

AN EXAMINATION OF QUANTIFIER SCOPE AMBIGUITY
IN TURKISH

A THESIS SUBMITTED TO
THE GRADUATE SCHOOL OF INFORMATICS
OF
THE MIDDLE EAST TECHNICAL UNIVERSITY

BY

KÜRŞAD KURT

IN PARTIAL FULFILLMENT OF THE REQUIREMENTS FOR THE DEGREE OF
MASTER OF SCIENCE
IN
THE DEPARTMENT OF COGNITIVE SCIENCE

SEPTEMBER 2006

Approval of the Graduate School of Informatics.

Assoc. Prof. Dr. Nazife Baykal
Director

I certify that this thesis satisfies all the requirements as a thesis for the degree of Master of Science.

Prof. Dr. Deniz Zeyrek
Head of Department

This is to certify that we have read this thesis and that in our opinion it is fully adequate, in scope and quality, as a thesis for the degree of Master of Science.

Assoc. Prof. Dr. H. Cem Bozşahin
Supervisor

Examining Committee Members

Prof. Dr. Deniz Zeyrek (METU, COGS) _____

Assoc. Prof. Dr. H. Cem Bozşahin (METU, CENG) _____

Assist. Prof. Dr. Ayşenur Birtürk (METU, CENG) _____

Assist. Prof. Dr. Annette Hohenberger (METU, COGS) _____

Assist. Prof. Dr. Bilge Say (METU, COGS) _____

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Name, Last Name: KÜRŞAD KURT

Signature :

ABSTRACT

AN EXAMINATION OF QUANTIFIER SCOPE AMBIGUITY IN TURKISH

Kurt, Kürşad

M.S., Department of Cognitive Science

Supervisor : Assoc. Prof. Dr. H. Cem Bozşahin

September 2006, 125 pages

This study investigates the problem of quantifier scope ambiguity in natural languages and the various ways with which it has been accounted for, some of which are problematic for monotonic theories of grammar like Combinatory Categorical Grammar (CCG) which strive for solutions that avoid non-monotonic functional application, and assume complete transparency between the syntax and the semantics interface of a language. Another purpose of this thesis is to explore these proposals on examples from Turkish and to try to account for the meaning differences that may be caused by word order and see how the observations from Turkish fit within the framework of CCG.

Keywords: Quantifier Scope Ambiguity, Combinatory Categorical Grammar, Information Structure, Turkish

ÖZ

TÜRKÇE'DEKİ NİCELEYİCİ ÇOKANLAMLILIĞI PROBLEMİNİN BİR İNCELEMESİ

Kurt, Kürşad

Yüksek Lisans, Bilişsel Bilimler Bölümü

Tez Yöneticisi : Doç. Dr. H. Cem Bozşahin

Eylül 2006, 125 sayfa

Bu çalışma doğal dillerdeki niceleyici çokanlamlılığı problemini ve bu problem ile ilgili olarak geliştirilmiş çeşitli çözümleri incelemektedir. Bu çözümlerden bazıları Ulamsal Dilbilgisi (UD) gibi sözdizimi ve anlam arasında tam saydamlık varsayan monotonik gramer teorileri ile uyumsuzdur. Bu tezin bir diğer amacı da bu çözümleri Türkçe üzerinde test etmek ve Türkçe'deki kelime sırasından kaynaklanabilecek olası anlam farklarını Ulamsal Dilbilgisi açısından incelemektir.

Anahtar Kelimeler: Niceleyici Alanı Çokanlamlılığı, Ulamsal Dilbilgisi, Bilgi Yapısı, Türkçe

To My Parents

ACKNOWLEDGMENTS

I would like to thank my advisor Cem Bozşahin for his guidance and support in conducting this research. His knowledge and insight led to overcome many problems that I faced through this thesis. More specifically, without his help it is clear that I wouldn't be able to sort out the data presented by the previous studies of the same subject and extract anything meaningful from them. For this reason, the fact that the title bears my name as the sole author shouldn't be taken too seriously.

I am also grateful to the jury members, for their help, criticism and comments.

I am also grateful to Bob Carpenter, Mark Steedman, L.T.F. Gamut, Patrick Blackburn and Johan Bos for their highly readable books on Logic and Linguistics, without which it would not be possible for me to get back to study after years of delay. They made me 'remember' many things which I had probably never learnt to begin with.

Finally, I would like to thank my family for their endless support and patience.

TABLE OF CONTENTS

ABSTRACT	iv
ÖZ	v
DEDICATION	vi
ACKNOWLEDGMENTS	vii
TABLE OF CONTENTS	viii
LIST OF TABLES	xi
LIST OF FIGURES	xii
LIST OF ABBREVIATIONS	xiv
CHAPTER	
1 INTRODUCTION	1
2 HISTORICALLY IMPORTANT APPROACHES TO QUANTIFIER SCOPE AMBIGUITY	9
2.1 Quantifying-In	11
2.2 Intensionality and Quantifying-In	12
2.3 Storage Methods	13
2.4 Underspecification	13
2.5 A Discourse Based Approach	14
2.6 Combinatory Categorical Grammar	16
2.7 Surface Compositional Scope Alternation	16
2.8 Summary	20
3 QUANTIFIER SCOPE AMBIGUITY IN TURKISH	23
3.1 A Short Literature Survey - Kural (1994), Kornfilt (1996), Göksel (1998) and Zwart (2002)	23
3.2 Kennelly (2003)	26
3.2.1 Locality or Discourse Structure	30

3.2.2	Accompany Type Predicates	31
3.3	Problems Caused by The Lack of Compositionality	32
3.4	Bringing Syntax Back to the Issue of Quantifier Scope in Turkish . .	32
3.5	Is There a Solution Independent of Discourse Structure?	34
4	AN APPLICATION OF STEEDMAN'S THEORIES OF INFORMATION STRUCTURE AND ALTERNATING QUANTIFIER SCOPE TO TURKISH	41
4.1	Information Structure of Turkish	41
4.2	Implications to Quantifier Scope Ambiguity in Turkish	42
4.2.1	The Argument Against Isolated Sentences	42
4.2.2	The Argument Against Intonation	45
4.2.3	The Argument Against Discourse Binding of 'Old Infor- mation' to a Single Entity	49
4.3	Focus and Background in Kennelly (2003)	50
4.4	An Argument from Incremental Processing of Subsequent Utterances	51
4.5	Possessive Constructions and Quantifiers	53
4.6	Prepositional Phrases and Quantifiers	55
4.7	Negation Operator and Quantifiers	56
4.8	Additional Notes and Concluding Remarks	58
4.9	Conclusion	60
5	CONCLUSION	61
5.1	Explanatory Remarks	63
	REFERENCES	68
	APPENDICES	
A	COMBINATORY CATEGORIAL GRAMMAR	75
A.1	Predicate Argument Structure	78
A.2	Combinatory Generalization	79
A.3	Resource Control	83
A.4	Information Structure	83
B	QUANTIFYING-IN AND STORAGE METHODS	85
B.1	Quantifying-In	85
B.2	Storage Methods	87
B.2.1	Cooper Storage	87
B.2.2	Keller Storage	91

C	SURFACE COMPOSITIONAL SCOPE ALTERNATION (STEEDMAN, 1999, 2000a, 2005)	95
C.1	Syntactic Constraints on Quantifier Scope in English	98
C.1.1	Donkey Sentences	101
D	INFORMATION STRUCTURE IN COMBINATORY CATEGORIAL GRAMMAR	107
D.1	Update Semantics	109
E	INFORMATION STRUCTURE OF TURKISH	112
E.1	Some Remarks on Özge and Bozşahin (2006)	115

LIST OF TABLES

Table 1.1	Compositionality, systematicity, direct compositionality and the accounts of quantifier scope ambiguity	5
Table 3.1	Lexical assignments for the quantifier <i>her</i>	37
Table 3.2	Lexical assignments for the quantifier <i>her</i> - Modified.	39

LIST OF FIGURES

Figure 2.1	A derivation for the sentence <i>Every boxer loves a woman</i>	20
Figure 2.2	Another derivation for the sentence <i>Every boxer loves a woman</i>	20
Figure 3.1	A derivation that succeeds	34
Figure 3.2	A derivation that fails	34
Figure 3.3	A derivation for the narrow scope reading of the indefinite in <i>Her doktor bir hasta-yı tedavi etti</i>	37
Figure 3.4	A derivation for the wide scope reading of the indefinite in <i>Her doktor bir hasta-yı tedavi etti</i>	38
Figure 3.5	A derivation for the (incorrectly available without being focused) narrow scope reading of the indefinite in <i>Bir hasta-yı her doktor tedavi etti</i>	38
Figure 3.6	A derivation for the wide scope reading of the indefinite in <i>Bir hasta-yı her doktor tedavi etti</i>	38
Figure 3.7	Lexically corrected version of the derivation of Figure 3.5	39
Figure 4.1	A derivation for the sentence (87)	54
Figure 4.2	A derivation for the sentence <i>Mia knows every owner of a hash bar</i>	55
Figure 4.3	A derivation for the sentence <i>Mia knows every owner of a hash bar</i>	55
Figure 4.4	A derivation for the sentence <i>Mia knows every flower in a garden</i>	56
Figure 4.5	A derivation for the sentence <i>Every boxer doesn't love a woman</i>	57
Figure 4.6	A derivation for the sentence <i>John doesn't love a woman</i>	57
Figure A.1	A CG derivation with minor syntactic features	77
Figure A.2	A CG derivation with minor syntactic features, viewed as a phrase structure	77
Figure A.3	The corresponding phrase structure with traditional node labels	77
Figure A.4	Use of lambda abstraction	79
Figure A.5	Forward B combinator example	80
Figure A.6	Backward B combinator example	81
Figure A.7	Backward crossed substitution example	82
Figure B.1	Derivation for <i>Every boxer loves her-3</i>	85
Figure B.2	Derivation for <i>Every boxer loves a woman</i>	86
Figure B.3	Derivation for <i>Every boxer loves a woman</i>	87
Figure B.4	Derivation for <i>Every student wrote a program</i>	89
Figure B.5	Derivation for <i>Every owner of a hash bar (NP)</i>	94
Figure B.6	Derivation for <i>Every owner of a hash bar (NP)</i>	94
Figure D.1	A derivation for the sentence (87)	110

Figure E.1 Parallel levels of surface and phonetic syntax that were assumed in Steed-	
man (1990)	116
Figure E.2 The derivation for the sentence (181)	122
Figure E.3 F0 curves	125

LIST OF ABBREVIATIONS

3SG	Agreement suffix third person singular	PTQ	‘The proper treatment of quantification in ordinary English.’ (Montague, 1973)
ACC	Accusative Case		
AGR	Agreement	QR	Quantifier Raising
AGRP	Agreement Phrase	RAS	Rheme Alternative Set
BE	to be	SKO	Skolem
CG	Categorial Grammar	TAS	Theme Alternative Set
CCG	Combinatory Categorial Grammar	UVC	Unbound Variable Constraint
		VP	Verb Phrase
DAT	Dative Case		
DC	Descriptive Content		
DP	Determiner Phrase		
DRT	Discourse Representation Theory		
FUT	Future Tense		
GB	Government and Binding		
GEN	Genitive Case		
IMP	Imperative		
LOC	Locative Case		
LF	Logical Form		
MP	Minimalist Program		
MUST	Obligation		
NEG	Negation Suffix		
NP	Noun Phrase		
PASS	Passive		
PAST	Past Tense		
PF	Phonetic Form		
PLU	Plural		
POSS	Possessive		
PP	Prepositional Phrase		
PRES	Present Tense		
PROG	Progressive Tense		

CHAPTER 1

INTRODUCTION

This thesis has two purposes: first, to study the problem of quantifier scope ambiguity in natural language, and the various ways that it has been accounted for, and second, to study the same problem on Turkish data comparatively with previous work on the same topic.

Sentences with scope ambiguities are often semantically ambiguous (that is, they have at least two non-equivalent logical representations) but fail to exhibit any syntactic ambiguity (that is, they have only one syntactic analysis). As the principle of compositionality requires semantic construction to be guided by syntactic structure, we face with an obvious problem: if there is no syntactic ambiguity, there should not be semantic ambiguity. This is not the case as the following example indicates:¹

- (1) Every student wrote some program.

There are two possible readings of the sentence (1), given below in (2)

- (2) a. **every**($\lambda x.$ **student**(x) $\rightarrow$ **some**($\lambda y.(\mathbf{program}(y) \wedge \mathbf{wrote}(x,y))$)
b. **some**($\lambda x.\mathbf{program}(x) \wedge \mathbf{every}(\lambda y.\mathbf{student}(y) \rightarrow \mathbf{wrote}(y,x))$)

Logically, generalized quantifiers can be characterized as functions from properties² to propositions³, with determiners acting as functions from properties into generalized quantifiers⁴. It remains to be explained how quantifiers can remain *in situ* yet take semantic scope around an arbitrary amount of surrounding material. As an illustration, consider the following sentences with quantified subjects and objects, which are adapted from Carpenter (1997, p.213).

¹ In this study, we shall use the simply typed lambda calculus. The notation of both higher order and first order logics shall be used where appropriate.

² A function from entities to propositions, i.e. type $\langle e, t \rangle$.

³ A statement that affirms or denies something and is either true or false, i.e. type t .

⁴ This statement proposes the type $\langle \langle e, t \rangle, t \rangle$ for generalized quantifiers, and the type $\langle \langle e, t \rangle, \langle \langle e, t \rangle, t \rangle \rangle$ to determiners.

- (3) a. Every kid played with some toy.
 b. Some kid broke every toy.

These sentences illustrate the simplest possible example of *scope ambiguity*. The first sentence has two readings, distinguished by the *scope* of the quantifiers **every**²(**kid**) and **some**²(**toy**).

- (4) a. **every**²(*kid*)(λx .**some**²(*toy*)(λy .**play**(*y*)(*x*)))
 b. **some**²(*toy*)(λy .**every**²(*kid*)(λx .**play**(*y*)(*x*)))

Under the first reading, (4a), every kid may have played with a different toy. In this reading, *every kid* is said to take *wide scope*, and *some toy* *narrow scope*. Under the second reading (4b), there must be some toy such that every kid played with that toy. Thus the second reading, (4b), where the object takes wide scope over the subject quantifier, entails the first reading. From this fact alone, it might be tempting to claim that (4a) is the only relevant reading, as its truth is entailed by the truth of the other proposed reading. But now notice that although there is also an entailment between the two readings of (3b), it is the subject wide-scope reading that entails the object wide-scope reading. Similarly, by negating the predicate, any entailment relations between quantifier scopings will be reversed. Furthermore, with other quantifiers, such as *many* and *most*, there is no entailment relationship between the two relative scopings. It is simply not sufficient to provide a quantifier narrow scope at the point at which it can be thought to act as an argument.

The problem of Quantifier Scope Ambiguity has been studied extensively in the literature, but the following list represents the main approaches that shall be investigated in Chapter (2).

1. *Montague's method (also called 'quantifying-in' or 'quantifier raising')*: Montague, in the article with the title "The proper treatment of quantification in ordinary English" (Montague, 1973), has developed a grammar that was capable of generating all of the possible scope alternations of an expression containing general noun phrases.
2. *Storage methods*: Storage methods are an extension of Montague's original method and they are basically a way of separating the considerations of scope ambiguity from syntactic issues and handling them separately. This study examines two of these approaches: *Copper Storage* and *Keller Storage*.
3. A discourse based approach (Farkas, 1997).

4. A recent study within the theory of Combinatory Categorical Grammar, which uses an interesting variant of underspecification called *skolemization*. (Steedman, 1999; Steedman, 2000a, p.70 – 85; and Steedman, 2005)

This list has been chosen for the following reasons: (1) and (2) are very traditional and important. (3) is referred by Kennelly (2003), which is an important study for the subject of quantifier scope ambiguity in Turkish and is a required reading to be able to understand Kennelly's ideas. (4) is the solution designed for CCG, and for that reason it is important.

Modern theories of grammar vary in the degree they adhere to the principle of compositionality. Some theories, which are described as being *monotonic* or *surface compositional* are more rigid in this respect than others, and do not allow operations such as movement or deletion; and do not make use of empty strings in the grammar either. For these theories, words and morphemes are basic building blocks and all syntactic and semantic properties of more complex expressions derive directly from the syntactic categories and the meaning of the building blocks recursively, with an explicit commitment in each step such that later structure changing operations such as movement or deletion are prohibited. According to this view, every word, morpheme and syntactic rule has an interpretation. The principle is also known as “principle of direct interpretation”, “strict compositionality”, “direct compositionality”, and “monotonicity”. Jacobson (1999) elaborates the principle as follows⁵.

Under this view, the combinatory syntactic operations combine expressions into larger ones, and each syntactic operation is coupled with a semantic operation which supplies a model-theoretic interpretation for the new expression built in the syntax. (Not all operations need combine more than one expression; there can also be unary rules which simply map expressions into new expressions, and I will in fact make heavy use of these. Such rules have generally gone under the rubric of “type-shift” operations.) ... This view makes no use of any extra level of LF as input to the model-theoretic interpretation and, consequently, it needs no rules mapping one structure into another. In fact, while the Montagovian view is often referred to as “direct interpretation of surface structures” or “surface compositionality,” this is somewhat misleading – it suggests that a structure is first built by the syntax and then “sent” to the semantics. This, however, is not the case: the semantics interprets as the syntax builds. I will thus call the present program “direct compositionality.” (Jacobson, 1999)

More information about strict/direct compositionality can be found in Montague (1973), Zadrozny (1992, 1994), Jacobson (1996, 1999) and Hausser (1984, 1999).

⁵ In addition to the notion of monotonicity, a theory of grammar is said to be *monostratal*, if it only assumes a single level of representation. One possibility is to reject any structural representation other than the derivation, as in the Minimalist Program (MP). Another possibility is to reject any structural representation other than the logical form as in the Combinatory Categorical Grammar (Steedman, 1996, p.7). MP is monostratal, but is not monotonic. CCG is both monotonic and monostratal.

The assumption of strict compositionality and the methodology it requires is in sharp contrast to transformational theories of grammar which make extensive use of structure changing operations. For example, in Government and Binding framework, scope ambiguities are handled in Logical Form (LF) with an approach similar to Quantifying-In (which is called *Quantifier Raising* in May (1985)), in such a way that a term A has scope over another term B if A c-commands B in LF. If LF derived from surface structure does not give the desired scope relation in terms of a c-command relation, then a movement rule is proposed which alters the scope relations between such terms and brings them to the appropriate positions at LF.

The approaches represented by (1), (2) assume acceptability of such structure changing operations in some form or the other, and therefore are incompatible with the assumptions of monotonic theories of grammar. In (1) (the Montagovian PTQ fragment), the movement rule is syntactic in nature and therefore is a part of the grammar, in (2) (the storage methods) it is a part of semantic interpretation. As for (3), while Farkas (1997) does not tell us how to compute the discourse representation structures she uses, it is known that any syntax-semantics mapping can be defined compositionally (in the worst case, in a case by case basis as shown in Zadrozny (1992)), and therefore the general notion of compositionality (that the meaning of a whole is a function of its parts) is considered to be ‘formally vacuous’. For these reasons this thesis is also concerned with the implications of the problem of Quantifier Scope Ambiguity with respect to the notion of strict compositionality, exemplified by the theory of Combinatory Categorical Grammar.

Table 1.1 summarizes the status of these approaches with respect to compositionality, together with the typical representatives of these notions⁶. While the notion of direct compositionality apparently originates with Montague (1973), when faced with the problem of scope ambiguity Montague has chosen a more indirect approach, and for that reason in Table 1.1 we have shown Montague’s PTQ as not being directly compositional. As the table also implies, according to Zadrozny (2000) ‘systematicity’ is a different notion from compositionality. Direct compositionality subsumes both compositionality and systematicity, and it also introduces restrictions on the methods of semantic construction allowed.

The following list presents the relations between compositionality, systematicity and direct compositionality.

1. Compositionality (that there is a homomorphism between syntax and semantics, (Partee

⁶ In Table 1.1, N.A. means ‘not applicable’.

et al., 1990).

2. Systematicity (that we need a systematic relation, not just a homomorphism, and that we can have a systematic relation without having compositional semantics, (Zadrozny, 1992, 1994, 2000).
3. No structure changing transformations and no empty strings in the grammar, (Jacobson, 1999).
4. Direct Compositionality $\iff$ (1), (2) and (3).

Turkish, which is a flexible word order language, presents further problems. For example, both Kural (1994) and Göksel (1998) have proposed that interpretation of sentences such as the following are dependent on a linearity relation (or linear order):

- (5) Her çocuk bir öğretmene çiçek verdi.
 every child a teacher-dat flower gave
 ‘Every child gave flowers to a teacher.’

a. Distributive reading: For every child there was a teacher, such that each child gave flowers to a (different) teacher.

b. Nondistributive reading: There was some teacher to whom every child gave flowers.

- (6) Bir çocuk her öğretmene çiçek verdi.
 a child every teacher-dat flower gave
 ‘A child gave flowers to every teacher.’

a. * Distributive reading: For every teacher there was a child, such that each teacher received flowers from a (different) child.

b. Nondistributive reading: There was some child who gave flowers to every teacher.

- (7) Bir hemşire bakıyor her hastaya.
 a nurse is seeing every patient-dat.

Table 1.1: Compositionality, systematicity, direct compositionality and the accounts of quantifier scope ambiguity

	Compositionality (Partee, 1990)	Systematicity (Zadrozny, 2000)	Direct Compositionality (Jacobson, 1999)
Quantifying-In	Yes	Yes	No (allows movement)
Storage Methods	No	Yes	No (allows movement)
Farkas (1997)	N.A.	N.A.	N.A.
Steedman (2005)	Yes	Yes	Yes

‘A nurse is seeing every patient’. (Distributive reading available.)

In (6), unlike (5), the narrow scope reading of the indefinite (so called ‘Distributive Reading’ in Göksel’s terms) is considered to be unavailable (unless it is focused as in (7)), and for this reason a linearity constraint on quantifier scope is proposed, but it is made sensitive to focus. On the other hand, it is known that for the following sentence (8) in English, which bears a similar linear order of surface constituents, the narrow scope reading of the indefinite *somebody* is known to be available, and we are not aware of a claim that states it is available only if *somebody* is focused.

(8) Somebody loves everybody.

Another issue that has been identified in Turkish is about the post verbal arguments. For example, Kural (1994) proposes that post-verbal quantification has wide scope with respect to pre-verbal positions (as cited in Kornfilt (1996)), and uses this claim to argue against Kayne’s Linear Correspondence Axiom (Kayne (1994)). This means that, for example, in the following sentence, the indefinite *üç kişiye* is not considered to have a narrow scope reading, and is restricted to a wide scope interpretation.

(9) Her besteci bu yıl eserlerini ithaf etmiş üç kişiye.

Every composer this year artwork-plu-3sg-acc dedicate-past three person-dat

Kennelly (2003) argues against Kural (1994) and Göksel (1998) and proposes that the linearity constraint is unnecessary because there is a one-to-one correspondence between the linear order and the default order of Turkish discourse structure, which has traditionally been accepted as being *topic-focus-predicate-backgrounded elements*, in linear order. Therefore, the narrow scope reading of the indefinite in (6) is considered to be unavailable not because of a linearity relation but rather because it is not focused. While lifting Kural and Göksel’s linearity constraint, Kennelly (2003) proposes two special constraints about quantifier scope: One of which concerns the locality (the predicate domain) and the other is a discourse structure constraint.

In this thesis we shall apply Steedman’s theories of Surface Compositional Scope Alteration (Steedman, 1999, 2000a, 2005) and Information Structure of English (Steedman, 1990, 2000a, 2000b, 2003) (a version of which is available for Turkish, due to Özge and Bozşahin (2006)) to Turkish and argue that the proposed constraints above (post-verbalness, linearity, locality and discourse structure constraints) are all not only unnecessary, but also incorrect. In

particular, we also reject the claim that the narrow scope reading of an indefinite is available *because* it is focused, a version of which could be mistakenly stated to be that it is available due to intonationally marked contrast in a contrastive context. While defending such a position was our initial intent, later we have realized that it isn't possible to support this idea with Steedman's theories; quite the contrary, the application of Steedman's theories above to Turkish predicts that the narrow scope reading is available regardless of such contrast, but a contrastive context nevertheless is helpful to establish a convincing example. The reason is that, when an underspecified skolem term is contrasted versus another underspecified skolem term of a similar form, (or the resulting specified skolem terms, if that matters), the scope relations themselves are not contrasted. Therefore, we claim that some extra-linguistic inference mechanism that uses the previously established scope relation in the discourse is involved and it reveals the narrow scope reading of the indefinite in (6), which otherwise is obscured due to the way such sentences are processed. An interesting theory (due to Morrill (2000)) considers the memory load during sentence processing and predicts left-to-right scope preference in sentences such as (8), among other linguistic phenomena, in spite of a competence grammar that does not impose such restrictions. For these reasons we believe that the narrow scope reading is more easily detectable in a contrastive context, but this result is only predictable if we were to also consider the characteristics of a possible cognitively plausible utterance processor and propose that previous context have an influence on the way succeeding utterances are processed.

The remainder of the thesis is organized as follows: Chapter (2) is a brief literature survey on the topic of quantifier scope ambiguity, which gives more information about the approaches listed above. Chapter (3) describes the approaches that were proposed in studies of Turkish and examines them from the point of view of monotonicity and strict compositionality and points out that Göksel (1998) is not compositional. Chapter (4) is an application of Steedman's theories of Surface Compositional Scope Alternation and Information Structure to Turkish, which shows that in CCG, the scope ambiguities in Turkish can be explained without extra stipulation or special constraints. The last chapter summarizes our conclusions. Appendix contains relevant background material on several topics, with the exception that Appendix (E) also investigates the question of whether Steedman's CCG model of English information structure as presented in Steedman (2000a, 2000b) is compatible with the findings of Özge and Bozşahin (2006) regarding Turkish information structure, and gives a positive answer to this question. This part of the thesis is presented in an appendix, because

the question has turned out to be irrelevant to the subject matter of the thesis, because our conclusion states that information structure only affects the discourse felicity of quantificational sentences, and that contextual dependencies are more complex than what can be formulated with the idea of discourse binding of “old information”.

CHAPTER 2

HISTORICALLY IMPORTANT APPROACHES TO QUANTIFIER SCOPE AMBIGUITY

The approaches to the problem of quantifier scope ambiguity mainly differ at the linguistic level(s) they seek to find solutions. Some approaches seek to find structural solutions, as in Montague's approach (1973) which is called *Quantifying-In*. Quantifying-In employs a movement rule and is compatible with the assumptions of the Government and Binding framework in that scope relations are considered to be arising from structural issues, with the crucial difference being that Montague had attempted to solve the problem at the level¹ of syntax by proposing a way to implement the idea of movement in his grammar fragment. On the other hand, proposals within the framework of GB (May (1985), for example) would attempt a solution at the level of LF instead, but using exactly the same idea, which is called Quantifier-Raising in May (1985). Similarly many others seek to separate the issues of scope and syntax. Underspecification based Storage methods essentially implement the same idea of NP movement, but unlike Montague's method, they do so without modifying the grammar, and propose to solve the problem at the semantic level, totally isolating the issue from that of syntax. Still others seek explanations at the discourse level, which include the Discourse Representation Theory (DRT) (Kamp and Reyle, 1993), and more recently Farkas (1997), among many others). On the other hand, the issue of quantifier scope is apparently affected by syntactic constraints, such as being sensitive to *island constraints*, in the sense of Ross (1967), as Park (1995), and Steedman (1999, 2000, 2005) argue.

With respect to word order, which is an important issue for a scrambling free word or-

¹ Since the principle of direct compositionality is implicit in Montague's PTQ fragment, actually there is only one level in the mentioned fragment, which is a logical grammar. On the other hand, the quantification rule that Montague has proposed is equivalent to a movement rule at LF, but specified as a part of the grammar, which justifies the claim that it is a syntactic rule.

der language like Turkish, the linguistic frameworks differ in their assumptions. In transformational grammar, the word order was initially a part of syntax, as in the Government and Binding theory. In Minimalist Program, word order is taken out of grammar and left to the PF. In non-transformational theories, word order usually is a part of the grammar. This is a crucial factor because the accounts of quantifier scope ambiguity, like any other linguistic account, need to adhere to the basic assumptions of the framework they build upon. For example, it is very natural for a linguist working within the framework of Minimalist Program to prefer discourse level accounts, as the word order is not considered to be a part of sentence level syntax. The issues arising from linear precedence relations cannot be explained at the sentential level. For theories that word order is a part of the grammar, like Government and Binding (GB) within the transformational tradition and almost all non-transformational theories, sentence level accounts are available, in addition to discourse level ones. Within GB, transformationalists sought to relate word order to structural factors, as in the work of Kayne (1994), whose (and his followers) research program can be summarized as the translation of precedence relations into c-command terms. According to this approach, the lower an element in a string, the lower it is in terms of hierarchical relations.

Needless to say, non-transformationalists also, need to adhere to their basic theoretical assumptions, while attempting to solve a linguistic problem. For example, according to Steedman's Combinatory Categorical Grammar, the solution must not involve transformations such as movement (due to the principle of monotonicity), and it must adhere to the principle of monostratality, which excludes approaches that map linear precedence relations to structural factors at the logical form, because to begin with CCG has its own very specific notion of logical form, called *the predicate argument structure*, which is the only level of representation that CCG assumes and it cannot afford to have yet another incompatible one. Storage methods are similarly excluded as being just another implementation of a transformation at the logical form. Within this framework, discourse based approaches to quantifier scope ambiguity are available, but the freedom that is available to the transformationalists aren't, again due to the very specific definition of logical form and the strict form of compositionality that CCG assumes. With respect to word order, CCG has the power to keep it as a part of the grammar via subcategorization, and we shall see in Chapter (4) that according to a research report by Özge and Bozşahin (2006), in Turkish there are prosodic restrictions on information structure, which in turn imposes several constraints on felicitous word ordering possibilities in a given discourse.

In this chapter we survey the two most historically significant approaches to quantifier scope ambiguity, Montague’s quantifying-in scheme and Storage mechanisms. We shall also examine a discourse based approach to quantifier scope, namely the work of Farkas (1997). We will conclude this chapter by discussing Steedman’s solution proposal within the framework of Combinatory Categorical Grammar, under the title *Surface Compositional Scope Alternation* in section (2.7), based on which we shall attempt to formulate our own theory of quantifier scope ambiguity in Turkish.

2.1 Quantifying-In

The most well known and widely studied approach to quantifier scoping phenomena is that of Montague, as embodied in his PTQ grammar (Montague, 1973) (for a more readable introduction to Quantifying-In than Montague’s original paper, see Blackburn and Bos (2005), a summary of which is available in Appendix (B)). Classical Montague semantics is the source of the *direct (surface) compositionality* principle that was discussed in Chapter (1). However, motivated by the problem of quantifier scope ambiguities, Montague introduced a rule of quantification that allowed a more indirect approach. The effect of the rule in practice is that it allows a permutation of NP arguments to predicates. For example, in the following example, which is from Hobbs and Shieber (1987) (as cited by Park (1995)), it is reported that there are three quantifiers and 6 different ways of ordering them.

(10) Two representatives of three companies say most samples. (Hobbs and Shieber, 1987)

Among these 6 different readings, Hobbs and Shieber (1987) reports that one of them is excluded due to the unbound variable constraint (UVC). However, while discussing Cooper Storage (which is another implementation of the idea of quantifier movement, with very similar characteristics with the exception of the modularity of grammar design), Blackburn and Bos (2005) points out that such readings with unbound (free) variables are generated because of the fact that Cooper Storage allows too much freedom in retrieving expressions from the store (and that Keller Storage solves this problem by taking into account the hierarchical structure of sentences). Thus, we may conclude that the unbound variable constraint, which Park (1995) wasn’t certain about the necessity of, can be dropped by stating that Quantifying-In allows too much freedom in reordering arguments to predicates. Thus, unlike Park (1995), we consider it safe to explicitly state that the UVC can be dropped, and that UVC does not

need to be kept once Quantifying-In and its variants are dropped, the reason being that it does not have an independent reason for its existence.

The proper way to get rid of quantifier movement is to recognize the referential nature of indefinites, according to which any indefinite can take wide scope because of its referential nature, as Park (1995) and Steedman (1999, 2000, 2005) argues, which is the subject matter of section (2.7) and Appendix (C).

2.2 Intensionality and Quantifying-In

As Park (1995) also points out, Quantifying-In was originally proposed as a technique to produce appropriate semantic forms for *de re* interpretations of NPs inside intensional operators, such as *believe*.

- (11) a. John believes that a Republican will win. (Park, 1995)
 b. $\exists r.\text{repub}(r) \wedge \text{believe}(\text{john}, \text{will}(\text{win}(r)))$ (The *de-re* reading.)
 c. $\text{believe}(\text{john}, \exists r.\text{repub}(r) \wedge \text{will}(\text{win}(r)))$ (The *de-dicto* reading.)

According to Park (1995), (11b) exemplifies the fact that *de re* interpretations of NPs are strongly related to referential NP semantics, which was what Montague had intended. The intended interpretation of expressions such (11b) and (11c) is as follows: The notation in (11c) is intended to imply that *the republican that John believes* is the one that exists in the (possible) world according to John's believes, and a scope relation between *believe* and the existential quantifier is taken as an indication of this fact. That is, the republican *r* in (11c) refers to an entity in the possible world of John's believes. Similarly, (11b) is an implicit statement of the fact that the republican in question has a referent in the actual world, and this reading is considered to be available due to a scope relation between the existential quantifier and the intensional operator *believe*, and the fact that *r* out-scopes *believe*.

Adequacy of this proposal can be questioned, though. For example, an intensional interpretation can be attributed to examples that are normally supposed to have a purely extensional interpretation. For example, in (12b), while the intended reading is an extensional one, and the rule of Quantifying-In can generate the two possible extensional readings as in (12a) and (12b), the sentence could also possibly be used to mean that *every man loves a woman in his dreams* as in (13), in a properly established context.

- (12) a. Every man loves a woman.
 b. $\forall x.\mathbf{man}(x) \rightarrow \exists y.(\mathbf{woman}(y) \wedge \mathbf{loves}(y,x))$
 c. $\exists y.\mathbf{woman}(y) \rightarrow (\forall x.\mathbf{man}(x) \wedge \mathbf{loves}(y,x))$

(13) In his dreams, every man loves a woman.

The additional prepositional phrase makes it clear that an intensional context can be established independently of the choice of the predicate. Moreover, recursively embedded intensional constructions are possible as in (14) and as in (15).

(14) In his dreams, every man believes a woman.

(15) Every man believes that a woman thinks that he is handsome.

2.3 Storage Methods

One of the important distinctions between Storage Methods and Montague's PTQ is due to a principle of grammar design: Modularity. Both can be considered as making an implementation of the idea of quantifier movement, but in Montague's PTQ it is an otherwise unjustified syntactic rule, whereas Storage Methods make it a rule of interpretation. Otherwise, they share similar properties, and have the same argument-reordering effect discussed in section (2.1). Cooper Storage can be considered as an implementation of quantifier movement which, unlike Montague's PTQ, keeps the syntactic rules clean. Keller Storage is a revised version which takes hierarchical factors into account, and it makes unbound variable constraint (UVC) of Hobbs and Shieber (1987) unnecessary (It's also possible to consider Keller Storage as an implementation of UVC.). Park (1995) discusses both Storage Methods and Montague's PTQ under the title *Quantifying-In*, because both are implementations of the same idea of quantifier movement, differing only from a grammar engineering point of view².

