

The Parallel Grammar Project

Miriam Butt

Cent. for Computational Linguistics
UMIST
Manchester M60 1QD GB
mutt@csl.stanford.edu

Helge Dyvik

Dept. of Linguistics
University of Bergen
N5007 Bergen NORWAY
helge.dyvik@lil.uib.no

Tracy Holloway King

Palo Alto Research Center
Palo Alto, CA 94304 USA
thking@parc.com

Hiroshi Masuichi

Corporate Research Center
Fuji Xerox Co., Ltd.
Kanagawa 259-0157, JAPAN
hiroshi.masuichi@fujixerox.co.jp

Christian Rohrer

IMS Universität Stuttgart
D-70174 Stuttgart GERMANY
rohrer@ims.uni-stuttgart.de

Abstract

We report on the Parallel Grammar (ParGram) project which uses the XLE parser and grammar development platform for six languages: English, French, German, Japanese, Norwegian, and Urdu.¹

1 Introduction

Large-scale grammar development platforms are expensive and time consuming to produce. As such, a desideratum for the platforms is a broad utilization scope. A grammar development platform should be able to be used to write grammars for a wide variety of languages and a broad range of purposes. In this paper, we report on the Parallel Grammar (ParGram) project (Butt et al., 1999) which uses the XLE parser and grammar development platform (Maxwell and Kaplan, 1993) for six languages: English, French, German, Japanese, Norwegian, and Urdu. All of the grammars use the Lexical-Functional Grammar (LFG) formalism which produces c(onstituent)-structures (trees) and f(unctional)-structures (AVMs) as the syntactic analysis.

LFG assumes a version of Chomsky's Universal Grammar hypothesis, namely that all languages are structured by similar underlying principles. Within LFG, f-structures are meant to encode a language universal level of analysis, allowing for cross-linguistic parallelism at this level of abstraction. Although the construction of c-structures is governed

by general wellformedness principles, this level of analysis encodes language particular differences in linear word order, surface morphological vs. syntactic structures, and constituency.

The ParGram project aims to test the LFG formalism for its universality and coverage limitations and to see how far parallelism can be maintained across languages. Where possible, the analyses produced by the grammars for similar constructions in each language are parallel. This has the computational advantage that the grammars can be used in similar applications and that machine translation (Frank, 1999) can be simplified.

The results of the project to date are encouraging. Despite differences between the languages involved and the aims and backgrounds of the project groups, the ParGram grammars achieve a high level of parallelism. This parallelism applies to the syntactic analyses produced, as well as to grammar development itself: the sharing of templates and feature declarations, the utilization of common techniques, and the transfer of knowledge and technology from one grammar to another. The ability to bundle grammar writing techniques, such as templates, into transferable technology means that new grammars can be bootstrapped in a relatively short amount of time.

There are a number of other large-scale grammar projects in existence which we mention briefly here. The LS-GRAM project (Schmidt et al., 1996), funded by the EU-Commission under LRE (Linguistic Research and Engineering), was concerned with the development of grammatical resources for nine European languages: Danish, Dutch, English, French, German, Greek, Italian, Portuguese, and Spanish. The project started in January 1994 and ended in July 1996. Development of grammatical resources was carried out in the framework of the

¹We would like to thank Emily Bender, Mary Dalrymple, and Ron Kaplan for help with this paper. In addition, we would like to acknowledge the other grammar writers in the ParGram project, both current: Stefanie Dipper, Jean-Philippe Marcotte, Tomoko Ohkuma, and Victoria Rosén; and past: Caroline Brun, Christian Fortmann, Anette Frank, Jonas Kuhn, Veronica Lux, Yukiko Morimoto, María-Eugenia Niño, and Frédérique Segond.

