara</i>
POTENTIAL		飲める <i>nomeru</i>	
PASSIV		飲まれる <i>nomareru</i>	
KAUSATIV		飲ませる <i>nomaseru</i>	
KAUSATIV-PASSIV		飲ませられる <i>nomaserareru</i>	
		飲まされる <i>nomasareru</i>	

乗る – *noru* – fahren, einsteigen v_{cr}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	乗る <i>noru</i>	乗らない <i>noranai</i>
	FORMELL	乗ります <i>norimasu</i>	乗りません <i>norimaseñ</i>
PERFEKT	NEUTRAL	乗った <i>notta</i>	乗らなかった <i>noranakatta</i>
	FORMELL	乗りました <i>norimasita</i>	乗りませんでした <i>norimaseñ desita</i>
PARTIZIP		乗って <i>notte</i>	乗らなくて <i>noranakute</i>
IMPERATIV		乗れ <i>nore</i>	乗るな <i>noru na</i>
VOLITIONAL	NEUTRAL	乗ろう <i>norou</i>	
	FORMELL	乗りましょう <i>norimasyou</i>	
PRÄSUMPTIV	NEUTRAL	乗るだろう <i>noru darou</i>	乗らないだろう <i>noranai darou</i>
	FORMELL	乗るでしょう <i>noru desyou</i>	乗らないでしょう <i>noranai desyou</i>
KONDITIONAL	NEUTRAL	乗れば <i>noreba</i>	乗らなければ <i>noranakereba</i>
		乗ったら <i>nottara</i>	乗らなかつたら <i>noranakattara</i>
	FORMELL	乗りましたら <i>norimasitara</i>	乗りませんでしたら <i>norimaseñ desitara</i>
POTENTIAL		乗れる <i>noreru</i>	
PASSIV		乗られる <i>norareru</i>	
KAUSATIV		乗らせる <i>noraseru</i>	
KAUSATIV-PASSIV		乗らせられる <i>noraserareru</i>	
		乗らされる <i>norasareru</i>	

取る – *toru* – nehmen v_{cr}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	取る <i>toru</i>	取らない <i>toranai</i>
	FORMELL	取ります <i>torimasu</i>	取りません <i>torimaseñ</i>
PERFEKT	NEUTRAL	取った <i>totta</i>	取らなかった <i>toranakatta</i>
	FORMELL	取りました <i>torimasita</i>	取りませんでした <i>torimaseñ desita</i>
PARTIZIP		取って <i>totte</i>	取らなくて <i>toranakute</i>
IMPERATIV		取れ <i>tore</i>	取るな <i>toru na</i>
VOLITIONAL	NEUTRAL	取ろう <i>torou</i>	
	FORMELL	取りましょう <i>torimasyou</i>	
PRÄSUMPTIV	NEUTRAL	取るだろう <i>toru darou</i>	取らないだろう <i>toranai darou</i>
	FORMELL	取るでしょう <i>toru desyou</i>	取らないでしょう <i>toranai desyou</i>
KONDITIONAL	NEUTRAL	取れば <i>toreba</i>	取らなければ <i>toranakereba</i>
		取ったら <i>tottara</i>	取らなかったら <i>toranakattara</i>
	FORMELL	取りましたら <i>torimasitara</i>	取りませんでしたら <i>torimaseñ desitara</i>
POTENTIAL		取れる <i>toreru</i>	
PASSIV		取られる <i>torareru</i>	
KAUSATIV		取らせる <i>toraseru</i>	
KAUSATIV-PASSIV		取らせられる <i>toraserareru</i>	
		取らされる <i>torasareru</i>	

