

**Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar**

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

The papers are published under a CC-BY license:
<http://creativecommons.org/licenses/by/4.0/>

Contents

Editor's note	3
 I Contributions to the Main Conference	 4
Liesbeth Augustinus: A construction-based analysis of Dutch verb clusters	6
Berthold Crysmann, Alain Kihm: Deriving reversal in Old French nominal inflection	20
Christian M. Curtis: Valence-changing morphology via lexical rule composition	36
Anke Holler, Markus Steinbach: Sign language agreement: A constraint-based perspective	49
Kristen Howell, Olga Zamaraeva, Emily M. Bender: Nominalized clauses in the Grammar Matrix	66
David Lahm: Plural in Lexical Resource Semantics	87
Nurit Melnik: Symmetry and asymmetry in the Hebrew copula construction	108
David Moeljadi, Francis Bond: HPSG analysis and computational implementation of Indonesian passives	127
Elizabeth Nielsen, Emily M. Bender: Modeling adnominal possession in multilingual grammar engineering	138
 II Contributions to the Workshop	 152
Antonio Machicao y Priemer, Paola Fritz-Huechante: Korean and Spanish psych-verbs: Interaction of case, theta-roles, linearization, and event structure in HPSG	153
Shûichi Yatabe, Kei Tanigawa: The fine structure of clausal right-node raising constructions in Japanese	174

Editor's note

The 25th International Conference on Head-Driven Phrase Structure Grammar (2018) was held at the University of Tokyo, Komaba Campus.

The conference featured 2 invited talks, 13 papers, and 3 posters selected by the program committee (Anne Abeillé, Doug Arnold, Daisuke Bekki, Olivier Bonami, Francis Bond, Gosse Bouma, George Broadwell, Rui Chaves, Philippa Cook, Berthold Crysmann, Kordula De Kuthy, Antske Fokkens, Petter Haugereid, Fabiola Henri, Anke Holler, Gianina Iordachoaia, Jong-Bok Kim, Jean-Pierre Koenig, Yusuke Kubota, David Lahm, Robert D. Levine, Yo Matsumoto, Nurit Melnik, Philip Miller, Stefan Müller, Tsuneko Nakazawa, Joanna Nykiel, Rainer Osswald, Gerald Penn, Frank Richter (chair), Manfred Sailer, Pollet Samvellian, Sanghoun Song, Frank van Eynde, Stephen Wechsler, Shûichi Yatabe, Eun-Jung Yoo).

There was a workshop on *The Clause Structure of Japanese and Korean* with three talks and two invited talks.

We want to thank the program committees for putting this nice program together.

Thanks go to Shuichi Yatabe, who was in charge of local arrangements, and his assistants Kei Tanigawa, Tsuneko Nakazawa, Morine Kondo, Mayu Kawakita, Takeshi Kishiyama, and Fuga Terasaki.

As in the past years the contributions to the conference proceedings are based on the five page abstract that was reviewed by the respective program committees, but there is no additional reviewing of the longer contribution to the proceedings. To ensure easy access and fast publication we have chosen an electronic format.

The proceedings include all the papers except the one by Anne Abeillé & Shrita Hassamal, Gabrielle Aguila"=Multner & Berthold Crysmann, Makoto Kanazawa, Yusuke Kubota, and Frank Van Eynde.

Part I

Contributions to the Main Conference

A construction-based analysis of Dutch verb clusters

Liesbeth Augustinus 

KU Leuven

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 6–19

Keywords: Dutch, verb clusters, HPSG

Augustinus, Liesbeth. 2018. A construction-based analysis of Dutch verb clusters. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 6–19. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.1.  CC-BY 4.0

Abstract

Dutch is well-known for the formation of verb clusters. A characteristic aspect of such constructions is that the order of the verbs may differ from the order in which they are selected. Across the Dutch language area verb clusters show different types of word order variation.

This paper proposes a constructivist account of word order variation in Dutch verb clusters. Linearization is not modelled in terms of the GVOR feature, after Kathol (2000). Instead, it crucially relies on the bidimensional phrase hierarchy initiated by Ginzburg & Sag (2000), which is extended for the analysis of constructions with verb clusters. This proposal accounts for the most common instances of word order variation in Dutch verb clusters, and it can be easily adapted to model a specific variety or dialect.

1 Introduction

In Dutch, verbs form a cluster in verb-final clauses with two or more verbs, as in (1), and in verb-initial clauses with three or more verbs, as in (2).¹

- (1) ... dat ze die wedstrijd heeft₁ gewonnen₂.
that she that competition has won
'... that he has won that competition.'
- (2) Ze zal die wedstrijd kunnen₁ winnen₂.
She will that competition can win
'She will be able to win that competition.'

The linear order of the verbs in a cluster canonically coincides with the order of selection, i.e. a verb selects its verbal complement to the right.² Alternative orders are possible though. In constructions with a past or passive participle, the participle may occupy any position in the cluster, but the order of the other verbs must be ascending:³

- (3) In de tussentijd zouden de twee belangrijkste getuigen ... moeten₁
in the meantime would the two most-important witnesses ... must
worden₂ gehoord₃.
be heard
'In the mean time the two most important witnesses would have to be heard.' [LASSY]

[†]I thank the audience of the HPSG 2018 conference (Tokyo) for their comments.

¹Verb-initial clauses comprise verb-first and verb-second clauses.

²The order of selection is indicated by subscripts, the hierarchically highest verb being 1.

³The examples in (3)-(8) are taken from the CGN treebank for spoken Dutch (Oostdijk et al., 2002) and the LASSY treebank for written Dutch (van Noord et al., 2013).

- (4) er zijn toch zo'n paar boeken die ge moet₁ gelezen₃ hebben₂
 there are actually such couple books that you must read have
 in uw leven.
 in your life
 'actually there are a couple of books that you should have read in your life.'
 [CGN]
- (5) Diversiteit in onze samenleving zou nog veel meer benadrukt₃
 diversity in our society should still much more focussed
 moeten₁ worden₂.
 must be
 'Diversity in our society should be much more focussed on.' [LASSY]

A second set of constructions that show word order variation are constructions with a substitute infinitive or *Ersatzinfinitiv*. In (6) the verb *kunnen* 'can' appears as an infinitive and not as the past participle *gekund* 'could'. For most speakers of Dutch the verbs always appear in the canonical ascending order, but in Belgian Dutch some speakers also allow the order in (7), in which the auxiliary of the perfect appears at the end of the cluster in verb-final clauses. Such constructions are also known as *Oberfeldumstellung*. In German it is obligatory in a number of cases, but in Dutch the phenomenon is always optional and not grammatical for all speakers.

- (6) Pas nu hebben we dat ook kunnen zien in de hersenen.
 only now have we that also can.IPP see in the brains
 'Only now we have been able to see that in the brains.' [LASSY]
- (7) ... terwijl dat 'k ik naar buiten gaan₂ kijken₃ ben₁.
 while that I I to outside go.IPP look am
 '... while I was going to look outside.' [CGN]

A third type of word order alternation includes two-verb clusters with a finite modal verb, such as (8).

- (8) ... om ervoor te zorgen dat dit nooit meer gebeuren₂ zal₁.
 to there-for to make-sure that this never again happen will
 '... in order to make sure that this will never happen again.' [LASSY]

This type of variation is only possible in verb-final clauses, as the finite verb needs to be part of the cluster. In longer verb clusters of this type, word order variation is not allowed (9), and also in constructions with non-modal finite verbs the descending order is ungrammatical (10):

- (9) * ... om ervoor te zorgen dat dit nooit meer kunnen₂ gebeuren₃
to there-for to make-sure that this never again can happen
zal₁.
will
intended: ‘... in order to make sure that this will never be able to happen again.’
- (10) * ... om ervoor te zorgen dat hij dit nooit meer gebeuren₂ laat₁.
to there-for to make-sure that he this never again happen let
intended: ‘... in order to make sure that he will never let this happen again.’

If the verbs in (9) and (10) are put in the canonical, ascending order, the sentences are well-formed.

In sum, Dutch syntax is marked by verb cluster formation, which shows word order variation that does not entail a change in meaning. There are different types of word order variation, depending on the form of the verbs in the cluster (infinitival, participial), and the type of the selecting verb (e.g. modal verb).

Section 2 discusses previous accounts of word order variation, while section 3 presents a new model. Section 4 concludes.

2 Previous models of word order variation

The literature on West Germanic verb clusters is vast. Some influential HPSG analyses of verb clusters are Hinrichs & Nakazawa (1994), Bouma & van Noord (1998), Kathol (2000), and Müller (2002).

In HPSG verb clusters are canonically treated as binary-branching structures modelled in terms of *argument inheritance*, i.e. the non-subject arguments of unsaturated verbal complements are treated in a similar way as raised subjects, cf the lexical constraint in (11), after Hinrichs & Nakazawa (1994). If $\boxed{A}$ is an empty list, the constraint is similar to the one for subject raising proposed in Ginzburg & Sag (2000, 22).

$$(11) \left[\text{ARG-ST } \langle \boxed{1} \rangle \oplus \boxed{A} \oplus \left\langle \left[\text{LOCAL} \mid \text{CAT } \begin{bmatrix} \text{HEAD } verb \\ \text{SUBJ } \langle \boxed{1} \rangle \\ \text{COMPS } \boxed{A} \end{bmatrix} \right] \right\rangle \right]$$

The application of (11) to (1) is illustrated in Figure 1. The unsaturated complement of *gewonnen* ‘won’ ($\boxed{2}$) is shared with the selecting verb *heeft* ‘has’, before it is propagated to the mother node. So both the verbal complement and the unsaturated complement appear on the COMPS lists of the selecting verb.

In order to model word order variation in German verb clusters, Hinrichs & Nakazawa (1994) employ the binary head feature FLIP. Kathol (2000) replaces

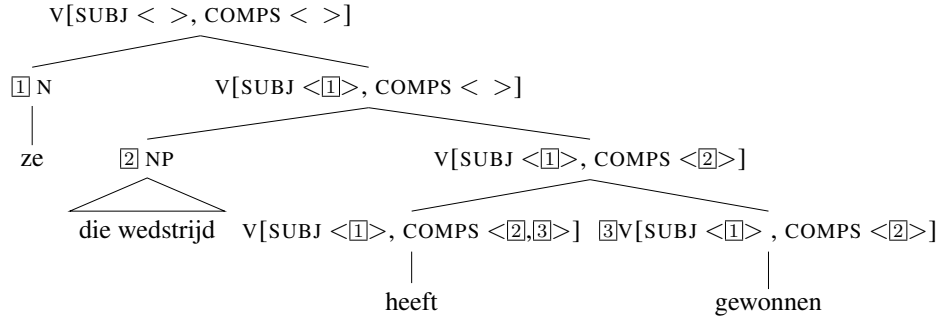


Figure 1: Argument Inheritance

Hinrichs and Nakazawa’s FLIP feature by the head feature $G(O)V(ERN)OR$ in order to model the word order of the verbs in the cluster. If a verb has the feature $[GVOR \rightarrow]$, its governor should appear to its right, while the governor of verbs with the feature $[GVOR \leftarrow]$ should appear to the left (e.g. in the case of German Oberfeldumstellung). Applied to the Dutch construction in (4) it yields the tree structure in Figure 2.

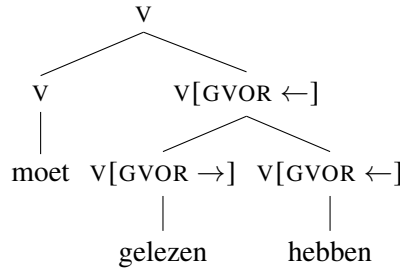


Figure 2: The GVOR feature

Gelezen ‘read’ is selected by *hebben* ‘have’ on the right, which is why it has the feature $[GVOR \rightarrow]$. *Hebben* on its turn is selected by the finite verb *moet* ‘must’ and has the feature $[GVOR \leftarrow]$. As GVOR is a head feature, $[GVOR \leftarrow]$ is shared with the mother node following the head feature principle.

While constructions with auxiliary flip pose no problem for a binary-branching treatment of verb clusters, constructions such as (5), in which the selecting verb does not appear next to its complement, do. In order to account for all linearization possibilities, Bouma & van Noord (1998) analyse verb clusters as flat tree structures. The downside of their approach is that they need additional features and complex word order constraints in order to avoid overgeneration compared to binary-branching analyses.

Kathol (2000) tackles the problem in a different way. He employs an additional feature in order to model the linear order of verb clusters, i.e. the $DOM(AIN)$ feature. The order of the elements in DOM may differ from the order of the elements of the tree structure. Also his approach overgenerates for Dutch. Kathol

assumes that the GVOR value of Dutch infinitival complements is underspecified as [GVOR *dir*]. In this way he deals with verbs that can select their complement to the left or to the right, e.g. Dutch *wil lezen* versus *lezen wil* ‘wants to read’ (Kathol, 2000, 199–200). For Dutch past participles, this is what you want, cf (2), but for infinitival complements, this assumption overgenerates. As mentioned in section 1, an infinitival complement may only precede its selector if it is selected by a finite modal. In clusters of more than two verbs, or if a non-modal verb selects the infinitive, the only grammatical order is the ascending one. An accurate model of word order variation in Dutch should take this into account.

In what follows, it will be illustrated that Dutch verb clusters can be modelled in a binary-branching analysis, in which the linear order of the verbs in the cluster is similar to the order in which they appear in the phrase structure tree.

3 A constructivist proposal

3.1 Complement raising

In the argument inheritance approach discussed in section 2, raised complements are treated in a similar way as raised subjects. In Van Eynde & Augustinus (2013) and Augustinus (2015) it is motivated that subject and complement raising are different phenomena.⁴ While subject raising is modelled using the canonical lexical constraint, a phrasal constraint is employed for raised complements. The *Complement Raising Principle* (CRP) in (12) states that in a headed phrase, the COMPS list of the non-head daughter is added to the COMPS list of the mother.⁵

$$(12) \quad hd-ph \Rightarrow \left[\begin{array}{l} \text{SYNSEM} \mid \text{LOC} \mid \text{CAT} \mid \text{COMPS } \boxed{A} \oplus \boxed{B} \\ \text{HD-DTR} \mid \text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{COMPS } \boxed{A} \\ \text{NONHD-DTR} \mid \text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{COMPS } \boxed{B} \end{array} \right]$$

Cancellation of elements from the COMPS list is modelled in the definition of phrases of type *head-complement*. The constraint is given in (13), after Sag et al. (2003, 96-97). As *head-complement phrase* is a subtype of *headed-phrase*, it follows that the COMPS list can expand and shrink at the same time.

The application of (12) and (13) to (1) is illustrated in Figure 3. In contrast to the argument inheritance approach Hinrichs-Nakazawa style in Figure 1, the unsaturated complement of *gewonnen* ‘won’ is not shared with the selecting verb *heeft* ‘has’, but it is directly propagated to the mother node. Only the verbal complement appears on the COMPS lists of the selecting verb.

⁴Arguments against the lexical constraint in (11) include the occurrence of complement raising without subject raising, interaction with the binding principles and the passive lexical rule.

⁵The CRP is a phrasal constraint and is, hence, a very powerful mechanism. In order to avoid overgeneration, complement raising is blocked in CPs, V-initial VPs, and P-initial PPs. For a detailed discussion, see Augustinus (2015).

$$(13) \quad hd-comp-ph \Rightarrow \left[\begin{array}{l} \text{SYNSEM} \mid \text{LOC} \mid \text{CAT} \mid \text{COMPS} \quad \boxed{A} \\ \text{HD-DTR} \mid \text{SYNSEM} \mid \text{LOC} \mid \text{CAT} \mid \text{COMPS} \quad \boxed{A} \oplus \langle \boxed{1} \rangle \\ \text{NONHD-DTR} \mid \text{SS} \quad \boxed{1} \text{ synsem} \end{array} \right]$$

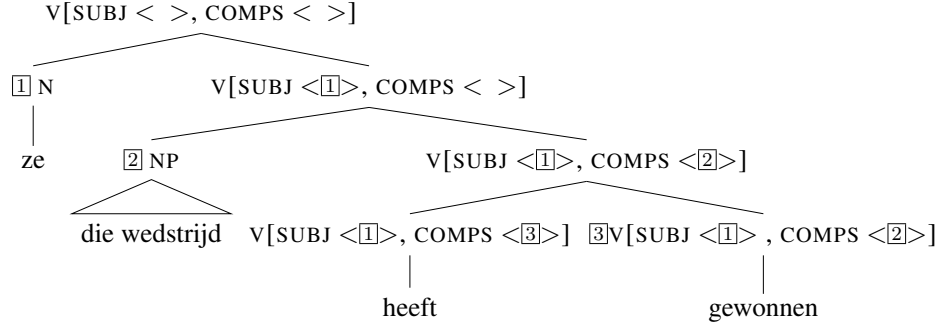


Figure 3: Complement raising

The complement raising analysis will be assumed for the verb clusters dealt with in this paper.

3.2 Word order variation

We discard the use of the GVOR feature for the analysis of word order variation in Dutch verb clusters for three reasons. First, most linearization possibilities depend on the VFORM of the verbs in the cluster (e.g. clusters with participles have different linearization possibilities compared to clusters without a participle). Second, the type of the selecting verb is important (e.g. substitute infinitives in constructions with *Oberfeldumstellung*). Third, the length of the constructions has an influence on certain linearization patterns (e.g. the constructions in which an infinitival complements precedes a finite modal verb). In order to account for all types of verb clusters in Dutch we opt for a constructivist account.

Ginzburg & Sag (2000) advocate a constructivist version of HPSG, in which they propose a bidimensional type hierarchy for phrase types. The main distinction concerns the difference between clausality and headedness, cf Figure 4.

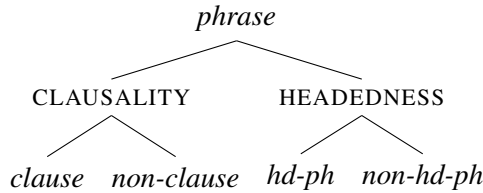


Figure 4: Phrase type hierarchy (Ginzburg & Sag, 2000)

The clausality dimension distinguishes clauses from non-clauses. The headedness dimension differentiates headed phrases (*hd-ph*), i.e. phrases with a head-daughter

such as head-complement phrases, from non-headed phrases (e.g. coordinate structures).

In order to accurately model Dutch verb clusters, we extend the phrase type hierarchy proposed in Ginzburg & Sag (2000, 38-45). The proposed extension deals with the non-clause type. It is given in Figure 5.

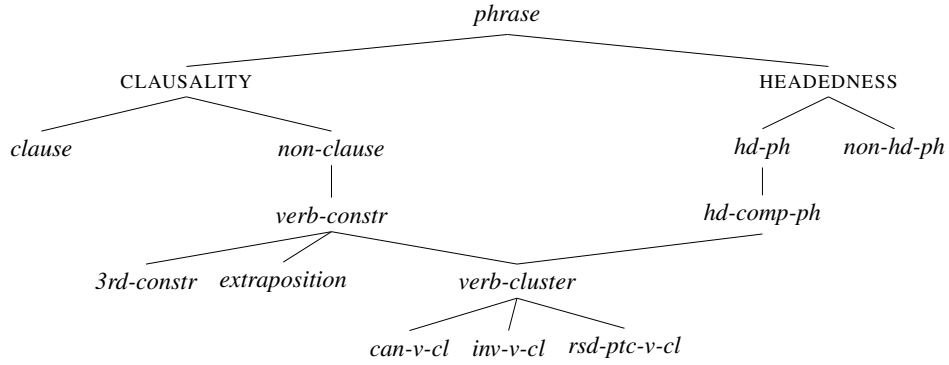


Figure 5: Extended phrase type hierarchy

The types that are relevant in this discussion are *verb-constr(uction)* and *verb-cluster*. The former includes phrases with a head daughter of type *verb* and a non-head daughter which has a nonfinite verb as its head, cf (14).

$$(14) \quad verb-constr \Rightarrow \left[\begin{array}{l} SS \mid LOC \mid CAT \mid HEAD \quad verb \\ NON-HD-DTR \mid SS \mid LOC \mid CAT \mid HEAD \quad \left[\begin{array}{l} verb \\ VFORM \quad nfin \end{array} \right] \end{array} \right]$$

The type *verb-constr* not only comprises instances of verb clusters, such as the examples in (1-8), but also extraposition (15) and the third construction (16).

- (15) ... net nu hij dacht een goede indruk te maken.
 just now he thought a good impression to make
 ‘... just as he thought to make a good impression.’ [CGN]
- (16) ... dat men daarin moet trachten het juiste evenwicht te zoeken.
 that one there-in has try the right balance to search
 ‘... and I think that one has to try to find the right balance in that.’ [CGN]

In (15) the verb *denken* ‘think’ selects its verbal complement *te maken* ‘to make’ in the Nachfeld. The same holds for *trachten* ‘try’ in (16), but in this construction the object *daarin* ‘in that’ belonging to the extraposed VP *te zoeken* ‘to search’ appears in the Mittelfeld.

The hierarchy in Figure 5 makes use of multiple inheritance. The type *verb-cluster* is a subtype from both *verb-constr* and *hd-comp-ph* and therefore inherits properties from those types. Its defining property is given in (17).

$$(17) \text{ verb-cluster} \Rightarrow [\text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{HEAD } cl\text{-verb}]$$

A verb cluster is a construction which has a clustering verb (*cl-verb*) as its head. Clustering verbs are verbs that may select another verb in a verb cluster, as opposed to non-clustering verbs (*non-cl-verb*), cf Figure 6.

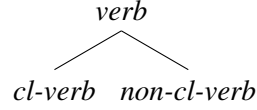


Figure 6: *verb* type hierarchy

In (3), for example, the verbs *moeten* ‘must’ and *worden* ‘be’ are clustering verbs. *Gehoord* ‘heard’ is also part of the verb cluster, but it is not a clustering verb since it does not select another verb.

The set of clustering verbs in Dutch is limited. Augustinus (2015) has identified different types of clustering verbs, such as modals, perception verbs, auxiliaries of the perfect etc. We introduce the feature *VTYP*E to differentiate between those types. Its values are presented in Figure 7.

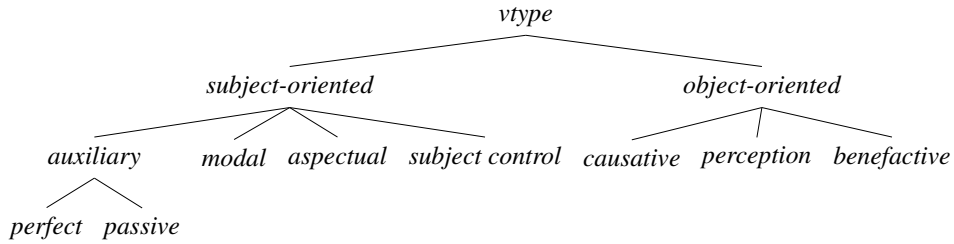


Figure 7: *vtype* type hierarchy

A characteristic aspect of Dutch clustering verbs is that they never appear as a participle, as participles cannot select another verb in the cluster. The formal definition of clustering verbs (*cl-verb*) is given in (18).

$$(18) \begin{bmatrix} cl\text{-verb} \\ \text{VFORM} & \neg ptc \\ \text{VTYP}E & vtype \end{bmatrix}$$

In order to model word order variation in verb clusters, three subtypes are introduced: *canonical verb clusters*, *inverted verb clusters*, and *raised participle verb clusters*.

3.2.1 Canonical verb clusters

The most general cluster type is the **canonical verb cluster** (*can-v-cl*). Its formal properties are given in (19).

$$(19) \quad \text{can-v-cl} \Rightarrow \begin{bmatrix} \text{PHON } \boxed{A} \oplus \boxed{B} \\ \text{HD-DTR} \mid \text{PHON } \boxed{A} \\ \text{NON-HD-DTR} \mid \text{PHON } \boxed{B} \end{bmatrix}$$

Canonical verb clusters inherit from (17) that they have a head daughter with a clustering verb as its head, and a nonfinite non-head daughter (which may be a word or a phrase). As indicated in PHON, the clustering verb ($\boxed{A}$) appears before its non-head daughter ($\boxed{B}$). (19) accounts for constructions with the canonical (ascending) word order such as (1), (2), and (3), repeated in (20-22).

- (20) ... dat ze die wedstrijd heeft₁ gewonnen₂.
 that she that competition has won
 ‘... that he has won that competition.’
- (21) Ze zal die wedstrijd kunnen₁ winnen₂.
 She will that competition can win
 ‘She will be able to win that competition.’
- (22) In de tussentijd zouden de twee belangrijkste getuigen ... moeten₁
 in the meantime would the two most-important witnesses ... must
 worden₂ gehoord₃.
 be heard
 ‘In the mean time the two most important witnesses would have to be
 heard.’ [LASSY]

3.2.2 Inverted verb clusters

The second cluster type is the **inverted verb cluster** (*inv-v-cl*):

$$(23) \quad \text{inv-v-cl} \Rightarrow \begin{bmatrix} \text{PHON } \boxed{B} \oplus \boxed{A} \\ \text{HD-DTR} \begin{bmatrix} \text{PHON } \boxed{A} \\ \text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{HEAD} \begin{bmatrix} \text{cl-verb} \\ \text{VTYPE } \text{auxiliary} \vee \text{modal} \end{bmatrix} \end{bmatrix} \\ \text{NON-HD-DTR} \mid \text{PHON } \boxed{B} \end{bmatrix}$$

(23) states that the head-daughter should be a clustering verb of type *auxiliary* or *modal*.⁶ The verbal non-head daughter appears in front of its head sister, as indicated in the PHON feature. In order to account for constructions such as (4), (7) and (8), we introduce three subtypes.

⁶Clustering verbs of type *auxiliary* include the auxiliaries of the perfect and the passive, cf Figure 7.

Participle-inverted clusters If the selecting verb is of type auxiliary, the non-head daughter should be a past or passive participle to form constructions in which the past participle occurs right in front of its selector, such as (24) and (4), repeated in (25).

- (24) ... toen ze voor de eerste keer rechtstreeks verkozen₂ werden₁.
 when they for the first time directly elected were
 ‘... when they were directly elected for the first time’ [CGN]
- (25) er zijn toch zo’n paar boeken die ge moet₁ gelezen₃ hebben₂
 there are actually such couple books that you must read have
 in uw leven.
 in your life
 ‘actually there are a couple of books that you should have read in your life.’
 [CGN]

The formal constraint accounting for such constructions is given in (26).

- (26) $ptc-inv-v-cl \Rightarrow \left[\begin{array}{l} HD-DTR \mid SS \mid LOC \mid CAT \mid HEAD \mid VTYPE \text{ auxiliary} \\ NON-HD-DTR \mid HEAD \mid VFORM \text{ ptc} \end{array} \right]$

Auxiliary-inverted clusters In order to deal with *Oberfeldumstellung*, the head daughter should be of type auxiliary, whereas the non-head daughter should be a canonical verb cluster to ensure the other verbs appear in the ascending order:

- (27) $aux-inv-v-cl \Rightarrow \left[\begin{array}{l} HD-DTR \mid SS \mid LOC \mid CAT \mid HEAD \mid VTYPE \text{ auxiliary} \\ NON-HD-DTR \text{ can-v-cl} \end{array} \right]$

The constraint in (27) yields constructions like (7), repeated in (28), and excludes ungrammatical orders such as * *kijken₃ gaan₂ ben₁*.

- (28) ... terwijl dat ’k ik naar buiten gaan₂ kijken₃ ben₁.
 while that I I to outside go.IPP look am
 ‘... while I was going to look outside.’ [CGN]

For variants of Dutch that do not accept *Oberfeldumstellung*, the non-head daughter of (23) should be of type *word*.

Modal-inverted clusters The third subtype of inverted verb clusters includes constructions in which a finite modal verb follows its infinitival complement, as in (8), repeated in (29).

- (29) ... om ervoor te zorgen dat dit nooit meer gebeuren₂ zal₁.
 to there-for to make-sure that this never again happen will
 ‘... in order to make sure that this will never happen again.’ [LASSY]

The formal definition of such verb clusters in (30) states that the selecting verb should be [VTYPE *modal*]. This excludes constructions in which another type of verb follows its infinitival complement, such as the causative verb *laten* ‘let’ in (10).

In addition the non-head daughter should be an infinitive of type *non-cl-verb*. This avoids the embedding of longer clusters, which would yield ungrammatical constructions such as **kunnen gebeuren zal* ‘will be able to happen’ in (9).

$$(30) \quad \text{mod-inv-v-cl} \Rightarrow \left[\begin{array}{l} \text{HD-DTR} \mid \text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{HEAD} \mid \text{VTYPE } \textit{modal} \\ \text{NON-HD-DTR} \mid \text{HEAD} \left[\begin{array}{l} \textit{non-cl-verb} \\ \text{VFORM } \textit{inf} \end{array} \right] \end{array} \right]$$

3.2.3 Raised participle verb clusters

The third cluster type is the **raised participle verb cluster** (*rsd-ptc-v-cl*). It deals with constructions in which the main verb does not appear next to its head, such as the construction in (5), repeated in (31).

- (31) Diversiteit in onze samenleving zou nog veel meer benadrukt₃
diversity in our society should still much more focussed
moeten₁ worden₂.
must be
‘Diversity in our society should be much more focussed on.’ [LASSY]

(31) is treated as a construction in which the participle *benadrukt* ‘focussed’ is raised. As the CRP in (12) does not put any restrictions on the type of complement that can be raised, it accounts for this kind of constructions. The formal specifications of clusters with a raised past participle are given in (32).

$$(32) \quad \text{rsd-ptc-v-cl} \Rightarrow \left[\begin{array}{l} \text{PHON } [\overline{B}] \oplus [\overline{A}] \\ \text{HD-DTR} \left[\begin{array}{l} \textit{can-v-cl} \\ \text{PHON } [\overline{A}] \end{array} \right] \\ \text{NON-HD-DTR} \left[\begin{array}{l} \textit{word} \\ \text{PHON } [\overline{B}] \\ \text{SS} \mid \text{LOC} \mid \text{CAT} \mid \text{HEAD} \left[\begin{array}{l} \textit{verb} \\ \text{VFORM } \textit{ptc} \end{array} \right] \end{array} \right] \end{array} \right]$$

(32) accounts for the combination of a participial non-head daughter with a *can-v-cl* head daughter. As only participles can occur in a raised position in Dutch verb clusters, the non-head daughter should be a participle. The requirement that the head daughter should be a canonical verb cluster accounts for the fact that the order of the verbs in the cluster is ascending. It furthermore avoids spurious ambiguity between the *rsd-ptc-v-cl* construction and the *inv-v-cl* construction in which a past participle occurs right in front of the selecting verb, as in (4).

Another reason that raised constructions need to be differentiated from inverted clusters with a past participle, is that some varieties of Dutch accept a raised participle construction, but not an inverted verb cluster in constructions with more than two verbs.⁷ For those varieties, one could restrict the inverted cluster construction with a past participle to finite constructions, in a similar way as the modal-inverted verb clusters discussed in section 3.2.2.

4 Conclusion

This paper proposes a constructivist account of word order variation in Dutch verb clusters. In this model the linear order of the verbs in the cluster is similar to the order in which they appear in the phrase structure tree. Linearization is not modelled in terms of the GVOR feature of the verbal complement. Instead, it crucially relies on the bidimensional phrase hierarchy initiated by Ginzburg & Sag (2000), which is extended for the analysis of constructions with verb clusters, cf Figure 8. This proposal accounts for the most common instances of word order variation in Dutch verb clusters, but it can be easily adapted in order to model a specific variety or dialect.

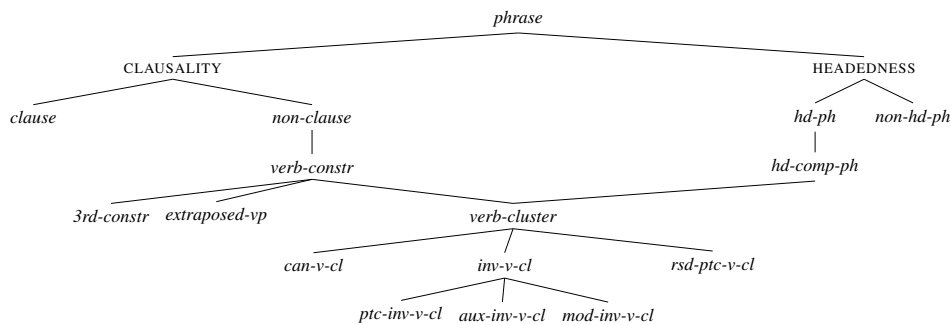


Figure 8: Extended phrase type hierarchy (bis)

References

- Augustinus, Liesbeth. 2015. *Complement Raising and Cluster Formation in Dutch. A Treebank-supported Investigation* LOT Dissertation Series 413. Utrecht: LOT.
- Barbiers, Sjef. 2005. Word order variation in three-verb clusters and the division of labour between generative linguistics and sociolinguistics. In Leonie Cornips & Karen Corrigan (eds.), *Syntax and Variation. Reconciling the Biological and the Social*, chap. 10, 233–264. Amsterdam/Philadelphia, PA: John Benjamins.

⁷For instance, 1-3-2 constructions are typically accepted in Belgian Dutch, whereas in the Netherlands they generally prefer 3-1-2 constructions (next to the canonical 1-2-3 constructions), see Barbiers (2005).

- Bouma, Gosse & Gertjan van Noord. 1998. Word Order Constraints on Verb Clusters in German and Dutch. In Erhard Hinrichs, Andreas Kathol & Tsuneko Nakazawa (eds.), *Complex Predicates in Nonderivational Syntax*, vol. 30 Syntax and Semantics, 43–72. New York: Academic Press.
- Ginzburg, Jonathan & Ivan Sag. 2000. *Interrogative investigations*. Stanford, CA: CSLI.
- Hinrichs, Erhard & Tsuneko Nakazawa. 1994. Linearizing AUXs in German Verbal Complexes. In John Nerbonne, Klaus Netter & Carl Pollard (eds.), *German in Head-Driven Phrase Structure Grammar*, 11–37. Stanford, CA: CSLI Publications.
- Kathol, Andreas. 2000. *Linear Syntax*. Oxford University Press.
- Müller, Stefan. 2002. *Complex predicates: Verbal complexes, resultative constructions, and particle verbs in German* (Studies in Constraint-Based Lexicalism 13). Stanford, CA: CSLI Publications.
- Oostdijk, Nelleke, Wim Goedertier, Frank Van Eynde, Lou Boves, Jean-Pierre Martens, Michael Moortgat & Harald Baayen. 2002. Experiences from the Spoken Dutch Corpus Project. In Manuel Gonzalez Rodriguez & Carmen Paz Saurez Araujo (eds.), *Proceedings of the Third International Conference on Language Resources and Evaluation (LREC2002)*, 340–347. Las Palmas.
- Sag, Ivan A., Thomas Wasow & Emily M. Bender. 2003. *Syntactic Theory. A Formal Introduction*. Stanford, CA: CSLI Publications 2nd edn.
- Van Eynde, Frank & Liesbeth Augustinus. 2013. Why and How to Differentiate Complement Raising from Subject Raising in Dutch. In Stefan Müller (ed.), *Proceedings of the 20th International Conference on Head-Driven Phrase Structure Grammar*, 222–242. Stanford, CA: CSLI Publications.
- van Noord, Gertjan, Gosse Bouma, Frank Van Eynde, Daniël de Kok, Jelmer van der Linde, Ineke Schuurman, Erik Tjong Kim Sang & Vincent Vandeghinste. 2013. Large Scale Syntactic Annotation of Written Dutch: Lassy. In Peter Spyns & Jan Odijk (eds.), *Essential Speech and Language Technology for Dutch: Results by the STEVIN programme*, 147–164. Springer.

Deriving reversal in Old French nominal inflection

Berthold Crysmann 

CNRS, Laboratoire de linguistique formelle

Alain Kihm 

CNRS, Laboratoire de linguistique formelle

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 20–35

Keywords: Morphological reversal, Old French, Declension, HPSG, IBM, Overabundance

Crysmann, Berthold & Alain Kihm. 2018. Deriving reversal in Old French nominal inflection. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 20–35. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.2.

 CC-BY 4.0

Abstract

In this paper, we study Old French declension, a system which exhibits the theoretically challenging phenomenon of morphological reversal (Baerman, 2007). Furthermore, the declension system of Old French only recognises a single exponent *-s*, which marks different case/number combinations in different paradigms, contrasting with the unmarked form. We show that reversal is only one of several syncretism patterns found in the language and propose that Old French declension is best understood in terms of two systematic syncretisms: a natural split between singular and plural for feminines, and a Paninian split for masculines that systematically marks the objective plural. Reversal, and other seemingly morphomic splits arise as a result of idiosyncrasy in the *NOM.SG* cell, comprising inflection class-specific *-s* marking, as well as stem alternation and overabundance. We provide a formal analysis in terms of Information-based Morphology (Crysmann & Bonami, 2016) that effortlessly captures the systematic splits, as well as the variation in the nominative singular. We suggest that the high degree of idiosyncrasy in this cell paired with the reduced frequency of overt nominative NPs when compared to objective NPs may serve to explain why the system was actually quite short-lived.

Among syncretism patterns, morphological reversals must certainly be regarded as one of the theoretically more challenging types (see Baerman, 2007, for a survey). In Old French, the majority of masculine nouns show a pattern where the distribution of unmarked and *s*-marked forms in the plural is reversed in the singular, as illustrated in Table 1. Historically, this pattern came about as a result of regular sound change from Latin via Late Latin to Old French: deletion of accusative singular /m/ was already lost in spoken Latin at the time of the Republic and subsequent deletion of unstressed vowels in the transition from Late Latin to Old French neutralised the contrast between accusative singular ($\emptyset < -u < -um$) and the nominative plural ($\emptyset < -i$) in the *o*-declension, as well as between nominative singular (*-s < -us*) and accusative plural (*-s < -ōs*).

	SG	PL
NOM	murs	mur
OBJ	mur	murs

Table 1: Reversal in Old French (Kihm, 2017, p. 41)

Reversals contrast with more well-behaved syncretism patterns such as motivated syncretism (see F1 in Table 3), which can easily be captured by underspecification,

[†]The research reported on in this paper has been partially carried out within the excellency cluster (LabEx) “Empirical Foundations in Linguistics”, supported by a public grant overseen by the French National Research Agency (ANR) as part of the “Investissements d’Avenir” program (reference: ANR-10-LABX-0083). We are especially grateful to Jean-Pierre Koenig for comments on an earlier version of this paper. Finally, we would like to thank the audience of the HPSG 2018 meeting in Tokyo, in particular, Anne Abeillé, Doug Arnold, and Nurit Melnik.

or so-called Paninian splits, where one or more cells are exceptional yet the remainder follows a default pattern: the M2 class noun *pere* in Table 2 may serve as an example. Reversals clearly involve the most unnatural classes, since in Table 1 the syncretic forms have neither case nor number values in common. While no morphological theory we are aware of is fully comfortable with reversals, it is clear that morpheme-based theories are probably the most hard-pressed (Kihm, 2017).

An important question in the study of reversals is to establish to what degree the reversal pattern has actually been generalised, i.e. whether or not it is truly symmetric (Baerman, 2007). For Old French, Kihm (2017) has argued that the system was actually quite unstable and disappeared after only a couple of centuries. This contrasts with more long-lived and more systematic reversals, as found e.g. in Neo-Aramaic (Baerman, 2007; Doron & Khan, 2012).

Another striking property of Old French is that nominal inflection only involves a single exponent *-s* to express distinctions of case and/or number.

In this paper, we shall investigate the exact nature of reversal in Old French and conclude that reversal has not been fully generalised, but is only one of several syncretism patterns. We shall see, however, that the distribution of *-s* in nominal paradigms follows some very regular patterns and show that the single cell that is characterised by massive idiosyncrasy is the nominative singular. We therefore argue that reversal in Old French is only apparent and propose a formal theory within Information-based Morphology (Crysmann & Bonami, 2016; Crysmann, 2017) that concisely captures the full range of syncretisms where reversals emerge by way of a combination of regular and idiosyncratic constraints on the distribution of *-s*.

1 Old French declension

Noun declension in Old French¹ exhibits three paradigms for masculine (given in Table 2) and equally three paradigms for feminine nouns (given in Table 3). The numbering of paradigms reflects overall productivity, i.e. the reversal pattern in M1 holds for the great majority of masculine nouns in the Old French lexicon. The majority pattern for feminines F1, by contrast, does not show any reversal, but instead displays a motivated split between unmarked singular and *s*-marked plural. As for adjectival declension (cf. Table 4), which is generally heteroclite, the by-far most common pattern A1 combines the most productive patterns for masculine and feminine noun declension (M1 and F1).

Compared to M1, the other two masculine inflection classes M2 and MAS only differ with respect to a single cell: in M2, nominative singular is unmarked, and in MAS, this very same cell is subject to both stem suppletion and optionality of *s*-marking, leading to overabundance (see Thornton, 2011, for an overview). One way to conceptualise this paradigm is in terms of underspecification of inflection class membership, i.e. MAS nouns can inflect according to M1 or M2. Likewise, feminine inflection classes F2 and FAS only minimally contrast with F1, and again they do so

¹We follow the nomenclature and empirical description given in Kihm 2017.

	M1		M2		MAS	
	SG	PL	SG	PL	SG	PL
NOM	chevaliers	chevalier	pere	pere	ber(s)	baron
OBJ	chevalier	chevaliers	pere	peres	baron	barons

Table 2: Old French masculine declensions (Kihm, 2017, p. 46–47)

	F1		F2		FAS	
	SG	PL	SG	PL	SG	PL
NOM	porte	portes	flors	flors	none	nonains
OBJ	porte	portes	flor	flors	nonain	nonains

Table 3: Old French feminine declensions (Kihm, 2017, p. 48–49)

in the same cell as masculines: F2 nouns are s-marked in the nominative singular (like M1 and unlike F1), and FAS nouns undergo stem alternation, but otherwise inflect like F1.

In terms of frequency, we should point out that M1 and F1 include the overwhelming majority of Old French nouns. M2 nouns are few and often aligned on M1 by supplying NOM.SG with *-s*, while F2 — often ‘regularised’ as F1 by not supplying NOM.SG with *-s* — gets some bulk from the fact that abstract nouns in *-té* (e.g. *beauté* ‘beauty’) fall into this class. Although not exactly insignificant in number, MAS and FAS nouns (especially the latter) constitute a small subset, progressively reduced by extending one stem to the whole paradigm, usually the OBJ.SG one. There are several types of MAS noun (see e.g. *emperere(s)/emperëor* ‘emperor’), but we cannot enter into that much detail here.

Turning to adjectives, all paradigms are heteroclite, i.e. they are mere combinations of the patterns we already observed for masculine and feminine nouns. While A1 is the combination of M1 and F1 where the feminine stems are affixed with *-e*, A2 does not show any independent gender marking. As for syncretism, A2 inflects just like M1 in the masculine, but it is overabundant in the feminine, patterning with both F1 and F2: again, the nominative singular is special, in that it is the locus of overabundance. AAS, which mainly contains comparatives, finally exhibits stem alternation, targeting again the nominative singular. Inflectional marking in the feminine follows the F1 pattern, like A1 adjectives do, but in the masculine we find again overabundance.

Although the vast majority of masculine nouns and adjectives indeed inflect according to the reversal pattern in Table 1, a look at the full range of paradigms reveals that reversal has not been fully generalised: As witnessed by the paradigms in Table 2 and 3, only two out of the six paradigms display a reversal pattern (M1 and MAS). Among the three paradigms where the nominative singular may bear the same formal

marking as the objective plural, the identical marking of the two cells is either not obligatory, as in the case of MAS (*ber(s)*), or identical marking is part of a larger syncretism pattern, as witnessed by the L-shaped pattern for F2 (*flors*), which singles out the objective singular (unmarked) vs. all other forms (marked with *-s*). In terms of syncretism patterns in the distribution of *-s*, we find four different patterns in total: reversal (M1), marked objective plural vs. unmarked default (M2), unmarked singular vs. marked plural (F1/FAS) and unmarked objective singular vs. all other cells marked by *-s* (F2). In terms of syncretism of the marker *-s*, MAS is overabundant in the nominative singular cell and can be considered as a mix of the syncretism patterns found with M1 and M2.

Looking at the entire set of Old French paradigms, we can establish, however, some straightforward generalisations that are independent of inflection class or syncretism pattern: first, objective singular is always unmarked, objective plural is always overtly marked with *-s*, and so is feminine plural. Second, nominative singular constitutes the one cell that is the domain of class-specific variation and even item-specific idiosyncrasy: while the realisation of nominative singular is clearly class-specific, distinguishing M1/F2 (marked by *-s*) from M2/F1/FAS (unmarked), the same cell is singled out as the locus of stem allomorphy, either idiosyncratic *ber/baron* or subregular *-e/-ain*. Finally, across all paradigms, this cell is the only one where overabundance can be observed, both for masculine nouns and feminine adjectives.

(a) A1 <i>buen(e)</i> ‘good’				
	MASC		FEM	
	SG	PL	SG	PL
NOM	buens	buen	buene	buenes
OBJ	buen	buens	buene	buenes

(b) A2 <i>grant</i> ‘big’				
	MASC		FEM	
	SG	PL	SG	PL
NOM	grant	grants	grant	grants
OBJ	grants	grant	grant(s)	grants

(c) AAS: <i>mieudre/meillor</i> ‘better’				
	MASC		FEM	
	SG	PL	SG	PL
NOM	mieudre(s)	meillor	mieudre	meillors
OBJ	meillor	meillors	meillor	meillors

Table 4: Old French adjectival declensions (Moignet, 1973, p. 26–31)

$$word \rightarrow \left[\begin{array}{l} \text{MPH} \quad \boxed{e_1} \circ \dots \circ \boxed{e_n} \\ \text{MS} \quad \boxed{0} \ (\boxed{m_1} \uplus \dots \uplus \boxed{m_n}) \\ \text{RR} \quad \left\langle \begin{bmatrix} \text{MPH} & \boxed{e_1} \\ \text{MUD} & \boxed{m_1} \\ \text{MS} & \boxed{0} \end{bmatrix}, \dots, \begin{bmatrix} \text{MPH} & \boxed{e_n} \\ \text{MUD} & \boxed{m_n} \\ \text{MS} & \boxed{0} \end{bmatrix} \right\rangle \end{array} \right]$$

Figure 1: Morphological wellformedness

2 Analysis

The analysis of the Old French data we are going to propose is formalised in terms of Information-based Morphology (=IbM; Crysmann & Bonami, 2016; Crysmann, 2017; Broadwell, 2017; Diaz et al., 2017), an inferential-realisation theory of inflection couched entirely in terms of inheritance hierarchies of typed feature structures. IbM differs from other inferential-realisation theories by adopting a morphous approach (Crysmann, 2003), which permits the treatment of the $m : n$ nature of the relationship between form and function at the most basic level, i.e. the individual rules. Furthermore, IbM systematically exploits inheritance, as well as cross-classification in the sense of Koenig (1999), to systematically establish vertical and horizontal generalisations over rules of exponence.

Realisation rules are pairings of a set of morphosyntactic properties to be expressed (MUD; = Morphology Under Discussion) with a list of exponents (MPH), possibly the empty list (cf. *zero-rr* in Figure 7). Members of MPH consist of a phonological description, paired with position class information. Since morphotactic information is now a first class citizen of rule descriptions, standard underspecification techniques of constraint-based grammar can be easily employed to extract generalisations about shape and position independently of each other (Crysmann & Bonami, 2016; Broadwell, 2017). The third top-level feature of every rule (MS) represents the entire morphosyntactic property set of the word and thus provides an easy way to address allomorphic conditioning (Crysmann, 2017; Diaz et al., 2017).

As depicted in Figure 1, a simple principle of completeness and coherence relates the MUD values of the rules to the morphosyntactic property set of the word. In essence, it requires that every member of the word’s MS set be licensed by some realisation rule. The word’s phonology is simply the concatenation of that of the morphs contributed by the rules, in the order of their positional indices, see Bonami & Crysmann (2013) for details. Since the relation between a word’s properties (MS/PH) to the realisation rules is entirely regulated by principle, grammatical specification of an individual inflectional system amounts to defining a signature of the properties themselves (features and appropriate values) and a hierarchy of realisation rules that pair them with the exponents that express these features.

2.1 Inflection classes in Old French

A recurrent observation about inflectional systems is that the exact choice of exponents is determined to one part by the properties being expressed, yet to another by lexically determined class membership. E.g., in Old French we do not just need to know that nominative singular can be expressed by *-s*, but we also need to know which classes of lexemes this rule applies to. Thus, before laying out the inflectional rules proper, we shall sketch how the nominal lexicon of Old French is partitioned into inflection classes, i.e. its morphomic properties (Aronoff, 1994).

In the previous section, we observed two fundamental levels of variation between paradigms: first, we found that nouns and adjectives contrast in using a single stem for all four (eight) cells of the paradigm, or else to use an alternate stem in the nominative singular (MAS, FAS, AAS). Second, both masculine and feminine nouns need to be distinguished as to their inflectional behaviour in the nominative singular, one class each that obligatorily takes the marker *-s* (M1, F2) and another that systematically refuses to do so (M2, F1). Regular adjectives (A1) are special in that they are heteroclite, following the productive pattern for masculine nouns (M1) in one part of the paradigm, yet that of feminine nouns (F1) in the other. What is more, some lexical classes (MAS, AAS, A2) display overabundance, being underspecified for inflection class in either the masculine (MAS,AAS) or the feminine (A2).