2.4 Underspecification

From our point of view, one of the most interesting properties of Storage methods is that they are one of the earliest example of *underspecified representations*. An underspecified

² These issues are discussed at some length in Blackburn and Bos (2005), with examples that illustrate the points, which are summarized in Appendix (B). In particular, the sentence (16) is being used to show that Cooper Storage generates ill-formed semantic expressions with unbound variables, without mentioning Hobbs and Shieber (1987)'s unbound variable constraint (UVC). On the other hand, we are of the opinion that it is safe to associate this particular problem with UVC.

(16) Mia knows every owner of a hash bar.

(Blackburn and Bos, 2005)

representation is a kind of information packaging structure that can possibly have more than one meaning, but it is not correct to consider such structures simply as lexical ambiguity. Underspecified representations can be as simple as *skolemization* as in Steedman (2005) or as complex as *hole semantics* (Blackburn and Bos, 2005, p.129). But even in the case of skolemization, we shall see that the structure being used is more than just an overloaded meaning assignment to a single lexical item – it comes with a well defined computational way of deriving the possible senses. That is, an underspecified representation is an overloaded information packaging which is defined together with an algorithm of unpacking it. In Storage methods, the unpacking operation is called *retrieval*. In skolemization, it is called *specification* of a skolem term. The algorithm can be deterministic as in Storage methods, or it can be nondeterministic, as in Steedman (2005). For details see Appendix (B) and Appendix (C).

Blackburn and Bos also shed some light upon the linguistic meta-theoretical status of underspecified representations, which can be summarized as the idea of taking underspecification very seriously as a genuine form of linguistic meaning representation, not just as a tool for computation. The following statement is theirs:

In the past, storage-style representations seem to have been regarded with some unease. They were (it was conceded) a useful tool, but they appeared to live in the conceptual no-man’s land – not really semantic representations, but not syntax either – that was hard to classify. The key insight that underlies the current approaches to underspecification is that it is both *fruitful* and *principled* to view representations more abstractly. That is, it is becoming increasingly clear that the level of representations is richly structured, that computing appropriate semantic representations demands deeper insight into this structure, and that – far from being a sign of some fall from semantic grace – semanticists should learn how to play with representations in more refined ways. (Blackburn and Bos, 2005, p.128)

2.5 A Discourse Based Approach

In his paper *Evaluation Indices and Scope* (Farkas, 1997), Donka Farkas proposes a theory of scope that is claimed to be less structure driven than the traditional Montagovian and GB approach. As explained above, the traditional view of scope is structural in the sense that the relative scope of two expressions is taken to be determined by their relative position at some level where hierarchical relations are encoded. Under this assumption, an expression e_1 is in the scope of e_2 iff e_2 commands e_1 at LF or at some other representational level.

In Farkas (1997) the relative scope of two expressions is considered to be a matter of possible dependencies between indices, which are Kaplan-style coordinates of evaluation.

According to Kaplan (1979), expressions are evaluated relative to an index I , which is a sequence of coordinates including a world w , a time t , a three-dimensional location, p , as well as a coordinate for the speaker and addressee. Coordinates of evaluation fix the parameters with respect to which the denotation of an expression is determined. In the terminology of Farkas, Kaplan's coordinates are called *indices of evaluation*.

Farkas (1997) is concerned with the scope of noun phrases (DPs), with respect to intensional operators (modals, the conditional operator), intensional predicates and nouns (such as *believe*, *belief*, *dream*), and quantificational DPs, that is, DPs whose determiner (D) contributes a proportional quantifier. In this theory, the semantic representation of a noun phrase includes at least a subscripted variable x_n , and a descriptive content, DC_n , in the form of a predicative expression on x_n . In addition, quantificational DPs, i.e. DPs whose D is a proportional quantifier such as *every*, *each* and *most*, induce a tripartite quantificational structure in which the quantifier is contributed by the determiner, and where the variable and the DC occur in the Restrictor, as in (17). In this case the quantificational force of the DP is determined by the quantifier it contributes. For example, (18) is assumed to have the quantificational structure in (19).

(17) Restrictor Q Nuclear Scope

(18) Every student left.

(19) Restrictor: $\forall x_1$ Nuclear Scope:
 x_1 **leave**(x_1)
student(x_1)

The main problem with this theory is that the so called 'possibilities' are still determined by structural factors (the theory does not allow a scope relation that isn't licenced structurally), and we free to choose an interpretation based on these possibilities. This in turn brings the question, does the theory have anything to say about scope ambiguities, other than just stating a philosophical position such as 'we see the problem as a matter of content instead of structure', and the answer is no, for the following reasons:

1. Farkas (1997) does not tell us how to construct the discoursal representation that is being used.
2. It does not tell if the theory brings any restrictions on scope relations that are not imposed structurally.

Given these facts, the main contribution of Farkas (1997) to the issue at hand can be considered as a clarification of some important philosophical assumptions. In other words, seeing scope ambiguity as a matter of content instead of as a structural relation does not solve the problem, it just clarifies its description and restates it more clearly. How to arrive at an interpretation of that proper content from a given structural description is not specified. Unless the idea of quantifier movement is abandoned, the unjustified argument-permutations mentioned in section (2.1) are still available, the solution of which requires realization of the referential nature of NPs and the adequacy of a surface compositional way of supplying arguments to predicates as Park (1995) points out.

2.6 Combinatory Categorical Grammar

When we introduce a monotonic and monostratal theory of grammar, such as *Combinatory Categorical Grammar*, the accounts above seem inappropriate in a number of ways: For example, Quantifying-In and Storage methods are essentially implementations of the idea of quantifier movement, but the principle of monotonicity disallows such structure changing operations.

A violation of the principle of monostratality is exemplified by Farkas (1997): The representation used by the theory, as represented by (17) cannot be the only representational level a linguistic theory assumes, because such expressions are built only if they are licenced by structural factors, which means that a representational level about the underlying syntactic structure must also be assumed. Another concern from the view point of CCG is that the *lf* that is assumed is incompatible with the predicate argument structure of CCG. On the other hand, there is no reason that the indexicality assumed by the theory cannot be migrated into CCG.

2.7 Surface Compositional Scope Alternation

As we have seen above, we have yet to see an account of quantifier scope that respects the hypothesis of direct compositionality. Montague's method involves a trick that implements the idea of movement syntactically (thus, it isn't strictly compositional; or formally we can state that it is strictly compositional but in an intuitively wrong manner), and Storage methods do the same thing by using an underspecified semantic representation; both of which allows too much freedom in reordering NP arguments to the predicates and discourse based methods

(as exemplified by Farkas (1997)) often assume the same structural treatment and thus do not introduce appropriate constraints upon such argument ordering devices.

Steedman's recent work (Steedman, 1999, 2005) is an attempt to find a strictly compositional solution, because this is what the basic principles of CCG requires. The solution must not involve any representational level other than the predicate argument structure (due to the principle of monostratality) and it must be strictly compositional (due to the principle of monotonicity). This is not to say that a theory of quantifier scope in CCG must be locked to the sentential level: There are discourse related phenomena that needed to be taken into account, but the whole theory does not need to be specified at the discourse level. For example, as in the indexical theory of Farkas, the problem of the selection of unique possible world among many available ones (so called the *scope of the DC* in Farkas's terms) can be handled at the level of discourse, where it is required, as exemplified by the dialogue in (20), but it can also be handled syntactically, where it is appropriate, as the sentence (11) illustrates.

(20) A. Every man loves a woman...

B. Really?

A. In his dreams, that is.

This raises a question: Since the predicate-argument structure assumed by CCG is not intended for representing a discourse, how do we keep the number of assumed representational levels to one, and still handle issue of discourse representation? For an answer to this particular question, see Kruijff (2001, p.36), in which a single representational level based on hybrid modal logic for both sentential and discourse level semantics is proposed.

Steedman (2005) handles the syntactic case exemplified by (11) by assigning a category that introduces intensional variables to the *environment* set, as in (21).

(21) $\text{seeks} := (\mathbf{S} \setminus \mathbf{NP}_{3\text{SG}}) / \mathbf{NP} : \lambda x. \lambda y. [\text{seek}'_{xy}]^{\{i_y\}}$ (Steedman, 2005)

Note that the category assignment in (21) explicitly introduces an intensional variable i_y to the environment set, which is available to the *skolem specification rule* and does not overload the interpretation a preexisting notation originally intended for nonintensional expressions, as (11) does.

We shall give a summary of this theory in Appendix (C), which explains the mentioned syntactic constraints on quantifier scope, introduces skolem terms, and explains theory specific concepts such as the *skolem specification rule* and *environment*. For the time being, we

shall only introduce some of the basic theoretical devices about the theory, as we shall be referring to these rather frequently.

In first order logic, the term skolemization refers to the elimination of existential quantifiers, so that a new *equisatisfiable* formula (i.e. a formula which can be satisfied under the same conditions) is obtained. Skolemization is an application of the equivalence:

$$(22) \quad \forall x \exists y M(x, y) \leftrightarrow \exists f \forall x M(x, f(x))$$

In (22), M refers to the matrix of the quantificational formula. We shall first explain this formula in a purely mathematical context, away from linguistic concerns, and afterwards we shall state its linguistic significance. Intuitively, the meaning of this formula is that if there is some y that satisfies the left hand side of (22), then there is some function called f , which defines y in terms of x (i.e. $y = f(x)$). In the worst case, f might be defined in a case by case basis (i.e. if $x = 1$, then $y = 4$, if $x = 2$, then $y = 5$, and so on). Alternatively, the value of y might be such that, f may map x to the same value, say 7, whatever the value of x may be.

Variables bound by existential quantifiers which are not inside the scope of universal quantifiers can simply be replaced by constants: that is, $\exists x(x < 3)$ can be changed to $(c < 3)$, in which c is a constant. In this case f is a constant function (i.e. it's the same for all values of x , that's why it can be replaced with the constant c). We shall be referring this usage by the term *skolem constant*, which is merely a convenient term since formally a skolem constant is actually a constant skolem function.

When the existential quantifier is inside a universal quantifier, the bound variable must be replaced by a skolem function of the variables bound by universal quantifiers. Thus $\forall x(x = 0 \wedge \exists y(x = y + 1))$ can be written as $\forall x(x = 0 \wedge x = f(x) + 1)$. We shall be referring this usage by the term *skolem function*. Again this is merely a convenient term, since formally a skolem function may be a constant function or a non-constant function.

When the kind of f is unknown, or underspecified, we shall be using the term *underspecified skolem term* and use the notation $skolem'x$ (or shortly $sko'x$), in which $skolem'$ or sko' refers to f . Since the function $skolem'$ may have one of the alternative meanings above (i.e. a skolem constant or a skolem function that is truly dependent on x), we can explicitly differentiate them with a more explicit notation. We shall be using $sk'x$ with the meaning *a skolem constant* and $sk^{\{x\}}x$ with the meaning *a skolem function that truly depends on x* . Differentiating between these two possible meanings of skolem terms is called *specification of an underspecified skolem term* and the set $\{x\}$ which is seen as a superscript on $sk^{\{x\}}x$ is the set

of variables upon which the skolem function depends on, which may possibly contain more than one variable (or alternatively it may be empty). If an underspecified or specified skolem constant/function needs to satisfy further properties, they are notated as follows, respectively: $skolem'(\lambda x.donkey'x)$, $sk_{\lambda x.donkey'x}$ and $sk_{\lambda x.donkey'x}^E$, where E is the environment set. Where there is no notational confusion, the lambda term will be omitted as in $skolem'(donkey'x)$, $sk_{donkey'x}$ and $sk_{donkey'x}^E$. When there is more than one skolem term in a logical formula, we shall distinguish them with different but arbitrary numerals, as in sk_{53} and sk_{95}^E .

We shall later see that calculation of the environment set is associated with the combinatory rules, which can be interpreted as an implementation of the mathematical fact that “the bound variable must be replaced by a skolem function of the variables bound by universal quantifiers”, which was stated above.

Linguistic significance of skolemization is that, according to Steedman (2005), existentials such as *a/an*, *some*, *at least*, and *most* are not true quantifiers, but underspecified skolem terms. Unlike true quantifiers, they do not invert scope, because they are not quantifiers to begin with. Once the fact that indefinites and other existentials are not quantifiers is recognized, it turns out that a movement analysis is not necessary to be able to account quantifier scope ambiguities. Indefinites are always considered to be referential (i.e. both the wide and the narrow scope readings of indefinites are considered referential). The wide scope reading is accounted by skolem constants, and the narrow scope reading which introduces a dependency to the bounding quantificational variables is accounted by skolem functions. For these reasons, indefinites are defined to be underspecified skolem terms in the lexicon, as follows:

$$(23) \quad a/an/some := \mathbf{NP}^\uparrow_{3SG} / \mathbf{N}_{3SG} : \lambda p.\lambda q.q(skolem'p)$$

(23) is a category schema for determiners like *a/an/some*, making them functions from nouns to type-raised noun phrases, schematizing over them using the $\mathbf{NP}^\uparrow$ abbreviation.

True quantifiers are defined as usual, but they contribute a variable to the environment set, as follows:

$$(24) \quad every, each := \mathbf{NP}^\uparrow_{3SG} \setminus \mathbf{N}_{3SG} : \lambda p.\lambda q.\lambda \dots \forall x[p x \rightarrow q x \dots]^{\{x\}} \quad (\text{Steedman, 2005})$$

(24) is a category schema for determiners like *every*, making them functions from nouns to type-raised noun phrases, schematizing over them using the $\mathbf{NP}^\uparrow$ abbreviation, where the schematized types are simply the syntactic types corresponding to a generalized quantifier (Steedman, 2005).

Environment calculation is simply associated with combinatory rules. When a combinatory rule combines two constituents, the environment set of the result is the union of the environment set of the functor category and the environment set of the argument category.

Figure 2.1 illustrates the theoretical apparatus introduced so far.

$$\begin{array}{c}
\begin{array}{ccccc}
\textit{Every} & \textit{boxer} & \textit{loves} & \textit{a} & \textit{woman} \\
\hline
(\mathbf{S}/(\mathbf{S}\backslash\mathbf{NP}))/\mathbf{N} & \mathbf{N} & (\mathbf{S}\backslash\mathbf{NP})/\mathbf{NP} & \mathbf{NP}/\mathbf{N} & \mathbf{N} \\
\hline
: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\} : \lambda x. \textit{boxer}'x : \lambda x. \lambda y. \textit{love}'xy : \lambda x. \textit{sco}'x : \textit{woman}'
\end{array} \\
\hline
\begin{array}{ccc}
(\mathbf{S}/(\mathbf{S}\backslash\mathbf{NP})) & & \mathbf{NP} \\
: \lambda q. \forall x [\textit{boxer}'x \rightarrow qx] : \{x\} & & : \textit{sco}'\textit{woman}'
\end{array} \\
\hline
\begin{array}{c}
\mathbf{S}\backslash\mathbf{NP} \\
\lambda y. \textit{love}'(\textit{sco}'\textit{woman}')y
\end{array} \\
\hline
\begin{array}{c}
\mathbf{S} \\
: \forall x [\textit{boxer}'x \rightarrow \textit{love}'(\textit{sco}'\textit{woman}')x] : \{x\}
\end{array} \\
\hline
: \forall x [\textit{boxer}'x \rightarrow \textit{love}'(\textit{sk}'_{\textit{woman}'}^{\{x\}})x] : \{x\} \quad \text{skolem specification}
\end{array}$$

Figure 2.1: A derivation for the sentence *Every boxer loves a woman*

In Figure 2.1, the execution of the skolem specification rule converts the underspecified skolem term $\textit{sco}'\textit{woman}'$ to a specified skolem term $\textit{sk}'_{\textit{woman}'}^{\{x\}}$ using the contents of the associated environment set. This results with a narrow scope reading of the indefinite *a woman*. Since the skolem specification rule can be executed at any time during derivation, it can also be executed at a time when the contents of the environment associated with $\textit{sco}'\textit{woman}'$ is empty, as in Figure 2.2, which yields a wide scope reading for the indefinite *a woman*.

$$\begin{array}{c}
\begin{array}{ccccc}
\textit{Every} & \textit{boxer} & \textit{loves} & \textit{a} & \textit{woman} \\
\hline
(\mathbf{S}/(\mathbf{S}\backslash\mathbf{NP}))/\mathbf{N} & \mathbf{N} & (\mathbf{S}\backslash\mathbf{NP})/\mathbf{NP} & \mathbf{NP}/\mathbf{N} & \mathbf{N} \\
\hline
: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\} : \lambda x. \textit{boxer}'x : \lambda x. \lambda y. \textit{love}'xy : \lambda x. \textit{sco}'x : \textit{woman}'
\end{array} \\
\hline
\begin{array}{ccc}
(\mathbf{S}/(\mathbf{S}\backslash\mathbf{NP})) & & \mathbf{NP} \\
: \lambda q. \forall x [\textit{boxer}'x \rightarrow qx] : \{x\} & & : \textit{sco}'\textit{woman}'
\end{array} \\
\hline
\begin{array}{c}
\mathbf{S}\backslash\mathbf{NP} \\
\lambda y. \textit{love}'(\textit{sco}'\textit{woman}')y
\end{array} \\
\hline
\begin{array}{c}
\mathbf{S} \\
: \forall x [\textit{boxer}'x \rightarrow \textit{love}'(\textit{sk}'_{\textit{woman}'}^{\{x\}})x] : \{x\}
\end{array}
\end{array}$$

Figure 2.2: Another derivation for the sentence *Every boxer loves a woman*

2.8 Summary

Montague's approach and Storage methods have been explicitly designed to allow scope alternation and generate the possible readings in a compositional way, though, not in a strictly compositional way. Both Montague's PTQ and Storage methods are implementations of the

idea of movement, but they place no linguistic restrictions on such movement. Not all the readings generated by these methods are linguistically justified.

For example, for the quantifier raising method, Hobbs (1983) exemplify a problem with the following sentence:

- (25) In most democratic countries most politicians can fool most of the people on almost every issue most of the time.

According to Hobbs, this sentence has 120 different distinct readings (or quantifier scopings) (which originates from the fact that there are $5! = 120$ different ways of ordering 5 arguments to a predicate.). While we have not examined all possible CCG derivations, the same sentence is guaranteed to have less than or equal to 8 different readings (due to the fact that there are 3 skolem terms yielding $2^3 = 8$ different configurations), which is a much more reasonable result³.

Quantifier Raising is a standard approach to quantifier scope ambiguity within the framework of Transformational Grammar and it's similar to Montague's PTQ with the departure that the movement operation is predicated at LF instead. The question is, however, is to find linguistic restrictions on such movement. One particular proposal is the unbounded variable constraint of Hobbs and Shieber (1987). On the other hand, the proper way of solving this problem is to abandon the idea of movement and to recognize the referential nature of indefinites, as Park (1995) argues.

Montague's PTQ considers movement as a syntactic rule. The merits of compositionality over such a grammar is limited, even though formally it may still be strictly compositional⁴. Storage methods solve this problem by a modular grammar design, making the movement rule a part of semantic interpretation, but the result is still an implementation of movement.

According to Lappin and Zadrozny (2000), storage methods are 'non-compositional but systematic'. The reason is that the interpretation rule involved (so called the 'retrieval step') makes it possible to use expressions to build logical formulae, but the expressions that are used to build a logical form do not syntactically combine in the corresponding phrase structure.

Farkas (1997) does not elaborate how to compute the mentioned discourse representation structures. Therefore, we cannot say whether the method is compositional or not.

³ Since Steedman (2005) does not address adverbial quantification we haven't counted the last *most* as a skolem term.

⁴ The reason is that in PTQ, movement is a syntactic rule – formally it is not possible to say that the account isn't strictly compositional, but obviously there is something wrong with it if we happen to think what it practically does instead of just thinking about what it formally does.

Steedman (2005), in line with the assumptions of CCG, does not make use of movement or other equivalent mechanisms like storage methods, because once the fact that indefinites are referential and never quantificational is recognized, a movement analysis turns out to be unnecessary. The method proposed can also impose syntactic constraints on scope relations as explained in Appendix (C).

Due to space limitations we have excluded most of the details of Farkas (1997) and Steedman (2005) and the range of phenomena they are intended to explain. For details, the reader is recommended to read the original papers by these authors.

CHAPTER 3

QUANTIFIER SCOPE AMBIGUITY IN TURKISH

The methods that we have examined in the last chapter (Montague’s approach, Storage methods, The Indexical Theory (Farkas, 1997) and Surface Compositional Scope Alternation (Steedman, 2005)) have all been explicitly designed to allow scope alternation. Turkish provides challenges to these approaches, as this chapter shall elaborate.

3.1 A Short Literature Survey - Kural (1994), Kornfilt (1996), Göksel (1998) and Zwart (2002)

To our knowledge, the study of quantifier scope ambiguity in Turkish dates back to Kural (1994), who claimed that S-structure arrangement¹ of two quantifiers determines their scopes with respect to each other. In this proposal Logical Form does not play any role, with the possible exception of post-verbal quantification, which is claimed to have wide scope with respect to pre-verbal positions (as cited in Kornfilt (1996)).

Kural’s study was an attempt to argue against the Linear Correspondence Axiom (LCA) of Kayne (1993,1994), according to which asymmetric c-command and linear sequence directly correlate; complements in all languages must follow their heads in underlying structure; there is no left-branching, nor rightward movement; “extraposed” constituents, therefore, are actually base-generated in their surface positions, and the material that precedes them is moved leftwards.

According to the Linear Correspondence Axiom (LCA), the later an element in a string, the lower it is in terms of hierarchical relations. However, according to Kural (1994) this is not the case for Turkish, and the post-verbal elements, which according to Kayne have to be c-commanded, may in fact be in a position to c-command. For example, Kural proposes that

¹ Also called *linearity* or word order.

post-verbal argument *üç kişiye* in the following sentence has wide scope with respect to the pre-verbal quantifier, and according to the traditional assumptions, it should be in a position to c-command the pre-verbal arguments.

- (26) Herkes bu yıl kitaplarını ithaf etmiş üç kişiye.
 everybody this year book-plu-3sg-acc dedicate-past three person-dat

Göksel (1998) proposes an alternative explanation for Kural's findings concerning the hierarchical supremacy of the post-verbal position which are based on examples containing (accusative) marked direct objects, where these have wide scope with respect to the other elements. According to Göksel, marked direct objects are specific in Turkish and that can be considered to be the reason why such post-verbal elements appear to have wide scope, rather than the claim that the post-verbal position itself is higher in the tree. For example, for a sentence like (27), the bare direct object appearing in the post-verbal position *necessarily* has narrow scope with respect to *herkes*, this constituting a counterexample to the claim that the post-verbal position is necessarily superior to any other positions.

- (27) Herkes görMÜŞ mü üç film? Göksel (1998)
 everyone-saw int. three films
 'Has everyone seen three films?'

Kornfilt (1998) argues against both Kural (1994) and Kayne's LCA, first by declaring the claim that post-verbal quantifiers have in fact wide scope with respect to pre-verbal ones to be doubtful according to her intuitions. Quoting Kornfilt (1996),

Kural claims that in such examples, the only reading possible is where the post-verbal quantifier takes wide scope over the pre-verbal quantifier. However, many speakers, myself included, actually allow this interpretation only as a secondary reading, much preferring an interpretation whereby the post-verbal quantifier has narrower scope than the pre-verbal one.

Kornfilt later gives support for the claim that post-verbal arguments are at a higher position in the phrase structure tree with respect to pre-verbal ones, by basing her arguments not on quantifier scope but on the status of the so called 'Focus-particle constructions'. She also gives an alternative proposal inspired from the 'Discourse-based properties of *Devrik cümle*' (Kornfilt, 1996), according to which she seeks to establish a movement analysis of post-verbal constituents as due to rightward movement, counter to a base-generation analysis found in Kayne.

Zwart (2002) also argues against Kural, but unlike Kornfilt, he does so in favor of Kayne (1994). Zwart, also unlike Kornfilt, however, does not consider Kural's claims with respect to the wide scope of post-verbal quantifiers to be doubtful, but attributes the situation not due to a configurational claim regarding the position of post-verbal arguments to be higher than pre-verbal ones, but due to the findings of Erguvanlı (1979, p.71), according to which "background information (represented in the post-predicate elements in Turkish) is material that is 'supplementary' to the communication of a linguistic expression". That is, post-verbal arguments are thought to convey information that is predictable or recoverable from previous discourse, or is backgrounded material. Zwart states that this is fully adequate to explain the wide scope of post-verbal arguments to with respect to the pre-verbal ones, and there is no need for configurational claims. Thus, according to Zwart, Turkish is perfectly compatible with the Linear Correspondence Axiom. Zwart also notes that Kural was aware of Erguvanlı (1979), but nevertheless chosen not to capitalize on it. He continues with further arguments from Dutch in favor of Kayne (1994).

Kennelly (2003) proposes a generalized discourse based account of quantifier scope ambiguity in Turkish, by making use of the one-to-one correspondence between the linear order and the conventionally assumed default discourse structure of Turkish, which is assumed to be *topic focus predicate and background elements* in linear order, and argues against Göksel (1998), declaring that only discourse structure is relevant.

The first study which took both linearity and discourse properties to be relevant and attempted to give a formalized account of it is Göksel (1998). The relevance of linearity is due to the following kind of sentences, in which, scope relations seem to be determined by the surface order:

- | | | | | | | | |
|------|--|-------|-------|-------------|--------|--------|------------------------------|
| (28) | Her | çocuk | bir | öğretmene | çiçek | verdi. | ($\forall E / \exists A$) |
| | every | child | a | teacher-dat | flower | gave | |
| | 'Every child have flowers to a teacher.' | | | | | | |
| | | | | | | | |
| (29) | Bir | çocuk | her | öğretmene | çiçek | verdi. | ($\forall E / ?\exists A$) |
| | a | child | every | teacher-dat | flower | gave | |
| | 'A child gave flowers to every teacher.' | | | | | | |

In (28) where the universally quantified expression *her çocuk* 'every child' precedes the indefinite *bir öğretmen* 'a teacher', and both wide and narrow scope readings of the indefinite

are available. On the other hand, in (29), where the indefinite *bir çocuk* ‘a child’ precedes the universally quantified expression *her öğretmen* ‘every teacher’ the only possible reading that is considered to be available (at least according to Göksel (1998)) is the wide scope reading of the indefinite (but for an exception, see (30) below). The narrow scope reading which one would expect by analogy to (28)), is not considered to be available.

Based on such examples Göksel concludes that linearity (word order) plays a role in the scope of quantifiers in Turkish. The exact reason to which the effect of linearity can be attributed is dependent on the assumptions of each linguistic framework, and Göksel (1998) proposes a natural deductive scheme based on Labelled Deductive Systems (Gabbay, 1994; Gabbay and Kempson, 1991, 1992).

In addition to linearity, the discourse structure is also taken to be relevant due to the following kind of sentence with a focused NP that precedes the quantifier, in which the narrow scope reading of the indefinite is considered to be available, in spite of the fact that the indefinite precedes the quantifier.

- (30) Bir *hemşire* bakıyor her hastaya. ($\forall E / \exists E$)
 a nurse is seeing every patient-dat
 ‘A nurse is seeing every patient’.

3.2 Kennelly (2003)

Motivated by the availability of one-to-one correspondence between the traditionally assumed default discourse² structure of Turkish and the linear order given in Göksel (1998), Kennelly (2003) proposes an alternative solution, in which only the discourse structure is considered essential. As mentioned above, Park (1995) and Steedman (1999, 2000a, 2005) show that such accounts in general have difficulties in explaining some important syntactic constraints on quantifier scope. Nevertheless, we need to examine this particularly interesting theory, because it is possible to incorporate the same idea within a theory that respects the hierarchical structure and still implements the proposed idea, particularly because Kennelly (2003) does not seem to assume a particular discourse based account of quantifier scope³, but restricts

² More appropriately, that should be called *Information Structure*, but the distinction does not seem to be recognized by the traditional accounts.

³ In a footnote Kennelly notes Farkas (1997), but we haven’t seen a strong dependency that requires the indexical theory of Farkas, other than that a discourse constraint can be seen as a non-structural theoretical claim. Also, both Kennelly and Farkas seem to agree upon the idea that a movement analysis of quantification is inappropriate.

her attention to the topic/focus distinction which can be embedded in a strictly compositional ‘syntax-aware’ theory of quantifier scope.

The data which Kennelly (2003) is intended to account for is adapted from Göksel (1998), and is briefly summarized here so that the details of Kennelly’s examples do not incorrectly interfere those of Göksel:

- (31) a. Her hasta-yı bir doktor tedavi etti. (ambiguous: one or multiple doctors)
b. Her hasta-ya bir doktor eşlik etti. (ambiguous: one or multiple doctors)

At this point we need some terminology clarification as everyone seems to have his or her own favorite terms. Kennelly calls the dependent DP *bir doktor* the *multiple* DP and *her hasta* the *Base Plural* (BP). *DP1* and *DP2* are used to indicate the left-to-right linear order of the arguments. For the salient reading in (31) the variation or *multiplication* of DP2 is said to be dependent on the variation within the plural DP1 and that dependency is defined by treat/accompany relations. Then (31) is held to be true iff for every patient there is a doctor who treats/accompanies that patient. In the less salient reading without *quantificational dependencies* (QDs), (31) is held to be true iff there is one doctor who treats/accompanies all the patients.

In (32), where the order of the arguments is inverse of that of (31), with the universal quantifier *her* on DP2, DP1 does not have a multiple interpretation. In Kennelly’s terminology, the multiplication of DP1 is called *inverse QDs*.

- (32) a. (Genç) bir doktor her hasta-yı tedavi etti. (not ambiguous: one doctor)
b. $\exists y(\text{doctor}(y) \wedge \forall x[\text{patient}(x) \rightarrow \text{treat}(y,x)])$.

It is held that any variation in the neutral stress pattern in Turkish, where the pitch accent lies on the immediately preverbal element, does not alter the judgements. (32) is held to be true iff there is only one doctor who has treated all the patients.

Kennelly builds upon the idea that there is one-to-one mapping between discourse structure and linearity in Turkish (Erk , 1982, 1983).

- (33) Topic = Given Information > Focus = New Information > Predicate

Given (33), Kennelly considers that it is problematic to determine if discourse structure or linearity is the crucial issue in (32) such that data that can distinguish between the two are needed. If there are examples of inverse QDs with multiple DP1 in Turkish, then that

might be taken to indicate that it is not linearity but discourse structure that is essential to the characterization of quantifier scope. The idea here is that, following van der Sandt (1992), Topic / Given Information can be associated with anaphora and be marked [+anaphoric]. Assuming that all quantified DPs contain a variable that needs to be bound, then the variable contained in a quantified Given DP must be bound under the identity relation with a DP that has occurred previously in the discourse⁴. The DP that is not Given is New Information, therefore [-anaphoric], with is associated with Focus. Its variable does not have a binder in the previous discourse.

Using this idea, Kennelly argues that linearity is not the crucial factor in QDs, claiming instead that (32a) does not conform to the appropriate organization of discourse structure for QDs to obtain. Because the quantified DP1 in (32) is Topic/Given Information and consequently [+anaphoric] it necessarily has a fixed quantity for the duration of the Speech Act, the set of utterances that constitute the discourse.⁵

A further part of the empirical evidence that is claimed to be true is that DP1 does multiply under inverse QDs in Turkish only if it is supported as Focus/New Information both by context *and* by the lexical choice of the predicate.

Finally the proposal is that QDs are a mapping from Given to New Information where the BP is Given and the multiple DP is New Information. It is also assumed that the definition of QDs in terms of functional predicate results in a locality constraint on multiple DP.⁶

The idea that it is discourse structure rather than linearity that determines QDs in Turkish is supported with the following example, in which availability of narrow scope reading of the indefinite is claimed:

(34) Genç bir doktor her hasta-ya eşlik etti.

Kennelly reports that a significant percentage of Turkish speakers, including Göksel, do accept an inverse QDs reading for (34). To able to explain the difference between (32) and (34), the following dialogue is introduced:

⁴ However, the idea that the binder must be fixed and cannot possibly be dependent on another entity seems to be unjustified.

⁵ Here it is necessary to point out that Göksel (1998) and Kennelly (2003) have different assumptions regarding Turkish discourse structure. According to Göksel, there is no canonical focus position in Turkish and any of the preverbal positions can be focused, and the fact that most natural position for focus being the immediately preverbal one is seemingly attributed to the fact that it is obligatorily stressed.

⁶ *Locality*, or *predicate domain*, is a concept that can be easily incorporated into the syntax of a CCG based account with subcategorization of predicates, in spite of the fact that locality was originally a predication within LF. Non-locality, then, is the Text-Level or global or so called ‘top-level’ readings in DRT.

- (35) a. Genç bir HEMŞİRE her hasta-ya eşlik etti.
b. Hayır, (genç) bir DOKTOR her hastaya eşlik etti. (ambiguous: one or multiple doctors).

Then it is assumed that it is discourse structure rather than linearity that is crucial to the multiple DP1 reading in (35b). The following sentences are given as evidence, in the dialogues that follow.

- (36) Kızgın bir polis her kapı-yı çalıyordu.
(ambiguous - highly marked: multiple policemen; salient reading: one policeman)
- (37) Dürüst bir çiftçi her toplan-tı-ya gitti.
(ambiguous - highly marked: multiple farmers; salient reading: one farmer)
- (38) Hoş bir kız her masa-da konuklar-ı ağırladı.
(ambiguous - highly marked: multiple girls; salient reading: one girl.)

The dialogues:

- (39) a. Kızgın bir POSTACI her kapı-yı çalıyordu.
b. Hayır, (kızgın) bir POLİS her kapıyı çalıyordu.
(ambiguous: one or multiple policeman)
- (40) a. Dürüst bir ÖĞRETMEN her toplan-tı-ya gitti.
b. Hayır, (dürüst) bir ÇİFTÇİ her toplantıya gitti.
(ambiguous: one or multiple farmers)
- (41) a. Hoş bir ADAM her masa-da konuklar-ı ağırladı.
b. Hayır, (hoş) bir KIZ her masada konukları ağırladı.
(ambiguous: one or multiple girls)

Parallel to (35), the multiple reading for DP is available. The examples (36)-(38) are presented as showing that the distinction between (32) and (34) has nothing to do with Case marking⁷. In the following discourse, where DP1 is [+anaphoric] Given Information, the multiple reading is not available at all.

⁷ The capital *Case* is used in the sense of *abstract case*, a concept within the framework of Chomskian theorizing.

- (42) a. Kızgın bir polis ne yapıyordu?
 b. (Kızgın bir polis) HER KAPIYI çalışıyordu. (one policeman)
- (43) a. Dürüst bir çiftçi nere-ye gitti?
 b. (Dürüst bir çiftçi) HER TOPLANTIYA gitti. (one policeman)
- (44) a. Hoş bir kız ne yaptı?
 b. (Hoş bir kız) KONUKLARI ağırladı. (one policeman)

In [(35), (39) - (41)] vs [(42) - (44)]. Contrastive focus / New Information DP1 multiplies while [+anaphoric] Given Information DP1 does not. Thus it is hypothesized that:

Definition 1 (Discourse Structure Constraint on QDs) *Under Quantificational Dependencies the multiple DP1 is Focus/New Information.*

Kennelly (2003) continues with a formalization of this idea, a part of which is a stipulation of a locality constraint. Here we consider it sufficient to mention the constraint, without mentioning the details:

Definition 2 (Locality Constraint on QDs) *A multiple DP under QDs is confined to the same predicate domain that of the base plural (BP).*

3.2.1 Locality or Discourse Structure

A part of the important arguments in Kennelly (2003) concerns whether it is locality or discourse structure that is primitive in QDs. The following judgement is hers:

In Turkish it is very difficult to get an inverse QDs reading while in English it is far easier. The difference between the two languages in this respect is not fully understood. Turkish is considered to be a discourse configurational language, and it is generally assumed that linearity plays a major role in determining discourse roles in Turkish, as well as prosody. English, on the other hand, uses linearity to determine grammatical structure, which necessarily leaves more flexibility in the structuring of discourse, determined mainly by prosody. At the same time it must be acknowledged that English also prefers sentence initial Topic(Given Information). In terms of locality the two languages behave the same but with respect to inverse QDs, the more flexible English also has greater flexibility for discourse structure. This is the empirical fact, so it is discourse structure as a primitive for QDs that is the more informative constraint.