ある - *aru* - sein v_{caru}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	ある <i>aru</i>	ない <i>nai</i>
	FORMELL	あります <i>arimasu</i>	ありません <i>arimaseñ</i>
PERFEKT	NEUTRAL	あった <i>atta</i>	なかった <i>nakatta</i>
	FORMELL	ありました <i>arimasita</i>	ありませんでした <i>arimaseñ desita</i>
PARTIZIP		あって <i>atte</i>	なくて <i>nakute</i>
IMPERATIV			(あるな) <i>aru na</i>
VOLITIONAL	NEUTRAL	(あろう) <i>arou</i>	
	FORMELL	(ありましょう) <i>arimasyou</i>	
PRÄSUMPTIV	NEUTRAL	あるだろう <i>aru darou</i>	ないだろう <i>nai darou</i>
	FORMELL	あるでしょう <i>aru desyou</i>	ないでしょう <i>nai desyou</i>
KONDITIONAL	NEUTRAL	あれば <i>areba</i>	なければ <i>nakereba</i>
		あったら <i>attara</i>	なかったら <i>nakattara</i>
	FORMELL	ありましたら <i>arimasitara</i>	ありませんでしたら <i>arimaseñ desitara</i>
POTENTIAL		(あれる) <i>areru</i>	
PASSIV		(あられる) <i>arareru</i>	
KAUSATIV		(あらせる) <i>araseru</i>	
KAUSATIV-PASSIV		(あらせられる) <i>araserareru</i>	
		[あさされる] <i>arasareru</i>	

だ – da – sein (<Partikelverb>)⁴⁴²

v_{da}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	だ <i>da</i>	
	FORMELL	です <i>desu</i>	
PERFEKT	NEUTRAL	だった <i>datta</i>	
	FORMELL	でした <i>desita</i>	
PARTIZIP		で <i>de</i>	
IMPERATIV			
VOLITIONAL	NEUTRAL	だろう <i>darou</i>	
	FORMELL	でしょう <i>desyou</i>	
PRÄSUMPTIV	NEUTRAL		
	FORMELL		
KONDITIONAL	NEUTRAL		
		だったら <i>dattara</i>	
	FORMELL	でしたら <i>desitara</i>	
POTENTIAL			
PASSIV			
KAUSATIV			
KAUSATIV-PASSIV			

⁴⁴² Das Verb だ (*da*; sein) wird im CJVG nicht aufgelistet.

なさる – *nasaru* – tun, machen v_{nasaru}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	なさる <i>nasaru</i>	なさらない <i>nasaranai</i>
	FORMELL	なさいます <i>nasaimasu</i>	なさいません <i>nasaimaseñ</i>
PERFEKT	NEUTRAL	なさった <i>nasatta</i>	なさらなかった <i>nasaranakatta</i>
	FORMELL	なさいました <i>nasaimasita</i>	なさいませんでした <i>nasaimaseñ desita</i>
PARTIZIP		なさって <i>nasatte</i>	なさらず <i>nasaranakute</i>
IMPERATIV		なされ <i>nasare</i>	なさるな <i>nasaru na</i>
VOLITIONAL	NEUTRAL	なさろう <i>nasarou</i>	
	FORMELL	なさいましょう <i>nasaimasyou</i>	
PRÄSUMPTIV	NEUTRAL	なさるだろう <i>nasaru darou</i>	なさらないだろう <i>nasaranai darou</i>
	FORMELL	なさるでしょう <i>nasaru desyou</i>	なさらないでしょう <i>nasaranai desyou</i>
KONDITIONAL	NEUTRAL	なされば <i>nasareba</i>	なさなければ <i>nasaranakereba</i>
		なさったら <i>nasattara</i>	なさなかったら <i>nasaranakattara</i>
	FORMELL	なさいましたら <i>nasaimasitara</i>	なさいませんでしたら <i>nasaimaseñ desitara</i>
POTENTIAL		なされる <i>nasareru</i>	
PASSIV		なさられる <i>nasarareru</i>	
KAUSATIV		なさらせる <i>nasaraseru</i>	
KAUSATIV-PASSIV		なさらせられる <i>nasaraserareru</i>	
		[なさられる] <i>[nasarasareru]</i>	