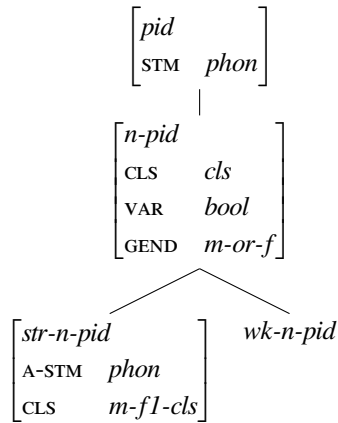


Figure 2: Signature of *pid* values

In IbM, lexically determined information, such as stem shapes or inflection class membership are interfaced with the inflection rule system via a distinguished feature structure (*pid*). We shall propose to represent the first property, i.e. availability of alternate stems, by a type hierarchy on *pid* values (cf. Figure 2), distinguishing *str(ong)-n-pid*, which has an appropriate feature for an alternate stem A-STM from the standard *w(ea)k-n-pid* which only has the STM appropriate of all *pid* values.

The second inflection class property pertains to the selection of paradigms proper: we introduce a feature CLS appropriate of *n-pid* that permits, inter alia, a systematic description of heteroclite and overabundant patterns, as given in Figure 3. At the bot-

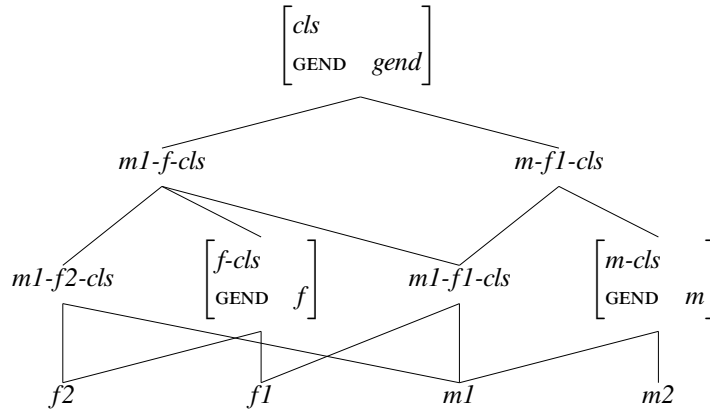


Figure 3: Nominal inflection classes

tom of the hierarchy, we find the four basic paradigm patterns *m1*, *m2*, *f1*, *f2*. The next level up represents three different abstractions: first, two gender types (*m-cl*, *f-cl*) with their appropriate *GEND* specifications, second, the representation of heteroclite regular adjectives (*m1-f1-cl*), and third, a type that singles out the paradigms taking *-s* as the exponent of nominative singular (*m1-f2-cl*). Even further up the hierarchy are the types for overabundance, which are underspecified w.r.t. paradigm membership either in the masculine (*m-f1-cl*), for MAS and AAS, or in the feminine *m1-f-cl*, for A2. Note that there is no abstraction of M2 independent of *m-cl*: this captures the fact that M2 does not serve as a model on its own for adjectival inflection. Furthermore, the inflectional patterning in the nominative singular of M2 corresponds to the unmarked case, such that independent targeting of e.g. F1 and M2 as a class is neither required nor desirable, but left to the elsewhere case.

Another piece of information that may be lexically specified is inherent gender for nouns: since gender is intimately tied to inflection class, we make it a feature appropriate of *cls*: the value of *GEND* will actually be narrowed down by the inflection class subtypes *m-cl* and *f-cl*, as depicted in Figure 3.

The last inflection class feature that we introduce via *pid* is *VAR*, a Boolean valued feature that controls whether or not adjectives have variable bases for masculine and feminine declension.

One generalisation about Old French is already captured at the level of the hierarchy of *pid* types: as depicted in Figure 2, stem alternation is correlated with a reduced set of class options *m-f1-cl*, capturing the fact that F2 stems do not undergo alternation.

The availability of inflectional patterns for any individual lexical item or word class is of course best captured by means of a hierarchy of lexical types. Owing to space considerations, we shall not give a full type hierarchy, but rather provide sample lexical specifications for the relevant nominal and adjectival classes (in Figures 4–6). Using Online Type Construction (Koenig & Jurafsky, 1994; Koenig, 1999), which is already assumed by IbM, extensional statements for subregular and irregular classes

can be cleanly separated from the underspecified description of regular and productive ones.

$$\begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{STM} \quad \text{chevalier} \\ \text{CLS} \quad m1\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(a) M1}
 \end{array}
 \quad
 \begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{STM} \quad \text{pere} \\ \text{CLS} \quad m2\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(b) M2}
 \end{array}$$

$$\begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{str-}n\text{-pid} \\ \text{STM} \quad \text{baron} \\ \text{A-STM} \quad \text{ber} \\ \text{CLS} \quad m\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(c) MAS}
 \end{array}$$

Figure 4: Sample entries of masculine nouns

$$\begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{STM} \quad \text{porte} \\ \text{CLS} \quad f1\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(a) F1}
 \end{array}
 \quad
 \begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{STM} \quad \text{flor} \\ \text{CLS} \quad f2\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(b) F2}
 \end{array}$$

$$\begin{array}{c}
 \left[\text{MS} \left\{ \left[\begin{array}{l} \text{str-}n\text{-pid} \\ \text{STM} \quad \text{nonain} \\ \text{A-STM} \quad \text{none} \\ \text{CLS} \quad f1\text{-cls} \end{array} \right] \right\} \right] \\
 \text{(c) FAS}
 \end{array}$$

Figure 5: Sample entries of feminine nouns

There are two aspects regarding the lexical representation of adjectives that deserve further elaboration, when compared to that of nouns: first, adjectives draw on the paradigms provided already for nouns, giving rise to heteroclisis between M1 and F1 (A1) and overabundance (A2: M1+F1+F2; AAS: M1+M2+F1). While the reliance on nominal patterns can be represented by drawing on the same hierarchy of inflection classes, we need to distinguish that gender is an inherent property for nouns, yet a morphosyntactic property for adjectives. As a consequence, we shall constrain adjectives to expose the value of the *GEND* feature contributed by the morphomic class as an inflectional property of its own, as shown in the sample entries in Figure 6. Second, regular productive adjectives (A1) undergo systematic gender inflection, using the productive M1 pattern in the masculine, whilst assimilating their feminine forms to the productive F2 pattern by affixation of *e* (/ə/). The other two

$$\begin{array}{c}
\left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{CLS} \left[\begin{array}{l} m1\text{-}f1\text{-cls} \\ \text{GEND } [g] \end{array} \right] \\ \text{STM} \text{ buen} \\ \text{VAR} + \end{array} \right] , [g] \right\} \right] \\
\text{(a) A1}
\end{array}
\quad
\begin{array}{c}
\left[\text{MS} \left\{ \left[\begin{array}{l} \text{wk-}n\text{-pid} \\ \text{STM} \text{ grant} \\ \text{CLS} \left[\begin{array}{l} m1\text{-}f\text{-cls} \\ \text{GEND } [g] \end{array} \right] \\ \text{VAR} - \end{array} \right] , [g] \right\} \right] \\
\text{(b) A2}
\end{array}$$

$$\begin{array}{c}
\left[\text{MS} \left\{ \left[\begin{array}{l} \text{str-}n\text{-pid} \\ \text{CLS} \left[\text{GEND } [g] \right] \\ \text{STM} \text{ meilleur} \\ \text{A-STM} \text{ mieudre} \\ \text{VAR} - \end{array} \right] , [g] \right\} \right] \\
\text{(c) AAS}
\end{array}$$

Figure 6: Sample entries of adjectives

patterns (A2 and AAS), however, do not show any direct gender marking. In order to distinguish the invariant patterns of A2 and AAS from the gender-inflected pattern exhibited by A1, we use a Boolean valued feature *VAR*.

2.2 Realisation rules

Now that we have provided a suitable representation of the more idiosyncratic morphomic information such as stem alternations and inflection class membership, we can move on to the core of the analysis, as given by the hierarchy of realisation rules in Figure 7.

As will become apparent shortly, our treatment of apparent reversal in Old French will essentially expose four empirical generalisations: first, the status of *-s* as the only non-stem exponent of case/number marking, and second, the fact that the distribution of this marker is highly regular, and third, that a single cell is the locus of all exceptions. Fourth, objective singular, which never undergoes any overt marking, should be regarded as an instance of the unmarked case.

The type hierarchy in Figure 7 depicts four classes of realisation rules, if understood in terms of *MUD* values: one class for stem realisation (*stem-rr*), one class for *s*-marking (*s-rr*), a third monadic class for feminine gender realisation (*f-rr*), appropriately restricted to a subclass of adjectives, and finally, default zero realisation (*zero-rr*).

The rule type *s-rr* mainly describes the shape and position of the morph *s*, while restricting its function to express some case/number combination. Subtypes of *s-rr* further constrain the *MUD* value. The right-hand subtype captures the fact that the marker may express plural, and its two subtypes further narrow down the conditions: the suffix *-s* can either mark plural in the objective case (true of all paradigms), or else

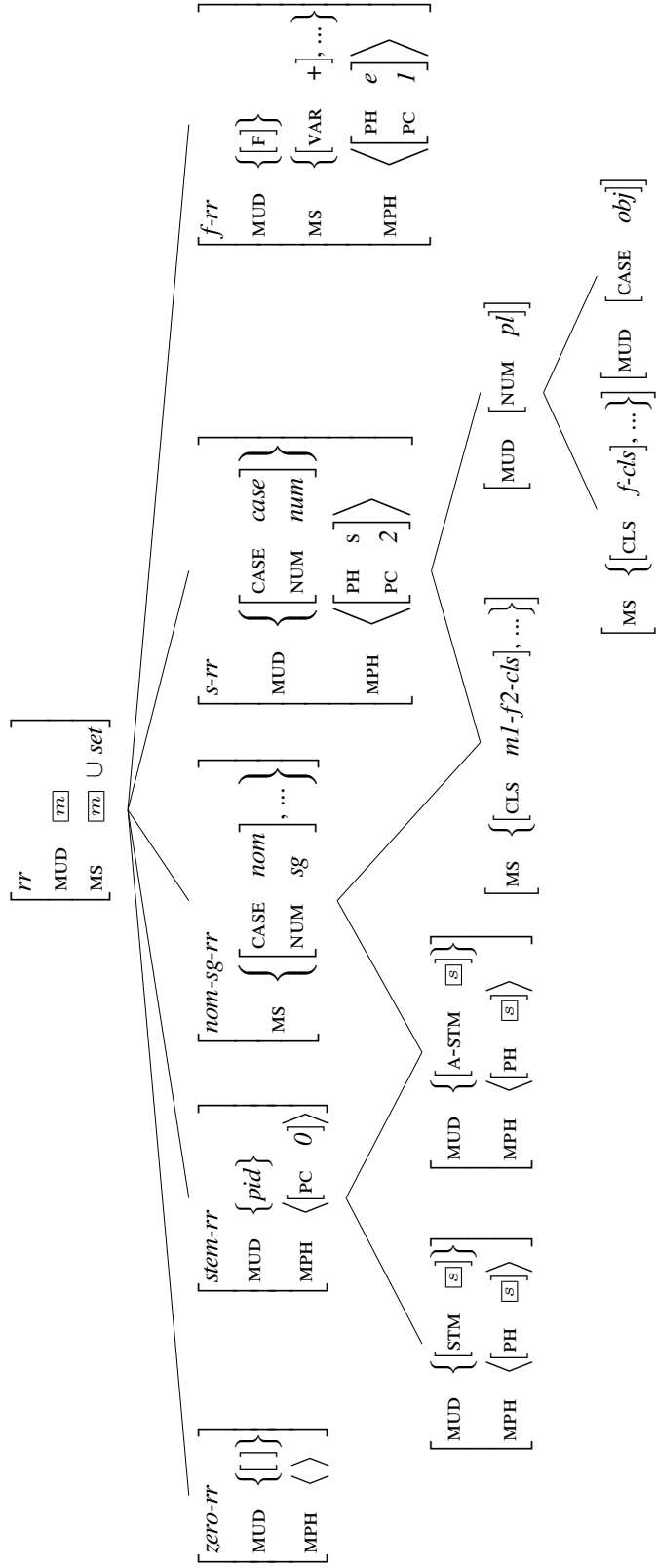


Figure 7: Realisation rules for Old French nominal inflection

$$\left[\begin{array}{l} \text{word} \\ \text{MS } \boxed{0} \left\{ \boxed{a} \left[\begin{array}{ll} \text{str-}n\text{-pid} \\ \text{STM} & \text{baron} \\ \text{A-STM} & \text{ber} \\ \text{CLS} & m\text{-cls} \end{array} \right], \boxed{b} \left[\begin{array}{ll} \text{CASE} & \text{nom} \\ \text{NUM} & \text{sg} \end{array} \right] \right\} \\ \text{RR } \left\{ \begin{array}{l} \text{stem-rr} \wedge \text{nom-sg-rr} \\ \text{MUD } \left\{ \boxed{a} \left[\begin{array}{ll} n\text{-str-pid} \\ \text{A-STM} & \boxed{s} \end{array} \right] \right\} \\ \text{MS } \boxed{0} \left\{ \left[\begin{array}{ll} \text{CASE} & \text{nom} \\ \text{NUM} & \text{sg} \end{array} \right], \dots \right\} \\ \text{MPH } \boxed{x} \left\langle \left[\begin{array}{ll} \text{PH} & \boxed{s} \\ \text{PC} & 0 \end{array} \right] \right\rangle \end{array} \right\}, \left\{ \begin{array}{l} s\text{-rr} \wedge \text{nom-sg-rr} \\ \text{MUD } \left\{ \boxed{b} \left[\begin{array}{ll} \text{CASE} & \text{nom} \\ \text{NUM} & \text{sg} \end{array} \right] \right\} \\ \text{MS } \boxed{0} \left\{ \left[\begin{array}{ll} \text{CLS} & m1\text{-f2-cl}\text{s} \end{array} \right], \dots \right\} \\ \text{MPH } \boxed{y} \left\langle \left[\begin{array}{ll} \text{PH} & s \\ \text{PC} & 2 \end{array} \right] \right\rangle \end{array} \right\} \\ \text{MPH } \boxed{x} \left\langle \left[\begin{array}{ll} \text{PH} & \boxed{s} \\ \text{PC} & 0 \end{array} \right] \right\rangle \circ \boxed{y} \left\langle \left[\begin{array}{ll} \text{PH} & s \\ \text{PC} & 2 \end{array} \right] \right\rangle \end{array} \right]$$

Figure 8: Derivation of *bers* ‘baron(M).NOM.SG’

it can mark the plural with feminine nouns or adjectives. While these two options are fully regular, *s-rr* caters for another subtype, constrained to class *m1-f2-cl*s, in order to accommodate lexically restricted nominative singular marking, by way of inheritance from *nom-sg-rr*.

Turning to stem selection, we find a similar pattern: *stem-rr* has a general subtype which selects the STM feature as an exponent of lexical identity, yet it also provides an alternate stem rule for the A-STM. The use conditions for this alternate stems are again the nominative singular, just as with the exceptional s-marking. The identity of condition is captured by inheritance from the common supertype *nom-sg-rr*.

Realisation of objective singular, or for that matter any unmarked cell, enjoy the status of a true default: since no rule description exists that is more specific, Paninian competition will license zero realisation (*zero-rr*).

By way of illustration, we shall provide sample derivations of the two possible realisations of the nominative singular of class MAS noun *ber(s)* ‘baron(M).NOM.SG’.

Figure 8 illustrates derivation of the s-marked variant *ber-s*. At the top of the word-level feature structure, we find the representation of the morphosyntactic property set MS, including lexemic information, the RR set of realisation rules, and finally, the word-level list of morphs on MPH. In correspondence with the principle of morphological wellformedness in Figure 1, the MS set of the *word* is exhausted by the MUD values of the rules in the RR set, as indicated by the co-reference tags $\boxed{a}$ and $\boxed{b}$. Likewise, the morphs contributed by the rules ($\boxed{x}$ and $\boxed{y}$) are shuffled together on the word’s MPH list, and finally, the entire MS set of the word ($\boxed{0}$) is distributed over the MS

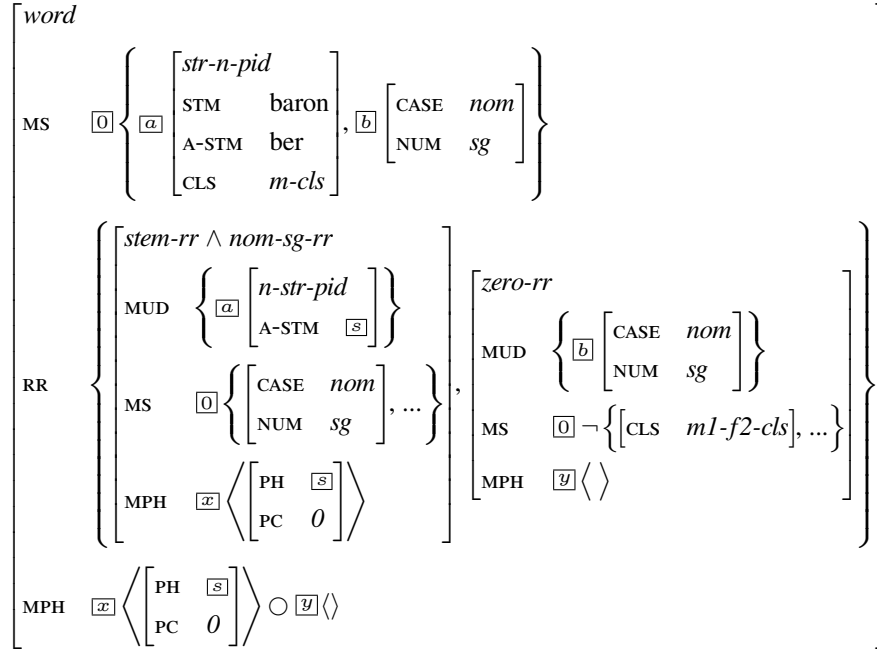


Figure 9: Derivation of *ber* ‘baron(M).NOM.SG’

features of the rules, making it possible for rules to impose allomorphic constraints.

Concretely, lexemic properties ($\boxed{a}$) are expressed by a stem selection rule, and more precisely by one that selects the alternate A-STM, the phonology of which is inserted in a morph ($\boxed{x}$). The morphosyntactic properties of case and number ($\boxed{b}$) are expressed by the morph $\boxed{y}$, with shape *s*. The realisation rule for *s*-marked nominative singular is restricted to *m1-f2-cl*s, unifying with lexemic *m-cl*s to *m1-cl*s. Conversely, the alternate stem selection rule is constrained to apply to nominative singular. Selection of the regular stem, however, is preempted by Paninian competition (see Crysmann, 2017).

Figure 9, moreover, illustrates derivation of the zero-marked variant *ber*. The main difference is with respect to expression of case/number inflection: here this morphosyntactic property is expressed by *zero-rr*, a rule that pairs the morphosyntactic property $\boxed{b}$ (an element of MUD, and hence MS) with the empty list of exponents ($\boxed{y}$). This rule is in Paninian competition with the *s*-marking nominative singular rule (by way of subsumption), so therefore its MS value is restricted to the complement of the more specific rule, yielding a negative existential on the class specification for *m1-f2-cl*s. In the case at hand, lexical underspecification (*m-cl*s) and the Paninian constraint ($\neg m1\text{-f2-cl}\text{s}$) will unify to *m2-cl*s.

Adjectival inflection in class A1 displays a systematic variation with respect to gender: as witnessed by the paradigm in Table 4, feminine forms are related to their masculine counterparts by suffixation of *-e* ([ə]), in addition to a shift of the inflectional pattern from M1 to F1. As detailed in Figure 7, feminine marking by *-e* (rule

f-rr) is restricted to [VAR +], effectively applying to A1 adjectives (cf. Figure 6a). Gender for A2 and AAS, by contrast, will be realised by rule *zero-rr*, just like masculine gender for all adjectives.

To summarise the analysis, apparent reversal in Old French emerges as the result of the combination of regular inflectional patterns that are true across all paradigms with class-specific realisation rules for the nominative singular. The formalisation in terms of inheritance hierarchies of realisation rules successfully captures what we take to be the two fundamental observations, namely that there is only a single affixal exponent for case/number distinctions in the entire declension system, and that the “problematic” cell is always the same, for affixation and stem selection alike. Finally, the observation that overabundance targets the same cell just falls out from the fact that this is the only cell where realisation depends on inflection class membership, such that lexical underspecification of class membership will suffice to ensure that MAS nouns, as well as A2 and AAS adjectives can undergo either default zero marking, or class-specific overt marking with *-s*. The treatment of overabundance in terms of lexical underspecification is furthermore fully in line with recent work on overabundance within IbM (Bonami & Crysmann, 2018).

3 Conclusion

In this paper we have looked at apparent reversals in Old French declension and shown that the reversal pattern, though frequent in the masculine, is only apparent. We have argued that with the exception of the nominative singular, Old French declension is highly regular across all paradigms and that the nominative singular cell is problematic in three respects: it is the locus of stem suppletion, class-specific marking with *-s*, and as a result to the availability of alternate inflection patterns, the locus of overabundance.

Furthermore, we have developed a formal analysis of the Old French system within the framework of Information-based Morphology (Crysmann & Bonami, 2016; Crysmann, 2017) that captures several salient facts about Old French concisely by means of underspecification in inheritance hierarchies of realisation rules: within the inflectional system proper, generalisations about exponence are factored out into a supertype, and so are the constraints on exceptional inflection and stem suppletion. On the lexical side, underspecification of inflection class effortlessly derives overabundance.

Finally, on the diachronic side, our analysis helps one understand what made Old French declension an unstable system and why it was as short-lived as it was. As we have shown, fragility was located in the *NOM.SG* cell of M1 nouns. As it became increasingly unmarked, not only did the case contrast in the singular collapse, but, more seriously, the number contrast in the nominative threatened to do so as well. The remedy consisted in doing away with case inflection entirely, keeping only the number contrast of the two formerly objective cells. Since F1 nouns never marked case to begin with, and given the various ‘regularisations’, the whole declension system

simply vanished.

References

- Aronoff, Mark. 1994. *Morphology by itself*. Cambridge: MIT Press.
- Baerman, Matthew. 2007. Morphological reversals. *Journal of Linguistics* 43. 33–61.
- Bonami, Olivier & Berthold Crysmann. 2013. Morphotactics in an information-based model of realisational morphology. In Stefan Müller (ed.), *Proceedings of the 20th International Conference on Head-Driven Phrase Structure Grammar, Freie Universität Berlin*, 27–47.
- Bonami, Olivier & Berthold Crysmann. 2018. Lexeme and flexeme in a formal theory of grammar. In Olivier Bonami, Gilles Boyé, Georgette Dal & Hélène Giraudo et Fiammetta Namer (eds.), *The lexeme in descriptive and theoretical morphology* (Empirically Oriented Theoretical Morphology and Syntax 4), 175–202. Berlin: Language Science Press.
- Broadwell, George Aaron. 2017. Parallel affix blocks in Choctaw. In Stefan Müller (ed.), *Proceedings of the 24th International Conference on Head-Driven Phrase Structure Grammar, University of Kentucky, Lexington*, 103–119. Stanford, CA: CSLI Publications. <http://cslipublications.stanford.edu/HPSG/2017/hpsg2017-broadwell.pdf>.
- Crysmann, Berthold. 2003. *Constraint-based coanalysis. Portuguese cliticisation and morphology–syntax interaction in HPSG* (Saarbrücken Dissertations in Computational Linguistics and Language Technology 15). Saarbrücken: Computational Linguistics, Saarland University and DFKI LT Lab.
- Crysmann, Berthold. 2017. Inferential-realisation morphology without rule blocks: an information-based approach. In Nikolas Gisborne & Andrew Hippisley (eds.), *Defaults in morphological theory*, 182–213. Oxford: Oxford University Press.
- Crysmann, Berthold & Olivier Bonami. 2016. Variable morphotactics in information-based morphology. *Journal of Linguistics* 52(2). 311–374.
- Diaz, Thomas, Jean-Pierre Koenig & Karin Michelson. 2017. Oneida prepronominal prefixes in Information-Based Morphology. Paper presented at the 24th International Conference on Head-Driven Phrase Structure Grammar, 7–9 July, Lexington, Kentucky.
- Doron, Edit & Geoffrey Khan. 2012. The typology of morphological ergativity in Neo-Aramaic. *Lingua* 122. 225–240.
- Kihm, Alain. 2017. Old French declension. In Nikolas Gisborne & Andrew Hippisley (eds.), *Defaults in morphological theory*, 40–72. Oxford: Oxford University Press.

- Koenig, Jean-Pierre. 1999. *Lexical relations*. Stanford: CSLI Publications.
- Koenig, Jean-Pierre & Daniel Jurafsky. 1994. Type underspecification and on-line type construction. In *Proceedings of WCCFL XIII*, 270–285.
- Moignet, Gérard. 1973. *Grammaire de l'ancien français*. Librairie Klincksieck.
- Thornton, Anna M. 2011. Overabundance (multiple forms realizing the same cell): A non-canonical phenomenon in Italian verb morphology. In Martin Maiden, John Charles Smith, Maria Goldbach & Marc-Olivier Hinzelin (eds.), *Morphological autonomy: Perspectives from romance inflectional morphology*, 358–381. Oxford: Oxford University Press.

Valence-changing morphology via lexical rule composition

Christian M. Curtis 

University of Washington

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 36–48

Keywords: HPSG, Valence Change, Grammar Matrix

Curtis, Christian M. 2018. Valence-changing morphology via lexical rule composition. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 36–48. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.3.  CC-BY 4.0

Abstract

I describe an analysis of valence-changing verbal morphology implemented as a library extending the LinGO Grammar Matrix customization system. This analysis is based on decomposition of these operations into rule components, which in turn are expressed as lexical rule supertypes that implement specific, isolatable constraints. I also show how common variations of these constraints can be abstracted and parameterized by their axes of variation. I then demonstrate how these constraints can be recomposed in various combinations to provide broad coverage of the typological variation of valence change found in the world's languages. I evaluate the coverage of this library on five held-out world languages that exhibit these phenomena, achieving 79% coverage and 2% overgeneration.

Introduction

In this paper, I describe a cross-linguistic analysis of valence-changing morphology that is implemented in a meta-grammar engineering system, the LinGO Grammar Matrix (Bender, Flickinger & Oepen, 2002)¹. The core of the Grammar Matrix is a collection of implemented analyses for cross-linguistic phenomena, developed by linguists and grammar writers over many years, in a framework that provides infrastructure and context for reuse in development of precision grammars. The Grammar Matrix customization system (Bender, Drellishak, Fokkens, Poulson & Saleem, 2010) combines a structured means of eliciting typological characteristics, validating responses for consistency, and using those choices to combine Matrix core grammar elements with stored analyses of various linguistic phenomena into a customized grammar. The Grammar Matrix uses Minimal Recursion Semantics (MRS; Copestake, Flickinger, Pollard & Sag 2005) as its semantic representation. MRS can be naturally expressed in terms of feature structures and so is integrated into its HPSG mechanisms and feature structures.

One category of linguistic phenomenon not previously covered by the Grammar Matrix customization system is valence change: verbal morphology that alters the argument structure, either increasing or decreasing the valency, and changing the relationship of realized arguments to syntactic roles. To extend the customization system to include these operations, I developed a library that implements valence-changing operations that can be selected as part of a customized grammar. In building this library, I separated each high-level operation into foundational rule components, or “building blocks,” which can then be composed as needed to implement valence change for a wide variety of the world's languages.

I first provide a brief typological survey of valence change (Section 1), followed by some examples of my analyses of these phenomena (Section 2); in particular, I illustrate the separation of these operations into rule components, as well as their re-composition into complete lexical rule types. I then describe my implementation of the library and evaluate its coverage of valence change variation (Sections 3-4).

¹<http://matrix.ling.washington.edu>

1 Typology of valence change

I describe the cross-linguistic range of these operations below: following the broad conceptual framework provided by Haspelmath & Müller-Bardey (2004) (henceforth H&MB), operations are grouped by whether they reduce or increase valency and whether they affect the subject or object. I also retain their focus on verbal valence-changing morphology, excluding e.g., periphrastic constructions.

1.1 Valence-reducing operations

Both anticausative and passive constructions remove the subject and move the former object into the subject position; the essential distinction between them is that the anticausative removes the subject argument entirely, while the passive merely moves it to the periphery (H&MB). Analogous to the anticausative, the object-removing operation where the object is completely removed is referred to as the deobjective Haspelmath & Müller-Bardey (2004) or the absolute antipassive (Dayley, 1989, as cited in H&MB). The deaccusative (H&MB) or antipassive (Dixon & Aikhenvald, 2000) is similar, but, instead of completely removing the underlying object argument, moves it out of the core to the periphery. The Turkish [tur] anticausative is illustrated in (1) as an example of a typical valence-reducing operation.

- (1) a. Anne-m kapı-yı aç-tı
mother-1SG door-ACC open-PAST(3SG)
‘My mother opened the door.’ [tur]
- b. Kapı aç-tı-dı
door open-ANTIC-PAST(3SG)
‘The door opened.’ [tur] (H&MB, p. 5)

1.2 Valence-increasing operations

Cross-linguistically the most common valence-changing category (Bybee, 1985), the causative adds a new subject, the causer of the event described by the verb. The addition of a causer to an intransitive verb can simply move the underlying subject into an object position. The situation with underlying transitive verbs is more complex, as there are different strategies for dealing with the underlying subject (causee), given the presence of an already-existing direct object. Other subject-adding constructions are structurally similar to the causative, such as the affective (‘indirect passive’) in Japanese [jpn]. A crucial aspect of the causative and similar constructions is the addition of a new elementary predication (EP) which functions as a scopal operator with respect to the verb’s own EP and also takes as an argument the added participant.

Object-adding constructions can collectively be grouped under the term ‘applicative,’ which subsumes a broad variation in potential roles for the added structural argument. The prototypical applicative is the benefactive; however, in many languages (such as the Bantu family), applicatives can serve many additional functions, including locative, possessor-raising, and instrumental variations. Unlike the causative, the applicative is non-scopal. An example of object adding is presented below in (2).

As this brief survey highlights, there is a broad variety of specific valence-changing lexical operations in the world’s languages, but strong threads of similarity also run through them. To cover this variety, I followed a “building block” approach in my analysis and implementation, as described in the following sections.

2 Analysis

The overall approach I followed was to decompose the high-level, linguistically-significant valence-changing operations into their component operations on feature structures. These individual component operations can then be selected by the customization system and composed to achieve the intended high-level result. The components I selected to analyze and implement included addition and removal of subjects and objects, case constraints and alternations, and argument reordering. In this section, as an illustrative example, I discuss my analysis of object addition, its breakdown into rule components, and the resulting effects on the feature structures.

Object addition covers the general category of the applicative, which subsumes a variety of different types of oblique roles cross-linguistically, including the instrumental and benefactive (H&MB, p. 7). In adding an argument, there are several underlying operations in my analysis: (a) adding an argument to the COMPS list;² (b) constraining the added argument, e.g. to be an np or pp (HEAD *noun* or *adp*), or applying a CASE constraint; (c) appending the new argument’s non-local dependencies to the rule mother’s list; (d) contributing an added EP (via C-CONT); (e) linking the new EP’s first argument to the daughter’s INDEX; and (f) linking the new EP’s second argument to the new argument’s INDEX.

The first two of these operations, (a) and (b), are directly grounded in the addition of a new argument and are straightforward; appending the new argument’s non-local dependencies simply maintains the threading analysis of Bouma, Malouf & Sag (2001) and is similarly straightforward to motivate.

The addition of a new EP to the rule output is not as straightforward, and requires some additional discussion. To motivate this analysis, consider the example of the benefactive from Indonesian in (2). In this example, the addition of the benefactive applicative suffix *-kan* in (2b) adds an argument position to the verb, which is filled by *ibunja* “his mother”.

²Note that, cross-linguistically, the added argument can be added either more- or less-obliquely to the verb’s existing dependencies (i.e., at the head or tail of the COMPS list)

- (2) a. Ali memi televisi untuk ibu-nja
 Ali TR.buy television for mother-his
 ‘Ali bought a television for his mother.’ [ind]
- b. Ali mem-beli-kan ibu-nja televisi
 Ali TR-buy-APPL mother-his television
 ‘Ali bought his mother a television.’ [ind]
- (Chung, 1976, in Wunderlich, 2015, p. 21)

Notionally, the benefactive is adding a third semantic argument to the verb, which would add a hypothetical third argument to the EP contributed by the verb;³ however, this would seem to violate the principles of semantic composition in Copestake et al. 2005 for Minimal Recursion Semantics (MRS), namely, that composition consists solely of concatenation of daughter RELS values, not modification. More concretely, there is no EP-modifying operation available within the algebra of Copestake, Lascarides & Flickinger (2001).⁴

The solution is to have the lexical rule contribute a new EP, which takes both the EP contributed by the verb and the additional syntactic argument as semantic arguments. The predicate value for this new EP will provide the particular species of applicative (e.g., benefactive, as here). This new EP contributes its own event, and takes as its arguments the respective indexes of the input and the added argument. In this analysis, I treat the relationship between the added EP and the verb as non-scopal, with no intervening handle relationships. This contrasts with my analysis of subject addition; in my survey of valence change, the scopal relationship appears to only arise with added subjects.

The introduction of a new event by this EP differs from the analysis of the benefactive presented by Müller (2018, p. 68); my analysis here makes the relation contributed by the EP potentially available for modification separately from the event of the main verb (as in the English [eng] periphrastic form *Kim read the book, probably for Sandy*).

The MRS resulting from this analysis is shown in (3):

³Or a lexical rule which has previously been applied to the input.

⁴Although this principle is generally strongly embraced by DELPH-IN grammars, it is not entirely settled whether this necessarily should be as strictly applied within lexical rules (see e.g. Copestake, Lascarides & Bender, 2016, and Bender, 2015); it also is not prohibited by the DELPH-IN joint reference formalism (Copestake, 2002). My analyses in this work are not frustrated by adhering to this principle, so I retain it as applying throughout a grammar.

$$(3) \left[\begin{array}{c} \text{RELS} \left\langle \begin{array}{c} \begin{array}{c} \text{memi_v_rel} \\ \text{ARGO } \boxed{4} \text{ event} \\ \text{ARG1 } \boxed{1} \\ \text{ARG2 } \boxed{2} \end{array}, \begin{array}{c} \text{Ali_n_rel} \\ \text{ARGO } \boxed{1} \end{array}, \begin{array}{c} \text{telefisi_n_rel} \\ \text{ARGO } \boxed{2} \end{array} \right. \\ \left. \begin{array}{c} \text{ibu_n_rel} \\ \text{ARGO } \boxed{3} \end{array}, \begin{array}{c} \text{benefactive_rel} \\ \text{ARGO event} \\ \text{ARG1 } \boxed{4} \\ \text{ARG2 } \boxed{3} \end{array} \right\rangle \end{array} \right]$$

With all these elements combined, a complete rule implementing the benefactive (with the argument being added less-obliquely, in this example) is illustrated in (4).

$$(4) \left[\begin{array}{c} \text{benefactive-lex-rule} \\ \\ \text{SYNSEM} \left[\begin{array}{c} \text{LOCAL} \mid \text{CAT} \mid \text{VAL} \mid \text{COMPS} \left\langle \boxed{1}, \right. \\ \left. \begin{array}{c} \text{LOCAL} \left[\begin{array}{c} \text{CAT} \left[\begin{array}{c} \text{HEAD } \textit{noun} \\ \text{VAL} \left[\begin{array}{c} \text{SPR } \langle \rangle \\ \text{COMPS } \langle \rangle \end{array} \right] \end{array} \right] \\ \text{CONT} \mid \text{HOOK} \mid \text{INDEX } \boxed{2} \end{array} \right] \\ \text{NON-LOCAL} \left[\begin{array}{c} \text{SLASH } \boxed{3} \\ \text{QUE } \boxed{4} \\ \text{REL } \boxed{5} \end{array} \right] \end{array} \right\rangle \end{array} \right] \\ \\ \text{NON-LOCAL} \left[\begin{array}{c} \text{SLASH } \boxed{7 \oplus 3} \\ \text{QUE } \boxed{8 \oplus 4} \\ \text{REL } \boxed{9 \oplus 5} \end{array} \right] \\ \\ \text{C-CONT} \left[\begin{array}{c} \text{RELS} \left\langle ! \begin{array}{c} \text{event-relation} \\ \text{PRED } \textit{"benefactive_rel"} \\ \text{ARG1 } \boxed{6} \\ \text{ARG2 } \boxed{2} \end{array} ! \right\rangle \end{array} \right] \\ \\ \text{DTR} \left[\begin{array}{c} \text{verb-lex} \\ \\ \text{SYNSEM} \left[\begin{array}{c} \text{LOCAL} \left[\begin{array}{c} \text{CAT} \mid \text{VAL} \mid \text{COMPS } \boxed{1} \\ \text{CONT} \mid \text{HOOK} \mid \text{INDEX } \boxed{6} \end{array} \right] \\ \text{NON-LOCAL} \left[\begin{array}{c} \text{SLASH } \boxed{7} \\ \text{QUE } \boxed{8} \\ \text{REL } \boxed{9} \end{array} \right] \end{array} \right] \end{array} \right] \end{array} \right]$$

This rule, however, in combining the distinct operations identified above, obscures common elements that can be reused for other similar object-adding operations. Reviewing the five operations, it is evident that they vary along several different axes, as shown in Table 1.

This leads to a simplification and optimization: these building-block operations can be treated as being parameterized along their axes of variation, and then

rule component	varies by
added argument	position (obliqueness), number of existing args
constraint on new argument	position (obliqueness), constraint (e.g. case, head)
non-local dependencies	position (obliqueness)
new EP's pred value	predicate
new EP's arg1	does not vary
new EP's arg2	position (obliqueness)

Table 1: Rule components

combined to make the final rule type. Concretely, taking these operations in turn, the first operation (adding the argument) needs to have variants for adding an argument: (a) to intransitive or transitive verbs, and (b) at the front or end of the COMPS list. That is, the lexical rule type implementing each of the component operations can be viewed as the output of a function: $f : tr \in \{intrans, trans\} \times pos \in \{front, end\} \rightarrow lrt$. To illustrate one variation, the rule type at (5) adds an argument at the *end* of the COMPS list for an *intransitive* verb,⁵ and the rule at the rule type shown in (6) adds an argument at the front of the COMPS list for a transitive verb and links the INDEX of that argument to its second semantic argument (ARG2).

$$(5) \left[\begin{array}{l} \text{added-arg2of2-lex-rule} \\ \\ \text{SYNSEM} | \text{LOCAL} | \text{CAT} | \text{VAL} | \text{COMPS} \quad \left\langle \left[\begin{array}{l} \text{LOCAL} \quad \left[\begin{array}{l} \text{CAT} | \text{VAL} \quad \left[\begin{array}{l} \text{SPR} \quad \langle \rangle \\ \text{COMPS} \quad \langle \rangle \end{array} \right] \right] \\ \text{CONT} | \text{HOOK} | \text{INDEX} \quad \boxed{1} \end{array} \right] \right] \right\rangle \\ \\ \text{C-CONT} | \text{RELS} \quad \langle ! \left[\text{ARG2} \quad \boxed{1} \right] ! \rangle \\ \\ \text{DTR} | \text{SYNSEM} | \text{LOCAL} | \text{CAT} | \text{VAL} | \text{COMPS} \quad \langle \rangle \end{array} \right]$$

$$(6) \left[\begin{array}{l} \text{added-arg2of3-lex-rule} \\ \\ \text{SYNSEM} | \text{LOCAL} | \text{CAT} | \text{VAL} | \text{COMPS} \quad \left\langle \left[\begin{array}{l} \text{LOCAL} \quad \left[\begin{array}{l} \text{CAT} | \text{VAL} \quad \left[\begin{array}{l} \text{SPR} \quad \langle \rangle \\ \text{COMPS} \quad \langle \rangle \end{array} \right] \right] \\ \text{CONT} | \text{HOOK} | \text{INDEX} \quad \boxed{1} \end{array} \right] \right] , \boxed{2} \right\rangle \\ \\ \text{C-CONT} | \text{RELS} \quad \langle ! \left[\text{ARG2} \quad \boxed{1} \right] ! \rangle \\ \\ \text{DTR} | \text{SYNSEM} | \text{LOCAL} | \text{CAT} | \text{VAL} | \text{COMPS} \quad \langle \boxed{2} \rangle \end{array} \right]$$

In a similar fashion, constraining the head type of the added argument can be isolated to an individual rule supertype, as in (7):

⁵Naturally, for an intransitive input there is no difference between the front and end of the comps list. The same rule would be generated for either specification; formally, $f(intrans, front) \equiv f(intrans, end)$.

$$(7) \left[\begin{array}{l} \text{added-arg2-np-head-lex-rule} \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{VAL} \mid \text{COMPS} \quad \langle \left[\text{LOCAL} \mid \text{CAT} \mid \text{HEAD} \quad \textit{noun} \right], [] \rangle \end{array} \right]$$

The non-local dependencies carried by the added argument, as analyzed in the threading analysis of Bouma et al. (2001), are implemented in the Grammar Matrix as difference-list append operations. This is normally handled in the Grammar Matrix by a lexical type’s mapping from argument structure to valence lists, with the additional difference-list append constraints provided via inheriting the appropriate lexical supertype (*basic-one-arg*, *basic-two-arg*, etc.). As this analysis operates outside these existing mechanisms, an additional constraint, parameterized on the position of the added argument, must be added to perform these appends.⁶ An example of this operation is shown in (8):

$$(8) \left[\begin{array}{l} \text{added-arg2of3-non-local-lex-rule} \\ \text{SYNSEM} \quad \left[\begin{array}{l} \text{LOCAL} \mid \text{CAT} \mid \text{VAL} \mid \text{COMPS} \quad \langle \left[\begin{array}{l} \text{NON-LOCAL} \quad \left[\begin{array}{l} \text{SLASH} \quad [4] \\ \text{REL} \quad [5] \\ \text{QUE} \quad [6] \end{array} \right] \end{array} \right] \rangle \\ \text{NON-LOCAL} \quad \left[\begin{array}{l} \text{SLASH} \quad [1 \oplus 4] \\ \text{REL} \quad [2 \oplus 5] \\ \text{QUE} \quad [3 \oplus 6] \end{array} \right] \end{array} \right] \\ \text{DTR} \mid \text{SYNSEM} \mid \text{NON-LOCAL} \quad \left[\begin{array}{l} \text{SLASH} \quad [1] \\ \text{REL} \quad [2] \\ \text{QUE} \quad [3] \end{array} \right] \end{array} \right]$$

The most variable, individualized component is the predicate (pred value) of the added semantic relation. For example, the benefactive and instrumental rules may be entirely common in structure, but would need to be distinguished by their predicate. A separate rule supertype, therefore, can be created as in (9) to constrain the PRED value appropriately.

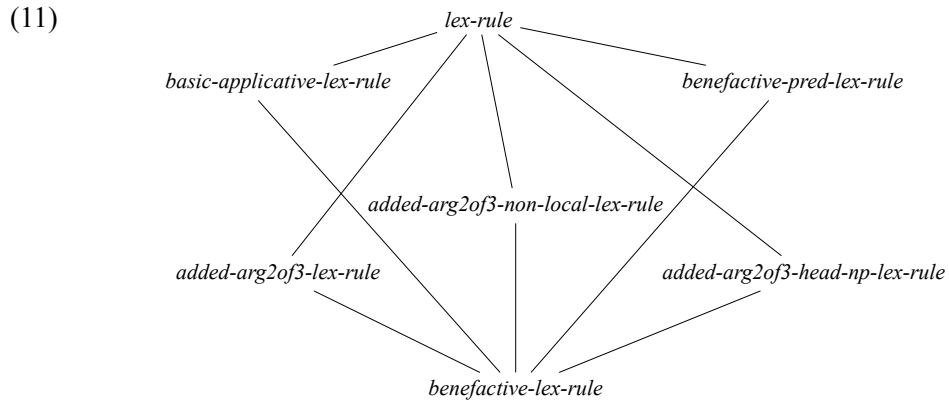
$$(9) \left[\begin{array}{l} \text{benefactive-pred-lex-rule} \\ \text{C-CONT} \mid \text{RELS} \quad \langle ! \left[\text{PRED} \quad \textit{“benefactive_rel”} \right] ! \rangle \end{array} \right]$$

Finally, the (invariant) core of the “generic” applicative can be isolated and analyzed as in (10). The function of this rule supertype is to contribute the new predication, and link its ARG1 to the daughter’s intrinsic variable (i.e., the ARGO).

⁶I have arbitrarily selected the ordering here of the added non-local dependencies; this analysis may need to be refined in the event that order of non-local dependency satisfaction becomes relevant.

$$(10) \left[\begin{array}{l} \textit{basic-applicative-lex-rule} \\ \text{C-CONT} \\ \text{DTR} \mid \text{SYNSEM} \mid \text{LOCAL} \mid \text{CONT} \mid \text{HOOK} \mid \text{INDEX} \quad \boxed{1} \end{array} \quad \left[\text{RELS} \quad \left\langle ! \left[\begin{array}{l} \textit{event-relation} \\ \text{ARG1} \quad \boxed{1} \end{array} \right] ! \right\rangle \right] \right]$$

These building-block rule component supertypes can then be assembled as inherited constraints on a complete applicative rule type, ready to be instantiated in a grammar. The partial inheritance tree showing these rule component supertypes for the original example full rule type in (4) is illustrated below in (11):



In the case of a subject-adding operation, such as the causative illustrated from Georgian [kat] in (12), I treat the added (“causing”) EP as a scopal predicate: it outscopes the underlying verb’s EP and so provides the `HOOK` feature values for the entire VP. The resulting MRS should be as shown in (13).

- (12) Mama-m Mzia-s daanteb-in-a cecxli
 father-ERG Mzia-DAT light-CAUS-AOR:3SG fire(ABS)
 ‘Father made Mzia light the fire.’ [kat] (Harris, 1981, in H&MB, p. 12)

$$(13) \left[\begin{array}{l} \text{HOOK} \left[\begin{array}{ll} \text{LTOP} & [5] \\ \text{INDEX} & [6] \\ \text{XARG} & [3] \end{array} \right. \\ \\ \text{RELS} \left\langle \begin{array}{l} \left[\begin{array}{ll} \text{_daanteb_v_light} \\ \text{LBL} & [8] \\ \text{ARGO} & [4] \textit{event} \\ \text{ARG1} & [1] \\ \text{ARG2} & [2] \end{array} \right] , \\ \left[\begin{array}{ll} \text{_mama_n_father} \\ \text{ARGO} & [3] \end{array} \right] , \\ \left[\begin{array}{ll} \text{_cecxli_n_fire} \\ \text{ARGO} & [2] \end{array} \right] \end{array} \right\rangle , \\ \\ \text{HCONS} \left\langle \begin{array}{l} \left[\begin{array}{ll} \textit{qeq} \\ \text{HARG} & [7] \\ \text{LARG} & [8] \end{array} \right] ! \end{array} \right\rangle \end{array} \right]$$

Note that, consistent with the strategy in Copestake et al. (2001), the scopal relationship is expressed by a handle constraint (hcons) rather than directly, representing equality modulo quantifiers ($=_q$). This handle constraint predicts that quantifiers can scope in between the EPs of the causative and embedded verb.

Similarly to my analysis of the applicative, the causative can also be decomposed into component operations, again parameterized along the axes of cross-linguistic variation. In the next section, I describe how these analyses were added to the Grammar Matrix.

3 Implementation in the Grammar Matrix

The Grammar Matrix customization system (Bender et al., 2010) combines a structured means of eliciting typological characteristics, validating responses for consistency, and using those choices to combine Matrix core grammar elements with stored analyses of various linguistic phenomena into a customized grammar. These stored analyses can include both static representations of cross-linguistically common phenomena as well as dynamically-generated implementations that embody language-specific variations. Elicitation is accomplished via a dynamic, iteratively-generated HTML questionnaire, which records the responses (while validating the consistency of both individual responses and their combination) in a structured choices file. This choices file is then processed by the customization script to produce the customized grammar.

My implementation of a library leverages the existing morphotactics machinery in the customization system (Goodman, 2013) by adding options to the questionnaire for grammar writers to attach valence-changing operations to lexical rule types, along with the relevant parameters (e.g., position of erstwhile subject)

necessary to generate the operations. My extensions to the grammar customization scripts, in turn, uses the selections in the choices file to generate the appropriate parameterized and common rule components, and then combine them into types to be instantiated.

While developing the library, two types of tests were used. Initially, I developed small, abstract pseudolanguages to exercise specific operations and combinations; I then attempted to model valence change in three natural languages, Lakota [lkt], Japanese [jpn], and Zulu [zul], and produced test suites of grammatical and ungrammatical examples. During this phase of development, I continued to revise my analyses and code to achieve full coverage of the examples. Once this phase was complete, I then froze library development and moved to the evaluation phase, described in the next section.

4 Evaluation

To evaluate the library as developed against a representative sample of the world's languages, I selected five held-out languages, from different familial and areal groups, that had not been used during development. Two languages were selected from descriptive articles intentionally held out, and the rest were selected by drawing randomly from a large collection of descriptive grammars, discarding those without valence changing morphology, until sufficient evaluation languages were collected.

I created test suites for each of these languages consisting of grammatical and ungrammatical examples of valence change, and attempted to model the corresponding phenomena using only the facilities available in the customization system questionnaire. I then attempted to parse the test suites using the customization system-generated grammars and recorded which grammatical examples were correctly parsed, which ungrammatical examples were erroneously parsed, and to what extent the parses generated spurious ambiguity. These results⁷ are summarized in Table 2.

On the test suites for the five held-out languages, this approach as implemented in my library achieved an overall coverage of 79% and an aggregate overgeneration rate of only 2%. The language with the poorest coverage (55%), Rawang [raw], suffered almost entirely due to a relatively rich system of reflexive and middle constructions; my library lacked the ability to fill a valence slot while coindexing with an existing argument and so these examples could not be modeled. The sole example of overgeneration, from Javanese [jav], was similarly due to the inability of the current library to apply a HEAD constraint to an already-existing argument. Neither of these failures appear to be particularly difficult to add to the library, which would significantly improve its flexibility and applicability.