Advanced Language Engineering Platform (ALEP). The coverage of the grammars implemented in LS-GRAM was, however, much smaller than the coverage of the English (Riezler et al., 2002) or German grammar in ParGram. An effort which is closer in spirit to ParGram is the implementation of grammar development platforms for HPSG. In the Verbmobil project (Wahlster, 2000), HPSG grammars for English, German, and Japanese were developed on two platforms: LKB (Copestake, 2002) and PAGE. The PAGE system, developed and maintained in the Language Technology Lab of the German National Research Center on Artificial Intelligence DFKI GmbH, is an advanced NLP core engine that facilitates the development of grammatical and lexical resources, building on typed feature logics. To evaluate the HPSG platforms and to compare their merits with those of XLE and the ParGram projects, one would have to organize a special workshop, particularly as the HPSG grammars in Verbmobil were written for spoken language, characterized by short utterances, whereas the LFG grammars were developed for parsing technical manuals and/or newspaper texts. There are some indications that the German and English grammars in ParGram exceed the HPSG grammars in coverage (see (Crysmann et al., 2002) on the German HPSG grammar).

This paper is organized as follows. We first provide a history of the project. Then, we discuss how parallelism is maintained in the project. Finally, we provide a summary and discussion.

2 Project History

The ParGram project began in 1994 with three languages: English, French, and German. The grammar writers worked closely together to solidify the grammatical analyses and conventions. In addition, as XLE was still in development, its abilities grew as the size of the grammars and their needs grew.

After the initial stage of the project, more languages were added. Because Japanese is typologically very different from the initial three European languages of the project, it represented a challenging case. Despite this typological challenge, the Japanese grammar has achieved broad coverage and high performance within a year and a half. The South Asian language Urdu also provides a widely spoken, typologically distinct language. Although it is of Indo-European origin, it shares many characteristics with Japanese such as verb-finality, relatively free word order, complex predicates, and the abil-

ity to drop any argument (rampant pro-drop). Norwegian assumes a typological middle position between German and English, sharing different properties with each of them. Both the Urdu and the Norwegian grammars are still relatively small.

Each grammar project has different goals, and each site employs grammar writers with different backgrounds and skills. The English, German, and Japanese projects have pursued the goal of having broad coverage, industrial grammars. The Norwegian and Urdu grammars are smaller scale but are experimenting with incorporating different kinds of information into the grammar. The Norwegian grammar includes a semantic projection; their analyses produce not only c- and f-structures, but also semantic structures. The Urdu grammar has implemented a level of argument structure and is testing various theoretical linguistic ideas. However, even when the grammars are used for different purposes and have different additional features, they have maintained their basic parallelism in analysis and have profited from the shared grammar writing techniques and technology.

Table (1) shows the size of the grammars. The first figure is the number of left-hand side categories in phrase-structure rules which compile into a collection of finite-state machines with the listed number of states and arcs.

(1)

Language	Rules	States	Arcs
German	444	4883	15870
English	310	4935	13268
French	132	1116	2674
Japanese	50	333	1193
Norwegian	46	255	798
Urdu	25	106	169

3 Parallelism

Maintaining parallelism in grammars being developed at different sites on typologically distinct languages by grammar writers from different linguistic traditions has proven successful. At project meetings held twice a year, analyses of sample sentences are compared and any differences are discussed; the goal is to determine whether the differences are justified or whether the analyses should be changed to maintain parallelism. In addition, all of the f-structure features and their values are compared; this not only ensures that trivial differences in naming conventions do not arise, but also gives an overview of the constructions each language covers and how

they are analyzed. All changes are implemented before the next project meeting. Each meeting also involves discussion of constructions whose analysis has not yet been settled on, e.g., the analysis of participles or proper names. If an analysis is agreed upon, all the grammars implement it; if only a tentative analysis is found, one grammar implements it and reports on its success. For extremely complicated or fundamental issues, e.g., how to represent predicate alternations, subcommittees examine the issue and report on it at the next meeting. The discussion of such issues may be reopened at successive meetings until a consensus is reached.