3.2.2 Accompany Type Predicates

An argument in Kennelly (2003) which we consider to be questionable is the one regarding the so called *accompany type predicates* which are assumed to carry an intrinsic reversal of discourse roles. This implies that argument structure also encodes a default discourse structure. The linear organization of discourse roles in Turkish, discounting background information and contrastive focus, patterns with the default SOV word order, as given in (33).

The default organization for discourse roles that might be built into the argument structure is that the subject is the default Topic(Given Information) and the object is the default Focus(New Information).

It is claimed that in Turkish, and to a lesser degree in English, the felicity of inverse QDs depends on the lexical selection of the predicate. In (45), inverse QDs are difficult, but not impossible to obtain.

- (45) a. A boy ate every pizza. (Steedman, 1999)
b. Few students read every book.

In contrast, in (46), the inverse QDs reading is readily available.

- (46) a. A computer sat on every desk.
b. A flag waved from every rooftop.
c. A teacher accompanied every student.
d. At least one farmer attended every meeting.

According to Kennelly's proposal, the fact that there is a lexical distinction between (45) and (46) argues against a covert LF movement or c-command analysis for QDs, even for English, and she explains the difference by claiming that accompany type predicate reverse the discourse structure.

While it is possible (and trivial) to model this behavior in CCG, we think that a simpler solution exists: It is logically possible that a single boy may eat every pizza, it is rather counter intuitive that a single computer may sit on every desk. Perhaps a very large computer would suffice to satisfy the truth conditions of the sentence without needing an argument about the reversal of discourse structure. For example, the physical size of a single ENIAC should be sufficient to satisfy the truth conditions of the sentence (46a), but that's not the way computers are built nowadays, so today's computers cannot sit on every desk. Therefore we consider the argument about accompany-type predicates to be an unreliable one.

3.3 Problems Caused by The Lack of Compositionality

The natural deductive scheme of Göksel (1998), employs skolem constants whose dependencies are freely chosen from the linearly preceding elements. It can allow for any expression, regardless of whether it is licenced by syntax or not (the only constraint is linearity) to be chosen as the entity to be dependent upon. This level of freedom in choosing skolem dependencies causes the same kind of problems that are associated with Quantifier Raising and other unnecessarily free implementations of the idea of movement, such as Cooper Storage. For example it is highly difficult for an account in which linearity and focus is the sole determining factor to explain the syntactic constraints upon quantifier scope as presented by Steedman (2005). The problem cannot be captured with a linearity relation: it needs to take into account the hierarchical structure determined by syntax, and therefore, needs a notion of compositionality that forbids operations such as movement (or brings a principled restriction to it) and similarly also forbids freely determinable dependents for skolem functions. The freedom in the way skolem dependencies are determined also implies that the method employed in Göksel (1998) is not compositional.

3.4 Bringing Syntax Back to the Issue of Quantifier Scope in Turkish

One possibility in which the strengths of a discourse based account can be combined with sentential level account is to claim that discourse related factors impose a filter upon the possibilities that are licensed by syntax. For example, the proposed *Locality Constraint* and *Discourse Structure Constraint* proposed by Kennelly (2003) have a very direct and straightforward implementation in CCG. All that is needed for a rather sketchy implementation is to alter the definition of *Skolem Specification Rule* of Steedman (2005) such that the *environment* set is used by the rule only if the noun phrase is marked [-anaphoric] in the sense of Kennelly (2003). If the noun phrase is marked [+anaphoric], the environment will be ignored, yielding a wide scope reading for the indefinite. The altered definition is given below:

Definition 3 (The specification of a skolem term (Modified Version)) *Skolem specification of a term t of the form $\text{skolem}'_n p$ in an environment E yields a generalized Skolem term $sk_{n,p}^E$ which applies a generalized Skolem functor $Sk_{n,p}$ (where n is the number unique to the noun phrase that gave rise to t , and p is a nominal property corresponding to the restrictor of that noun phrase), to the tuple E (defined as the environment of t at the time of specification),*

which constitutes the arguments of the generalized Skolem term. The environment in the definition is assumed to be empty regardless of its actual content if the noun phrase denoted by the skolem term is marked discourse anaphoric.

The locality constraint is imposed by the compositional nature of the account, and the default discourse structure can be lexically generated by subcategorizing the transitive verbs as follows⁸.

$$(47) \quad \text{okumak(read)} : (S \backslash NP_{nom,topic,+dlink}) \backslash NP_{acc,focus}$$

Of course, this says nothing about the case where the predicate itself is the focus, but we can safely assume that the *default* information structure is a case where the predicate isn't focused.

For the post verbal *backgrounded* arguments, Rightward Contraposition⁹ rule of Bozşahin (2002) can be modified such that:

$$(48) \quad \text{Rightward Contraposition}(> XP): \quad X : a \Rightarrow \quad S_{-t} \backslash (S \backslash X_{+dlink}) : \lambda f.f[a] \\ S_{-t} \backslash (S_{-t} \backslash X_{+dlink}) : \lambda f.f[a]$$

In other words, we can assume that any noun phrase which does not bear the *focus* feature is discourse anaphoric. For that reason the following alternative rule could also be proposed:

$$(49) \quad \text{okumak(read)} : (S \backslash NP_{nom,topic}) \backslash NP_{acc,focus,-dlink}$$

With this rule, we could state that any noun phrase which is not marked [-dlink] is [+dlink] and therefore discourse anaphoric.

To be able to implement this approach, we can assume that the determiner *bir* is lexically ambiguous between a [+dlink] and [-dlink] versions:

$$(50) \quad -a/an := \text{bir} - \mathbf{NP}_{-dlink} / \mathbf{N} : \quad \lambda x.sk_o'(x) \\ -a/an := \text{bir} - \mathbf{NP}_{+dlink} / \mathbf{N} : \quad \lambda x.sk_n x$$

⁸ Here we could name this feature *+anaphoric* or *+ana* (as Kennelly calls it) instead of *+dlink* but Steedman (1996) makes use of the syntactic feature *+ana* for lexically specified binding of reflexives and reciprocals, and the two should not be confused, therefore we have chosen to use *+dlink* instead.

⁹ Leftward Contraposition rule is not modified because it's not involved in backgrounding. The subscript '-t' in the formula stands for 'detopicalization'. These rules later have been modified by Dr.Cem Bozşahin as follows (quoted from the lecture notes):

$$\text{Leftward Contraposition } (< \mathbf{T}_x): \quad \mathbf{NP}:a \Rightarrow \quad \mathbf{S}/(\mathbf{S}/_{>}\mathbf{NP}_{+top}): \lambda f.f a \\ \text{Rightward Contraposition } (> \mathbf{T}_x): \quad \mathbf{NP}:a \Rightarrow \quad \mathbf{S}/(\mathbf{S}/_{<}\mathbf{NP}_{-top}): \lambda f.f a \\ ((< \mathbf{T}_x) \text{ is topicalization, and } (> \mathbf{T}_x) \text{ is detopicalization/backgrounding.})$$

In the formula above $sk_n x$ is a skolem constant, and $sko'(x)$ is an underspecified skolem term. As an example of how this works, consider the example in Figure 3.1. Feature checking with respect to [+dlink] ensures that the correct version of the determiner *bir* is used. As a contrast, the derivation in Figure 3.2 cannot continue because of the feature mismatch¹⁰.

<i>Her</i>	<i>doktor</i>	<i>bir hasta</i>	<i>yı</i>	<i>tedavi etti</i>
$(S / (S \setminus NP)) / N$: $\lambda p. \lambda q. \forall x (px \rightarrow qx) : \{x\} : \lambda x. doc'x$	$\overset{b}{\triangleleft} N$: $\lambda x. doc'x$	$\overset{n}{\triangleleft} N_{-dlink}$: $sko'pat'$	$\overset{c}{\triangleleft} N \setminus \overset{o}{\triangleleft} N$: $\lambda x. x$	$(\overset{v}{\triangleleft} S \setminus NP_{nom}) \setminus NP_{acc, -dlink}$: $\lambda x. \lambda y. treat'xy$
$\xrightarrow{>}$	$\xrightarrow{>}$	$\xrightarrow{<}$	$\xrightarrow{<}$	$\xrightarrow{<}$
$S / (S \setminus NP)$: $\lambda q. \forall x (doc'x \rightarrow qy) : \{x\}$		$\overset{c}{\triangleleft} N_{-dlink}$: $sko'pat'$		
		$\xrightarrow{\top}$		
		$(S \setminus NP) / ((S \setminus NP) \setminus NP_{acc, -dlink})$: $\lambda f. f[sko'pat']$		
		$\xrightarrow{<}$		
		$(S \setminus NP)$		
		$: \lambda y. treat'(sko'pat')y$		
		$\xrightarrow{>}$		
		$S : \forall x [doc'x \rightarrow treat'(sko'pat')x] : \{x\}$		

Figure 3.1: A derivation that succeeds

<i>Her</i>	<i>doktor</i>	<i>bir hasta</i>	<i>yı</i>	<i>tedavi etti</i>
$(S / (S \setminus NP)) / N$: $\lambda p. \lambda q. \forall x (px \rightarrow qx) : \{x\} : \lambda x. doc'x$	$\overset{b}{\triangleleft} N$: $\lambda x. doc'x$	$\overset{n}{\triangleleft} N_{+dlink}$: $sk'_n pat'$	$\overset{c}{\triangleleft} N \setminus \overset{o}{\triangleleft} N$: $\lambda x. x$	$(\overset{v}{\triangleleft} S \setminus NP_{nom}) \setminus NP_{acc, -dlink}$: $\lambda x. \lambda y. treat'xy$
$\xrightarrow{>}$	$\xrightarrow{>}$	$\xrightarrow{<}$	$\xrightarrow{<}$	$\xrightarrow{<}$
$S / (S \setminus NP)$: $\lambda q. \forall x (doc'x \rightarrow qy) : \{x\}$		$: sk'_n pat'$		
		$\xrightarrow{\top}$		
		$(S \setminus NP) / ((S \setminus NP) \setminus NP_{acc, +dlink})$: $\lambda f. f[sk'_n pat']$		
		$\xrightarrow{<}$		

Figure 3.2: A derivation that fails

These category assignments and the modified rightward contraposition rule directly implement both of the locality and discourse constraints of Kennelly (2003), and are much shorter and easier to understand than Kennelly's own formalization of the same idea. In the next section, we shall try to answer the question whether the same problem can be solved without making use of the discourse structure and discourse related concepts.

3.5 Is There a Solution Independent of Discourse Structure?

The account given in Steedman (2005) does not fully explain the data regarding the data presented in Göksel (1998) and Kennelly (2003) regarding scope ambiguities observed in Turkish, which, as far as the subject matter of this chapter is concerned, will be assumed to be correct. In this section, we shall show what happens if we compare the predictions of

¹⁰ The lexicon design that is implicitly assumed in these figures and the ones that follow is the morpho-syntactic account of Bozşahin (2002).

Steedman (2005) with the data examined so far.

The mismatches between the data and the predictions of the theory originates from the fact that, as a matter of principle, surface structure is not a level of representation in CCG. As a result sentences such as *Her doktor bir hasta-yı tedavi etti* and *Bir hasta-yı her doktor tedavi etti* which are variants of each other at the surface level are mapped to the same logical form in the predicate argument structure. It is important to notice that this fact originates from the principles of CCG and is not a coincidental result. The way a particular logical form is derived is not a representational level either. A derivation may in some sense look like a phrase structure tree, but no theoretical status is assigned to such derivations.

Depending on the time skolem specification rule is executed, one can arrive at two possible readings for the sentence *Her doktor bir hasta-yı tedavi etti*. If the skolem specification rule is executed before the indefinite enters the scope of the quantifier at the logical form, the wide scope reading of the indefinite is generated. If, on the other hand, the skolem specification rule is executed after the indefinite enters the scope of the quantifier, the narrow scope reading of the indefinite is generated. Figure 3.3 shows the case where a narrow scope reading is generated and Figure 3.4 shows the case of generating a narrow scope reading. In Figure 3.5, an incorrectly available narrow scope reading is derived. In Figure 3.6, the wide scope reading of the same sentence is correctly derived.

The problem here is that, unless we look for a discourse-based solution, we need to find a way to block the derivation in Figure 3.5 without breaking the other correct derivations. Since we cannot block that derivation altogether, the only possibility is to change the nature of the skolem specification rule, and to enforce its execution at certain times. That is, we may introduce a *skolem specification trigger* and associate it with a lexical category or combinatory rule, such that the skolem specification rule is executed. However, it turns out that this cannot be done without introducing side effects that block other correct derivations. For example, if we associate the skolem specification trigger with the type raising rule, the logical form in (3.3) would be incorrectly modified. Similarly we cannot associate the skolem specification rule with the accusative case marker *-yı* either, for exactly the same reason. What is needed is to be able to place a skolem specification trigger¹¹ right at the point of the star sign in the sentence (51) without affecting other example sentences.

(51) *Bir hasta-yı * her doktor tedavi etti.*

¹¹ “Skolem specification trigger” is an concept invented for this thesis and does not refer to a concept defined earlier in the literature.

For example, the same skolem specification trigger must not be available in the sentence (52) at the star position.

(52) Her doktor bir hasta-yı * tedavi etti.

The problem, then, is to find such a skolem specification trigger, and make sure that it is available in (51) but not in (52). The mentioned skolem specification trigger must be of the category defined below:

(53) skolem specification trigger := * – NP \ NP : $\lambda x.trigger'x$

The function *trigger'* is assumed to be a syntactic formula modifier which modifies¹² skolem terms and is executed immediately, without any delayed application.

If we assume that an intonation contour is associated with this trigger, then it might be possible to block the incorrectly available reading in Figure 3.5.

The problem is, even if we can find such an intonation contour, the solution then becomes curiously similar to the discourse based approach discussed in the earlier section, because in CCG it is assumed that it is intonation contours that determines the information structure. This, in turn, is nothing other than returning back to the solution in the section (3.4), and formulating the same idea in another way.

The category assignments for the quantifier *her* regarding the discussion so far can be seen in Table 3.1.

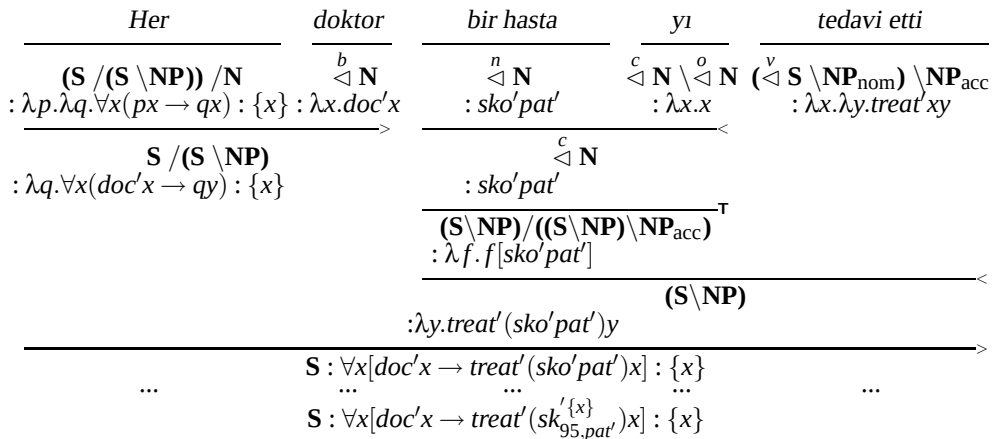
Meanwhile, there is yet another solution that can be considered which involves subcategorization of the quantifier *her*. In this approach, the quantifier *her* is assumed to be lexically ambiguous, and the idea is based on distinguishing the quantifiers occurring in DP1 position from those in the DP2 position, and postulating that only those in the DP1 position introduce variables to the environment set. While there is no theoretical motivation for doing that, it is clear that modelling the observed behavior is within the expressive power of CCG at the sentential level. Whether the solution is explanatorily adequate or not, is another concern. The table in Table 3.2 shows the possible lexical assignments, which is a modified version of the table in Table 3.1.

The argument in this section shows that no explanatorily adequate sentential level solution that can explain the data at hand is available. While the account of Kennelly (2003) can be

¹² Note that such modification operations can be considered to be a violation of the principle of direct compositionality. For example, see Jacobson (1999), in which the logical form is considered to be inaccessible to the process of derivation.

Table 3.1: Lexical assignments for the quantifier *her*

Name	Category Assignment	Sample Sentence
<i>her</i> ₁	$([S / (S \setminus NP_{nom})] / N_{nom})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her</i> adam uyu
<i>her</i> ₂	$([S_t / (S \setminus NP_{nom})] / N_{nom})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	uyu <i>her</i> adam
<i>her</i> ₃	$([S_t / (S_t \setminus NP_{acc})] / N_{acc})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	oku adam <i>her</i> kitap ı
<i>her</i> ₄	$([S_t / (S_t \setminus NP_{nom})] / N_{nom})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	oku kitap ı <i>her</i> adam
<i>her</i> ₅	$((S / (S \setminus NP_{acc})) / N_{acc})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her</i> kitap ı adam oku
<i>her</i> ₆	same as <i>her</i> ₁	<i>her</i> adam kitap ı oku
<i>her</i> ₇	$((S_t / (S_t \setminus NP_{acc})) / N_{acc})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her</i> kitap ı oku adam
<i>her</i> ₈	$([S_t / (S_t \setminus NP_{nom})] / N_{nom})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her</i> adam oku kitap ı
<i>her</i> ₉	$((S / ((S \setminus NP) \setminus NP_{acc})) \setminus N_{nom}) / N_{acc}$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w] : \{x\}$	adam <i>her</i> kitap ı oku
<i>her</i> ₁₀	$((S / ((S \setminus NP) \setminus NP_{nom})) \setminus N_{acc}) / N_{nom}$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w] : \{x\}$	kitap ı <i>her</i> adam oku
<i>her</i> ₁₁	$((S \setminus NP_{nom}) / (S \setminus NP_{acc})) / N_{nom}$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w] : \{x\}$	kitap ı <i>her</i> adam oku
<i>her</i> ₁₂	$([S_t / (S \setminus NP_{acc})] / N_{acc})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	adam oku <i>her</i> kitap ı
<i>her</i> ₁₃	$([S_t / (S \setminus NP_{nom})] / N_{nom})$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	kitap ı oku <i>her</i> adam
<i>her</i> ₁₄	$([(S \setminus NP_{nom}) / ((S \setminus NP_{acc}) \setminus NP_{nom})] / N_{acc})$ $: \lambda p. \lambda q. \lambda y. \forall c [pc \rightarrow (qy)c] : \{c\}$	<i>her</i> adam <i>her</i> kitap ı oku
<i>her</i> ₁₅	$([(S \setminus NP) / ((S \setminus NP) \setminus NP_{acc})] / N_{nom})$ $: \lambda p. \lambda q. \lambda y. \forall c [pc \rightarrow (qy)c] : \{c\}$	<i>her</i> kitap ı <i>her</i> adam oku


 Figure 3.3: A derivation for the narrow scope reading of the indefinite in *Her doktor bir hasta-yı tedavi etti*

[illegible]

Figure 3.4: A derivation for the wide scope reading of the indefinite in *Her doktor bir hasta-yı tedavi etti*

<i>Bir hasta</i>	<i>y1</i>	<i>her</i>	<i>doktor</i>	<i>tedavi etti</i>
$\overset{n}{\triangleleft} \mathbf{N}$: <i>sko'pat'</i>	$\overset{c}{\triangleleft} \mathbf{N} \setminus \overset{o}{\triangleleft} \mathbf{N}$: $\lambda x.x$	$((\mathbf{S} \setminus \mathbf{NP}) / ((\mathbf{S} \setminus \mathbf{NP}) \setminus \mathbf{NP}_{acc})) / \mathbf{N}_{nom}$: $\lambda p.\lambda q.\lambda y.\forall x[p x \rightarrow (q y)x] : \{x\}$	$\overset{b}{\triangleleft} \mathbf{N}$: $\lambda x.doc'x$	$\overset{v}{\triangleleft} \mathbf{S} \setminus \mathbf{NP}_{nom}) \setminus \mathbf{NP}_{acc}$: $\lambda x.\lambda y.treat'xy$
$\xleftarrow{\quad}$		$\xrightarrow{\quad}$		
$\overset{c}{\triangleleft} \mathbf{N}$: <i>sko'pat'</i>		$((\mathbf{S} \setminus \mathbf{NP}) / ((\mathbf{S} \setminus \mathbf{NP}) \setminus \mathbf{NP}_{acc}))$: $\lambda q.\lambda y.\forall x[doc'x \rightarrow (q y)x] : \{x\}$		
$\xrightarrow{\quad}$		$\xrightarrow{\quad}$		
$\mathbf{S} / (\mathbf{S} \setminus \mathbf{NP})$: $\lambda f.f[sko'pat']$		$(\mathbf{S} \setminus \mathbf{NP})$: $\lambda y.\forall x[doc'x \rightarrow treat'yx] : x$		
$\xrightarrow{\quad}$		$\xrightarrow{\quad}$		
...	...	...	...	...
		$\mathbf{S} : \forall x[doc'x \rightarrow treat'(sko'pat')x] : \{x\}$		
		$\mathbf{S} : \forall x[doc'x \rightarrow treat'(sk'_{95.pat'}^{\{x\}})x] : \{x\}$		

Figure 3.5: A derivation for the (incorrectly available without being focused) narrow scope reading of the indefinite in *Bir hasta-yı her doktor tedavi etti*

<i>Bir hasta</i>	<i>y1</i>	<i>her</i>	<i>doktor</i>	<i>tedavi etti</i>
$\overset{n}{\triangleleft} \mathbf{N}$ $: sko'pat'$	$\overset{c}{\triangleleft} \mathbf{N} \setminus \overset{o}{\triangleleft} \mathbf{N}$ $: \lambda x.x$	$((\mathbf{S} \setminus \mathbf{NP}) / ((\mathbf{S} \setminus \mathbf{NP}) \setminus \mathbf{NP}_{acc})) / \mathbf{N}_{nom}$ $: \lambda p. \lambda q. \lambda y. \forall x [px \rightarrow (qy)x] : \{x\}$	$\overset{b}{\triangleleft} \mathbf{N}$ $: \lambda x.doc'x$	$(\overset{v}{\triangleleft} \mathbf{S} \setminus \mathbf{NP}_{nom}) \setminus \mathbf{NP}_{acc}$ $: \lambda x. \lambda y. treat'xy$
$\overset{c}{\triangleleft} \mathbf{N}$ $: sko'pat'$ $\vdots$ $: sk'_{35.pat'}$	$\dots$	$((\mathbf{S} \setminus \mathbf{NP}) / ((\mathbf{S} \setminus \mathbf{NP}) \setminus \mathbf{NP}_{acc}))$ $: \lambda q. \lambda y. \forall x [doc'x \rightarrow (qy)x] : \{x\}$ $\dots$	$\dots$	$\dots$
$\mathbf{S} / (\mathbf{S} \setminus \mathbf{NP})$ $: \lambda f. f[sk'_{35.pat'}]$		$(\mathbf{S} \setminus \mathbf{NP})$ $: \lambda y. \forall x [doc'x \rightarrow treat'yx] : x$		
		$\mathbf{S} : \forall x [doc'x \rightarrow treat'(sk'_{35.pat'})x] : \{x\}$		

Figure 3.6: A derivation for the wide scope reading of the indefinite in *Bir hasta-yı her doktor tedavi etti*

Table 3.2: Lexical assignments for the quantifier *her* - Modified.

Name	Category Assignment	Sample Sentence
<i>her</i> ₁	$(([S / (S \setminus \mathbf{NP}_{\text{nom}})] / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her adam uyu</i>
<i>her</i> ₂	$(([S_{-t} \setminus (S \setminus \mathbf{NP}_{\text{nom}})] / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx]$	<i>uyu her adam</i>
<i>her</i> ₃	$(([S_{-t} \setminus (S_{-t} \setminus \mathbf{NP}_{\text{acc}})] / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx]$	<i>oku adam her kitap ı</i>
<i>her</i> ₄	$(([S_{-t} \setminus (S_{-t} \setminus \mathbf{NP}_{\text{nom}})] / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx]$	<i>oku kitap ı her adam</i>
<i>her</i> ₅	$(([S / (S \setminus \mathbf{NP}_{\text{acc}})] / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her kitap ı adam oku</i>
<i>her</i> ₆	same as <i>her</i> ₁	<i>her adam kitap ı oku</i>
<i>her</i> ₇	$(([S_{-t} / (S_{-t} \setminus \mathbf{NP}_{\text{acc}})] / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her kitap ı oku adam</i>
<i>her</i> ₈	$(([S_{-t} / (S_{-t} \setminus \mathbf{NP}_{\text{nom}})] / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\}$	<i>her adam oku kitap ı</i>
<i>her</i> ₉	$((([S / ((S \setminus \mathbf{NP}) \setminus \mathbf{NP}_{\text{acc}})] \setminus \mathbf{N}_{\text{nom}}) / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w]$	<i>adam her kitap ı oku</i>
<i>her</i> ₁₀	$((([S / ((S \setminus \mathbf{NP}) \setminus \mathbf{NP}_{\text{nom}})] \setminus \mathbf{N}_{\text{acc}}) / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w]$	<i>kitap ı her adam oku</i>
<i>her</i> ₁₁	$(([S \setminus \mathbf{NP}_{\text{nom}}] / (S \setminus \mathbf{NP}_{\text{acc}})) / \mathbf{N}_{\text{nom}})$ $: \lambda p. \lambda w. \lambda q. \forall x [px \rightarrow (qx)w]$	<i>kitap ı her adam oku</i>
<i>her</i> ₁₂	$(([S_{-t} \setminus (S \setminus \mathbf{NP}_{\text{acc}})] / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx]$	<i>adam oku her kitap ı</i>
<i>her</i> ₁₃	$(([S_{-t} \setminus (S \setminus \mathbf{NP}_{\text{nom}})] / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \forall x [px \rightarrow qx]$	<i>kitap ı oku her adam</i>
<i>her</i> ₁₄	$((([S \setminus \mathbf{NP}_{\text{nom}}] / ([S \setminus \mathbf{NP}_{\text{acc}}] \setminus \mathbf{NP}_{\text{nom}}]) / \mathbf{N}_{\text{acc}}))$ $: \lambda p. \lambda q. \lambda y. \forall c [pc \rightarrow (qy)c]$	<i>her adam her kitap ı oku</i>
<i>her</i> ₁₅	$((([S \setminus \mathbf{NP}] / ([S \setminus \mathbf{NP}] \setminus \mathbf{NP}_{\text{acc}}]) / \mathbf{N}_{\text{nom}}))$ $: \lambda p. \lambda q. \lambda y. \forall c [pc \rightarrow (qy)c]$	<i>her kitap ı her adam oku</i>

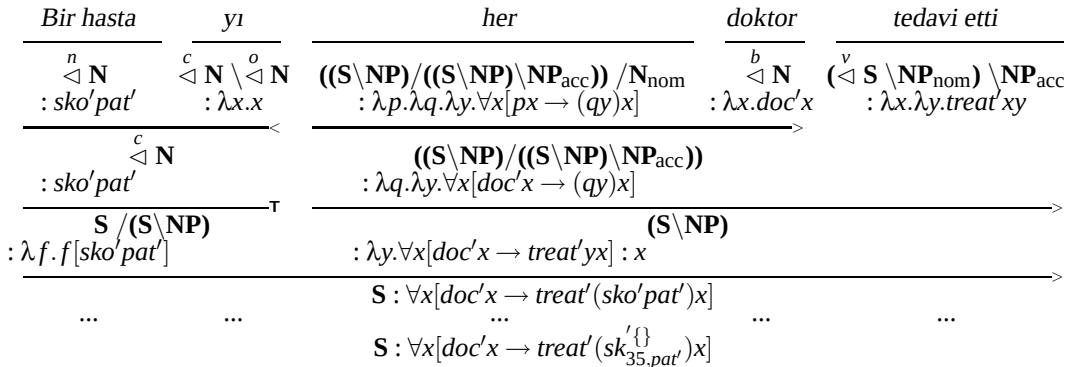


Figure 3.7: Lexically corrected version of the derivation of Figure 3.5

integrated into CCG, as in (47), this commits us to an assumption of default information structure, an assumption that is not supported by current theories of information structure in CCG (Steedman, 2000a, 2000b, 2003). Therefore, we need to have a look at what can be done with the standard assumptions in CCG, which is the subject matter of Chapter (4).

CHAPTER 4

AN APPLICATION OF STEEDMAN’S THEORIES OF INFORMATION STRUCTURE AND ALTERNATING QUANTIFIER SCOPE TO TURKISH

In the previous chapter we have made a literature survey of the studies of Quantifier Scope Ambiguity in Turkish, and argued that no explanatory sentential level solution that can explain the data at hand is available, and a discourse-based account seemed to be necessary. However, unlike Kennelly (2003), in CCG such an account can be provided without extra stipulation or special rules about quantifier scope, which is the subject matter of this chapter. Another conclusion of this chapter is that the proposed constraints on quantifier scope are actually incorrect.

4.1 Information Structure of Turkish

The traditional view in Turkish linguistics is that Turkish is a verb final language with free word order, and the variation of word order serves the purpose of assigning discourse related functions to certain parts of an utterance. For example, *focus* is considered to be associated with the immediately preverbal position, and the post-verbal positions are reserved for *background* material. Recently Özge and Bozşahin (2003, 2006) have proposed a tune based account of Turkish Information Structure in which word order variations are considered to be an epiphenomenal side effect caused by prosodic constraints, and the sole factor determining the information structure is considered to be intonation contours that the phrases of an utterance bears. Their paper argues that for an explanatory account of Turkish information structure, there is no need for any predications over sentence positions. A tune-based perspective where prosody is considered to be the sole structural determinant of information structure

is instead proposed. Also it is noted that, Turkish prosody imposes precedence constraints on certain intonational contours that are responsible for the realization of information structural units. Word order variation therefore, is considered epiphenomenal, rather than being a determinant in attaining the right information structure required by discourse context. In this thesis, we shall assume Özge and Bozşahin’s 2006 account of Turkish information structure, which is summarized in Appendix (E). Steedman’s update semantics for theme and rheme is also relevant for the subject matter, and that can be found in section (D.1) of Appendix (D).

4.2 Implications to Quantifier Scope Ambiguity in Turkish

In this section, we shall use Steedman’s theory of English Information Structure and Özge and Bozşahin’s 2006 account of Turkish Information Structure to give a negative answer to the proposal that quantificational dependencies are characterized by discourse binding of old information.

The idea of default discourse structure permits researchers to formulate ideas about discourse structure to examine the meaning of single utterances. First, we shall start by arguing that this practice is mistaken for the following reason: If there is no supplied context, any context can be accommodated¹. Second, we shall argue against the idea that contrastive focus and intonation contrast somehow affect quantificational dependencies. Third, we shall argue that contextual dependencies are more complex than what can be captured with the idea of discourse binding of old information. We shall start with our first argument.

4.2.1 The Argument Against Isolated Sentences

As we have stated above, it is a dangerous practice to think about the meaning of single utterances such as the following in isolation, because some contexts might be easier to accommodate than others².

(54) Bir hastayı her doktor tedavi etti.

For example, Niv (1994) proposes a “Psycholinguistically Motivated Parser” for CCG. One of the stated facts in his paper is that there is overwhelming evidence that ambiguity

¹ The relevance of context to the issue of quantifier scope ambiguity was pointed out by Dr. Cem Bozşahin and the idea of working on examples with contrastive focus is also his.

² Regarding ‘null context’, Crain and Steedman (1985) states the following: ‘The fact that the experimental situation in question makes a null contribution to the context does not mean that the context is null. It is merely not under the experimenter’s control. ... the so-called null context is simply an *unknown* context.’

is resolved within a word or two of the arrival of disambiguating information. While the resolution of syntactic ambiguity is not directly relevant, the point is that the human sentence processor is a bit *greedy*³ in nature, in the sense that it makes decisions as quickly as possible without exhaustive search. It follows that if there is no supplied context, the accommodated one is likely to be one that is decided using the information at hand, and decided as quickly as possible. That is, one might imagine a context in which there is a single patient much more easily than existence of multiple patients, each of which treated by different doctors, when faced by a question regarding possible scope ambiguities in (54), because at the time *bir hastayı* is heard, its sense is very likely to be determined before the word *her* is heard. The implication is that the presented data in the literature regarding single utterances such as (54) should be taken with a grain of salt. Therefore, as far as we are only working in a purely linguistic approach, we shall always examine utterances in some established context, and ignore the past literature data if it is presented about sentences in isolation, because we think that such data needs a psycholinguistic explanation. In particular, the claimed difference between (55a) and (55b) in the case of ‘null context’ is simply a misunderstanding due to incorrect interpretation of the data. What’s required is not a linguistic explanation, but a psycholinguistic one.

- (55) a. Bir hastayı her doktor tedavi etti. ($\exists\forall$)
 b. Her doktor bir hastayı tedavi etti. ($\forall\exists$) or ($\exists\forall$)

For these reasons, we think that the natural-deductive scheme in Göksel (1998), which has been designed to explain such differences regarding linear order is unnecessary. Similarly we also think that the parallelism between the linear order and so called default information structure (such as topic-focus-predicate), as assumed by Kennelly (2003) is an argument which needs to be re-evaluated by taking into account two issues: First, we reject the idea of a default information structure, and other word order based accounts of information structure, and assume the tune based account of information structure as presented by Özge and Bozşahin (2006). When we are discussing information structure, we shall also consider intonation marking, since the sentences bearing no intonation marking, if possible at all, are information structurally ambiguous. Second, we reject the idea that there are special constraints about quantifier scope: Thus, the locality constraint and the discourse constraint in

³ In the sense of *greedy algorithms*, which work and attempt to find a solution without making an exhaustive search.

Kennelly (2003) are superfluous, we do not think that we need special constraints to explain scope ambiguities in Turkish, because the presented data seems to be highly unreliable for reasons stated above. Given unreliability of such data, we consider the predictions of Steedman's relevant theories (which are independently verified for English) to be more reliable than rather weak intuitions regarding the meaning of such sentences. So, the subject matter of this section can be literally stated as follows: "What can (and/or cannot) possibly be defended about the scope ambiguities in Turkish, given the assumption that Steedman (2000a, 2000b, 2005) and Özge and Bozşahin (2006) are correct?"

Similarly we also think that there is no need for a theory like Kayne's Linear Correspondence Axiom to be able to explain anything about post-verbal constituents, if there is anything to be explained to begin with (to see why, see Kornfilt (1996), who simply rejects the validity of Kural's intuitions, casting yet another doubt about the validity of the presented data.). The proposed prosodic constraint in Özge and Bozşahin (2006) which states that H* LL% causes pitch flooring to its right has an immediate consequence of implying there can be no post-verbal kontrast. For example, Kural (1994) proposes that the post-verbal argument *üç kişiye* in (56) has wide scope with respect to the pre-verbal quantifier, and according to the traditional assumptions, it should be in a position to c-command the pre-verbal arguments. Therefore, Kural concludes that Kayne's Linear Correspondence argument must be false.