⁷None of the test suites generated spurious ambiguity.

Language	Family	examples		performance		
		pos	neg	parses	coverage	overgen.
Tsez [ddo]	NE Caucasian	11	8	10	91%	0%
West Greenlandic [kal]	Eskimo-Aleut	15	14	12	73%	0%
Awa Pit [kwi]	Barbacoan	7	7	5	71%	0%
Rawang [raw]	Sino-Tibetan	11	6	6	55%	0%
Javanese [jav]	Austronesian	13	8	12	92%	13%
Total		57	43	45	79%	2%

Table 2: Test languages test summary and performance

5 Conclusion

In this work I have presented an HPSG analysis of valence-changing verbal morphology, implemented in the LinGO Grammar Matrix, which I evaluated against several held-out languages. The results appear to support the hypothesis that a “building-block” based approach is an effective way to provide significant typological coverage of valence change. By developing and implementing this analysis within the larger Grammar Matrix project, these elements of valence change can be combined and recombined in different ways to test linguistic hypotheses and compare modeling choices, including the interactions of valence change with other phenomena. Although the scope of this work was limited to valence change expressed through verbal morphology, future work might include determining whether this approach can be extended to other phenomena, including, for example, periphrastic valence-changing constructions.

References

- Bender, Emily M. 2015. (Un-)orthodoxy in use of HOOK features, semantic algebra compliance. MRS representation issues [discussion], presented at the Eleventh DELPH-IN Summit. <http://moin.delph-in.net/SingaporeSemanticAlgebraCompliance>.
- Bender, Emily M., Scott Drellishak, Antske Fokkens, Laurie Poulson & Safiyyah Saleem. 2010. Grammar customization. *Research on Language & Computation* 8(1). 23–72.
- Bender, Emily M., Dan Flickinger & Stephan Oepen. 2002. The Grammar Matrix: An Open-Source Starter-Kit for the Rapid Development of Cross-linguistically Consistent Broad-coverage Precision Grammars. In John Carroll, Nelleke Oostdijk & Richard Sutcliffe (eds.), *Proceedings of the Workshop on Grammar Engineering and Evaluation at the 19th International Conference on Computational Linguistics*, 8–14. Taipei, Taiwan.
- Bouma, Gosse, Robert Malouf & Ivan A. Sag. 2001. Satisfying constraints on

- extraction and adjunction. *Natural Language & Linguistic Theory* 19(1). 1–65.
- Bybee, Joan L. 1985. *Morphology: A study of the relation between meaning and form*. Amsterdam: John Benjamins Publishing.
- Chung, Sandra. 1976. An object-creating rule in Bahasa Indonesia. *Linguistic Inquiry* 7. 41–87.
- Copestake, Ann. 2002. Definitions of typed feature structures. In Stephan Oepen, Dan Flickinger, Jun-ichi Tsujii & Hans Uszkoreit (eds.), *Collaborative language engineering*, 227–230. Stanford, CA: CSLI Publications.
- Copestake, Ann, Dan Flickinger, Carl Pollard & Ivan A. Sag. 2005. Minimal recursion semantics: An introduction. *Research on Language and Computation* 3(2). 281–332.
- Copestake, Ann, Alex Lascarides & Emily M. Bender. 2016. Algebra replacement. Revising the algebra [discussion], presented at the Twelfth DELPH-IN Summit. <http://moin.delph-in.net/StanfordAlgebraAdditions>.
- Copestake, Ann, Alex Lascarides & Dan Flickinger. 2001. An algebra for semantic construction in constraint-based grammars. In *Proceedings of the 39th Annual Meeting of the Association for Computational Linguistics ACL '01*, 140–147. Stroudsburg, PA, USA: Association for Computational Linguistics.
- Dayley, Jon Philip. 1989. *Tümpisa (Panamint) Shoshone grammar*, vol. 115. Univ of California Press.
- Dixon, R. M. W. & Alexandra Y. Aikhenvald. 2000. *Changing valency*. Cambridge University Press.
- Goodman, Michael W. 2013. Generation of machine-readable morphological rules from human-readable input. *University of Washington Working Papers in Linguistics* 30.
- Harris, Alice C. 1981. *Georgian syntax: A study in relational grammar*. Cambridge: Cambridge University Press.
- Haspelmath, Martin & Thomas Müller-Bardey. 2004. Valence change. *Morphology: A handbook on inflection and word formation* 2. 1130–1145.
- Müller, Stefan. 2018. *A lexicalist account of argument structure: Template-based phrasal LFG approaches and a lexical HPSG alternative* (Conceptual Foundations of Language Science 2). Berlin: Language Science Press.
- Wunderlich, Dieter. 2015. Valency-changing word-formation. In Peter O. Müller, Ingeborg Ohnheiser, Susan Olsen & Franz Rainer (eds.), *Word-formation*, vol. 3, 1424–1466. Berlin/Boston: De Gruyter Mouton.

Sign language agreement: A constraint-based perspective

Anke Holler 

University of Göttingen

Markus Steinbach 

University of Göttingen

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 49–65

Keywords: HPSG, DGS, Agreement, PAM

Holler, Anke & Markus Steinbach. 2018. Sign language agreement: A constraint-based perspective. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 49–65. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.4.

 CC-BY 4.0

Abstract

The paper addresses verbal agreement in German Sign Language from a constraint-based perspective. Based on Meir’s Agreement Morphology Principles it presents an HPSG analysis of plain, regular and backwards agreement verbs that models the interaction between phonological (manual) features and syntactico-semantic relationships within a verbal sign by well-defined lexical restrictions. We argue that a sign-based declarative analysis can provide an elegant approach to agreement in sign language since it allows to exploit cross-modular constraints within grammar, and hence permits a direct manipulation of all relevant phonological features of a verb depending on its syntactic and semantic properties.

1 Introduction

Agreement between a verb and two of its arguments is one of the best studied areas in sign language linguistics (Lillo-Martin & Meier 2011). The range of analyses varies from gesturally oriented approaches via semantic, i.e. thematic, accounts up to purely syntactic implementations.¹ In the present paper, we argue for a constraint-based modeling of sign language agreement because it allows for a combination of the insights of both semantic and syntactic approaches. As we will show below, a constraint-based account has the noteworthy advantage that (manual) phonological features of verbs that inflect for agreement, such as beginning and end point of path movement as well as hand orientation, can be manipulated in a direct way.

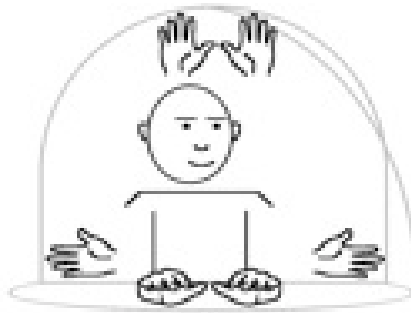


Figure 1: Signing space

Agreement in sign languages is locus agreement, which means that it is expressed in the signing space by a manipulation of phonological features. The relevant phonological features of the verb agree with or depend on the

[†]We thank the reviewers and the audience of the HPSG 2018 conference in Tokyo for discussion and valuable comments.

¹For a deeper discussion of these analyses, see Salzmann et al. 2018.

referential locations (R-loci) the discourse referents of the subject and object are linked to in the signing space (Steinbach & Onea 2016). These R-loci are either actual locations of present referents (i.e. deictic locations) or locations that are assigned for non-present referents on the horizontal plane of the signing space (i.e. anaphoric locations, cf. figure 1). Non-present discourse referents can be localized in various ways. One major strategy is the use of determiner-like signs such INDEX_X and POSS_X . The first referent is typically assigned to the ipsilateral area of the signing space and the second one to the contralateral area (cf. figure 2).

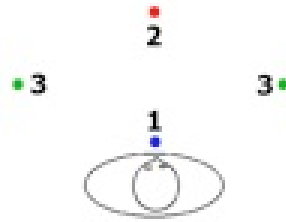


Figure 2: Localization of referents

To give an example: In the first sentence of (1) below, the first discourse referent MARIA is localized with the pointing sign INDEX_{3a} at location 3a, which is the ipsilateral area of the signing space, i.e. the right side for right-handed signers (cf. figure 2). Similarly, the second discourse referent NEW TEACHER is localized at the contralateral area of signing space, i.e. 3b. This R-locus is then used to pronominalize NEW TEACHER in the second sentence.²

- (1) M-A-R-I-A INDEX_{3a} TEACHER NEW INDEX_{3b} LIKE. INDEX_{3b} SMART.
‘Maria likes the new teacher. S/he is smart.’

The two R-loci introduced in the first sentence can also be used to express agreement on the verb GIVE by moving from the R-locus associated with the subject to the R-locus associated with the object. This is illustrated by (2) below. Hence, sign languages, just like spoken languages, use similar means for pronominalization and agreement. However, unlike spoken languages, sign languages do not use sequential agreement affixes but express referential indices of the subject and object simultaneously on the verb. (cf. Aronoff et al. 2005).

- (2) YESTERDAY INDEX_{3a} BOOK $_{3a}\text{GIVE}_{3b}$
‘Yesterday she gave him a book.’

We address verbal agreement in German Sign Language (DGS) from a constraint-based perspective in this paper. In particular we aim at modeling

²This is a second revised version of the originally published paper because of an obvious erratum in referencing to example (1). There are no other differences to the first version instead of replacing this erroneous paragraph [date of correction: 15-10-2019].

well-known restrictions on agreement in sign languages in such a way that the interaction between phonological (manual) features and syntactico-semantic relationships can be adequately described. We show that a constraint-based approach offers an elegant analysis of sign language agreement since it permits a direct manipulation of the relevant phonological features of the verb sign.

This article is organized as follows: In the following section, we describe the general properties of agreement in sign languages and introduce the three most important verb classes, i.e. regular agreement verbs, backwards agreement verbs and plain verbs. Section 3 introduces so-called agreement auxiliaries that are used to mark agreement manually in the case of plain verbs which are not able to express agreement. In section 4 we present and discuss a constraint-based analysis of sign language agreement couched in standard HSPG theory, that is, our analysis does not depend on any modality-specific assumptions or modifications of HPSG.

2 Agreement in sign languages

The huge amount of studies on agreement in many different sign languages has shown that agreement in the visual-gestural modality differs in several respects from agreement in the auditory-oral modality (Lillo-Martin & Meier 2011; Mathur & Rathmann 2012; Salzmann et al. 2018).

First of all, it is well documented that not all verbs in a sign language are able to realize verbal agreement overtly. In addition to so-called agreement verbs such as GIVE in example (2) above, sign languages also have so-called plain verbs such as LIKE in the first sentence of example (1) above, which cannot be inflected for agreement. A third class of verbs are so-called spatial verbs, whose beginning and endpoints are not determined by arguments of the verb (or grammatical functions) but by topographic referents. Like agreement verbs, spatial verbs can be spatially modified but the controller of the agreement is not a locus with a referential interpretation but a locus with a topographic interpretation (e.g. the village on the left, the house on the right, ...). Examples of typical DGS verbs for each of these three verb classes are listed in (3). In the following, we ignore spatial verbs since the topographic relations expressed are not agreement relations in the strict sense but descriptions of the location or movement of an entity in the real world.

- (3) a. Agreement verbs: GIVE, HELP, TEACH, ASK, VISIT, SHOW, ...
- b. Plain verbs: LIKE, KNOW, WAIT, THINK, BUY, ...
- c. Spatial verbs: MOVE, PUT, STAND, LIE, BE-AT, ...

Secondly, verbs in sign languages express agreement with their arguments directly in the signing space by path movement and/or orientation of the

hands (i.e. palm orientation or orientation of the fingertips, cf. Meir 1998). With the DGS verb GIVE in (4a), path movement begins at x, the R-locus associated with the discourse referent of the subject, and ends at y, the R-locus associated with the discourse referent associated with the object. By contrast, the DGS verb INFLUENCE does not only express agreement by path movement but also by orientation of the hands. In (4b) the fingertips are oriented towards the location associated with the object, i.e. y.

- (4) a. x GIVE y
 ‘to give something to someone’
 b. x INFLUENCE y
 ‘to influence someone’

The examples in (4) also illustrate another property of agreement in sign languages: Verbs in DGS do not only agree with the subject (first argument) but also with the object (second argument). Subject and object agreement is the standard case not only in DGS but also in many different unrelated sign languages.

A fourth important property of agreement in sign languages is that it affects directly the phonological form of the verb. Agreement is expressed through the manipulation or specification of the two phonological features hand orientation and path movement of the corresponding agreement verb. Consequently, phonological properties of the verb may block the overt realization of agreement. This is the case with plain verbs: Agreement with subject and object is prohibited because hand orientation and the beginning and endpoint of path movement are lexically specified. Consider, for instance, the plain verb LIKE in example (1) above. Path movement always involves a downward movement of the dominant hand in front of the signers chest. Therefore, this movement cannot be modified and adapted to the R-loci that subject and object are linked to. Even with agreement verbs, agreement may sometimes be blocked by phonological constraints. In some varieties of DGS, verbs like TRUST only agree with first person subjects and non-first person objects because the beginning of the path movement is lexically specified (i.e. the forehead of the signer). In other varieties of DGS, however, the verb TRUST also inflects with non-first person subjects and first person objects, which means that it has already been developed into a full subject-object agreement verb. In these varieties, the inflected form in (5b) would be grammatical.

- (5) a. $_1$ TRUST $_2$
 ‘I trust you.’
 b. $*_2$ TRUST $_1$
 ‘You trust me.’

A fifth unique property of sign language agreement, which is highly relevant for each analysis, is the distinction between two different kinds of agreement verbs: regular and backwards agreement verbs. Regular agreement verbs follow the pattern described above. The path movement starts at the R-locus associated with the subject and ends at the R-locus of the object. By contrast, backwards agreement verbs such as INVITE in (6) below show the reverse pattern. The path movement begins at the position of the object and ends at the position of the subject. Interestingly, the hand is always oriented towards the object, even with backwards agreement verbs. We will see that the general distinction between regular and backwards agreement verbs (i.e. the difference in movement direction) can be derived from thematic restrictions discussed in Meir (1998, 2002). By contrast, the specification of the hand orientation follows from syntactic restrictions on the COMPS list.

- (6) $_2\text{INVITE}_1$
 ‘I invite you.’

The following figure 3 gives an overview of the agreement picture described in this section. Note that these modality-specific properties of agreement in sign languages and the specific verb classes follow from the spatial nature (path movement and orientation) of sign language agreement and its gestural origins (transfer of an entity). This does, however, not mean that sign language agreement is not part of the linguistic system (for a more detailed discussion, cf. Salzmann et al. 2018).

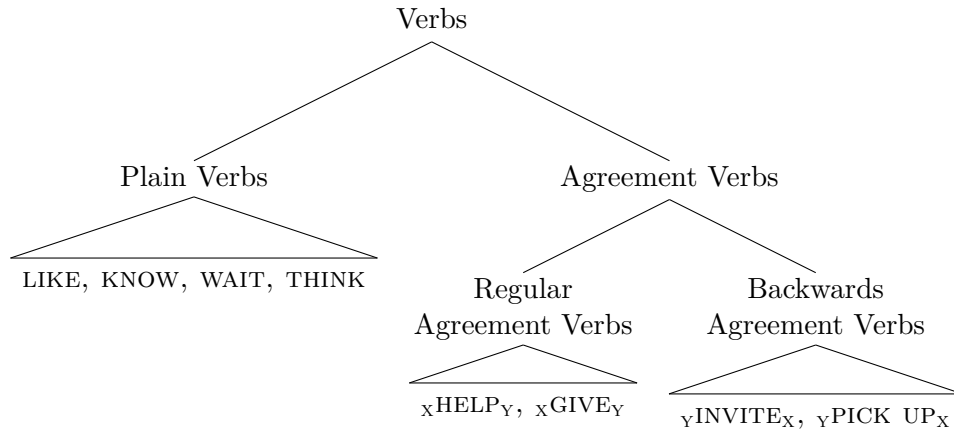


Figure 3: Verb classes in German Sign Language

3 Agreement auxiliaries

In the previous section, we have shown that plain verbs such as LIKE in example (1) cannot be inflected for agreement. Interestingly, many sign languages have developed various grammatical means to overcome the agreement gap caused by plain verbs. These sign languages make either use of a specific class of auxiliaries (so-called agreement auxiliaries) or they use non-manual markers such as eye gaze and head tilt to express the agreement relations with plain verbs (Steinbach & Pfau 2007; Sapountzaki 2012; Neidle et al. 2000; Thompson et al. 2006).

In the following, we only focus on agreement auxiliaries since DGS belongs to the group of sign languages that make use of manual agreement markers. Like agreement verbs, agreement auxiliaries express subject and object agreement by means of path movement and hand orientation. Agreement auxiliaries in sign languages differ from typical spoken language auxiliaries in that they are not used to mark tense, aspect, modality, or voice (so-called TAM auxiliaries) but “only” to mark agreement with the subject and the object. Genuine agreement auxiliaries seem to be rare in spoken languages. The German auxiliary *tun* (‘to do’) in (7a), which is frequently used in colloquial variants of German and in many German dialects, might be an exception to this generalization. Unlike other auxiliaries in German, *tun* is not a TAM marker, it is not restricted to certain semantic contexts (the corresponding sentence without *tun* in example (7b) is functionally identical to its counterpart in (7a)) and its use seems to be functionally very similar to agreement auxiliaries in sign languages (Erb 2001; Steinbach & Pfau 2007).

- (7) a. Sie tu-t ein Buch les-en.
 She do-3.sg a book read-inf
 b. Sie lies-t ein Buch.
 She read-3.sg a book
 ‘She is reading a book.’

The auxiliary *tun* seems to be some kind of dummy auxiliary that is only used to express morphosyntactic features such as present and past tense and agreement. Note that these features can always be optionally expressed by the main verb as illustrated in (7b). Hence, *tun* resembles the DGS agreement auxiliary, the Person Agreement Marker PAM (cf. Rathmann 2003; Steinbach & Pfau 2007).

The source of the DGS agreement auxiliary PAM is the noun PERSON as demonstrated by figure 4. Contrary to PAM, PERSON does not exhibit a directional movement but only a simple downward movement. The agreement auxiliary PAM, however, expresses the agreement relation by a manipulation of the phonological features path movement and hand orientation and behaves in this respect just like regular agreement verbs.

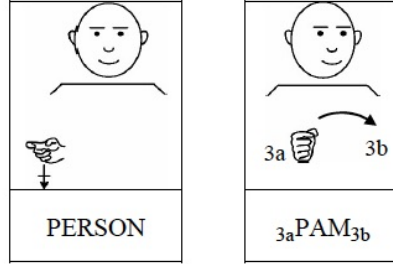


Figure 4: From noun to auxiliary in DGS

PAM can be used with plain verbs as in (8a), with adjectival predicates as in (8b), and with verbs like TRUST, which cannot be inflected for non-first person subject agreement and first person object agreement as in (8c).

- (8) a. MOTHER INDEX_{3a} NEIGHBOR NEW INDEX_{3b} LIKE _{3a}PAM_{3b}
 ‘(My) mother likes the new neighbor.’
 b. INDEX₁ POSS₁ BROTHER INDEX_{3a} PROUD ₁PAM_{3a}
 ‘I am proud of my brother.’
 c. INDEX₂ TRUST ₂PAM₁
 ‘You trust me.’

Note that there seems to be some variation in the syntactic position of PAM. In Southern German variants, PAM is preferably inserted in pre-verbal position (even before the object) as can be seen in (9), whereas in most variants of DGS, PAM is usually inserted in post-verbal position as illustrated in example in (8a) above (Rathmann 2003; Macht 2016; Macht & Steinbach, to appear). In example (9), ‘H-A-N-S_{3a}’ means that the name ‘Hans’ is fingerspelled at the location 3a, i.e. fingerspelled names can be directly linked to R-loci.

- (9) H-A-N-S_{3a} _{3a}PAM_{3b} M-A-R-I-E_{3b} LIKE
 ‘Hans likes Maria.’

Interestingly, PAM can also be combined with uninflected agreement verbs. Although this combination seems to be less acceptable than the version with inflected agreement verb without PAM, it reveals interesting insights in the inflectional pattern of PAM. With uninflected backwards verbs like INVITE in (10), PAM does not follow the inflectional pattern of the backwards verb but moves from the position of the subject to the position of the object, i.e. even in the context of backwards agreement verbs, PAM inflects like a regular agreement verb. Hence, the semantically empty agreement auxiliary PAM generally expresses agreement with subject and object, no matter of the thematic structure of the corresponding main verb.

- (10) INDEX_{3a} INDEX_{3b} INVITE _{3a}PAM_{3b}
 ‘S/he invites him/her.’

Consequently, PAM is not subject to any semantic restriction and can be used with all kinds of plain verbs.

- (11) DGS plain verbs that express agreement by means of PAM:
 BE-PROUD, BE-ANGRY, KNOW, LIKE, TRUST, WAIT, BE-INTERESTED-
 IN, LAUGH, ...

Note finally that PAM can also be productively used to extend the argument structure of the main verb. Since PAM is a transitive agreement marker, it can be used as a transitivizer in DGS.

- (12) a. INDEX₁ LAUGH ₁PAM₂
 ‘I laugh at you.’
 b. INDEX₁ LETTER WRITE ₁PAM₂
 ‘I write a letter to you.’

4 A lexical analysis of agreement in HPSG

The specific phonological and semantic properties of the three different verb classes discussed in the previous section and the interaction between their formal (phonological and syntactic) and semantic (thematic) properties call for a constraint-based lexical treatment of verbal agreement in sign language. Such an approach not only enables the formulation of cross-modular restrictions within grammar but also allows for a direct relation of phonological and argument structural information within a sign. In particular, the thematic conditions and the interaction with phonological features highlighted in the previous section can explicitly be stated in the lexical entry of a verb. Such an approach perfectly meets with the insights formulated in the thematic approach in Meir (1998, 2002) and the HPSG account sketched in Cormier et al. (1999). In this section, we build on these two approaches and develop an HPSG analysis of (regular and backwards) agreement and plain verbs on the one hand and the agreement auxiliary PAM on the other.

4.1 Basic assumptions for lexical signs

A lexical item in sign language structurally differs from a lexical sign in spoken language as it includes a description of the manual phonology, which consists of a particular handshape, a location, a movement, and a hand orientation as well as a description of the non-manual phonology (whose lexical aspects we mainly ignore in the following). Thus, the phonological component of a sign language is much more complex than the corresponding phonological component of a spoken language. This has its reflex in

the structure of the PHON value. Following Safar & Marshall (2004), we assume that PHON represents relevant aspects of non-manual phonology such as the face, especially the brows, and the mouthing as well as comprises fine-grained information about the hand(s) with respect to shape, orientation and movement. A partial description of PHON adapted from Safar & Marshall (2004) is given in (13). The most important part for our analysis of agreement is, of course, the manual features movement and orientation. Movement of the hand(s) is defined by two positions in the signing space which mark the beginning and the end point of the movement. Orientation is defined by palm and finger orientation.

$$(13) \quad \left[\begin{array}{c} \text{PHON} \\ \left[\begin{array}{cc} \text{FACE|BROW} & \text{brow} \\ \text{MANUAL} & \left[\begin{array}{cc} \text{HANDSHAPE} & \text{handshape} \\ \text{ORIENTATION} & \left[\begin{array}{cc} \text{PALM} & \text{palm} \\ \text{FINGER|INDEX|LOC} & \text{locus} \end{array} \right] \\ \text{MOVEMENT} & \left\langle \left[\begin{array}{cc} \text{BEGIN|INDEX|LOC} & \text{locus} \\ \text{END|INDEX|LOC} & \text{locus} \end{array} \right] \right\rangle \end{array} \right] \\ \text{MOUTH|PICTURE} & \text{picture} \end{array} \right] \end{array} \right]$$

As discussed in the previous sections discourse referents (and thus indices) in sign languages are linked to R-loci in the signing space. In order to represent these R-loci, we have to redefine the INDEX value as is also illustrated in (13). To account for the differences between spoken and sign languages with respect to their index values we suggest to define two new subtypes of the type *index*, called *categorical_index* and *positional_index* as is depicted in (14). This accounts for our general assumption that the type *index* can be thought of as an HPSG analog of a reference marker in Discourse Representation Theory (cf. Kamp & Reyle 1993).

$$(14) \quad \begin{array}{ll} \text{a.} & \left[\text{INDEX} \left[\begin{array}{c} \text{categorical_index} \\ \text{PERSON } \textit{person} \\ \text{NUMBER } \textit{number} \\ \text{GENDER } \textit{gender} \end{array} \right] \right] \\ \text{b.} & \left[\text{INDEX} \left[\begin{array}{c} \text{positional_index} \\ \text{LOCUS } \textit{locus} \end{array} \right] \right] \end{array}$$

The INDEX value of type *categorical_index* is exploited for spoken languages and represents the usual morpho-syntactic features like person, number and gender. However, for sign languages, we follow Cormier et al. (1999) in stipulating a *positional_index* which refers to specific loci in the signing space. These are represented by a LOCUS value. For the LOCUS feature a type *locus* is appropriate which is further partitioned into the subtypes

speaker, *addressee* and *other*, where *other* subsumes a set of variables, *i*, *j*, *k*, etc., representing possible indices.

Next, we come back to the observations concerning agreement in sign languages discussed in the previous section. We will develop an analysis that accounts for the two basic verb classes in DGS, i.e. agreement verbs as well as plain verbs.

4.2 Agreement verbs

We follow Meir (1998, 2002) in the distinction between two kinds of agreement in sign language, (i) thematic agreement, and (ii) syntactic agreement, as formulated in the Agreement Morphology Principles (AMP). Below, we implement the AMP directly into our HPSG analysis to take up the generalization that thematic agreement marks the direction of the path movement (see 15a) whereas syntactic agreement is responsible for the orientation of the hand(s) (see 15b).

- (15) Agreement Morphology Principles (AMPs):
- a. The direction of the path movement of agreement verbs is from source to goal [...]
 - b. The facing of the hand(s) is towards the object of the verb.

The Agreement Morphology Principles account for both, regular and backwards verbs, which share the facing of the hands but differ in the direction of the path movement. According to (15a), the direction of the path movement is controlled by the thematic roles source and goal which could be mapped on the arguments of FROM and TO in Jackendoff's (1990) componential analysis. The facing of the hands, on the other hand, is controlled by the indirect object which is comparable with the dative object in spoken language.

To account for the agreement facts of DGS and to model Meir's principles in a constraint-based way, we develop a type-based representation of the existing classes of agreement verbs. In a first step genuine agreement verbs are distinguished from plain verbs by stipulating two verbal subtypes, called *plain_verb* and *agr(eement)_verb* respectively. Secondly, the type *agreement_verb* is further partitioned by two subtypes which are called *reg(ular)_agreement_verb* and *back(wards)_agreement_verb* in accordance with the analysis of Cormier et al. (1999). Additionally, there is a transitive and ditransitive variant of both subtypes. This is illustrated by the resulting signature in figure 5.

With this type hierarchy at hand, we can now define appropriate lexical constraints that restrict verbal agreement with respect to a certain verbal class.

Based on the usual HPSG practice to model agreement as structure-sharing between INDEX values, we analyze syntactic agreement in DGS by

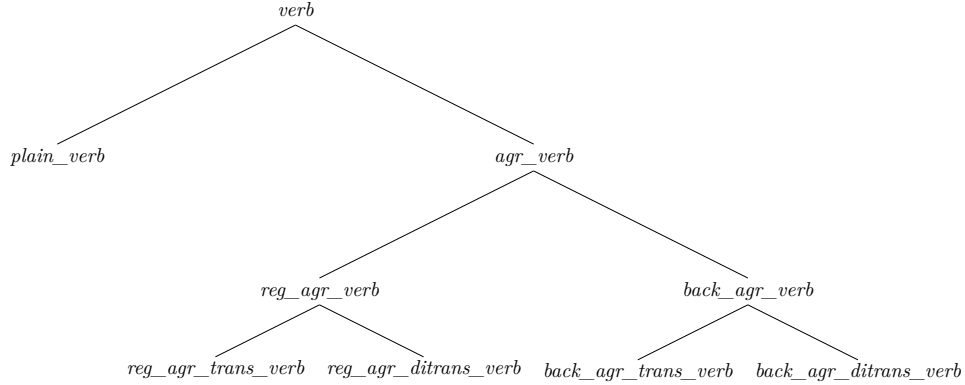


Figure 5: Partition of type *verb*

manipulating the ORIENTATION value of MANUAL and structure-sharing its INDEX value with the INDEX value of the indirect object on the COMPS-list of the respective verb. This accounts for Meir’s definition of syntactic agreement in (15b), where the facing of the hands is syntactically controlled by the respective object in the argument structure. (16) shows the partial description of the PHON value that we assume for all verbs of type *agreement_verb* in the lexicon. The analysis is built on the analysis developed in Safar & Marshall (2004).

$$(16) \left[\begin{array}{l} \text{agreement_verb} \\ \text{PHON|MANUAL} \\ \text{SYNSEM|LOC|CAT|COMPS} \end{array} \left[\begin{array}{l} \text{ORIENT|FINGER|INDEX } \boxed{1} \text{ locus} \\ \text{MOVEMENT} \left\langle \begin{array}{l} \text{BEGIN|INDEX locus} \\ \text{END|INDEX locus} \end{array} \right\rangle \\ < \dots, \text{NP}_{\boxed{1}}, \dots > \end{array} \right] \right]$$

To implement Meir’s first clause of the Agreement Morphology Principles, which expresses her observation on thematic agreement, we add a relation of type *transfer* to the CONTENT value for all verbs of type *agreement_verb*. This relation comes with three arguments: SOURCE, GOAL and SOA. Our implementation of thematic agreement relies on the manipulation of the MOVEMENT value: the BEGIN value of MOVEMENT is identified with the SOURCE value of the transfer relation and the END value with GOAL value of the same relation. This accounts for agreement as path movement in sign language.

$$(17) \left[\begin{array}{l} \textit{agreement_verb} \\ \text{PHON|MANUAL} \left[\text{MOVEMENT} \left\langle \begin{array}{l} \text{BEGIN|INDEX } \boxed{2} \\ \text{END|INDEX } \boxed{1} \end{array} \right\rangle \right] \\ \text{SYNSEM|LOCAL} \left[\text{CONT} \left[\begin{array}{l} \text{RELATION } \textit{transfer} \\ \text{SOURCE } \boxed{2} \\ \text{GOAL } \boxed{1} \\ \text{SOA } \textit{qfpsoa} \end{array} \right] \right] \end{array} \right]$$

Note that the semantics of any agreement verb is introduced by the SOA value of the *transfer* relation. This is necessary to prevent the prediction of a semantic hierarchy in which all semantic relations that are expressed by agreement verbs are at the same time subcases of a general transfer relation. Cognitively, this might be correct but in this paper, we do not argue for such a strong assumption and our analysis does not hinge on it.

The main difference between regular and backwards agreement verbs basically concerns the direction of the path movement which is mediated by the argument structural properties of the respective verbs. Following Meir's insights on thematic agreement, we assume that the path movement begins at the position of the subject (SOURCE) and ends at the position of the object (GOAL) in case of regular agreement verbs. By contrast, with backwards agreement verbs, path movement works the other way around. In this case, the movement starts at the position of the object (GOAL) and ends at the position of the subject (SOURCE). Again, this is realized as structure-sharing of positional INDEX values as can be seen in (18) and (19) respectively.

$$(18) \left[\begin{array}{l} \textit{regular_agreement_verb} \\ \text{SYNSEM|LOCAL} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{SUBJ } \langle \text{NP}_{\boxed{2}} \rangle \\ \text{COMPS } \langle \dots, \text{NP}_{\boxed{1}}, \dots \rangle \end{array} \right] \\ \text{CONT} \left[\begin{array}{l} \text{RELATION } \textit{transfer} \\ \text{SOURCE } \boxed{2} \\ \text{GOAL } \boxed{1} \\ \text{SOA } \textit{qfpsoa} \end{array} \right] \end{array} \right] \end{array} \right]$$

$$(19) \left[\begin{array}{c} \text{backwards_agreement_verb} \\ \text{SYNSEM|LOCAL} \end{array} \left[\begin{array}{c} \text{CAT} \left[\begin{array}{c} \text{SUBJ} <\text{NP}_{\boxed{2}}> \\ \text{COMPS} <\dots, \text{NP}_{\boxed{1}}, \dots> \end{array} \right] \\ \text{CONT} \left[\begin{array}{c} \text{RELATION } \textit{transfer} \\ \text{SOURCE} \boxed{1} \\ \text{GOAL} \boxed{2} \\ \text{SOA} \textit{qfsoa} \end{array} \right] \end{array} \right] \right]$$

The only difference between the restrictions of both verb classes consists in the assignment of the indices. With regular agreement verbs the object is identified with the goal of the transfer relation, whereas with backwards agreement verbs, the subject is identified with the goal of the transfer relation. This is illustrated by the following structures. The full lexical specifications of the transitive regular agreement verb *HELP* and the ditransitive regular agreement verb *GIVE* are depicted in (20) and (21). By contrast, (22) exemplifies the reverse specification of the transfer relation for the backwards agreement verb *INVITE*.

$$(20) \left[\begin{array}{c} \text{regular_agreement_trans_verb} \\ \text{PHON|MANUAL} \\ \text{SYNSEM|LOCAL} \end{array} \left[\begin{array}{c} \left[\begin{array}{c} \text{ORIENT|FINGER|INDEX} \boxed{1} \\ \text{MOVEMENT} \left\langle \left[\begin{array}{c} \text{BEGIN|INDEX} \boxed{2} \\ \text{END|INDEX} \boxed{1} \end{array} \right] \right\rangle \end{array} \right] \\ \text{CAT} \left[\begin{array}{c} \text{SUBJ} <\text{NP}_{\boxed{2}}> \\ \text{COMPS} <\text{NP}_{\boxed{1}}> \end{array} \right] \\ \text{CONT} \left[\begin{array}{c} \text{RELATION } \textit{transfer} \\ \text{SOURCE} \boxed{2} \\ \text{GOAL} \boxed{1} \\ \text{SOA} \left[\begin{array}{c} \textit{help_rel} \\ \text{HELPER} \boxed{2} \\ \text{HELPEE} \boxed{1} \end{array} \right] \end{array} \right] \end{array} \right] \right]$$

$$(21) \left[\begin{array}{c} \text{regular_agreement_ditrans_verb} \\ \text{PHON|MANUAL} \left[\begin{array}{c} \text{ORIENT|FINGER|INDEX } \boxed{1} \\ \text{MOVEMENT } \langle \begin{array}{c} \text{BEGIN|INDEX } \boxed{2} \\ \text{END|INDEX } \boxed{1} \end{array} \rangle \\ \text{CAT} \left[\begin{array}{c} \text{SUBJ } \langle \text{NP}_{\boxed{2}} \rangle \\ \text{COMPS } \langle _, \text{NP}_{\boxed{1}} \rangle \end{array} \right] \\ \text{SYNSEM|LOCAL} \left[\begin{array}{c} \text{RELATION } \textit{transfer} \\ \text{SOURCE } \boxed{2} \\ \text{GOAL } \boxed{1} \\ \text{CONT} \left[\begin{array}{c} \textit{give_rel} \\ \text{GIVER } \boxed{2} \\ \text{GIFT } \boxed{3} \\ \text{GIVEE } \boxed{1} \end{array} \right] \\ \text{SOA} \end{array} \right] \end{array} \right] \end{array} \right]$$

$$(22) \left[\begin{array}{c} \text{backwards_agreement_trans_verb} \\ \text{PHON|MANUAL} \left[\begin{array}{c} \text{ORIENT|FINGER|INDEX } \boxed{1} \\ \text{MOVEMENT } \langle \begin{array}{c} \text{BEGIN|INDEX } \boxed{1} \\ \text{END|INDEX } \boxed{2} \end{array} \rangle \\ \text{CAT} \left[\begin{array}{c} \text{SUBJ } \langle \text{NP}_{\boxed{2}} \rangle \\ \text{COMPS } \langle \text{NP}_{\boxed{1}} \rangle \end{array} \right] \\ \text{SYNSEM|LOCAL} \left[\begin{array}{c} \text{RELATION } \textit{transfer} \\ \text{SOURCE } \boxed{1} \\ \text{GOAL } \boxed{2} \\ \text{CONT} \left[\begin{array}{c} \textit{invite_rel} \\ \text{INVITER } \boxed{2} \\ \text{INVITEE } \boxed{1} \end{array} \right] \\ \text{SOA} \end{array} \right] \end{array} \right] \end{array} \right]$$

4.3 Plain verbs

Recall that in the case of plain verbs such as KNOW and LIKE, phonological properties of the verb block the overt realization of agreement. This means that agreement with subject and object is prohibited because the beginning and endpoint of path movement and hand orientation are already lexically specified. As is illustrated in (23) the respective LOCUS values are instantiated by fixed values (i.e. lexically specified loci in the signing space) expressed by the variables i , j and k representing indexical reference points.

$$(23) \left[\begin{array}{l} \textit{plain_verb} \\ \text{PHON}|\text{MANUAL} \\ \text{SYNSEM}|\text{LOCAL}|\text{CAT} \end{array} \left[\begin{array}{l} \text{ORIENT}|\text{FINGER}|\text{INDEX}|\text{LOC } k \\ \text{MOVEMENT} \left\langle \left[\begin{array}{l} \text{BEGIN}|\text{INDEX}|\text{LOC } i \\ \text{END}|\text{INDEX}|\text{LOC } j \end{array} \right] \right\rangle \\ \left[\begin{array}{ll} \text{HEAD} & \textit{verb} \\ \text{SUBJ} & \textit{nelist} \\ \text{COMPS} & \textit{nelist} \end{array} \right] \end{array} \right] \right]$$

As discussed above, sign languages have developed different means to overcome the agreement gap caused by plain verbs. DGS, for instance, makes use of the agreement auxiliary PAM, which, just like regular agreement verbs, marks agreement manually by means of hand orientation and path movement. Therefore, PAM insertion is a practicable option to express agreement overtly with plain verbs. Since the relevant phonological features ORIENTATION and MOVEMENT are already lexically specified with plain verbs and hence not available for agreement inflection, PAM can be used to agree with the subject and object of the plain verb and realize the corresponding features overtly.

In principle, there are different HPSG analyses available that have been proposed to account for several kinds of auxiliaries in spoken language and could be used to account for PAM. One option is a lexical analysis of auxiliaries as proposed by Ackerman & Webelhuth (1998). Following this account, PAM would be added to the lexical entry of a plain verb. An alternative option would be that PAM is subcategorized for a plain verb, and attracts all relevant arguments which are necessary to express agreement from this verb. This analysis accommodates the construction of verbal clusters in German by argument composition (cf. Hinrichs & Nakazawa 1989, 1994; Müller 2007). It ensures that the agreement auxiliary PAM may exploit path movement and hand orientation to express subject and object agreement. Hence, PAM does not differ from regular agreement verbs in this respect. Nevertheless, PAM acts as a syntactic marker only as it makes no use of the transfer relation as defined for regular and backwards agreement verbs.

The partial description in (24) gives the lexical specification of PAM. It depicts that PAM selects a verb of type *plain_verb* and attracts all arguments of this verb, which comprises the subject marked by tag [3] and the whole COMPS list marked by tag [4]. Since the indices of the plain verb's subject and object are structure shared with the beginning and the end point of the MOVEMENT feature of PAM, agreement is expressed purely syntactically.

$$(24) \left[\begin{array}{l} \textit{personal_agreement_marker} \\ \text{PHON|MAN} \left[\begin{array}{l} \text{ORIENT|FINGER|INDEX } [2] \\ \text{MOVEMENT } \langle \begin{array}{l} \text{BEGIN|INDEX } [1] \\ \text{END|INDEX } [2] \end{array} \rangle \end{array} \right] \\ \text{SS|LOC} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{HEAD } \textit{verb} \\ \text{SUBJ } [3] \\ \text{COMPS } [4] \oplus \langle V[\textit{plain}, \text{SUBJ } [3] \langle \text{NP}_{[1]} \rangle, \text{COMPS } [4] \langle \dots, \text{NP}_{[2]}, \dots \rangle \rangle] \end{array} \right] \end{array} \right] \end{array} \right]$$

In order to account for the variation in the positioning of PAM mentioned above, lexical precedence rules are needed that regulate whether PAM has to be positioned pre- or postverbally. In cases where PAM is used to extend the argument structure of the selected verb, an analysis is conceivable that adds an argument to the COMPS list of PAM in dependence of a featural specification that marks that the corresponding main verb is one that qualifies in principle for argument structural extensions. Alternatively, one might argue that the COMPS list of PAM is inherently specified for an object, which is then added to the COMPS list of a one-place main verb and triggers a corresponding transitive interpretation.

5 Conclusion

In sum, the HPSG analysis of agreement in DGS developed in this paper illustrates that a constraint-based lexical approach offers an elegant account of the modality-specific properties of sign language agreement. In particular, the interdependence of phonological, syntactic, and semantic properties of the verb and the simultaneous realization of agreement can be implemented in a straightforward way using cross-modular constraints on syntactic and thematic agreement in DGS. Moreover, we can account for agreement in DGS without assuming additional morphosyntactic features or specific agreement morphemes since the agreement principles directly operate on phonological locus features of the verb and its arguments. Finally, our analysis correctly predicts the distribution of the agreement auxiliary PAM in DGS.

Nominalized clauses in the Grammar Matrix

Kristen Howell 

University of Washington

Olga Zamaraeva 

University of Washington

Emily M. Bender 

University of Washington

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 66–86

Keywords: HPSG, Typology, Grammar Engineering, Nominalized clauses, Nominalization, Grammar Matrix

Howell, Kristen, Olga Zamaraeva & Emily M. Bender. 2018. Nominalized clauses in the Grammar Matrix. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 66–86. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.5.

 CC-BY 4.0

Abstract

We present an analysis of clausal nominalization developed in the context of the LinGO Grammar Matrix (Bender et al., 2002, 2010) to support the addition of subordinate clauses to the grammar customization framework. In particular, we examine the typological variation of nominalized clausal complements and nominalized clausal modifiers. To account for the range of variation in nominalized clauses across the world’s languages and to support linguists in exploring alternative analyses, we propose a flexible library of analyses, allowing nominalization of the clause to occur at the V, VP or S level.

1 Introduction

Languages differ in the range of means they provide for expressing embedded propositions (propositions that serve as a dependent of some predicate). One prominent strategy in the world’s languages is nominalization: a morphological or syntactic means of ‘wrapping’ a verbal constituent inside a nominal projection. This paper presents a cross-linguistic analysis of nominalized clauses in the context of a broader cross-linguistic grammar implementation project, namely the LinGO Grammar Matrix (Bender et al., 2002, 2010). The Grammar Matrix is a starter-kit for creating broad-coverage implemented precision grammars in HPSG (Pollard & Sag, 1994) which map between surface strings and Minimal Recursion Semantics (MRS; Copestake et al., 2005) representations. It includes a shared core grammar as well as a series of libraries extending that core with analyses for cross-linguistically variable phenomena. The analysis of nominalization presented here was developed in the context of our work on expressions of embedded propositions more generally, including as complements of verbs (Zamaraeva et al., to appear) and as modifiers of verbal projections (Howell & Zamaraeva, 2018). Typological surveys of these phenomena including Noonan 2007 show that clausal nominalization is a common strategy for embedded clauses, so we develop an analysis for nominalized clauses with these types of clausal subordination in mind.

As is typical for Grammar Matrix libraries, our analysis is intended to account for a broad range of typological possibilities as well as to give the user analytical freedom in modeling those possibilities. In particular, we allow for nominalization at different levels in the parse tree:

- Low: the nominal constituent is built out of a lexical verb (V)
- Mid: the nominal constituent is built out of a VP constituent comprising the verb and its complement
- High: the nominal constituent is built out of a full S: a verb plus all of its dependents

We also provide options on the semantic side, allowing high nominalization to be either strictly a syntactic phenomenon or one with semantic effects. A linguist

using the customization system can test alternative analyses in combination with analyses for other phenomena against text from their language to explore which best models the data.

We begin by describing in more detail the particular phenomena we are analyzing (§2) and briefly reviewing previous approaches (§3). We present our cross-linguistic analysis in §4, which includes the three levels of nominalization and two possible semantic representations. Finally we describe our implementation in the Grammar Matrix (§5) and how we evaluated the robustness of our analysis (§6). We conclude with a discussion of areas in which this work can be extended (§7).

2 Nominalized Subordinate Clauses

Nominalization is a common strategy for subordination in the world’s languages (Noonan, 2007). To illustrate the difference between a verbal clause and a nominalized clause, consider the following data from Rukai, which contrasts the non-finite verb *amo-dhaace* ‘leaving’ in (1) with the nominalized *to’a-dhaac-ae* ‘the reason for leaving’ in (2).

- (1) *amo-dhaace* = *lrao*
IRR-DYN.NFIN:leave = 1SG.NOM
‘I am leaving’ [dru] (adapted from Zeitoun 2007)
- (2) *to’a-dhaac-ae* = *li*
REAS.NMZ-DYN.NFIN:leave-REAS.NMZ = 1SG.GEN
ma-lrakas-iae
STAT.FIN-dislike-1SG.OBL
‘The reason why I’m leaving is because I dislike being here’ [dru] (adapted from Zeitoun 2007)

In contrast with the non-nominalized form *amo-dhaace*, the nominalized verb *to’a-dhaac-ae* is marked with a nominalization circumfix which is specific to reason adverbial clauses. It also co-occurs with a genitive (rather than nominative) subject clitic. Nominalization morphemes and case frame change are common markers of nominalized clauses cross-linguistically, as shown in the following examples from Uzbek (3) and Irish (4). In fact, the Irish example demonstrates that case frame change for nominalized verbs is possible on on objects as well.¹

- (3) *Xotin bu odam-niŋ joŋa-ni oŋirla-š-i-ni istandi*
woman this man-GEN chicken-OBJ steal-NMZ-3.SG-OBJ want.PST.3SG
‘The woman wanted the man to steal the chicken.’ [uzb] (adapted from Noonan 2007)

¹In this example the subject is not overt in the nominalized clause. The genitive NP is the object.

- (4) Is ionadh liom Seán a bhualadh Thomáis
 COP surprise with.me John COMP hit.NMZ Thomas.GEN
 ‘I’m surprised that John hit Thomas.’ [gle] (adapted from Noonan 2007)

The characteristics of nominalized clauses in examples (2)–(3) may reflect the level at which nominalization occurred. The following examples of English gerunds (adapted from Malouf 2000) suggest a hierarchy of nominalization types for increasingly nominal properties of the phrase’s internal distribution.

- (5) a. The DA was shocked that Pat illegally destroyed the evidence.
 b. The DA was shocked that she illegally destroyed the evidence.
 (6) a. The DA was shocked by Pat having illegally destroyed the evidence.
 b. The DA was shocked by her having illegally destroyed the evidence.
 (7) a. The DA was shocked by Pat’s having illegally destroyed the evidence.
 b. The DA was shocked by her having illegally destroyed the evidence.
 (8) a. The DA was shocked by Pat’s illegal destroying of the evidence.
 b. The DA was shocked by her illegal destroying of the evidence.
 (9) a. The DA was shocked by Pat’s illegal destruction of the evidence
 b. The DA was shocked by her illegal destruction of the evidence
 (adapted from Malouf 1998)

Malouf (1998) notes that (5) has no internal properties of an NP and is a fully verbal phrase: the *destroyed* is modified by an adverb and its subject’s and object’s case markings are consistent with those of English verbs. On the other hand, (9) has all of the properties of an NP and is a deverbal noun: *destruction* is modified with an adjective and its subject and object are both marked with different cases than those of the verb in (5).² The remaining examples illustrate the range between fully verbal and fully nominal expressions.

We take this variation in verbal and nominal properties as an indication of the level at which the verbal projection took on the properties of nominal projections, or put another way, at what level the clause was nominalized. In §4, we propose an analysis based on this observation, such that high nominalization (at S) allows adverbial modifiers and does not allow case change on subjects or objects; mid nominalization (at VP) allows adverbial modifiers and only allows case change on subjects; and low nominalization (at V) allows adjectival modifiers and case change on both subjects and objects.

3 Previous Approaches

Malouf (1998) provides a thorough review of previous approaches to clauses with both nominal and verbal characteristics. Here, we summarize his review as well as his own approach in order to situate our analysis within this body of work.

²Here we take *of* to be a kind of case-marking preposition.

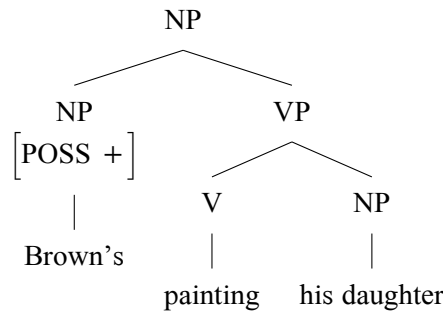


Figure 1: Pullum's approach (Malouf, 1998)

Pullum (1991) presents an analysis for English gerunds in which the VP headed by the verbal gerund combines with a possessive NP to form a larger NP constituent, the nominalized clause, as illustrated in figure 1. Lapointe (1993), on the other hand, takes a different approach, proposing a dual lexical category $\langle X|Y \rangle$ such that X determines the external distribution and Y the internal structure. Thus in the case of gerunds or nominalized clauses, the underlying lexical type would be $\langle N|V \rangle$. Malouf notes that neither approach accounts for gerunds with accusative subjects, e.g. *her having illegally destroyed the evidence* in (6) or adjective modification, as in *my wicked leaving my father's house*, as seen in old English. Furthermore, while Pullum's approach violates the principle of endocentricity by positing a head daughter which does not have the same distribution as the phrase, Lapointe's approach could generalize to other mixed categories that do not occur in the world's languages.