Even within a given linguistic formalism, LFG for ParGram, there is usually more than one way to analyze a construction. Moreover, the same theoretical analysis may have different possible implementations in XLE. These solutions often differ in efficiency or conceptual simplicity and one of the tasks within the ParGram project is to make design decisions which favor one theoretical analysis and concomitant implementation over another.

3.1 Parallel Analyses

Whenever possible, the ParGram grammars choose the same analysis and the same technical solution for equivalent constructions. This was done, for example, with imperatives. Imperatives are always assigned a null pronominal subject within the f-structure and a feature indicating that they are imperatives, as in (2).

- (2) a. Jump! Saute! (French)
 Spring! (German) Tobe! (Japanese)
 Hopp! (Norwegian) kuudoo! (Urdu)

- b.
$$\left[\begin{array}{ll} \text{PRED} & \text{'jump<SUBJ>} \\ \text{SUBJ} & \left[\begin{array}{ll} \text{PRED} & \text{'pro'} \end{array} \right] \\ \text{STMT-TYPE} & \text{imp} \end{array} \right]$$

Another example of this type comes from the analysis of specifiers. Specifiers include many different types of information and hence can be analyzed in a number of ways. In the ParGram analysis, the c-structure analysis is left relatively free according to language particular needs and slightly varying theoretical assumptions. For instance, the Norwegian grammar, unlike the other grammars, implements the principles in (Bresnan, 2001) concerning the relationship between an X'-based c-structure and the f-structure. This allows Norwegian specifiers to be analyzed as functional heads of DPs etc.,

whereas they are constituents of NPs in the other grammars. However, at the level of f-structure, this information is part of a complex SPEC feature in all the grammars. Thus parallelism is maintained at the level of f-structure even across different theoretical preferences. An example is shown in (3) for Norwegian and English in which the SPEC consists of a QUANT(ifier) and a POSS(essive) (SPEC can also contain information about DETERminers and DEMONstratives).

- (3) a. alle mine hester (Norwegian)
 all my horses
 'all my horses'

- b.
$$\left[\begin{array}{ll} \text{PRED} & \text{'horse'} \\ \text{SPEC} & \left[\begin{array}{ll} \text{QUANT} & \left[\begin{array}{ll} \text{PRED} & \text{'all'} \end{array} \right] \\ \text{POSS} & \left[\begin{array}{ll} \text{PRED} & \text{'pro'} \\ \text{PERS} & 1 \\ \text{NUM} & \text{sg} \end{array} \right] \end{array} \right] \end{array} \right]$$

Interrogatives provide an interesting example because they differ significantly in the c-structures of the languages, but have the same basic f-structure. This contrast can be seen between the German example in (4) and the Urdu one in (5). In German, the interrogative word is in first position with the finite verb second; English and Norwegian pattern like German. In Urdu the verb is usually in final position, but the interrogative can appear in a number of positions, including following the verb (5c).

- (4) Was hat John Maria gegeben? (German)
 what has John Maria give.PerfP
 'What did John give to Mary?'

- (5) a. jon=nee marii=koo kyaa diiyaa? (Urdu)
 John=Erg Mary=Dat what gave
 'What did John give to Mary?'

- b. jon=nee kyaa marii=koo diiyaa?

- c. jon=nee marii=ko diiyaa kyaa?

Despite these differences in word order and hence in c-structure, the f-structures are parallel, with the interrogative being in a FOCUS-INT and the sentence having an interrogative STMT-TYPE, as in (6).

$$(6) \left[\begin{array}{ll} \text{PRED} & \text{'give<SUBJ,OBJ,OBL>'} \\ \text{FOCUS-INT} & \left[\begin{array}{ll} \text{PRED} & \text{'pro'} \\ \text{PRON-TYPE} & \text{int} \end{array} \right]_1 \\ \text{SUBJ} & \left[\text{PRED} \text{ 'John'} \right] \\ \text{OBJ} & []_1 \\ \text{OBL} & \left[\text{PRED} \text{ 'Mary'} \right] \\ \text{STMT-TYPE} & \text{int} \end{array} \right]$$

In the project grammars, many basic constructions are of this type. However, as we will see in the next section, there are times when parallelism is not possible and not desirable. Even in these cases, though, the grammars which can be parallel are; so, three of the languages might have one analysis, while three have another.