- (56) Herkes bu yıl kitaplarını ithaf etmiş üç kişiye.
 everybody this year book-plu-3sg-acc dedicate-past three person-dat

In present terms, there are two things that can be said about this claim. First, this sentence exemplifies the use of 'null context' data in a linguistic study. As we have stated above, we wouldn't consider such sentences in isolation and generalize the results into theories about the competence grammar (it's quite possible to establish a context in which Kural's claims are proven incorrect - see section (4.2.3) for example.). Second, since *üç kişiye* is a post-verbal argument, it cannot have a pitch accent due to the obligatory existence of LL% boundary on *etmiş*; therefore it follows that *üç kişiye* cannot be contrasted. For that reason (56) would be infelicitous in a discourse like (57).

(57) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.
- b. #Hayır, herkes bu yıl kitaplarını ithaf etti üç kişiye.

For (57a), either wide or narrow scope interpretation of the indefinite *iki kişiye* is appropriate (theoretical sentence-level semantics of the sentence, as described in the previous chapter, allows both), and the sense of the word *herkes* might possibly be restricted to those who have written a book; but (57b) is infelicitous. Of course, we are not claiming that infelicitous sentences cannot be interpreted. Rather, the fact that the sentence is infelicitous is the only conclusion that can be drawn.

4.2.2 The Argument Against Intonation

The sentence (56), if wasn't infelicitous, would exemplify what is sometimes called a *contrastive focus*. For example, in (58b), both narrow and wide scope interpretations of the indefinite are available ⁴.

(58) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.
- b. Hayır, herkes bu yıl kitaplarını ÜÇ kişiye ithaf etti.

H* LL%

If we had changed the word order in (58), as in (59), there would be no difference, and the narrow scope interpretation of the indefinite is still available.

(59) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.
- b. Hayır, ÜÇ kişiye ithaf etti herkes bu yıl kitaplarını.

H* LL% < – F – >

What should be noted here is that, if the same sentence was uttered with an intonation marker on *kadına* instead of an intonation marker on *üç*, it would be infelicitous according to Steedman's theory of information structure. Thus, (60b) is infelicitous.

(60) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.

⁴ In this thesis we shall be using Pierrehumbert's notation for intonation marking, information about which can be found in Pierrehumbert (1980) and Özge and Bozşahin (2003).

b. #Hayır, herkes bu yıl kitaplarını üç KADINA ithaf etti.

H* LL%

To be able to make the utterance (60b) felicitous, one either needs to move the intonation marker to *üç* as in (61b) or change the word *üç* to *iki*, as in (62b).

(61) Situation: There are several people in a room, some of which have written a book.

a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.

b. Hayır, herkes bu yıl kitaplarını ÜÇ KADINA ithaf etti.

H* !H* LL%

(62) Situation: There are several people in a room, some of which have written a book.

a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.

b. Hayır, herkes bu yıl kitaplarını iki KADINA ithaf etti.

H* LL%

In present terms, it follows that prosody of Turkish may or may not be able to distinguish between the contrast *ÜÇ KADINA* from *ÜÇ kadına*, and the contour H* LL% may also be ambiguous in this respect, in addition to the theme/rheme ambiguity regarding H* LL% as stated in Özge and Bozşahin (2006).

In these examples ((58), (59), (60), (61), (62)), for every case the utterance in (a) can be considered to be composed of a theme corresponding to what we shall informally call *everybody dedicating his/her books* and a rheme *to two people*. This theme establishes a rheme alternative set which can be represented by (63), and the rheme *to two people* restricts this rheme alternative to one that can be represented by (64).

(63) $\lambda y. \diamond \forall x. \text{dedicate}'y \text{ of}'(\text{sko}'\text{books}'x)x.$

As a result, in these examples the utterance (a) establishes an underspecified theme, represented by the formula (64).

(64) $\diamond \forall x. \text{dedicate}'\text{sko}'(\text{people}' ; \lambda n. |n| = 2) \text{ of}'(\text{sko}'\text{books}'x)x.$

The rheme alternative set represented by (64) is as follows.

(65)

$$\left\{ \begin{array}{l} \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^{\{x\}} of'(sk'_{\text{books}'x}^{\{x\}})x \\ \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^{\{x\}} of'(sk'_{\text{books}'x}^0) x \\ \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^0 of'(sk'_{\text{books}'x}^{\{x\}})x \\ \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^0 of'(sk'_{\text{books}'x}^0) x \end{array} \right\}$$

This set is the available permutations of possible skolem specifications. Some of the set members are logically incoherent because the functor of' implies a dependency that a skolem constant such as $sk'_{\text{books}'}^0$ can not satisfy. For that reason, the rheme alternative set in (65) can be reduced to the one in (66). This intricacy is not directly relevant to our present concerns, we will return to it later in section (4.5). For the present argument, let's assume that the rheme alternative set is the one in (66).

(66)

$$\left\{ \begin{array}{l} \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^{\{x\}} of'(sk'_{\text{books}'x}^{\{x\}})x \\ \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n|=2}^0 of'(sk'_{\text{books}'x}^{\{x\}})x \end{array} \right\}$$

Now consider the discourse (61). The answer (61b), which is uttered at a time when the discourse database contains (66), introduces a compatible theme what can be informally thought as *everybody dedicating his/her books*. This theme can be represented by (67).

(67) $\lambda y. \diamond \forall x. \text{dedicate}' y of'(sko' \text{books}' x)x.$

The rheme 'ÜÇ KADINA' in (61b), restricts the rheme alternative set represented by (67) to the one represented by (68).

(68) $\diamond \forall x. \text{dedicate}' * sko'(\text{women}'; \lambda n. |n|=3) of'(sko' \text{books}' x)x.$

(69)

$$\left\{ \begin{array}{l} \diamond \forall x. \text{dedicate}' sk'_{\text{women}'; \lambda n. |n|=3}^{\{x\}} of'(sk'_{\text{books}'x}^{\{x\}})x \\ \diamond \forall x. \text{dedicate}' sk'_{\text{women}'; \lambda n. |n|=3}^0 of'(sk'_{\text{books}'x}^{\{x\}})x \end{array} \right\}$$

According to Steedman's update semantics, (64) is retracted from the discourse database, and (68) is asserted. For that reason, the new rheme alternative set is (69), which is represented by (68). As a consequence we predict that both wide and narrow scope interpretations of the indefinite 'ÜÇ KADINA' in (61b) are available. The other indefinite ('kitaplarını') is restricted to narrow scope due to the semantics of the predicate of' .

If we had considered other example discourses, the same scope relations would be predicted for the other felicitous replies, including (59b), in which the indefinite appears before the quantifying NP *herkes* in the surface order. This would be the case even if the scope relation in the question was restricted to wide scope as in (70). Note that the wide scope interpretation restriction is due to the fact that *Mehmet* is a proper name, and therefore is not a skolem term. This is a case in which it is possible to contrast scope relations. Notice that this is possible because *Mehmet* is a proper name, not an indefinite.

(70) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını Mehmet’e ithaf etti.
- b. Hayır, herkes bu yıl kitaplarını ÜÇ KADINA ithaf etti.

H* !H* LL%

This happens because the new theme completely replaces the old one, together with its scope relations, and sentence level semantics for (70) predicts both narrow and wide scope interpretations of the indefinite ‘ÜÇ KADINA’. The theory similarly predicts that (71b) is also ambiguous, thus both narrow and wide scope interpretations of the indefinite ‘ÜÇ kişiye’ are available. This happens because, Steedman’s Surface Compositional Scope Alternation theory does not make distinctions regarding linear order, and according to the information structure of Turkish, (71b) is a felicitous answer to (71a).

(71) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını Mehmet’e ithaf etti.
- b. Hayır, ÜÇ KİŞİYE ithaf etti herkes bu yıl kitaplarını.

H* LL% < – F – >

The theme established by (71a) is (72), and the part that enforces a wide scope interpretation (that is *Mehmet*) can be fully overwritten by the kontrast *ÜÇ KİŞİYE*. This happens because (71b) introduces a compatible theme which can be represented by (73), which when restricted by the rheme in (74) yields (75). The rheme alternative set corresponding to (75) is (76). Note that in (74) and (75), the contrastive diacritic ‘*’ applies to the whole skolem term *sko’(people’;λn. |n| = 3)*.

(72) $\diamond \forall x. \text{dedicate}' \text{ mehmet}' \text{ of}' (\text{sko}' \text{ books}' x) x.$

(73) $\lambda y. \diamond \forall x. \text{dedicate}' y \text{ of}' (\text{sko}' \text{books}' x) x.$

(74) $*(\text{sko}'(\text{people}'; \lambda n. |n| = 3)).$

(75) $\diamond \forall x. \text{dedicate}' * \text{sko}'(\text{people}'; \lambda n. |n| = 3) \text{of}' (\text{sko}' \text{books}' x) x.$

(76)

$$\left\{ \begin{array}{l} \diamond \forall x. \text{dedicate}' sk'_{\text{people}'; \lambda n. |n| = 3}^{\{x\}} \text{of}' (sk'_{\text{books}' x}^{\{x\}}) x \\ \diamond \forall x. \text{dedicate}' sk'^0_{\text{people}'; \lambda n. |n| = 3} \text{of}' (sk'_{\text{books}' x}^{\{x\}}) x \end{array} \right\}$$

As seen in these examples, when a proper context is supplied, intonation has no affect on quantificational dependencies. The reason is that according to Steedman's theory of Information Structure, contrast applies to lexical items. When the lexical item kontrasted is underspecified, intonation kontrast has no disambiguation capability.

4.2.3 The Argument Against Discourse Binding of 'Old Information' to a Single Entity

Kennelly (2003) states that under quantificational dependencies, the multiple DP1 is Focus/New Information. Compatible with that claim, so far the examples we have examined had the 'multiple DP' as a rheme. Implicit in her argument is the claim that if DP1 is not 'focused', it would have wide scope with respect to the quantifier. We think that this claim is true only if the previous discourse restricts the reference of DP1 to a single entity. Otherwise, narrow scope reading for the indefinite would be available even if DP1 is not a rheme, but a part of the theme. The reason is that, the rheme alternative set would contain both the wide/narrow scope readings. For example, consider the following dialogue:

(77) Situation: There are several books, notebooks and pens on a table.

a. İki kalemle yazıldı her kitap değil mi?

b. Hayır, iki kalemle her DEFTER yazıldı.

H*

LL%

In (77b) narrow scope reading of the indefinite *iki kalemle* is available, even if it is in the DP1 position, and is not a rheme or theme kontrast either. The reason is that (77a) establishes a theme which yields a rheme alternative set that contains both the wide and narrow scope readings of the indefinite *iki kalemle*. (77b) replaces this theme by a new one without contrasting *iki kalemle*. Moreover, the NP *iki kalemle* is a part of the theme in the information

structure of (77b). In Kennelly's terms, it is not possible to say that it is focused; but still, it has a narrow scope reading.

The following example illustrates that it is possible to construct examples that cannot be explained by Kennelly's idea of discourse binding of old information.

(78) a. Bu gruptaki herkes bir elma aldı.

b. Aynı şekilde, [diğer gruptaki herkes]_{focus} de aldı [bir elma]_{backgrounded}.

According to the theory of discourse structure that Kennelly assumes, this is the discourse roles assigned. What's wrong is that not everyone can possibly buy the same apple. This violates the constraint that 'old information' needs to be bound to a single entity.

On the other hand, this example gives a clue about the nature of the kind of contextual dependencies involved: What's needed is a resolution of several constraints, some of which might be about real world knowledge. This idea recasts the issue as an instance of a constraint satisfaction problem: A combination of linguistic constraints (i.e. what grammar has to say), possibly binding constraints (if any that applies), logical constraints, and constraints imposed by real world knowledge is the minimum of what's required to be able to understand an utterance, and quantificational sentences are not an exception.

What's wrong with Kennelly's discourse constraint is that, indefinites are more complex than pronouns. Because indefinites are underspecified skolem terms, in a properly established context they can refer to more than one entity, because while a skolem constant refers to a single entity, a skolem function may refer to a *set* of entities, each member of which involves a quantificational dependency to the restrictor of the quantifier *her*. The discourse binding constraint can be violated, because the truth conditions of the sentence does not require a single referent for the indefinite *bir elma*. The conclusion is that results obtained from limited data does not always generalize to universal linguistic constraints.

Kennelly's constraints create further problems, which are the subject matter of the next section.

4.3 Focus and Background in Kennelly (2003)

In Kennelly (2003), the referential nature of non-focused (i.e. background) elements are restricted to be able to propose an alternative explanation for the data in Göksel (1998). Quoting Kennelly (2003),

Only New/Focus information DPs that are locally bound can multiply while Given information DPs with text level binding have a fixed quantity for the duration of the Speech Act.

It is clear that Kennelly’s *text level bound* elements correspond to Steedman’s *background*. For that reason, in present terms this claim means that backgrounded elements, when they include an underspecified skolem term within the scope of a universal quantifier cannot have a narrow scope with respect to the universal quantifier and therefore their interpretation is restricted to wide scope, which simply isn’t true.

To be able to implement this incorrect idea, Kennelly (2003) proposes two special constraints on quantifier scope:

(79) Discourse Structure Constraint on QDs: Under Quantificational Dependencies the multiple DP1 is Focus/New Information.

(80) Locality Constraint on QDs: A multiple DP under QDs is confined to the same predicate domain that of the BP.

These two constraints are not only proposed to be able to implement an incorrect idea, but as Kennelly herself was also aware, the locality constraint creates a problem with relative clauses of the following kind, which are analyzed under the title ‘Intermediate Readings’ in Steedman (2005).

(81) a. Every farmer who owns a donkey that she likes feeds it.

b. $\forall x[(farmer'x \wedge own'sk_{\lambda y.(donkey'y \wedge like'y(pro'x))}^{(x)})x] \rightarrow feed'(pro'sk_{\lambda y.(donkey'y \wedge like'y(pro'x))}^{(x)})x]$

Since the very reason that these constraints are proposed is misguided by Göksel’s misinterpretation of sentences such as (30), the problems that are associated with them are easy to solve. Simply pointing out that the notion of focus was incorrectly related to the issue of quantifier scope is adequate. Such sentences have wide or narrow scope interpretations depending on the context, and the fact that the indefinite may also be focused (*kontrasted* in Steedman’s terms) is irrelevant.

4.4 An Argument from Incremental Processing of Subsequent Utterances

So far we have presented arguments about what does *not* characterize the nature of quantificational dependencies, since the idea that all of the sentential level readings should be available

in a properly established context is a rather weak conclusion and is not saying much about the nature of quantification. Since anyone would naturally wonder if a stronger conclusion is possible or not, the subject matter of this section is investigate this possibility, and to determine whether the interpretation of two subsequent utterances can be affected by the contents of theme/rheme alternative sets, because it is known that such sentences usually share the same theme in a proper discourse. We shall start with an example, in which the second utterance in the discourse does not have its own theme.

- (82) a. Her hastayı kim tedavi etti?
b. Bir doktor.

In a discourse like (82), it is quite clear that the logical form of (82a) determines what is being understood, and (82b) simply supplies the missing information.

The theme established by (82) is $\lambda y.\forall x.[patient'x \rightarrow treat'xy]$. The rheme alternative set established by this theme can be enumerated as follows:

$$(83) \quad \left\{ \begin{array}{l} \diamond ahmet' \\ \diamond mehmet' \\ \diamond skolem' doctor' \\ \diamond skolem' nurse' \\ \dots \end{array} \right\}$$

Which could yield a new theme as in (84).

$$(84) \quad \forall x.[patient'x \rightarrow treat'x(skolem' doctor')^{\{x\}}]$$

At this point, seems like it is possible to say that the only option is to execute the skolem specification rule, and get the reading in (85), which is a narrow scope reading of the indefinite *bir doktor*.

$$(85) \quad \forall x.[patient'x \rightarrow treat'x(sk_{doctor'}^{\{x\}})^{\{x\}}]$$

As long as there doesn't exist a strong contextual factor that gives a strong influence against this kind of analysis, it can be said that quantificational dependencies are affected by the theme shared by subsequent utterances. The result can also be generalized to discourses like (86), in which the themes of (86a) and (86b) are the same, and for that reason (86b) can be interpreted as having the same effect as (82b).

- (86) a. Her hastayı kim tedavi etti?
 b. Bir doktor her hastayı tedavi etti.

This kind of reasoning associates quantificational dependencies across sentences via themes (i.e. via rheme alternative sets).

4.5 Possessive Constructions and Quantifiers

The following sentence, which is adopted from Bozşahin (2002) exemplifies the scope relation between the possessive marker *-ı* (of) and the quantifier *her* (every).

- (87) Her çalışan-in bazı hak-lar-ı vardır
 every worker-gen3 some right-plu-poss3s exists

This sentence can be parsed with the following category assignments.

- (88) her $:= ((S \setminus NP) / (\overset{o}{\bowtie} N \setminus \overset{o}{\bowtie} N)) / (\overset{o}{\bowtie} N \setminus \overset{o}{\bowtie} N) \quad : \lambda w. \lambda p. \lambda f. \forall x (wx \rightarrow f(px))$
 (89) in(-gen) $:= (\overset{o}{\bowtie} N \setminus \overset{o}{\bowtie} N) \setminus \overset{n}{\triangleleft} N \quad : \lambda x. x$
 (90) ı(-poss) $:= (\overset{o}{\bowtie} N \setminus \overset{o}{\bowtie} N) \setminus \overset{n}{\triangleleft} N \quad : \lambda x. \lambda y. of'yx$
 (91) bazı $:= \overset{b}{\triangleleft} N \setminus \overset{b}{\triangleleft} N \quad : \lambda x. sko'x$
 (92) çalışan $:= \overset{b}{\triangleleft} N \quad : \lambda x. worker'x$
 (93) hak $:= \overset{b}{\triangleleft} N \quad : right'$

A derivation for (87) can be seen in Figure 4.1. Note that in Figure 4.1 the lattice diacritics have been suppressed to save page width space.

The underspecified skolem terms in the resulting formula in Figure 4.1 ($\forall x (worker'x \rightarrow exists'(of'(sko'rights)x)$, that is) can be specified to yield the following results.

- (94) a. $\forall x (worker'x \rightarrow exists'(of'(sk_{rights}^{\{x\}})x))$
 b. $\forall x (worker'x \rightarrow exists'(of'(sk_{rights}^{\emptyset})x))$

Availability of (94b) seems to present a problem. *Rights* can be considered to be a property that can be owned by a certain individual, or the semantics of the predicate *of'* can be considered to incompatible with shared ownership. Following Walid Saba (1995), we think

that (94b) can be eliminated by a cognitive process similar to lexical disambiguation. Similar to the task of lexical disambiguation, this is an inference problem and (94b) is eliminated by an extra-linguistic cognitive process. In section (4.8), the sentence (106) exemplifies a similar disambiguation task.

The syntactic category in (88) seems to imply that such possessive constructions involve higher order (second order, to be specific) quantification. This implies quantification over properties, instead of entities. While this is compatible with the fact that category in (92) is a property, the situation is not unique to quantification with possessives. For example, a normal quantificational sentence such as the following also involves a noun (in this case, *man*) which is used as a property function; but in this case the functor *her* (every) combines with a syntactic category of type N , not $N \setminus N$. But, since in both cases we use the notation of first order logic which does not have higher order quantification, the notation can be misleading.

(95) Her adam uyudu.

$\forall x(man'x \rightarrow slept'x),$

but not $\forall x((x = man') \rightarrow slept'x)$

Note that $N \setminus N$ type in the argument of the category assigned to quantifier *her* can be justified with the following example from English, which exemplifies a problem that arise even an argument type of N is used. In Figure 4.2, the derivation cannot continue because *of'owner'* (*sko'hashbar'*) has no free variables left over which *every* could quantify.

However, with an argument type of $N \setminus N$, the sentence can be parsed successfully.⁵

⁵ Note that in Turkish just about every noun is also a property (Dr.Cem Bozşahin, personal communication).

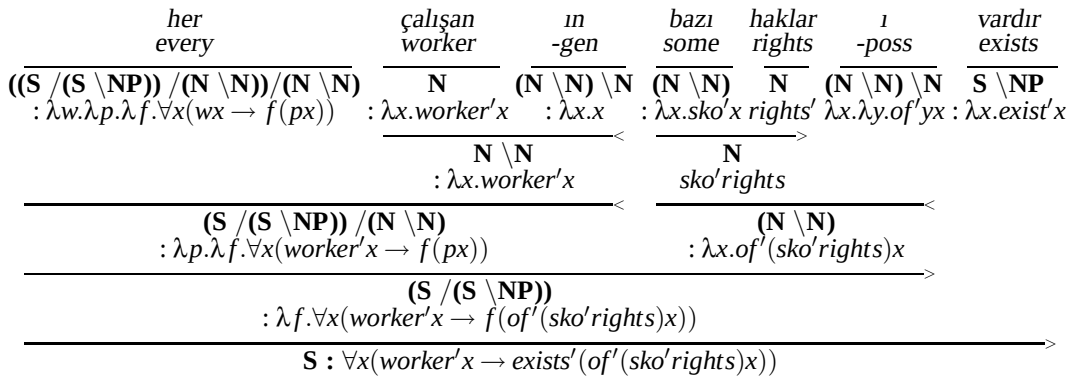


Figure 4.1: A derivation for the sentence (87)

4.6 Prepositional Phrases and Quantifiers

The sentence (16), which is repeated here exemplifies a potential problem:

(96) Mia knows every owner of a hash bar.

The theory in Steedman (2005) predicts a narrow scope reading of *a hash bar* with respect to the quantifying NP *every owner*. The problem can be seen in Figure 4.3. Again, we may propose that the narrow scope reading is actually available, but is eliminated with an inference task.

The case is not unique to possessive constructions. For example consider (97), which can be parsed as in Figure 4.4. Again, an otherwise unjustified narrow scope reading of the indefinite (in this case *a garden*) with respect to the quantifying NP (*every flower*) is predicted.

(97) Mia knows every flower in a garden.

The semantics of such prepositional phrases are beyond the scope of this thesis. Interested readers are referred to Francez and Steedman (2006), in which a detailed study of them can be found.

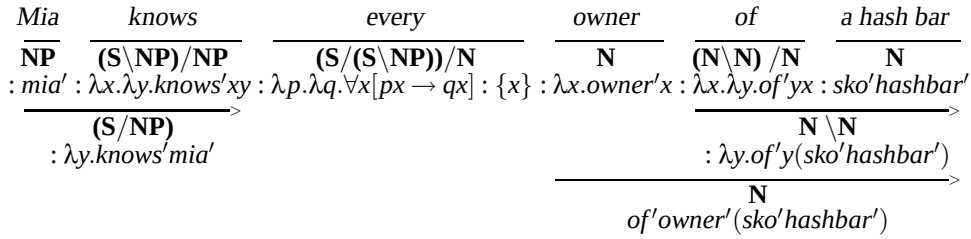


Figure 4.2: A derivation for the sentence *Mia knows every owner of a hash bar*

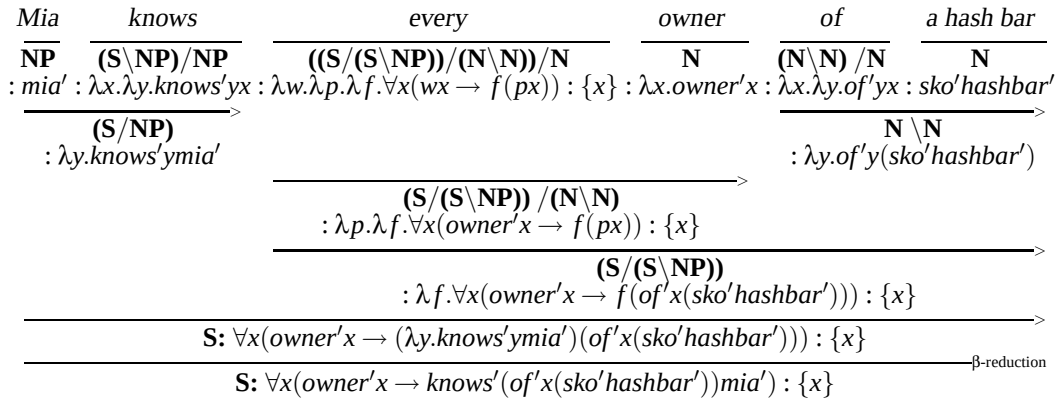


Figure 4.3: A derivation for the sentence *Mia knows every owner of a hash bar*

4.7 Negation Operator and Quantifiers

While Steedman (2005) does not have a concluding statement about the interaction of negation with scope ambiguity, it seems like there is a need to add it to the list of phenomena that needs to be explained, as in (98).

(98) Every boxer doesn't love a woman. $\neg\forall x/\neg\exists y/\neg\forall x/\neg\exists y/\neg\forall x/\neg\exists y$

(99) Every boxer doesn't love a woman. $\neg\forall x/\neg\exists y/\neg\forall x/\neg\exists y/\neg\forall x/\neg\exists y$

While my personal intuitions prefer (99), native speakers of English do not seem to agree. For example, according to Carpenter (1997, p.244), it is possible to say that such sentences have four readings (Carpenter says that they have *two* readings, without considering the scope of the quantifier with respect to the indefinite, but it is implicit that when the scope of the negation operator with respect to the quantifier and the scope of the indefinite with respect to the quantifier *every* is considered together, we can arrive the conclusion that there are *four* readings.).

Two of these readings is due to the wide scope reading of the negation operator with respect to the quantifier *every* both of which can be paraphrased as “it is not true that every boxer loves a woman”, (that is “Some boxers love a woman and some boxers do not love a woman’), as shown in (100).

(100) $\neg(\forall x(\text{boxer}'x \rightarrow \text{loves}'(\text{sko}'\text{woman}')x))$

The other two readings, as shown in (101) does not entail form (100), because the truth conditions of (101) is such that none of the boxers love a woman.

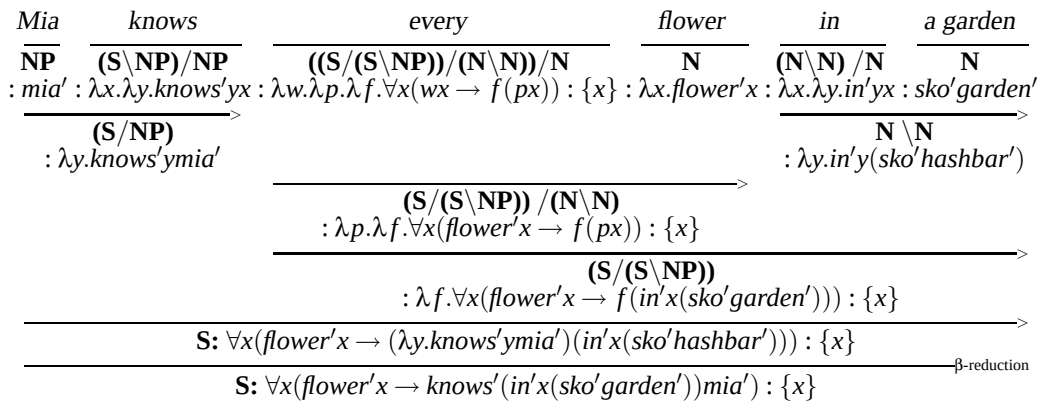


Figure 4.4: A derivation for the sentence *Mia knows every flower in a garden*

$$(101) \quad \forall x(\text{boxer}'x \rightarrow \neg \text{loves}'(\text{sko}'\text{woman}')x)$$

The problem can be solved by assigning the negation operator a value-raising category in (102), in addition to the usual lexical assignment shown in (103):

$$(102) \quad \text{doesn't} \quad := ((S \setminus (S / (S \setminus NP))) / (S \setminus NP)) \quad : \lambda g. \lambda f. \neg(fg)$$

$$(103) \quad \text{doesn't} \quad := (S \setminus NP) / (S \setminus NP) \quad : \lambda f. \neg f$$

A derivation for the sentence (98) that makes use of the value-raising category in (102) can be seen in Figure 4.5.

$$\begin{array}{c}
\begin{array}{c} \text{Every} \\ \hline (S / (S \setminus NP)) / N \\ \hline : \lambda p. \lambda q. \forall x [px \rightarrow qx] : \{x\} \end{array} \quad \begin{array}{c} \text{boxer} \\ \hline N \\ \hline : \lambda x. \text{boxer}'x \end{array} \quad \begin{array}{c} \text{doesn't} \\ \hline ((S \setminus (S / (S \setminus NP))) / (S \setminus NP)) \\ \hline : \lambda g. \lambda f. \neg(fg) \end{array} \quad \begin{array}{c} \text{love} \\ \hline (S \setminus NP) / NP \\ \hline : \lambda x. \lambda y. \text{love}'xy : \lambda x. \text{sko}'x : \text{woman}' \end{array} \quad \begin{array}{c} a \\ \hline NP / N \\ \hline : \lambda x. \text{sko}'x : \text{woman}' \end{array} \quad \begin{array}{c} \text{woman} \\ \hline N \\ \hline : \lambda x. \text{sko}'x : \text{woman}' \end{array} \\
\hline
\begin{array}{c} (S / (S \setminus NP)) \\ \hline : \lambda q. \forall x [\text{boxer}'x \rightarrow qx] : \{x\} \end{array} \quad \begin{array}{c} NP \\ \hline : \text{sko}'\text{woman}' \end{array} \\
\hline
\begin{array}{c} S \setminus NP \\ \hline \lambda y. \text{love}'(\text{sko}'\text{woman}')y \end{array} \\
\hline
\begin{array}{c} S \setminus (S / (S \setminus NP)) \\ \hline : \lambda f. \neg(f[\lambda y. \text{love}'(\text{sko}'\text{woman}')y]) \end{array} \\
\hline
S : \neg(\forall x [\text{boxer}'x \rightarrow \text{love}'(\text{sko}'\text{woman}')x]) : \{x\}
\end{array}$$

Figure 4.5: A derivation for the sentence *Every boxer doesn't love a woman*

$$\begin{array}{c}
\begin{array}{c} \text{John} \\ \hline NP \\ \hline : \text{john}' \end{array} \quad \begin{array}{c} \text{doesn't} \\ \hline ((S \setminus (S / (S \setminus NP))) / (S \setminus NP)) \\ \hline : \lambda g. \lambda f. \neg(fg) \end{array} \quad \begin{array}{c} \text{love} \\ \hline (S \setminus NP) / NP \\ \hline : \lambda x. \lambda y. \text{love}'xy : \lambda x. \text{sko}'x : \text{woman}' \end{array} \quad \begin{array}{c} a \\ \hline NP / N \\ \hline : \lambda x. \text{sko}'x : \text{woman}' \end{array} \quad \begin{array}{c} \text{woman} \\ \hline N \\ \hline : \lambda x. \text{sko}'x : \text{woman}' \end{array} \\
\hline
\begin{array}{c} (S / (S \setminus NP)) \\ \hline : \lambda f. f[\text{john}'] \end{array} \quad \begin{array}{c} NP \\ \hline : \text{sko}'\text{woman}' \end{array} \\
\hline
\begin{array}{c} S \setminus NP \\ \hline \lambda y. \text{love}'(\text{sko}'\text{woman}')y \end{array} \\
\hline
\begin{array}{c} S \setminus (S / (S \setminus NP)) \\ \hline : \lambda f. \neg(f[\lambda y. \text{love}'(\text{sko}'\text{woman}')y]) \end{array} \\
\hline
S : \neg(\text{love}'(\text{sko}'\text{woman}')\text{john}')
\end{array}$$

Figure 4.6: A derivation for the sentence *John doesn't love a woman*

According to my personal intuitions⁶, the negation operator should always out-scope the quantifier, but since native speakers of English do not seem to agree, we are left with the conclusion that the negation operator is lexically ambiguous. Notice that the value raising category in (102) is also compatible with a sentence such as *John doesn't love a woman*, as can be seen in Figure 4.6.

⁶ This may be because I am a native Turkish speaker, and in Turkish the negation operator is at the end of the sentence, and always seems to negate the whole sentence, as in 'Her boksör bir kadını sevmiyor', but in a properly established context, this intuition may also be proven incorrect.

4.8 Additional Notes and Concluding Remarks

The prosodic constraints on information structure mentioned in section (E) have an affect on in which contexts a particular word order is likely to be selected, but as we have seen in section (4.2), this has no affect on possible scope relations. Corrective tunes are only able to contrast and correct lexically specifiable meanings. In a dialogue as in (59), the scope relations are not lexically contrasted. On the other hand, Turkish provides words for that purpose, as in (104).

- (104) a. Her doktor aynı hastayı tedavi etti.
b. Hayır, her doktor BİRER hastayı tedavi etti.

H*

LL%

A particular word order may cause a certain context to be accommodated only if the context is not clearly established by the discourse, and such accommodation can be considered to be an extra-linguistic cognitive process, which may involve parser preferences, memory and time restrictions on cognitive processes, and maybe, simply imagination.

Steedman (2005) mentions a similar word order related scope alternation restriction due to the Japanese quantifier *daremo*, which is often translated to English as *everyone*. However, in contrast to the English scope inverting example *Someone loves everyone*, the following Japanese example is said to be unambiguous, and fails to invert scope.

- (105) Dareka-ga daremo-o aisitei-ru.
Someone-NOM everyone-ACC loves

‘Someone loves everyone’

($\exists A / * \forall A$)

Since we have found that similar Turkish sentences exhibit some ambiguity when an appropriate context is established, we predict that in some context, this Japanese sentence may also exhibit scope ambiguity. While this statement is not an conclusive argument against the Japanese data presented, the Japanese sentence above reminds us the fact that such data should always be taken with a grain of salt. In such cases, the proper act is not always to question the theory. Sometimes the data itself should be questioned.

As there are cognitive processes involved in such extra-linguistic performance phenomena there are also cognitive processes that have a role on disambiguation of quantifier scope ambiguities. For example, Walid Saba (1995) argues that disambiguation of quantifier scope ambiguities (in his terms *disambiguation of quantifiers* only, not *scope* ambiguities), falls under

the general problem of “lexical disambiguation”, which is essentially an inference problem. For example consider (106).

- (106) a. Every student in CS404 received a grade.
 b. Every student in CS404 received a course outline.

The syntactic structures in (106) are identical, and thus they should have the same scope relation between the quantifier and the indefinite. Both have scope ambiguity that would be resolved by general knowledge of the domain: typically students in the same class receive the same outline, but different grades. For (106b), such logical inference does not disambiguate the sentence.

The idea that such disambiguation is akin to lexical disambiguation seems to be compatible with Steedman’s 2005 account of quantifier scope, because the disambiguation process would involve choosing between the two different senses of skolem terms. But unlike a typical lexical disambiguation task, this one includes syntactic constraints: the narrow scope reading of the indefinite is allowed if it is syntactically licenced, which is subject to the surface-compositional constraints.

Michael Hess (1985) has interesting notes on the kind of data that linguists study, and the way natural languages make use of quantification. He particularly criticizes data used by non-computational linguists as follows.

Non-computational linguists do not very often use real-world examples in their investigations; they create their own example sentences to make a certain point. Everything which is not in the primary focus of their interest is made so explicit as to become largely self-explanatory. They tend, for instance, to create only sentences where quantification is explicit. Computational linguists, on the other hand, have to use real world texts. They have to face certain nasty facts of life which they, too, would prefer to ignore. One of them concerns the way in which Natural Language quantifies.

The paper continues to argue that, unlike the traditional way quantification has been studied (by making use of sample sentences which contain explicit quantifiers such as “every”, “all”, etc.), natural language has implicit ways of quantification. Such implicit quantification involve sentences such as the following.