Bresnan (1997) proposes a 'change-over' approach, wherein the verbal constituent changes to a nominal constituent, as illustrated in figure 2. In doing so, the gerund will have the properties of a verb until the change over occurs, and then will take on the properties of a noun. Malouf notes that in addition to violating the principle of endocentricity like Pullum's approach, this analysis also doesn't correctly account for adverb position. In particular, the gerund is the daughter of NP, so an adverb would attach after the gerund, not before. This incorrectly predicts *Pat's watching avidly movies* and incorrectly rules out *Pat's avidly watching movies*.

Finally, Malouf (1998) presents a mixed category analysis, positing a gerund head value, modeled with multiple inheritance, as shown in the hierarchy in figure 3. This allows gerunds to interact with phrase structure rules sometimes like verbs and sometimes like nouns. He pairs this with a lexical rule that derives the valence properties of the gerunds and shows how a similar approach can work for a variety of languages, including English, Arabic, Boumaa Fijian, and Dagaare.

Malouf argues against 'change-over' approaches (e.g. that of Bresnan 1997, inter alia), because they don't constrain what kinds of change overs are available.

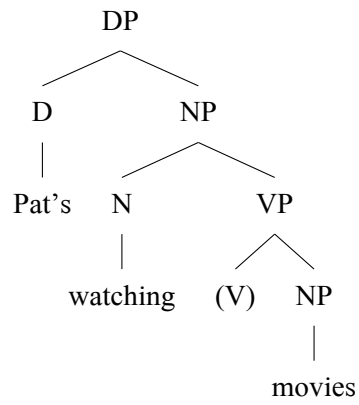


Figure 2: Bresnan's change-over approach (Malouf, 1998)

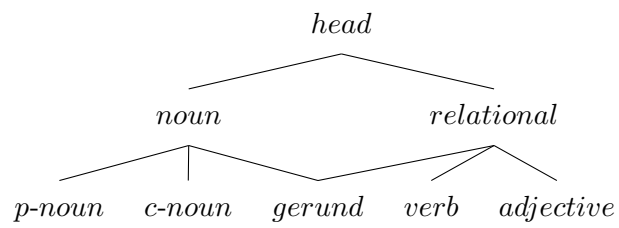


Figure 3: Malouf's multiple inheritance hierarchy (Malouf, 1998)

His mixed-category approach, combined with language-specific versions of the Head-Specifier and Head-Subject rules, elegantly accounts the mixed behavior of verbal gerunds. However, given the goal of grammar customization and the context of the Grammar Matrix code base, we take a change-over approach as it integrates more easily with the other libraries providing the phrase structure rules. On our analysis, the change over can happen at the S, VP or V levels. High nominalization (at S) allows adverbial modifiers and does not allow case change on subjects; mid nominalization (at VP, as in (7)) allows adverbial modifiers and only allows case change on subjects; and low nominalization (at V, as in (8)) allows adjectival modifiers and case change on both subjects and objects.³

In the next section, we present an analysis akin to that of Bresnan 1997, in that we take a change-over approach, using unary rules to transform verbal projections in to nominal projections. It differs in that it also includes lexical rules. Accordingly the change over of HEAD value need not correlate with the changes to constraints on arguments. This avoids some of the problems that Malouf (1998) finds with change-over approaches.⁴ Acknowledging that this approach violates the principle of endocentricity, our goal in the Grammar Matrix is to facilitate modeling grammars, rather than narrowing the class of possible languages. Furthermore, we find that this change-over approach allows us to account for case-frame changes as well as adjective and adverb attachment effectively, in order to model the range of nominalization strategies discussed in the previous section.

4 A Cross-linguistic Analysis

In this section we present three distinct analyses for nominalization to account for the variation described in §2. We begin by introducing the NMZ feature in §4.1. This is followed by a description of three analyses for high (§4.2), mid (§4.3) and low (§4.4) nominalization, which we motivate using the data presented in §2. We discuss the additional work necessary to accommodate case frame changes in §4.5 and propose two possible semantic representations of nominalized clauses in §4.6.

4.1 The NMZ feature

Our analysis allows for the disassociation of the nominalization morphology from the actual change of the HEAD value from *verb* to *noun*. To facilitate this, we propose a Boolean HEAD feature NMZ, which we use to distinguish verbs inflected with a nominalization morpheme (but not yet nominalized) from other verbs. We also use this feature to differentiate between nominal constituents built from nominalized verbs and other (lexical) nouns. Nouns and verbs in the lexicon are constrained to be [NMZ −] and this constraint is changed to [NMZ +] only by

³Neither (5) nor (9) involve nominalization of the type we are concerned with; the former because the constituent is verbal at all levels, and the latter because the clause has no verbal properties.

⁴We leave to future work the project of ensuring that our analysis can account for all the data presented in Malouf 1998.

nominalization lexical rules. The low nominalization analysis changes the HEAD value from *verb* to *noun* in the lexical rule. However, the mid and high nominalization analyses employ a unary rule to change the HEAD value and that unary rule has [NMZ +] on both the daughter and mother. These processes are illustrated in detail in figures 4–6 below. Under our analysis, complementizers, subordinators and clausal verbs that require nominalized clausal complements constrain their complement to be both [NMZ +] to prevent selection of a lexical noun and [HEAD *noun*] to prevent selection of a verb that has gone through the lexical rule, but not the corresponding unary rule.

4.2 High Nominalization

Our first nominalization analysis involves nominalization at the S level, such that the constituent maintains verbal properties until all arguments are picked up (including the subject) and only then is the nominal constituent built. We have not found clear evidence that this option is attested in the world’s languages: such evidence would involve a case language with nominalized clauses and no case change on the subject. Nevertheless, we provide this analysis as an option to linguists who may wish to test it against their data.

To accommodate clauses that remain verbal until all valence features are satisfied and then undergo nominalization, we posit two rules: a lexical rule that puts a morpho-syntactic marker on the verb and a unary phrase structure rule that builds a nominal constituent out of a verbal one. This is illustrated with the hypothetical example *Pat destroying the evidence*, where we pretend that Pat is a nominative subject (contrary to the facts of English), in figure 4.

The lexical rule is shared with the analysis for mid nominalization (§4.3), and accordingly is named *high-or-mid-nominalization-lex-rule*. This rule, defined in (10), adds [NMZ +] to the mother and identifies the INDEX of the daughter’s subject with the INDEX of the mother’s subject. We constrain only the subject’s INDEX in order to accommodate case change under the mid-nominalization analysis. However, for high nominalization, a sub-type of this rule identifies the entire subject between the mother and daughter.⁵

$$(10) \left[\begin{array}{l} \text{high-or-mid-nominalization-lex-rule} \\ \text{SYNSEM} \mid \text{LOCAL} \quad \left[\begin{array}{l} \text{CAT} \mid \text{HEAD} \quad [\text{NMZ} \quad +] \\ \text{VAL} \mid \text{SUBJ} \quad \langle \text{INDEX } \boxed{\alpha} \rangle \end{array} \right] \\ \text{DTR} \mid \text{SYNSEM} \mid \text{LOCAL} \quad \left[\text{CAT} \mid \text{VAL} \mid \text{SUBJ} \quad \langle \text{INDEX } \boxed{\alpha} \rangle \right] \end{array} \right]$$

Once the morpho-syntactic marker NMZ + has been added to the verb and its

⁵The AVMs shown in this paper are abbreviated in order to focus on features of interest. The lexical rules produced by the Grammar Matrix customization system also have many constraints that serve to copy information from daughter to mother. The reader can assume that all features are copied from daughter to mother unless otherwise specified. Grammars that exemplify these constraints can be checked out from revision 41825 here: svn://lemur.ling.washington.edu/shared/matrix/trunk/gmcs/regressiontests

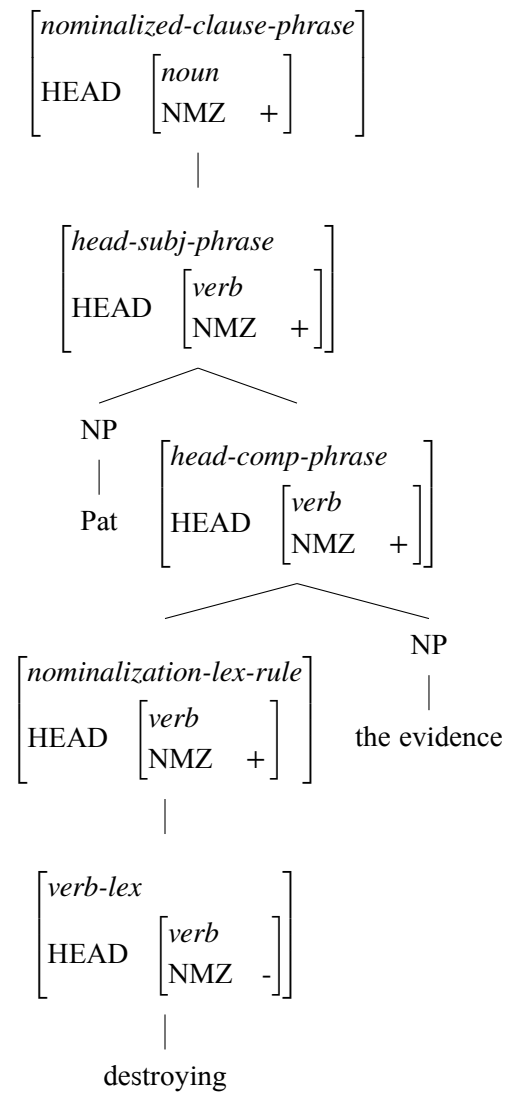


Figure 4: High nominalization

valence requirements have been satisfied, the clause can serve as the daughter of the *nominalized-clause-phrase* unary rule, defined in (11).

$$(11) \left[\begin{array}{l} \text{nominalized-clause-phrase} \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{l} \text{CAT} \mid \text{HEAD} \left[\begin{array}{l} \textit{noun} \\ \text{NMZ} \quad + \end{array} \right] \\ \text{VAL} \left[\begin{array}{l} \text{COMPS} \quad \langle \rangle \\ \text{SUBJ} \quad \langle \rangle \end{array} \right] \end{array} \right] \\ \text{ARGS} \left\langle \begin{array}{l} \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \textit{verb} \\ \text{NMZ} \quad + \end{array} \right] \\ \text{VAL} \left[\begin{array}{l} \text{COMPS} \quad \langle \rangle \\ \text{SUBJ} \quad \langle \rangle \end{array} \right] \end{array} \right] \\ \text{CONT} \left[\text{HOOK} \mid \text{LTOP} \quad \textcircled{0} \end{array} \right] \end{array} \right] \end{array} \right\rangle \\ \text{C-CONT} \left[\begin{array}{l} \text{RELS} \left\langle \begin{array}{l} \text{PRED} \quad \textit{nominalization_rel} \\ \text{LBL} \quad \textcircled{1} \\ \text{ARG0} \quad \textcircled{2} \\ \text{ARG1} \quad \textcircled{3} \end{array} \right\rangle \\ \text{HCONS} \left\langle \begin{array}{l} \textit{qeq} \\ \text{HARG} \quad \textcircled{3} \\ \text{LARG} \quad \textcircled{0} \end{array} \right\rangle, \left\langle \begin{array}{l} \textit{qeq} \\ \text{HARG} \quad \textcircled{2} \\ \text{LARG} \quad \textcircled{1} \end{array} \right\rangle \end{array} \right] \end{array} \right]$$

We constrain both the SUBJ and COMPS lists to be empty on the mother and daughter, so that this rule will only select clauses which are valence saturated. This rule effects the syntactic change from verbal to nominal projection, changing the HEAD type to *noun*. The unary rule also adds the necessary semantic constraints for the nominalized verb to be represented as a noun. This is accomplished by adding *nominalization_rel* to the C-CONT (constructional-content) list and linking the ARG1 of that predication to the daughter's LTOP.⁶ This has the effect of ‘wrapping’ a nominal predication around the proposition built by the verb. The resulting MRS representation will be discussed in more detail in §4.6.

4.3 Mid Nominalization

Our next analysis involves the nominalization of verb phrases, i.e. verbal projections with empty COMPS lists. This analysis is motivated by examples such as (6) and (7), repeated here as (12) and (13).

- (12) a. The DA was shocked by Pat having illegally destroyed the evidence.
b. The DA was shocked by her having illegally destroyed the evidence.
- (13) a. The DA was shocked by Pat's having illegally destroyed the evidence.

⁶This connection is mediated by an ‘equal modulo quantifiers’ constraint (*qeq*) given in the value of HCONS. These constraints are part of the MRS analysis of quantifier scope ambiguity (Copestake et al., 2005) and introducing one here allows quantifiers in the nominalized clause to have the option of scoping below the embedding predicate, as desired.

- b. The DA was shocked by her having illegally destroyed the evidence.

These examples exhibit hybrid properties: In (12) and (13) the verb is modified by an adverb and its complement bears its canonical case, i.e. within the VP constituent we see verbal properties. However, the subject appears with a non-canonical case, genitive or accusative.

Our mid nominalization analysis is very similar to the high nominalization analysis in that a morpho-syntactic marker is added by the *high-or-mid-nominalization-lex-rule* and the projection is changed from verbal to nominal by a unary rule higher in the tree, as illustrated in figure 5.

The lexical rule in (10) is also used for mid nominalization. As discussed in the previous section, this rule only identifies the INDEX of the subject, allowing the case value of the subject to be changed.⁷ This process is described in more detail in §4.5. This analysis also uses a unary rule change the projection from verbal to nominal. The *mid-nominalized-clause-phrase* rule in (14) differs from the rule in (11) in only one way: instead of an empty subject list, the subject list of the daughter is constrained to be non-empty and identified with the the subject list of the mother.

$$(14) \left[\begin{array}{l} \text{mid-nominalized-clause-phrase} \\ \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{l} \text{CAT} \mid \text{HEAD} \left[\begin{array}{l} \text{noun} \\ \text{NMZ} \quad + \end{array} \right] \\ \text{VAL} \left[\begin{array}{l} \text{COMPS} \quad \langle \rangle \\ \text{SUBJ} \quad \boxed{0} \end{array} \right] \end{array} \right] \\ \text{ARGS} \left\langle \begin{array}{l} \text{SYNSEM} \mid \text{LOCAL} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \text{verb} \\ \text{NMZ} \quad + \end{array} \right] \\ \text{VAL} \left[\begin{array}{l} \text{COMPS} \quad \langle \rangle \\ \text{SUBJ} \quad \boxed{0} \end{array} \right] \end{array} \right] \\ \text{CONT} \left[\text{HOOK} \mid \text{LTOP} \quad \boxed{1} \end{array} \right] \end{array} \right\rangle \\ \text{C-CONT} \left[\begin{array}{l} \text{RELS} \left\langle \begin{array}{l} \text{!} \left[\begin{array}{l} \text{PRED} \quad \text{nominalization_rel} \\ \text{LBL} \quad \boxed{2} \\ \text{ARG0} \quad \boxed{3} \\ \text{ARG1} \quad \boxed{4} \end{array} \right] \text{!} \end{array} \right\rangle \\ \text{HCONS} \left\langle \begin{array}{l} \text{!} \left[\begin{array}{l} \text{qeq} \\ \text{HARG} \quad \boxed{4} \\ \text{LARG} \quad \boxed{1} \end{array} \right], \left[\begin{array}{l} \text{qeq} \\ \text{HARG} \quad \boxed{3} \\ \text{LARG} \quad \boxed{2} \end{array} \right] \text{!} \end{array} \right\rangle \end{array} \right] \end{array} \right]$$

⁷Under our analysis mid nominalization without case change is allowed. While it is typologically unlikely that a language would have VP nominalization without case change on the subject (hypothetically exemplified by an adjective modifier above VP but below the subject), it is possible that a user developing a grammar for a language without a case system would want to avoid adding the additional case-change-related constraints to their grammar.

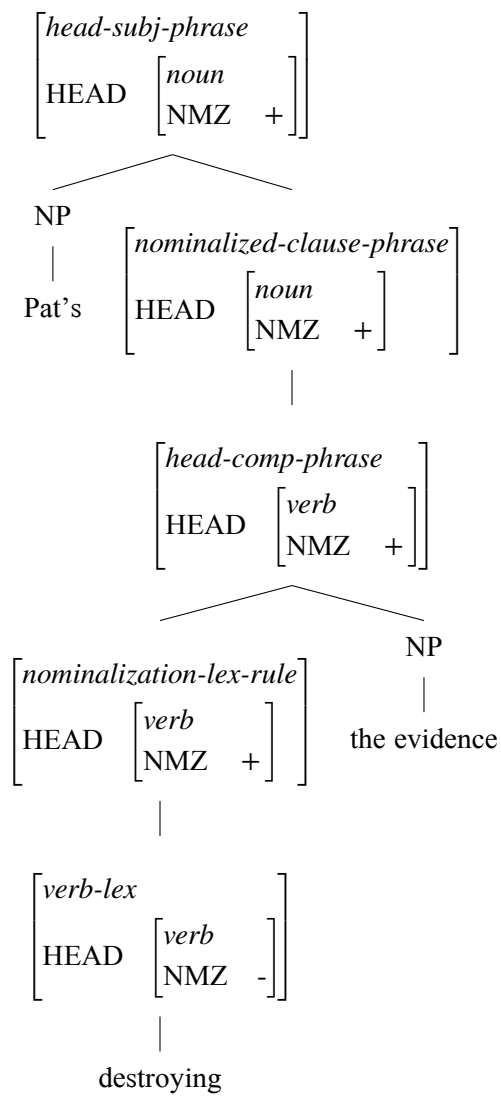


Figure 5: Mid nominalization

4.4 Low Nominalization

Our final analysis involves nominalization at the lexical level, before the underlying verb combines with any of its arguments. Although this analysis for nominalization occurs on the lexical level, we do not claim that it extends to all deverbal nouns. In particular, it is only appropriate for productive morphology which furthermore results in event nominalization (as opposed to e.g. agent nominalization). This analysis is appropriate for examples where the nominalized verb is modified by a low-attaching adjective and/or the case on the verb's complement differs from that found in its ordinary (non-nominalized) use. (8), repeated here as (15), falls into this category:

- (15) a. The DA was shocked by Pat's illegal destroying of the evidence.
 b. The DA was shocked by her illegal destroying of the evidence.

It may be that low nominalization is also motivated by changes to the CASE or HEAD value required of the complement. Under our analysis, these are actually always handled low (in the lexical rule), but linguists may prefer to analyze them as co-incident with the change of the HEAD and INDEX on the nominalized constituent itself.

Under our analysis of low nominalization, the lexical rule that provides the nominalization morpheme and the morpho-syntactic marker also directly changes the verb to a noun, as illustrated in figure 6. This rule, shown in (16), specifies [HEAD *noun*] and [NMZ +] on the mother. The lexical rule also adds the predication nominalization_rel to the MRS and links its first argument with the daughter (via a *qeq* constraint).

$$(16) \left[\begin{array}{l} \text{low-nominalization-lex-rule} \\ \\ \text{SYNSEM} \mid \text{LOCAL} \\ \\ \text{DTR} \mid \text{SYNSEM} \mid \text{LOCAL} \\ \\ \text{C-CONT} \end{array} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{HEAD} \left[\begin{array}{l} \textit{noun} \\ \text{NMZ} \mid + \end{array} \right] \\ \text{VAL} \mid \text{SUBJ} \langle \text{INDEX} \mid \boxed{0} \rangle \end{array} \right] \\ \\ \text{CAT} \left[\begin{array}{l} \text{VAL} \mid \text{SUBJ} \langle \text{INDEX} \mid \boxed{0} \rangle \\ \text{CONT} \left[\text{HOOK} \mid \text{LTOP} \mid \boxed{1} \right] \end{array} \right] \\ \\ \text{RELS} \left\langle ! \left[\begin{array}{l} \text{PRED} \mid \text{nominalization_rel} \\ \text{LBL} \mid \boxed{2} \\ \text{ARG0} \mid \boxed{3} \\ \text{ARG1} \mid \boxed{4} \end{array} \right] ! \right\rangle \\ \\ \text{HCONS} \left\langle ! \left[\begin{array}{l} \textit{qeq} \\ \text{HARG} \mid \boxed{4} \\ \text{LARG} \mid \boxed{1} \end{array} \right], \left[\begin{array}{l} \textit{qeq} \\ \text{HARG} \mid \boxed{3} \\ \text{LARG} \mid \boxed{2} \end{array} \right] ! \right\rangle \end{array} \right] \right]$$

The lexical rule in (16) is a somewhat underspecified supertype that is further constrained depending on the specifications given by a user for a particular

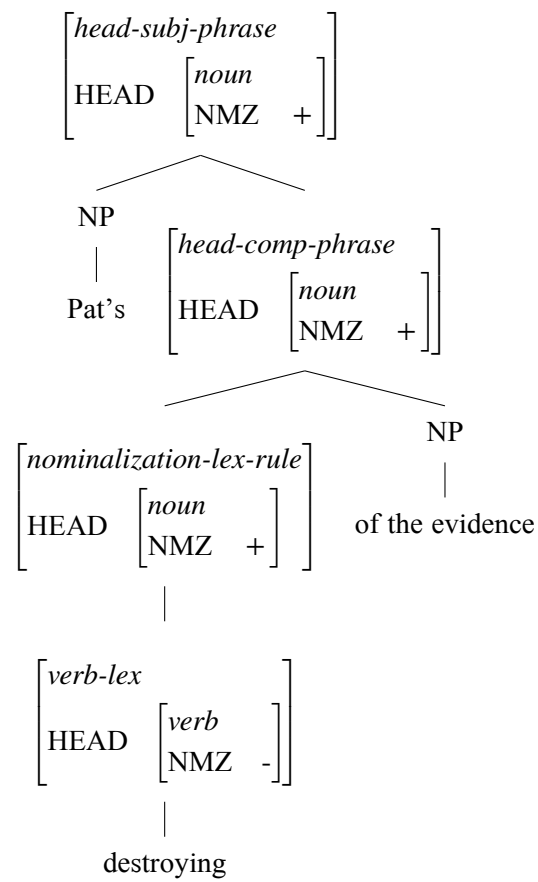


Figure 6: Low nominalization

language. It identifies the INDEX of the mother's subject with the INDEX of the daughter's subject. If the case on the subject changes upon nominalization, this constraint is sufficient (in combination with constraints on case discussed in §4.5 below). However, if case frame change does not occur, then we create a subtype of this rule that identifies the entire subject, rather than just the INDEX. Similarly, we add constraints to subtypes of this rule based on whether or not the object's case is changed. If the case on the object changes, a constraint to identify the complement's INDEX⁸ between mother and daughter is added, whereas if the object's case does not change or the verb is intransitive, the entire complements list is identified between daughter and mother.⁹

4.5 Accommodating Case Frame Changes

In §4.3 and §4.4 we noted that the nominalization lexical rule supertypes only identify the indices of subjects and complements and that work remains to be done if the case frame of the nominalized verb differs from that of a non-nominalized verb. Subtypes of these rules are used to make changes to the HEAD features, including both the case and the associated head type.

In particular, when a user of the Grammar Matrix defines a morphological rule associated with nominalization, they may also indicate the case of the subject and/or object if they differ from the standard verbal case-frame. These case constraints are then added to the nominalization lexical rules by the Grammar Matrix customization system. Because certain cases may be associated with particular HEAD types in the language, the customization system has built-in functions for detecting the head types that are compatible with given case. We use these functions to identify the appropriate HEAD type and add that constraint to the lexical rule as well. Thus a hypothetical language in which nominalized verbs require genitive subjects and genitive case is marked by a preposition would have the following rule, inheriting from the *low-nominalization-lex-rule*.

$$(17) \left[\begin{array}{l} \text{low-intransitive-nominalization-lex-rule} \\ \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{VAL} \left[\begin{array}{l} \text{SUBJ} \quad \langle \left[\begin{array}{l} \text{INDEX } \boxed{0} \\ \text{HEAD } \left[\begin{array}{l} \text{prep} \\ \text{CASE gen} \end{array} \end{array} \right] \rangle \\ \text{COMPS} \quad \langle \rangle \end{array} \right] \end{array} \right] \\ \text{DTR} \mid \text{SYNSEM} \mid \text{LOCAL} \mid \text{CAT} \mid \text{VAL} \left[\begin{array}{l} \text{SUBJ} \quad \langle \text{INDEX } \boxed{0} \rangle \\ \text{COMPS} \quad \langle \rangle \end{array} \right] \end{array} \right]$$

⁸Currently ditransitive verbs are not supported by the Grammar Matrix, so our analysis only accounts for one complement.

⁹For languages with case change on the object, we use two separate rules, one for transitive verbs which identifies the object's INDEX and one for intransitive verbs that identifies the entire COMPS list.

4.6 Semantic Representations

We provide two possible representations for nominalized clauses, using Minimal Recursion Semantics (MRS; Copestake et al., 2005). On the one hand, in many languages it can be argued that a nominalized subordinate clause has a different meaning than a fully verbal subordinate clause. At the very least, there must be a nominal predication in the semantic representation to which adjectives, like that in (8), repeated here as (18), and quantifiers can attach.

- (18) a. The DA was shocked by Pat’s illegal destroying of the evidence.
b. The DA was shocked by her illegal destroying of the evidence.

On the other hand, a linguist modeling a language in which nominalization is the only strategy for subordination might argue that there is no difference in meaning between nominalized subordinate clauses in that language and subordinate clauses in other languages. Therefore, we provide both options for our high nominalization analysis: one with a `nominalization_rel` and one without. At this time we do not allow a representation without a `nominalization_rel` for low and mid nominalization as this would prevent adjective modification of those clauses. This option may be appropriate to add, but only in languages which never allow adjectival modification of the low or mid nominalized structures.

For an example like the Turkish sentence in (19) with a nominalized clausal complement, the analyses described earlier in this section result in the MRS semantic representation in (20).¹⁰

- (19) *senin sinema-ya gel-me-n-i isti-yor-um*
2SG.GEN cinema-DAT come-NMZ-2SG-ACC want-PROG-1SG
“I want you to come to the movies.” [tur] adapted from Kornfilt (1997, p. 48)

$$(20) \quad \langle h_1, e_2, \left[\begin{array}{l} h_3:\text{pron_rel}(\text{ARG0 } x_4), \\ h_5:\text{exist_q_rel}(\text{ARG0 } x_4, \text{RSTR } h_6, \text{BODY } h_7), \\ h_8:\text{_cinema_n_rel}(\text{ARG0 } x_9), \\ h_{10}:\text{exist_q_rel}(\text{ARG0 } x_9, \text{RSTR } h_{11}, \text{BODY } h_{12}), \\ h_{13}:\text{come_v_rel}(\text{ARG0 } e_1, \text{ARG1 } x_4, \text{ARG2 } x_9), \\ h_{15}:\text{nominalization_rel}(\text{ARG0 } x_{17}, \text{ARG1 } h_{16}), \\ h_{18}:\text{exist_q_rel}(\text{ARG0 } x_{17}, \text{RSTR } h_{19}, \text{BODY } h_{20}), \\ h_{23}:\text{pron_rel}(\text{ARG0 } x_{22}), \\ h_{24}:\text{exist_q_rel}(\text{ARG0 } x_{22}, \text{RSTR } h_{25}, \text{BODY } h_{26}), \\ h_{21}:\text{want_v_rel}(\text{ARG0 } e_2, \text{ARG1 } x_{22}, \text{ARG2 } x_{17}) \end{array} \right] \{ h_6 =_q h_3, h_{11} =_q h_8, h_{16} =_q h_{13}, h_{19} =_q h_{15}, h_{25} =_q h_{23} \} \rangle$$

¹⁰Note that while Turkish this is not an example of high nominalization, all three analyses presented in this section produce the same semantic representation.

This semantic structure contains the predication `nominalization_rel` and the verb is the first argument of this predication.¹¹ The intrinsic argument (`ARG0`) of the `nominalization_rel` is of type x for individual, rather than e for event, so as to be a suitable argument for adjectival modifiers and bound variable for quantifiers. Because the low and mid analysis allow for the attachment of adjectives and quantifiers syntactically, this must be accounted for in the semantics as well.

However, we provide the user with analytical freedom regarding the semantic structure, by developing an option for nominalization that is purely syntactic. In this case the unary rule changes the `HEAD` value to `noun` and creates a direct semantic identity between the mother and daughter without adding `nominalization_rel`, resulting in MRSs like the one shown in (21).¹²

$$(21) \quad \left\langle h_1, e_2, \begin{array}{l} h_3:\text{pron_rel}(\text{ARG0 } x_4), \\ h_5:\text{exist_q_rel}(\text{ARG0 } x_4, \text{RSTR } h_6, \text{BODY } h_7), \\ h_8:\text{cinema_n_rel}(\text{ARG0 } x_9), \\ h_{10}:\text{exist_q_rel}(\text{ARG0 } x_9, \text{RSTR } h_{11}, \text{BODY } h_{12}), \\ h_{13}:\text{come_v_rel}(\text{ARG0 } e_1, \text{ARG1 } x_4, \text{ARG2 } x_9), \\ h_{23}:\text{pron_rel}(\text{ARG0 } x_{22}), \\ h_{24}:\text{exist_q_rel}(\text{ARG0 } x_{22}, \text{RSTR } h_{25}, \text{BODY } h_{26}), \\ h_{21}:\text{want_v_rel}(\text{ARG0 } e_2, \text{ARG1 } x_{22}, \text{ARG2 } h_{27}) \end{array} \right. \\ \left. \{ h_6 =_q h_3, h_{11} =_q h_8, h_{16} =_q h_{13}, h_{25} =_q h_{23}, h_{27} =_q h_{13} \} \right\rangle$$

4.7 Summary

This section has presented our cross-linguistic analysis of nominalization. As is typical for Grammar Matrix libraries, the analysis encompasses a range of options. These options accommodate both cross-linguistic variation in the underlying phenomenon and analytic variation, facilitating the exploration of different analyses within implemented grammars.

5 Implementation in the Grammar Matrix

We implemented the analyses described in §4 in the Grammar Matrix, such that the user can define multiple nominalization strategies that can be accessed by the subordinate clause libraries, including Clausal Complements (Zamaraeva et al., to appear) and Clausal Modifiers (Howell & Zamaraeva, 2018). The user can give each nominalization strategy a name and select the level and desired semantic representation for that strategy. This strategy can then be associated with morphological rules (corresponding to nominalization affixes) and clausal complement and

¹¹This relationship is mediated by a so-called *qeq* constraint. See note 6.

¹²As Turkish does not in fact have high nominalization, this MRS would not be produced for Turkish. We provide it here for comparison with the one in (20) only.

clausal modifier strategies that require nominalization. The relevant portion of the Grammar Matrix web questionnaire is illustrated by figure 7.

Nominalized Clauses [\[documentation\]](#)

If your language uses nominalization in the context of clausal complements and/or clausal modifiers, define the nominalization strategies here. They will then be available on the Clausal Complements, Clausal Modifiers, and Morphology pages.

▼ ns1

☐ **Nominalization Strategy 1:**

Nominalization Strategy Name:

The nominalization of the clause happens:

☐ at V ☐ at VP ☐ at S

Is the nominalization syntactic only or should it also be reflected in the semantics?
 (Note: for mid or low nominalization, currently you must say that it is reflected in the semantics).

☐ Nominalization is syntactic only
☐ Nominalization should be reflected in the semantics

[Add a Nominalization Strategy](#)

Figure 7: Snippet of Grammar Matrix questionnaire for nominalization library

6 Testing, Evaluation, and Error Analysis

Following typical practice in the development of Grammar Matrix libraries (Bender et al., 2010), we evaluated our implementation of this analysis by creating grammar fragments for a number of languages. This allows us to verify both that the analyses generalize to languages we didn’t directly consider during library development and that the analyses in the library interact appropriately with other libraries.

We do initial verification using both artificial ‘pseudolanguages’ designed to test each combination of nominalization level and semantic representation and real languages. In both cases, we first develop testsuites including grammatical and ungrammatical examples, and then create choices files describing those languages. We feed the choices files to the Grammar Matrix customization system and use the resulting grammars to parse the testsuites using the LKB software (Copestake, 2002). Undergeneration, overgeneration, spurious ambiguity, or incorrect parses of testsuite items will indicate errors in the analysis or its implementation, which we fix during the development process.

We developed pseudolanguage choices files and testsuites for each level of nominalization and each semantic representation for both nominalized clausal complements and clausal modifiers, resulting in a total of 8 pseudolanguages. For our

real language verification tests, we used Russian [rus] and Turkish [tur] for clausal complements, and Rukai [dru] for clausal modifiers. We refined our implementation until we achieved full coverage (all grammatical sentences correctly parsed) and no overgeneration (no ungrammatical sentences parsed) over the development testsuites. While the 8 pseudolanguage testsuites were targeted at nominalization, the real language testsuites contained examples for clausal complements or clausal modifiers in general, so not all examples were relevant to nominalization. The following table identifies both the overall results and those relevant specifically to nominalization.¹³

Language	Total		Nominalized Clause	
	Coverage	Overgen.	Coverage	Overgen.
Russian [rus]	6/6	0/11	6/6	0/11
Turkish [tur]	7/7	0/9	6/6	0/8
Rukai [dru]	2/2	8/8	2/2	8/8

Table 1: Results for development languages¹⁴

Finally, we tested our analysis on languages that we had not previously considered in order to evaluate how well it generalizes cross-linguistically. We consider evaluation to be extrinsic as it was evaluated as part of our evaluation for clausal complements (Zamaraeva et al., to appear) and clausal modifiers (Howell & Zamaraeva, 2018).¹⁵ We evaluated our analysis in complement clauses in Yakima Sahaptin [yak] and Paresi-Haliti [pab], as well as in clausal modifiers in Basque [eus]. The results are presented in Table 2, again differentiating between the total number of examples and just those relevant to nominalization.

Language	Total		Nominalized Clause	
	Coverage	Overgen.	Coverage	Overgen.
Paresi-Haliti [pab]	5/5	0/6	3/3	0/4
Yakima Sahaptin [yak]	10/10	0/6	10/10	0/6
Basque [eus]	13/16	0/10	5/8	0/3

Table 2: Results for held-out languages¹⁶

The error analysis revealed one error (affecting three sentences in the test-suite for Basque), which was not directly related to the analysis presented in this

¹³We define “relevant” here as examples either containing a nominalized verb, or negative examples that are ungrammatical because they lack a nominalized verb.

¹⁴Russian, Turkic and Rukai are from the Indo-European, Altaic and Austronesian language families, respectively.

¹⁵More detailed discussion of the evaluation for those libraries beyond that which is relevant to nominalized clauses can be found in their respective papers.

¹⁶Paresi-Haliti, Yakima Sahaptin and Basque are from the Arawaken, Penutian, and Basque language families, respectively.

paper, but revealed an interaction with another analysis stored in the Grammar Matrix. The Argument Optionality library (Saleem, 2010) adds phrase structure rules to grammar fragments that facilitate argument dropping. As this library was created before nominalized verbs were supported, these rules constrained the head-daughter to be [HEAD *verb*], thereby prohibiting subject dropping for nominalized verbs in Basque. We were able to confirm that these sentences would otherwise parse by adding a subject dropping rule to the grammar that allowed a nominalized verb to be the head daughter.

7 Conclusion

In this paper we present a cross-linguistic analysis of nominalization, designed to support analyses of this phenomenon as it appears in both clausal complements and clausal modifiers. The analysis is implemented in the form of a Grammar Matrix library and its interoperability with libraries for not just clausal complements and clausal modifiers but also other libraries including argument optionality, case, and word order is tested according to the standard Grammar Matrix evaluation methodology. We provided an analysis that allows nominalization to occur at three different levels in the syntax and provided two semantic representations. We plan to look at a wider range of languages as part of future work to determine the usefulness of the high nominalization analysis. We are also considering extending the option to omit nominalization from the semantics to the mid and low analyses, if we find evidence to do so. Our evaluation so far suggests that our analysis provides sufficient flexibility to handle both the typologically attested range of variation in this phenomenon and to provide a degree of analytical freedom to the linguist, while still maintaining comparability across language types.

Acknowledgments

We thank Edith Aldridge and the audience at HPSG 2018 for helpful comments and discussion. This material is based upon work supported by the National Science Foundation under Grant No. BCS-1561833. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

References

- Bender, Emily M., Scott Drellishak, Antske Fokkens, Laurie Poulson & Safiyyah Saleem. 2010. Grammar customization. *Research on Language & Computation* 8. 1–50.
- Bender, Emily M., Dan Flickinger & Stephan Oepen. 2002. The Grammar Matrix: An open-source starter-kit for the rapid development of cross-linguistically con-

- sistent broad-coverage precision grammars. In John Carroll, Nelleke Oostdijk & Richard Sutcliffe (eds.), *Proceedings of the workshop on grammar engineering and evaluation at the 19th International Conference on Computational Linguistics*, 8–14. Taipei, Taiwan.
- Bresnan, Joan. 1997. Mixed categories as head sharing constructions. In Miriam Butt & Tracy Holloway King (eds.), *Proceedings of the LFG'97 Conference*, Stanford: CSLI Publications.
- Copestake, Ann. 2002. Definitions of typed feature structures. In Stephan Oepen, Dan Flickinger, Jun-ichi Tsujii & Hans Uszkoreit (eds.), *Collaborative language engineering*, 227–230. Stanford, CA: CSLI Publications.
- Copestake, Ann, Dan Flickinger, Carl Pollard & Ivan A Sag. 2005. Minimal recursion semantics: An introduction. *Research on Language and Computation* 3(2-3). 281–332.
- Howell, Kristen & Olga Zamaraeva. 2018. Clausal modifiers in the Grammar Matrix. In *Proceedings of the 27th International Conference on Computational Linguistics*, 2939–2952.
- Kornfilt, Jaklin. 1997. *Turkish*. London: Routledge.
- Lapointe, Steven. 1993. Dual lexical categories and the syntax of mixed category phrases. In *Proceedings of the 10th Eastern States Conference of Linguistics*, 199–210. CLC Publications Ithaca, NY.
- Malouf, Robert P. 1998. *Mixed categories in the hierarchical lexicon*. Palo Alto, CA: Stanford University dissertation.
- Noonan, Michael. 2007. Complementation. In Timothy Shopen (ed.), *Language typology and syntactic description. volume 2: Complex constructions*, 52–150. Cambridge University Press.
- Pollard, Carl & Ivan A Sag. 1994. *Head-Driven Phrase Structure Grammar*. University of Chicago Press.
- Pullum, Geoffrey K. 1991. English nominal gerund phrases as noun phrases with verb-phrase heads. *Linguistics* 29(5). 763–800.
- Saleem, Safiyyah. 2010. *Argument optionality: A new library for the Grammar Matrix customization system*. University of Washington MA thesis.
- Zamaraeva, Olga, Kristen Howell & Emily M. Bender. to appear. Modeling clausal complementation for a grammar engineering resource. To appear in *Proceedings of the 2nd meeting of the Society for Computation in Linguistics*.
- Zeitoun, Elizabeth. 2007. *A grammar of Mantauran (Rukai)*. Institute of Linguistics, Academia Sinica Taipei.

Plural in Lexical Resource Semantics

David Lahm

University of Frankfurt

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 87–107

Keywords: plural, lexical resource semantics, cumulation, cumulative predication, polyadic quantification, underspecified semantics, semantic glue, maximalisation, covers

Lahm, David. 2018. Plural in Lexical Resource Semantics. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 87–107. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.6.

 CC-BY 4.0

Abstract

The paper shows how the plural semantic ideas of (Sternefeld, 1998) can be captured in Lexical Resource Semantics, a system of underspecified semantics. It is argued that Sternefeld’s original approach, which allows for the unrestricted insertion of pluralisation into Logical Form, suffers from a problem originally pointed out by Lasersohn (1989) with respect to the analysis offered by Gillon (1987). The problem is shown to stem from repeated pluralisation of the same verbal argument and to be amenable to a simple solution in the proposed lexical analysis, which allows for restricting the pluralisations that can be inserted. The paper further develops an account of maximalisation of pluralities as needed to obtain the correct readings for sentences with quantifiers that are not upward monotone. Such an account is absent in the original system in (Sternefeld, 1998). The present account makes crucial use of the possibility to have distinct constituents contribute identical semantic material offered by LRS and employs it in an analysis of maximalisation in terms of polyadic quantification

1 Introduction

We propose a treatment of plural semantics in Lexical Resource Semantics (LRS) (Richter & Sailer, 2004; Kallmeyer & Richter, 2007) by developing a lexical implementation of the analysis proposed by Sternefeld (1998). Sternefeld (1998) proposes to treat pluralisation as semantic glue freely insertible into logical forms, an approach that will be referred to as *Augmented Logical Form* (ALF). ALF allows for straightforward derivations of a wide range of conceivable sentence meanings.

An approach that allows for freely inserting semantic material in the derivation of a sentence is *prima facie* at odds with a basic tenet of LRS, namely that every part of an utterance’s meaning must be contributed by some lexical element in that utterance. But the combinatory system of LRS will be seen to be flexible enough to achieve very similar results by purely lexical means.

The resulting approach will then be seen to allow for a straightforward solution of an overgeneration problem of ALF. The approach predicts meanings for certain sentences that Lasersohn (1989) discusses as problems for the approach of (Gillon, 1987). Gillon’s approach predicts sentence (1) to be true if each of three TAs received \$7,000.

- (1) The TAs were paid exactly \$14,000.

The same prediction is made by the system developed by Sternefeld (1998). Its reformulation in LRS will however allow for the formulation of a straightforward lexical constraint that rules it out. All that is required is to rule out more than one

[†]The research reported here was funded by the German Research Foundation (DFG) under grant GRK 2016/1. I thank Frank Richter, Cécile Meier, Manfred Sailer, Sascha Bargmann and Andy Lücking and three anonymous reviewers for valuable comments and discussion. I thank my colleagues at HeBIS IT for valuable support and encouragement.

pluralisation of the same argument position of any verb. In the present proposal this can easily be achieved lexically, while implementing the same idea in ALF should require constraints on logical forms of a highly non-local nature.

The paper furthermore shows how Sternefeld’s original proposal can be extended with the maximalisation operations needed to deal with other than upward-monotone quantifiers. Sternefeld’s original system predicts the equivalence of *exactly three girls slept* and *at least three girls slept*, which can be overcome by demanding the existence of a set of girls who slept that is a maximal set of girls who slept. Since maximalisations of pluralities introduced by different NPs should sometimes not take scope with respect to one another an analysis is proposed that harnesses the ability of LRS to fuse meaning components of different constituents into a single polyadic quantifier in order to prevent this unwanted result.

The paper will proceed as follows: Section 2 gives an introduction to Sternefeld’s analysis. Section 3 introduces the problematic data, which are shown to be actual problems for the ALF account in section 4. Section 5 develops the basic LRS analysis and section 6 extends it with maximalisation operations. Section 8 concludes the paper.

2 Cumulative Predication and Augmented Logical Form

2.1 The System of Cumulative Predication

Sternefeld (1998) employs the pluralisation operations $*$, familiar from the work of Link, and $**$. The definitions are given in (2).¹

$$(2) \quad \begin{array}{ll} \text{a.} & *S = \{\bigcup X \mid X \in \mathcal{P}(S) \setminus \{\emptyset\}\} \\ \text{b.} & **R = \left\{ \langle X, Y \rangle \mid \begin{array}{l} X = \bigcup \{U \subseteq X \mid \exists V \subseteq Y : R(U, V)\} \wedge \\ Y = \bigcup \{V \subseteq Y \mid \exists U \subseteq X : R(U, V)\} \end{array} \right\} \end{array}$$

Basic decisions underlying the system and adhered to in the present paper are that pluralities are represented as sets of non-empty subsets of the universe of discourse D (i.e. subsets of $*D$) and individual *urlements* are counted as pluralities by assuming $\{x\} = x$ for $x \in D$. There is a distinction between sets in the sense of elements of $*D$ and expressions of type $\langle e, t \rangle$. In particular, all elements of $*D$ have type e . (These assumptions are identical with those in (Schwarzschild, 1996)).

According to (2-a), $*S$ is the set of all unions over sets of subsets of S . Given the equality $\{x\} = x$, $*S = \mathcal{P}(S) \setminus \{\emptyset\}$ if S is a set of individuals, i.e. if it does not contain any non-singleton sets. So $*\text{sleep}$, for instance, is the set of all nonempty sets of sleepers, given that *sleep* contains only individuals. $\llbracket \text{the children} \rrbracket \in *\text{sleep}$ thus expresses that the children slept, i.e. are a set of sleepers.

But $*$ also works for sets of non-singleton sets, which are taken to be the extensions of collective verbs like *gather*. The extension of *gather* consists of sets

¹These original definitions were recursive and did not play well with infinite sets. The present definitions are suitable for the general case.

of individuals that gathered as a single group; different members of the verb extension represent different groups. Then $\llbracket \text{the children} \rrbracket \in \text{*gather}$ expresses that the children gathered, but without any presumption that there was only one group, because $\llbracket \text{the children} \rrbracket$ may be the union of an arbitrary set of such groups. By leaving the pluralisation out, a reading can be expressed that expresses that the children gathered as a group. $\llbracket \text{the children} \rrbracket \in \text{gather}$.

The somewhat involved definition of * expresses that $\text{*}R(X, Y)$ holds if (i) the subsets U of X such that $R(U, V)$ holds for some subset V of Y cover X , i.e. every member of X is present in some U , and (ii) the same holds *vice versa*. So every member of X belongs to a subset of X that is related by R to some subset of Y , and *vice versa*.

Sentences of the kind of (3-a), as discussed in (Scha, 1981), can now conveniently be represented as in (3-b). Each of the 500 firms belongs to a (probably singleton) subset of the set of firms that is related to a subset of the 2000 computers, and likewise for each of the 2000 computers.

- (3) a. 500 Dutch firms use 2000 Japanese computers.
b. $\exists X(\text{500}(X) \wedge \text{*DF}(X) \wedge \exists Y(\text{2,000}(Y) \wedge \text{*JC}(Y) \wedge \langle X, Y \rangle \in \text{*}U))$

As intended by Sternefeld (1998), definition * also is robust with regard to collective verbs. This can be seen by replacing *use* in (3-a) with *own*. While using arguably takes place on the individual level, computers can certainly be jointly owned by more than one firm. But then this firm will be a member of a subset of the firms which jointly own the computer. So definition (2-b) and the formalisation in (3-b) can handle this case without further ado.

2.2 Augmented Logical Form

In (Sternefeld, 1998), the operators introduced in the previous section can freely be inserted into logical forms without needing to be contributed by some lexical item. More precisely, while pluralised nouns typically carry * as the semantic contribution of the plural ‘morpheme’, morphological pluralisation of verbs (which, in English, is mostly redundant anyway and only realises agreement with a single argument) is, as such, semantically vacuous. Argument slots of verbs are instead pluralised by inserting * or ** as ‘semantic glue’, which may happen in any place, given that the types permit it.

For a sentence like (4), among others, the readings illustrated in (5) are thus predicted.²

- (4) Five men lifted two pianos.

²To enhance legibility, I follow Sternefeld (1998) in using $x \in S$ as a notational variant of $S(x)$ and in using uncapitalised letters for variables that are subject to a pluralisation operation. But capitalisation has no bearing on the identity of variables, i.e. X and x are merely notational variants of the same variable. Variables are (also following Sternefeld (1998)) reused ‘after’ pluralisation, i.e. $\text{*}\lambda x.\phi$ typically will be applied to the variable x (then written X) again.

- (5) a. $\exists X(5(X) \wedge *M(X) \wedge \exists Y(2(Y) \wedge *P(Y) \wedge L(X, Y)))$
 b. $\exists X(5(X) \wedge *M(X) \wedge X \in *\lambda x.\exists Y(2(Y) \wedge *P(Y) \wedge L(x, Y)))$
 c. $\exists X(5(X) \wedge *M(X) \wedge \exists Y(2(Y) \wedge *P(Y) \wedge Y \in *\lambda y.L(X, y)))$
 d. $\exists X(5(X) \wedge *M(X) \wedge \exists Y(2(Y) \wedge *P(Y) \wedge X \in *\lambda x.L(x, Y))$
 e. $\exists X(5(X) \wedge *M(X) \wedge \exists Y(2(Y) \wedge *P(Y) \wedge X \in *\lambda x.Y \in *\lambda y.L(x, y)))$
 f. $\exists X(5(X) \wedge *M(X) \wedge \exists Y(2(Y) \wedge *P(Y) \wedge \langle X, Y \rangle \in **\lambda x.\lambda y.L(x, y)))$

(5-a) means that five men, together, lifted two pianos, at once. Generally, an unpluralised lexical predicate is supposed to relate only objects that stand in some given relation, e.g. lifting or being lifted, *together*. (5-b) means that there are five men who can, in some way, be split into subgroups, each of which, together, lifted two pianos at once. (5-c) means that there are five men and two pianos and that the men, together, lifted the pianos, either at once or separately. (5-d) means that the five men can be divided into subgroups, all of which lifted the same two pianos at once. According to (5-e), there are five men and two pianos and subsets of the five men exist who lifted the pianos together, but perhaps not all of them at once. (5-f) means that the five men lifted the two pianos in some arbitrary configuration. All that is required is that every man took part in a lifting and that every piano was lifted.

These examples illustrate that, on the verb, the numbers and types of pluralisation operations may vary freely (while plural noun denotations always involve $*$). Furthermore their scope need not be the verb meaning alone but may also involve arguments of the verb, as shown by (5-b). This is achieved by the treatment of pluralisation as ‘semantic glue’: it is not a part of the lexical meanings of plural verbs but inserted into logical forms in appropriate places.

Sternefeld (1998) points out that, while there is significant overlap and even entailment between the different propositions expressed by the readings shown in (5), all of these readings are conceivable and there is thus no harm in assuming their existence.