買う – *kau* – kaufen

v_{cw}

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	買う <i>kau</i>	買わない <i>kawanai</i>
	FORMELL	買います <i>kaimasu</i>	買いません <i>kaimaseñ</i>
PERFEKT	NEUTRAL	買った <i>katta</i>	買わなかった <i>kawanakatta</i>
	FORMELL	買いました <i>kaimasita</i>	買いませんでした <i>kaimaseñ desita</i>
PARTIZIP		買って <i>katte</i>	買わなくて <i>kawanakute</i>
IMPERATIV		買い <i>kae</i>	買うな <i>kau na</i>
VOLITIONAL	NEUTRAL	買おう <i>kaou</i>	
	FORMELL	買いましょう <i>kaimasyou</i>	
PRÄSUMPTIV	NEUTRAL	買うだろう <i>kau darou</i>	買わないだろう <i>kawanai darou</i>
	FORMELL	買うでしょう <i>kau desyou</i>	買わないでしょう <i>kawanai desyou</i>
KONDITIONAL	NEUTRAL	買えば <i>kaeba</i>	買わなければ <i>kawanakereba</i>
		買ったら <i>kattara</i>	買わなかったら <i>kawanakattara</i>
	FORMELL	買いましたら <i>kaimasitara</i>	買いませんでしたら <i>kaimaseñ desitara</i>
POTENTIAL		買える <i>kaeru</i>	
PASSIV		買われる <i>kawareru</i>	
KAUSATIV		買わせる <i>kawaseru</i>	
KAUSATIV-PASSIV		買わせられる <i>kawaserareru</i>	
		買わされる <i>kawasareru</i>	

来る – *kuru* – kommen $v_{i_{kuru}}$

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	来る <i>kuru</i>	来ない <i>konai</i>
	FORMELL	来ます <i>kimasu</i>	来ません <i>kimaseñ</i>
PERFEKT	NEUTRAL	来た <i>kita</i>	来なかった <i>konakatta</i>
	FORMELL	来ました <i>kimasita</i>	来ませんでした <i>kimaseñ desita</i>
PARTIZIP		来て <i>kite</i>	来なくて <i>konakute</i>
IMPERATIV		来い <i>koi</i>	来るな <i>kuru na</i>
VOLITIONAL	NEUTRAL	来よう <i>koyou</i>	
	FORMELL	来ましょう <i>kimasyou</i>	
PRÄSUMPTIV	NEUTRAL	来るだろう <i>kuru darou</i>	来ないだろう <i>konai darou</i>
	FORMELL	来るでしょう <i>kuru desyou</i>	来ないでしょう <i>konai desyou</i>
KONDITIONAL	NEUTRAL	来れば <i>kureba</i>	来なければ <i>konakereba</i>
		来たら <i>kitara</i>	来なかったら <i>konakattara</i>
	FORMELL	来ましたら <i>kimasitara</i>	来ませんでしたら <i>kimaseñ desitara</i>
POTENTIAL		来れる <i>koreru</i>	
		来られる <i>korareru</i>	
PASSIV		来られる <i>korareru</i>	
KAUSATIV		来させる <i>kosaseru</i>	
KAUSATIV-PASSIV		来させられる <i>kosaserareru</i>	
		来さされる <i>kosasareru</i>	