3.2 Justified Differences

Parallelism is not maintained at the cost of misrepresenting the language. This is reflected by the fact that the c-structures are not parallel because word order varies widely from language to language, although there are naming conventions for the nodes. Instead, the bulk of the parallelism is in the f-structure. However, even in the f-structure, situations arise in which what seems to be the same construction in different languages do not have the same analysis. An example of this is predicate adjectives, as in (7).

- (7) a. It is red.
 b. Sore wa akai. (Japanese)
 it TOP red
 'It is red.'

In English, the copular verb is considered the syntactic head of the clause, with the pronoun being the subject and the predicate adjective being an XCOMP. However, in Japanese, the adjective is the main predicate, with the pronoun being the subject. As such, these receive the non-parallel analyses seen in (8a) for Japanese and (8b) for English.

- (8) a. $\left[\begin{array}{ll} \text{PRED} & \text{'red<SUBJ>'} \\ \text{SUBJ} & \left[\text{PRED} \text{ 'pro'} \right] \end{array} \right]$
 b. $\left[\begin{array}{ll} \text{PRED} & \text{'be<XCOMP>SUBJ'} \\ \text{SUBJ} & \left[\text{PRED} \text{ 'pro'} \right]_1 \\ \text{XCOMP} & \left[\begin{array}{ll} \text{PRED} & \text{'red<SUBJ>'} \\ \text{SUBJ} & []_1 \end{array} \right] \end{array} \right]$

Another situation that arises is when a feature or construction is syntactically encoded in one language, but not another. In such cases, the information is only encoded in the languages that need it. The equivalence captured by parallel analyses is not, for example, translational equivalence. Rather, parallelism involves equivalence with respect to grammatical properties, e.g. construction types. One consequence of this is that a typologically consistent use of grammatical terms, embodied in the feature names, is enforced. For example, even though there is a tradition for referring to the distinction between the pronouns *he* and *she* as a gender distinction in English, this is a different distinction from the one called gender in languages like German, French, Urdu, and Norwegian, where gender refers to nominal agreement classes. Parallelism leads to the situation where the feature GEND occurs in German, French, Urdu, and Norwegian, but not in English and Japanese. That is, parallelism does not mean finding the same features in all languages, but rather using the same features in the same way in all languages, to the extent that they are justified there. A French example of grammatical gender is shown in (9); note that determiner, adjective, and participle agreement is dependent on the gender of the noun. The f-structure for the nouns *crayon* and *plume* are as in (10) with an overt GEND feature.

- (9) a. Le petit crayon est cassé. (French)
 the-M little-M pencil-M is broken-M.
 'The little pencil is broken.'
 b. La petite plume est cassée. (French)
 the-F little-F pen-F is broken-F.
 'The little pen is broken.'

$$(10) \left[\begin{array}{ll} \text{PRED} & \text{'crayon'} \\ \text{GEND} & \text{masc} \\ \text{PERS} & 3 \end{array} \right] \quad \left[\begin{array}{ll} \text{PRED} & \text{'plume'} \\ \text{GEND} & \text{fem} \\ \text{PERS} & 3 \end{array} \right]$$

F-structures for the equivalent words in English and Japanese will not have a GEND feature.

A similar example comes from Japanese discourse particles. It is well-known that Japanese has syntactic encodings for information such as honorification. The verb in the Japanese sentence (11a) encodes information that the subject is respected, while the verb in (11b) shows politeness from the writer (speaker) to the reader (hearer) of the sentence. The f-structures for the verbs in (11) are as in

(12) with RESPECT and POLITE features within the ADDRESS feature.