3 Lasersohn’s criticism of (Gillon 1987)

Having introduced the system of (Sternefeld, 1998), I now turn to the examples that Lasersohn (1989) put forth against the plural semantics advocated by Gillon (1987). It will turn out that Lasersohn’s criticism is also applicable to Sternefeld’s system. The reason will turn out to be that (Sternefeld, 1998) allows for pluralising the same verbal argument position more than once, which is a direct consequence of the treatment of pluralisation as semantic glue.

Gillon (1987) argues that the readings of a plural sentence like (6) correspond to the *minimal covers* of its subject.

- (6) The TAs wrote papers.

A minimal cover of a set S is a subset C of $\mathcal{P}(S)$ such that $\bigcup C = S$ and for no $C' \subset C$ is it the case that $\bigcup C' = S$.³ So a minimal cover of S splits S into (not necessarily disjoint) groups such that every member of S is in some group and there is no redundant group that could be dispensed with while retaining all members of S as members of one of the remaining groups.

Given a set of TAs comprising Alice, Bob and Ludwig, Gillon's proposal then predicts, among others, the following readings.

- $\{\{a, b, l\}\}$: All TAs wrote papers together.
- $\{\{a, b\}, \{l\}\}$: Alice and Bob wrote papers together and Ludwig wrote papers alone.
- $\{\{a\}, \{b\}, \{l\}\}$: Each of the TAs wrote papers alone.

Clearly, the examples given do not exhaust the possible readings of (6) under the analysis advocated by Gillon (1987), and it is clear that the number of readings will grow exponentially with the cardinality of the subject's extension.

The criticism put forth by Lasersohn (1989) is twofold: for one thing, he claims that Gillon's very concept of a *reading* is misguided: the number of readings should not be inflated in the manner indicated and, most importantly, not in a way that makes the readings that exist depend on contingent facts about the world. This criticism seems well justified; in a world in which Bob is not a TA or there is a fourth TA, the class of readings assigned to (6) would not only be different from the one assigned to that sentence in the situation considered above, but the classes would actually be disjoint. So in two utterance situations in which there are different sets of TAs, the meanings of utterances of (6) would have *nothing in common at all*, it seems.

More importantly, regarding our present concerns, Lasersohn (1989) points out that the account of the ambiguity of plural sentences offered by Gillon (1987) predicts that the sentences in (7) all have true readings in a situation in which each of three TAs got paid (exactly) \$7,000. (7-a) is true under a distributive reading and (7-b) under a collective reading in this situation. But (7-c) should not have a true reading.

- (7)
- a. The TAs were paid exactly \$7,000.
 - b. The TAs were paid exactly \$21,000.
 - c. The TAs were paid exactly \$14,000.

But if the TAs are again Alice, Bob and Ludwig, then $\{\{a, b\}, \{a, l\}\}$ is a minimal cover. But then each element of this cover fulfills *were paid exactly \$14,000* and *the TAs* thus also should, contrary to fact.

While one might guess that the non-empty intersection of the elements of the cover is to blame and that partitions should be used instead of covers, allowing non-empty intersections is actually a feature of Gillon's analysis, motivated by (8).

³ $\subset$ denotes *proper* subsethood.

(8) The men wrote musicals.

If *the men* is taken to refer to the set comprising Rodgers (*r*), Hammerstein (*hs*) and Hart (*ht*), it seems that the sentence would be judged true by those familiar with these men. But none of them wrote any musicals alone and likewise no musical was written by all three of them collaboratively. The minimal cover that corresponds to the true reading of (8) is $\{\{r, ht\}, \{r, hs\}\}$, the set of subsets of these men who collaboratively wrote musicals. This cover has just the same shape as the minimal cover of the TAs that proved problematic above. Lasersohn (1989) suggests that a meaning postulate like (9) may be used to guarantee the truth of (8) in the pertinent situation.

$$(9) \quad W(x, y) \ \& \ W(u, v) \Rightarrow W(x \cup u, y \cup v)$$

This clearly defies Gillon's aim to treat (8) as ambiguous between a collaborative (i.e. simple collective) and a distributive reading and further ones that are neither fully distributive nor collective. While the scepticism expressed by Lasersohn (1989) regarding the readings licensed by Gillon's account seems very justified, obliterating the distinction between collective and non-collective for the predicate *write* might be going too far.

4 Applying Sternefeld's semantics

(8) can be analysed in Sternefeld's system as in (10).

$$(10) \quad \llbracket the \ men \rrbracket \in {}^*\mathbf{WM}$$

Given that $\{\{r, ht\}, \{r, hs\}\} \subseteq \mathbf{WM}$, $\bigcup\{\{r, ht\}, \{r, hs\}\} = \{r, ht, hs\} \in {}^*\mathbf{WM}$.

Since $*$ need not be inserted into the logical form, analysis (11) is also possible. As a set is an element of an unpluralised predicate if its members fulfill the predicate as a group, this reading would only be true if the three composers had collaborated, which is not the case.

$$(11) \quad \llbracket the \ men \rrbracket \in \mathbf{WM}$$

The sentences in (7) can receive the representations in (12).

$$(12) \quad \begin{array}{ll} \text{a.} & \mathbf{TA} \in {}^*\lambda x. \exists Y ({}^*\$(Y) \wedge \mathbf{7,000}(Y) \wedge Y \in {}^*\lambda y. \mathbf{PAID}(x, y)) \\ \text{b.} & \exists Y ({}^*\$(Y) \wedge \mathbf{21,000}(Y) \wedge \langle \mathbf{TA}, Y \rangle \in {}^{**}\mathbf{PAID}) \\ \text{c.} & \exists Y ({}^*\$(Y) \wedge \mathbf{14,000}(Y) \wedge \langle \mathbf{TA}, Y \rangle \in {}^{**}\mathbf{PAID}) \end{array}$$

As things stand, all of these will be true. But (12-c) only is because the meaning of *exactly* cannot be correctly represented so far, which requires that *Y* be the unique maximal set of dollars that fulfills the scope (i.e. what follows the part stating that *Y* is a certain amount of dollars), and issue that will be taken up in section 6 below. When this maximization operation is put in place, (12-c) will come out false in the

pertinent situation. But there is a further possible rendering of (7-c), shown in (13).

$$(13) \quad \mathbf{TA} \in \lambda x. \exists Y (*\$ (Y) \wedge \mathbf{14,000}(Y) \wedge \langle x, Y \rangle \in **\mathbf{PAID})$$

In (13), the expression beginning with λx denotes the set of all sets of individuals who received \$14,000 in total. In the situation considered above, both $\{a, b\}$ and $\{a, l\}$ – for instance – are such sets. Pluralising $\lambda x. \dots$ then yields a set that contains all unions of sets of this kind, and the set $\mathbf{TA} = \bigcup \{\{a, b\}, \{a, l\}\}$ is such a set. Thus it appears that (13) is a predicted reading of (7-c) that should be true in the situation considered, while in fact (7-c) is not true. This parallels the situation found in (Gillon, 1987): under both accounts, finding groups who received a total of \$14,000 is enough to make (7-c) true.⁴

In Sternefeld’s system, it is essential for (13) to be obtained that the subject position of **PAID** be pluralised twice, once using ****** and then again using *****. Leaving out the latter operation yields (14). It is easily seen that this formula – also a predicted reading of (7-c) under Sternefeld’s approach – is not true in the given situation. It expresses that there is a sum of (exactly) \$14,000 that the TAs received, without any implications as to who of them received how much.

$$(14) \quad \mathbf{TA} \in \lambda x. \exists Y (*\$ (Y) \wedge \mathbf{14,000}(Y) \wedge \langle x, Y \rangle \in **\mathbf{PAID})$$

Since the ALF framework allows for free insertion of pluralisation, it is not clear how it could rule out (13), which is just (14) with an additional pluralisation operator inserted, without imposing restrictions of a decidedly non-local nature on LF. In the next section, a solution to this problem is developed in LRS.

5 Recasting the system in Lexical Resource Semantics

The present proposal addresses the problem discussed above by capturing the essential ideas of (Sternefeld, 1998) about *where* pluralisation should be insertible while taking a different stance with respect to *how* pluralisation should be inserted. The locus of pluralisation will be strictly lexical. At the same time, pluralisation can occur in different places, not directly tied to the core meaning of the verb, i.e. with material contributed by other expressions intervening. This is achieved using Lexical Resource Semantics.

⁴In a reply to (Lasnik, 1989), Gillon (1990) describes a situation in which two departments employ two TAs each. \$14,000 are paid for each pair of TAs, which they may divide among themselves as they deem fit. It then seems that (i-a) would be judged true. But – disregarding the role of “their” and ignoring the temporal modifier – this can be formalised as in (i-b) under the present approach.

- (i) a. The TAs were paid their \$14,000 last year.
- b. $\mathbf{TA} \in * \lambda x. \exists Y (*\$ (Y) \wedge \mathbf{14,000}(Y) \wedge \mathbf{PAID}(x, Y))$

Since each pair of TAs was paid as a team, the sets of respective team members will each be related to \$14,000, but the individual members will not. Under these circumstances, (i-b) is true.

5.1 Lexical Resource Semantics

LRS (Richter & Sailer, 2004; Kallmeyer & Richter, 2007) is a flavour of underspecified semantics that makes use of the descriptive means of HPSG and uses its constraint language, which is assumed here to be Relational Speciate Reentrant Language (RSRL; Richter, 2004), as the locus of underspecification. Disregarding the treatment of local semantics (Sailer, 2004a), the semantic representation connected to a sign (i.e. a syntactic object) is an object of a sort *lrs* to which three features are appropriate: INCONT, EXCONT and PARTS. For each word, the value of the INCONT feature is this word’s scopally lowest semantic contribution, i.e. that part of its semantics over which every other operator in the word’s maximal projection takes scope. The EXCONT value roughly corresponds to the meaning of the maximal projection of a word. Both INCONT and EXCONT project strictly along head lines.

The value of PARTS is a list that contains the lexical resources that a sign contributes. For words, they are lexically specified. For phrases, they always are the concatenation of the PARTS lists of the daughters. In an utterance, – an unembedded sign – each element of the PARTS list must occur in the EXCONT value and everything that occurs in it must be on the PARTS list. The EXCONT value of an utterance is regarded as its meaning.

The values of the three attributes are related by a small set of core constraints. In addition to these, the SEMANTICS PRINCIPLE provides further more or less construction-specific constraints that ensure that they are also related in a way that correctly represents how meaning is composed in different syntactic configurations. For example, in *every dog*, the INCONT of *dog*, $D(x)$, must be a subexpression of the restrictor of the universal quantifier. Since the NP contains no further material that combines with *dog*, this will actually result in identity. Similarly, if a quantified NP combines with a verb, the verb’s INCONT must be found in the NP quantifier’s scope. *Every dog barks* thus receives the desired interpretation $\forall x(D(x), B(x))$.

5.2 The analysis

Almost everything that needs to happen for the present approach to work happens on the PARTS list. Manipulating this list gives the opportunity to furnish lexical items with semantic material that must occur in the utterance they are used in, but the places in which this material can occur are not subject to any restriction that is not explicitly stated. By placing pluralisations on this list, it is thus possible to achieve an effect similar to their treatment as freely insertible glue in (Sternefeld, 1998). At the same time, lexically constraining the PARTS list to disallow repeated pluralisation of the same variable makes it possible to rule out readings like (13).

To illustrate the approach, it will be shown how the system accounts for the readings of *five men lifted two pianos* in (5).

The general shape of the LRS semantics of verbs like *lift* is as follows.⁵

$$\left[\begin{array}{ll} \text{INCONT} & \boxed{1} \mathbf{L}(y)(x) \\ \text{EXCONT} & \boxed{1} \\ \text{PARTS} & \langle \boxed{1}, (\mathbf{L}y), \mathbf{L} \rangle \oplus \boxed{2} \end{array} \right]$$

The PARTS list contains the INCONT (as required by a fundamental principle of LRS) and those of its subexpressions that *lift* needs to contribute as lexical resources (the variables are contributed by the NPs). In addition, it contains all elements of the list $\boxed{2}$, which is where pluralisation operations enter. $\boxed{2}$ is subject to the following conditions.^{6,7}

- (15) a. Every variable that is associated with a plural nominal argument of the verb may be subject to at most one pluralisation operation on $\boxed{2}$.
b. Only variables that are associated with a plural nominal argument of the verb may be subject to a pluralisation operation on $\boxed{2}$.

In the sense of (15), a variable x_i , $1 \leq i \leq n$, is subject to a pluralisation operation on $\boxed{2}$ if $\boxed{2}$ contains ${}^n\lambda x_1 \dots \lambda x_n.\phi$.⁸ It is the restriction (15-a) that will prevent the unwanted reading of Lasersohn's example sentence. Since at most one pluralisation is allowed, the kind of double pluralisation that was identified as problematic above is ruled out. Formalising (15) in RSRL is a tedious but straightforward task.

Now consider again the lexical entry for *lift*. $\boxed{2}$ may be empty. This will give the reading of (4) in (5-a). Five men jointly lifted two pianos at once. Further admissible lists are shown in (16).⁹

⁵Tags like $\boxed{1}$ are variables used to indicate token identity. So $\boxed{1}$ in the entry for *lift* always denotes the expression $\mathbf{L}(y)(x)$. $\boxed{1}$ may be any expression that contains $\boxed{1}$.

⁶Readers have raised questions regarding independent motivation of this constraint. But I think that (while it would of course be welcome) demanding such makes the possibility of repeated pluralisation, as in ALF, the null hypothesis. But since being allowed to enrich meanings with unlimited amounts of material is not the established standard in semantics. I fail to see any better motivation for allowing multiple pluralisations of the same argument than for not doing so, especially if the latter approach makes more accurate predictions.

⁷It must be possible to isolate $\boxed{2}$ from the idiosyncratic contributions of the verb, i.e. the parts on the list that $\boxed{2}$ is appended to. The most straightforward solution is to introduce a new attribute PLURALISATIONS whose value is $\boxed{2}$. Then the following AVM would describe all verbal lexical items.

$$\left[\begin{array}{ll} \text{PARTS} & \langle \dots \rangle \oplus \boxed{2} \\ \text{PLURALISATIONS} & \boxed{2} \end{array} \right]$$

Constraining the pluralisations that may be introduced is then possible by formulating the appropriate constraints with regard to the value of PLURALISATIONS.

⁸I.e. an operator of n stars applied to an n -place relation. In this paper, $n \leq 2$.

⁹The dots on the list stand for lexical resources that are needed in addition to those explicitly shown (namely parts of these). For list (16-a), for instance, these would be ${}^*\lambda x.\boxed{1}$ and $\lambda x.\boxed{1}$.

- (16) a. $\langle (*\lambda x. [\boxed{1}])(x), \dots \rangle$
 b. $\langle (*\lambda y. [\boxed{1}])(y), \dots \rangle$
 c. $\langle (*\lambda x. [\boxed{1}])(x), (*\lambda y. [\boxed{1}])(y), \dots \rangle$
 d. $\langle (*\lambda x. \lambda y. [\boxed{1}])(x)(y), \dots \rangle$

It is easily seen that all lists in (16) conform to (15): In (16-a), only x is subject to pluralisation and only pluralised once. The same is true for y regarding list (16-b). In (16-c), both are pluralised once, independently of each other. In (16-d), both variables are pluralised together once using $**$.

Importantly, now, $[\boxed{1}]$ in the expressions above may stand for any expression that has $[\boxed{1}]$ (the verb's INCONT) as a subexpression.¹⁰ All that is thus said about the scope of the pluralisations is that they contain $L(y)(x)$. Unless constrained further, they may thus occur anywhere in the meaning representation of a complete sentence, provided that the types fit. While the number of possible pluralisations is thus limited and while they enter into the semantics from the lexicon, their distribution will in other respects be as under the ALF approach.

As remarked above, plural nouns are always pluralised using $*$. For *pianos*, e.g., the semantics is as follows.¹¹

$$\left[\begin{array}{l} \text{INCONT} \quad [\boxed{3}]^* \mathbf{P}(X) \\ \text{EXCONT} \quad [\boxed{2}] \exists X ([\boxed{3}], [\boxed{4}]) \\ \text{PARTS} \quad \langle [\boxed{2}], [\boxed{3}], * \mathbf{P}, \mathbf{P}, X \rangle \end{array} \right]$$

For cardinals, an analysis as higher-order intersective modifiers is assumed here, where intersective modifiers are analysed along the lines outlined in (Sailer, 2004b), although nothing hinges on this decision. A cardinal like *three* will have the following LRS semantics.

$$\left[\begin{array}{l} \text{INCONT} \quad [\boxed{5}] 3(X) \\ \text{EXCONT} \quad [\boxed{6}] ([\boxed{5}] \wedge [\boxed{7}]) \\ \text{PARTS} \quad \langle [\boxed{5}], 3, [\boxed{6}] \rangle \end{array} \right]$$

The INCONT expresses that X is a set of (at least) three individuals. It is required to be a part of the first conjunct of the EXCONT, which conjoins it with an underspecified expression represented as $[\boxed{7}]$. The clauses of the SEMANTICS PRINCIPLE responsible for adjective-noun-constructions will then require the INCONT of the nominal head to be a part of this second conjunct. The variable X is embedded in the agreement index that is shared between noun and adjective: the adjective modifies a noun with an INDEX value token-identical with its own. So the adjectival and nominal INCONT involve the same variable. Hence *three pianos*

¹⁰Each occurrence of $[\boxed{1}]$, even on the same list, may stand for a different such expression.

¹¹The existential quantifier is now taken to combine with the variable it binds and two expressions of type t and to state that the sets formed by abstracting over the bound variable have a non-empty intersection. This slightly reduces syntactic complexity since $\wedge$ can be eliminated. The representation of the existential quantifier will be revised further below.

will contribute the part $[3(X)] \wedge [*P(X)]$, and in the absence of additional material, this will resolve to $3(X) \wedge *P(X)$.

The combinatorial behaviour of plural NPs is as dictated by the SEMANTICS PRINCIPLE and discussed above: the INCONT of a verbal projection they combine with needs to be a part of their scope, i.e. ④ above.

The system can be illustrated by comparing (5-b) and (5-d).

(5-b) $\exists X(5(X) \wedge *M(X), (*\lambda x.\exists Y(2(Y) \wedge *P(Y), L(Y)(x)))(X))$

(5-d) $\exists X(5(X) \wedge *M(X), \exists Y(2(Y) \wedge *P(Y), (*\lambda x.L(Y)(x))(X)))$

Both expressions are predicted to represent possible meanings of *five men lifted two pianos*. They are only distinguished by the place in which x is pluralised. The variable y is not pluralised. As required by the SEMANTICS PRINCIPLE,

- ④ (i.e. $L(y)(x)$) is in the scope of both quantifiers in both (5-b) and (5-d),
- the INCONT values of the nouns are in the restrictors of the existential quantifiers,
- the INCONT values of the nouns further are the second conjuncts of the conjunctions in these restrictors and
- the INCONT values of the adjectives are the first conjuncts of these conjunctions.

The basic requirements for components other than pluralisation are thus fulfilled. The pluralisation remains to be considered.

Since only x is pluralised, the pluralisation list (16-a) needs to be assumed in both cases, but with different expressions as values of $[\textcircled{1}L(y)(x)]$. In (5-b), this expression includes the expression $\exists Y \dots$, in (5-d) it is identical with ④. Both conform to the requirement of (16-a) that ④ be in the scope of the pluralisation operator. Since both readings fulfill all pertinent constraints, they both are predicted to be possible, as desired.

The problematic reading (13) of (7-c) is ruled out since it would require a list like (17) in order for the two pluralisations found in that formula to be available as lexical resources.

(17) $* \langle \textcircled{1}, (*\lambda x. \textcircled{1})(x), (*\lambda x. \lambda y. \textcircled{1})(x)(y), \dots \rangle$

But since x is pluralised twice, (17) violates (15-a) and is hence not a possible pluralisation list. (14) only requires list (16-d) and thus remains a possible reading.

6 Maximalisation

As remarked above, the system as it stands cannot deal with non-upward-monotone quantifiers. The predicted reading (18-b) of (18-a) is true even if 10 children were at the party. ($=2$ denotes the set of all sets of entities of cardinality *exactly* 2.)

- (18) a. Exactly two children were at the party.
 b. $\exists X (=2(X) \wedge *C(X), *P(X))$

If there were ten children at the party, it is possible to single out a set of exactly two of them, which is enough to make (18-b) true, under the plausible assumption that **P** itself is a set of individuals, i.e. having been at the party is lexically a property of individuals, not pluralities (cf. van Benthem, 1986, 52f.). This issue can be addressed by requiring the set of exactly two children to be a maximal set of children who were at the party, as shown in (19).

$$(19) \quad \exists X (=2(X) \wedge *C(X), \max(X)(\lambda X. =2(X) \wedge *C(X))(\lambda X. *P(X)))$$

The meaning of *max* is defined in (20).

$$(20) \quad \max(X)(P)(Q) := X \in Q \wedge \forall Y \in *P \cap Q: X \not\subset Y$$

$\max(X)(P)(Q)$ is true if X is in Q (the extension of *were at the party* in the case of (18-a)) and no set that is in both $*P$ and Q is a superset of X . In this second condition, pluralising P is necessary in order to cancel out the cardinality restriction that the elements of P obey, like being sets of exactly two children. If this cardinality restriction remained in place (if P were used instead of $*P$), then each set considered in (18-b) would be maximal in the sense of *max* because $\forall$ would only quantify over sets of exactly two children. Hence, even if there were ten children at the party, (18-a) would still come out as true because the set of ten children would be no element of the restrictor. The use of $*$ ascertains that the quantification is over the set of all sets of at least two children. So (18-b) is true if there is a set of exactly two children which is a maximal set of children who were at the party.

The meaning representations can be simplified somewhat. Note that *max* needs to know the restrictor of the existential quantifier in order to determine the correct subclass of sets in which to maximise – (18-a) is not false if there were twenty adults at the party in addition to the two children. But then *max* can also state itself that the quantified variable takes a value from the restrictor. At this point, then, the meaning of *max* can also be incorporated directly into the existential quantifier, defining

$$(21) \quad \begin{aligned} \text{a. } & \max(X)(P)(Q) := X \in P \cap Q \wedge \forall Y \in *P \cap Q: X \not\subset Y \\ \text{b. } & \exists^\circ := \lambda R. \lambda S. \exists X: \max(X)(R)(S) \end{aligned}$$

$\exists^\circ$ will be called a *maximalisation quantifier*. Using this quantifier, the representation of the meaning of (18-a) becomes (22).

$$(22) \quad \exists^\circ (\lambda X. =2(X) \wedge *C(X)) (\lambda X. *P(X))$$

What are the predictions of this account in cases involving more than one plural noun phrase?

For (23-a), the reading (23-b) is predicted, among others.¹² Readings without pair maximalisation are ruled out since they will not be able to use the *max_{me}* expression contributed by the verb.

- (23) a. Exactly three boys invited exactly four girls.
b. $\exists^\circ(\lambda X. =3(X) \wedge *B(X))(\lambda X. \exists^\circ(\lambda Y. =4(Y) \wedge *G(Y))(\lambda Y. **I(X, Y)))$

This formula is verified by fig. 1, where black circles stand for boys and white circles stand for girls: there is a set of three boys X (those on the left) for whom a set of four girls Y exists such that $**I(X, Y)$ and such that no larger set of girls exists such that the same holds. X also is the largest set of this kind, i.e. no superset of X is also related to a (maximal) set of four girls.

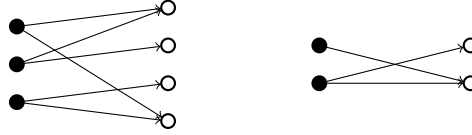


Figure 1: A situation that verifies (23-b).

This prediction will be discussed below. First note that there is a problem with the current approach due to the fact that one maximalisation quantifier must outscope above the other. This is illustrated by fig. 2. Considering sentence (24), does fig. 2 verify it or not?

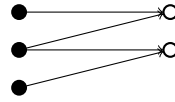


Figure 2: Figure that verifies (24).

- (24) Exactly one boy invited exactly one girl.¹³

This depends on the scope relations. If the scope is as in (25-a), then (24) is true: there is a set of one boy such that there is a set of one girl that is a maximal set of girls he invited; the sets may be any set containing a boy who invited just one girl and the girl he invited, respectively.

But (25-b) is false in the same situation: there is no set of one girl such that the maximal set of boys who invited her has just one member.

- (25) a. $\exists X(1B(X) \wedge \max(X)(\lambda X. \exists Y(1G(Y) \wedge \max(Y)(\lambda Y. **I(X, Y))))))$

¹²Predications are written as $R(x, y)$ instead of $R(y)(x)$ from now on. This is only a shorthand of no further significance.

¹³There is no claim made here that a sentence with singular noun phrases should really be analysed in the way indicated. The only purpose for using this sentence is that it allows to keep the illustration small. In fact, the more complex fig. 1 could also be used to illustrate the same fact.

$$\text{b. } \exists Y(1\mathbf{G}(Y) \wedge \max(Y)(\lambda Y.\exists X(1\mathbf{B}(X) \wedge \max(X)(\lambda X.**\mathbf{I}(X, Y))))))$$

According to my intuitions, there is no such ambiguity involved in sentences like (24) (or actual plural sentences of the same kind) and (24) is not verified by the situation in fig. 2. So it seems that the approach is in need of modification.

Before such a modification is actually introduced, let us return to the prediction about (23-a) and fig. 1; the former was predicted to be true in the situation depicted in the latter. This prediction differs from that of the theory laid out in (Landman, 2000). According to (Landman, 2000), (23-a) should be false for fig. 1 since all boys who invited girls and all girls invited by boys are taken into account by determining the maximal event of boys inviting girls. Only if this event has three boys as its agent and four girls as its patient will (23-a) come out true, but in fig. 1 the agent of this event would consist of five boys and the patient of six girls.

Robaldo (2010, 260), offers evidence against the approach advocated by Landman (2000) and an analysis according to which (23-a) in fact has a reading that is true given fig. 1. Note that this implies that (23-a) is not incompatible with e.g. *exactly five boys invited exactly six girls*, which is also verified by fig. 1. My own (non-native) intuitions on the issue are equivocal, but I tend to side with the predictions of (Robaldo, 2010). This also holds for the corresponding sentence in my native German.

(26) Genau drei Jungen haben genau vier Mädchen eingeladen.

Unlike the account developed so far, that of (Robaldo, 2010) does not exhibit false scope ambiguities. Taking its departure from (Sher, 1997), it is based on maximalisation of pairs of sets similarly to what is shown in (27).¹⁴ Robaldo (2010) argues that this approach should be used for all plural quantificational expressions, regardless of their monotonicity properties, giving examples corroborating the approach even for downward monotone quantifiers.¹⁵

$$\begin{aligned} (27) \quad & \text{MAX}(\langle P, Q \rangle, N_1, N_2, R) := \\ & **R(P, Q) \wedge \forall P' Q' \\ & ((P \subseteq P' \wedge P' \in *N_1 \wedge **R(P', Q) \rightarrow P' \subseteq P) \wedge \\ & (Q \subseteq Q' \wedge Q' \in *N_2 \wedge **R(P, Q') \rightarrow Q' \subseteq Q)) \end{aligned}$$

A pair is maximal if no component of it can be made any larger while keeping the other fixed.¹⁶ The representation of (23-a) then becomes

¹⁴The formulation in Robaldo (2010) employs quantification over covers to achieve the effect of **, and a contextually determined cover variable as advocated by Schwarzschild (1996). Furthermore, Robaldo's definition does not incorporate the restrictors N_1 and N_2 . This omission is an error: if three children watched two movies and one of their grandparents also watched one of the movies, *exactly three children watched exactly two movies* will be predicted to be false due to the larger set that includes the grandparent.

¹⁵Robaldo (2010) suggests that apparent failures of the approach for such quantifiers in other cases should be explained by pragmatics. In this paper, I follow these assumptions.

¹⁶Specifying a game with two players where, for each $i \in \{1, 2\}$, outcome $f_i(\langle x_1, x_2 \rangle) \geq$

$$(28) \quad \exists X (=3(X) \wedge X \in *B \wedge \exists Y (=4(Y) \wedge Y \in *G \wedge MAX(\langle X, Y \rangle, *B, *G, I)))$$

This sentence is true in the situation depicted by fig. 1: the three boys in the group to the left are the maximal set of boys that cumulatively invited the four girls in that group and these four girls are the maximal group of girls they invited. Likewise, (24) is now false in the situation depicted by fig. 2, since for no set of girls is there a set of just one boy that is a maximal set of boys who cumulatively invited the girl. In the sequel I will adopt the idea proposed by Robaldo (2010) and show how it can be implemented in LRS. The implementation does not require adopting Robaldo's proposal, though. With a different definition of the *max* operator below, the essential ideas of (Landman, 2000) could likewise be implemented.

7 Implementation of maximalisation in LRS

In order to implement the proposal by Robaldo (2010), it is necessary to be able to maximalise pairs of sets instead of only one set at a time. In addition, in order to actually rule out the unwanted readings, maximalisation of pairs instead of sets also needs to be enforced. Each of these requirements is addressed in turn.

Maximalisation of pairs

Maximalisation of pairs is achieved by exploiting one of the most notable features of LRS, namely that distinct expressions may contribute identical parts of the semantics, which allows for meaning components to be 'fused'. This feature was put to use in an analysis employing polyadic quantifiers in (Iordăchioaia & Richter, 2015). In the present approach, the maximalisation quantifiers contributed by distinct noun phrases can turn out to be one and the same. This is achieved by analysing maximalisation quantifiers categorically, instead of as variable binders (cf. Richter, 2016) and ascertaining that the contributions of two distinct NPs can be fused into a single semantic expression that employs a polyadic quantifier to express the desired pair maximalisation.

In order to express pair maximalisation, *max* is renamed *max*¹ and in addition, pair maximalisation *max*² is introduced as defined in (29).¹⁷

$$(29) \quad \begin{aligned} &max_{\langle \langle e, t \rangle, \langle \langle e, t \rangle, \langle \langle e, t \rangle, \langle \langle e, t \rangle, \langle \langle e, \langle e, t \rangle \rangle, t \rangle \rangle \rangle \rangle}^2 := \\ &\lambda X. \lambda N. \lambda Y. \lambda M. \lambda R. \\ &X \in N \wedge Y \in M \wedge R(X, Y) \wedge \\ &\forall X' \subseteq X: (*N(X') \wedge R(X', Y) \rightarrow X' \subseteq X) \wedge \\ &\forall Y' \subseteq Y: (*M(Y') \wedge R(X, Y') \rightarrow Y' \subseteq Y) \end{aligned}$$

$f_i(\langle x'_1, x'_2 \rangle)$ iff $x_i \in *N_1$, $\langle x_1, x_2 \rangle \in R$ and $x_i \supseteq x'_i$, $MAX(\langle x_1, x_2 \rangle, N_1, N_2, R)$ states that $\langle x_1, x_2 \rangle$ is a Nash Equilibrium.

¹⁷If needed, *max*³ or *max*ⁿ for even larger *n* could of course also be introduced. Also note that individuals and pluralities both are of type *e* in the present system.

```

me
  max_me
    max_application FUNCTOR max_me
     $\exists^{\circ 1}$ 
     $\exists^{\circ 2}$ 
  application
    max_application
    non_max_application

```

Figure 3: Additions to the sort hierarchy for LRS expressions.

(29) encodes the meaning of *MAX* above: X and Y are the sets provided by the existentially bound variables and N and M are the sets denoted by the corresponding nominal expressions. As before, the assertion that the sets belong to the noun denotations and the verbal scope is also encoded in *max*.

$\exists^{\circ}$ is likewise renamed $\exists^{\circ 1}$ and $\exists^{\circ 2}$ is defined analogously:

$$(30) \quad \exists^{\circ 2} := \lambda R_1. \lambda R_2. \lambda S. \exists XY : \max^2(X)(R_1)(Y)(R_2)(S)$$

Next, the sort hierarchy is extended slightly by introducing a new subsort *max_me* of the sort *me* of meaningful expressions.¹⁸ This will make it possible to talk about the quantifiers $\exists^{\circ 1}$ and $\exists^{\circ 2}$ without knowing whether they are the primitive $\max^1$ or $\max^2$ or the result of applying $\max^2$ to some of its arguments.

The sort *max_me* has as its subsorts *max_application*, which is also a subsort of *application*, and $\exists^{\circ 1}$ and $\exists^{\circ 2}$.¹⁹ ‘Normal’ application of non *max_me* functions is now of the sort *non_max_application*, which is not a subsort of *max_me*. The maximally specific subsorts $\exists^{\circ 1}$ and $\exists^{\circ 2}$ of *max_me* represent the monadic and polyadic maximalisation quantifier, respectively. *max_application* represents the results of applying an expression of sort *max_me* to an argument. Since it is a subsort of *application*, the constraints that regulate the wellformedness of application expressions with regard to typing affect it as well. But so far, nothing necessitates using *max_application* in applications of *max_me* expressions. The sort hierarchy only rules out using *max_application* to apply anything that is not of this sort. To enforce the converse as well, the following constraint is introduced:

¹⁸The type-logical language used in LRS is built from the same kind of graph structures that are used to model natural language. All expressions of the formal language have the sort *me* to which the attribute TYPE is appropriate. Constants and variables have the subsorts *constant* and *variable* respectively and are identified by an INDEX value (a natural number, itself encoded in the same way). Complex expressions are represented by structures of, e.g. sort *application* with appropriate attributes FUNCTOR and ARG. Suitable constraints guarantee that, for instance, the TYPE value of an *application* is the type of the FUNCTOR value applied to the ARG value. So if the FUNCTOR value of an *application* object represents an expression ϕ of type $\langle \tau, \sigma \rangle$ and its ARG value represents an expression α of type τ , then the *application* object itself represents $\phi(\alpha)$ of type σ . See (Penn & Richter, 2004) for a concise formal statement.

¹⁹The quantifiers must of course be further constrained to have the appropriate types.

$$(31) \quad [application, \text{FUNCTOR } max_me] \rightarrow [max_me]$$

This constraint states that every application of a *max_me* functor to an argument must again result in a *max_me*. The modified sort hierarchy and constraint (31) ascertain that max^1 , max^2 and whatever results from iteratively applying them to their arguments is a *max_me* expression and that nothing else is.

With these preliminaries in place, it is possible to conveniently refer to expressions of sort *max_me* without needing to know about their exact shape in a way that would straightforwardly generalise to quantifiers with even more than two restrictors. This will be the key to fusing the quantifiers contributed by different NPs into a single polyadic quantifier. The next subsection specifies the syntax-semantics interface that will allow for these fusions and enforce them where required.

Fusing maximalisations

The lexical entries of plural nouns are now given the following shape.

$$\left[\begin{array}{l} \text{INCONT } \boxed{3}^* \mathbf{P}(X) \\ \text{EXCONT } \boxed{2} \phi(\boxed{6} \lambda X. \boxed{3}) \\ \text{PARTS } \langle \boxed{3}, * \mathbf{P}, \mathbf{P}, X, \boxed{6}, \boxed{2}, \phi \rangle \end{array} \right]$$

Where ϕ is of sort *max_me*.

ϕ is a *max_me* expression, i.e. one of the primitive quantifier symbols or the polyadic quantifier applied to the restrictor that comes with some other noun. This is all that is required to allow quantifiers contributed by distinct noun phrases to fuse.

Note that ϕ itself is contributed by the noun on its PARTS list, even if it is of the shape $\exists^{\circ 2}(\dots)$. One might suspect this fact to result in the possibility of smuggling in arbitrary meaning parts, since such an expression contains parts that are not contributed by the noun itself. But precisely the fact that the noun itself does not contribute these parts prevents such unwelcome results: if the noun itself does not contribute the components of something on its PARTS list, something else needs to – in this case, another noun. Also, leaving ϕ out is not an option since the primitive $\exists^{\circ 1}$ or $\exists^{\circ 2}$ needs to be contributed somewhere, and this is precisely what ϕ will need to be on the PARTS list of at least one noun.

Note that, unlike in the entries above, the lexical entries of nouns do not mention the verbal scope anymore. The EXCONT of a noun now is a quantifier that still needs to be applied to the verbal scope. This does not only bring the present LRS analysis more in line with mainstream semantics but also allows for enforcing the fusion of quantifiers: the application of the quantifier to its scope will be enforced in the lexical entry of the verb itself. Verbal lexical entries still look as shown below.

$$\left[\begin{array}{l} \text{INCONT } \boxed{1} \mathbf{L}(x, y) \\ \text{EXCONT } \boxed{1} \\ \text{PARTS } \langle \boxed{1}, (\mathbf{L}y), \mathbf{L} \rangle \oplus \boxed{2} \end{array} \right]$$

But now the list $\overline{\mathbb{E}}$ of a transitive verb with arguments pluralised by $**$ should look as shown in (32).

$$(32) \quad \langle \overline{\mathbb{E}}(**\lambda x.\lambda y.\overline{\mathbb{I}})(x, y), \phi(\lambda x.\lambda y.\overline{\mathbb{E}}) \dots \rangle, \text{ where } \phi \text{ is of sort } \textit{max_me}.$$

ϕ is a *max_me* expression that is applied to the verbal scope. The pluralisation is contained in the argument to the maximalisation operator and in turn contains the INCONT of the verb. What if a list that contains a $**$ pluralisation is required to also contain such an expression? Then ϕ must result from the application of $\textit{max}^2$ to the semantic material of two noun phrases. This is ascertained by the fact that this is the only kind of *max_me* expression that can consume an argument of type $\langle e, \langle e, t \rangle \rangle$. Note that the only part of the *max_me* expression ϕ that the verb contributes is this expression itself and the verbal scope. All its subexpressions need to be collected from somewhere else, so for ϕ to actually appear in the meaning of a full utterance, they must be contributed by appropriate noun phrases.

To guarantee in a principled way that pluralisation lists in fact have the shape in (32), the constraint in (33) is imposed on them.

$$(33) \quad \text{For each pluralisation } p \text{ on } \overline{\mathbb{E}} \text{ there needs to be a maximalisation on } \overline{\mathbb{E}} \text{ that maximalises exactly the variables pluralised by } p, \text{ in the same order.}$$

The constraint guarantees that $(**\lambda x.\lambda y.\overline{\mathbb{I}})(x)(y)$, a pluralisation of x and y , is flanked by a maximalisation quantifier like $\overline{\mathbb{E}}(\lambda x.\lambda y.\overline{\mathbb{E}})$ of the same variables. A single-star pluralisation $*\lambda x.\phi$ will accordingly need to be flanked by a maximalisation quantifier $\exists^{\circ 1}(\lambda x.\phi)$. As a consequence, whenever two variables are pluralised separately, they also need to be maximised separately. The empirical consequences of this fact merit further investigation but are beyond the scope of the present paper.

There still is need for one more constraint: the restrictors of $\exists^{\circ 2}$ must be prevented from swapping places: $\exists^{\circ 2}(\lambda x.\phi)(\lambda y.\psi)(\lambda y.\lambda x.\theta)$ must be disallowed. The outer abstraction in the third argument needs to abstract over the same variable that is abstracted over in the first argument and the inner abstraction needs to abstract over that abstracted over in the second. Such a well-formedness constraint is easy to state.

The system now predicts (34) as a reading of (23-a), as desired.

$$(34) \quad \exists^{\circ 2}(\lambda X.\overline{=3}(X) \wedge *B)(\lambda Y.\overline{=5}(Y) \wedge *G)(\lambda X.\lambda Y.*I(X, Y))$$

By the same token, the correct reading is now predicted for (7-c), paralleling (34). This reading is no more true in the situation considered in section 3 above. While \$21,000 have plenty of subsets of \$14,000, none of these is a maximal set of dollars the TAs were cumulatively paid. Of course, the problematic reading of sentence (7-c) discussed in section 4 remains unlicensed, as it would still require pluralising the same variable twice.

8 Conclusion

It has been argued that (Sternefeld, 1998) suffers from the same problem of over-generation that Lasersohn (1989) points out with respect to the analysis proposed by Gillon (1987). The source of the problem was identified as inherent to the syntax-semantics interface Sternefeld (1998) employs. His approach allows for multiple pluralisations of a single verbal argument position. Without this possibility, the overgeneration disappears. A lexicalist reformulation of Sternefeld's system was then offered that puts verbal argument pluralisation into the lexical semantics of the verb. This allowed for restricting the number of pluralisations on any argument to one. The account employs Lexical Resource Semantics (LRS), thereby offering the first approach plural semantics in this framework.

It was further demonstrated that LRS allows for a straightforward implementation of maximalisation operations that, while absent in (Sternefeld, 1998), are needed to get correct results for quantifiers that are not upward monotone. The analysis relies on the possibility of the semantic contributions of distinct constituents to be the same in LRS. This feature was used to fuse quantifiers associated with different plural NPs into a single polyadic quantifier stating the existence of a maximal pair of sets. This way, scoping of maximalisations over each other is avoided and the correct truth conditions for sentences like *exactly three boys invited exactly four girls* are derived.

References

- van Benthem, Johan. 1986. Quantifiers. In *Essays in Logical Semantics* (Studies in Linguistics and Philosophy 29), Dordrecht: Reidel 1st edn.
- Gillon, Brendan S. 1987. The readings of plural noun phrases in English. *Linguistics and Philosophy* 10(2). 199–219.
- Gillon, Brendan S. 1990. Plural noun phrases and their readings: A reply to Lasersohn. *Linguistics and Philosophy* 13(4). 477–485.
- Iordăchioaia, Gianina & Frank Richter. 2015. Negative concord with polyadic quantifiers. *Natural Language & Linguistic Theory* 33(2). 607–658.
- Kallmeyer, Laura & Frank Richter. 2007. Feature logic-based semantic composition: a comparison between LRS and LTAG. In *1st International Workshop on Typed Feature Structure Grammars (TFSG'06)*, 31–83.
- Landman, Fred. 2000. *Events and plurality: The Jerusalem lectures* (Studies in Linguistics and Philosophy 76). Dordrecht: Springer Netherlands.
- Lasersohn, Peter. 1989. On the readings of plural noun phrases. *Linguistic Inquiry* 20(1). 130–134.

- Penn, Gerald & Frank Richter. 2004. Lexical Resource Semantics: From theory to implementation. In Stefan Müller (ed.), *11th International Conference on Head-Driven Phrase Structure Grammar (HPSG)*, 423–443. Stanford: CSLI Publications.
- Richter, Frank. 2004. *A mathematical formalism for linguistic theories with an application in Head-Driven Phrase Structure Grammar*: Eberhard-Karls-Universität Tübingen Phil. dissertation (2000).
- Richter, Frank. 2016. Categorical uninterpretable polyadic quantifiers in lexical resource semantics. In Doug Arnold, Miriam Butt, Berthold Crysmann, Tracy Holloway King & Stefan Müller (eds.), *Joint 2016 Conference on Head-driven Phrase Structure Grammar and Lexical Functional Grammar*, 599–619. Stanford: CSLI Publications.
- Richter, Frank & Manfred Sailer. 2004. Basic concepts of Lexical Resource Semantics. In Arnold Beckmann & Norbert Preining (eds.), *ESSLLI 2003 – Course Material I*, vol. 5 Collegium Logicum, 87–143. Kurt Gödel Society Wien.
- Robaldo, Livio. 2010. On the maximalization of the witness sets in independent set readings. In *19th European Conference on Artificial Intelligence (ECAI 2010)*, vol. 215 Frontiers in Artificial Intelligence and Applications, 1133–1134. Amsterdam: IOS.
- Sailer, Manfred. 2004a. Local semantics in Head-driven Phrase Structure Grammar. In Olivier Bonami & Patricia Cabredo Hofherr (eds.), *Empirical Issues in Syntax and Semantics 5*, vol. 5 Empirical Issues in Formal Syntax and Semantics, 197–214.
- Sailer, Manfred. 2004b. Propositional Relative Clauses in German. In Stefan Müller (ed.), *11th conference on Head-Driven Phrase Structure Grammar (HPSG)*, 223–243. Stanford: CSLI Publications.
- Scha, Remko. 1981. Distributive, collective and cumulative quantification. In J. A. G. Groenendijk, T. M. V. Janssen & M. B. J. Stokhof (eds.), *Formal methods in the study of language*, vol. 2, 483–512. Amsterdam: Mathematisch Centrum.
- Schwarzschild, Roger. 1996. *Pluralities*. Dordrecht: Springer.
- Sher, Gila. 1997. Partially-ordered (branching) generalized quantifiers: A general definition. *Journal of Philosophical Logic* 26. 1–43.
- Sternefeld, Wolfgang. 1998. Reciprocity and cumulative predication. *Natural Language Semantics* 6(3). 303–337.

Symmetry and asymmetry in the Hebrew copula construction

Nurit Melnik 

The Open University of Israel

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 108–126

Keywords: HPSG, copula, agreement, inversion, Information Structure, Hebrew

Melnik, Nurit. 2018. Symmetry and asymmetry in the Hebrew copula construction. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 108–126. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.7. 
CC-BY 4.0

Abstract

The copula construction in Hebrew has received much attention in the linguistic literature. Nevertheless, one non-canonical variant has been largely neglected. In this variant the copula, flanked by two NPs, exhibits agreement with the post-copular NP, contrary to the canonical variant, where the agreement controller is the initial NP. This phenomenon challenges the notion of *subject* and its relation to agreement. The current corpus-based study investigates the word order and agreement patterns exhibited by the Hebrew copular construction and shows that their distribution is largely motivated by information structure considerations. The proposed analysis accounts for the syntactic symmetry and semantic asymmetry between the two NPs.

1 Overview

The copula construction in Hebrew has received much attention in the linguistic literature. Nevertheless, one non-canonical variant has been largely neglected. In this variant the copula, flanked by two NPs, exhibits agreement with the post-copular NP, contrary to the canonical variant, where the agreement controller is the initial NP. This construction, often referred to in the literature as ‘copula inversion’, poses challenges to the notion of *subject* and its relation to agreement in various and diverse languages.

This study proposes that two mechanisms are responsible for the licensing of the Hebrew NP–NP copula construction. First, alongside the general argument realization principle, a copula-specific rule reverses the mapping of ARG-ST members to VALENCE categories and allows for both NPs to function as either subject or complement. Second, copula inversion is argued to be an instance of a general V2 construction in Hebrew, where a clause-initial constituent triggers subject–verb inversion. This construction is shown to be motivated by information structure considerations. The two mechanisms account for the apparent symmetry between the two NPs. Nevertheless, there is no symmetry with respect to semantics; each NP maintains its semantic function as subject or predicate regardless of its linear position or syntactic role.

2 Background

The standard data items which appear in the literature on the Hebrew copula construction are given in (1).

- (1) dani (hu) more / nexmad / ba-xacer.
Danny Pron.3SM teacher.SM nice.SM in.the-yard
‘Danny is a teacher/nice/in the yard.’

The predicates consist of NPs, AdjPs, and PPs. The copula linking the subject and the predicate is homophonous with the 3rd person pronoun (hence the gloss) and agrees with the subject. The pronominal copula is only used in present tense, and is sometimes optional. In past and future tense an inflected form of the verb *haya* ‘be’ is obligatorily used (2). The present tense form of *haya* is missing from the paradigm.

- (2) dina hayta / tihiye [mora / nexmada / ba-xacer].
 Dina was.3SF will.be.3SF teacher.SF nice.SF in.the-yard
 ‘Dina was/will be a teacher/nice/in the yard.’

AdjP predicates obligatorily exhibit number-gender agreement with their subjects (e.g., *nexmad/nexmada* ‘nice’ in (1) & (2), respectively). With NP predicates, however, agreement is not imposed by the grammar. Rather, the agreement between the animate subject and NP predicate *more/mora* ‘teacher’ in (1) & (2) is due to sortal restrictions. This point is often overlooked, due to the preponderance of examples with animate (human) subjects in the literature. In (3), for example, there are agreement mismatches between the subject and two alternative predicates.

- (3) ha-sfarim ha-’ele hem matana / matanot
 the-books.PM the-these.PM Pron.3PM present.SF presents.PF
 mi-xaveray.
 from-my.friends
 ‘These books are a present/presents from my friends.’

The focus of this paper is on a different agreement domain, namely the agreement properties exhibited by the pronominal and verbal copulas. In an overwhelming majority of cases the copula agrees with the (clause-initial) subject.¹

- (4) [ha-merivot ha-kolaniyot ve-ha-mexo’arot] hen
 the-fights.PF the-loud.PF and-the-ugly.PF Pron.3PF
 [*ha-davar* ha-yaxid ha-me’anyen].
 the-thing.SM the-only.SM the-interesting.SM
 ‘The loud and ugly fights are the only thing that is interesting.’

¹Throughout this paper, the two NPs appear in square brackets, with the agreement controller underlined and the head of the other NP in italics.

There are, however, instances where the post-copular NP controls the agreement.

- (5) [*merivot* beyn axim] hi [*derex* me'ula
 fights.PF between siblings.PM Pron.3SF way.SF excellent.SF
 lehitkonen la-xayim].
 to.prepare to.the-life
 'Fights between siblings are an excellent way to prepare for life.'

This construction, often referred to in the literature as 'copula inversion', challenges the notion of *subject* and its relation to agreement: Is the post-copular NP the subject or is there non-subject agreement? As I explore this issue I refer to the two constituents by using the linear terms *NP1* and *NP2*.

3 Copula inversion

3.1 The copular construction

The copular construction is a clause type in which the predicate is not a verb, but rather an NP, AdjP or PP, and is often linked to the subject by a copula. Following Higgins's (1979) taxonomy, Mikkelsen (2005) illustrates three types of copular constructions:

- (6) PREDICATIONAL
 Ingrid Bergman is the lead actress in that movie.
- (7) SPECIFICATIONAL
 The lead actress in that movie is Ingrid Bergman.
- (8) EQUATIVE
 She is Ingrid Bergman.

Broadly speaking, in predicational clauses the predicate expresses a property of the referent of the subject. As such, subjects of predicational clauses are referring expressions. Conversely, in specificational clauses the post-copular expression is a referring expression which identifies the referent of the denotation of the syntactic subject (i.e., it answers the question of who is the lead actress). Equatives involve two referring expressions which are equated (the referent of *she* is the same individual denoted by *Ingrid Bergman*).