する – *suru* – tun, machen

$v_{i\text{suru}}$

		AFFIRMATIV	NEGIERT
PRÄSENS	NEUTRAL	する <i>suru</i>	しない <i>sinai</i>
	FORMELL	します <i>simasu</i>	しません <i>simaseñ</i>
PERFEKT	NEUTRAL	した <i>sita</i>	しなかった <i>sinakatta</i>
	FORMELL	しました <i>simasita</i>	しませんでした <i>simaseñ desita</i>
PARTIZIP		して <i>site</i>	しなくて <i>sinakute</i>
IMPERATIV		しろ <i>siro</i>	するな <i>suru na</i>
VOLITIONAL	NEUTRAL	しよう <i>siyou</i>	
	FORMELL	しましょう <i>simasyou</i>	
PRÄSUMPTIV	NEUTRAL	するだろう <i>suru darou</i>	しないだろう <i>sinai darou</i>
	FORMELL	するでしょう <i>suru desyou</i>	しないでしょう <i>sinai desyou</i>
KONDITIONAL	NEUTRAL	すれば <i>sureba</i>	しなければ <i>sinakereba</i>
		したら <i>sitara</i>	しなかったら <i>sinakattara</i>
	FORMELL	しましたら <i>simasitara</i>	しませんでしたら <i>simaseñ desitara</i>
POTENTIAL		できる <i>dekiru</i>	
PASSIV		される <i>sareru</i>	
KAUSATIV		させる <i>saseru</i>	
KAUSATIV-PASSIV		させられる <i>saserareru</i>	

Bibliographie

Abkürzungen

- CJVG *The Complete Japanese Verb Guide*, Ishii et al. 1989.
DBJG *A Dictionary of Basic Japanese Grammar*,
Makino und Tsutsui 1986.
IJDW *IKUBUNDO Japanisch-Deutsches Wörterbuch*,
Tomiyama et al. 1996.
JM *Japanische Morphosyntax*, Rickmeyer 1995.
JPSG *Japanese Phrase Structure Grammar*, Gunji 1987.
LIJC *The Lexical Integrity of Japanese Causatives*,
Manning et al. 1999.

Literaturverzeichnis

- Anne Abeillé und Danièle Godard 1994: The Syntax of French Auxiliaries. Unpublished manuscript Université Paris 7 and CNRS.
- Peter Aczel, David Israel, Yosuihiro Katagiri und Stanley Peters (Herausgeber) 1993: *Situation Theory and Its Applications*, Band 3. CSLI Publications, Stanford.
- Rolf Backofen und Christoph Weyers 1994: *UDiNe*. A Feature Constraint Solver with Distributed Disjunction and Classical Negation. Technischer Bericht, Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI), Saarbrücken, Germany. Teil der Dokumentation des PAGE-Systems.
- Jon Barwise, Jean Mark Gawron, Gordon Plotkin und Syun Tutiya (Herausgeber) 1991: *Situation Theory and Its Applications*, Band 2. CSLI Publications, Stanford.
- Jon Barwise und John Perry 1983: *Situations and Attitudes*. Bradford Books, MIT Press, Cambridge, Massachusetts.
- H. U. Block und S. Schachtl 1992: Trace & Unification Grammar. In *Proc. of the 14th COLING*, Seiten 87–93.
- Dietrich Bollmann 1996: Das japanische Schriftsystem und seine Transkription. Unveröffentlichtes Manuskript.