- (11) a. sensei ga hon wo oyomininaru.
teacher Nom book Acc read-Respect
'The teacher read the book.' (Japanese)
- b. seito ga hon wo yomimasu.
student Nom book Acc read-Polite
'The student reads the book.' (Japanese)

- (12) a.
$$\left[\begin{array}{ll} \text{PRED} & \text{'yomu<SUBJ,OBJ>} \\ \text{ADDRESS} & \left[\begin{array}{ll} \text{RESPECT} & + \end{array} \right] \end{array} \right]$$
- b.
$$\left[\begin{array}{ll} \text{PRED} & \text{'yomu<SUBJ,OBJ>} \\ \text{ADDRESS} & \left[\begin{array}{ll} \text{POLITE} & + \end{array} \right] \end{array} \right]$$

A final example comes from English progressives, as in (13). In order to distinguish these two forms, the English grammar uses a PROG $\pm$ feature within the tense/aspect system. (13b) shows the f-structure for (13a.ii).

- (13) a. John hit Bill. i. He cried.
ii. He was crying.
- b.
$$\left[\begin{array}{ll} \text{PRED} & \text{'cry<SUBJ>'} \\ \text{SUBJ} & \left[\begin{array}{ll} \text{PRED} & \text{'pro'} \end{array} \right] \\ \text{TNS-ASP} & \left[\begin{array}{ll} \text{TENSE} & \text{past} \\ \text{PROG} & + \end{array} \right] \end{array} \right]$$

However, this distinction is not found in the other languages. For example, (14a) is used to express both (13a.i) and (13a.ii) in German.

- (14) a. Er weinte. (German)
he cried
'He cried.'
- b.
$$\left[\begin{array}{ll} \text{PRED} & \text{'weinen<SUBJ>'} \\ \text{SUBJ} & \left[\begin{array}{ll} \text{PRED} & \text{'pro'} \end{array} \right] \\ \text{TNS-ASP} & \left[\begin{array}{ll} \text{TENSE} & \text{past} \end{array} \right] \end{array} \right]$$

As seen in (14b), the German f-structure is left underspecified for PROG because there is no syntactic reflex of it. If such a feature were posited, rampant ambiguity would be introduced for all past tense forms in German. Instead, the semantics will determine whether such forms are progressive.

Thus, there are a number of situations where having parallel analyses would result in an incorrect analysis for one of the languages.

3.3 One Language Shows the Way

Another type of situation arises when one language provides evidence for a certain feature space or type of analysis that is neither explicitly mirrored nor explicitly contradicted by another language. In theoretical linguistics, it is commonly acknowledged that what one language codes overtly may be harder to detect for another language. This situation has arisen in the ParGram project. Case features fall under this topic. German, Japanese, and Urdu mark NPs with overt case morphology. In comparison, English, French, and Norwegian make relatively little use of case except as part of the pronominal system. Nevertheless, the f-structure analyses for all the languages contain a case feature in the specification of noun phrases.

This "overspecification" of information expresses deeper linguistic generalizations and keeps the f-structural analyses as parallel as possible. In addition, the features can be put to use for the isolated phenomena in which they do play a role. For example, English does not mark animacy grammatically in most situations. However, providing a ANIM + feature to known animates, such as people's names and pronouns, allows the grammar to encode information that is relevant for interpretation. Consider the relative pronoun *who* in (15).

- (15) a. the girl[ANIM +] who[ANIM +] left
b. the box[ANIM +] who[ANIM +] left

The relative pronoun has a ANIM + feature that is assigned to the noun it modifies by the relative clause rules. As such, a noun modified by a relative clause headed by *who* is interpreted as animate. In the case of canonical inanimates, as in (15b), this will result in a pragmatically odd interpretation, which is encoded in the f-structure.