- (107) a. Dogs eat meat.
 b. A man who loves a woman is happy.
 c. A man who loves a woman will give her a ring.

d. A text editor makes modifications to a text file.

Such sentences are known by the term *generics*, and Steedman (2005) explicitly states that his theory is not intended to cover generics.

When the claims of Michael Hess (1985) and Walid Saba (1995) are considered together, it may be possible to come up with an account for generics. From Saba comes the idea that quantifiers are lexically disambiguated. From Hess comes the idea that natural language makes implicit use of quantifiers. Combining both, we may claim that even the simplest nouns are actually lexically ambiguous between being a noun in the ordinary sense and being a generalized quantifier. The generalized quantifier sense is triggered by several factors. One is the availability of a restrictor such as ‘a man who loves a woman’. Another trigger maybe the tense of the predicate, as in ‘Dogs eat meat’. Under these conditions, it may be possible to claim that the lexical category N is suppressed and a syntactic category compatible with that of generalized quantifiers are used instead.

4.9 Conclusion

In this chapter we have applied Steedman’s theories of English Information Structure (Steedman, 2000a, 2000b), Steedman’s account of quantification (Steedman, 2005), and Özge and Bozşahin’s theory of Turkish Information Structure (Özge and Bozşahin, 2006) to the problem of quantifier scope ambiguity in Turkish and compared the results with previous studies of the same subject. We have pointed out how the data should be interpreted, and where and why we disagree with the results of the previous studies of the same subject in Turkish. The additional remarks (sections (4.5, 4.6, 4.7, and 4.8)) clarify some of the relevant subjects. Further discussion is deferred to Chapter (5).

CHAPTER 5

CONCLUSION

In this thesis we have examined the problem of quantifier scope ambiguity in Turkish and the various ways in which it has been accounted. The most difficult problems that we have faced during the study wasn't about the theories which we have examined, but about the misleading data presented by previous studies of the same problem. For example, given a weak intuition about the meaning of sentences such as the following, a more appropriate reaction is to examine what a particular linguistic theory or framework would say about it, instead of trying to create a separate theory for each such question.

- (108) a. Her hastayı bir doktor tedavi etti.
b. Bir doktor her hastayı tedavi etti.
c. Her hastayı tedavi etti bir doktor.

In this thesis we have examined what Steedman's relevant theories predict about the meaning of such sentences in properly established contexts. The conclusion of this study is that unlike some of the previous studies of quantifier scope ambiguity in Turkish, in CCG the problem can be accounted for without extra stipulation, and we couldn't see a need to propose a new theory just to give an explanation of rather weak intuitions about the meaning of such constructions. We had several related arguments which can be summarized as follows.

In section (4.2.1) we have argued against the use of the data extracted from 'null context' experiments to develop theories about the competence grammar, and argued that such data only needs a psycholinguistic explanation, not a linguistic one. More specifically, we have argued that word-order has no effect on availability of scope ambiguities, but rather just on discourse felicity of sentences. Some word orders may make certain scope relations more preferable than others but we consider such factors to be a performance phenomena related to memory/time constraints on processing (perhaps in the sense of Morrill (2000) or perhaps

due to immediate NP evaluation, for details see below), unrelated to the competence grammar. A question for further research might be exactly how the Type Logical approach of Morrill (2000) can be redefined in a way that is applicable to CCG. One possible solution might be to simulate CCG within Type Logical Grammar, as in Kruijff and Baldridge (2000) and to test whether the same results still hold. Another possible solution might involve proposing a cognitively plausible parser with some psycholinguistic support, as in Niv (1994).

In section (4.2.2), we have argued that prosodic marking and intonation does not yield a difference in the propositional content of the way an utterance such as (109b) is interpreted, because in such cases the intonationally kontrasted term (*ÜÇ kişiye* in this example) is an existential, and existentials are the underspecified terms. Such underspecified terms are inherently ambiguous, and kontrasting an ambiguous term does not clarify one of its possible senses. What is needed is a word with an explicit unambiguous meaning, such as *birer*, *aynı*, etc. A question for further research might be to look for if there is any data that can falsify this claim. In such data existence of intonational marking on existentials needs to be able to make a propositional content difference, and not just a discourse felicity difference.

(109) Situation: There are several people in a room, some of which have written a book.

- a. Herkes bu yıl kitaplarını iki kişiye ithaf etti.
- b. Hayır, herkes bu yıl kitaplarını ÜÇ kişiye ithaf etti.

H* LL%

In section (4.2.3) we have argued against the Kennelly (2003), stating that discourse binding of old information does not characterize quantificational dependencies either, and that Kennelly's constraints are not only unnecessary, but also incorrect, yielding further problems pointed out in section (4.3).

In section (4.4), we have argued that in an incremental processing of the utterances that make up a discourse, the concept of *theme* serves as bridge between subsequent utterances yielding a shared interpretation for scopally ambiguous quantificational utterances (that is, if the previous utterance is interpreted as having a narrow scope for an indefinite, the next one will also be interpreted as having a same narrow scope reading. Such utterances may still have differences in their propositional content (i.e. an indefinite may change from *iki kişiye* to *ÜÇ kişiye*, but the sentence still will have the scope relation of the way the previous utterance is interpreted, because of the shared discourse theme.

5.1 Explanatory Remarks

The idea that interpretation of an utterance is dependent on context is hardly new. Context is known to affect not only the task of lexical disambiguation, but also parsing of syntactically ambiguous sentences. For example Bever (1970) observes that naive subjects typically fail to find any grammatical analysis at all for “garden path” sentences such as (110a) while not having any difficulty with syntactically equivalent ones like (110b):

- (110) a. The doctor sent for the patient arrived.
b. The flowers sent for the patient arrived.

Crain and Steedman (1985) and Altmann and Steedman (1988) shows that establishing a proper context for the related sentences eliminates the garden path effect. Accordingly, (1) the proposed cognitively plausible parser processes sentences left-to-right, (2) the syntactic analyses are developed in parallel with semantic analyses, and (3) semantically most plausible syntactic analysis according to the current context is preferred. The idea can be made specific with the following principle:

- (111) *The Principle of Parsimony* (Steedman, 2000a, p.238)

The analysis whose interpretation carries fewest unsatisfied but accommodatable pre-suppositions will be preferred.

According to Steedman’s analysis of quantifier scope ambiguity (Steedman, 1999, 2005), the issue has both lexical and syntactic aspects. Lexically, indefinites are underspecified skolem terms, and according to our interpretation that can be considered as a special kind of lexical ambiguity. Syntactically, Steedman (1999, 2005) formulates a particular way such underspecified terms may get specified (in other words, disambiguated) during derivation. The method that he proposes explicitly assumes an interaction between lexical semantics and syntax, and gives a formal model of that interaction. As stated above, context is known to affect both lexical disambiguation and syntax. Since Steedman’s account has both lexical and syntactic aspects, it is hard to point out exactly what kind of effect context has – i.e. lexical, syntactic or both, but we can safely state that the easier detectability of narrow scope reading of the indefinite in Kennelly’s DP1 position is due to context, because (112a) already establishes a narrow scope reading in the rheme alternative set of discourse context, and that reading is available to the hearer at (112b). In present terms, there is no need to look for any other explanation. Context, by itself, is satisfactory.

(112) Situation: There are several people in a room, some of which have written a book.

a. Herkes bu yıl kitaplarını iki adama ithaf etti.

b. Hayır, ÜÇ KADINA ithaf etti herkes bu yıl kitaplarını.

H* LL% ⟨ – F – ⟩

To complete the understanding of the subject material we also need to find an explanation for the ‘null context’ case, as implicitly assumed in the previous studies of quantifier scope ambiguity in Turkish. As we shall see, the issue has connections to the subject of *interactivity* of syntax and semantics, and the question of *autonomy of syntax*. Crain and Steedman (1985) distinguishes between several senses of ‘autonomy’ with respect to syntax, which we shall summarize as follows¹.

1. Syntax and semantics can be distinguished *in theory*. Autonomy in this sense is called *formal autonomy*. There is no alternative to formal autonomy in a theory of language, as no theory can deny the existence of a phenomena called syntax.
2. In a second sense, autonomy considered to be *representational autonomy*, and refers to the extend to which, at some level of analysis, purely syntactic representations are built, which are later translated into semantic representations.
3. The alternative to (2), which can be *radical nonautonomy*, according to which the semantic interpretation is assembled directly, as is assumed in CCG. In such a theory, the rules of syntax describe what a processor *does* while assembling a semantic interpretation. The difference from more standard theories is that the rules do not describe a class of structures that are built (for example, the combinatory derivation which can be considered as a phrase structure is not a level of representation). On the other hand, semantic representation, as distinct from the process of its evaluation, must be represented *somehow* and it might be appropriate to think of that representation as a structure (for example, in CCG, the predicate argument structure is a level of representation). However, according to the radical version of the doctrine of representational nonautonomy, it is a structure that neither can be inspected or changed, nor needs to be, in order to produce an object that can be evaluated, but it can permit the evaluation of subexpressions while syntactic processing of higher expressions remains incomplete (thus, it leaves the

¹ Examples in parentheses are ours.

door open to possibilities such as making use of (possibly evaluated) semantic features in syntactic derivation and lexical disambiguation, and makes the possibility of such interaction easier to formulate and formalize, as exemplified in the case of environment calculation in Steedman (2005), with the difference that in Steedman (2005) the environment set is just calculated, not evaluated, but the resulting formula after skolem specification can be evaluated.).

4. The term *autonomy* has also been used to the possibility that interaction of syntax and semantics during local ambiguity resolution is possible. Theories are free to be either entirely noninteractive (interaction is forbidden before the syntactic analysis of the sentence is complete), or partially interactive at some level such as the clause or phrase, or entirely interactive, even at the level of words or morphemes. The choice of having interaction during local ambiguity resolution is independent of the question of representational autonomy. Representationally nonautonomous models may be interactive or noninteractive. Representationally autonomous theories are also free to be interactive or noninteractive.

While (4) states that the issue of interactivity and representational nonautonomy of syntax are independent questions, it has hard to see how one can formulate a theory like Steedman (2005) and still assume representational autonomy of syntax. As explained in (3), CCG is not only a radically lexicalist theory with lexically specified syntax, semantics and derivational control, but also a radically nonautonomous one with respect to the question of syntax being a representational level, and it makes use of the possibilities that this understanding gives whenever appropriate. In such a theory, the idea that context can affect lexical interpretation and syntactic derivation is highly natural: Syntax is just the trace of what the processor does under certain circumstances. Accordingly, under this view competence grammar and performance grammar² is one and the same, but the processor is ‘incomplete’ in the sense that it

² According to Carpenter (1997), the line between competence grammar and performance grammar can be drawn as follows:

The primary goal of theoretical linguistics, as opposed to psycholinguistics, is to formulate a theory of language itself, rather than the human ability to process it. Chomsky (1965) drew a distinction between linguistic *competence*, on the one hand, and *performance* on the other. Chomsky believed that our competence comprised a system of rules for the construction of utterances and their meanings. Chomsky further assumed that the knowledge of such rules is innate, as is the ability to use the rules. Linguistic performance is subject to the limitations of all human cognitive activities, such as attention span, alertness, distractions, and so on. Abstracting away from performance issues, Chomsky took it to be the job of linguistics to construct a competence model (Carpenter, 1997).

My personal interpretation of the subject is as follows: ‘Competence’ is the knowledge of language and it’s a theoretical construct. It can be consciously known, or be studied in a linguistics class. It’s an idealized linguistic

does not follow all syntactic possibilities available, but only semantically plausible ones (Note that ‘semantic plausibility’ is defined as a parsimony principle on the number of accommodated presuppositions, not in the general sense of what can possibly happen in the current context according to the world knowledge). It should also be noted that radical lexicalism is a corollary of nonautonomy of syntax: meaning matters not just for interpretation but also for syntax, and this fact can only be lexically handled – not by grammar rules.

Under the light of the information from Crain and Steedman (1985) regarding the nature of the cognitive utterance processor, we can have further predictions about the processing of sentences such as (108b). It is possible to speculate that when the Principle Of Parsimony (111) is considered together with the assumption of incremental processing, representational nonautonomy and the possibility of immediate evaluation, it directly predicts that in so called ‘null contexts’, the most easily accommodatable context is the one that involves a wide scope interpretation of the indefinite³. Consider a sentences such as (108b), repeated here:

(113) Bir doktor her hastayı tedavi etti.

△

At the instant that the hearer has heard the phrase *bir doktor*, but not the rest of the utterance, the assumption of incremental interpretation together with the Principle of Parsimony

knowledge of an ideal speaker. On the other hand, a performance grammar is the declarative statement of what a language processor can possibly do, and in what way it does it. It has a specific cognitive claim about the human processor. Ideas from studies of artificial languages (computer programming languages, for example) and processors designed for them also casts the issue in another way: Sometimes, the grammars that we have on paper are not suitable for a specific kind of processor. In that case, if we still want the processor to be able to parse the strings of that grammar, the competence grammar must be converted to some weakly equivalent grammar suitable for the nature of the processor. Accordingly, the strict competence hypothesis (that the competence grammar and the performance grammar is one and the same) can be understood as follows: We have the correct grammar on paper, which is strongly equivalent to a declarative statement of the actions that can be possibly taken by the processor. In this view, there is a one-to-one correspondence between processor actions and the grammar rules we have on the paper, and the competence grammar does not just generate the correct strings, it generates them in the same way as the cognitive processor does. This is known to simplify linguistic theories. For example, in CCG the surface constituents and the intonational constituents are the same, and do not require complex mapping structures, contrary to theories by Selkirk (Selkirk, 1972, 1984, 1986, 1990, 1995) who assumed the grammars designed in the generative paradigm and run into complex problems about the mismatch between empirically supported intonational constituents and theoretically assumed surface constituents. The reason is that, the kind grammars assumed by the generative paradigm is not strongly equivalent to the grammar used by the processor (probably, they have never intended to be, because according to Chomsky the job of linguistics is to study the language itself, not the human processing of it (at least that seems to be the initial purpose; apparently the purpose has changed later during the 1970s as mentioned in Chomsky (2005), but still it did not yield a correct grammar, possibly because of false initial assumptions)). Therefore, grammars designed in the generative paradigm just generate the correct strings (i.e. they are weakly equivalent to the actual performance grammar.).

³ Note that this argument should not be confused with the skolem trigger idea discussed in Chapter (3). Also note that this argument refers to an ‘online processing task’, i.e. the subject is not given the time to study the sentence and examine its meaning on several contexts he/she might imagine, in which case he/she could possibly find out all the readings licenced by the competence grammar.

directly predicts that the hearer will accommodate a wide scope reading of the NP ‘bir doktor’, because at the instant marked by the triangle ($\triangle$), the quantifier *her* is not heard yet, so a context in which more than one doctor is available is not accommodatable (as it would require a yet-unjustified presupposition), and after the rest of the sentence is processed, the previously accommodated context is not corrected, unless there is a reason to do so. Since this reasoning path predicts the scope preference indicated by Morrill (2000) and explains it in a more intuitive way, it should be considered as the preferred way to explain the issue at hand. For details, see Crain and Steedman (1985) and Altmann and Steedman (1988) which present further details and interesting arguments in favor of the assumed properties of the cognitive processor. In particular, Altmann and Steedman (1988) further argues that such interactive accounts of utterance processing does not undermine Fodor’s (1983) modularity criterion.

REFERENCES

- Ajdukiewicz, K. (1935) "Die syntaktische Konnexitat." In Storrs McCall, ed., *Polish Logic 1920-1939*, 207-231. Oxford: Oxford University Press. Translated from *Studia Philosophica*, 1, 1-27.
- Altmann, G., and Steedman, M. (1988) "Interaction with Context During Human Sentence Processing." *Cognition*, 30, 191-238.
- Austin, J. L. (1961) "Performative utterances." In J. L. Austin, *Philosophical Papers*, J. O. Urmson and G. J. Warnock, editors, Oxford: Oxford University Press.
- Austin, J. L. (1962) *How to Do Things with Words*. Revised second edition, 1975. Oxford: Oxford University Press.
- Baldrige, J. (2002) *Lexically Specified Derivational Control in Combinatory Categorical Grammar*. Ph.D. thesis, University of Edinburgh, Edinburgh.
- Bar-Hillel, Y. (1953) "A Quasi-Arithmetical Notation for Syntactic Description." In *Language*, 29, 47-48.
- Bever, T. (1970) "The Cognitive Basis for Linguistic Structures." In John Hayes, ed., *Cognition and the Development of Language*, 279-362. New York: Wiley.
- Blackburn, P. and Bos, J. (2005) *Representation and Inference for Natural Language: A First Course in Computational Semantics*. CSLI Publications.
- Bozsahin, C. (2002) "Combinatory Morphemic Lexicon." In *Computational Linguistics* 28:2. 145 - 186.
- Cahn, J. (1995) "The Effect of Pitch Accenting on Pronoun Referent Resolution." In *Computational Linguistics*, 33, 290-292.

- Calhoun, S., Nissim M., Steedman M., and Brenier J. (2005) "A Framework for Annotating Information Structure in Discourse." In *ACL 2005 Workshop: Frontiers in Corpus Annotation II: Pie in the Sky*.
- Carnap, R. (1952) "The meaning postulates." In *Philosophical Studies* 3:65-73.
- Carpenter, B. (1997) *Type-Logical Semantics*. The MIT Press.
- Crain, S., and Mark Steedman. (1985) "On Not Being Led up the Garden Path: the Use of Context by the Psychological Parser." In Lauri Karttunen David Dowty and Arnold Zwicky, eds., *Natural Language Parsing: Psychological, Computational and Theoretical Perspectives*, 320-358. Cambridge: Cambridge University Press.
- Cook, V. J. Newson, M. (1996) *Chomsky's Universal Grammar: An Introduction*. Second Edition. Malden: Blackwell Publishing, 1996.
- Chomsky, N. (1965) *Aspects of the Theory of Syntax*. The MIT Press.
- Chomsky, N. (2005) "Three Factors in Language Design." In *Linguistic Inquiry*, 36:1, 1-22. The MIT Press.
- Church, A. (1951) "The need for abstract entities in semantic analysis." In *American Academy of Arts and Science Proceedings*, 80:100-113. Reprinted as "Intensional semantics" in A. P. Martinich, editor, *The Philosophy of Language*, second edition, 1990, 40-47. Oxford: Oxford University Press.
- Cooper, R. (1983) *Quantification and Syntactic Theory*. Dordrecht:Reidel.
- Erguvanlı, E. E. (1979) *The Function of Word Order in Turkish Grammar*, Ph.D. dissertation, UCLA. Published version 1984, University of California Press, Berkeley.
- Erk , F. (1982) *Topic, comment, and word order in Turkish*. Minnesota Papers in Linguistics and Philosophy of Language, 30-38.
- Erk , F. (1983) *Discourse pragmatics and word order in Turkish*. Dissertation, University of Minnesota.
- Farkas, D. (1997) "Evaluation Indices and Scope." In Anna Szabolsci, ed., *Ways of Scope Taking*. Dordrecht: Kluwer, 1997.

- Fodor, J. A. (1983) *The Modularity of Mind*. Cambridge, MA: MIT Press.
- Frege, G. (1892) “Über Sinn and Bedeutung.” *Zeitschrift für Philosophie und Philosophische Kritik* 100:25-50. Translated by H. Feigl as “Sense and nominatum” in H. Feigl and W. Sellars, editors, *Readings in Philosophical Analysis*, 85-102.
- Gabbay, D. Kempson, R., Meyer-Viol, W. (1994) “Utterance interpretation: a truth theoretic persiective.” Ms. (To appear in Kempson, Ruth and Lappin, Shalom (editors), *Bulletin of IGPL*.).
- Gabbay, D. and Kempson, R. (1991, 1992). “Natural language content: a proof-theoretic perspective.” In Göksel and Parker (editors) 1991/1992: 85-116.
- Geach, Peter. (1972) “A Program for Syntax.” In Donald Davidson and Gilbert Harman, eds., *Semantics of Natural Language*, 483–497. Dordrecht: Reidel.
- Göksel, A. (1998) “Linearity, focus and the postverbal position in Turkish.” In L. Johanson, et.al., eds., *The Maiz Meeting: Proceedings of the Seventh International Conference on Turkish Linguistics, August 3-6, 1994*, 85-106, Wiesbaden:Harrassowitz, 1998.
- Grice, H. P. (1965) “Logic and conversation.” In P. Cole and J. L. Morgan, editors, *Speech Acts, Syntax and Semantics*, 3:41-58. New York: Academic Press.
- Grosz, B., A. Joshi and S. Weinstein, (1995) “Centering: A Framework for Modeling the Local Coherence of Discourse.” In *Computational Linguistics*, 21, 203-225.
- Grosz, B., and C. Sidner, (1986) “Attention, Intention and the Structure of Discourse.” In *Computational Linguistics*, 12, 175-204.
- Hess, M. (1985) “How does natural language quantify?” In *Second Conference of the European Chapter of the Association for Computational Linguistics*, University of Geneva, Geneva, Switzerland.
- Hausser, R. (1984) *Surface Compositional Grammar*. München: Wilhelm Fink.
- Hausser, R. (1999) *Foundations of Computational Linguistics: Man-Machine Communication in Natural Language*. Berlin, New York: Springer.
- Hobbs, J. R. (1983) “An improper treatment of quantification in ordinary English”. In *Proceedings of 21st Annual Meeting of the Association for Computational Linguistics, 15-17 June 1983, Massachusetts Institute of Technology Cambridge, Massachusetts*,

57-63.

Hobbs, J. R. and Shieber, S. M. (1987) "An algorithm for generating quantifier Scopings." In *Computational Linguistics*, 13:47- 63.

Kaplan, D. (1964) *Foundations of intensional logic*. Ph.D. dissertation, University of California, Los Angeles.

Kaplan, D. (1979) "On the logic of demonstratives." In *Journal of Philosophical Logic*. 8(1):81-98.

Kamp, H. Reyle, U. (1993) *From Discourse To Logic: Introduction to Model-Theoretic Semantics of Natural Language, Formal Logic and Discourse Representation Theory*. Dordrecht: Kluwer, 1993.

Kayne, R. (1994) *The antisymmetry of syntax*. Cambridge: MIT Press.

Keller, W. R. (1988) "Nested Cooper Storage: The proper treatment of quantification in ordinary noun phrases." In U. Reyle and C. Rohrer, eds., *Natural Language Parsing and Linguistic Theories*, pages 432-447. Dordrecht: Reidel.

Kennelly, S. D. (2003) "The implications of quantification for the role of focus in discourse structure." In *Lingua*, 113:1055-1088.

Kornfilt, J. (1996) "On rightward movement in Turkish." In L. Johanson, et.al., eds., *The Mainz Meeting: Proceedings of the Seventh International Conference on Turkish Linguistics, August 3-6, 1994*, 107-123, Wiesbaden:Harrassowitz, 1998.

Kruijff, Geert-Jan M. (2001) *A Categorical-Modal Logical Architecture of Informativity: Dependency Grammar Logic and Information Structure* Ph.D. thesis, Charles University, Prague, Czech Republic.

Kruijff, Geert-Jan M. and Baldridge, Jason M. (2000) "Relating categorial type logics and CCG through simulation." Unpublished manuscript, University of Edinburgh.

Kruijff, Geert-Jan M. and Vasishth, S. (1997) *Competence and Performance Modelling of Free Word Order*. ESSLI 2001.

Kural, M. (1994) *Postverbal Constituents in Turkish*. Ms. UCLA.

- May, R. (1985) *Logical Form*. The MIT Press.
- Montague, R. (1973) "The proper treatment of quantification in ordinary English." In Paul Portner and Barbara H. Partee, eds., *Formal semantics: The Essential Readings*. Oxford: Blackwell Publishers Ltd, 2002.
- Morrill, G. V. (1994) *Type Logical Grammar: Categorical Logic of Signs*. Dordrecht: Kluwer, 1994.
- Morrill, G. V. (2000) "Incremental Processing and Acceptability." In *Computational Linguistics*, 26:3. The MIT Press.
- Niv, M. (1994) "A Psycholinguistically Motivated Parser for CCG." In *Proceedings of the 32th Annual Meeting of the Association for Computational Linguistics, New Mexico State University Las Cruces, New Mexico, USA*, 125-132. San Francisco, CA: Morgan Kaufmann.
- Özge, U. and Bozşahin, C. (2003) *A Tune-based Account of Turkish Information Structure*. Ms. METU.
- Özge, U. and Bozşahin, C. (2006) *A Tune-based Account of Turkish Information Structure and Surface Word Order*. Unpublished manuscript, METU.
- Park, J. C. (1995) "Quantifier Scope and Constituency." In *33rd Annual Meeting of the Association for Computational Linguistics, Massachusetts Institute of Technology Cambridge, Massachusetts, USA*, 205-212. San Francisco, CA: Morgan Kaufmann.
- Partee, B. H., Meulen, A. T., and Wall, R. E. (1990) *Mathematical Methods in Linguistics*. Dordrecht, The Netherlands: Kluwer.
- Ross, J.R. (1967) *Constraints on Variables in Syntax*. Ph.D. thesis, MIT. Published as *Infinite Syntax!*, Ablex, Norton, NJ, 1986.
- Searle, J. R. (1965) "What is a speech act?" In M. Black, editor, *Philosophy in America*, 221-239. Ithaca, New York: Cornell University Press.
- Searle, J. R. (1975) "Indirect speech acts." In P. Cole and J.L. Morgan, editors, *Speech Acts, Syntax and Semantics*, 3:58-82. New York: Academic Press.
- Selkirk, E. (1972) *The Phrase Phonology of English and French*, Ph.D. thesis, MIT. Garland: New York, 1980.

- Selkirk, E. (1984) *Phonology and Syntax*, Cambridge, MA: MIT Press.
- Selkirk, E. (1986) 'Derived Domains in Sentence Phonology', In *Phonology*, 3:371-405.
- Selkirk, E. (1990) 'On the Nature of Prosodic Constituency.', In *Papers in Laboratory Phonology*, 1:179-200.
- Selkirk, E. (1995) 'Sentence Prosody: Intonation, Stress, and Phrasing.', In Goldsmith, ed. *The Handbook of Phonological Theory*, 550-569. Oxford: Blackwell.
- Stabler, E. (1997) "Derivational Minimalism." In C. Retore(Ed.), *Logical aspects of computational linguistics*. 68-95. Springer Verlag.
- Stabler, E. (1999) "Remnant Movement and Complexity." In Bouma, E. Hinrichs, G.-J. Kruij, D. Oerhle (eds.), *Constraints and Resources in Natural Language Syntax and Semantics*. 68-95. CSLI.
- Steedman, M. (1990) "Structure and Intonation in Spoken Language Understanding." In *Proceedings of the 28th Annual Meeting of the Association for Computational Linguistics, 6-9 June 1990 University of Pittsburgh Pittsburgh, Pennsylvania, USA*, 9-16. Madison, Wisconsin: Omnipress.
- Steedman, M. (1996) *Surface Structure and Interpretation*. Cambridge, MA: MIT Press
- Steedman, M. (1999) "Quantifier Scopep Alternation in CCG." In *Proceedings of the 37th Annual Meeting of the Association for Computational Linguistics, 20-26 June 1999 University of Maryland College Park, Maryland, USA*, 301-308. Madison, Wisconsin: Omnipress.
- Steedman, M. (2000a) *The Syntactic Process*. Cambridge, MA: MIT Press
- Steedman, M. (2000b) "Information Structure and the Syntax-Phonology Interface." In *Linguistic Inquiry*, 34, 649-689.
- Steedman, M. (2003) "Information-Structural Semantics for English." Paper to LSA Summer Institute Workshop on Topic and Focus, Santa Barbara July 2001.
- Steedman, M. (2005) "Surface Compositional Scope Alternation Without Existential Quantifiers." Unpublished manuscript, University of Edinburgh.

- Steedman, M., Calhoun, S., Nissim M., and Brenier J. (2005) "A Framework for Annotating Information Structure in Discourse." In *Proceedings of the Workshop on Frontiers in Corpus Annotation II: Pie in the Sky*, 45-52.
- Francez, N. Steedman, M. (2006) "Categorial Grammar and the Semantics of Contextual Prepositional Phrases." To appear in *Linguistics and Philosophy*, 29.
- Steedman, M. and Korbayova, I. K. (2003) "Discourse and Information Structure." In *Journal of Logic, Language and Information*, 12, 249-259.
- Strube, M. and Hahn, U. (1999) "Functional Centering – Grounding Referential Coherence in Information Structure." In *Computational Linguistics*, 25:1, 309-344.
- Tarski, A. (1935) "Der Wahrheitsbegriff in den formalisierten Sprachen." *Studia Philosophica* 1:261–405. English translation, "The concept of truth in formalized languages" in A. Tarski, *Logic, Semantics, Metamathematics*. Oxford: Clarendon Press, 1956.
- Saba, W. (1995) "Towards a Cognitively Plausible Model for Quantification." In *Proceedings of the 33th Annual Meeting of the Association for Computational Linguistics, 26-30 June 1995 Massachusetts Institute of Technology, Cambridge, Massachusetts, USA*, 323-325. San Francisco, CA: Morgan Kaufmann.
- van der Sandt, R.A. (1992) "Presupposition projection as anaphora resolution." In *Journal of semantics*, 9:333-377.
- Walker, M., A. K. Joshi, and E. F. Prince (eds.), (1998) *Centering Theory in Discourse*. Oxford University Press.
- Zadrozny, W. (1992) "On Compositional Semantics." In *Proceedings of the fifteenth International Conference on Computational Linguistics*, 1, 260-266.
- Zadrozny, W. (1994) "From Compositional to Systematic Semantics." In *Linguistics and Philosophy*, 17, 329–342.
- Zadrozny W., Lappin S. (2000) "Compositionality, Synonymy, and the Systematic Representation of Meaning." Unpublished manuscript, submitted to *Linguistics and Philosophy*.
- Zwart, J.W. (2002) "The Antisymmetry of Turkish." In *Generative Grammar in Geneva* 3:23-36, 2002.

APPENDICES

APPENDIX A. COMBINATORY CATEGORIAL GRAMMAR

Combinatory Categorical Grammar has a very distinctive advantage of having a simple, clear, and easy to understand mathematical model instead of a long wordy presentation of principles, which is in sharp contrast to the tradition of transformational grammar. In this section we shall briefly summarize Steedman's Combinatory Categorical Grammar (CCG) (Steedman, 1996, 2000a). For details, the reader is referred to Steedman (1996, 2000a) and Baldridge (2002).

CCG is a generalization of Pure Categorical Grammar (hereafter 'CG', but sometimes abbreviated as 'AB') of Ajdukiewicz (1935) and Bar-Hillel (1953). CG is a lexicalist approach, like other lexicalist approaches it puts most of the information that can be defined with rewriting rules to the lexicon. The grammatical signs of CG are categories which are atomic elements (primitive categories such as N, NP, S, and so on) or functions which specify the linear direction in which they seek arguments, as in the following examples¹.

book	:=	N
red	:=	N /N
Turkey	:=	NP
the	:=	NP /N
sleep	:=	S \NP
eats	:=	(S \NP)/NP

The slash notation defines *book* to be an atomic element which is a noun, and *eats* to be a function category. The forward slash '/' defines the first argument of this functor to be an **NP** which is to be found in the right side of the functor *eats*, and its result to be a predicate of type

¹ In this thesis we shall use Steedman's argument rightmost notation for categories, which always places the result category on the left, contrary to Lambek's notation in which functions that look for an argument in their left side are notated such that the argument category is placed in the left of the result category. In Lambek's notation the lexical entry for *eats* would be $(np \backslash s)/np$.

$S \backslash NP$, which is a function category that is looking for an NP in its left, due to the backward slash ‘\’, to yield a result type of S .

In pure categorial grammar (CG), only two rules are defined, with which combination takes place:

(114) Functional application

- a. $X/Y \quad Y \quad \Longrightarrow X$ ($>$)
- b. $Y \quad X \backslash Y \quad \Longrightarrow X$ ($<$)

These rules correspond to a binary context free phrase structure rule schemata, and in fact this pure form of CG is just context free grammar written in the accepting, rather than the producing direction. Due to the fact that categories such as X , Y and X / Y are retrieved from the lexicon, the definition of grammar is transferred from the phrase structure rules to the lexicon, which is the sole source of information about the grammar.

Practical language understanding systems also need to introduce minor syntactic features like number, gender and person agreement. For example, the verb *eats* can be defined as follows, in which *eats* is specified only for number and person agreement, and “underspecified” for gender:

(115) $eats := (S \backslash NP_{3SG}) / NP$

Such minor syntactic features can be checked and computed with the process of unification, and CG augmented with such minor syntactic features is a lexicalized context free grammar with attributes².

While at first glance this idea roughly corresponds to the ‘Case Theory’ in the minimalist version of Chomskian theorizing, according to which Case features of Determiner Phrases (DPs) are inserted from the lexicon, and that DPs move to the specifier position of an ‘Agreement Phrase’ (AGRP) to check their Case features (Cook and Newson, 1996, p.329), notice that CG with minor syntactic features can handle such checking simply with unification and without making use of movement and the related stipulation of any specific syntactic position, whether AGRP or something else.

Figure A.1 illustrates the theoretical apparatus introduced so far.

This derivation can be seen as a binary branching phrase structure tree turned upside down as shown in Figure A.2, which is adapted from Baldridge (2002, p.16). Apart from the

² Context free grammars with attributes are still context free grammars, because it is possible to write the same grammar without using attributes, but the resulting grammars can be too cumbersome, and that can be interpreted as a loss of explanatory power even if they may be strongly equivalent to the original grammar with attributes.

usage of more informative node labels, this is reminiscent of the traditional phrase structure representation, given in Figure A.3.

While the trees in Figure A.2 and Figure A.3 seem very similar, the usage of more informative node labels goes a long way:

1. Subcategorization is directly encoded in functor categories rather than through the use of new symbols such as $V_{intrans}$, V_{trans} and $V_{ditrans}$ (Baldrige, 2002, p. 16)). In addition, when you encode useful information in a usable way, you can not only computationally process it more elegantly, but also that opens the door for more interesting possibilities, such as the availability of ‘nonstandard’ constituents as discussed in section (A.2).
2. There is a systematic³ correspondence between notions such as *intransitive* and *transi-*

³ The importance of this systematicity is that transitive verbs can be converted to intransitive verb phrases by supplying an argument, and the resulting intransitive ‘constituent’ can be used in coordinate sentences as follows:
 (116) Burak can [read a book]_{iv} and [sleep]_{iv} at the same time.

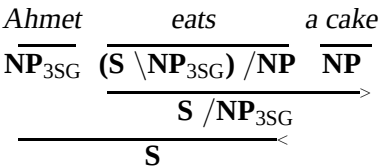


Figure A.1: A CG derivation with minor syntactic features

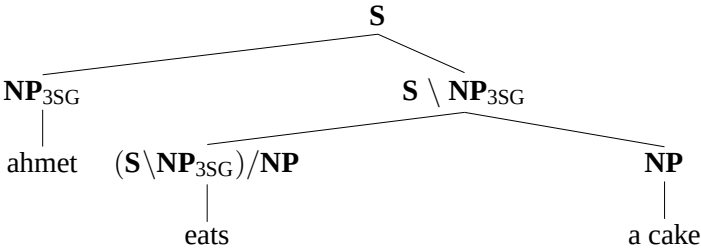


Figure A.2: A CG derivation with minor syntactic features, viewed as a phrase structure

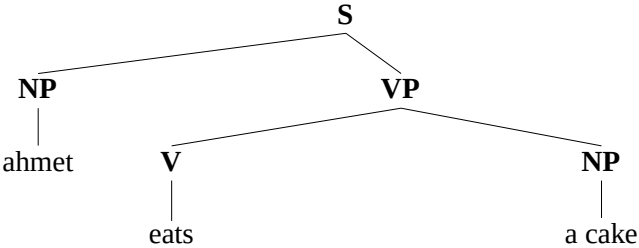


Figure A.3: The corresponding phrase structure with traditional node labels

tive (Baldrige, 2002, p. 16)).