- Dietrich Bollmann 1997: Beschreibung der japanischen Verbmorphologie durch Vererbungshierarchien und Defaults. Studienarbeit, Technische Universität Berlin.
- Johan Bos, Elsbeth Mastenbroek, Scott McGlashan, Sebastian Millies und Manfred Pinkal 1994: The Verbmobil Semantic Formalism (Version 1.3). Technischer Bericht, Universität des Saarlandes.
- Gosse Bouma und Gertjan van Noord 1994: The Scope of Adjuncts and the Processing of Lexical Rules. In *Proceedings of Coling, 1994, Kyoto Japan*.
- Bob Carpenter 1992: *The Logic of Typed Feature Structures with Applications to Unification-based Grammars, Logic Programming and Constraint Resolution*, Band 32, *Cambridge Tracts in Theoretical Computer Science*. Cambridge University Press, New York.
- Bob Carpenter 1993: Skeptical and Credulous Default Unification with Applications to Templates and Inheritance. In T. Briscoe, A. Copestake und V. de Paiva (Herausgeber), *Default Reasoning and Lexical Organization*, Cambridge. Cambridge University Press.
- Bob Carpenter und Gerald Penn 1999: ALE, The Attribute Logic Engine, User's Guide. Technischer Bericht. Version 3.2 Beta.
- Rudolf Caspari und Ludwig Schmid 1994: Parsing und Generierung in TrUG. Verbmobil-Report 40, Siemens AG.
- Noam Chomsky 1957: *Syntactic Structures*. Mouton & Co, La Haye.
- Noam Chomsky 1981: *Lectures on Government and Binding*. Studies in Generative Grammar 9. Foris, Dordrecht.
- Noam Chomsky 1995: *The Minimalist Program*. Current Studies in Linguistics. MIT Press, Cambridge, Massachusetts.
- Noam Chomsky und Morris Halle 1968: *The sound pattern of English*. Harper and Row, New York.
- Robin Cooper, Kuniaki Mukai und John Perry (Herausgeber) 1990: *Situation Theory and Its Applications*, Band 1. CSLI Publications, Stanford.
- Ann Copestake, Dan Flickinger und Ivan A. Sag 1997: Minimal Recursion Semantics: An Introduction. Unpublished.
- Ferdinand de Saussure 1916: *Cours de linguistique générale*. Lausanne, Paris. postum hrsg. von Charles Bally und Albert Sechehaye.

- Martin Christoph Emele 1991: Unification with Lazy Non-Redundant Copying. In *Proceedings of the 29th Annual Meeting of the Association for Computational Linguistics*, Seiten 323–330, Berkeley. Association for Computational Linguistics.
- Martin Christoph Emele 1993: TFS – The Typed Feature Structure Representation Formalism. In H. Uszkoreit (Herausgeber), *Proceedings of the EAGLES workshop on implemented formalisms, DFKI report 1993*. (vgl. auch Emele 1994).
- Martin Christoph Emele 1994: The Typed Feature Structure Representation Formalism. In *Proceedings of the International Workshop on Sharable Natural Language Resources*, Ikoma, Nara, Japan. (vgl. auch Emele 1993).
- Martin Christoph Emele und Rémi Zajac 1990a: A Fix-Point Semantics for Feature Type Systems. In *Proceedings of the 2nd International Workshop on Conditional and Typed Rewriting Systems, CTRS'90, Montréal, June 1990*.
- Martin Christoph Emele und Rémi Zajac 1990b: Typed Unification Grammars. In *Proceedings of the 13th International Workshop Conference on Computational Linguistics, COLING-90, Helsinki, August 1990*, Helsinki.
- Ulrich Engel 1977: *Syntax der deutschen Gegenwartssprache*.
- Ulrich Engel 1994: *Syntax der deutschen Gegenwartssprache*. Erich Schmidt Verlag, Berlin. 3., völlig neu bearb. Aufl. (vgl. Engel 1977).
- Wolfgang Finkler und Günter Neumann 1986: MORPHIX: Ein hochportabler Lemmatisierungsmodul für das Deutsche. Technical Report Memo Nr. 8, German Research Center for AI, DFKI GmbH, Saarbrücken, Germany. DFKI report R:S90-092.
- Wolfgang Finkler und Günter Neumann 1988: MORPHIX: A Fast Realization of a Classification-Based Approach to Morphology. In Harald Trost (Herausgeber), *4. Österreichische Artificial-Intelligence-Tagung. Wiener Workshop — Wissensbasierte Sprachverarbeitung. Proceedings*, Seiten 11–19, Berlin. Springer. DFKI report R:S90-045.
- Friedrich Ludwig Gottlob Frege 1879: *Begriffsschrift. Eine der arithmetischen nachgebildete Formelsprache des reinen Denkens*. Halle.
- Takao Gunji 1987: *Japanese Phrase Structure Grammar*. Reidel, Dordrecht.