The relationship between predicational and specificational sentences is subject to much debate in the literature, primarily because they look like mirror images of each other. Indeed, this has raised the question of whether specificational sentences are instances of predicate raising; thus associating the role of subject with the referential argument. Such a role-reversal analysis in the context of the current discussion can naturally account for the

phenomenon of copular inversion; the agreement relation between the copula and the post-copular NP is a manifestation of subject–verb agreement, albeit in a non-canonical configuration. Nevertheless, agreement triggering may not be a necessary nor sufficient condition for subjecthood.

The following sections briefly present the phenomenon of copula inversion in a number of languages, specifically highlighting the questions raised above, namely the relationship between subjecthood, agreement, and word order.

3.2 Copula inversion in Catalan

Alsina & Vigo (2014) focus on copula inversion and non-subject agreement in Catalan and related languages (e.g., Spanish and Italian) and provide the following examples (Alsina & Vigo’s exx. 1&2).

- (9) a. [Els impostos] són [el problema].
the.PL taxes.PL be.PRES.3P the.SG problem.SG
‘The taxes are the problem.’
- b. [el problema] són [els impostos].
the.SG problem.SG be.PRES.3P the.PL taxes.PL
‘The problem is taxes.’
- c. * [el problema] és [els impostos].
the.SG problem.SG be.PRES.3S the.PL taxes.PL
- d. * [Els impostos] és [el problema].
the.PL taxes.PL be.PRES.3S the.SG problem.SG

As is suggested by these examples, agreement remains with the plural NP regardless of its position.

The analysis which Alsina & Vigo (2014) propose to account for the agreement patterns exhibited above is couched within their novel LFG approach to subject–verb agreement. Under their proposal, the agreement properties defined for a verb are not associated with a particular grammatical function, but defined in a special AGR feature. The values of this feature are unified with a grammatical function in the sentence, whose identity is determined by OT-like ranked constraints that implement a Person-Number hierarchy. This grammatical function may coincide with the subject, as is illustrated in (9a) but this is not necessarily so. In the copular inversion case, illustrated by (9b), the subject is NP1 and yet NP2 controls the agreement, since as a plural NP it is ranked higher in the hierarchy.

3.3 Reversed Equative *be* in English

Post-copular agreement is also found in English. Kay & Michaelis (2017a,b) discuss the Reversed Equative *be* construction where plural NP2s (optionally) control the agreement properties of the copula.

- (10) a. [My biggest *worry*] are [the injury risks].
 b. [My worst *nightmare*] were [the soups she would make for dinner].
 (Kay & Michaelis, 2017b, ex. 0.6)

Kay & Michaelis (2017a,b) argue that the Reversed Equative *be* construction is a subtype of the more general Split Subject construction, which includes constructions such as the various *there* constructions, Deditic Inversion (e.g., *Here comes the bus*), and Presentational Inversion (e.g., *On the porch stood marble pillars.*). These constructions combine special grammatical form with special discourse pragmatics. Grammatically, Kay & Michaelis argue, subject properties are split between the preverbal and postverbal arguments. While the postverbal NP controls verb agreement, the preverbal NP occupies the subject position and can undergo raising. From a discourse-pragmatic perspective, the postverbal constituent in all Split Subject constructions is in focus.

More technically, Kay & Michaelis's (2017b) formalization of this analysis in the Sign-Based Construction Grammar framework involves the distinction between the External Argument (XARG) and the Agreement Source. The Reversed Equative *be* construction is a subtype of the Split-Subject Construction and its single daughter is the Equative *be* Listeme with a plural XARG and a singular second ARG-ST list member. This derivational construction reverses the order of its daughter's ARG-ST list members, associates XARG with the first member of the new list, and retains the AGR specifications of the original *be* listeme. Consequently, NP1 is identified as the XARG and NP2 controls the agreement properties exhibited by the copula.

3.4 Hebrew non-canonical copula constructions

Doron (1983) in her comprehensive analysis of verbless predicates in Hebrew discusses a number of non-canonical copula constructions. One construction is the predicate-first construction, which is the mirror image of the canonical example given in (1).

- (11) nexmad / more hu dani. (Doron, 1983, ex. 51)
 nice teacher Pron.3SM Danny
 'Danny is nice/a teacher.'

In Doron's (1983) (transformational) system, this construction is derived by the predicate moving to adjoin INFL and the subject moving to an appositive (A') position (to satisfy the θ -criterion). Importantly, the agreement controller is the post-copular NP subject. This in essence is the gist of the predicate-raising analysis of specificational copular clauses mentioned in Section 3.1.

In addition, Doron (1983), citing Rubinstein (1968, p.137), discusses cases where the copula exhibits variable agreement. As an illustration of the two agreement options, she provides the following example (due to Emmon Bach).

- (12) [ma še-dekart katav] hu / hi [ha-hoxaxa
 what that-Descartes wrote Pron.3SM Pron.3SF the-proof.SF
 le-kiyumo]. (Doron, 1983, ex. 43)
 to-his.existence
 ‘What Descartes wrote was the proof of his existence.’

She claims that with NP1-agreement the sentence can be paraphrased as ‘what Descartes wrote proves his existence’,² whereas with NP2-agreement there is only an identity reading. More generally, Doron (1983, p.91) suggests that “AGR in nominal sentences agrees with the subject or the predicate, depending on which is ‘more referring’”. Nevertheless, she does not provide an analysis of the NP2-agreement variant.³

3.5 Interim summary

The cursory presentations of NP2-agreement phenomena in Catalan, English and Hebrew revealed different licensing conditions. In Catalan, when NP1 and NP2 differ in their number property the verb agrees with the plural NP, regardless of its position (Alsina & Vigo, 2014). In English, on the other hand, where NP2-agreement is licensed, canonical NP1-agreement is also possible. Nevertheless, NP2-agreement is pragmatically motivated; the copula may exhibit agreement with the postcopular NP provided that the NP is plural and focal (Kay & Michaelis, 2017a,b). Finally, in Hebrew, Doron (1983) suggests that agreement depends on referentiality; AGR agrees with the more referring NP.

The phenomenon of NP2-agreement in the copular construction certainly challenges the unmarked alignment between subjects and agreement controllers. Indeed, the analyses proposed by Alsina & Vigo (2014) and by Kay & Michaelis (2017a,b) explicitly involve the disassociation of subjecthood and agreement; the agreement controller in their systems is not necessarily the syntactic subject. While this phenomenon is not in the focus of Doron (1983), she too suggests that AGR in nominal sentences may agree with the subject *or* the predicate.

²Doron refers to the NP1-agreement paraphrase as ‘specificational’ yet the paraphrase she proposes is predicational.

³Hebrew has an additional pronominal copula, *ze*, which alternates between exhibiting agreement with NP2 or appearing in default form (Sichel, 1997, among others). A discussion of this construction is not in the scope of this paper.

4 Copula inversion in Hebrew: Corpus data

The discussions of the Hebrew copula construction in the literature are mostly based on made-up examples (e.g., 1-3). A corpus investigation revealed a much richer dataset with a non-negligible number of non-canonical constructions.⁴ Nevertheless, it is important to emphasize that NP2-agreement is the more marked variant; in each of the following examples an NP1-agreeing copula is the unmarked option.

4.1 NP2-agreement and cardinality

The English Reversed Equative *be* and the Catalan copula inversion construction were found to be sensitive to the *number* feature of the NPs. In the two languages NP2-agreement is restricted to cases where NP2 is plural. Hebrew, however, exhibits more variability; NP2-agreement occurs with plural NP2s (13), but also with singular NP2s, where NP1 is plural (14). The latter is claimed to be an ungrammatical configuration in English and Catalan. In fact, all four agreement options illustrated in (9) for Catalan are possible in Hebrew.

- (13) [*makor* *tov* *le-sidan*] *hem* [*mucarey* *he-xalav*
 source.SM good.SM for-calcium Pron.3PM products.PM the-milk
 ha-šonim].
 the-different.PM
 ‘A good source of calcium is the different milk products.’

- (14) [*nehagim* *ayefim*] *hi* [*be'aya* *globalit* *xamura*].
 drivers.PM tired.PM Pron.3SF problem.SF global.SF serious.SF
 ‘Tired drivers are a serious global problem.’

4.2 NP2-agreement and reference

The choice between the two agreement patterns is attributed by Doron (1983) to semantics. She predicts that NP2-agreement occurs when NP2 is the more referring argument. This is indeed the case with (15), where NP2 is a proper noun, but not with (16), where the post-copular agreement controller is predicationa (and indefinite). Thus, we find NP2-agreement with both specificational and predicationa sentences.

⁴All the examples in the following sections are retrieved from *heTenTen 2014*, a billion-token web-crawled Hebrew corpus (Baroni et al., 2009).

- (15) [*dugma* le-tocar šel ha-tkufa] hu [beyt
example.SF of-product.SM of the-era.SF Pron.3SM house.CS.SM
akiva be-rexov hercel].
Akiva in-street Herzl
‘An example of a product of this era is Akiva House on Herzl Street.’
- (16) omnam naxon ha-davar ki [*ha-xumca* ha-hyaluronit]
indeed true the-thing that the-acid.SF the-hyaluronic.SF
hu [muca ha-mešameš ke-xomer miluy...]
Pron.3SM product.SM that-used.SM as-substance filling
‘Indeed it is true that hyaluronic acid is used as a filling substance...’

4.3 NP2-agreement and semantic functions

Syntactically, the NP–NP copula clause exhibits full symmetry: each NP can appear in either position and the copula can agree with either NP. This is not the case with the NPs’ semantic functions: regardless of word order, it is always the same NP that is the predicate of the other.⁵ This asymmetry is evident when the *consider*-test is applied.

Consider the copular construction in (17). Its NP1 and NP2 can feature as the two complements of a *consider*-like Hebrew construction (18). Nevertheless, unlike the copular construction, the order of the complements of *ro’a* ‘see’ is fixed: the semantic subject must precede the semantic predicate; the reversed order is ungrammatical. Thus, agreement in (17) is with NP2, which is the semantic predicate.

- (17) eclenu ba-mišpaxa [*haskala*] haya [davar hexraxi
for.us in.the-family education.SF was.3SM thing.SM essential.SM
ve-bsisi].
and-basic.3SM
‘For us in my family education was an essential and basic thing.’
- (18) a. ani ro’a be- [*haskala*] [davar hexraxi
I see.SF in- education.SF thing.SM essential.SM
ve-bsisi].
and-basic.SM
‘I consider education an essential and basic thing.’
- b. *ani ro’a be- [davar hexraxi ve-bsisi]
I see.SF in- thing.SM essential.SM and-basic.SM
[*haskala*].
education.SF

⁵Equative sentences with two referential NPs (e.g., *Cicero is Tully* or *Danny is Mr. Cohen*) are not easy to find in a corpus.

The *consider* test applied to the example in (15) above, also an instance of NP2-agreement, reveals that in this case the copula agrees with the post-copular semantic subject.

- (19) a. *ani ro'a be- [dugma le-tocar šel ha-tkufa]
 I see.SF in- example.SF of-product.SM of the-era.SF
 [beyt akiva].
 house.CS.SM Akiva
- b. ani ro'a be- [beyt akiva] [dugma
 I see.SF in- house.CS.SM Akiva example.SF
 le-tocar šel ha-tkufa].
 of-product.SM of the-era.SF
 'I consider Akiva House an example of a product of this era.'

As for the rest of the NP2-agreement examples presented above, the *consider* test shows that NP2 is the semantic subject in (13) and the semantic predicate in (5), (14) & (16).

4.4 Symmetry and asymmetry

Corpus-based data regarding the distribution of NP2-agreement in the Hebrew copular construction suggest that this construction is not subject to the constraints identified for its English and Catalan counterparts. First, cardinality does not seem to play a role in the licensing of NP2-agreement. Moreover, instances of copula agreement with NP2 were attested with referring and non-referring arguments. Syntactically, the NP–NP copula clause exhibits full symmetry: each NP can appear in either position and the copula can agree with either NP. Nevertheless, from a semantic perspective, the relationship between the two NPs is asymmetrical: regardless of word order or agreement pattern, it is always the same NP that is the predicate of the other.

5 Triggered inversion and copula clauses

I propose that NP2-agreement clauses are instances of a construction referred to in the literature as *triggered inversion* (Shlonsky & Doron, 1992). Although the unmarked word of Hebrew clauses is SV(O), the language also licenses a construction in which, similarly to V2 constructions in other languages, a clause-initial constituent triggers subject–verb inversion.⁶ A corpus example of a triggered inversion construction is given in (20a) and its constructed SVO counterpart is (20b).

⁶It should be noted that subject–verb inversion in this cases is not obligatory.

- (20) a. [et ha-toxnit] movil ha-misrad le-haganat
 ACC the-project.SF leading.SM the-ministry.SM for-protection
 ha-sviva.
 the-environment
- b. ha-misrad le-haganat ha-sviva movil
 the-ministry.SM for-protection the-environment leading.SM
 [et ha-toxnit].
 ACC the-project.SF
 ‘The ministry of environmental protection is leading the project.’

The SVO variant is clearly the unmarked option, whereas the inverted example is only felicitous in a context where a particular project is salient in the discourse. The NP *et ha-toxnit* ‘the project’ is proposed to form a link to this discourse. The new information contributed by the sentence is the identity of the leader of the project. In accordance with the principle of “new information comes last”, the NP which denotes this participant is inverted to appear post-verbally.

5.1 New information comes last

Many instances of NP2-agreement exhibit the same information structure properties that characterize the triggered inversion construction discussed in the previous section. In these instances NP1 serves as a link to the previous discourse and NP2 provides the new information. Indeed, in isolation, NP2-agreement clauses do not always sound perfectly grammatical. Some speakers would even label them as performance errors or instances of extra-grammatical “attraction”. Yet, these clauses appear in written (possibly proofread and/or edited) texts of diverse registers. Moreover, in many cases of NP2-agreement the distance and material between the head of NP1 and the copula are not substantial enough to cause distraction or accidental mismatches.

A discourse excerpt illustrating the licensing conditions of this construction is given in (21).

- (21) a. *aval anaxnu lo ro'im ba-mitnaxalim et šoreš*
 but we not see.PM in.the-settlers ACC root
ha-be'aya...
the-problem...
 'But we don't consider the settlers the root of the problem...'
- b. ...[*šoreš* *ha-be'aya*] *hem* [*gufey* *ha-šilton*
 root.CS.SM the-problem Pron.3PM bodies.PM the-regime
ha-yisra'elim še-menahalim et ha-mediniyut ha-zu].
 the-israeli that-maintain ACC the-policy the-this
 '...The root of the problem is the Israeli governing bodies who
 maintain this policy.'

The copular inversion sentence in (21b) is felicitous due to the information contributed by sentence (21a), which precedes it. In the copular construction the semantic predicate *šoreš ha-be'aya* 'the root of the problem' is preposed to a clause-initial position and functions as a link to the topic of the previous discourse, namely what is the root of the problem. The subject, which constitutes the new information is postposed (or, in other words, inverted with the copula). Similarly to all instances of triggered inversion in Hebrew, the post-verbal argument is the agreement controller.

Corpus searches retrieve many instances of copular clauses with NP2-agreement where the head of NP1 is modified by the adjective *nosaf* 'additional'. Two examples are give in (22).

- (22) a. [*dugma* *nosefet*] *hu* [*ha-mešorer* *yicxak*
 example.SF additional.SF Pron.3SM the-poet.SM Yitzhak
la'or].
 Laor
 'An additional example is the poet Yitzhak Laor.'
- b. [*bonus* *nosaf*] *hem* [*ha-kisim* *šel ha-simla*].
 bonus.SM added.SM Pron.3PM the-pockets.PM of the-dress
 'An added bonus is the pockets of the dress.'

Expressions such as *additional example* or *additional bonus* can only be felicitous in a context where other examples or bonuses were already mentioned. Thus, their preposing is well motivated. Moreover, here too, the new information is supplied by NP2, which in this case is the semantic subject and the agreement controller.

While many instances of NP2-agreement with *additional* NP1 were found in the corpus the alternative pattern where the copula agrees with the *additional* NP1 were also found.⁷ One such example is given in (23).

⁷A quantitative assessment of this distribution as well as the distribution of other alternations is left for future research.

- (23) [dugma nosefet] hi [*ha-moniyot* ha-carfatiyot
 example.SF additional.SF Pron.3SF the-taxi.PF the-French.PF
 ha-'atikot be-kahir].
 the-antique.PF in-Cairo
 'An additional example is the antique French taxis in Cairo.'

As will be discussed in Section 6, the availability of the two agreement patterns is a particularly challenging aspect of this construction.

5.2 Contrastive focus

An additional function which triggered inversion constructions fulfill is the expression of *contrastive focus*. Consider the example in (24).

- (24) [et ha-tik šela] macati be-megirat ha-garbayim aval [et
 ACC the-bag of.her found.1S in-drawer the-socks but ACC
 ha-maclema] bal'a ha-'adama.
 the-camera swallowed.3SF the-ground.SF
 'I found her bag in the sock drawer but the camera vanished (literally:
 the ground swallowed the camera).'

The speaker contrasts the results of his/her search for two items: a bag and a camera. The NPs denoting the two items are preposed to the clause-initial position of their respective clauses. The subject of the first conjunct is *pro*-dropped, whereas the second clause is an instance of triggered inversion: the subject, *ha-'adama* 'the ground' appears post-verbally.

Similar contrastive pairs are also found in the copular construction, whereby the contrasted element is fronted and the copula exhibits NP2-agreement. Consider the example in (25).

- (25) [*ha-tokfan*] hem [mimšelet yisra'el u-mimšal
 the-aggressor.SM Pron.3PM government Israel and-regime
 xamas ve-šutafav be-aza] ve'ilu [*ha-korban*]
 Hamas and-its.partners.PM in-Gaza whereas the-victim.SM
 hem [tošavey aza ve-tošavey medinat
 Pron.3PM inhabitants.PM Gaza and-inhabitants.PM state
 yisra'el].
 Israel
 'The aggressor is the Israeli government and the Hamas regime and
 its partners in Gaza, whereas the victim is the inhabitants of Gaza
 and the inhabitants of the state of Israel.'

The sentence clearly contrasts the aggressor with the victim. The contrast is expressed by fronting the NPs denoting each "role" to their respective clause-initial position and inverting the subject and copula, while maintaining their agreement relationship.

The post-copular NPs in (25) are the agreement controllers and the semantic subjects. There are, however, also instances of contrastive focus with NP2-agreement where the preposed NP is the semantic subject and NP2 is the semantic predicate. One such example is given in (26).

- (26) *[sidur ha-šulxan] hi [ha-teritya ha-bil'adit*
 setting.CS.SM the-table.SM Pron.SF the-territory.SF the-exclusive
šeli]. [be-noga'a la-tafrit] le'umat zot ani menahelet
 my with-regards to.the-menu contrastively I hold
diyunim nokvim im modi.
 discussions profound with Modi
 'Setting the table is my exclusive territory. With regards to the
 menu, on the hand hand, I hold profound discussions with Modi.'

In this case the speaker contrasts duties related to the organization of a dinner: setting the table and deciding on the menu. The speaker assumes sole responsibility over the former, while asserting that she shares the responsibility over the latter with another person named Modi. In the two clauses the contrasted items are preposed.

6 Formalization

The analysis proposed here assumes that the copula in an NP–NP clause selects an NP subject and an NP complement.⁸ However, unlike the “standard” HPSG raising analysis of the copula (Pollard & Sag, 1994, p. 147), predication in this case is only semantic. The semantic predicate does not select the semantic subject as its syntactic subject and does not “pass” this requirement to the copula. An abbreviated description of the argument structure of the copula is given in (27).

- (27) Canonical argument realization of the copula
- $$\left[\begin{array}{l} \textit{canonical-cop} \\ \text{VAL} \quad \left[\begin{array}{l} \text{SUBJ} \quad \langle [1] \rangle \\ \text{COMPS} \quad \langle [2] \rangle \end{array} \right] \\ \text{ARG-ST} \quad \left\langle [1]\text{NP} \left[\text{INDEX } [3] \right], [2]\text{NP} \left[\begin{array}{l} \text{CAT} | \text{HEAD} | \text{PRED} + \\ \text{CONT} | \text{RELS} \quad \langle [\text{ARG1 } [3]] \rangle \end{array} \right] \right\rangle \end{array} \right]$$

The canonical copular construction is structured similarly to transitive clauses. The copula combines with its complement (the semantic predicate)

⁸This analysis is not compatible with a previous HPSG analysis of nonverbal predicates in Hebrew (Haugereid et al., 2013).

in a *hd-comp* phrase and this phrase, in turn, combines with the semantic and syntactic subject in a *hd-subj* phrase. Subject–copula agreement is constrained by general principles regarding the *hd-subj* phrase type.

Let us illustrate this by considering a pair of examples. Example (22b), repeated here as (28a) is an inverted construction. Its constructed SVO counterpart is given in (28b).

- (28) a. [*bonus* *nosaf*] hem [ha-kisim šel ha-simla].
 bonus.SM added.SM Pron.3PM the-pockets.PM of the-dress
- b. [ha-kisim šel ha-simla] hem [*bonus* *nosaf*].
 the-pockets.PM of the-dress Pron.3PM bonus.SM added.SM
- ‘An added bonus is the pockets of the dress.’

An abbreviated analysis of the SVO variant in (28b) is given in Figure 1. The inverted construction in (28a) is licensed by a *hd-filler* phrase in which the filler-daughter is the syntactic complement of the copular. An abbreviated tree representation of (28a) is given in Figure 2. Note that the syntactic structure of this construction is identical to that of the productive triggered inversion construction (e.g., 20a & 24).

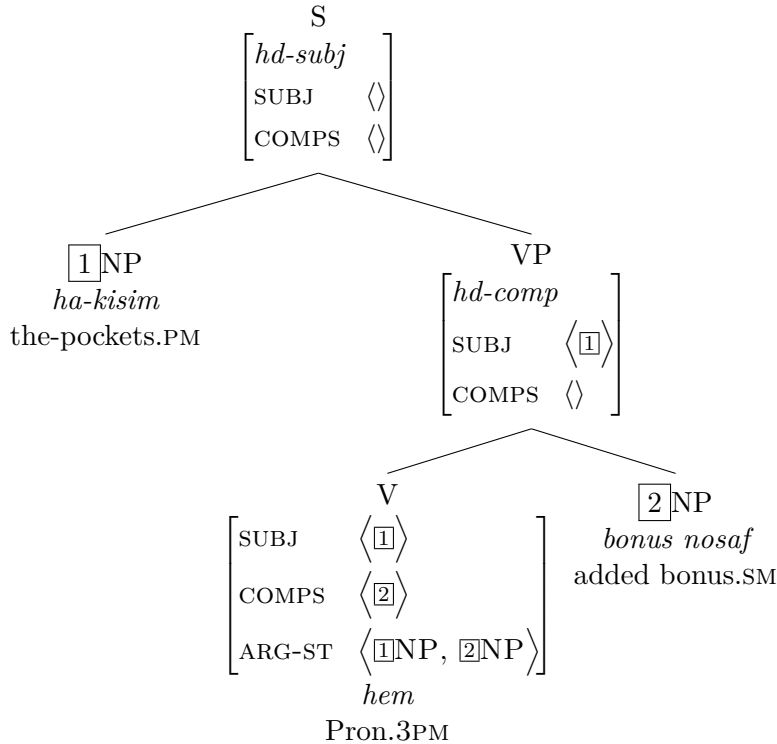


Figure 1: Canonical copular construction

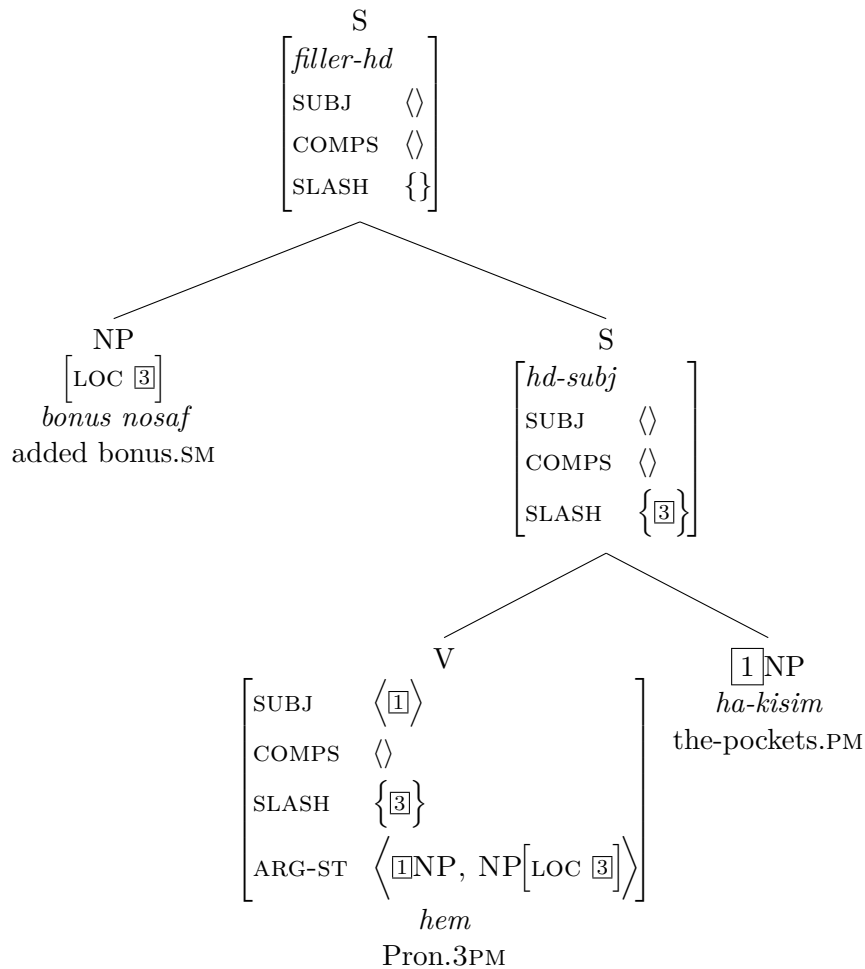


Figure 2: Inverted copular construction

The analyses sketched above account for cases where the copula agrees with the first element in ARG-ST, which is both the syntactic subject and the semantic subject. The data, however, revealed that in the NP-NP copular construction the copula may also agree with the semantic predicate, regardless of its position. Table 1 summarizes the four attested word order and agreement patterns, along with reference to an example sentence of each pattern.

The canonical argument realization of the copula described in (27) above, along with the optional triggered inversion construction account for the patterns described in the first row of the table: NP1-agreement with the semantic subject is the unmarked pattern for all (SVO) clauses, and NP2-agreement occurs when the predicate/complement is preposed and the subject is inverted with the copula. There is nothing surprising about these

	NP1-agreement	NP2-agreement
Semantic Subject	(4)	(22b)
Semantic Predicate	(23)	(5)

Table 1: Word order and agreement patterns

patterns.

The second row, however, poses a challenge. As was shown above, the copula was found to exhibit agreement with the semantic predicate in its clause-initial position as well as when it appears post-verbally. Moreover, a similar information structure function, namely the expression of contrastive focus, was shown to motivate the preposing of the semantic subject in (25) and the semantic predicate in (26).

To resolve this conflict I distinguish between syntactic and semantic predication by allowing NPs which are the semantic predicates to function as the syntactic subjects. This, I suggest, is due to the special status of NPs, which are compatible with the two functions. In formal HPSG terms, a lexical rule reverses the “default” mapping between ARG-ST and VALENCE list members, so that the semantic predicate is mapped to SUBJ and the semantic subject to COMP (29). This rule is conceptually similar to the derivation construction proposed by Kay & Michaelis (2017a,b) for the Reversed Equative *be* in English.

(29) Non-canonical argument realization of the copula

$$\left[\begin{array}{l} \textit{non-canonical-cop} \\ \text{VAL} \left[\begin{array}{ll} \text{SUBJ} & \langle \boxed{2} \rangle \\ \text{COMPS} & \langle \boxed{1} \rangle \end{array} \right] \\ \text{ARG-ST} \left\langle \boxed{1}\text{NP} \left[\text{INDEX } \boxed{3} \right], \boxed{2}\text{NP} \left[\begin{array}{l} \text{CAT} \mid \text{HEAD} \mid \text{PRED} + \\ \text{CONT} \mid \text{RELS } \langle \boxed{1}\text{NP} \left[\text{INDEX } \boxed{3} \right] \rangle \end{array} \right] \right\rangle \end{array} \right]$$

The non-canonical argument realization preserves the semantic relation between the two arguments while building on existing mechanisms for licensing subject–verb agreement and inverted constructions. Thus, when a non-canonical copula heads a canonical SVO clause the copula exhibits subject–verb agreement with NP1, which is the semantic predicate (e.g., 23). Conversely, the non-canonical copula can also head a triggered inversion construction. In this case, too, the copula agrees with the semantic predicate, which is its syntactic subject (e.g., 5). Thus, the canonical and non-canonical argument realization rules, together with the two alternative clause structures account for the four patterns exhibited by the data and summarized in Table 1.

To summarize, triggered inversion coupled with two alternative mappings of argument structure elements account for the different variations of the copula construction and capture the syntactic symmetry and semantic asymmetry between the two NPs. Moreover, an information-structure account explicates the motivation behind these variations.

References

- Alsina, Alex & Eugenio M Vigo. 2014. Copular inversion and non-subject agreement. In Miriam Butt & Tracy Holloway King (eds.), *Proceedings of the LFG2014 conference*, 5–25. Stanford, CA: CSLI Publications.
- Baroni, Marco, Silvia Bernardini, Adriano Ferraresi & Eros Zanchetta. 2009. The WaCky wide web: A collection of very large linguistically processed web-crawled corpora. *Language Resources and Evaluation* 43(3). 209–226.
- Doron, Edit. 1983. *Verbless predicates in Hebrew*. University of Texas at Austin dissertation.
- Haugereid, Petter, Nurit Melnik & Shuly Wintner. 2013. Nonverbal predicates in Modern Hebrew. In Stefan Müller (ed.), *Proceedings of the 20th international conference on head-driven phrase structure grammar, Freie Universität Berlin*, 69–89. <http://cslipublications.stanford.edu/HPSG/2013/hmw.pdf>.
- Higgins, Roger Francis. 1979. *The pseudo-cleft construction in English*. New York: Garland.
- Kay, Paul & Laura Michaelis. 2017a. Partial inversion in english. In Stefan Müller (ed.), *Proceedings of the 24th international conference on head-driven phrase structure grammar, University of Kentucky, Lexington*, 212–216. Stanford, CA: CSLI Publications. <http://cslipublications.stanford.edu/HPSG/2017/hpsg2017-kay-michaelis.pdf>.
- Kay, Paul & Laura A. Michaelis. 2017b. Partial inversion in English. Unpublished ms., Stanford University and University of Colorado Boulder.
- Mikkelsen, Line. 2005. *Copular clauses: Specification, predication and equation*, vol. 85. Amsterdam/Philadelphia: John Benjamins Publishing Company.
- Pollard, Carl & Ivan A. Sag. 1994. *Head-driven Phrase Structure Grammar*. University of Chicago Press and CSLI Publications.
- Rubinstein, Eliezer. 1968. *Ha-mishpat ha-shemani (The nominal sentence)*. Merhavia: Ha-Kibbutz Ha-Me’uxad. In Hebrew.

- Shlonsky, Uri & Edit Doron. 1992. Verb second in Hebrew. In D. Bates (ed.), *Proceedings of the tenth west coast conference on formal linguistics*, 431–446. Stanford, CA: CSLI Publications.
- Sichel, Ivy. 1997. Two pronominal copulas and the syntax of Hebrew nonverbal sentences. In R.C. Blight & M.Y. Moosally (eds.), *Texas linguistics forum 38: The syntax and semantics of predication*, vol. 38, 295–306. Dept. of Linguistics, University of Texas at Austin.

HPSG analysis and computational implementation of Indonesian passives

David Moeljadi 

Nanyang Technological University

Francis Bond 

Nanyang Technological University

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 127–137

Keywords: HPSG, Indonesian, passives, INDRA

Moeljadi, David & Francis Bond. 2018. HPSG analysis and computational implementation of Indonesian passives. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 127–137. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.8.

 CC-BY 4.0

Abstract

This study aims to analyze and develop a detailed model of syntax and semantics of passive sentences in standard Indonesian in the framework of Head-Driven Phrase Structure Grammar (HPSG) (Pollard & Sag, 1994; Sag et al., 2003) and Minimal Recursion Semantics (MRS) (Copestake et al., 2005), explicit enough to be interpreted by a computer, focusing on implementation rather than theory. There are two main types of passive in Indonesian, following Sneddon et al. (2010, pp. 256-260) and Alwi et al. (2014, pp. 352-356), called ‘passive type 1’ (P1) and ‘passive type 2’ (P2). Both types were analyzed and implemented in the Indonesian Resource Grammar (INDRA), a computational grammar for Indonesian (Moeljadi et al., 2015).

1 Introduction

A passive is a semantically transitive (two-participant) clause. Typically, the agent is either omitted or demoted to an oblique role, the other core participant possesses all properties of subjects, and the verb possesses formal properties of intransitive verbs (Payne, 2008, p. 204). Passive constructions are far more frequent in Indonesian than in English; an Indonesian passive is often naturally translated into English by an active construction (Sneddon et al., 2010, pp. 256, 263-264). Passive constructions in Indonesian are used in imperatives and for politeness, as well as in relative clauses which can only relativize subjects on defining relative clauses.

Research on Indonesian passives has been done by many linguists, such as McCune (1979), Voskuil (2000), Arka & Manning (2008), Cole et al. (2008), and Nomoto (2013). There has been a lot of linguistic work on Indonesian voice, in particular the status of passive-like structures in Indonesian and Austronesian languages (Musgrave, 2001; Riesberg, 2014). However, to the best of our knowledge, no research on Indonesian passives has been done in the HPSG framework. Our analysis is implemented in the Indonesian Resource Grammar (INDRA), a computational grammar for Indonesian (Moeljadi et al., 2015), which can parse and generate sentences.

There are two main types of passive in Standard Indonesian,¹ following Sneddon et al. (2010, pp. 256-260) and Alwi et al. (2014, pp. 352-356). They are called ‘passive type 1’ (P1) and ‘passive type 2’ (P2).² Both types are available for monotransitive and ditransitive verbs. They promote an object to subject. If there are two objects in an active ditransitive clause, only the one immediately following the verb (which has semantic role as patient or recipient) can be promoted to subject of the passive (Sneddon et al., 2010, p. 260).³ P1 and P2 are in (near) complementary

¹Indonesian is a diglossic language. This paper only deals with the ‘High’ variety of Indonesian, also known as the standard or formal Indonesian.

²Other types such as passives with prefix *ter-* and circumfix *ke-...-an* have not been analyzed and implemented in INDRA. They are for future work.

³This study only describes passives for monotransitive verbs. However, the analysis proposed here can be applied to ditransitive verbs as well.

distribution. P1 takes only a third person agent, while P2 may take first, second, and third person agent. P1 and P2 overlap with respect to the third person agent (Sneddon et al., 2010, p. 256).

2 Basic Data

2.1 Passive type 1

The verb in P1 is morphologically built by attaching a prefix *di-* to a transitive verbal stem (lexeme) in the lexicon. The subject (which usually has semantic role as agent) in the active sentence becomes an optional complement, immediately follows the passive verb (post-verbal), and it is optionally marked by a semantically empty preposition *oleh* ‘by’. Its PERNUM is third person, i.e. pronoun *dia* ‘3SG’, *mereka* ‘3PL’, enclitic *=nya* ‘3SG’, (common) noun, or proper name (Sneddon et al., 2010, p. 256-257). The position of the components of the predicate, such as auxiliaries and temporal markers, as well as the negative word *tidak* ‘NEG’ remain unchanged, i.e. they immediately precede the verb predicate both in active and passive voice.

Example (1a) shows a transitive sentence in active voice.⁴ An aspect marker *sudah* ‘PRF’ immediately precedes the active voice verb *menjemput* ‘ACT-pick.up’. Its corresponding P1 constructions are shown in Example (1b) to (1e). The position of the aspect marker is the same in all example sentences in (1). Example (1b), (1c), and (1d) show the optional preposition *oleh* ‘by’. Example (1c) and (1d) show that the enclitic *=nya* ‘3SG’ can attach directly to the passive verb or to the preposition *oleh* ‘by’. Example (1e) shows that a P1 construction may occur without a complement.

- (1) a. *Dia sudah menjemput Budi.*
 3SG PRF ACT-pick.up Budi
 ‘He has met Budi.’ (lit. ‘He has picked Budi up.’) (based on Sneddon et al., 2010, p. 256)
- b. *Budi sudah dijemput (oleh) dia.*
 Budi PRF PASS-pick.up by 3SG
 ‘Budi has been picked up by him.’ (based on Sneddon et al., 2010, p. 257)
- c. *Budi sudah dijemputnya.*
 Budi PRF PASS-pick.up=3SG
 ‘Budi has been picked up by him.’ (based on Sneddon et al., 2010, p. 257)

⁴A number of nasalization (sound changes) or morphology process occurs when *meN-* ‘ACT’ combines with stems, listed up in Moeljadi et al. (2015).

- d. Budi sudah dijemput olehnya.
 Budi PRF PASS-pick.up by=3SG
 'Budi has been picked up by him.' (based on Sneddon et al., 2010, p. 257)
- e. Budi sudah dijemput.
 Budi PRF PASS-pick.up
 'Budi has been picked up.' (based on Sneddon et al., 2010, p. 257)

In a coordinative construction with two or more passive verbs, the agent (both full forms and the bound form or enclitic =*nya*) can appear only once, following the last passive verb, as shown in (2).

- (2) Budi sudah ditunggu dan dijemputnya.
 Budi PRF PASS-wait and PASS-pick.up=3SG
 'Budi has been waited and picked up by him.'

2.2 Passive type 2

The verb in P2 is morphologically built by not attaching any affixes to a transitive verb lexeme in the lexicon. The verbs appear in bare stem form. Different from P1, the subject (agent) in the active sentence becomes an obligatory complement (argument), immediately preceding the verb (pre-verbal), without any prepositions such as *oleh* 'by'. The agent is a pronoun such as *aku* '1SG', *engkau* '2SG', *dia* '3SG' etc. or 'pronoun substitute', i.e. kinship terms such as *bapak* 'father', *ibu* 'mother', and personal names which can refer to the addressee, meaning 'you', or to the speaker, meaning 'I' (Sneddon et al., 2010, pp. 257, 259). No other component of the clause, such as negative and temporal marker, can come between the NP agent and the P2 verb (Sneddon et al., 2010, p. 258). They must occur before the agent.

Example (3) shows the corresponding P2 construction of Example (1a). The aspect marker *sudah* 'PRF' precedes the agent *dia* '3SG'.

- (3) Budi sudah dia jemput.
 Budi PRF 3SG pick.up
 'Budi has been picked up by him.' (based on Sneddon et al., 2010, p. 257)

If the agent is *aku* '1SG' or *engkau* '2SG', the bound forms (also called as 'proclitics' by some grammarians) *ku-* '1SG' and *kau-* '2SG' usually occur (Sneddon et al., 2010, p. 258), as shown in (4).

- (4) Budi sudah kujemput.
 Budi PRF 1SG-pick.up
 'I have met Budi.' (lit. 'Budi has been picked up by me.') (based on Sneddon et al., 2010, p. 257)

$$(6) \quad \left[\begin{array}{l} \text{INPUT} \quad \left\langle X, \left[\begin{array}{l} \text{lexeme (tr-verb-lex)} \\ \text{ARG-ST} \left\langle \left[\dots \right], \left[\dots \right] \right\rangle \end{array} \right] \right\rangle \\ \\ i\text{-rule :} \\ \text{OUTPUT} \quad \left\langle (\text{di-/ku-/kau-})X, \left[\begin{array}{l} \text{word (passive-transitive-lex-item)} \\ \text{ARG-ST} \left\langle \left[\begin{array}{l} \text{INDEX} \quad \boxed{1} \\ \text{ICONS-KEY} \quad \boxed{3} \end{array} \right], \left[\begin{array}{l} \text{INDEX} \quad \boxed{2} \\ \text{ICONS-KEY} \quad \boxed{4} \end{array} \right] \right\rangle \\ \text{LKEYS.KEYREL} \quad \left[\begin{array}{l} \text{ARG1} \quad \boxed{2} \\ \text{ARG2} \quad \boxed{1} \end{array} \right] \\ \text{ICONS} \left\langle \left[\begin{array}{l} \text{!} \quad \boxed{3} \quad \left[\begin{array}{l} \text{focus-or-topic} \\ \text{IARG2} \quad \boxed{1} \end{array} \right] \end{array} \right], \left[\begin{array}{l} \boxed{4} \quad \left[\begin{array}{l} \text{non-topic} \\ \text{IARG2} \quad \boxed{2} \end{array} \right] \end{array} \right] \right\rangle \end{array} \right] \right\rangle \end{array} \right]$$

In a coordinative construction with two or more passive verbs, the bound forms usually occur before each passive verb, as shown in (5a).

- (5) a. Budi sudah kutunggu dan kujemput.
 Budi PRF 1SG-wait and 1SG-pick.up
 ‘Budi has been waited and picked up by me.’
 b. ??Budi sudah kutunggu dan jemput.
 Budi PRF 1SG-wait and pick.up

3 Analysis

We treat passive as an inflectional rule, as shown in (6). The input is a lexeme, of type *tr-verb-lex*, which has two arguments. The output is a word, of type *passive-transitive-lex-item* which adds the semantic information for passives, i.e. its ARG1 is coindexed with the ARG0 of the complement (agent) and its ARG2 with the subject. The prefix *di-*, *ku-*, or *kau-* may be attached. Following Song (2017, pp. 211-214), we added information in the ICONS. The promoted argument or the subject is marked as *focus-or-topic*, while the demoted argument is marked as *non-topic*.

We treat *ku-* ‘1SG’, *kau-* ‘2SG’, and *=nya* ‘3SG’ differently because of the difference in their occurrence in coordinative constructions and their optionality. Following Zwicky & Pullum (1983) who distinguish clitics from inflectional affixes, we tokenize *=nya*, treating it as a word which belongs to a type *encl-3pers*. One of the reasons is because *=nya* can attach both to the verb or to a preposition. On the other hand, we do not tokenize *ku-* and *kau-* and treat them as inflectional affixes.

We made four lexical rules for P1 and P2, as shown in Figure 1. The first rule is for P1 (having an optional complement) without *oleh* ‘by’ and the second one

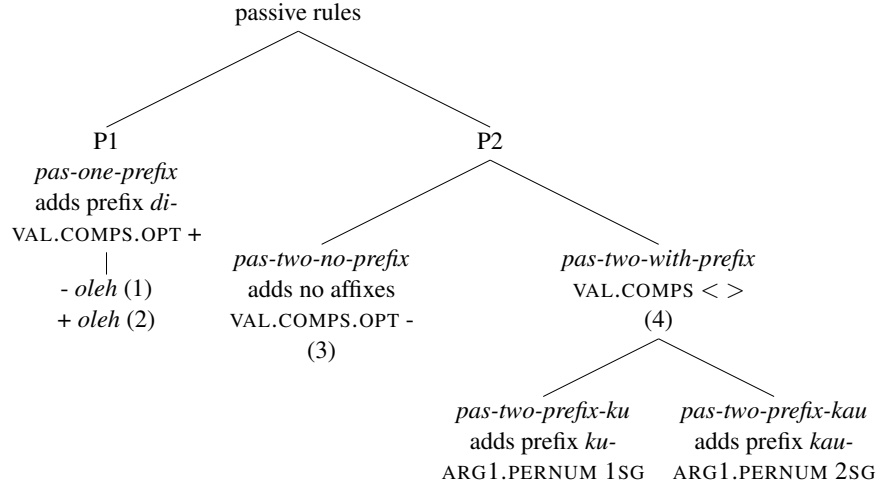


Figure 1: Type hierarchy for passive lexical rules

is with *oleh*. The third rule is for P2 (having an obligatory complement) without affixes and the fourth one is for P2 with a saturated complement and a prefix *ku-* or *kau-*. The details of each rule will be discussed in the next section.

3.1 Passive type one

We define a rule for P1, called *pas-one-prefix*. It is a rule which adds a prefix *di-* ‘PASS’. It inherits from the inflectional rule in (6). The output is a word, of type *passive-one-verb-lex*, with an optional complement, which inherits from *passive-transitive-lex-item*. It contributes the HEAD value, which is of type *pass1*. The COMPS has one item as its value. It has a feature POSTHEAD whose value is plus, as shown in (7). The COMPS’s HEAD is of type *pass1agent*. Its type hierarchy is shown in Figure 2.

$$(7) \left[\begin{array}{l} \text{HEAD } \textit{pass1} \\ \text{VAL.COMPS } \langle \boxed{1} \rangle \\ \text{ARG-ST } \left\langle \left[\begin{array}{l} \text{HEAD } \textit{subj-noun} \\ \text{VAL } \left[\begin{array}{l} \text{SPR } \langle \rangle \\ \text{COMPS } \langle \rangle \end{array} \right] \end{array} \right], \boxed{1} \left[\begin{array}{l} \text{HEAD } \textit{pass1agent} \\ \text{POSTHEAD } + \\ \text{VAL } \left[\begin{array}{l} \text{SPR } \langle \rangle \\ \text{COMPS } \langle \rangle \end{array} \right] \end{array} \right] \right\rangle \end{array} \right]$$

The parse tree of (1c) is shown in Figure 3. It shows the *pas-one-prefix* rule changes the lexeme *jemput* ‘pick.up’ to an inflected passive word *dijemput* ‘PASS-pick.up’. The inflected passive word is combined with its optional complement

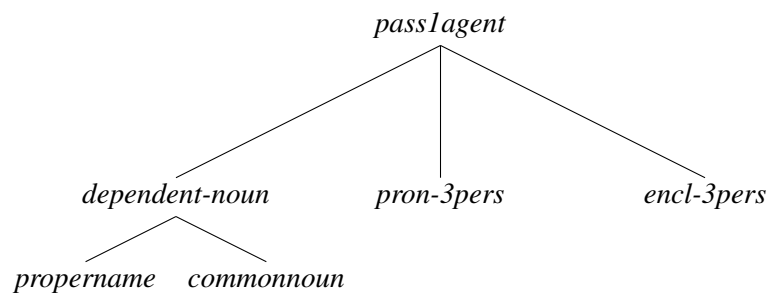


Figure 2: Type hierarchy for P1 agent

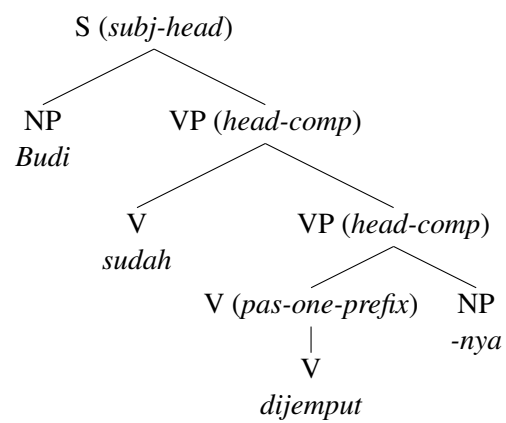


Figure 3: Parse tree of *Budi sudah dijemputnya*

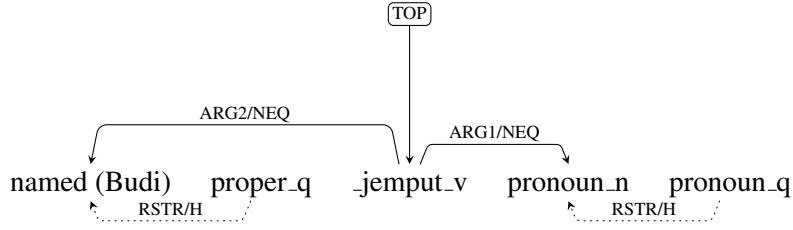


Figure 4: DMRS of *Budi sudah dijemputnya*

via *head-comp* rule. The semantics of the passive sentences in examples (1b) to (1d) look very much like the semantics of their active sentence counterpart in (1a), as shown in Figure 4, with additional information on the information structure. The ARG1 is linked to the optional agent complement and the ARG2 linked to the subject.

We treat *oleh* ‘by’ as a semantically empty preposition. It adds nothing to the meaning except the information that the COMPS of the passive verb is coindexed with the one of *oleh*, as shown in (8). The semantics of the PP headed by *oleh* is identical to that of *oleh*’s NP complement.

$$(8) \left[\begin{array}{c} \text{HEAD} \left[\begin{array}{c} \text{oleh-adp} \\ \text{MOD.LOC.CAT} \left[\begin{array}{c} \text{HEAD} \quad \text{pass1} \\ \text{VAL.COMPS} \quad \langle \boxed{1} \rangle \end{array} \right] \end{array} \right] \\ \text{VAL.COMPS} \quad \langle \boxed{1} \rangle \end{array} \right]$$

3.2 Passive type two

We made a rule *pas-two-no-prefix*, which adds no affixes for P2. It inherits from the same inflectional rule and the output is a word, of type *passive-two-verb-lex*, with an obligatory complement. Its AVM is shown in (9). It takes two saturated noun phrase arguments: the first argument is the subject whose HEAD’s value is of type *subj-noun* and the second argument is the sole item in the COMPS whose HEAD’s value is of type *pass2agent* and it has a feature POSTHEAD whose value is minus, i.e. it must occur before the head verb. The type hierarchy for *pass2agent*, which is the head type for agent in P2, is shown in Figure 5.⁵ The type *pass1agent* (see Figure 2) and *pass2agent* have *propername* and *pron-3pers* as their subtypes.