3. Feature unification implements case-checking and agreement without use of movement, which is possible because of the more informative node labels.

In CCG, such phrase structures are simply records of the derivation process, and unlike the standard generative tradition, they do not constitute a level of representation, and we are not interested in keeping these records for future use, because nothing is predicated over them. In fact, a cognitively plausible incremental parser may not even ‘remember’ all of the derivation history.

The lexical entries for verbs such as *sleep* and *eats* as defined above, also provide a natural definition of the notion of “domain of locality” (relevant for the binding theory), in terms of the primitive clause containing the verb and its arguments. The assumed binding theory is roughly the same of the one developed within GB, but predicated at the level of ‘predicate argument structure’ instead, using a notion similar to Chomsky’s “c-command”.

A.1 Predicate Argument Structure

Categories also include semantic interpretations. Interpretations can be considered as saturated or unsaturated predicate argument structures, logical forms in the logician’s sense of the term.

There are a number of ways in which interpretations can be made explicit in the notation of categories. One possibility is the following, which makes use of unification for this purpose:

$$(117) \text{ eats} := (\mathbf{S} : \text{eat}'_{xy} \setminus \mathbf{NP}_{3\text{SG}} : y) / \mathbf{NP} : x$$

Another possibility is to make use of lambda abstraction.

$$(118) \text{ eats} := (\mathbf{S} \setminus \mathbf{NP}_{3\text{SG}}) / \mathbf{NP} : \lambda x. \lambda y. \text{eat}'_{xy}$$

In this notation, the application rules can be written as follows:

(119) Functional application

- a. $X/Y : f \quad Y : a \quad \Longrightarrow X : fa$
- b. $Y : a \quad X \setminus Y : f \quad \Longrightarrow X : fa$

Such rules are subject to the Principle of Combinatory Transparency, which can be defined as follows:

(120) *The Principle of Combinatory Transparency*: The interpretation of a syntactic combinatory rule must be the one that would result from the equivalent combinatorily transparent unification-based interpretation of the rule. (Steedman, 1996)

This principle says that the syntactic form of these rules entirely determines their semantics, which explicitly assumes the principle of direct (surface) compositionality (Montague, 1973; Jacobson 1996, 1999; Hausser 1984, 1999). In other words, Combinatory rules are prohibited from accessing or manipulating the predicate argument structure in any way. Movement, deletion, and other familiar transformations are not allowed.

Interpretations and predicate-argument structures are such that all predicates are “curried” (that is, are functions from their first argument into a function over their next argument, and so on.) They are written using a convention under which the application of such a function (like *eat'*) to an argument (like *apples'*) is represented by concatenation (as in *eat'apples'*), and the application associates to the left, i.e. *eat'apples'keats'* is equivalent to *(eat'apples')keats'*. Assumption of left-associativity also means that such predicate argument structures are equivalent to binary trees, which preserve traditional dominance and command (which is called lf-command in CCG literature) but do not preserve the linear order of the string. For these reasons, the traditional binding theory and the related notions such as the obliqueness order of predicate arguments can be preserved.

Figure A.4 illustrates the use of lambda abstraction to generate the predicate argument structure. Note that, in the resulting logical form *ahmet'* lf-commands *a'cake'*.

$$\begin{array}{c}
 \text{Ahmet} \quad \text{eats} \quad \text{a cake} \\
 \hline
 \text{NP}_{3\text{SG}} \quad (\text{S} \setminus \text{NP}_{3\text{SG}}) / \text{NP} \quad \text{NP} \\
 : \text{ahmet}' \quad : \lambda x. \lambda y. \text{eat}' xy \quad : \text{a'cake}' \\
 \hline
 \text{S} / \text{NP}_{3\text{SG}} \quad \rightarrow \\
 \hline
 \text{S} \quad \leftarrow \\
 : \text{ate}'(\text{a'cake}') \text{ahmet}'
 \end{array}$$

Figure A.4: Use of lambda abstraction

A.2 Combinatory Generalization

To extend categorial grammars to cope with coordination we need a rule schema like the following⁴:

⁴ Instead of this rule, in Jason Baldridge’s Multi-Modal CCG, categorial assignments like “and := $(X \setminus X) / X$ ” can also be used because lexically specified derivational control avoids certain undesirable combinations available

(121) Coordination(ϕ^n)

$$X \text{ CONJ } X \implies \phi^n X$$

Because X may be any category including functor categories of any valency, the rule has to be schematized semantically for such types:

(122) Coordination(ϕ^n)

$$X : f \text{ CONJ } : b \ X : g \implies \phi^n X : \lambda \dots b(f \dots)(g \dots)$$

The rules of function composition generate nonstandard surface components which account for sentences like (124) as seen in Figure A.5. This is required because in pure categorial grammar there is no way to combine the subject with the verb until the verb first combines with its object.

(123) Forward composition ($>\mathbf{B}$)

$$X/Y : f \quad Y/Z : g \implies X/Z : \lambda x. f(gx)$$

(124) Keats cooked and might eat apples.

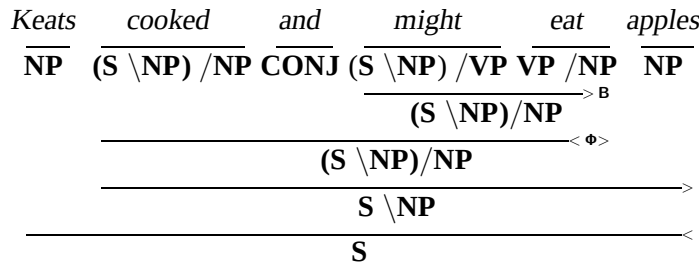


Figure A.5: Forward $\mathbf{B}$ combinator example

Similarly backward composition has been found to be necessary in certain derivations (Steedman, 2000a, p46), as seen in Figure A.6.

(125) Backward composition ($<\mathbf{B}$)

$$Y \setminus Z : g \quad X \setminus Y : f \implies X \setminus Z : \lambda x. f(gx)$$

The forward composition rule can be generalized as follows, allow composition into higher valency functors.

(126) Generalized forward composition ($>\mathbf{B}^n$)

$$X/Y : f \quad (Y/Z) / \$: \dots \lambda z. gz \dots \implies_{>\mathbf{B}^n} (X/Z) / \$: \dots \lambda z. f(gz \dots)$$

due to such categorial assignments when a categorial system more powerful than CG (such as CCG) is employed (Baldrige, 2002, p. 17, p.97).

Similarly we can have a generalized backward composition rule. In the formula above, $\$$ is a shorthand defined as below:

Definition 4 (\$ convention) For a category α , $\{\alpha\ \$\}$, (respectively $\{\alpha/\ \$\}$, $\{\alpha\backslash \$\}$) denotes the set containing α and all functions (respectively, leftward functions, rightward functions) into a category in $\{\alpha\ \$\}$ (respectively $\{\alpha/\ \$\}$, $\{\alpha\backslash \$\}$).

Unbracketed $\alpha\ \$$, $\alpha\backslash \$$ and, $\alpha/\ \$$ denote a single member of such sets⁵. Furthermore, subscripts such as $S\backslash \$_1$ are used with the meaning that the $\$$ category with the same numeral subscript is the same member of that set.

In order to capture a number of further phenomena related to coordination and unbounded dependency, it is desirable to regard categories such as NP and PP as functors, obtained by applying type-raising rules (61) to the original argument category⁶:

(127) Type Raising (**T**)

$$a. \quad X : a \implies_T T/(T\backslash X) : \lambda f.f a$$

$$b. \quad X : a \implies_T T\backslash(T/X) : \lambda f.f a$$

where X is an argument type.

Another rule, which permits analysis of sentences such as (129) is Backward Crossed Substitution, is defined below.

(128) Backward Crossed Substitution ($< \mathbf{S}_X$)

$$Y/Z : g \quad (X/Y)\backslash Z : f \implies_S X/Z : \lambda x.f x(gx)$$

(129) Mary will copy and file without reading these articles. (Steedman, 2000a, p.52)

Additional variants of these rules can be defined, but all combinatory rules are subject to the following principles:

⁵ For example, $S/\ \$$ refers to a single member of the set $\{S, S/NP, (S/NP)/NP, \dots\}$

⁶ Note that the abbreviation $NP^\uparrow$ is often used with the meaning ‘type raised NP’.

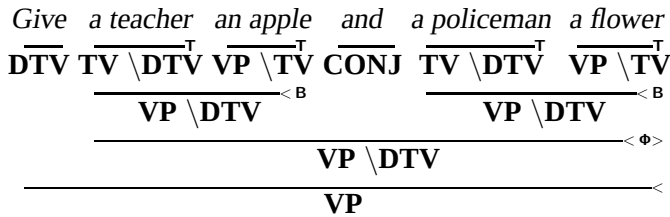


Figure A.6: Backward **B** combinator example

(130) *The Principle of Adjacency*: Combinatory rules may apply to finitely many phonologically realized and string-adjacent entities. (Steedman, 2000a, p.54)

(131) *The Principle of Consistency*: All syntactic combinatory rules must be consistent with the directionality of the principal function. (Steedman, 2000a, p.54)

(132) *The Principle of Inheritance*: If the category that results from the application of a combinatory rule is a function category, then the slash defining directionality for a given argument in that category will be the same as the one(s) defining directionality for the corresponding argument(s) in the input functions(s). (Steedman, 2000a, p.54)

The Principle of Adjacency is self explanatory. The Principle of Consistency excludes the following kind of rule:

(133) $X \backslash Y \ Y \not\Rightarrow X$

The Principle of Inheritance excludes rules like the following instance of composition:

(134) $X / Y \ Y / Z \not\Rightarrow X \backslash Z$

In other words, these principles simply state that the combinatory rules may not contradict the directionality specified in the lexicon.

Their relation to syntax aside, the combinatory rules employed are the ones defined by Curry and Feys (1958):

- (135) a. $\mathbf{B}fg \equiv \lambda x. f(gx)$
 b. $\mathbf{T}x \equiv \lambda f. fx$
 c. $\mathbf{S}fg \equiv \lambda x. fx(gx)$

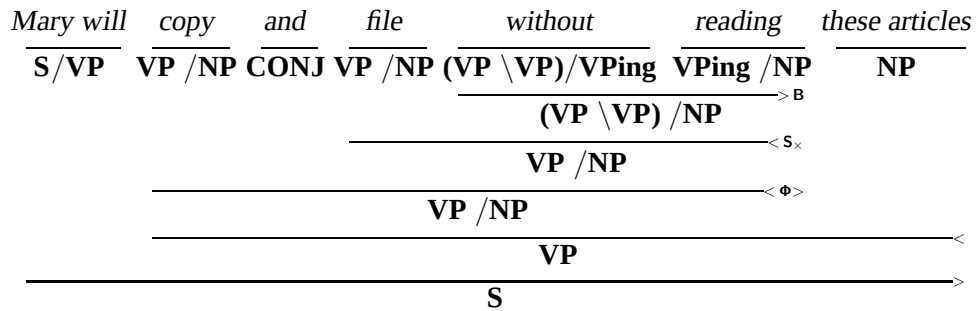


Figure A.7: Backward crossed substitution example

The combinator **B** composes a function f with its argument g before g has applied to its own argument. The result is a new function that applies its argument to the embedded function g . The combinator **T** turns an argument x into a function whose argument f is a function that x applies itself to. The combinator **S** is similar to **B**, except that the function it creates applies to its argument x to both f and g (Baldrige, 2002, p.24).

A.3 Resource Control

Recently CCG has been modified to implement a lexically controlled application of combinatory rules. Following Baldrige (2002), function categories specified in the lexicon may restrict the combinatory rules applicable to them via slashes typed with four basic modalities: $\star$, $\times$, $\diamond$, and $\cdot$. $\star$ modality is the most restricted and allows only the basic applicative rules. $\diamond$ permits order preserving associativity, $\times$ allows limited permutation, and $\cdot$ is the most permissive, allowing all rules.

For example, a lexical category that replaces the conjunction rule (which isn't one of Curry's combinators) can be expressed as follows:

$$(136) \text{ and} := (X \backslash_{\star} X) /_{\star} X$$

The $\star$ modality on the slashes of this category means that the only rules permitted are the functional application rules.

With the exception of type raising, the rules that are permitted by these modalities should be self explanatory. Type raising is redefined in the following way:

(137) Type Raising (**T**)

$$a. \quad X : a \implies_T T /_i (T \backslash_i X) : \lambda f. fa$$

$$b. \quad X : a \implies_T T \backslash_i (T /_i X) : \lambda f. fa$$

where X is an argument type.

The subscript i on the slashes mean that they both have the same modality as whatever function the new type raised category applies to.

A.4 Information Structure

Compatible with the assumption of monostratality and the linguistic trend of minimalism, Steedman (2000a, 2000b) defines a single level of linguistic representation, which is called

‘Information Structure’, and it subsumes the predicate argument structure of CCG, augmented with the notions of theme/rheme and background/kontrast. In this thesis, we shall be frequently referring to concepts from Steedman’s theory of English Information Structure (Steedman, 2000a, 2000b), and we shall be using Steedman’s terminology, with the exception that when we refer to the previous literature we shall be using the terminology assumed by the referred paper⁷. A summary of Steedman’s theory of English Information Structure can be found in Appendix (D).

⁷ Steedman, and Korbayova (2003) notes that the terminology describing Information Structure and its semantics is both diverse and under-formalized, but all definitions seem to make at least one of the following distinctions: (i) a “topic/comment” or “theme/rheme” distinction between the part of the utterance that relates it to the discourse purpose, and the part that advances the discourse. (ii) a “background/kontrast” or “given/new” distinction, between parts of the utterance (more specifically, words), which contribute to distinguishing its actual content from alternatives the context makes available. Özge and Bozşahin (2006) also presents a literature review before presenting their own theory of Turkish Information Structure, which adopts the terminology assumed by Steedman. Also note that, the ‘default discourse structure’ that is traditionally assumed in Turkish linguistics does not seem to make the distinction of (i) and (ii) above. The concept *topic* can be understood to be referring to idea ‘the part that relates the current utterance to the discourse purpose’, but its meaning is also overloaded in that it also carries the sense ‘old information’. The concept *focus* is also similarly overloaded: it means both ‘new information’ and ‘the part that advances the discourse’, without distinguishing between these two concepts. The term *background* is being used with the sense ‘old information’. We shall avoid this terminology confusion by adopting the terminology used by Mark Steedman (Steedman 2000a, 2000b), but when referring to the previous studies of quantifier scope ambiguity in Turkish, we shall be using the terms used by the referred paper(s), which can be understood as above.

APPENDIX B. QUANTIFYING-IN AND STORAGE METHODS

This chapter contains a summary of Blackburn and Bos (2005)’s introduction to Quantifying-In and Storage Methods; included here for convenience. This chapter is not intended to be a detailed introduction to Quantifying-In and Storage methods. For details, please see Blackburn and Bos (2005), which is one of the most readable books on the subject.

B.1 Quantifying-In

The basic idea is that instead of directly composing syntactic entities with quantifying noun phrase we are interested in, we are permitted to choose an ‘indexed pronoun’ and to combine the syntactic entity with the index pronoun instead. Intuitively, such indexed pronouns are ‘placeholders’ for the quantifying noun phrase. When this placeholder is at an high enough position in the tree to give us the scoping we are interested in, we are permitted to replace it by the quantifying NP of interest.

As an example, let’s consider how to analyse *Every boxer loves a woman*. Here is the first part of the tree we need:

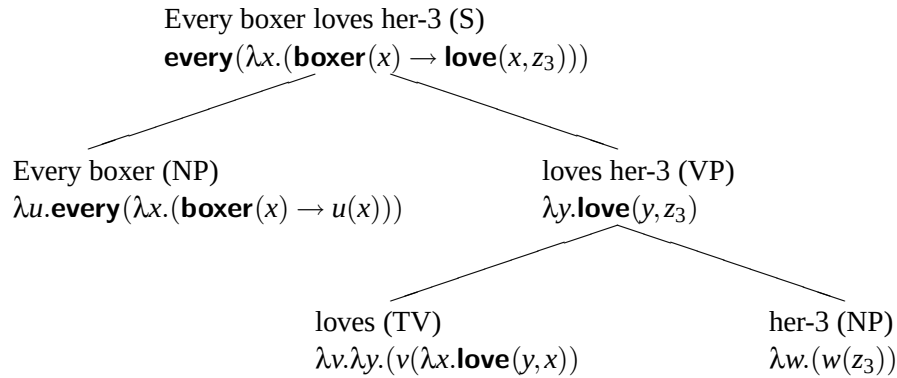


Figure B.1: Derivation for *Every boxer loves her-3*

Instead of combining *loves* with the quantifying term *a woman*, we have combined it with the placeholder pronoun *her-3*. This pronoun bears an *index*, namely the numeral ‘3’. The placeholder pronoun is associated with a ‘semantic placeholder’, namely $\lambda w.w(z_3)$. As we shall see, it is the semantic placeholder that does most of the real work. Note that the

pronoun's index appears as a subscript on the free variable in the semantic placeholder. From a semantic perspective, choosing an index pronoun really amounts to opting to work with the semantic placeholder (instead of the semantics of the quantifying NP) and stipulating which free variable the semantic placeholder should contain.

To be able to ensure that *a woman* out-scopes *every boxer*, we have delayed introduction of *a woman* into the tree. But *every boxer* is now firmly in place, we replace *her-3* by *a woman* and get the desired scoping relation. Predictably, there is a rule that lets us do this: given a quantified NP, and a sentence containing a placeholder pronoun, we are allowed to construct a new sentence by substituting the quantifying NP for the placeholder. In short, we are allowed to extend the previous tree as follows:

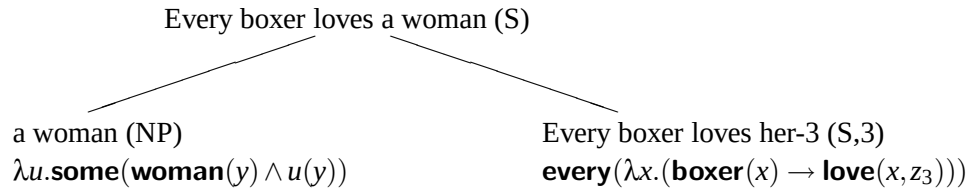


Figure B.2: Derivation for *Every boxer loves a woman*

What we intend to achieve semantically is that we want the formula $\text{some}(\lambda y.(\text{woman}(y) \wedge \text{every}(\lambda x.(\text{boxer}(x) \rightarrow \text{love}(y, x))))$ to be assigned to the parent node. To be able to do so, we want *a woman* to take wide scope over *Every boxer* semantically. For that reason, we should use the semantic representation associated with *a woman* as a function and apply it to the semantic representation of *Every boxer loves her-3*. This is because right at the bottom of the tree we have used the semantic placeholder $\lambda w.(w(z_3))$. When we raised this placeholder up the tree using functional application, we were essentially ‘recording’ what the semantic representation of *a woman* would have encountered if we had used it directly. The formula $\text{every}(\lambda x.(\text{boxer}(x) \rightarrow \text{love}(x, z_3)))$ is what we have at the end of this process. When we are ready to use it, we lambda abstract with respect to z_3 and use the resulting expression as an argument to the semantic representation of *a woman*. β -reduction will give us the result as in the figure (B.3).

As seen, Montague’s method requires additional syntactic rules to derive the scoping relations needed, such as introducing placeholder pronouns and for eliminating placeholder pronouns in favour of quantifying noun phrases. The grammar rules are there to tell about syntactic structure – but now we are using them to manipulate mysterious-looking placeholder

entities in a rather ad-hoc looking attempt to reduce scope issues to syntax.¹

The resulting grammar is not monotonic, because the introduced syntactic rules that introduce placeholder pronouns and later eliminate them are effectively an implementation the idea of movement, and such transformations are not allowed if we are to obey the principle of strict compositionality.

B.2 Storage Methods

Storage methods are a modular way of handling the problem of quantifier scope ambiguities. These methods make use of the basic ideas that Montague had proposed as described above, but they separate the issue of grammar and the issue of quantifier scope. They are also the first examples of more advanced underspecification based methods.

B.2.1 Cooper Storage

Cooper Storage was developed by Robin Cooper (1975, 1979, 1983). Cooper Storage is historically the first method involving underspecification intended to solve the problem of quantifier scope ambiguity. Like Montague’s approach Cooper Storage also has a two stage semantic construction, making it incompatible with the monotonic theories of grammar which require semantic representations to be built once and not to be modified later, but unlike Montague’s solution, Cooper Storage allows semantic construction to be handled in a modular way, without postulating extra rules that are needed to be added to the grammar. The idea is associate each node of a parse tree with a store, which contains a core semantic representation in addition to the quantifiers associated with nodes lower in the tree. After a sentence is parsed, the store is used to generate the scoped representations. Therefore it is a two stage process: The first step is to parse the sentence and the second step is to use the quantifier

¹ While discussing Steedman’s *skolemization*, we shall see that it is possible to reduce scope issues to syntax without such ad-hoc tricks.

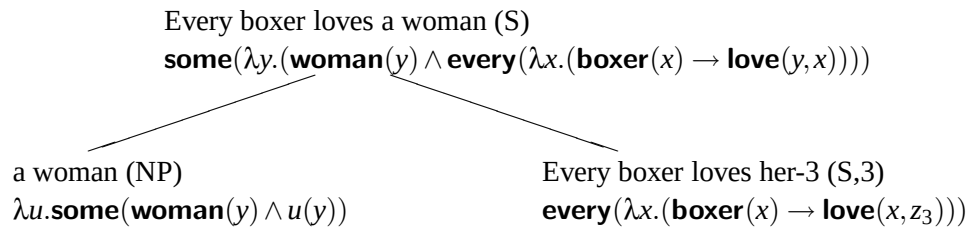


Figure B.3: Derivation for *Every boxer loves a woman*

storage information and the core semantic representations that has been built during parsing to generate the final semantic representation. The order in which the stored quantifiers are retrieved from the store and combined with the core representation determines the different scope assignments.

Formally, a store is a n-place sequence which is assigned to every parse tree node. The first item of the sequence is the core semantic representation and it's simply a lambda expression. The subsequent items in the sequence (if any) are pairs (β, i) , where β is the semantic representation of an NP (that is an another lambda expression), and i is the index. An index is a label which picks out a free variable in the core semantic representation. This index has the same purpose as the indexes that are used in Montague's solution. The pairs (β, i) are called indexed binding operators. The semantic construction with the data structure described above is implemented as follows: We start with the quantified noun phrases (that is, the noun phrases that contain a determiner). Quantified noun phrases add new information to the store, and other sorts of NPs do not. In other words, with quantified noun phrases we are free to use the following rule, but for other NP's the rule below does not apply:²

Definition 5 (Cooper Storage Rule) *If a store $\langle \phi, (\beta, j), \dots, (\beta', k) \rangle$ is a semantic representation of a quantified NP, then the store $\langle \lambda w.(w@z_i), (\phi, i), (\beta, j) \dots (\beta', k) \rangle$, where i is some unique index, is also a representation for that NP.*

At this point, it is important to notice that the index associated with ϕ is identical with the subscript on the free variable in $\lambda w.(w@z_i)$.

During the parsing of a sentence, if we see a quantified NP, we have a choice: We can either pass on $\langle \phi, (\beta, j), \dots, (\beta', k) \rangle$ straight up the tree, or we can decide to pass on the expression $\langle \lambda w.(w@z_i), (\phi, i), (\beta, j) \dots (\beta', k) \rangle$ instead.

It is important to note that application of the rule is not recursive. It offers a choice once: either use and pass on the ordinary representation or use and pass on the modified one. We are not allowed to apply the rule again to the modified representation.

The following example demonstrates how this works on the sample sentence *Every student wrote a program*.

To explain the first step of cooper storage method, we need to explain how the representations associated with each node of the tree above were assigned. We know that the lambda

² Note that we use the notation $\alpha@ \beta$ to make functional application explicit to avoid confusion with other usages of parentheses.

representation associated with an NP like ‘a program’ is $\lambda u.\text{some}(\lambda y.(\text{program}(y) \wedge u@y))$.

With the data structure that cooper storage uses, that corresponds to the 1 place store:

$$< (\lambda u.\text{some}(\lambda y.(\text{program}(y) \wedge u@y)) >$$

While that’s a legitimate interpretation for the NP ‘a program’, we have a second choice due to the cooper storage rule, and the rule states that we can introduce a new free variable, namely z_7 in the tree above, and use the expression $\lambda w.(w@z_7)$ as our new core representation, and keep the previous form of the lambda expression in our store, but by remembering the index of the newly introduced free variable associated with it:

$$(\lambda u.\text{some}(\lambda y.(\text{program}(y) \wedge u@y)), 7)$$

As a result, the semantic representation $< \lambda w.(w@z_7), (\lambda u.\text{some}(\lambda y.(\text{program}(y) \wedge u@y)), 7 >$ is associated with the node corresponding the phrase ‘a program’. The same rule application is also used for the NP ‘Every student’, and the new free variable z_7 has been associated with it.

To calculate the semantic representation for higher level nodes, the first step of cooper storage method uses functional application followed by β -conversion as usual, but only for the first element of each sequence functional application is used, and the rest of the sequences are simply appended to form the new sequence. That is, the first element of each sequence is the active part, and the rest – which is also called the freezer - is simply a list that keeps track of the original expressions that were replaced by the cooper storage rule.

The second part of cooper storage, which is called retrieval, is the task of generating ordinary first order representations from the result of the analysis above. Retrieval removes one of the index binding operators from the freezer list and combines it with the core representation to form a new core representation. (If the freezer is empty, then the store associated with the

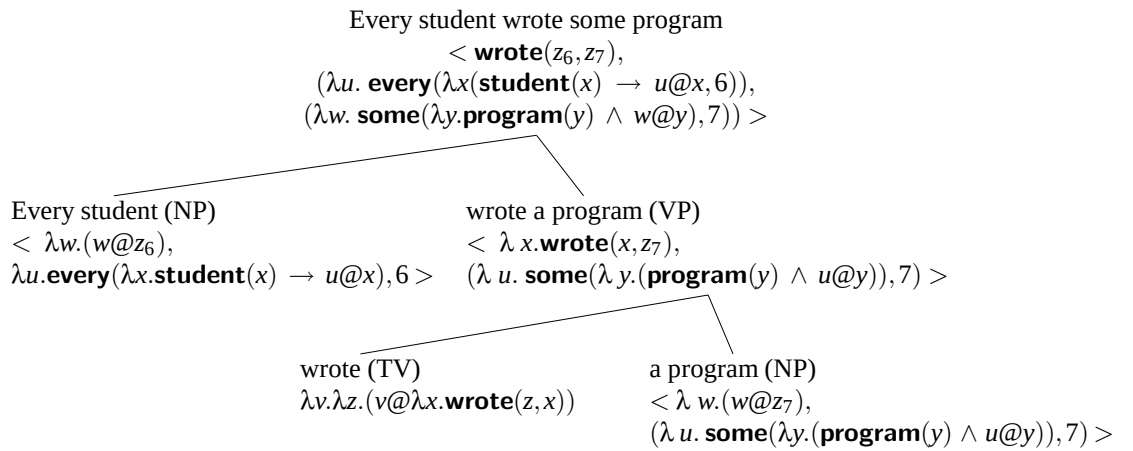


Figure B.4: Derivation for *Every student wrote a program*

S node must already be a 1-place sequence that contains a first order formula.) The retrieval process continues until all the indexed binding operators are used. The last core representation obtained in this way will be the desired semantic formula with the appropriate scoped semantic representation. The mentioned combination operation during retrieval involves the following: Assume that the retrieval process removed a binding operator (β, i) from the store $\langle \phi, (\beta, i), \dots (\beta', k) \rangle$. The process first lambda abstracts the core semantic representation ϕ with respect to the index variable z_i , obtains an expression of the form $\lambda z_i. \phi$. Then the process functionally applies β to $\lambda z_i. \phi$. The result is the new core semantic representation which replaces the old one, ϕ .

Another way of defining the retrieval process is the following:

Definition 6 (Cooper Retrieval Rule) *Let s_1 and s_2 to be (possibly empty) sequences of binding operators and the store $\langle \phi, s_1, (\beta, i), \dots, s_2 \rangle$ is associated with an expression of category S , then the store $\langle \beta @ \lambda z_i. \phi, s_1, s_2 \rangle$ is also associated with this expression.*

As an example, consider how to generate the first order formulas from the storage representation of the sample sentence every student wrote a program, which is:

$\langle \mathbf{wrote}(z_6, z_7),$
 $(\lambda u. \mathbf{every}(\lambda x. (\mathbf{student}(x) \rightarrow u @ x, 6)),$
 $(\lambda w. \mathbf{some}(\lambda y. (\mathbf{program}(y) \wedge w @ y)), 7) \rangle.$

This store contains two indexed binding operators, namely, $(\lambda u. \mathbf{every}(x(\mathbf{student}(x) \rightarrow u @ x, 6))$ and $(\lambda w. \mathbf{some}(\lambda y. (\mathbf{program}(y) \wedge w @ y)), 7) \rangle$. The retrieval rule allows either of them to be removed first (without any order preference) and be combined with the core representation, which is $\mathbf{wrote}(z_6, z_7)$. Let's assume that we have chosen the indexed binding operator that corresponds to the NP 'every student', namely $(\lambda u. \mathbf{every}(x(\mathbf{student}(x) \rightarrow u @ x, 6))$. Then, the retrieval rule requires the following store the replace the old one:

$\langle (\lambda u. \mathbf{every}(\lambda x(\mathbf{student}(x) \rightarrow u @ x) @ (\lambda z_6. \mathbf{wrote}(z_6, z_7))),$
 $(\lambda w. \mathbf{some}(\lambda y(\mathbf{program}(y) \wedge w @ y)), 7) \rangle$

Which can be simplified using β – conversion and the following representation is obtained: $\langle \lambda x(\mathbf{student}(x) \rightarrow \mathbf{wrote}(x, z_7)), (\lambda w. \mathbf{some}(\lambda y(\mathbf{program}(y) \wedge w @ y)), 7) \rangle$

There is another indexed binding operator (namely, $\lambda w. \mathbf{some}(\lambda y(\mathbf{program}(y) \wedge w @ y)), 7)$ in the store above, and when it is retrieved in the same way we obtain:

$\langle \mathbf{some}(\lambda y(\mathbf{program}(y) \wedge \mathbf{every}(\lambda x(\mathbf{student}(x) \rightarrow \mathbf{wrote}(x, y)))) \rangle$ which is the same as the first order formula in (4b), where 'a program' has scope over 'every student'.

The other possible reading, (4a) is obtained when the indexed binding operators are retrieved in the other possible order.

B.2.2 Keller Storage

Cooper Storage allows a great deal of freedom in retrieving information from the store. It allows quantifiers to be obtained in any order, and as explained below, this level of freedom is not always safe. (Blackburn, P. & Bos, J. (2005), 122.) In the above example we have not run into a problem, but we have examined only one kind of scope ambiguity: a sentence containing a transitive verb with quantifying NPs in subject and object position. However, there are other syntactic structures where quantifier scope ambiguities occur. One example is relative clauses (138), and another example is prepositional noun phrases (139).

(138) Every piercing that is done with a gun goes against the entire idea behind it.

(Ambiguous: Every piercing that is done with (possibly different) guns or every piercing that is that with the same gun.)

(139) Mia knows every owner of a hash bar.

(Ambiguous: Mia knows all owners of (possibly different) hash bars or Mia knows all owners that own one and the same hash bar.)

Both examples contain nested NPs and this is where Cooper Storage is not good enough: it completely ignores the hierarchical structure of the NPs. Let's shall examine the problems caused by (139). When the first stage of Cooper Storage is completed, the following store is obtained:

$$\begin{aligned} < \text{know}(mia, z_2), \\ & (\lambda u. \text{every}(\lambda y(\text{owner}(y) \wedge of(y, z_1) \rightarrow u@y)), 2), \\ & (\lambda w. \text{some}(\lambda x(\text{hbar}(x) \wedge w@x), 1)) > \end{aligned}$$

There are two indexed binding operators, and therefore two ways to perform retrieval: The first possibility is to remove the first one first and the second one next. The other possibility is to remove the second one first and the first one next. Assuming that we have chosen to remove the first one first, we have the following sequence of retrieval operations (functional application and β -reduction operations are included for clarity):

$$\begin{aligned}
& (\lambda u. \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow u@y)]) @ (\lambda z_2. \text{know}(\text{mia}, z_2)) \\
& \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow (\lambda z_2. \text{know}(\text{mia}, z_2)) @ y)] \\
& \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow \text{know}(\text{mia}, y))] \\
& (\lambda w. \text{some}[\lambda x. (\text{hbar}(x) \wedge w@x)] @ (\lambda z_2. \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow \text{know}(\text{mia}, y))])) \\
& \text{some}[\lambda x. (\text{hbar}(x) \wedge (\lambda z_2. (\text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow \text{know}(\text{mia}, y))])) @ x)] \\
& \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, x) \rightarrow \text{know}(\text{mia}, y))])]
\end{aligned}$$

The result states that there is a hash bar which Mia knows every owner of. If we had chosen the other option and had removed the second quantifier first from the store, we would have the following sequence of retrieval operations.

$$\begin{aligned}
& (\lambda w. \text{some}[\lambda x. (\text{hbar}(x) \wedge w@x)] @ (\lambda z_1. \text{know}(\text{mia}, z_1)) \\
& \text{some}[\lambda x. (\text{hbar}(x) \wedge (\lambda z_1. \text{know}(\text{mia}, z_1)) @ x)] \\
& \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{know}(\text{mia}, z_1))] \\
& (\lambda u. \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow u@y)] @ (\lambda z_2. \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{know}(\text{mia}, z_2))])) \\
& \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow (\lambda z_2. \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{know}(\text{mia}, z_2))])) @ y)] \\
& \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow (\text{some}[\lambda x. (\text{hbar}(x) \wedge \text{know}(\text{mia}, y))]))]
\end{aligned}$$

The problem with the result above is that it still contains a free variable, namely z_1 . The source of the problem is that Cooper Storage ignores the hierarchical structure of the NPs. The sub-NP ‘a hash bar’ contributes the free variable z_1 . However, this free variable does not exist in the core representation $\text{know}(\text{mia}, z_2)$. When the NP ‘every owner of a hash bar’ is processed during the first stage of Cooper Storage, the variable z_1 is removed from the core representation and put on the freezer list. For that reason, lambda abstracting the core representation with respect to z_1 does not take into account the contribution that z_1 makes, because z_1 makes its contribution indirectly via the stored quantifier. The result is fine if we retrieve this quantifier first, because it has the effect of restoring z_1 to the core representation, but if we use other retrieval option we fail to obtain a correct representation for the sentence.

Keller Storage (due to Bill Keller) solves this problem by allowing nested stores. The nesting structure of these stores is intended to keep track of the nesting structure of NPs. To allow this, the following rule replaces the previously given Cooper Storage rule.