- Takao Gunji 1991: An Overview of JPSG: A Constraint-Based Descriptive Theory for Japanese. In R. Mazuka und N. Nagai (Herausgeber), *Proceedings of Japanese Syntactic Processing Workshop*, Duke University. Lawrence Erlbaum.
- Takao Gunji 1996a: On Lexicalist Treatments of Japanese Causatives. In Takao Gunji (Herausgeber), *Studies on the Universality of Constraint-Based Phrase Structure Grammars*, Seiten 61–89. Osaka University.
- Takao Gunji (Herausgeber) 1996b: *Studies on the Universality of Constraint-Based Phrase Structure Grammars*. Osaka University.
- Takao Gunji 1999: On Lexicalist Treatments of Japanese Causatives. In Robert Levine und Georgia Green (Herausgeber), *Studies in Contemporary Phrase Structure Grammar*. Cambridge University Press, Cambridge and New York. (vgl. Gunji 1996a).
- Gerhard Helbig und Joachim Buscha 1998: *Deutsche Grammatik. Ein Handbuch für den Ausländerunterricht*. Verlag Enzyklopädie / Langenscheidt, Leipzig. 18. Auflage.
- Masayo Iida, Christopher D. Manning, Patrick O'Neill und Ivan Sag 1994: The Lexical Integrity of Japanese Causatives. Presented at the LSA 1994 Annual Meeting, Boston. Unpublished ms., Stanford University.
- K. Ishii, M. Kono, Y. Miyazawa, N. Nagano, H. Watanabe und N. Watanabe 1989: *Hiroo Japanese Center: The Complete Japanese Verb Guide*. Charles E. Tuttle, Tokyo.
- Bernd Kiefer und Thomas Fettig 1995: FEGRAMED—An Interactive Graphics Editor for Feature Structures. Research Report RR-95-06, German Research Center for Artificial Intelligence (DFKI), Saarbrücken, Germany.
- Bernd Kiefer und Oliver Scherf 1994: Gimme More HQ Parsers. The Generic Parser Class of DISCO. Technischer Bericht, Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI), Saarbrücken, Germany. Teil der Dokumentation des PAGE-Systems.
- Jong-Bok Kim und Ivan A. Sag 1995: The Parametric Variation of English and French Negation. In *Proceedings of the Fourteenth West Coast Conference on Formal Linguistics*, Band 14, Stanford University. CSLI Publications/SLA.
- Kimmo Koskenniemi 1983: Two-level Morphology: A General Computational Model for Word-Form Recognition and Production. Publications 11, University of Helsinki, Department of General Linguistics, Helsinki.

- Kimmo Koskeniemi 1985a: A System for Generating Finnish Inflected Word-Forms. Publications 13, University of Helsinki, Department of General Linguistics, Helsinki. Computational Morphosyntax: Report on Research 1981-84.
- Kimmo Koskeniemi 1985b: An Application of the Two-level Model to Finnish. Publications 13, University of Helsinki, Department of General Linguistics, Helsinki. Computational Morphosyntax: Report on Research 1981-84.
- Kimmo Koskeniemi und Kenneth W. Church 1988: Two-level Morphology and Finnish. In Denes Vargha (Herausgeber), *Proceedings of the 12th International Conference on Computational Linguistics, COLING-88, Budapest*, Seiten 335–339.
- Hans-Ulrich Krieger und Ulrich Schäfer 1994a: *TDL*—A Type Description Language for Constraint-Based Grammars. In *Proceedings of the 15th International Conference on Computational Linguistics, COLING-94, Kyoto, Japan*.
- Hans-Ulrich Krieger und Ulrich Schäfer 1994b: *TDL*—A Type Description Language for HPSG. Part 1: Overview. Research Report RR-94-37, Deutsches Forschungszentrum für Künstliche Intelligenz, Saarbrücken, Germany.
- Hans-Ulrich Krieger und Ulrich Schäfer 1994c: *TDL*—A Type Description Language for HPSG. Part 2: User Manual. DFKI Document D-94-14, Deutsches Forschungszentrum für Künstliche Intelligenz, Saarbrücken, Germany.
- Alex Lascarides und Ann Copestake 1999: Default Representation in Constraint-based Frameworks. In *Computational Linguistics, 25.1*, Seiten 55–105. MIT Press.
- Bruno Lewin 1990: *Abriss der japanischen Grammatik*. Harrassowitz, Wiesbaden. 3., verbesserte Auflage.
- John Wylie Lloyd 1984: *Foundations of Logic Programming*. Springer-Verlag, Berlin.
- Seiichi Makino und Michio Tsutsui 1986: *A Dictionary of Basic Japanese Grammar*. The Japan Times, Tōkyō.
- Christopher D. Manning, Ivan A. Sag und Masayo Iida 1996: The Lexical Integrity of Japanese Causatives. In Takao Gunji (Herausgeber), *Studies on the Universality of Constraint-Based Phrase Structure Grammars*, Seiten 9–37. Osaka University.