⁵ Another approach is to analyze P2 agents as “lite” pronouns (Abeillé & Godard, 2001) because they must be adjacent to the P2 verbs but can be coordinated, like *aku atau dia* ‘me or him/her’ or modified, like *aku sendiri* ‘me alone’. At present, we are still analyzing the possibility of this approach.

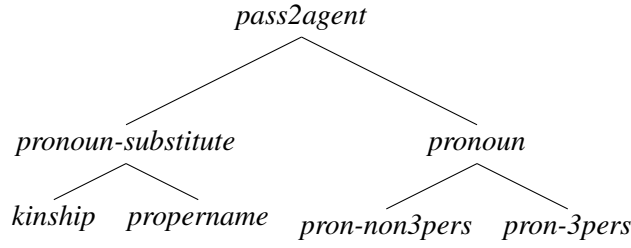


Figure 5: Type hierarchy for P2 agent

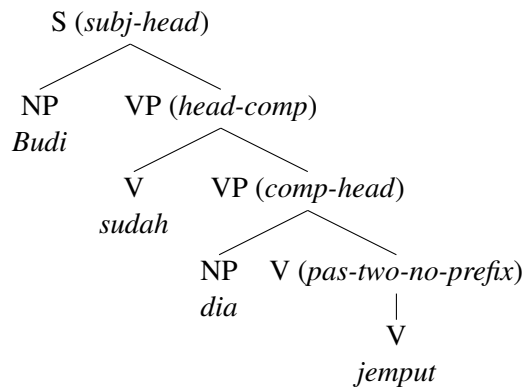


Figure 6: Parse tree of *Budi sudah dia jemput*

$$(9) \left[\begin{array}{l} \text{HEAD } \textit{passive-two} \\ \text{VAL.COMPS } \langle \boxed{1} \rangle \\ \text{ARG-ST } \left\langle \left[\begin{array}{l} \text{HEAD } \textit{subj-noun} \\ \text{VAL } \left[\begin{array}{l} \text{SPR } \langle \rangle \\ \text{COMPS } \langle \rangle \end{array} \right] \end{array} \right], \boxed{1} \left[\begin{array}{l} \text{HEAD } \textit{pass2agent} \\ \text{POSTHEAD -} \\ \text{VAL } \left[\begin{array}{l} \text{SPR } \langle \rangle \\ \text{COMPS } \langle \rangle \end{array} \right] \end{array} \right] \right\rangle \end{array} \right]$$

In addition, we made a new phrase rule called *complement-head* rule, which is constrained to lexical P2 head only. The HEAD value of its HEAD-DTR is of type *passive-two*. Parse tree of (3) is shown in Figure 6. The complement (agent) and P2 verb are combined by *complement-head* rule, the result is combined with the aspect marker by *head-complement* rule. Its semantics is similar to the one shown in Figure 4.

For P2 with *ku-* ‘1SG’, we made a rule *pas-two-prefix-ku* which adds *ku-*. It adds the semantic information that the PERNUM of the ARG1 is first person singular. The COMPS is saturated.

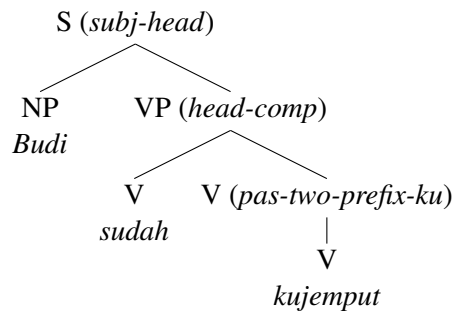


Figure 7: Parse tree of *Budi sudah kujemput*

The result is a passive verb with *ku-* whose COMP's value is empty (saturated) but still needs a subject. The verb's ARG2 is coindexed with the INDEX of the subject, whose HEAD's value is of type *subj-noun*.

Parse tree of (4) is shown in Figure 7. It shows *pas-two-prefix-ku* rule makes the lexeme *jemput* 'pick.up' become *kujemput* '1SG-pick.up'. The result is the verb *kujemput* '1SG-pick.up' which has *aku* '1SG' in the semantics, coindexed with the ARG1 of the verb. This verb is then combined with an aspect marker *sudah* 'PRF' by *head-complement* rule. Its semantics is similar to the one shown in Figure 4.

For P2 with *kau-*, we treat it similarly as for P2 with *ku-*. We made a rule *pas-two-prefix-kau* which adds *kau-* with the PERNUM of ARG1 is second person singular.

4 Conclusion

We made four rules for two types of passive (P1 and P2) and type hierarchies for the complement nouns (agent). Due to the optionality of the complements in coordinative constructions, the bound pronouns *-nya* '3SG', *ku-* '1SG', and *kau-* '2SG' are treated differently: *-nya* is treated as a word, while *ku-* and *kau-* are treated as affixes. We made a *complement-head* rule which combines a complement with a P2 verb without affixes.

References

- Abeillé, Anne & Daniele Godard. 2001. A Class of 'Lite' Adverbs in French. *Romance syntax, semantics and their L2 acquisition* 9–26.
- Alwi, Hasan, Soenjono Dardjowidjojo, Hans Lapoliwa & Anton M. Moeliono. 2014. *Tata Bahasa Baku Bahasa Indonesia*. Jakarta: Balai Pustaka 3rd edn.
- Arka, I Wayan & C. D. Manning. 2008. Voice and grammatical relations in Indone-

- sian: a new perspective. In *Voice and Grammatical Relations in Austronesian Languages*, 45–69. Stanford: CSLI Publications.
- Cole, Peter, Gabriella Hermon & Yanti. 2008. Voice in Malay/Indonesian. *Lingua* 118(10). 1500–1553.
- Copestake, Ann, Dan Flickinger, Carl Pollard & Ivan A. Sag. 2005. Minimal Recursion Semantics: An Introduction. *Research on Language and Computation* 3(4). 281–332.
- McCune, Keith. 1979. Passive Function and the Indonesian Passive. *Oceanic Linguistics* 18(2). 119–169.
- Moeljadi, David, Francis Bond & Sanghoun Song. 2015. Building an HPSG-based Indonesian Resource Grammar (INDRA). In *Proceedings of the GEAF Workshop, ACL 2015*, 9–16. <http://aclweb.org/anthology/W/W15/W15-3302.pdf>.
- Musgrave, Simon. 2001. *Non-subject arguments in Indonesian*. Melbourne: The University of Melbourne PhD dissertation.
- Nomoto, Hiroki. 2013. On the optionality of grammatical markers: A case study of voice marking in Malay/Indonesian. *Voice variation in Austronesian languages of Indonesia* 121–143.
- Payne, Thomas E. 2008. *Describing Morphosyntax: A Guide for Field Linguists*. New York: Cambridge University Press.
- Pollard, Carl & Ivan A. Sag. 1994. *Head Driven Phrase Structure Grammar*. Chicago: University of Chicago Press.
- Riesberg, Sonja. 2014. *Symmetrical Voice and Linking in Western Austronesian* Pacific Linguistics 646. Berlin: Mouton de Gruyter.
- Sag, Ivan A., Thomas Wasow & Emily M. Bender. 2003. *Syntactic Theory: A Formal Introduction*. Stanford: CSLI Publications 2nd edn.
- Sneddon, James Neil, Alexander Adelaar, Dwi Noverini Djenar & Michael C. Ewing. 2010. *Indonesian Reference Grammar*. New South Wales: Allen & Unwin 2nd edn.
- Song, Sanghoun. 2017. *Modeling information structure in a cross-linguistic perspective*. Berlin: Language Science Press.
- Voskuil, Jan E. 2000. Indonesian voice and A-bar movement. In *Formal issues in austronesian linguistics*, 195–213. Springer.
- Zwicky, Arnold M. & Geoffrey K. Pullum. 1983. Cliticization vs. Inflection: English N'T. *Language* 59(3). 502–513.

Modeling adnominal possession in multilingual grammar engineering

Elizabeth Nielsen

University of Washington

Emily M. Bender 

University of Washington

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 138–151

Keywords: adnominal possession, grammar engineering, computational typology, Grammar Matrix

Nielsen, Elizabeth & Emily M. Bender. 2018. Modeling adnominal possession in multilingual grammar engineering. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 138–151. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.9.

 CC-BY 4.0

Abstract

In this paper we describe insights gained from building an extension to the LinGO Grammar Matrix customization system to cover adnominal possessive phrases. We show how the wide range of such constructions attested in the world’s languages can be handled with the typical major phrase types used in HPSG and discuss the value of feature bundling in the multilingual grammar engineering context.

1 Introduction

This paper presents observations drawn from an implemented, typologically grounded, cross-linguistic, HPSG analysis of adnominal possession. This particular phenomenon is an interesting target for such an analysis because it is likely to occur in most if not all languages, has an interesting range of typological parameters each of which has a tractable number of possible options. Our analysis is developed and implemented in the context of the LinGO Grammar Matrix (Bender et al., 2002, 2010) and we draw conclusions about the range of phrase structure rule types required for these expressions and the value of bundling information together within ancillary types.

The LinGO Grammar Matrix is an open-source project that allows user-linguists to jump-start the creation of implemented HPSG grammars. The Grammar Matrix consists of a core grammar and a customization system. The core grammar is a set of grammatical type definitions which can be used to model various realizations of typologically widespread phenomena; the customization system consists of a web interface that elicits typological information from the user-linguist via a questionnaire (Bender et al., 2010), and Python-based back-end code that draws from and adds to the core grammar in order to produce the implemented grammar for a given language. Since the Grammar Matrix project has always had the goal of being able to model the attested typological variation within the various linguistic phenomena that it covers, it functions not only as a tool for grammar engineers, but also as a set of typological generalizations and predictions, in a testable and internally consistent format (Bender et al., 2010; Bender, 2016).

We extended the current Grammar Matrix customization system by adding a library to model adnominal possession. This paper relates two of the typological and theoretical generalizations that were arrived at in the process of developing this extension to the Grammar Matrix. We begin by giving some background on the phenomenon — adnominal possession — that the library was intended to cover and the way in which we broke down this typological space. Second, we discuss one generalization we arrived at in the process of library creation, namely the suitability of major phrase types already in existence in the Grammar Matrix (head-specifier, head-complement, head-modifier) to model possessive phrases. We demonstrate that all marked possessive constructions can be modeled without requiring specific

additional binary phrasal constructions in any language. Lastly, we discuss another discovery, namely the implications of the decision made within the Matrix to bundle person, number, and gender features under a single feature called PNG. This bundling of features turns out to be very beneficial in the context of multilingual grammar engineering, since it allows a consistent way of dealing with these features in languages with disparate ways of dealing with person, number, and gender.

2 Describing the typological space

The goal in constructing this library was to be able to model all attested adnominal possessive constructions—that is, constructions involving two noun phrases whose referents participate in a possessive relation—based on a minimal amount of information from the user-linguist. To that end, in this section, we lay out the attested typological variation in terms of a few binary- or ternary-valued features that distinguish possible types of adnominal possessive phrases and define the boundaries of the typological space under discussion. The majority of typological variation in adnominal possessive phrases can be captured by the following features:

- Order: possessum–possessor, possessor–possessum
- Presence and type of marker: ϕ , affix, clitic, word
- Location of marker: possessum, possessor, both
- Syntactic relation: modifier-like, specifier-like
- Agreement: with possessum, with possessor, both
- Possessor type: full NP, pronoun

We briefly describe each of these in turn, bearing in mind that any given language can have multiple different possessive constructions.

Order We observe constructions in which the possessor always precedes the possessum, and the reverse:

- (1) Komi, possessum-final:

kyf kor-jas
birch leaf-PL
birch’s leaves [kom] (Grashchenkov 2005:29)

- (2) Maltese, possessum-first:

bin is-sultn
son DEF-king
the king’s son [mlt] (Grashchenkov 2005:29)

Presence and type of marker In a given possessive construction, overt markers of possession may or may not appear. If no such markers exist, a possessive phrase may simply consist of a pair of juxtaposed nouns, as in the following example from Yoruba:

- (3) Yoruba, no marking:

íwè baba
book father
'father's book' [yor] (Grashchenkov, 2005, 28)

In cases where these markers appear, they may take the form of an affix, a clitic, or an independent word:

- (4) Imbabura Quechua, affix:

José-paj wasi
José-POSS house
'José's house' [qvi] (Grashchenkov, 2005, 34-35)

- (5) Basque, clitic:

neska gazte-a=ren edertasuna
girl young-DEF=POSS beauty
'the beauty of the young girl' [baq] (Grashchenkov, 2005, 33)

- (6) Bulgarian, independent word:

lah na proletta
breath POSS spring
'the breath of spring' [bul] (Grashchenkov, 2005, 31)

Location of markers In possessive constructions which are marked by an overt morpheme, those morphemes can also be described in terms of where they occur: markers of possession may appear on the possessor, on the possessum, or in both locations. For example, in Yucatec Maya, possession is marked by inflection on the possessum,¹ while in Malagasy, it is marked by inflection on the possessor:

- (7) Yucatec Maya, possessum-marking:

u=k'àaba' le x-ch'up-pàal-a'
POSS.3=name DEF FEM-woman-child-D1
'the name of that girl' [yua] (Grashchenkov, 2005, 36-37)

- (8) Malagasy, possessor-marking:

zana d-rabe
child POSS-Rabe
'the child of Rabe' [mlg] (Grashchenkov, 2005, 34-35)

¹In this example, the possessum-marking inflection also carries agreement information, indicating that it agrees with a third-person possessor; examples of possessum-marking inflection without agreement are rare (Grashchenkov, 2005).

Syntactic relation In the typological literature on possessive phrases, a distinction is often drawn between specifier-like possessors and modifier-like possessors. The English 's-genitive is a classic example of a construction with a specifier-like possessor, since the possessor fills the same slot as a specifier, blocking the possessum from taking a determiner:

- (9) English, specifier-like possessor:
 the father's house
 * the the father's house [eng]

By contrast, Ancient Greek possessive pronouns are more like modifiers, in that they occur alongside the possessum's determiner:

- (10) Ancient Greek, modifier-like possessor:
 he: to patròs oikía
 the.F.SG.NOM the.M.SG.GEN father(M).SG.GEN house(F)SG.NOM
 'the father's house' [gre] (Goodwin, 1894)

Agreement There are languages in which the possessor agrees in person, number, and/or gender with the possessum, such as Romani, shown below. There are also languages, such as Yucatec Maya, as shown in (7) above, or Finnish, illustrated in (12), where the possessum agrees with the possessor.

- (11) Romani, possessor agreement:
 e manús-es-quiri buzni
 the:OBL.M.SG man-OBL.SG.M-GEN:F.SG.NOM goat(F)
 'the man's goat' [rom] (Koptjevskaja-Tamm, 2001, 962)
- (12) Finnish, possessum agreement:
 heidän ystävä-nsä
 their friend-3POSS
 'their friend' [fin] (Toivonen, 2000, 585)

Noun type All languages allow the expression of possessives with both full NP and pronominal possessors. In some languages, the same possessive constructions are used for both. In others, pronominal possessors are treated differently. A special case of this is when the pronominal possessives are just the affixes that would attach to the possessum to indicate agreement with the possessor, in the absence of any overt possessor.

This brief summary of the typological space under consideration provides the background for our crosslinguistic analysis of possessives. In the next section, we will give a brief overview of the analysis we put forward in this library for possessive phrases.

3 Phrase types used to model possession

In this section, we give an overview of the semantic and syntactic structures that we posit to model adnominal possessive phrases. For a more detailed presentation of this analysis, see Nielsen 2018. We will focus in particular on one important subset of possessive phrases — namely marked possessive phrases — and show that they can be modelled in terms of phrase types that already exist in the Grammar Matrix.

The Grammar Matrix produces grammars which map from strings to Minimal Recursion Semantics (MRS; Copestake et al. 2005) representations. A full description of the elements of MRS is beyond the scope of the present work. For the purposes of this discussion, it is sufficient to note that the various entities in a possessive phrase such as *the dog’s cat* each correspond to a predication element in the MRS representation (see (13)). For example, the predication *dog_n_rel* corresponds to the noun *dog*, and so forth. The possessive relation itself is likewise encoded by means of a predication, namely a possessive relation (called *poss_rel*). This relation takes two arguments, which correspond to possessor (*x3*) and posses-
sum (*x2*):

$$(13) \quad \left\langle \begin{array}{l} h_{13}, \\ h_4:\text{cat_n_rel}(\text{ARG0 } x_2), \\ h_6:\text{dog_n_rel}(\text{ARG0 } x_3), \\ h_{11}:\text{exist_q_rel}(\text{ARG0 } x_3, \text{RSTR } h_7, \text{BODY } h_{12}), \\ h_8:\text{exist_q_rel}(\text{ARG0 } x_2, \text{RSTR } h_5, \text{BODY } h_9), \\ h_4:\text{poss_rel}(\text{ARG0 } e_{10}, \text{ARG1 } x_2, \text{ARG2 } x_3) \\ \{ h_5 =_q h_6, h_7 =_q h_4 \} \end{array} \right\rangle$$

Any given possessive construction must include some element that introduces this *poss_rel*, in order to ensure that the final phrase has the correct possessive semantics. In this section, we outline briefly the approach taken to solving this problem in two cases: unmarked and marked possessive constructions. In unmarked possessive phrases, we introduce a unique binary phrase structure rule to model possessive phrases; in the case of marked possessive phrases, we have demonstrated that all possessive phrases can be modeled in terms of existing phrase structure types in the Grammar Matrix.

3.1 Unmarked possessive phrases

In marked possessive constructions (see §3.2), the *poss_rel* is introduced on our analysis by the overt marker of possession. Unmarked possessives represent the same meaning, but there is no such marker to pin the semantics on. Accordingly, we introduce a new binary phrase type, called *poss-phrase*, to license the juxtaposition of possesum and possessor and introduce the possessive semantics (*poss_rel*). One variant of this phrase rule is shown as an AVM in (14) below.² It inherits from one of two supertype phrase structure rules which are defined in the matrix

²Some constraints not relevant to this discussion are omitted due to space constraints.

core grammar: *head-initial* or *head-final*, which introduce the appropriate ordering constraints. The rule also varies depending on other properties of the possessive construction. The version shown in (14) is appropriate for the case where the possessor is in a modifier-like relationship to the possessum. Accordingly, the SPR value is shared between mother and daughter. If the possessor fills the specifier role for the possessum, the SPR value on the mother will be the empty list and the *poss-phrase* will also contribute a quantifier for the possessum through its C-CONT. For further details on variants of this rule, see Nielsen 2018.

$$(14) \left[\begin{array}{l} \text{poss-phrase} \\ \\ \text{SYNSEM|LOCAL|CAT} \left[\begin{array}{l} \text{HEAD } \boxed{1} \\ \text{VAL} \left[\begin{array}{l} \text{COMPS } \langle \rangle \\ \text{SUBJ } \langle \rangle \\ \text{SPEC } \langle \rangle \\ \text{SPR } \boxed{7} \end{array} \right] \end{array} \right] \\ \\ \text{HEAD-DTR|SYNSEM|LOCAL} \left[\begin{array}{l} \text{CAT} \left[\begin{array}{l} \text{HEAD } \boxed{1} \left[\begin{array}{l} \text{noun} \\ \text{PRON } - \end{array} \right] \\ \text{VAL.SPR } \boxed{7} \langle X \rangle \end{array} \right] \\ \text{CONT|HOOK } \boxed{2} \left[\begin{array}{l} \text{LTOP } \boxed{5} \\ \text{INDEX } \boxed{3} \end{array} \right] \end{array} \right] \\ \\ \text{NON-HEAD-DTR|SYNSEM|LOCAL} \left[\begin{array}{l} \text{CONT|HOOK|INDEX } \boxed{4} \\ \text{CAT} \left[\begin{array}{l} \text{VAL|SPR } \langle \rangle \\ \text{HEAD noun} \end{array} \right] \end{array} \right] \\ \\ \text{C-CONT} \left[\begin{array}{l} \text{HOOK } \boxed{2} \\ \text{RELS } \left\langle \left[\begin{array}{l} \text{PRED } \text{poss_rel} \\ \text{LBL } \boxed{5} \\ \text{ARG1 } \boxed{3} \\ \text{ARG2 } \boxed{4} \end{array} \right] \right\rangle \\ \text{HCONS } \langle \rangle \end{array} \right] \end{array} \right]$$

This phrase structure rule allows the correct possessive relationship to be modeled between possessor and possessum in the absence of any overt markers of possession. In the next section, we discuss how we model possessive phrases which do include an overt marker of possession.

3.2 Marked possessive phrases

In the literature on adnominal possession, both within the HPSG framework and beyond, it is common to discuss possessive phrases as being one manifestation of highly general phrase types. In one classic example, Lyons (1986) draws a distinction between ‘adjective-genitives’ and ‘determiner-genitives’, suggesting that, modulo some inflectional morphology, possessors are essentially just another kind of specifier or modifier, no different from any other. Within the HPSG literature, there are examples of analyses of possessive phrases being described as instances of head-modifier phrases (e.g. Beerman and Ephrem, 2007) or head-specifier phrases (e.g. Kolliakou, 1995).

Though there are challenges in modeling possessive phrases in terms of these major phrase types, we demonstrate that the head-specifier, head-modifier, and head-complement³ phrase structure rules can adequately model all attested possessive phrase types. This serves to validate the practice of referring to possessive phrases as subtypes of these major phrase types.

Using major phrase types to model possessive phrases does present several challenges. Most pronounced of these are the challenges involved in using the head-specifier construction to model possessive phrases with specifier-like possessors. As constituted in the Grammar Matrix with its implementation of Minimal Recursion Semantics (MRS; Copestake et al., 2005), the head-specifier rule is non-head-compositional — that is, semantic information used for further composition (the information in HOOK) from the non-head daughter is ‘passed up’ to the mother (as shown in (15)). Given the nature and goals of the Grammar Matrix, this formulation of the head-specifier rule is not merely a convenient implementation choice, but a cross-linguistic analytical claim (Bender et al., 2002).

$$(15) \left[\begin{array}{l} \textit{basic-head-spec-phrase} \\ \text{NON-HEAD-DTR} | \text{SYNSEM} | \text{LOCAL} | \text{CONT} | \text{HOOK } \boxed{1} \\ \text{C-CONT} | \text{HOOK } \boxed{1} \end{array} \right]$$

In a typical head-specifier construction, such as a noun phrase consisting of a determiner and a noun, the determiner identifies its own INDEX with the INDEX of

³The head-complement phrase structure rule is used to model some modifier-like possessive phrases. This is because the head-modifier rule only includes one-way selection — the modifier selects for its head, but not vice versa. In order for the possessive semantics to work out correctly, the element that carries the *poss_rel* must have access to the semantic information of both possesum and possessor, since this relation takes both possessor and possesum as arguments. This circumstance only obtains for the selecting element. This could always be modeled by simply having the possessive semantics appear on the selecting element, regardless of whether or not it is the marked element. This is a perfectly acceptable solution. However, we chose to keep the possessive semantics on the element that carries overt marking of possession. This means that in some cases, such as in Hungarian, it will be the case that the selecting element in a head-modifier construction will not be the marked element. In these cases, we use the head-complement rule in order to make the marked element the selecting element, allowing us to construct the same semantic representation. For further detail, see Nielsen 2018.

the noun (through its SPEC value), so the INDEX of the head-specifier phrase is still identified with the INDEX of the noun.

$$(16) \left[\begin{array}{l} \text{basic-determiner-lex} \\ \text{SYNSEM|LOCAL|CAT|VAL|SPEC} \left\langle \left[\text{LOCAL|CONT|HOOK|INDEX } \boxed{1} \right] \right\rangle \\ \text{LKEYS|KEYREL} \left[\text{ARG0 } \boxed{1} \right] \end{array} \right]$$

However, this approach does not work for modeling specifier-like possessive constructions. Take for example the scenario where the possessor is marked with a possessive affix. If we were to attempt a similar approach, the lexical rule for the possessive affix would look something like (17) (in abbreviated form), where the overall index of the lexical rule ($\boxed{1}$) is identified with the index of the possessum noun, much like in the lexical type for determiners.

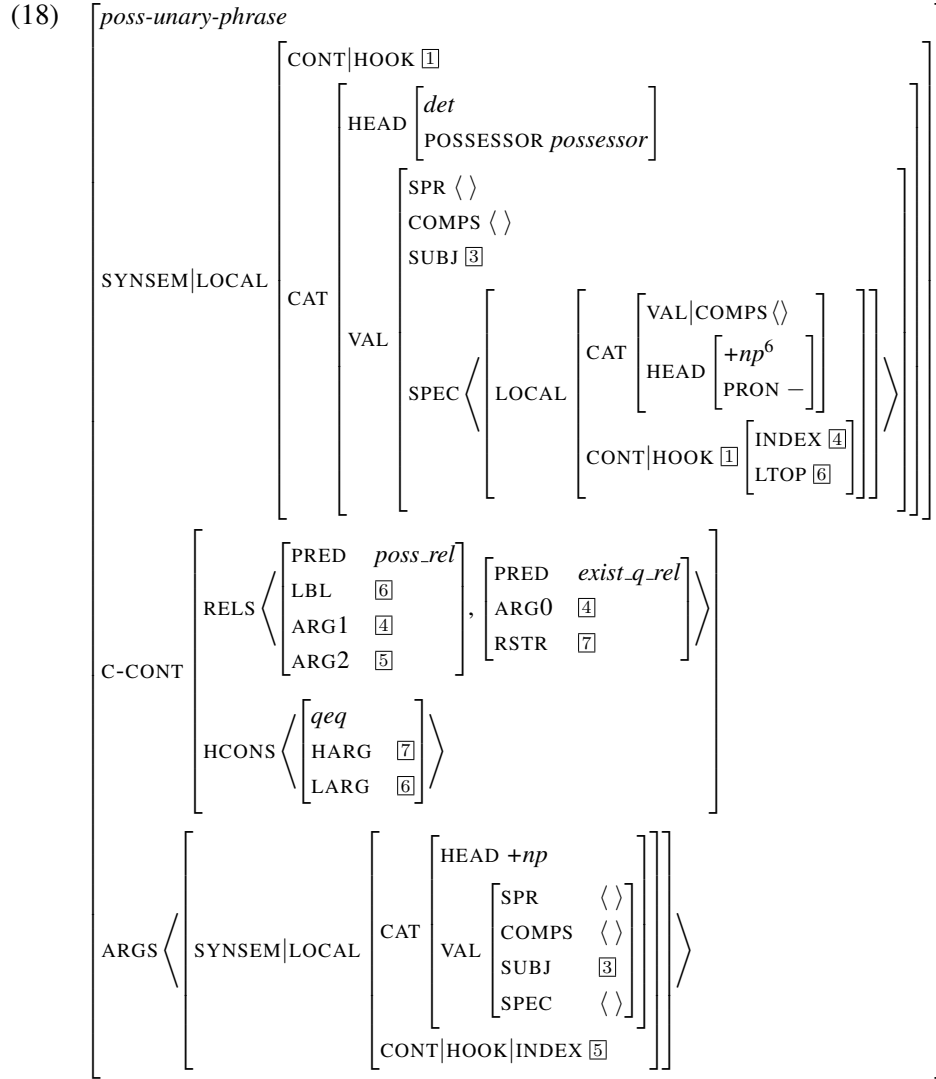
$$(17) \left[\begin{array}{l} \text{possessor-lex-rule (hypothetical)} \\ \text{SYNSEM|LOCAL|CAT|VAL|SPEC} \left\langle \left[\text{LOCAL|CONT|HOOK|INDEX } \boxed{1} \right] \right\rangle \\ \text{DTR|SYNSEM|LOCAL|CONT|HOOK|INDEX } \boxed{2} \\ \text{C-CONT} \left[\begin{array}{l} \text{HOOK|INDEX } \boxed{1} \\ \text{RELS} \left\langle \left[\begin{array}{ll} \text{PRED} & \text{poss_rel} \\ \text{ARG1} & \boxed{1} \\ \text{ARG2} & \boxed{2} \end{array} \right], \left[\begin{array}{ll} \text{PRED} & \text{exist_q_rel} \\ \text{ARG0} & \boxed{4} \\ \text{RSTR} & \boxed{7} \end{array} \right] \right\rangle \\ \text{HCONS} \left\langle \left[\begin{array}{ll} \text{qeq} & \\ \text{HARG} & \boxed{7} \\ \text{LARG} & \boxed{6} \end{array} \right] \right\rangle \end{array} \right] \end{array} \right]$$

Problems arise with this analysis because the possessor noun must still participate in some constructions as a typical noun would, but it has partially adopted the semantics of the possessum. For example, when a determiner attaches to this possessor noun, it will serve as the quantifier for the possessum, rather than for the possessor.

We solve this in our library in the following way: the *possessor-lex-rule* is pared down to a rule that simply adds a HEAD feature [POSSESSOR *possessor*].⁴ A unary phrase rule (shown in (18) below) then takes the NP consisting of the possessor (and any determiner and/or modifiers it may take) as its daughter, introduces the possessive predication (*poss_rel*), and produces a constituent whose INDEX is identified with the possessum, as shown in (18). This allows the possessor to be a

⁴The feature POSSESSOR has the values *possessor* and *nonpossessive*. Similarly, there exists a POSSESSUM feature with values *possessum* and *nonpossessive*.

semantically typical noun within its own NP, and then to take on the necessary specialized semantic behavior when interacting with the rest of the possessive phrase. This analysis is used for all specifier-like possessive constructions.⁵



Though possessive phrases are challenging for the established major phrase types in the Grammar Matrix, it is ultimately still possible to assimilate them to

⁵This analysis differs from the analysis put forth for the English 's-possessive in Sag et al. (2003) and Flickinger (2002). Since this construction features a specifier-like possessor, these previous accounts have analyzed 's as a determiner. The semantics of determiners make the unary rule discussed here unnecessary. However, this analysis is only possible for specifier-like possessives where the possessive marker is an independent word. Since this is only one of many construction types that must be covered by the adnominal possession library, we chose the more general solution put forward here.

⁶This is an abbreviation used in the Grammar Matrix for a supertype that includes the HEAD values *adp(osition)* and *noun*.

these types (though at the cost of adding minor phrase types that are specific to possessives). This analysis supports the widespread claim that possessive phrases are instances of head-specifier, head-modifier, and head-complement phrases.

4 Feature bundling

In this section, we discuss the analysis developed for agreement between possessor and possessum, focusing on how bundling together certain features is particularly useful in multilingual grammar engineering. The phenomenon of either the possessum or the possessor agreeing with the other element of the possessive phrase is observed in many languages, as shown in (11) and (12) and above, reproduced as (19) and (20) below:

(19) Romani:

e manús-es-quiri buzni
the:OBL.M.SG man-OBL.SG.M-GEN.:F.SG.NOM goat(F)
‘the man’s goat’ [rom] (Koptjevskaja-Tamm, 2001, 962)

(20) Finnish:

heidän ystävä-nsä
their friend-3POSS
‘their friend’ [fin] (Toivonen, 2000, 585)

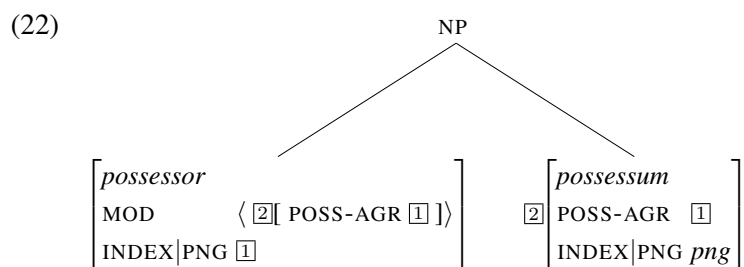
When the possessor is specifier-like, this is easy enough to account for: an agreeing possessum constrains the relevant person, number, and gender features of the possessor, which appears on its SPR list. Since the head and non-head daughters both select for each other in the head-specifier schema (Pollard and Sag, 1994, Ch. 9), this analysis works equally well in the case where the possessor agrees with the possessum. However, when the possessor is modifier-like, possessor and possessum are joined by a head-modifier rule which has no such mutual selection. The possessor can constrain the features of the possessum, which appears on its MOD list, but the possessum has no access to its possessor’s features. In order to fully cover the possible typological space, agreement in both directions should be possible whether the possessor is modifier-like or specifier-like.

Indeed, that full typological space is attested in the world’s languages. Hungarian provides an example of the scenario where the possessor is the modifier of the possessum, but we still see agreement markers on the possessum, as illustrated in (21).

(21) Hungarian:

az én kalap-ja-i-m
the I hat-POSS-PL-1SG
‘my hats’ [hun] (Laczko, 2007)

Since the possessum cannot select its modifier, instead the possessum must somehow ‘publish’ the person, number or gender features it agrees with, so that the possessor can select for a possessum with the correct agreement features. This means it is necessary for the possessum to carry two sets of agreement features: the inherent person, number, and gender features it has as a noun; and the person, number, and gender features that it agrees with. The former are found (as usual) at INDEX.PNG, while the latter are in the new head feature we posit, called POSS-AGR. The possessum can then do the work of identifying the possessor’s agreement features with its own features, as sketched in the tree in (22):



Adding this second set of agreement features has the potential to be difficult in the multilingual grammar engineering context. While some languages have separate person, number, and gender features, others lack one of these three, or are better analyzed as having a combined PERNUM feature (Drellishak, 2009). Our library needs to be interoperable with all of these options. Given just three possible features, which may or may not appear, or which may be combined, there are a dozen possible features sets available. Creating different variants of the possessor-possessum agreement constraints in each of these cases would amount to redundantly reproducing the work of Drellishak’s PNG library.

Fortunately, and for independent reasons, Drellishak bundled all person, number, and gender features as features of the type *png*. This turns out to be very beneficial for us: We simply reuse the type *png* as the value of our new feature POSS-AGR. This allows our library to abstract away from the specifics of how person, number, and gender work. Thus we see that in addition to providing efficiency as a monolingual level (Flickinger, 2000), types also add efficiency to cross-linguistic grammar engineering.

5 Conclusion

The process of implementing an analysis for any phenomenon frequently leads to theoretical insights or analytical refinements. In the context of multilingual grammar engineering, the added constraint of harmonizing analyses for hundreds of possible variations on the phenomenon crosslinguistically provides all the more

opportunity for finding such insights and refinements. In this paper, we have detailed two ways in which the crosslinguistic perspective on modeling adnominal possession is beneficial, namely the confirmation of the applicability of major phrase types to modeling possessive phrases and the insight into the advantages of bundling person, number, and gender features under a single type. This analysis has been tested by constructing testsuites for ten typologically and genetically diverse languages, half of which weren't considered during library development and then creating grammars using the augmented customization system to evaluate against those testsuites. The results of these tests can be found in Nielsen 2018. Possible directions for future work include extending the library to fully cover inalienable possession (which is currently only partially covered by the library) and handling agreement in features such as case, where currently only person, number, and gender are handled.

References

- Dorothee Beerman and Binyam Ephrem. The definite article and possessive markings in Amharic. In Frederick Hoyt, Nikki Seifert, Alexandra Teodorescu, and Jessica White, editors, *Proceedings of the Texas Linguistics Society IX Conference*, pages 21–32. CSLI Publications, 2007.
- Emily M. Bender. Linguistic typology in natural language processing. *Linguistic Typology*, 20:645–660, 2016.
- Emily M. Bender, Dan Flickinger, and Stephan Oepen. The Grammar Matrix: An open-source starter-kit for the rapid development of cross-linguistically consistent broad-coverage precision grammars. In John Carroll, Nelleke Oostdijk, and Richard Sutcliffe, editors, *Proceedings of the Workshop on Grammar Engineering and Evaluation at the 19th International Conference on Computational Linguistics*, pages 8–14, Taipei, Taiwan, 2002.
- Emily M. Bender, Scott Drellishak, Antske Fokkens, Laurie Poulson, and Safiyyah Saleem. Grammar customization. *Research on Language and Computation*, 8 (1):23–72, 2010.
- Ann Copestake, Dan Flickinger, Carl Pollard, and Ivan Sag. Minimal recursion semantics: An introduction. *Research on Language and Computation*, 3(2): 281–332, July 2005.
- Scott Drellishak. *Widespread but not universal: Improving the typological coverage of the Grammar Matrix*. PhD thesis, University of Washington, 2009.
- Dan Flickinger. On building a more efficient grammar by exploiting types. *Natural Language Engineering*, 6 (1) (Special Issue on Efficient Processing with HPSG): 15 – 28, 2000.

- Dan Flickinger. On building a more efficient grammar by exploiting types. In Jun'ichi Tsujii Stephan Oepen, Dan Flickinger and Hans Uszkoreit, editors, *Collaborative Language Engineering*, pages 1–17. CSLI Publications, Stanford, 2002.
- William W. Goodwin. *A Greek Grammar*. Macmillan & Co., London, 1894.
- P. Grashchenkov. Typology of possessive constructions. *Voprosy yazykoznaniya*, 55(3):25–54, 2005.
- Dimitra Kolliakou. *Definites and possessives in Modern Greek*. PhD thesis, University of Edinburgh, 1995.
- Maria Koptjevskaja-Tamm. Adnominal possession. In W. Oesterreicher Haspelmath M., E. Konig and W. Raible, editors, *Language Typology and Language Universals*, volume 2, pages 960–970. Walter de Gruyter, Berlin, 2001.
- Tibor Laczko. Revisiting possessors in Hungarian DPs: A new perspective. In Miriam Butt and Tracy Holloway King, editors, *Proceedings of the LFG07 Conference*. CSLI Publications, 2007.
- Elizabeth Nielsen. Modeling adnominal possession in the lingo grammar matrix. Master's thesis, University of Washington, 2018.
- Carl Pollard and Ivan A. Sag. *Head-Driven Phrase Structure Grammar*. Studies in Contemporary Linguistics. The University of Chicago Press and CSLI Publications, Chicago, IL and Stanford, CA, 1994.
- Ivan A. Sag, Thomas Wasow, and Emily Bender. *Syntactic Theory: A Formal Introduction*. CSLI lecture notes; no. 152. Center for the Study of Language and Information, Stanford, CA, 2nd ed. edition, 2003.
- Ida Toivonen. The morphosyntax of Finnish possessives. *Natural Language and Linguistic Theory*, 18(3):579–609, 2000.

Part II

Contributions to the Workshop

Korean and Spanish psych-verbs: Interaction of case, theta-roles, linearization, and event structure in HPSG

Antonio Machicao y Priemer 

Humboldt-Universität zu Berlin

Paola Fritz-Huechante 

Humboldt-Universität zu Berlin

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 153–173

Keywords: Spanish, Korean, Psych-Verbs, Theta-Roles, Neo-Davidsonian, Linking

Machicao y Priemer, Antonio & Paola Fritz-Huechante. 2018. Korean and Spanish psych-verbs: Interaction of case, theta-roles, linearization, and event structure in HPSG. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 153–173. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.10.

 CC-BY 4.0

Abstract

In this paper, we argue that by making a more detailed distinction of theta-roles, while at the same time investigating the correlation of case marking, theta-role assignment, and eventuality types, we can describe different psych-verb subclasses and explain their alignment patterns in Spanish and Korean. We propose a neo-Davidsonian treatment of psych-verbs in HPSG that allows us to account for the underspecification of theta-roles which are modeled in an inheritance hierarchy for semantic relations. By assuming linking properties modeled lexically, we can constrain the properties for psych-verbs that shows the mapping of semantic arguments (i.e. *experiencer*, *stimulus-causer*, *subject matter* and *target*) to the elements in the argument structure. The type hierarchy and lexical rules proposed here capture the alternation in case marking not only of the experiencer (as traditionally assumed in the literature), but also of the stimulus. This analysis leads us to a new fourfold classification of psych-verbs for both languages.

1 Introduction

Psychological verbs (henceforth psych-verbs), such as English *frighten*, *worry*, *anger*, have caused large interest due to their particular properties and their implications for the theory of argument structure (cf. Belletti & Rizzi, 1988; Grimshaw, 1990; Pesetsky, 1995; Landau, 2010, a.o.). The configuration of these verbs contains two arguments: (a) an EXPERIENCER (EXP), which is an animate individual affected by a psychological eventuality; and (b) a STIMULUS (STM), which refers to an animate or inanimate entity that triggers the psychological state in the EXP (cf. Pesetsky, 1995). The literature classifies these verbs into two classes according to their argument and event structures: (a) experiencer-subject (ES) verbs, e.g. *love* and *fear* (1a); and (b) experiencer-object (EO) verbs, e.g. *frighten* and *worry* (1b).

- (1) a. Clara_{EXP} loves David_{STM}.
- b. David_{STM} frightens Clara_{EXP}.

The EO class has been further divided into those verbs that only assign dative to the experiencer (e.g. Spanish *gustar* ‘like’ cf. (6)), and those that alternate the experiencer between accusative (the structural case for objects)

[†]We want to thank many colleagues for their valuable comments (chronologically and alphabetically): the participants of the Syntax-Semantik Kolloquium at the Humboldt-Universität zu Berlin – specially Elisabeth Verhoeven, Elodie Winckel, Marc Felfe, and Stefan Müller; the participants of the 5th European Workshop in HPSG and of the HPSG Conference 2018 – specially Anne Abeillé, Berthold Crysmann, Doug Arnold, Frank Van Eynde, Manfred Sailer and Sang-Hee Park; the participants of the DeMiNeS Workshop – specially Athina Sioupi and Berry Claus. This paper was partly funded by the German Research Association (DFG; project VE 570/1–3). All remaining errors are ours.

and dative (henceforth DAT/ACC alternation), e.g. Spanish *asustar* ‘frighten’ (cf. (3) and (2), respectively) (cf. Van Voorst, 1992; Arad, 1998). There is a general agreement that for the ES class, verbs denote states (Grimshaw, 1990). However, this is not the case for the EO class, which are categorized as accomplishments (cf. Van Voorst, 1992), causative state/events (cf. Arad, 1998; Pykkänen, 2000), and recently as inchoative states (cf. Bar-el, 2005; Marín & McNally, 2011). In addition, EO verbs also show exceptional syntactic properties; one of those being linearization. Empirical studies have demonstrated that the preferred word order in dative (DAT) structures is that of EXP-DAT > STM-DAT (2); whereas in accusative (ACC) constructions, the preferred word order is that of STM-NOM > EXP-ACC (3) for a number of languages (for Spanish cf. Fábregas et al., 2017; for German, Greek, Hungarian, and Korean cf. Temme & Verhoeven, 2016; for English and Polish cf. Jiménez-Fernández & Rozwadowska, 2016).

- (2) [A Clara]_{EXP} le asusta [David / el reporte]_{STM}.
to Clara CL.DAT frightens David the report
‘David / the report frightens Clara.’
- (3) [David / el reporte]_{STM} (la) asusta [a Clara]_{EXP}.
David the report CL.ACC frightens to Clara
‘David / the report frightens Clara.’

Less attention has been paid to the class of ES psych-verbs, which show a canonical word order as in SUBJ_{EXP-NOM} > OBJ_{STM-ACC}; and contrary to the EO psych class that presents a DAT/ACC alternation of the EXP, it has been claimed to have no alternation in case marking of the stimulus-object (cf. Belletti & Rizzi, 1988). However, there is data showing that this is not the case, at least for languages such as Spanish, where verbs as *temer* ‘fear’ and *admirar* ‘admire’ normally assign dative.¹

- (4) David le teme / admira a Clara.
David CL.DAT fears admires to Clara
‘David fears / admires (something about) Clara.’

In this paper, we focus on two typologically different languages, namely Spanish (SVO) and Korean (SOV). We address the challenging issue of the languages’ unmarked word order in association with grammatical functions, theta-roles, case and eventualities in the sentence structure, which matches the prominence relation of these features. We model psych-verbs in HPSG by means of a typed inheritance hierarchy and lexical rules (LRs). We propose a more detailed division in the psych domain for both languages,

¹Since Spanish shows differential object marking w.r.t. full NPs, the ACC/DAT distinction is sometimes blurred (cf. Machicao y Priemer 2014 for more details). Hence, we are making the distinction more clear using clitics.

capturing the fact that not only the EXP alternates in case marking in EO structures, but also the STM alternates with respect to case in ES constructions. Furthermore, our data suggest a differentiation of theta-roles similar to Pesetsky’s (1995) proposal. We assume a STM role which can be further specified as *stimulus-causer*, *subject matter* or *target*. This division correlates with the different subtypes of psych-verbs proposed and the unmarked word order.

2 Properties of psych-verbs in Spanish and Korean

Since Belletti & Rizzi’s (1988) work on Italian psych predicates, these verbs have been seen as a threefold classification: (a) class I (e.g. *temere* ‘fear’): a stative ES structure; (b) class II (e.g. *preoccupare* ‘worry’): a stative/eventive EO construction; and class III (e.g. *piacere* ‘please’): a stative EO structure where the EXP is only assigned dative case and generally appears in pre-verbal position. In addition, Alexiadou et al. (2004) argue that the Italian verb classes II and III are unaccusative and that the mapping of theta-roles to syntactic positions is indeed guided by UTAH (Baker, 1998). The authors claim that psych-verbs have different underlying representations, and at D-structure, the EXP is projected higher than the STM. In terms of their semantic structure, Pesetsky (1995) provides a more detailed analysis of the verbs with respect to their arguments’ theta-roles, where: the subject of ES verbs is the EXP and the object is seen as a *target/subject matter*; while EO verbs have a *causer* as the subject, and thus expanding the thematic hierarchy as in (5).

- (5) Causer > Experiencer > Target/Subject Matter

Based on these ideas, the next section attains to a description of the properties of basic psych-verb constructions in the target languages. Spanish and Korean present different morphological structures in terms of argument alternations and directionality: Spanish derives intransitive ES verbs from more basic transitive EO verbs (e.g. *asustar* ‘frighten’) by means of reflexivization (e.g. *asustarse* ‘get frightened’); whereas Korean derives transitive EO items from more basic intransitive ones (e.g. *mwusepta* ‘scary’) by means of a periphrastic causative operation (e.g. *mwusepkey hata* ‘frighten’). In this paper, we focus on the basic psych-verbs constructions (leaving aside their derivations) and their case alternation patterns, linearization, theta-roles and event structure; providing a more detailed classification of the predicates.

2.1 Spanish

Starting with the EO verb class (class II in Belletti & Rizzi’s 1988 work), the alternation of the EXP between ACC and DAT is generally associated with

the eventuality of the verbs, where dative experiencers appear in stative constructions and accusative experiencers in eventive ones (cf. Arad, 1998; Marín, 2015). This distinction is clear for Spanish. For instance, sentence (2) is stative, with no change of state (CoS) in the EXP (cf. Marín, 2015). Following Pesetsky (1995), the STM bears the theta-role of the subject matter (SM): a non-agentive argument which provokes an emotional response in the EXP, but does not cause the emotion directly. The interpretation of sentence (2) is that the experiencer Clara is frightened by something about David/the report, but not the STM volitionally frightening Clara. This stative structure is associated with an unmarked OVS word order in all-focus sentences (cf. Fábregas et al., 2017; Jiménez-Fernández & Rozwadowska, 2016). On the contrary, accusative constructions, such as (3), are eventive and entail a CoS (Fábregas et al., 2017); the external argument is generally perceived as a volitional animate interpreted as a causer (CSR). As in Landau (2010), these structures are considered bi-eventive and the unmarked word order in all-focus sentences is SVO (cf. Fábregas et al., 2017).

There is data showing that there is a correlation between the DAT/ACC alternation of the EXP and the theta-role of the STM, where the SM appears in dative stative structures and the CSR in accusative eventive ones. In fact, verbs such as *gustar* (class III) that only assign DAT to their EXP are no distinct from the dative alternant of class II in that they are stative non-agentive constructions, with no CoS (cf. Landau, 2010; Reinhart, 2002), and the STM is perceived as the SM (6).

- (6) [A Clara]_{DAT.EXP} (le) gusta [David / el reporte]_{NOM.STM}.
to Clara CL.ACC likes David the report
‘Clara likes David / the report.’

In order to have a clearer mapping of roles and case marking, a more detailed distinction of theta-roles needs to be made (cf. Fig. 1). Psych-verbs in their causative eventive constructions can present two different sources of emotion: (a) an *animate stimulus-causer* (e.g. *David* in (3)), who has control over the event and directly causes a psychological state in the EXP; and (b) an *inanimate stimulus-causer* (e.g. *the report* in (3)) that directly triggers the emotion in the EXP (Pesetsky, 1995).²

In addition, data from Spanish show that there is also an alternation of the STM in ES structures, and this alternation is related to the interpretation of the target (TG) vs. SM distinction. Traditionally, it has been said that class I stative predicates assign ACC to their objects, as in the case of *amar* ‘love’ in (7). However, there are lexical items in this class that are more frequently

²As in Alexiadou & Iordachioaia (2014), we separate the *agent* from the *causer*. We further differentiate the *stimulus* of psych predicates, which includes a *stimulus-causer*, from that of a *pure-causer* occurring in non-psych-verb constructions (e.g. *Peter broke the vase*) (see Fig. 1).

found in DAT structures such as *temer* ‘fear’ and *admirar* ‘admire’ in (4). In addition, *amar* ‘love’ also appears in more marked DAT sentences like (8). The same is true for *temer* ‘fear’ items, with a more marked ACC alternant as (9) shows.

- (7) [David]_{EXP} (la) ama [a Clara]_{TG}.
David CL.ACC loves to Clara
‘David loves Clara.’
- (8) [David]_{EXP} le ama [las manos]_{TG.ACC} [a Clara]_{SM.DAT}.
David CL.DAT loves the hands to Clara
‘David loves (something about) Clara, her hands.’
- (9) [David]_{EXP} la teme [a Clara]_{SM.ACC}.
David CL.DAT fears to Clara
‘David fears Clara.’