Teasing apart these different phenomena crosslinguistically poses a challenge that the ParGram members are continually engaged in. As such, we have developed several methods to help maintain parallelism.

3.4 Mechanics of Maintaining Parallelism

The parallelism among the grammars is maintained in a number of ways. Most of the work is done during two week-long project meetings held each year.

Three main activities occur during these meetings: comparison of sample f-structures, comparison of features and their values, and discussions of new or problematic constructions.

A month before each meeting, the host site chooses around fifteen sentences whose analysis is to be compared at the meeting. These can be a random selection or be thematic, e.g., all dealing with predicatives or with interrogatives. The sentences are then parsed by each grammar and the output is compared. For the more recent grammars, this may mean adding the relevant rules to the grammars, resulting in growth of the grammar; for the older grammars, this may mean updating a construction that has not been examined in many years. Another approach that was taken at the beginning of the project was to have a common corpus of about 1,000 sentences that all of the grammars were to parse. For the English, French, and German grammars, this was an aligned tractor manual. The corpus sentences were used for the initial f-structure comparisons. Having a common corpus ensured that the grammars would have roughly the same coverage. For example, they all parsed declarative and imperative sentences. However, the nature of the corpus can leave major gaps in coverage; in this case, the manual contained no interrogatives.

The XLE platform requires that a grammar declare all the features it uses and their possible values. Part of the Urdu feature table is shown in (16) (the notation has been simplified for expository purposes). As seen in (16) for QUANT, attributes which take other attributes as their values must also be declared. An example of such a feature was seen in (3b) for SPEC which takes QUANT and POSS features, among others, as its values.

- (16) PRON-TYPE: { pers poss null }.
 PROPER: { date location name title }.
 PSEM: { locational directional }.
 PTYPE: { sem noseml }.
 QUANT: { PRED QUANT-TYPE
 QUANT-FORM }.

The feature declarations of all of the languages are compared feature by feature to ensure parallelism. The most obvious use of this is to ensure that the grammars encode the same features in the same way. For example, at a basic level, one feature declaration might have specified GEN for gender while the others had chosen the name GEND; this divergence in naming is regularized. More interesting cases arise

when one language uses a feature and another does not for analyzing the same phenomena. When this is noticed via the feature-table comparison, it is determined why one grammar needs the feature and the other does not, and thus it may be possible to eliminate the feature in one grammar or to add it to another.

On a deeper level, the feature comparison is useful for conducting a survey of what constructions each grammar has and how they are implemented. For example, if a language does not have an ADEGREE (adjective degree) feature, the question will arise as to whether the grammar analyzes comparative and superlative adjectives. If they do not, then they should be added and should use the ADEGREE feature; if they do, then the question arises as to why they do not have this feature as part of their analysis.

Finally, there is the discussion of problematic constructions. These may be constructions that already have analyses which had been agreed upon in the past but which are not working properly now that more data has been considered. More frequently, they are new constructions that one of the grammars is considering adding. Possible analyses for the construction are discussed and then one of the grammars will incorporate the analysis to see whether it works. If the analysis works, then the other grammars will incorporate the analysis. Constructions that have been discussed in past ParGram meetings include predicative adjectives, quantifiers, partitives, and clefts. Even if not all of the languages have the construction in question, as was the case with clefts, the grammar writers for that language may have interesting ideas on how to analyze it. These group discussions have proven particularly useful in extending grammar coverage in a parallel fashion.

Once a consensus is reached, it is the responsibility of each grammar to make sure that its analyses match the new standard. As such, after each meeting, the grammar writers will rename features, change analyses, and implement new constructions into their grammars. Most of the basic work has now been accomplished. However, as the grammars expand coverage, more constructions need to be integrated into the grammars, and these constructions tend to be ones for which there is no standard analysis in the linguistic literature; so, differences can easily arise in these areas.