- Christopher D. Manning, Ivan A. Sag und Masayo Iida 1999: The Lexical Integrity of Japanese Causatives. In Robert Levine und Georgia Green (Herausgeber), *Studies in Contemporary Phrase Structure Grammar*. Cambridge University Press. (vgl. Iida et al. 1994).
- Christopher D. Manning und Hinrich Schütze 1999: *Foundations of Statistical Natural Language Processing*. MIT Press, Cambridge, Massachusetts.
- Yōkō Matsuoka McClain 1981: *Handbook of Modern Japanese Grammar*. Hokuseidō Press, Tōkyō.
- Akira Mikami 1960: *Zō wa hana ga nagai*. Kuroshio Shuppan, Tōkyō.
- Yutaka Mitsuishi, Kentaro Torisawa und Jun-ichi Tsujii 1998: HPSG-Style Underspecified Japanese Grammar with Wide Coverage. In *Proceedings of the 36th Annual Meeting of the Association for Computational Linguistics (COLING-ACL '98)*. Association for Computational Linguistics.
- Manfred Pinkal 1995: Radical Underspecification. In Paul Dekker und Martin Stokhof (Herausgeber), *Proceedings of the 10th Amsterdam Colloquium*. University of Amsterdam.
- Manfred Pinkal 1999: On Underspecification. In *5th International Workshop on Computational Semantics*, Tilburg. ITK, Tilburg University.
- Carl Pollard und Ivan A. Sag 1987: *Information-based Syntax and Semantics, Vol. 1*. Number 13 in Lecture Notes. CSLI Publications, Stanford University. University of Chicago Press.
- Carl Pollard und Ivan A. Sag 1994: *Head-Driven Phrase Structure Grammar*. University of Chicago Press, Chicago.
- Jens Rickmeyer 1977: *Kleines japanisches Valenzlexikon*. Helmut Buske Verlag, Hamburg.
- Jens Rickmeyer 1983: *Morphosyntax der japanischen Gegenwartssprache*, Band 2, *Deutsch und Japanisch im Kontrast*. Groos, Heidelberg.
- Jens Rickmeyer 1984: Prolegomena zur "Morphosyntax der japanischen Gegenwartssprache". In Hans Link, Rudolf Herzer und Werner Sasse (Herausgeber), *Bochumer Jahrbuch zur Ostasienforschung*, Band 7, Seiten 57–112, Bochum. Studienverlag Dr. Norbert Brockmeyer.
- Jens Rickmeyer 1994: Ein taxonomisch-dependentielles Modell für eine deskriptive Grammatik der japanischen Sprache. In *Bochumer Jahrbuch zur Ostasienforschung*, Band 18, Seiten 249–58, München. Iudicium Verlag.