Definition 7 (Keller Storage Rule) *If the (nested) store $\langle \phi, s \rangle$ is an interpretation for an NP, then the nested store $\langle \lambda u. (u@z_i), (\langle \phi, s \rangle, i) \rangle$, for some unique index i , is also an interpretation for this NP.*

The following example (Figure B.5) shows how this works:

Definition 8 (Keller Retrieval Rule) *Let s, s_1 and s_2 be (possibly empty) sequences of binding operators. If the nested store $\langle \phi, s_1, (\langle \beta, s \rangle, i), s_2 \rangle$ is an interpretation for an expression of category S , then $\langle \beta @ \lambda z_1. \phi, s_1, s, s_2 \rangle$ is also an interpretation of the same expression.*

Keller Retrieval Rule ensures that any binding operators stored while processing β become accessible for retrieval only after β itself has been retrieved. The following example –which is the same example that wasn’t correctly analysed with Cooper Storage– shows how Keller Retrieval works in practice:

$$\begin{aligned} & \langle \text{know}(\text{mia}, z_2), \\ & \quad (\langle \lambda u. \text{every}[\lambda y. (\text{owner}(y) \wedge \text{of}(y, z_1) \rightarrow u@y)], \\ & \quad \quad (\langle \lambda w. \text{some}[\lambda x. (\text{hbar}(x) \wedge w@x)] \rangle, 1) \rangle, 2) \rangle \end{aligned}$$

In this example there is only one way to perform retrieval: removal of the universal quantifier, followed by removal of the existential quantifier. Since this is the only possibility, the unwanted reading generated by Cooper Storage is not produced.

The second issue to understand about Keller Storage is how it generates the reading where Mia knows all owners of possibly different hash bars. Remember that, the application of the storage rule is optional. All that we need to do to generate the mentioned reading is to choose not to apply the storage rule for the NP ‘a hash bar’ during construction of the parse tree, as shown in Figure B.6.

For the full sentence ‘Mia knows every owner of a hash bar’ this leads to the following semantic representation:

$$\begin{aligned} & \langle \text{know}(\text{mia}, z_2), \\ & \quad (\langle \lambda u. \text{every}[\lambda y. (\text{owner}(y) \wedge \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{of}(y, x))] \rightarrow u@y)] \rangle, 2) \rangle \end{aligned}$$

There is only one binding operator in the store, and retrieving it generates the correct logical formula:

$$\text{every}[\lambda y. (\text{owner}(y) \wedge \text{some}[\lambda x. (\text{hbar}(x) \wedge \text{of}(y, x))] \rightarrow \text{know}(\text{mia}, y))]$$

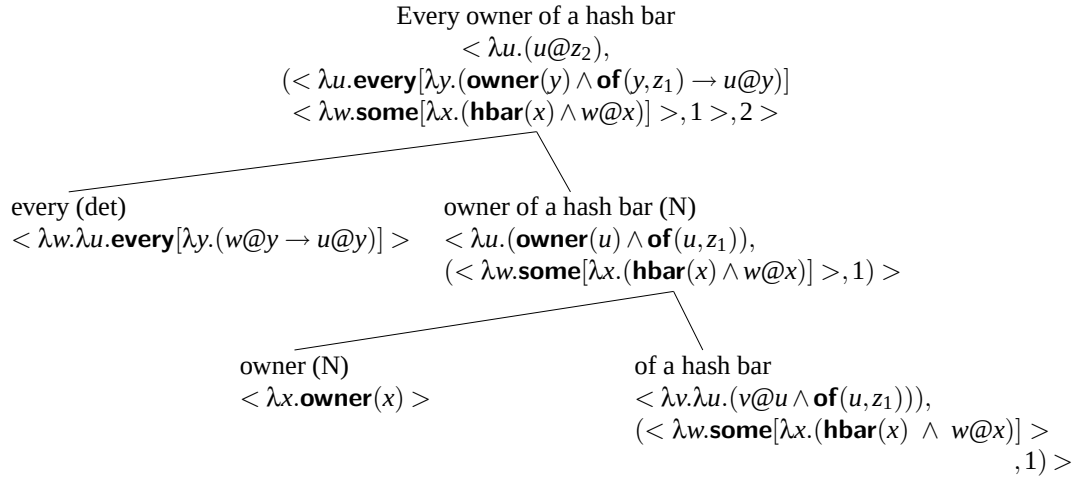


Figure B.5: Derivation for *Every owner of a hash bar (NP)*

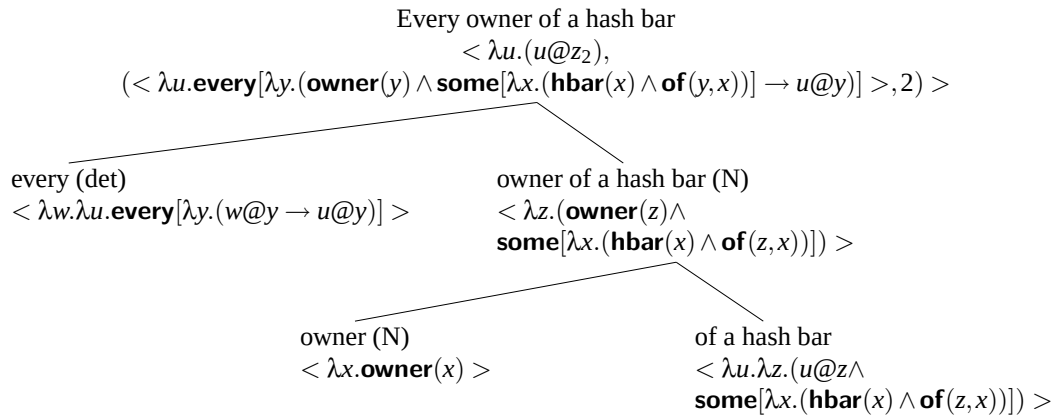


Figure B.6: Derivation for *Every owner of a hash bar (NP)*

APPENDIX C. SURFACE COMPOSITIONAL SCOPE

ALTERNATION (STEEDMAN, 1999, 2000a, 2005)

This chapter briefly introduces Steedman’s 2005 account of quantifier scope ambiguity in CCG. There is no way to describe the theory in adequate detail given the space and time limitations, but since the theory is highly relevant to the subject matter of this thesis, it is summarized here without the details.

Steedman’s analysis of quantifier scope ambiguity is an example of an underspecification based theory based on the idea of *skolemization* that unlike other underspecification based methods that we have examined (i.e. the storage methods) it does not make use of a retrieval step, and the required logical representations are built as a result of a single stage process, which is parsing. But before explaining the way Steedman accounts for quantifier scope ambiguity, we should first explain the term skolemization as it is a mathematical term which can found in purely mathematical contexts whose subject matter is not linguistics. The following information about skolemization can be found in textbooks about formal logic:

In first order logic, the term skolemization refers to the elimination of existential quantifiers, so that a new *equisatisfiable* formula (i.e. a formula which can be satisfied under the same conditions) is obtained. Skolemization is an application of the equivalence:

$$(140) \quad \forall x \exists y M(x, y) \leftrightarrow \exists f \forall x M(x, f(x))$$

Where M refers to the matrix of the quantificational formula. The essence of skolemization is the observation that if a formula in the form:

$$(141) \quad \forall x_1 \dots \forall x_n \exists y M(x_1, \dots, x_n, y)$$

is satisfiable in some model, then for every possible assignment of the variables $x_1 \dots x_n$ that makes $M(x_1, \dots, x_n)$ true, there must exist some function $f(x_1, \dots, x_n)$ which makes

$$(142) \quad \forall x_1 \dots \forall x_n M(x_1, \dots, x_n, f(x_1, \dots, x_n))$$

true. The function f is called the *skolem function*.

Variables bound by existential quantifiers which are not inside the scope of universal quantifiers can simply be replaced by constants: that is, $\exists x(x < 3)$ can be changed to $(c < 3)$, in which c is a constant.

When the existential quantifier is inside a universal quantifier, the bound variable must be replaced by a skolem function of the variables bound by universal quantifiers. Thus $\forall x(x = 0 \wedge \exists y(x = y + 1))$ can be written as $\forall x(x = 0 \wedge x = f(x) + 1)$.

Before returning back to the Steedman's account of quantifier scope ambiguity, there is one more point that needs to be explained, and that is to take the idea of direct surface composition as literally as possible, and to see what happens when we try to account for quantifier scope ambiguity using combinatorics of syntax over the lexical elements. If we take that approach, the following are the appropriate CCG categories for the quantifiers *every* and *some*:

- (143) *every* $:= (T / (T \backslash NP)) / N$ $: \lambda p. \lambda q. \forall x(px \rightarrow qx)$
every $:= (T \backslash (T / NP)) / N$ $: \lambda p. \lambda q. \forall x(px \rightarrow qx)$
some $:= (T / (T \backslash NP)) / N$ $: \lambda p. \lambda q. \exists x(px \wedge qx)$
some $:= (T \backslash (T / NP)) / N$ $: \lambda p. \lambda q. \exists x(px \wedge qx)$

This approach links syntactic derivation and scope as simply and as directly as possible, but has an effect of implying that in sentences where there is no syntactic ambiguity, there should be no scope ambiguity either, which is obviously wrong. Nevertheless, have a look at the following example (which is from Steedman (2000a)) to see how these category assignments work in practice, as Steedman keeps one of them (the one assigned to the quantifier 'every') and eliminates the need for the other (the one assigned to the existential quantifier 'some') using skolemization.

- (144)
- | | | | | |
|---|--------------------------------|---|--|--------------------------------|
| <i>Every</i> | <i>boy</i> | <i>admires</i> | <i>some</i> | <i>saxophonist.</i> |
| $\frac{(T / (T \backslash NP)) / N}{: \lambda p. \lambda q. \forall y(py \rightarrow qy) : \lambda x. boy'x}$ | $\frac{N}{: \lambda x. boy'x}$ | $\frac{(S \backslash NP) / NP}{: \lambda x. \lambda y. adm'xy}$ | $\frac{(T' \backslash (T' / NP)) / N}{: \lambda p. \lambda q. \exists x(px \wedge qx) : \lambda x. sax'x}$ | $\frac{N}{: \lambda x. sax'x}$ |
| $\xrightarrow{\quad}$ | | | $\xrightarrow{\quad}$ | |
| $\frac{T / (T \backslash NP)}{: \lambda q. \forall y(boy'y \rightarrow qy)}$ | | | $\frac{(T' \backslash (T' / NP))}{: \lambda q. \exists x(sax'x \wedge qx)}$ | |
| | | $\xrightarrow{\quad}$ | | |
| | | $\frac{S \backslash NP}{: \lambda y. \exists x(sax'x \wedge adm'xy)}$ | | |
| | | $\xrightarrow{\quad}$ | | |
| $S : \forall y(boy'y \rightarrow \exists x(sax'x \wedge adm'xy))$ | | | | |

A second reading is available and CCG is able to generate that reading due to the fact that the availability functional composition rule **B** which generates a nontraditional syntactic constituent¹ in the example below and correctly generates the desired semantic reading:

¹ which in the example effectively corresponds to the context free rule $VP \rightarrow NP_{subject}V$

$$\begin{array}{c}
(145) \quad \begin{array}{ccccc}
\text{Every} & \text{boy} & \text{admires} & \text{some} & \text{saxophonist.} \\
\hline
\text{(T/(T\NP))/N} & \text{N} & \text{(S \NP)/NP} & \text{(T'/(T'/NP))/N} & \text{N} \\
: \lambda p. \lambda q. \forall y (py \rightarrow qy) & : \lambda x. \text{boy}'x : \lambda x. \lambda y. \text{adm}'xy & : \lambda p. \lambda q. \exists x (px \wedge qx) & : \lambda x. \text{sax}'x & \\
\hline
\text{T/(T\NP)} & & & \text{(T'/(T'/NP))} & \\
: \lambda q. \forall y (\text{boy}'y \rightarrow qy) & & & : \lambda q. \exists x (\text{sax}'x \wedge qx) & \\
\hline
& \text{S /NP} & & & \\
: \lambda x. \forall y (\text{boy}'y \rightarrow \text{adm}'xy) & & & & \\
\hline
& \text{S : } \exists x (\text{sax}'x \wedge \forall y (\text{boy}'y \rightarrow \text{adm}'xy)) & & &
\end{array}
\end{array}$$

The fact that the above example has been correctly accounted due to the availability of the nontraditional constituent does not mean that this approach is generic enough to account every possible scope ambiguity, as exemplified by the following examples:

(146) Some saxophonist, every boy admires. [Topicalization]

(147) Every boy admires, and every boy detests, some saxophonist. [Object right node raising]

According to Steedman (2000a), both of these sentences have a narrow scope reading in which every individual has some attitude toward some saxophonist, but not necessarily the same saxophonist. But the following derivation implies that every boy admires the same saxophonist. Steedman (2005) states that the universal quantifiers every and each can take scope over c-commanding indefinites, that is, there is another reading where every boy out scopes some saxophonist but this is derivationally impossible, even with the nontraditional syntactic constituents that CCG generates:

$$\begin{array}{c}
(148) \quad \begin{array}{ccccc}
\text{Some} & \text{saxophonist} & \text{every} & \text{boy} & \text{admires.} \\
\hline
\text{(T/(T\NP))/N} & \text{N} & \text{(T/(T\NP))/N} & \text{N} & \text{(S \NP)/NP} \\
: \lambda p. \lambda q. \exists x (px \wedge qx) & : \lambda x. \text{sax}'x & : \lambda p. \lambda q. \forall y (py \rightarrow qy) & : \lambda x. \text{boy}'x & : \lambda x. \lambda y. \text{adm}'xy \\
\hline
\text{T/(T\NP)} & & & \text{(T/(T\NP))} & \\
: \lambda q. \exists x (\text{sax}'x \wedge qy) & & & : \lambda q. \forall y (\text{boy}'y \wedge qy) & \\
\hline
& & \text{S /NP} & & \\
& & : \lambda x. \forall y (\text{boy}'y \rightarrow \text{adm}'xy) & & \\
\hline
& \text{S : } \exists x (\text{sax}'x \wedge \forall y (\text{boy}'y \rightarrow \text{adm}'xy)) & & &
\end{array}
\end{array}$$

Similarly, $\exists x (\text{sax}'x \wedge (\forall y (\text{boy}'y \rightarrow \text{adm}'xy) \wedge \forall y (\text{boy}'y \wedge \text{detests}'xy)))$ is the only available reading for the second sentence using this approach, but as we shall later there are more.

While this evidence seems to be against the hypothesis of direct compositionality (at least as far as scope ambiguities are concerned, that is.), there are also strong arguments in favor of it. For example, the sentence *Every boy admires, and every boy detests, some saxophonist*

has an interesting property first observed by Geach (1972) which makes it seem that scope phenomena are strongly restricted by syntax. Although the sentence has a reading where all the boys and girls have some feeling towards the same saxophonist, and another reading where all feelings are directed at possibly different saxophonists, it does not have a reading where the saxophonist has wide scope with respect to *every boy* but narrow scope with respect to *every girl*, that is, where all the boys all admire the same saxophonist but all girls detest possibly different saxophonists. There is not a reading involving separate wide scope saxophonists respectively taking scope over boys and girls, either – for example, here the boys all admire the saxophonist A but the girls all detest saxophonist B.

Such observations require extra stipulation such as a ‘parallelism constraint’ when we attempt to explain them using theories that allow movement at the level of logical form as in Montague’s approach and its underspecification based followers. For example, it is not clear why some saxophonist should have the same scope in both conjuncts, the reason being the fact that the rule of conjunction duplicates the NP *some saxophonist* in the phrase structure at LF, and both of these NPs can move independently.

Such concerns indicate a need for a theory of scope ambiguity that explains various observed scope ambiguities that arise in natural language, but without allowing room for introducing arbitrary rules that introduce more freedom than necessary. That is, the theory should explain the observed phenomena without introducing arbitrary rules, and should not generate unwanted readings, again without requiring arbitrary rules. Skolemization introduced in Steedman (1999, 2000a, 2005) is a theory intended to achieve that.

C.1 Syntactic Constraints on Quantifier Scope in English

The data that the theory is intended to explain is as follows (quoted from Steedman (2005)):

1. The universal quantifiers *every* and *each* can take scope over c-commanding indefinites, as in (149a). Such scope inversion of universals is both unbounded² as in (149b) and sensitive to island constraints, as in (149c,d), where scope alternation over the matrix subject is inhibited, parallel to the extractions in (150).

² Beyond the local domain of the predicate.

- (149) a. Some member of our group attended every rally $(\forall\exists/\exists\forall)$
 b. Some member of our group proved that the candidate had attended every rally. $(\forall\exists/\exists\forall)$
 c. Some member of our group met a candidate who attended every rally. $(\#\forall\exists/\exists\forall)$
 d. Some member of our group said that every candidate will win.
- (150) a. The rally that some member of our group attended
 b. The rally that some member of our group proved that the candidate had attended
 c. #The rally that some member of our group met a candidate who attended
 d. #The candidate that some member of our group said that will win.

What is more, such quantifier ‘movement’ appears to be subject to a constraint reminiscent of William’s 1978 ‘Across-the-Board’ exception to the Coordinate Structure Constraint upon Wh-movement of Ross (1967), in examples like the following, as first noted by Geach (1972). (The scope possibilities of the sentence was explained above.)

- (151) Every boy admires, and every girl detests, some saxophonist.

2. Existential quantifiers like *some*, *a*, and *at least/at most/exactly three* appear to be able to take wide scope over unboundedly c-commanding quantifiers, and (unlike the universals) are not sensitive to island boundaries in this respect.

- (152) a. Exactly half of the boys in the class kissed some girl. $(\frac{1}{2}\exists/\exists\frac{1}{2})$
 b. Every member of our group attended some rally. $(\forall\exists/\exists\forall)$
 c. Every member of our group proved that the candidate had attended some rally. $(\forall\exists/\exists\forall)$
 d. Every member of our group met a candidate who attended some rally. $(\forall\exists/\exists\forall)$
 e. Every member of our group said that some candidate will win. $(\forall\exists/\exists\forall)$

3. However, existentials in general cannot truly invert scope, in the strong sense of distributing over a structurally-commanding existential:

- (153) a. Some member of our group attended exactly tree rallies. $(\#3\exists/\exists3)$
 b. Exactly half the boys in the class kissed tree girls. $(\#3\frac{1}{2}/\frac{1}{2}3)$

The theory is founded upon the idea of distinguishing true generalized quantifiers and other purely referential categories. For example, in order to capture the narrow scope object

reading for Geach's (object) right node raised sentence in whose CCG derivation the object must command everything else, a single non-quantificational interpretation of some saxophonist is proposed. The skolem terms that are introduced by inference rules like Existential Elimination in proof theories of first-order predicate calculus are of interest because they directly express dependency on other entities in the model, namely on the variables bound by the universal quantifiers.

As we had seen in the beginning of this section, Skolem terms are obtained by replacing all occurrences of a given existentially quantified variable by an application of a unique functor to all variables found by a universal quantifier in whose scope the existentially quantified variable falls. If there are no such universal quantifiers, then the Skolem term is a constant. Thus, the two interpretations of the sentence *everybody loves somebody* can be expressed as follows:

- (154) a. $\forall x[\text{person}'x \rightarrow (\text{person}'(sk'_{53}x) \wedge \text{loves}'(sk'_{53}x)x)]$
 b. $\forall x[\text{person}'x \rightarrow (\text{person}'(sk'_{95}x) \wedge \text{loves}'sk'_{95}x)]$

In (154a) sk'_{53} is a Skolem function. In (154b) sk'_{95} is a Skolem constant. (154a) means that every person loves the thing that the Skolem function sk'_{53} maps them onto – their own specific dependent person. (154b) means that every person loves the person identified by the Skolem constant sk'_{95} . The interesting thing about this alternative to the more usual logical forms below in (155) is that the two formulas are identical, apart from the details of the Skolem terms themselves, which capture the meaning distinction in terms of whether or not the referent of someone is dependent upon the individuals quantified over by everyone. Notice that the fact that x is a parameter of the Skolem function sk'_{53} directly models the quantificational dependency involved.

- (155) a. $\forall x[\text{person}'x \rightarrow \exists y[\text{person}'y \wedge \text{loves}'yx]]$
 b. $\exists y[\text{person}'y \rightarrow \forall x[\text{person}'x \rightarrow \text{loves}'yx]]$

Now that the need for the existential quantifier (*exists*) has been eliminated using skolemization, Steedman (2005) continues to argue that:

1. The only determiners in English that are associated with traditional Generalized Quantifiers, and take scope including inverse scope, distributing over structurally commanding indefinites as in (149) are the universals *every*, *each* and their relatives;
2. All indefinite determiners are associated with Skolem terms, which are interpreted *in situ* at the level of logical form (lf), forcing parallel interpretations in coordinate sentences like

(151).

3. The appearance of indefinites taking wide scope arises from flexibility as to which bound variables (if any) the Skolem term involves;

4. Indefinites never distribute over structurally commanding indefinites, because their interpretations are never quantificational.

Thus, the evidence (or data) presented in (149) ... (153) is explained. Steedman (2000a) and (2005) also mention independent support for the theory, due to so called the donkey sentences, whose analysis is reported to be puzzling and unsettled before introduction of skolem terms (which are therefore found to be useful in analyzing phenomena other than scope ambiguities).

C.1.1 Donkey Sentences

Sentences like the following has been investigated in modern semantic studies since Geach 1962, who attributes them to even earlier sources:

(156) Every farmer who owns a donkey_{*i*} feeds it_{*i*}. (Geach, 1962)

(157) Everybody who has a face mask wears it. (Steedman, 2005)

These sentences are interesting for the following reasons. The existence of preferred readings in which each person feeds the donkey he owns makes the pronoun *it* seem as though it might be a variable bound by an existential quantifier associated with a donkey. However, no purely combinatoric analysis in terms of classical quantifiers allow this, since the existential cannot both remain within the scope of the universal and come to *lf*-command the pronoun, as is required for true bound pronominal anaphora, of the kind exemplified by:

(158) Every man_{*i*} in the bar thinks that he_{*i*} is a genius.

Donkey sentences have been extensively analyzed in the literature (as cited in Steedman (2005)), it may seem that it is unlikely that there may be anything new to say about them, or any need for yet another account. However, the existing theories are pulled in different directions by a pair of problems called the proportion problem and the uniqueness problem, with which we shall be concerned later. dealing with these problems has caused very considerable complications to the theories, variously including recategorization of indefinites as universals, dynamic generalizations of the notion of scope itself, exotic varieties of pronouns

including choice-functional interpretations, local minimal situations, and various otherwise unmotivated syntactic transformations. Even if some or all of these accounts cover the empirical observations completely, there seems to be a room for a simpler theory.

The theory (Steedman, 2005) claims that in spite of the need for DRT-style dynamic semantics to capture the asymmetric processes of pronominal reference itself, the compositional semantics of sentences like (156) over such referents can be captured with standard statically-scoped models.

As cited in Steedman (2005), many researchers have pointed out that donkey pronouns look in many respects more like non-bound-variable or discourse-bound pronouns, in examples like the following, than like the bound variable pronoun in (158)

(159) Everyone who meets Monbodd_{*i*} likes him_{*i*}.

For example, the pronouns in (156) and (159) can be replaced by epithets, whereas true bound variable pronouns like that in (158) cannot :

- (160) a. Everyone who meets Monbodd_{*i*} likes the follow_{*i*}.
 b. Every farmer who owns a donkey_{*i*} feeds the lucky beast_{*i*}.
 c. *Every professor_{*i*} in the department thinks the old dear_{*i*} is a genius.

This observation suggests that the pronoun here is simply a discourse-bound pronoun, and that it is the donkey to which it refers in (156) that we should concentrate our attention on. In particular, we should consider the possibility that the latter (the donkey, that is) may translate as a referential (or referent-introducing) expression rather than as a generalized quantifier.

For that reason, Steedman proposes that ‘a donkey’ is a skolem term, to which the pronoun is simply discourse-anaphoric rather than bound variable anaphoric.

It is important to realize that an indefinite NP like ‘a donkey’ translates at predicate-argument structure as a Skolem term, to which the pronoun is simply discourse-anaphoric rather than bound-variable anaphoric.

It is also important to realize that the way this translation is done is different from standard skolemization of the kind just illustrated. Skolem terms in Steedman (2005) are elements of the logical form in their own right, initially unspecified as to their bound variables, if any. They appear as arguments of the verb translation, in order to bring them within the logical form scope of any quantifiers that may eventually determine their specification.

The latter requirement removes the need to separately introduce property predicates like *person*’*x* over Skolem terms, as in (161):

- (161) a. $\forall x[\text{person}'x \rightarrow (\text{person}'(\text{sk}'_{53}x) \wedge \text{loves}'(\text{sk}'_{53}x)x)]$
 a. $\forall x[\text{person}'x \rightarrow (\text{person}'\text{sk}'_{95} \wedge \text{loves}'\text{sk}'_{95}x)]$

We must instead associate such nominal properties with the Skolem term itself. We therefore write the underspecified translation of a donkey as $\text{skolem}'_n \text{donkey}'$. (The subscript n is a number uniquely assigned to the noun phrase from which the term originates, and distinct from any other occurrence of a donkey. When there is no other occurrence of a given noun phrase in the context, n may be suppressed.).

The noun property involved may also be arbitrarily complex. For example, to obtain the interpretation of the noun phrase a fat donkey in the sentence every farmer who owns a donkey feeds it, we must associate the property $\lambda y.\text{donkey}'y \wedge \text{fat}'y$ with the underspecified term. Such properties may recursively include other such referential terms, for example, a farmer who owns a donkey, some farmers who own a donkey, at most three farmers who own a donkey, and most farmers who own a donkey, which are represented as follows:

(162) a. a fat donkey	$\text{skolem}'(\lambda y.\text{donkey}'y \wedge \text{fat}'y)$
b. a farmer who owns a donkey	$\text{skolem}'(\lambda y.\text{farmer}'y \wedge \text{own}'(\text{skolem}'\text{donkey}'y))$
c. some farmers who own a donkey	$\text{skolem}'(\lambda y.\text{farmer}'y \wedge \text{own}'(\text{skolem}'\text{donkey}'y);$ $\lambda s. s > 1)$
d. at most three farmers who own a donkey	$\text{skolem}'(\lambda y.\text{farmer}'y \wedge \text{own}'(\text{skolem}'\text{donkey}'y);$ $\lambda s. s \leq 3)y)$
e. most farmers who own a donkey	$\text{skolem}'(\lambda y.\text{farmer}'y \wedge \text{own}'(\text{skolem}'\text{donkey}'y);$ $\lambda s. s > \frac{1}{2} \text{all}'n)$

(The connective ‘;’ used at above examples constructs a pair $p ; c$ consisting of a nominal property p and a cardinality property c . Where the cardinality property is trivial, it is suppressed in the notation. These properties are separately interpreted in the model theory developed.)

Specification of an underspecified Skolem term of the form $\text{skolem}'p$ is defined as an ‘anytime’ operation that can occur at literally any point in a grammatical derivation, to yield a term which associates a standard Skolem term made up of a Skolem functor and an environment consisting of some bound variables with a nominal property. We shall call such terms ‘generalized Skolem terms’. They are obtained during a derivation as defined by the rule The specification of a skolem term, below.

Definition 9 (The specification of a skolem term) *Skolem specification of a term t of the form $skolem'_n p$ in an environment E yields a generalized Skolem term $sk_{n,p}^E$, which applies a generalized Skolem functor $Sk_{n,p}$ (where n is the number unique to the noun phrase that gave rise to t , and p is a nominal property corresponding to the restrictor of that noun phrase), to the tuple E (defined as the environment of t at the time of specification), which constitutes the arguments of the generalized Skolem term.*

The environment mentioned above is defined as follows:

Definition 10 (The environment of a skolem term) *The environment E of an unspecified skolem term T is a tuple comprising all variables bound by a universal quantifier or some other operator in whose structural scope T has been brought at the time of specification, by the derivation so far.*

We can usually suppress the Skolem number, identifying such a Skolem term as sk_p^E . Where there are distinct Skolem terms arising from identical noun phrases with the same restrictor p and environment E , they will be distinguished as (say) $sk_{53,p}^E$, $sk_{95,p}^E$, etc.

The skolem specification rule (defined above) implies that $skolem'$ is a function from properties like $\lambda y. donkey'y$ to functions from environments like $\{x\}$ to generalized Skolem terms like $sk_n^{\{x\}}$, $\lambda y. donkey'y$. Here are a few more examples:

$skolem'(\lambda y. donkey'y \wedge fat'y)$	$sk_{\lambda y. donkey'y \wedge fat'y}^E$
$skolem'(\lambda y. farmer'y \wedge own'(skolem'donkey'y))$	$sk_{\lambda y. farmer'y \wedge own'(skolem'donkey'y)}^E$
$skolem'(\lambda y. farmer'y \wedge own'(skolem'donkey'y);$ $\lambda s. s > 1)$	$sk_{\lambda y. farmer'y \wedge own'(skolem'donkey'y); \lambda s. s > 1}^E$
$skolem'(\lambda y. farmer'y \wedge own'(skolem'donkey'y);$ $\lambda s. s \leq 3)y)$	$sk_{\lambda y. farmer'y \wedge own'(skolem'donkey'y); \lambda s. s \leq 3)y}^E$
$skolem'(\lambda y. farmer'y \wedge own'(skolem'donkey'y);$ $\lambda s. s > \frac{1}{2} all'n)$	$sk_{\lambda y. farmer'y \wedge own'(skolem'donkey'y); \lambda s. s > \frac{1}{2} all'n }^E$

If there is more than one occurrence of an underspecified Skolem term t derived from the same noun phrase in the same environment E , as in the following interpretation for *Giles owns and operates some donkey*, the above definitions mean that they will necessarily be specified as the same generalized Skolem term sk_p^E .

$$(163) \text{ own}'sk'donkey'_{giles} \wedge \text{operate}'sk'_{donkey}giles'$$

In verifying interpretations involving generalized Skolem terms of the form sk_p^E against a model, it is necessary to unpack them, reinstating the nominal property p as a predication

over a traditional Skolem term, as in a traditional Skolemized formula like (161). However, as far as the grammatical semantics and the compositional derivation of logical form goes, expressions like $sk_{\lambda y. donkey' y}^{\{x\}}$ are unanalyzed identifiers, and this part of the responsibility for building logical forms is transferred to interpretation.

Environment features are deterministically passed down from the operator nodes in their c- or lf- commanding domain, and a specified generalized Skolem term is deterministically bound to all scoping universals in the relevant intensional scope at the point in the derivation at which it is specified. The available readings for a given sentence are thereby determined by the combinatorics of syntactic derivation and the logical forms that result.

It is also assumed that the pronouns like *it* translate as uninterpreted constants, which we might as well write *it'*, distinguishing by subscripts where necessary. Such pronoun interpretations, including those in donkey sentences, are assumed to be replaced via a DRT-like mechanism and since this mechanism is not part of the account such pronouns are just marked with ‘pro-terms’ of the form *pro'* x instead, where x is a discourse referent taking the form of a copy of the antecedent expression. Thus the donkey sentence (156) (*every farmer who owns a donkey feeds it*) yields the following interpretation:

$$(164) \quad \forall x [farmer'x \wedge own'sk_{donkey'x}^{\{x\}} \rightarrow feed'(pro'sk_{donkey'x}^{\{x\}})]$$

Steedman (2005) also assumes that a pronoun translation *it'* that has been brought by the derivation into an environment E is obligatorily bound to a variable in it via the following rule:

Definition 11 (Bound-variable pronoun specification) *Bound-variable pronoun specification of a term t of the form him' , her' or it' in an environment E yields a pro-term of the form $pro'x$, where $x \in E$.*

This rule is important for the correct analysis of the interaction of coordination and bound variable pronoun anaphora, thus we get the bound reading (165b) for (165a).

$$(165) \quad \begin{array}{l} \text{a. Every man}_i \text{ loves a woman who loves him}_i. \\ \text{b. } \forall x [man'x \rightarrow love'(skolem'\lambda y. woman'y \wedge love'(pro'x)y)x] \end{array}$$

Bound-variable pronoun specification, like Skolem specification, is an any-time operation. Logical forms are subject to the standard binding conditions in Steedman (2000a).

A number of predictions follow from definitions (9) (skolem specification) and (11) (bound variable pronouns specification). Since skolem term specification as defined is an ‘anytime’ operation – that is, since it can apply at any point in a derivation, wide and narrow scope readings of indefinites can be explained in the following way: If the skolem specification rule applies as soon as an indefinite NP is derived, before combination with any quantified matrix, then it necessarily yields a specific individual- denoting constant, behaving as if ‘has scope everywhere’, and giving *somebody* in the following sentence the appearance of an existential taking inverse scope, without the involvement of a true quantifier.

(166) Everybody loves somebody.

If on the other hand, skolem specification rule is applied after S is derived, the above sentence yields a reading where the quantifier has scope over the indefinite.

For examples regarding how this works, see the relevant CCG derivations in Chapter (3).

APPENDIX D. INFORMATION STRUCTURE IN COMBINATORY CATEGORIAL GRAMMAR

Compatible with the assumption of monostratality and the linguistic trend of minimalism, Steedman (2000a, 2000b) defines a single level of linguistic representation, which is called ‘Information Structure’, and it subsumes the predicate argument structure of CCG, augmented with the notions of theme/rheme and background/kontrast.

Theme / rheme distinction determines how parts of an utterance relate to the discourse structure: if it relates back, it is *thematic*; if it advances the discourse it is *rhematic*. In this theory intonational phrases can mark information units as theme or rheme, although not all boundaries are realized and a unit may contain more than one phrase. (Calhoun, Nissim, Steedman and Breiner, 2005, p.48).

The Information Structure can be considered as a representation of how parts of a single utterance is related to the discourse context. For example, the theme can be thought of as denoting what the speaker assumes to be *the question under discussion* and the rheme can be thought as what the speaker believes to be *the part of the utterance that advances the discussion*. A theme can be represented with a functional abstraction as exemplified by (167), and it is natural to think of it a set of possible alternative answers to a question.

(167) $\lambda x.marry'x anna'$ (Steedman, 2000a)

Therefore the notion of theme is associated with the set of propositions among all those supported by the conversational context that could possibly satisfy an existential proposition. In the context set by the question in (168), the mentioned set can be denoted by (169), in which $\diamond$ indicates possibility, and can be enumerated as in (170).

(168) Q: Well, what about ANNA? Who did SHE marry? (Steedman, 2000a)

A: (ANNA married) (MANNY).

L+H* LH% H* LL%

(169) $\exists x.\diamond marry'x anna'$ (Steedman, 2000a)

(170)

$$\left\{ \begin{array}{l} \diamond \text{marry}' \text{alan}' \text{anna}' \\ \diamond \text{marry}' \text{fred}' \text{anna}' \\ \diamond \text{marry}' \text{manny}' \text{anna}' \\ \dots \end{array} \right\}$$

Such alternative sets are called *rheme alternative sets* and are assumed to not to be known exhaustively by the hearers, and in practice one would want to compute with something more like the quantified expression in (169) or the λ -term itself.

In semantic terms the theme and rheme can therefore be characterized as follows:

- (171) a. The Theme *presupposes* the rheme alternative set. (Steedman, 2000a)
b. The Rheme *restricts* the rheme alternative set.

The sense in which a theme “presupposes” the rheme alternative set is much the same as that in which a definite expression presupposes the existence of its referent. That is to say, there is a pragmatic presupposition that the relevant alternative set is available in the contextual “mental model” or database. The presupposition may be “accommodated”, that is be added by the hearer after the fact of utterance to a contextual model that is consistent with it. Since a discourse context may contain more than one theme, the notion of the *theme alternative set* is also proposed, and the discourse structure is considered to be composed of a theme alternative set (TAS) and several rheme alternative sets, among other structures that a more elaborated theory of discourse structure may assume to exist. Such structures are updated when an utterance is processed, the details of which is given in section (D.1) of Chapter (4), because the details of this process is highly relevant to the subject matter of Chapter (4).