4 Conclusion

The experiences of the ParGram grammar writers has shown that the parallelism of analysis and implementation in the ParGram project aids further grammar development efforts. Many of the basic decisions about analyses and formalism have already been made in the project. Thus, the grammar writer for a new language can use existing technology to bootstrap a grammar for the new language and can parse equivalent constructions in the existing languages to see how to analyze a construction. This allows the grammar writer to focus on more difficult constructions not yet encountered in the existing grammars.

Consider first the Japanese grammar which was started in the beginning of 2001. At the initial stage, the work of grammar development involved implementing the basic constructions already analyzed in the other grammars. It was found that the grammar writing techniques and guidelines to maintain parallelism shared in the ParGram project could be efficiently applied to the Japanese grammar. During the next stage, LFG rules needed for grammatical issues specific to Japanese have been gradually incorporated, and at the same time, the biannual ParGram meetings have helped significantly to keep the grammars parallel. Given this system, in a year and a half, using two grammar writers, the Japanese grammar has attained coverage of 99% for 500 sentences of a copier manual and 95% for 10,000 sentences of an eCRM (Voice-of-Customer) corpus.

Next consider the Norwegian grammar which joined the ParGram group in 1999 and also emphasized slightly different goals from the other groups. Rather than prioritizing large textual coverage from the outset, the Norwegian group gave priority to the development of a core grammar covering all major construction types in a principled way based on the proposals in (Bresnan, 2001) and the inclusion of a semantic projection in addition to the f-structure. In addition, time was spent on improving existing lexical resources (> 80,000 lemmas) and adapting them to the XLE format. Roughly two man-years has been spent on the grammar itself. The ParGram cooperation on parallelism has ensured that the derived f-structures are interesting in a multilingual context, and the grammar will now serve as a basis for grammar development in other closely related Scandinavian languages.

Thus, the ParGram project has shown that it is possible to use a single grammar development plat-

form and a unified methodology of grammar writing to develop large-scale grammars for typologically different languages. The grammars' analyses show a large degree of parallelism, despite being developed at different sites. This is achieved by intensive meetings twice a year. The parallelism can be exploited in applications using the grammars: the fewer the differences, the simpler a multilingual application can be (see (Frank, 1999) on a machine-translation prototype using ParGram).

References

- Joan Bresnan. 2001. *Lexical-Functional Syntax*. Blackwell.
- Miriam Butt, Tracy Holloway King, María-Eugenia Niño, and Frédérique Segond. 1999. *A Grammar Writer's Cookbook*. CSLI Publications.
- Ann Copestake. 2002. *Implementing Typed Feature Structure Grammars*. CSLI Publications.
- Berthold Crysmann, Anette Frank, Bernd Keifer, St. Müller, Günter Neumann, Jakub Piskorski, Ulrich Schäfer, Melanie Siegel, Hans Uszkoreit, Feiyu Xu, Markus Becker, and Hans-Ulrich Krieger. 2002. An integrated architecture for shallow and deep parsing. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics, University of Pennsylvania*.
- Anette Frank. 1999. From parallel grammar development towards machine translation. In *Proceedings of MT Summit VII*, pages 134–142.
- John T. Maxwell, III and Ron Kaplan. 1993. The interface between phrasal and functional constraints. *Computational Linguistics*, 19:571–589.
- Stefan Riezler, Tracy Holloway King, Ronald Kaplan, Dick Crouch, John T. Maxwell, III, and Mark Johnson. 2002. Parsing the wall street journal using a lexical-functional grammar and discriminative estimation techniques. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics, University of Pennsylvania*.
- Paul Schmidt, Sibylle Rieder, Axel Theofilidis, and Thierry Declerck. 1996. Lean formalisms, linguistic theory, and applications: Grammar development in alep. In *Proceedings of COLING*.
- Wolfgang Wahlster, editor. 2000. *Verbmobil: Foundations of Speech-to-Speech Translation*. Springer.