- Jens Rickmeyer 1995: *Japanische Morphosyntax*. Groos, Heidelberg. (Früher u.d.T.: Morphosyntax der japanischen Gegenwartssprache, Rickmeyer 1983).
- Melanie Siegel 1995: Problems of Automatic Translation of Japanese Dialogues into German. In Walther V.Hahn, Susanne Jekat und Ilona Maleck (Herausgeber), *Machine Translation and Machine Interpretation, Proceedings of the VERBMOBIL Workshop at the University of Hamburg, Computer Science Department*. University of Hamburg, Computer Science Department.
- Melanie Siegel 1996a: Definiteness and Number in Japanese to German Machine Translation. In Dafydd Gibbon (Herausgeber), *Natural Language Processing and Speech Technology*, Berlin. Mouton de Gruyter.
- Melanie Siegel 1996b: Die japanische Syntax im Verbmobil-Forschungsprototypen. Verbmobil-Report 133, German Research Center for AI, DFKI GmbH, Saarbrücken, Germany.
- Melanie Siegel 1997: *Die maschinelle Übersetzung aufgabenorientierter japanisch-deutscher Dialoge. Lösungen für Translation Mismatches*. Logos Verlag, Berlin. Dissertation (1996).
- Melanie Siegel 1998: Japanese Particles in an HPSG Grammar. Verbmobil-Report 200, Universität des Saarlandes.
- Melanie Siegel 1999: The Syntactic Processing of Particles in Japanese Spoken Language. In Jhing-Fa Wang und Chung-Hsien Wu (Herausgeber), *Proceedings of the 13th Pacific Asia Conference on Language, Information and Computation, Taipei 1999*.
- Leon Sterling und Ehud Y. Shapiro 1986: *The Art of Prolog: Advanced Programming Techniques*. MIT Press, Cambridge, Massachusetts.
- Lucien Tesnière 1959: *Éléments de Syntaxe structurale*. Paris.
- Yoshimasa Tomiyama, Yukio Miura und Kazuo Yamaguchi (Herausgeber) 1996: *IKUBUNDO Japanisch-Deutsches Wörterbuch*. Ikubundo, Tokyo. 3., völlig neu bearbeitete und erweiterte Auflage.
- Kentaro Torisawa, Takaki Makino, Takashi Ninomiya, Yutaka Mitsuishi, Yuka Tateishi, Kenji Nishida, Yusuke Miyao, Minoru Yoshida und Jun-ichi Tsujii 1998: Practical HPSG Parsers. In *JSPS-HITACHI Workshop on New Challenges in Natural Language Processing and its Application*, Tokyo, Japan.

- Jun-ichi Tsujii 1998: The JSPS Project on Natural Language Understanding and Processing. In *JSPS-HITACHI Workshop on New Challenges in Natural Language Processing and its Application*, Tokyo, Japan.
- Mariko Udo 1982: Syntax and Morphology of the Japanese Verb: A Phrase Structural Approach. Unpublished M.A. Thesis, University College London.
- Hans Uszkoreit, Rolf Backofen, Stephan Busemann, Abdel Kader Diagne, Elizabeth A. Hinkelman, Walter Kasper, Bernd Kiefer, Hans-Ulrich Krieger, Klaus Netter, Günter Neumann, Stephan Oepen und Stephen P. Spackman 1994: DISCO—an HPSG-based NLP system and its application for appointment scheduling. In *Proceedings of COLING-94, Kyoto, Japan*.
- Wolfgang Wahlster 1993: Verbmobil: Translation of Face-To-Face Dialogs. In O. Herzog, T. Christaller und D. Schütt (Herausgeber), *Grundlagen und Anwendungen der künstlichen Intelligenz (17. Fachtagung)*, Seiten 393–402. Springer, Berlin, Heidelberg.
- Wolfgang Wahlster 1997: VERBMOBIL: Erkennung, Analyse, Transfer, Generierung und Synthese von Spontansprache. Technical Report Memo Nr. 198, German Research Center for AI, DFKI GmbH, Saarbrücken, Germany. DFKI report 198.
- Rémi Zajac 1992: Inheritance and Constraint-Based Grammar Formalisms. In *Computational Linguistics 18*, Seiten 159 – 180.