The other two concepts, *background* and *kontrast*, correspond to the distinction between given (old) information versus new information in a discourse, and are marked by intonational contours as in (168). Kontrast / Background distinction is a word level feature whereas Theme / Rheme distinction is a feature borne by a larger phrases, and such phrases can possibly be involved in coordinate structures, and can be intonationally marked or unmarked. The theme can also be continuous or discontinuous in an utterance.

Steedman and Korbayova (2003) proposes to situate this theory within Grosz and Sidner’s (1986) computational model of Discourse Structure, and Information Structure (IS) is considered to be a part of “Linguistic Structure” of Grosz and Sidner’s (1986) and its discourse

semantics are related to the Grosz and Sidner’s conception of “Attentional State”. In particular the notion of a center of attention (Grosz et al., 1995; Walker et al., 1998) is considered to be related to the IS notion of theme.

D.1 Update Semantics

Update semantics for theme and rheme is highly relevant to the subject matter of section (4.2), so we shall describe it here in some detail.

According to Steedman’s account of English Information Structure (including that of 1990, 2000a, 2000b and 2003), the functors θ' and ρ' in Figure E.2 are identity functions, but have a side effect of updating the discourse database. Since they are identity functions, the resulting formula does not contain them, but it is assumed that the discourse database is updated as follows.

It is assumed that the discourse database consists of a set of themes, which is called the *theme alternative set*, each member of which yields a *rheme alternative set*. For example, if the theme alternative set contains a theme such as (172), in which the λ -operator is interpreted as an existential quantifier as in (173), it yields a rheme alternative set as in (174), which are possible propositions which satisfy the existential quantificational formula. (Note that the $\diamond$ modality is interpreted as *it is possible that*.)

$$(172) \lambda x. \diamond \text{marry}'x \text{anna}'$$

$$(173) \exists x. \diamond \text{marry}'x \text{anna}'$$

$$(174)$$

$$\left\{ \begin{array}{l} \diamond \text{marry}'\text{alan}'\text{anna}' \\ \diamond \text{marry}'\text{fred}'\text{anna}' \\ \diamond \text{marry}'\text{manny}'\text{anna}' \end{array} \right\}$$

The discourse function of a rheme, then, is to restrict such rheme alternative sets, for example to a single one, as in (175).

$$(175) \diamond \text{marry}'\text{manny}'\text{anna}'$$

The idea can be characterized as causing one or more existing referents or “facts” such as $(\theta'(\lambda f.f[*m']))$ in Figure E.2), where θ marks the λ -term as a theme, to be *retracted* or removed from the theme alternative set, and causing a new theme to be *asserted* or added.

The rheme, ($\rho'(\lambda y. ate' * ap'y)$) in Figure E.2), is also thought of as updating the discourse database with a similar type of referent (more specifically, it modifies one of the members of the theme alternative set), and it may also become the theme of a subsequent utterance. However, the rheme does not require a preexisting referent or cause any existing thematic referents to be retracted (i.e. it can be asserted without needing to have a prior compatible thematic referent), although it may have other effects on the database, via the entailments and implicatures. Kontrast ('*') is also available for rhemes and similarly it is used to mark incompatible parts, which are felicitous only when contrasted.

Q: I know that Mary envies the man who wrote the musical. But who does she ADMIRE?

Figure D.1: A derivation for the sentence (87)

(176) $\exists x. \diamond *admires'x\ mary'$

(177) $\exists x. \diamond *like'x\ mary'$

The set of alternative themes in this case is the following:

(178)

$$\left\{ \begin{array}{l} \exists x. \diamond \textit{admires}'x \textit{ marry}' \\ \exists x. \diamond \textit{like}'x \textit{ marry}' \end{array} \right\}$$

The set of alternative themes can also be thought as a “question alternative set”, since a theme is informally considered to be “what is established by a wh-question”.

The rheme alternative set presupposed by the theme is therefore a set of propositions about Mary admiring various people. The rheme is *the woman who directed the musical*, where only the word *directed* is contrasted.

It is important to note that it is all and only the material marked by the pitch accent(s) that is contrasted. Anything not so marked, including the material between the pitch accent(s) and the boundary, is background.

APPENDIX E. INFORMATION STRUCTURE OF TURKISH

The traditional view in Turkish linguistics is that Turkish is a verb final language with free word order, and the variation of word order serves the purpose of assigning discourse related functions to certain parts of an utterance. For example, *focus* is considered to be associated with the immediately preverbal position, and the post-verbal positions are reserved for *background* material. Recently Özge and Bozşahin (2003, 2006) have proposed a tune based account of Turkish Information Structure in which word order variations are considered to be an epiphenomenal side effect caused by prosodic constraints, and the sole factor determining the information structure is considered to be intonation contours that the phrases of an utterance bears. Their paper argues that for an explanatory account of Turkish information structure, there is no need for any predications over sentence positions. A tune-based perspective where prosody is considered to be the sole structural determinant of information structure is instead proposed. Also it is noted that, Turkish prosody imposes precedence constraints on certain intonational contours that are responsible for the realization of information structural units. Word order variation therefore, is considered epiphenomenal, rather than being a determinant in attaining the right information structure required by discourse context.

Assuming an account similar to that of Steedman's account of information structure of English, Özge and Bozşahin (2006) determine intonation markers used by Turkish speakers as follows.

H*, which marks a maximum in pitch that is accompanied with some stress, is considered to be a prosodic functor that takes the boundary tone LL%, to yield a rheme marker H* LL%. This accent can be borne by phrases larger than a single element as in (180), or sometimes by just a single word as in (179). In (179) there is a slight pause between *maymun* and *yedi*, whereas in (180), there is not. Another difference of (180) is that the fall in pitch accent is not as abrupt as it is in (179).

(179) (mayMUN) (elma-yı ye-di.)
H* LL%

(180) (mayMUN ye-di.)
H* LL%

They also note that H* is phrase-initial and it always starts a new prosodic phrase, and is considered obligatory¹. Other identified accent types differ from H* in two respects: (i) they are not obligatory, and (ii) they are bound to phrase-final elements. These second group of accents include L*H-, L+H* L-, and H*+L H-. The following sentence exemplifies L* H-.

- (181) (maymun) (elma-YI ye-di.)
L*H- H* LL%

The prosodic boundary between *maymun* and *elmayı* is realized by a rise in pitch, which is notated as H-. Given the phrase finality of the accent, they state that it is not clear whether H- acts as an independent boundary event or is a part of the pitch accent. Furthermore, they distinguish this type of boundary from the LL% boundary – LL% is more abrupt than boundaries designated by ‘-’. They also identify phrase boundaries associated by a fall in pitch as L-, using the same notation. The other identified pitch accents are exemplified by examples (182), (183).

- (182) (mayMUN) (elma-YI ye-di.)
L+H*L- H* LL%

- (183) (MAYmun) (elma-YI ye-di.)
H*+L H- H* LL%

Regarding the traditional claims about inability of post verbal arguments to bear stress, an alternative solution is proposed by claiming that LL% causes pitch flooring to its right. For example, prosodic marking is not possible after the matrix predicate in the sentence (184). Pitch flooring is marked with $\langle -F - \rangle$.

- (184) (DÜN gece) (Ali Aynur-u yemeğ-e götür-dü).
H* LL% $\langle -F - \rangle$

(185) To summarize their findings:

- (i) H* LL% causes flooring to its right.
- (ii) H* LL% is a pitch contour associated with Steedman’s rheme.

¹ In section (E.1), we shall propose a slightly different possibility: that LL% is obligatory due to a communicative necessity (that every sentence should eventually end and other boundaries seem to create a sense of sentence continuation) and that LL% is strongly associated with H* due to physical reasons involved in the task of detection of an intonation contour from the lack of one.

- (iii) H* LL% is required in complete declarative utterances (because all such utterances have a rheme).
- (iv) H* is prosodic phrase-initial.
- (v) In a complete declarative utterance, the matrix predicate cannot appear in a prosodic phrase marked an intonation contour other than H* LL%, or alternatively it is floored.
- (vi) H* LL% is information structurally ambiguous between being a complete rheme marker, and a theme-rheme marker which can be disambiguated by a slight pause after the theme.

Similar to Steedman's theory of English information structure, intonation markers are associated with *kontrast*. They also mark a prosodic phrase as *theme/rheme*. The intonation contours associated with theme are L* H-, H*+L H- and L+H* L-.

Regarding prosodic constraints on Turkish information structure, the following are mentioned in addition to (185).

(186) Turkish Information Structural Organization:

Theme-kontrast precedes the rheme: The rheme contour causes flooring, which renders the placement of thematic-kontrast to its right impossible.

Some consequences of the prosodic constraints in (185) and the information structural constraint in (186) render the matrix predicate position and its neighboring syntactic positions, and therefore the surface word order, epiphenomenal, as follows:

(187) Turkish Epiphenomenal Word Order:

- a. An utterance which consists of a rheme only is possible if there is no contrast.
- b. A matrix predicate is either part of the rheme or it must follow the rheme.
- c. Any verb-initial order has the matrix predicate as the rheme-kontrast.
- d. Rheme-kontrast precedes the rheme-background. This follows from a prosodic constraint, rather than an information structural constraint. Since H* pitch accent is phrase-initial, it follows that rheme-kontrast must precede the rheme-background.

- e. Discontinuous themes are possible, because theme-background and rheme background are not related by precedence.
- f. By (d), discontinuous rhemes are not possible. This also seems to be a prosodic constraint, and needs further inquiry.

They also note that none of the theme marking accents in Turkish are available to the matrix predicate, therefore it is impossible to restrict a Theme Alternative Set (TAS) consisting of lambda expressions that differ only in their matrix predicates.

E.1 Some Remarks on Özge and Bozşahin (2006)

As stated above, Özge and Bozşahin consider the pitch marker H^* to be a prosodic functor. While the usage of the term ‘functor’ is just a descriptive label, this is nevertheless reminiscent of Steedman’s 1990 account of English information structure (Steedman, 1990), in which, pitch accents, not boundaries, were considered to be functors. Later, Steedman (2000a, 2000b) presented a modified account, in which pitch accents are considered to be feature markers that operate by unification, and boundary tones are considered to be part of the string which are functors over such marked (or sometimes unmarked) utterance parts that yield prosodic phrases.

Özge and Bozşahin (2006) also state that the H^- marker in (181) might be a boundary tone or it might be a part of the pitch accent, and further phonological research is required to decide between these alternatives.

Given the statement about H^* being a prosodic functor and uncertainty concerning some prosodic markers such as H^- , it appeared to me that Steedman’s earlier proposal might be more appropriate for Turkish. While personal communication with Dr. Cem Bozşahin has revealed that this wasn’t their claim and that the term ‘functor’ was simply intended as a descriptive label, nevertheless we shall briefly examine this possibility.

Steedman’s 1990 paper with the title ‘Structure and Intonation in Spoken Language Understanding’ defines the following pitch accents to be functors over boundary tones, as follows:

- (188) $L+H^*$ $:=$ Theme /Bh
 H^* $:=$ Rheme /Bl

Bh and Bl represent the category of the boundary tones. $L+H^*$ and H^* are defined to be functors over such tones, to yield categories Theme or Rheme, respectively.

The following categories are assigned to the boundary tones:

- (189) LH% := Bh
 LL% := Bl
 LL := Bl

Interpolation of tunes over strings are modelled using a polymorphic category (instead of unification as in Steedman (2000a, 2000b)), and represented as the tone ‘0’.

- (190) 0 := X / X

Syntactic combination then is made subject to the following restriction:

- (191) *The Prosodic Constituent Condition*: Combination of two syntactic categories via a syntactic combinatory rule is only allowed if their prosodic categories *also* combine. (Steedman, 1990)

This principle has the effect of excluding certain derivations for spoken utterances that would be allowed for the equivalent written sentences². For example, in Figure E.1, which is adopted from Steedman (1990), the rule of forward composition is allowed to apply to the words *Fred* and *ate*, because the prosodic categories can combine (by functional application)³.

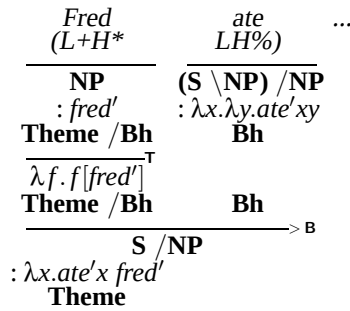


Figure E.1: Parallel levels of surface and phonetic syntax that were assumed in Steedman (1990)

The stipulation in (191) seems to have lead Steedman to an alternative formulation, and finally abandon the idea of pitch accents as prosodic functors, resulting in the later account in Steedman (2000a, 2000b). It is also possible to criticize the 1990 account by stating that the assumption of monostratality requires a single level of representation, but in Steedman

² Steedman (2000a, 2000b) implements this restriction via unification without having an additional syntactic level of prosody.

³ The functional application of phonetic categories is not shown in the figure but is assumed.

(1990), there is a violation of the principle of monostratality stronger than an extra level of representation. It has an extra grammar, which we may call the *prosodic grammar*⁴, which functions as a constraint on the possible derivations licenced by the main grammar. On the other hand, not everyone working on CCG is comfortable with Steedman's 2000 account either. For example, regarding this later account, Kruijff G.J. and Vasishth S. (1997) state the following:

Steedman subsequently makes a rather inelegant move, and models boundary tones as empty strings, reminiscent of transformational grammar's empty categories. A boundary tone has a functional category, with no realization, that combines with pitch accents. Important here is that the composition of a boundary tone category with a pitch accent category allows for the Theme / Rheme distinction to be projected from pitch accents onto prosodic phrases.

On the other hand, it may be possible to argue that there is empirical evidence regarding these boundary tones to be empty strings. For example, Özge and Bozşahin (2006) state that in Turkish H* LL% contour is prosodically ambiguous leading to an information structural ambiguity. They also state that a slight pause after the theme resolves the ambiguity. This evidence can be used as an example that shows that boundary tones sometimes actually *are* empty strings, and sometimes they are not (if we don't consider an audible pause to be an empty string). Therefore, it can be said that Turkish provides support for Steedman's 2000 account by having an intonation contour that has an appropriate *physical* property, not some theory-specific construct.

Whatever the truth may be, the difference between these two alternative formalizations of information structure have no bearing to the initially intended subject matter of this thesis. The relevance of Özge and Bozşahin (2006) to this thesis is that there are prosodic restrictions on the information structure of Turkish. The actual realization or formalization details are immaterial. Also, the ideas presented in Özge and Bozşahin (2006), whose CCG formalization wasn't given can be formulated in a way compatible with Steedman (2000a, 2000b): Boundary tones like LL% instead of pitch accents such as H* can be considered to be functors and boundary tones such as LL% (and possibly others such as H- in (181), which may actually be a boundary tone) can be considered as a part of the surface string, and this would not prevent us modelling prosodic constraints such as the idea of the main functor causing flooring to its

⁴ Every published version of Steedman's theory of Information Structure assume that the syntax and intonation structures of English are identical, and have the same grammar; but nevertheless, the way it is formulated in Steedman (1990) makes use of a separate lexicon of intonation contours. The principle of monostratality is therefore nevertheless violated by having an extra grammar which has its own representational level.

right. For example, following Steedman (2000a, 2000b), let's consider H^* to be a pitch accent marking the word it falls on to be a rheme kontrast by making use of a minor syntactic feature called INFORMATION. The subscript p is meant to supply a value of *rheme* to this minor syntactic feature.

$$(192) \quad \text{mayMUN} \quad := \overset{b}{\triangleleft} N_p : *maymun' \\ H^*$$

In (192), the diacritic ‘*’ marks Steedman’s *kontrast*, which distinguishes the parts of an utterance that are contrasted with respect to the given (background) information in the discourse. Unlike theme and rheme, kontrast is a word-level⁵ feature.

Again following Steedman (2000a, 2000b), unmarked words are considered to be under-specified for INFORMATION, as in (193).

$$(193) \quad \text{maymun} \quad := \overset{b}{\triangleleft} N : maymun'$$

Again, following Steedman (2000a, 2000b), these are not actually meant to be separate lexical assignments for the word *maymun*, but rather considered to be the results of a speech processing algorithm which takes words and morphemes from the lexicon and modifies their categories such that formulas as in (192) results, as appropriate intonation markers are detected together with the detection of words or morphemes, which can be considered to be interdependent tasks.

The boundary tone LL% can be defined as in (194), or we could possibly restrict this boundary tone be a functor that applies to phrases that are marked as a rheme, as in (195).

$$(194) \quad \text{LL\%} \quad := S\$_{\phi} \backslash S\$_{\eta} : \lambda f. \eta' f$$

$$(195) \quad \text{LL\%} \quad := S\$_{\phi} \backslash S\$_p : \lambda f. \rho' f$$

The need to associate the boundary tones with syntactic categories of the form $S\$$ emerges from the fact that INFORMATION features that are supplied by pitch accents becomes an obstacle to unification when the final syntactic category of S will be derived. The replacement value, ϕ , remedies this problem, because it is the same for both theme marked and rheme marked categories. These lexical assignments, which associate prosodic boundaries with syntactic constituents that can possibly coordinate are also safe from a theoretical point of view,

⁵ In a morphemic lexicon as in Bozşahin (2002) it would be applicable for morphemes, too.

because such a strong parallelism between syntax and prosodic structure is exactly what is assumed to begin with.

Neither we nor Steedman himself is claiming that these lexical assignments are actually specified in the lexicon. The lexical assignment symbol ($:=$) and the power of CCG is being used to model the behavior of the speech processor. Whether the speech processor might have its own lexicon of intonation markers (and therefore existence of a *grammar of prosody*⁶) is an interesting question, though. We do not give a yes or no answer to this question, other than just noting that it is within the generative power of CCG and that it is a reasonable way to integrate the problem of speech recognition and that of parsing, into a single spoken language understanding problem, which is one of the main goals of Steedman (2000a, 2000b)⁷.

The identified theme markers L^* H-, H^*+L H- and $L+H^*$ L- can be divided into two parts, one being the pitch accent related to theme (L^* , H^*+L and $L+H^*$), and the other being a compatible boundary event, such as H- or L-.

For example, the word *maymun* bearing a theme marker can be defined as in (196), and the boundaries H- and L- can be defined as in (197), or we can explicitly restrict these boundaries to phrases marked as theme, as in (198):

$$\begin{array}{ll}
 (196) \quad \text{mayMUN} & := \overset{b}{\triangleleft} N_\theta : *maymun' \\
 & L^* \\
 & \text{mayMUN} \quad := \overset{b}{\triangleleft} N_\theta : *maymun' \\
 & H^*+L \\
 & \text{mayMUN} \quad := \overset{b}{\triangleleft} N_\theta : *maymun' \\
 & L+H^*
 \end{array}$$

$$(197) \quad H-, L- \quad := S\$_\phi \backslash S\$_\eta : \lambda f. \eta' f$$

⁶ Steedman (1990) actually *did* have such a separate grammar of prosody, which was abandoned in Steedman (2000a, 2000b), and a unification-based mechanism is employed instead. In Steedman (2000a, 2000b) the pitch accents such as H^* are feature markers that contribute features to the unification mechanism, and the boundary tones are part of the input string and *are* lexically specified functors which take an incomplete intonationally marked (or unmarked) phrase to yield a complete prosodic phrase (which is either a rheme or a theme).

⁷ In Steedman (2000a), it is explicitly stated that the problem of spoken language understanding would be an easier problem if speech recognition and parsing is considered to be a single integrated problem, which explains why we have more than one lexical assignment for a single word like *maymun* differing only in intonation markers. The reason, we think, is that detecting whether a word bears stress (i.e. the $*$ marker in H^*) or not would involve considering harmonics above the fundamental frequency (F0), which also incorporates information about vowels. The implication is that detecting stress on a word and detecting the word itself is an interdependent problem. H or L marker on the other hand, only refers to the F0 level. Consonants can be considered as ‘noise’ in the signal, unrelated to F0, which may also have different physical realizations depending on whether the word or morpheme bears stress or not.

$$(198) \quad H-, L- \quad := S\$_{\phi} \backslash S\$_{\theta} : \lambda f. \theta' f$$

Predictions of these categorial assignments would be that these phonetic markers can be borne by and be interpolated by strings larger than a single word. While Özge and Bozşahin (2006) states that these theme markers are prosodic phrase final, it is not explicitly stated that they cannot be borne by surface constituents larger than a single word. It should also be noted that the assumed lexicon is morphemic (as in Bozşahin (2002)), and thus the concept of *word* has no particularly distinguished status from say, a case marker, such as -ACC or some other inflectional suffix, and in Özge and Bozşahin (2006) there are example sentences in which such a thematic marker is interpolated over a string bearing more than one lexical category. One such sentence is (199), in which L* H- is interpolated over *gel-di*.

$$(199) \quad \begin{array}{ccccccc} \text{(ev-e} & \text{gel-di),} & \text{(üst-ü-nü} & \text{çıkard-ı),} & \text{(DUŞ-A} & \text{gir-di).} \\ & L^* H- & & L^* H- & H^* & LL\% \end{array}$$

It is possible to restrict the application of these boundary tones to words by requiring them to apply to surface constituents bearing a lattice diacritic marker equal to or less than n-case (c) or s-person (s), as in (200), both of which are lower than the diacritic borne by type raised nouns (*f*) in the lattice structure; and type raised nouns are either complete words, or more complex surface constituents. The idea that the constituents marked by H- and L- cannot have the lattice diacritic *f* restricts their application to words.

$$(200) \quad \begin{array}{ll} H-, L- & := S\$_{\phi} \backslash S\$_{\eta}^{\leq c} : \lambda f. \eta' f \\ H-, L- & := S\$_{\phi} \backslash S\$_{\eta}^{\leq s} : \lambda f. \eta' f \end{array}$$

In (200) the variable *S\$* ranges over the set of categories including *S* and all functions into members of *S\$* – that is, it includes *S*, *S/NP*, and all verbs and type-raised arguments of verbs⁸, but not nouns and the like. Of course, there is no n-case diacritic that applies to a category of type *S*, so the $S\$_{\eta}^{\leq c}$ in (200) is meant to be a feature of the corresponding NP when the macro *S\$* is expanded.

On the other hand, a *theme* can be a surface constituent more complex than a word, and therefore this approach is not sustainable. For that reason, we stick to the category assignment in (197), or (198).

⁸ For example, type raised nouns are included.

The idea that the boundary tone LL% causes flooring to its right would be easier to implement if we had a separate phonetic grammar as in Steedman (1990), but in the present terms, we might assume a feature called ENDOFPLUT, with the meaning ‘End of Planned Utterance’⁹ in addition to the INFORMATION feature, and associate this value with the LL% boundary. In (201), the symbol $\flat$ is used to indicate an assignment of *true* to the feature ENDOFPLUT. We also assume that every pitch accent including H* (that is, H*, L*, H+L* and H*+L) and the boundaries except LL% implicitly have the value *false* assigned to ENDOFPLUT, albeit this shall be suppressed by the notation that we shall use. As an example we may define H* as in (202), $\flat$ meaning an assignment of *false* to ENDOFPLUT, but this is a detail that can safely be suppressed in the notation used without creating a risk of misunderstanding. When the category in (201) is combined with the one in (202), involved feature assignments can unify because $\flat$ is the value of the result category in (201), but H* is an argument to LL%. If there are words after LL% in the utterance, they cannot unify with $\flat$ if they bear a phonetic marker, any of which, including the boundaries H- and L- (but not LL%) have the assignment $\flat$. It follows that in this scheme the words after LL% can still have the LL% marker at the end of the utterance, but this does not seem to be a problem that we need to avoid. If, nevertheless, we want to block such a derivation, then (203) can be used instead of (201), which has an additional checking due to $\flat$ on the argument category.

$$(201) \quad \text{LL\%} \quad \quad \quad := S\$_{\phi, \flat} \backslash S\$_p : \lambda f. \rho' f$$

$$(202) \quad \text{mayMUN} \quad \quad \quad := \overset{b}{\triangleleft} N_{p, \flat} : *maymun'$$

H*

$$(203) \quad \text{LL\%} \quad \quad \quad := S\$_{\phi, \flat} \backslash S\$_{p, \flat} : \lambda f. \rho' f$$

As an example of how these category assignments work, consider Figure E.2¹⁰, which corresponds to the sentence (181).

⁹ The feature ENDOFPLUT, which we propose specifically for Turkish, is not present in any of Steedman’s accounts of English information structure. We associate this feature with the end of the intended utterance, while considering the rest to be unintended but later discovered additional material. Since Turkish is a verb-final language, we think that it is safe to consider post-verbal arguments to be a result of poorly planned utterance generation. While we have not studied the question of what kind of psycholinguistic process might explain the proposal that the end of a planned utterance (which is marked by LL%) causes pitch flooring. Perhaps it’s just that prosodic markers are only available for well-planned parts of an utterance. Since prosodic markers carry functions relevant to the discourse structure (such as theme, rheme and contrast), it might be plausible to claim that any part of an utterance that has a related part in the discourse would be in the well-planned parts of an utterance, because there are no lack of information about such utterance parts. If there is no lack of information, probably there is no reason as to way such parts should be poorly planned or remembered later after the LL% boundary.

¹⁰ In Figure E.2, the meaning of the abbreviations are as follows: *m'* stands for *monkey*, *ap'* stands for *apple*, *no'* stands for *nominative* case and *ac'* stands for *accusative* case.

With the category assignments mentioned in this section, the phonetic restriction regarding the fact that H* LL% causes flooring to its right (i.e. (185i)) is a stipulation that emerges from the category assignment of pitch accents and boundary tones, that is, this restriction is specified in the lexicon. Similarly (185ii) is also lexically specified. Regarding (185iv), there is nothing in the present lexical assignments that forces H* to be prosodic phrase-initial. (185v) could be lexically specified, but in the present terms it is not. The ambiguity regarding H* LL% can be modelled by assuming that some boundaries are indistinguishable from empty strings.

The claim in (186) (that in Turkish theme-contrast precedes the rheme) is obtained without further stipulation, and is already implemented with the present lexical assignments, because it directly follows from the fact that the rheme contour (H* LL%) causes flooring, and the behavior of H* LL% with respect to flooring is implemented by the present lexical assignments.

Regarding (187) (Turkish Epiphenomenal Word Order), the following facts can be noted: (187a) (that an utterance which consists of rheme only is possible if there is no kontrast) follows from the ambiguity of H* LL%, (187b) (that a matrix predicate is either part of the rheme or it must follow the rheme) is implemented by the present lexical assignments without further stipulation, (187c) (that any verb initial order has the matrix predicate as the rheme-kontrast) is not implemented, because the idea that H* is obligatory is not implemented, but it might be strongly associated with LL% for physical reasons (see below). An alternative

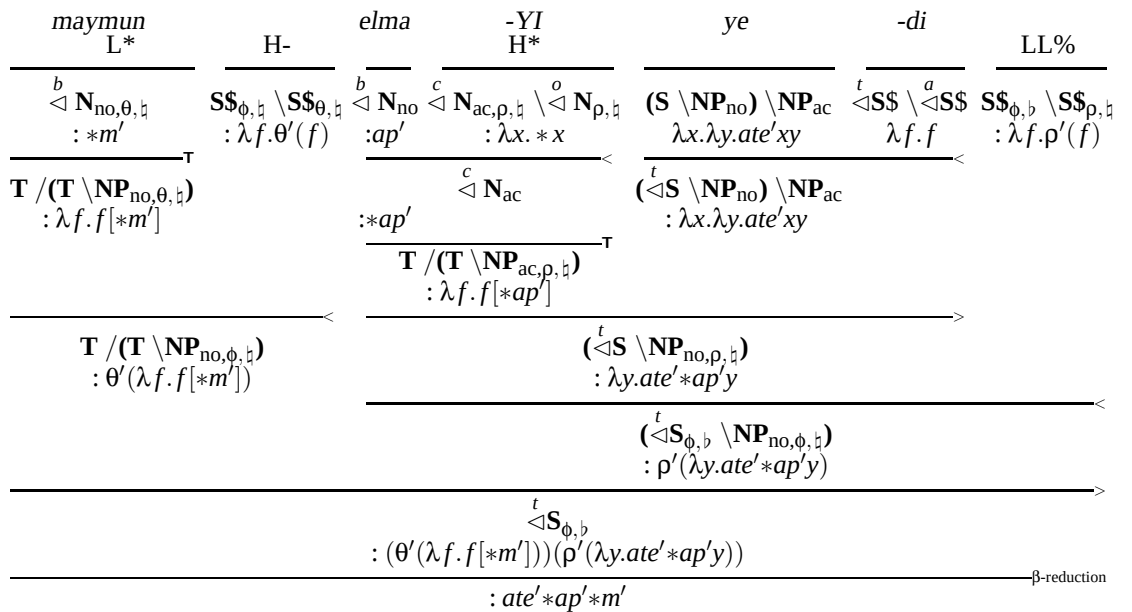


Figure E.2: The derivation for the sentence (181)

formulation might state that instead of H*, LL% is obligatory and that LL% is a functor looking for H* (ρ-marker) in its left. The idea that LL% is obligatory might be associated with the meaning of other boundaries. For example H- creates a sense of continuation (as in (199)) even when it is borne by a matrix predicate. If a sentence is supposed to end at all, it must have a LL% somewhere. (187d) (that rheme contrast precedes rheme background) follows from the proposal that LL% is a functor that is looking for a ρ-marked constituent in its left. (187e) (that discontinuous themes are possible) needs no further stipulation because it wasn't a stipulation in Özge and Bozşahin (2006) either. (187f) (that discontinuous rhemes are not possible) can be obtained if LL% is associated with ρ instead of η, and H- and L- is associated with θ instead of η, in (195), instead of (194), and as in (198) instead of (197) because in that case the fact that discontinuous rhemes are not possible would follow from the fact that no prosodic marking after LL% is possible, and no rheme before H* is possible either, because the boundaries H- and L- would not be able to combine with the rheme marker H* to yield a prosodic phrase.

What remains to be explained is why LL% needs to have a strong association with H*: that might follow from the fact that contours such as L* LL% and H* H- are indistinguishable (or harder to distinguish) from the lack of a contour. All audible contours involve a change in pitch, as in L* H-, H*+L H-, L+H* L- and H* LL%. If the audible difference between H* and L* is involved in distinguishing between a rheme and a theme, it might be natural to have a boundary tone that has a different fundamental frequency following such pitch accents. Steedman (2003) associates the boundaries L, LL%, HL% with speaker responsibility (modality [S]) and the boundaries H, HH%, and LH% with hearer responsibility (modality [H]). Such lexical assignments might create the impression that all pitch accent / boundary combinations are possible (because each has a different role in determining information structure) – but, if that were the case, there wouldn't be three different boundaries associated with modality [S], and three different boundaries associated with [H]; one would suffice.

The argument in the last paragraph creates the question as to why the pitch marker L+H* isn't also associated with LL%, but with L-, (which in turn yields a theme marker (L+H* L-) identified in Özge and Bozşahin (2006)) as the fundamental frequency of H* in L+H* is audibly different enough from that of LL%. Özge and Bozşahin (2006) states that LL% is more abrupt than boundaries marked with '-'. That is, the difference between LL% and L- is the speed of the change of F0: In LL%, F0 drops more rapidly than in L-. In Figure 1 of Özge and Bozşahin (2006) (which is reprinted here as Figure E.3) it is also seen that the rise in

L+H* before H* is also a slow one, but nevertheless much faster than some of the curves that distinguish H* from its environment (see Figure 1.c and 1.d of their paper). It follows that if H* is marked by such a slow rise in F0, only a rapid F0 fall as in LL% could make it audible. That associates H* with LL% by physical necessity, and existence of LL% in every sentence might be a communicative necessity, as argued above, because other possible boundaries such as H- create a sense of utterance continuation.

While the account described here has the spirit of a typical solution proposal to an engineering problem, namely that of integration of speech processing with parsing, rather than a purely linguistic account, it should be remembered that human brain is be able to solve the same engineering problem.

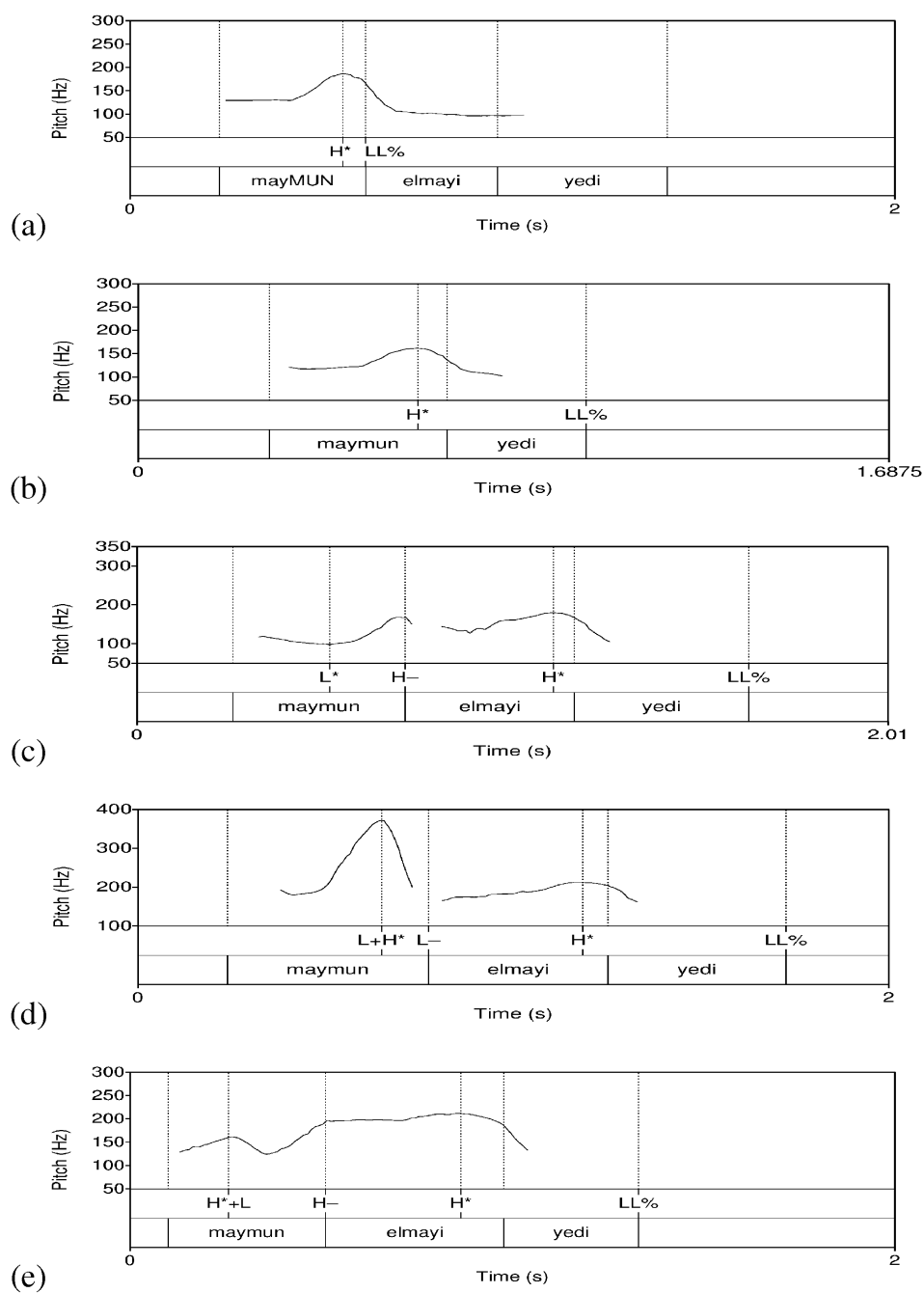


Figure E.3: F0 curves