As pointed out before, there is a correlation between DAT structures and the SM. In (8) and (9), Clara is the SM. The interpretation that is obtained is that David (constantly) loves/fears something about Clara (there is no CoS in the EXP). The other argument (i.e. TG) corresponds to what is being loved or feared by David, which in this case is ‘the hands’. Consequently, we understand the TG in lines of Seres & Espinal (2018): an individual entity, familiar to the EXP, with no abstract reference, where the emotion is targeted to. The presence of the SM in ES sentences implies that there is another argument that is not compelled to be realized in the syntax, but it is semantically implied (i.e. TG). The contrary is not possible, i.e. a TG semantically implying the existence of the SM.

The interaction of theta-roles and the distinct case marking of both EXP and STM has an impact in linearization yielding different unmarked word orders and further specifying the sub-classes proposed by Belletti & Rizzi (1988). As seen in (3), the transitive configuration of the psych-verb sentences resembles the default (canonical) linearization of verbs with an agent subject and a patient object (SUBJ_{AG-NOM} > OBJ_{PAT-ACC}). However, (2) deviates from that configuration placing the EXP in fronting position. This word order has been attributed to the subject-like properties of the EXP able to bind an anaphoric element (cf. Reinhart, 2002; Temme & Verhoeven, 2016), to show non-canonical passivization (cf. Grimshaw, 1990; Landau, 2010), and to accept extraction from direct objects (cf. Belletti & Rizzi, 1988). As a result, EO verbs can be distinguished into two classes: (a) class 1, which subsumes verbs such as *gustar* ‘like’ and the DAT alternant of *asustar* ‘frighten’ in one group, placing the EXP in fronting position; and (b) class 2, that only contains psych-verbs in eventive structures, and hence, yielding a preferable STMCSR-NOM > EXP-ACC alignment (cf. Tab. 1). On the contrary, ES verbs always place the EXP in fronting position and the STM

as the object. Furthermore, due to their stimulus DAT/ACC alternation, this class can be divided into: (a) class 3, with a more prototypical ACC marking of the STM (e.g. *amar* ‘love’, *odiar* ‘hate’); and (b) class 4, with a more prototypical DAT marking of the object (e.g. *temer* ‘fear’, *admirar* ‘admire’). Table 1 summarizes the properties previously described for Spanish psych-verbs yielding a new fourfold classification.

example	type	θ -role & case		eventuality	unmarked WO	class
		STM	EXP			
<i>gustar</i>	EO	SM-NOM	DAT	state (−CoS)	EXP-DAT > SMNOM	1
<i>asustar</i>	EO	SM-NOM	DAT	state (−CoS)	EXP-DAT > SM-NOM	1
		STMCSR-NOM	ACC	event (+CoS)	STMCSR-NOM > EXP-ACC	2
<i>amar</i>	ES	TG-ACC	NOM	state (−CoS)	EXP-NOM > TG-ACC	3
		SM-DAT ³	NOM	state (−CoS)	EXP-NOM > SM-DAT	4
<i>temer</i>	ES	TG-ACC ³	NOM	state (−CoS)	EXP-NOM > TG-ACC	3
		SM-DAT	NOM	state (−CoS)	EXP-NOM > SM-DAT	4

Table 1: Properties of Spanish psych-verbs

2.2 Korean

In the case of Korean, ES psych-verbs participate in double nominative (NOM-NOM) stative constructions, where both the EXP and the STM are assigned nominative case (Kim & Choi, 2004). Linearization is strict in NOM-NOM sentences with the EXP preceding the STM (i.e. word order freezing effects), as in (10)⁴.

- (10) [Mina-ka/-nun]_{EXP} [khun soli-ka / Minho-ka]_{STM} mwusepta.
Mina-NOM/-TOP big noise-NOM Minho-NOM is.scary
‘Mina is scared of the big noise / Minho.’

Corpus studies and elicitation tasks have shown that NOM-NOM constructions are more limited in the psych domain, and that the preferred structure is that of the EXP being assigned the topic (TOP) marker (Kim, 2008).⁵ In addition, double nominative sentences are subject to participate in case marking alternation and are also considered stative (Kim, 2008). In

³As mentioned previously, verbs like *amar* ‘love’ and *temer* ‘fear’ show the same kind of alternation, hence belonging to classes 3 and 4. However, the former prototypically assigns ACC to its object, whereas the latter normally assigns DAT to its object. This distinction leaves the SM-DAT for *amar* ‘love’ and TG-ACC for *temer* ‘fear’ more marked, but nevertheless possible. In addition, Spanish has ES psych-verbs that only assign ACC to their objects (e.g. *compadecer* ‘feel sorry for’), and those that only assign DAT (e.g. *codiciar* ‘covet’).

⁴We use the Yale Romanization for the examples in Korean.

⁵According to Yoon (2004), both NOM and TOP are structural case markers. We follow Yoon (2004) and treat NOM and TOP as variants of the first case assigned by the Case Principle (cf. Section 4).

terms of which argument alternates in case, Nam (2015) groups the verbs into two classes according to what she calls “causing sub-events”, where: (a) agentive experiencer predicates (AEP) alternate the EXP between NOM and DAT; while (b) patientive experiencer predicates (PEP) alternate the STM between NOM and DAT.⁶ We propose, however, that this alternation has to do with the event structure of the verbs and the theta-roles assigned to the STM, instead of a classification of causing sub-events.

Recent studies propose that a subclass of state has to be distinguished, namely inchoative states (Bar-el, 2005). For Korean, Choi & Demirdache (2014) and Choi (2015) claimed that there are two types of stative predicates: (a) *pure (typical) states*, which are atelic; and (b) *inchoative states*, items which entail a CoS due to a zero affixation of a BECOME operator in the lexical item. In the psych domain, this corresponds to (a) ES pure states consisting of verbs/adjectives (e.g. *mwusepta* ‘scary’), and (b) ES inchoative psych-verbs comprising inherently inchoative verbs (e.g. *ccacungnanta* ‘get irritated’). Looking at the Korean data, the distinction between SM/TG proposed here for Spanish is also productive in this language. In a sentence like (11) with pure state verbs, the STM is perceived as a SM; i.e. Mina is scared of something about the big noise. However, in sentences like (12) with inchoative psych-verbs, the STM is considered a TG; i.e. Mina directs her emotion of being irritated towards Minhø, a known entity by the experiencer.

- (11) [khun soli-ka/-nun]_{SM} [Mina-eykey]_{EXP} mwusepta.
big noise-NOM/-TOP Mina-DAT is.scary
‘(Something about) the big noise is scary to Mina.’
- (12) [Mina-ka/-nun]_{EXP} [Minho-eykey]_{TG} ccacungnanta.
Mina-NOM/-TOP Minhø-DAT gets.irritated
‘Mina gets irritated at Minhø.’

As in Spanish, the Korean data show that there is case marking alternation for both the EXP and STM between NOM and DAT case, but contrary to Spanish, both Korean pure states and inchoative psych-verbs do not allow for the co-occurrence of the SM and TG in the same structure (cf. (13) vs. (8), (9)).

- (13) [Minho-ka]_{SM} [*sengkyek-ul]_{TG} [Mina-eykey]_{EXP} mwusepta
Minho-NOM character-ACC Mina-DAT is.scary
‘Minho his character is scary to Mina.’

In terms of linearization, again the interaction of theta-roles, case marking and event structure plays a role in the different unmarked word order

⁶According to Nam (2015), the EXP plays the role of agent in the experiential causing sub-event in AEP structures; while in PEP, the EXP plays the role of patient or theme.

alignments. In Korean, double nominative constructions present word order freezing effects. However, the alternation of one of the arguments in DAT case allows for free word order. Correlating Nam’s (2015) classification of AEP with pure states and PEP with inchoative psych-verbs w.r.t. case marking alternations; we observe that Korean shows the following unmarked word order: (a) pure states prefer the EXP-DAT argument placed in object position and the SM-NOM in fronting position (cf. (11)), whereas (b) inchoative psych-verbs place the EXP-NOM in fronting position while the TG-DAT is the object in the sentence (cf. (12)). Parallel to Spanish, this leads us to have a fourfold classification of psych-verbs, as presented in Table 2.

example	type	θ role & case		eventuality	unmarked WO	class
		STM	EXP			
<i>mwusepta</i>	ES	SM-NOM	NOM	state (−CoS)	EXP-NOM > SM-NOM	1
	EO	SM-NOM	DAT	state (−CoS)	SM-NOM > EXP-DAT	2
<i>ccacungnata</i>	ES	TG-NOM	NOM	inch (+CoS)	EXP-NOM > TG-NOM	3
	ES	TG-DAT	NOM	inch (+CoS)	EXP-NOM > TG-DAT	4

Table 2: Properties of Korean psych-verbs

3 Restructuring predicates in HPSG

Similar to Koenig (1999) and Davis & Koenig (2000), we are not assuming a hierarchy based approach to theta-roles and linking along the lines of Baker (1998), Pesetsky (1995), a.o. Moreover, we are providing a constraint-based analysis of theta-roles and linking. In contrast to the classic treatments of predicates in HPSG, we are proposing two main changes that helps us to achieve a more elegant analysis: we model predications in a neo-Davidsonian style and theta-roles not as attributes, but as types.

3.1 A neo-Davidsonian treatment in HPSG

In HPSG, the treatment of theta-roles is typically Davidsonian (cf. Davidson, 1967), i.e. a predicate is seen as a relation between an event and its arguments. For instance, the semantics of the verb *to love* is represented as the CONT value in (14). It introduces a relation (of type *love-rel(ation)*) between three arguments: an event and two theta-roles. The arguments are modeled as attribute-value pairs such that ARG0 takes an event ($\boxed{1}$), and the STM and the EXP⁷ take indices as values. The value of ARG0 is structure-shared with the value of IND(EX), i.e. the verb *to love* denotes an event(uality) (cf. fn. 17).

⁷In different HPSG-accounts, arguments of relations have been modeled in different ways: as very predicate-specific attributes, e.g. LOVER and LOVEE (Pollard & Sag, 1987); as non-specific attributes, e.g. ARG1 and ARG2 (Copestake et al., 2005); as proto-role-like attributes, e.g. ACTOR and UNDERGOER (Davis & Koenig, 2000).

$$(14) \left[\begin{array}{c} \text{CONT} \\ \text{RELS} \end{array} \left[\begin{array}{cc} \text{IND} & \boxed{1} \text{ event} \\ \left\langle \begin{array}{cc} \text{ARG0} & \boxed{1} \\ \text{STM} & \text{index} \\ \text{EXP} & \text{index} \end{array} \right\rangle & \text{love-rel} \end{array} \right] \right]$$

The problem of the (strict)⁸ Davidsonian approach is that it does not allow for the manipulation of arguments. That is to say, we cannot simply add arguments to the relation or delete them without assuming a new predicate. For instance, the verb *to kick* in (15) realizes two syntactic arguments: *Luise*, interpreted as the agent, and *Jacob*, interpreted as the patient. In (16), the verb *to kick* could be interpreted in two different ways, cf. (16a) and (16b).

(15) Luise kicked Jacob.

(16) Luise kicks very elegantly.

- a. Luise kicks some person *x*, *x* is semantically implied, but syntactically not realised.
- b. Luise strikes out with her foot – without implying the existence of a target of the kick – e.g., doing martial arts.

For the interpretations intended in (15) and (16a), one single relation (cf. (17)) can be proposed. The difference between them can be modeled treating the object of *kick* as syntactically optional, but as present in the semantics of the predicate, hence semantically implied. For (16b) though, a different relation (cf. (18)) must be assumed, since (17) is defined for three arguments, and for (16b) no object is semantically implied.

$$(17) \begin{bmatrix} \text{ARG0} & \text{event} \\ \text{AG} & \text{index} \\ \text{PAT} & \text{index} \\ \text{kick}^1\text{-rel} \end{bmatrix}$$

$$(18) \begin{bmatrix} \text{ARG0} & \text{event} \\ \text{AG} & \text{index} \\ \text{kick}^2\text{-rel} \end{bmatrix}$$

Since the verb predication in (15) and (16) is actually the same, the only interpretative difference being the (non-)implication of a patient-argument, it would be desirable to have a semantic representation that avoids the necessity of two different *kick*-relations, i.e. (17) and (18). Thus, we are proposing a neo-Davidsonian approach⁹ along the lines of Parsons (1990) (cf. (19) and (20)), that allows us to manipulate the arity of predicates without having to assume different predicates (e.g. *kick*¹ and *kick*² in (17)

⁸Some Davidsonian analyses allow to add but not to delete arguments from a relation, see e.g. the analysis of benefactives in Müller (2018, 69).

⁹A neo-Davidsonian approach for HPSG has also been proposed in Copestake (2006) for independent reasons.

and (18)). In other words, the *kick-rel* tells us something about the kind of eventuality denoted by the predicate and the intension of the verb. The theta-roles related to the predicate are included in the RELS list as single elementary predications (EPs), linked to the main predicate via the value of the ARG1 attributes of the theta-roles.

$$(19) \left[\text{RELS} \left\langle \left[\begin{array}{cc} \text{ARG0} & \boxed{1} \\ \text{kick-rel} & \text{event} \end{array} \right], \left[\begin{array}{cc} \text{ARG0} & \text{index} \\ \text{ARG1} & \boxed{1} \\ \text{agent} & \end{array} \right], \left[\begin{array}{cc} \text{ARG0} & \text{index} \\ \text{ARG1} & \boxed{1} \\ \text{patient} & \end{array} \right] \right\rangle \right]$$

$$(20) \lambda y \lambda x \lambda e. \text{kick}(e) \wedge \text{agent}(x)(e) \wedge \text{patient}(y)(e)$$

Handling theta-roles as EPs is a further change we are proposing (cf. (19) and (17)) since we model theta-roles as relations between events and individuals in the spirit of neo-Davidsonian approaches (cf. (20)). For instance, the *kick-rel* is a predication of type *event(uality)* and the elements interpreted as agent and patient of the predication are objects of type *index*. The *agent* and *patient* types are relations between the value of ARG0 and the value of ARG1.

3.2 Underspecification of theta-roles

In line with the previous neo-Davidsonian approach, we are analyzing theta-roles as types.¹⁰ These types are modeled along an inheritance hierarchy for semantic relations (*sem-rels*). In this hierarchy (cf. Fig. 1),¹¹ theta-roles (*θ-role*) and predicates (*pred*) are subtypes of *sem-rels*. This reflects the way they are being modeled in the RELS list (cf. (19)), i.e. as conjoined EPs of the same (super-)type (i.e. *sem-rels*).

Modeling theta-roles as in Fig. 1 allows us to establish commonalities and differences among them by means of (multiple) inheritance. This classification is needed for theoretical as well as for empirical reasons. For instance, theoretically, it allows us to define psych-predicates as an eventuality involving an experiencer (*exp*) and a stimulus (*stm*), although the stimuli can be differentiated into: subject matter (*sm*), target (*tg*), and stimulus-causer (*stmcsr*). As it has been shown in (7)–(9), these different classes of stimuli are empirically needed in order to have, for instance, a more appropriate account for word order and case assignment for psych-predicates.

¹⁰Davis & Koenig (2000, 70–71) and Van Eynde (2015, 109–113) also treat theta-roles as types, but with a Davidsonian approach. That is, it is not the (definition of the) theta-role itself that is more specific along the inheritance hierarchy, but the Davidsonian relation, i.e. their EP gets more attribute-value pairs (representing theta-roles) along the hierarchy. In our approach, the hierarchy of type *θ-role* reflects an *ontology* of theta-roles.

¹¹Figure 1 depicts by no means an exhaustive representation of theta-roles. For the time being, we are focusing only on the relevant theta-roles for psych-predicates, i.e. only on the types *experiencer* and *stimulus*.

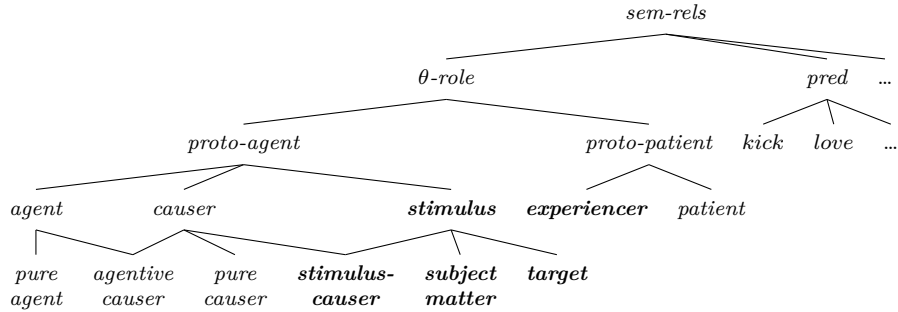


Figure 1: Type hierarchy for *semantic-relations*

Furthermore, this hierarchy allows for the modeling of generalizations of theta-roles and verb classes (e.g. psych-verbs). By means of multiple inheritance, we can account for entities that have the properties of causers as well as the properties of stimuli (e.g. *stimulus-causer*) without having to choose whether we are dealing with a causative or a psych-verb, since it could be both. Therefore, it is expected that some generalizations concerning stimuli will also affect some subset of causers, and some generalizations applying to causers will affect some subset of stimuli.

In a further state of the theory, this approach enables us to define theta-roles by means of constraints assigning semantic properties to their subtypes. This is one of the main differences between our proposal and e.g. Davis & Koenig (2000). They define theta-roles by means of (disjunctive) sets of characteristic entailments (Davis & Koenig, 2000, 72) and work mostly with proto-roles, similar to Dowty (1991). In their analysis, characteristic entailments are model-theoretic constraints, which do not belong to the descriptive language of the grammar. Therefore, “their satisfaction cannot be checked by looking at the metalanguage [...] use[d] in our descriptions” (Davis & Koenig, 2000, 72–73). Characteristic entailments are thus not properly part of the (described) grammatical system, but rather of some kind of meta-grammar. Davis & Koenig’s approach is mostly concerned with linking and word classes modeled through constraints in an inheritance hierarchy. We follow their approach to linking in many respects, but the empirical data in the psych domain force us to assume a different treatment of theta-roles (as specific neo-Davidsonian types) in order to achieve a more fine grained distinction of the verbs. To some extent, we take advantage of the analyses of proto-roles (Dowty, 1991), of (proto-)theta-roles as characteristic entailments, of linking as constraints in an inheritance hierarchy (Davis & Koenig, 2000; Van Eynde, 2015), and of hierarchy-based modeling of theta-roles (Baker, 1998; Belletti & Rizzi, 1988; Pesetsky, 1995).

Our analysis reflects the idea of proto-roles¹² via different levels of ab-

¹²Proto-roles are divided into proto-agent vs. proto-patient or actor vs. undergoer

straction encoded in the inheritance hierarchy. As such, a type *proto-agent* could be proposed as having e.g. *agent*, *causer*, and *stimulus* as subtypes, and being less constrained than its subtypes. As far as the empirical data suggest –i.e. some generalizations apply to this kind of (proto-)supertype– the assumption of such proto-roles is descriptively well-founded. For our current goal –the analysis of linking relations in the psych domain– and due to lack of space only two types (and their subtypes) will be considered: *stimulus* and *experiencer*.

4 Analysis of Spanish psych-verbs

As pointed out in Section 2.1, the data demand a fourfold classification for Spanish psych-verbs.¹³ For the issue in question, we are assuming that the linking properties can be modeled lexically by means of an inheritance hierarchy (cf. Fig. 2) constraining the properties of different types of lexemes (cf. Manning & Sag, 1998, 124–125; Davis & Koenig, 2000, 67; Van Eynde, 2015, 115).

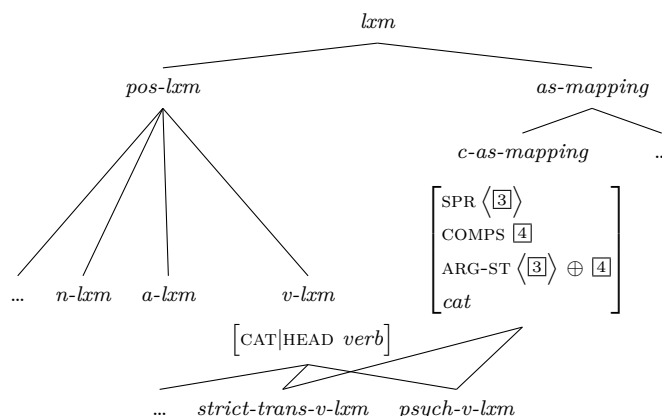


Figure 2: Type hierarchy for *lexeme*

The type *lexeme* (*lxm*) has two subtypes: *part-of-speech lexeme* (*pos-lxm*) and *argument-structure mapping* (*as-mapping*). The type *pos-lxm* constrains the HEAD value of lexemes, i.e. for *verb lexemes* (*v-lxm*) the HEAD value is of type *verb*. The *as-mapping* type¹⁴ constrains the correspondence between

in Dowty (1991) or Davis & Koenig (2000), respectively. We are using the former denomination.

¹³Due to space issues, we cannot provide a complete theory of linking in this paper. Our main goal here is just to provide a descriptive and more adequate treatment of Spanish and Korean psych-verbs, their properties and subclasses.

¹⁴Our *lxm* hierarchy is similar to the one proposed in Van Eynde (2015, 115). One difference we would like to point out here is that our *as-mapping* type only resembles Van Eynde’s *linking* type. We consider “linking” the relation between semantic and syn-

elements in the ARG-ST list and elements in the valence features (for Spanish: SPR and COMPS). Its subtype *canonical-as-mapping* (*c-as-mapping*) constrains the “canonical” correspondence for verbs in Spanish, thus passing its constraint (by means of multiple inheritance) to *strict-transitive verb lexeme* and – as we will see later – also to *psych-verb lexeme*.

As already mentioned, psych-verbs can be divided into two subclasses: ES and EO psych-verbs, each of which can be subdivided into two further subclasses: ES with accusative object, ES with dative object, EO with case alternation and EO without alternation (cf. Fig. 3).¹⁵ The *psych-v-lxm* type constrains the mapping of semantic arguments to the elements in the ARG-ST list (cf. the *linking* type in Van Eynde, 2015). The elements in the ARG-ST list are normally ordered according to their prominence w.r.t. case, binding, extraction, etc. (cf. Manning & Sag, 1998, 111; Koenig, 1999, 29; Müller, 2016, 295; a.o.).

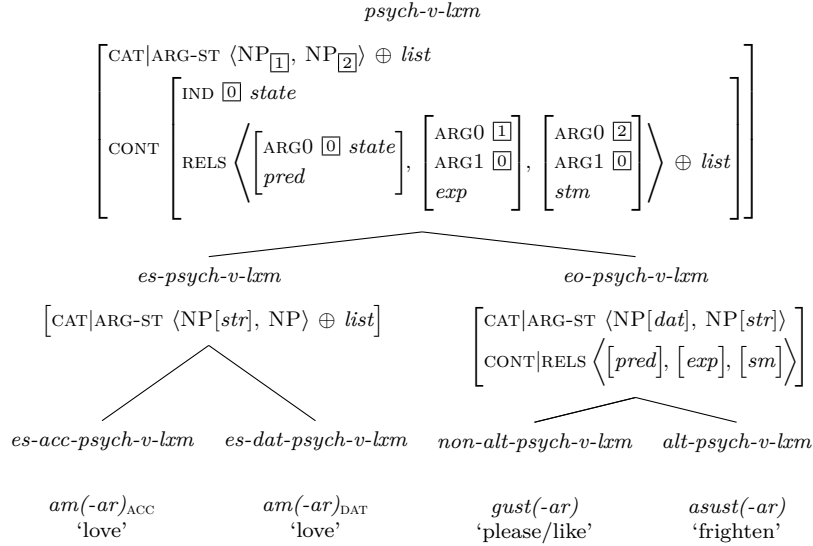


Figure 3: Type hierarchy for *psych-v-lxm* in Spanish (neo-Davidsonian)

In our analysis, the experiencer of psych-verbs is linked to the first element in the ARG-ST list, and the stimulus to the second element.¹⁶ Further elements could be considered (cf. $\oplus \text{list}$); however this is not required, since *list* could be further specified as being of type *empty list* (*e-list*), e.g. for objects of type *eo-psych-v-lxm*. In addition, *psych-v-lxm* constrains psych-verbs

tactic arguments of lexemes (cf. Machicao y Priemer, 2018). The constraint relating the elements of the ARG-ST list to the elements of the valence features –our *as-mapping*– is only *a part* of the whole linking concept. Furthermore, Van Eynde’s *linking* type is different in that it relates the semantic arguments to the elements in the ARG-ST list.

¹⁵A further class will be derived by means of a lexical rule (cf. Fig. 5).

¹⁶Having the experiencer as the first element of the ARG-ST list reflects the psych effects of the experiencer seen as a quirky subject.

as being of eventualities of type *state*.¹⁷ The type *es-psych-v-lxm* constrains the first element of the ARG-ST list (the experiencer) as having structural case, while *eo-psych-v-lxm* constrains the experiencer as a dative object and the stimulus as an NP bearing structural case. Moreover, *eo-psych-v-lxm* limits the ARG-ST list as having only these two arguments.

The two subtypes *es-acc-psych-v-lxm* and *es-dat-psych-v-lxm* add the further constraints needed in order to differentiate between *amar* ‘to love’ with accusative and dative (cf. Fig. 4).

$$\begin{aligned}
es-acc-psych-v-lxm &\Rightarrow \left[\begin{array}{l} \text{CAT|ARG-ST } \langle \text{NP}, \text{NP}[\textit{str}] \rangle \\ \text{CONT|RELS } \langle [\textit{pred}], [\textit{exp}], [\textit{tg}] \rangle \end{array} \right] \\
es-dat-psych-v-lxm &\Rightarrow \left[\begin{array}{l} \text{CAT|ARG-ST } \langle \text{NP}, \text{NP}[\textit{dat}] \rangle \oplus \langle \text{NP}[\textit{str}]_{\boxed{6}} \rangle \\ \text{CONT|RELS } \left\langle \left[\begin{array}{c} \text{ARG0 } \boxed{0} \\ \textit{pred} \end{array} \right], [\textit{exp}], [\textit{sm}] \right\rangle \oplus \left\langle \left[\begin{array}{c} \text{ARG0 } \boxed{6} \\ \text{ARG1 } \boxed{0} \\ \textit{tg} \end{array} \right] \right\rangle \end{array} \right]
\end{aligned}$$

Figure 4: Constraints for ES verbs

For ES verbs with an accusative object, the second element in the ARG-ST list gets also structural case and the theta-role of the stimulus is further specified as being a target. For ES verbs with a dative object, the second element in the ARG-ST list is specified as bearing dative and its theta-role is specified as subject matter. For NPs with structural case, case assignment follows the Case Principle (cf. Meurers, 1999, 204; Przepiórkowski, 1999, 93–94; a.o.), i.e. the first element in the ARG-ST list with structural case gets nominative, while further elements with structural case get accusative. A further important distinction between *es-acc-psych-v-lxm* and *es-dat-psych-v-lxm* is that the latter has an additional optional object (cf. (8)). This object is interpreted as a target and bears structural case, i.e. accusative.

Lexemes of type *eo-psych-v-lxm* are divided into two subtypes: a non-alternating type *non-alt-psych-v-lxm* for lexemes such as *gust(-ar)* ‘to like’ and an alternating one *alt-psych-v-lxm* for lexemes such as *asust(-ar)* ‘to frighten’. The alternation shown in (2)–(3) can be modeled by means of the LR in Figure 5.

This LR takes stative predicates with an experiencer-dative and a subject matter-nominative as input (to be more precise: elements of type *alt-psych-v-lxm*, see also (2)).¹⁸ The output of the LR represents an object in which the experiencer $\boxed{1}$ is realized with structural accusative, the aforementioned subject matter argument is deleted (represented in the LR-input as *nelist*),

¹⁷As a working hypothesis, we assume an ontology of eventualities similar as the one proposed by Bach (1986) with *state* as a subtype of *eventuality*.

¹⁸The distinction between *non-alt-psych-v-lxm* and *alt-psych-v-lxm* is important, since the LR takes only elements of the latter type as input, even if no other differences can be stated between these two types, yet.

$$\begin{array}{c}
\left[\begin{array}{c} \text{CONT} | \text{RELS } \boxed{8} \oplus \text{nelist} \\ \text{alt-psych-v-lxm} \end{array} \right] \mapsto \\
\left[\begin{array}{c} \text{CAT} | \text{ARG-ST } \langle \text{NP}[\text{str}]_{\boxed{5}}, \text{NP}[\text{str}]_{\boxed{1}} \rangle \\ \text{CONT} \left[\begin{array}{c} \text{IND } \boxed{4} \\ \text{RELS } \boxed{8} \left\langle \begin{array}{c} \left[\begin{array}{c} \text{ARG0 } \boxed{0} \\ \text{pred} \end{array} \right], \left[\begin{array}{c} \text{ARG0 } \boxed{1} \\ \text{exp} \end{array} \right] \end{array} \right\rangle \oplus \left\langle \begin{array}{c} \left[\begin{array}{c} \text{ARG0 } \boxed{4} \text{ hpng} \\ \text{ARG1 } \boxed{0} \\ \text{begin-pred} \end{array} \right], \left[\begin{array}{c} \text{ARG0 } \boxed{5} \\ \text{ARG1 } \boxed{4} \\ \text{csr} \end{array} \right] \end{array} \right\rangle \end{array} \right] \\ \text{cause-psych-v-lxm} \end{array} \right]
\end{array}$$

Figure 5: LR for case alternation for *alt-psych-v-lxm*

and a new semantic argument – a causer $\boxed{5}$ – is added to the RELS list. The causer is mapped to the first element of the ARG-ST list and is realized with structural nominative (see $\boxed{5}$). This new arrangement in the ARG-ST list has consequences for the mapping to SPR and COMPS, i.e. in the unmarked word order the experiencer is not going to precede the other arguments anymore (cf. Fig. 6), see e.g. (3). Moreover, the output of the LR is an eventuality of a different subtype; i.e. it is not a state $\boxed{0}$ anymore – as the input of the LR – but a happening $\boxed{4}$ (cf. fn. 17). Therefore, *cause-psych-v-lxm* is not a subtype of *psych-v-lxm* (cf. Fig. 3).

$$\left[\begin{array}{c} \text{CAT} \left[\begin{array}{c} \text{SPR } \langle \boxed{2} \rangle \\ \text{COMPS } \langle \boxed{3} \rangle \\ \text{ARG-ST } \langle \boxed{2} \text{ NP}[\text{str}]_{\boxed{5}}, \boxed{3} \text{ NP}[\text{str}]_{\boxed{1}} \rangle \end{array} \right] \\ \text{CONT} \left[\begin{array}{c} \text{IND } \boxed{4} \\ \text{RELS } \left\langle \begin{array}{c} \left[\begin{array}{c} \text{ARG0 } \boxed{0} \text{ state} \\ \text{pred} \end{array} \right], \left[\begin{array}{c} \text{ARG0 } \boxed{1} \\ \text{exp} \end{array} \right], \left[\begin{array}{c} \text{ARG0 } \boxed{4} \text{ hpng} \\ \text{ARG1 } \boxed{0} \\ \text{begin-pred} \end{array} \right], \left[\begin{array}{c} \text{ARG0 } \boxed{5} \\ \text{ARG1 } \boxed{4} \\ \text{csr} \end{array} \right] \end{array} \right\rangle \end{array} \right] \end{array} \right]$$

Figure 6: *asustar* with ACC

5 Analysis of Korean psych-verbs

For Korean, the inheritance hierarchy for the type *lxm* is similar to the one shown for Spanish (cf. Fig. 2), but since Korean is an SOV language (allowing scrambling), it is not necessary to assume a SPR attribute (cf. Müller, 2016, 293–295). Hence, canonically all elements in the ARG-ST list are mapped in the same order to the COMPS list of the lexeme, as shown in Figure 7.

The type *psych-v-lxm* links – as in Spanish – the experiencer with the first

$$c-as-mapping \Rightarrow \begin{bmatrix} \text{COMPS } \boxed{1} \\ \text{ARG-ST } \boxed{1} \\ \text{cat} \end{bmatrix}$$

Figure 7: Constraints for *c-as-mapping* for Korean

element of the ARG-ST list and the stimulus with the second one (cf. Fig. 8). Furthermore, the first element of the ARG-ST list is constrained as bearing structural case, i.e. nominative¹⁹ qua Case Principle. Contrary to the Spanish class, Korean psych-verbs are not constrained as stative in general, since this class can be divided into stative psych-verbs (type: *state-psych-v-lxm*) and inchoative psych-verbs (type: *inch-psych-v-lxm*).

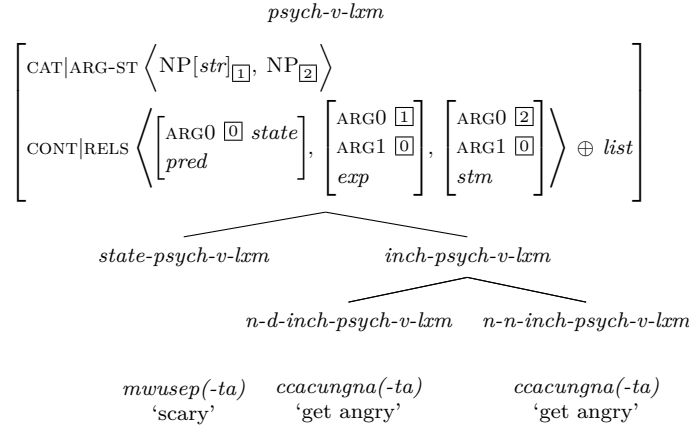


Figure 8: Type hierarchy for *psych-v-lxm* in Korean (neo-Davidsonian)

For elements of type *inch-psych-v-lxm* a further eventuality (i.e. a *begin-predication*) is introduced. This is an eventuality of type *happening* that takes the stative predication as argument (cf. value of ARG1 of *begin-pred* in Fig. 9). The IND value of *inch-psych-v-lxm* is the happening ($\boxed{5}$), not the state ($\boxed{0}$). Moreover, the stimulus argument is further specified as target (cf. (12)). In contrast, for elements of type *state-psych-v-lxm*, the theta-role of the stimulus is further specified as subject matter bearing lexical nominative, and the eventuality type of the predication is identified as a state (cf. Fig. 9).

The case alternation for lexemes of type *inch-psych-v-lxm* can be constrained by means of two types: *n-d-inch-psych-v-lxm* assigning dative to the target, and *n-n-inch-psych-v-lxm* assigning lexical nominative to it (cf. Fig. 10). The distinction between these two types concerns only case marking, neither theta-roles nor eventuality type are different.

¹⁹For Korean, we assume that the first element of the ARG-ST list with structural case gets nominative or topic case, being both just variants.

$$\begin{aligned}
state-psych-v-lxm &\Rightarrow \left[\begin{array}{c} \text{CAT|ARG-ST } \langle \text{NP}, \text{NP}[lnom] \rangle \\ \text{CONT} \left[\begin{array}{c} \text{IND } [0] \\ \text{RELS } \left\langle \left[\begin{array}{c} \text{ARG0 } [0] \text{ state} \\ \text{pred} \end{array} \right], [exp], [sm] \right\rangle \end{array} \right] \end{array} \right] \\
inch-psych-v-lxm &\Rightarrow \left[\begin{array}{c} \text{CONT} \left[\begin{array}{c} \text{IND } [5] \\ \text{RELS } \left\langle \left[\begin{array}{c} \text{ARG0 } [0] \text{ state} \\ \text{pred} \end{array} \right], [exp], [tg] \right\rangle \oplus \left\langle \left[\begin{array}{c} \text{ARG0 } [5] \text{ hpng} \\ \text{ARG1 } [0] \\ \text{begin-pred} \end{array} \right] \right\rangle \end{array} \right] \end{array} \right]
\end{aligned}$$

Figure 9: Constraints for stative and inchoative verbs in Korean

$$\begin{aligned}
n-d-inch-psych-v-lxm &\Rightarrow \left[\text{CAT|ARG-ST } \langle \text{NP}, \text{NP}[dat] \rangle \right] \\
n-n-inch-psych-v-lxm &\Rightarrow \left[\text{CONT|ARG-ST } \langle \text{NP}, \text{NP}[lnom] \rangle \right]
\end{aligned}$$

Figure 10: Constraints for NOM-DAT and NOM-NOM verbs in Korean

For the alternation applying to elements of type *state-psych-v-lxm*, we need a LR (cf. Fig. 11) that makes changes in case assignment and word order, cf. (10) vs. (11). With respect to case marking, the experiencer [1], which bears *str* in the input, takes dative in the output. The stimulus [2], bearing *lnom* in the input, takes *str* in the output. Additionally, with respect to unmarked word order, the mapping of ARG-ST and COMPS in the output does not follow the *c-as-mapping* in Figure 7, i.e. we do not have an experiencer first structure anymore. Instead, the NP interpreted as stimulus [6] precedes the NP interpreted as experiencer [5].

$$\left[\begin{array}{c} \text{CAT|ARG-ST } \langle \text{NP}_{[1]}, \text{NP}_{[2]} \rangle \\ state-psych-v-lxm \end{array} \right] \mapsto \left[\begin{array}{c} \text{CAT} \left[\begin{array}{c} \text{COMPS } \langle [6], [5] \rangle \\ \text{ARG-ST } \langle [5]\text{NP}[dat]_{[1]}, [6]\text{NP}[str]_{[2]} \rangle \end{array} \right] \\ n-d-state-psych-v-lxm \end{array} \right]$$

Figure 11: LR for case alternation for *state-psych-v-lxm*

6 Conclusion

The main goal of this paper was to give a detailed description of psych-verbs in Spanish and Korean. The different lexemes that can be subsumed under the label *psych-verb* show diverging characteristics as well as commonalities. We have focused mostly on the correlations between case marking, theta-role assignment, and eventuality types in order to describe the distinct psych-verb subclasses in the languages at hand.

We have proposed a neo-Davidsonian treatment of the predications in

order to be able to account for the underspecification of theta-roles that is needed for a proper description of case alternation in Spanish, and eventuality distinction in Korean. Furthermore, the presented psych-verb hierarchies reflect a possible ontology for the psych domain based on commonalities and differences between the psych-verb subclasses. This allows us, on the one hand, to localise connections between the psych domain and other verb-classes (e.g. between *strict-trans-v-lxm* and *es-acc-psych-v-lxm* in Spanish) that could be modelled by means of multiple inheritance – something that we cannot work out here due to lack of space. On the other hand, it shows the diversity of subclasses within the psych domain and illustrates the complexity of the psych-verb class.

Certainly, some aspects of our analysis have to be worked out in more detail. For instance, more work on the inheritance hierarchy of theta-roles is needed to find out on what basis theta-roles can be constrained and which further subclasses are needed. Moreover, the assumption of an inheritance hierarchy based approach on theta-roles has further theoretical consequences for the so called *Theta-Criterion* in the generative literature. The idea that “[e]ach argument bears one and only one theta-role, and each theta-role is assigned to one and only one argument” (cf. Chomsky, 1981, 36) should be reconsidered in the light of underspecified roles.

References

- Alexiadou, Artemis, Elena Anagnostopoulou & Martin Everaert. 2004. *The unaccusativity puzzle: Explorations of the syntax-lexicon interface*. New York: Oxford University Press.
- Alexiadou, Artemis & Gianina Iordachioaia. 2014. The psych causative alternation. *Lingua* 148. 53–79.
- Arad, Maya. 1998. Psych-notes. *UCL Working Papers in Linguistics* 10. 1–22.
- Bach, Emmon. 1986. The algebra of events. *Linguistics and Philosophy* 9(1). 5–16.
- Baker, Mark. 1998. *Incorporation: A theory of grammatical function changing*. Chicago: University of Chicago Press.
- Bar-el, Leora. 2005. *Aspectual distinctions in Skwxwú7mesh*. University of British Columbia PhD thesis.
- Belletti, Adriana & Luigi Rizzi. 1988. Psych-verbs and θ -theory. *Natural Language and Linguistic Theory* 6(3). 291–352.
- Choi, Jiyoung. 2015. On the universality of aspectual classes: Inchoative states in Korean. In Emmanuelle Labeau & Qiaochao Zhang (eds.), *Taming the TAME systems*, 123–135. Leiden: Brill.
- Choi, Jiyoung & Hamida Demirdache. 2014. Reassessing the typology of

- states: Evidence from Korean (degree) inchoative states. Paper presented at Workshop on the Ontology and the Typology of States.
- Chomsky, Noam. 1981. *Lectures on Government and Binding*. Dordrecht: Foris Publications.
- Copestake, Ann. 2006. Robust Minimal Recursion Semantics. Ms. University of Cambridge. <https://www.cl.cam.ac.uk/~aac10/papers.html>.
- Copestake, Ann, Daniel P. Flickinger, Carl J. Pollard & Ivan A. Sag. 2005. Minimal Recursion Semantics: An introduction. *Research on Language and Computation* 3(4). 281–332.
- Davidson, Donald. 1967. The logical form of action sentences. In Nicholas Resher (ed.), *The logic of decision and action*, 81–95. Pittsburgh: University of Pittsburgh Press.
- Davis, Anthony & Jean-Pierre Koenig. 2000. Linking as constraints on word classes in a hierarchical lexicon. *Language* 76(1). 56–91.
- Dowty, David R. 1991. Thematic proto-roles and argument selection. *Language* 67(3). 547–619.
- Fábregas, Antonio, Ángel Jiménez-Fernández & Mercedes Tubino. 2017. What’s up with dative experiencers. *Romance Languages and Linguistic Theory 12: Selected Papers from the 45th Linguistic Symposium on Romance Languages, Campinas, Brazil* 30–47.
- Grimshaw, Jane. 1990. *Argument structure*. Cambridge: MIT Press.
- Jiménez-Fernández, Ángel & Bożena Rozwadowska. 2016. The information structure of dative experiencer psych verbs. In Bożena Cetnarowska, Marcin Kuczok & Marcin Zabawa (eds.), *Various dimensions of contrastive studies*, 100–121. Katowice: University of Silesia Press.
- Kim, Ilkyu. 2008. *On the NOM-DAT alternation of experiencer in Korean: A conceptual semantics approach*. Hankuk University of Foreign Studies PhD thesis.
- Kim, Jong-Bok & Incheol Choi. 2004. The Korean case system: A unified, constraint-based approach. *Language Research* 40(4). 885–921.
- Koenig, Jean-Pierre. 1999. *Lexical relations*. Stanford: CSLI Publications.
- Landau, Idan. 2010. *The locative syntax of experiencers*. London: MIT Press.
- Machicao y Priemer, Antonio. 2014. Differentielle Objektmarkierung: Spezifität und Akkusativ im Spanischen. In Antonio Machicao y Priemer, Andreas Nolda & Athina Sioupi (eds.), *Zwischen Kern und Peripherie*, 103–130. Berlin: De Gruyter.
- Machicao y Priemer, Antonio. 2018. Argumentstruktur. In Stefan Schierholz & Pál Uzonyi (eds.), *Grammatik: Syntax* 1.2, Berlin: De Gruyter.
- Manning, Christopher D. & Ivan A. Sag. 1998. Argument structure, valence, and binding. *Nordic Journal of Linguistics* 21(2). 107–144.
- Marín, Rafael. 2015. Los predicados psicológicos: Debate sobre el estado de la cuestión. In Rafael Marín (ed.), *Los predicados psicológicos*, 11–50. Madrid: Visor.

- Marín, Rafael & Louise McNally. 2011. Inchoativity, change of state, and telicity: Evidence from Spanish reflexive psychological verbs. *Natural Language and Linguistic Theory* 29(2). 467–502.
- Meurers, Walt Detmar. 1999. Raising spirits (and assigning them case). *Groninger Arbeiten zur Germanistischen Linguistik (GAGL)* 43. 173–226. <http://irs.ub.rug.nl/dbi/475d4439a8e3f>.
- Müller, Stefan. 2016. *Grammatical theory: From transformational grammar to constraint-based approaches*. Berlin: Language Science Press. <http://langsci-press.org/catalog/book/25>.
- Müller, Stefan. 2018. *A lexicalist account of argument structure: Template-based phrasal LFG approaches and a lexical HPSG alternative*. Berlin: Language Science Press. <http://langsci-press.org/catalog/book/163>.
- Nam, Seungho. 2015. Lexical semantics: Lexicon-syntax interface. In Lucien Brown & Jaehoon Yeon (eds.), *The handbook of Korean linguistics*, 157–178. Hoboken: John Wiley & Sons.
- Parsons, Terence. 1990. *Events in the semantics of English: A study in subatomic semantics*. Cambridge: MIT Press.
- Pesetsky, David. 1995. *Zero syntax: Experiencers and cascades*. Cambridge: MIT Press.
- Pollard, Carl J. & Ivan A. Sag. 1987. *Information-based syntax and semantics. Volume 1: Fundamentals*. Stanford: CSLI Publications.
- Przepiórkowski, Adam. 1999. *Case assignment and the complement/adjunct dichotomy: A non-configurational constraint-based approach*. Eberhard-Karls-Universität Tübingen PhD thesis.
- Pylkkänen, Liina. 2000. On stativity and causation. In Carol Tenny & James Pustejovsky (eds.), *Events as grammatical objects: The converging perspectives of lexical semantics, logical semantics and syntax*, 417–442. Stanford: CSLI Publications.
- Reinhart, Tanya. 2002. The theta system: An overview. *Theoretical Linguistics* 28(3). 229–290.
- Seres, Daria & M. Teresa Espinal. 2018. Psychological verbs and their arguments. *Borealis* 7(1). 27–44.
- Temme, Anne & Elisabeth Verhoeven. 2016. Verb class, case, and order: A cross-linguistic experiment on non-nominative experiencers. *Linguistics* 54(4). 769–814.
- Van Eynde, Frank. 2015. *Predicative constructions: From the Freagean to a Montagovian treatment*. Stanford: CSLI Publications.
- Van Voorst, Jan. 1992. The aspectual semantics of psychological verbs. *Linguistics and Philosophy* 15(1). 338–345.
- Yoon, James. 2004. Non-nominative (major) subjects and case stacking in Korean. In Peri Bhaskararao & Karumuri Subbarao (eds.), *Non-nominative subjects*, 265–314. Amsterdam: John Benjamins.

The fine structure of clausal right-node raising constructions in Japanese

Shûichi Yatabe 

University of Tokyo

Kei Tanigawa 

University of Tokyo

Proceedings of the 25th International Conference on
Head-Driven Phrase Structure Grammar

University of Tokyo

Stefan Müller, Frank Richter (Editors)

2018

Stanford, CA: CSLI Publications

pages 174–194

Keywords: the lexicalist hypothesis, non-constituent coordination

Yatabe, Shûichi & Kei Tanigawa. 2018. The fine structure of clausal right-node raising constructions in Japanese. In Stefan Müller & Frank Richter (eds.), *Proceedings of the 25th International Conference on Head-Driven Phrase Structure Grammar, University of Tokyo*, 174–194. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2018.11.

 CC-BY 4.0

Abstract

We examine the fine structure of clausal right-node raising constructions in Japanese, and argue that there are sentences in which a tensed verb is right-node-raised out of coordinated tensed clauses as well as sentences in which a verb stem is right-node-raised out of coordinated tenseless phrases. In the latter case, the tense morpheme has to be assumed to take a tenseless complement clause, and we note that the existence of such a structure contradicts the so-called lexicalist hypothesis, according to which a verb stem and the tense morpheme immediately following it always form a morphosyntactic constituent.

1 Introduction

The aim of this paper is to determine the details of the syntactic structure of Japanese sentences like the following, which involves right-node raising (RNR).

- (1) [Hanako ga] [yama e], [Masao ga] [kawa e] itta.
[Hanako NOM] [mountain to] [Masao NOM] [river to] go-PAST
‘Hanako went to the mountain and Masao went to the river.’

In this paper, we assume that the HPSG-based analysis of right-node raising advocated in works such as Yatabe & Tam (2017) is on the right track; in other words, we assume that a sentence like this involves coordination of two normal constituents out of which something is dislocated. Even on that assumption, there remain several possibilities as to what types of syntactic constituent are coordinated in a sentence like (1), and that is the question that will be addressed in this paper.

Before we embark on the main discussion, however, we will briefly consider the following question. Can the sentence above be an instance of some grammatical phenomenon other than right-node raising? Is it not analyzable as an instance of gapping or argument-cluster coordination, for example?

We regard a sentence like (1) as a case of right-node raising rather than a case of gapping (a phenomenon in which a complete clause appears to be coordinated with another clause-like expression in which some expressions appear to have been elided), for the following two reasons. First, an example like (2) indicates that the clause-final expression that seems to be shared by multiple conjuncts in a sentence like (1) belongs (or, at least, *can* belong) syntactically and semantically not just to the final clause but also to the non-final clause(s) as well.

- (2) [Hanako ga] [yama e], [Masao ga] [kawa e], sorezore itta.
[Hanako NOM] [mountain to] [Masao NOM] [river to] individually go-PAST
‘Hanako went to the mountain and Masao went to the river, the two of them acting individually.’

[†]We thank the three anonymous reviewers who commented on the extended abstract and the audience at the HPSG 2018 conference.

The adverb *sorezore* ‘individually’ has the effect of emphasizing the distinctness of the multiple events being described by the clause involved, and cannot be used in front of a verb describing a single event, as shown in (3).

- (3) *[Hanako ga] [yama e] sorezore itta.
 [Hanako NOM] [mountain to] individually go-PAST
 ‘Hanako went to the mountain individually.’

Thus, the fact that *sorezore* can be used in (2) shows that the sentence-final verb expresses (or at least *can* express) not just the event of Hanako going to the mountain but also the event of Masao going to the river. That in turn means that the sentence-final verb belongs to both conjuncts simultaneously, as predicted by the RNR analysis but not by the gapping analysis. Second, the kind of apparent ellipsis that we see in the first conjunct in a sentence like (1) takes place only at the right edge of such a conjunct. This is illustrated by the following examples.

- (4) [Masao wa] ashita, (soshite) [Hanako wa] asatte [nani
 [Masao TOP] tomorrow (and) [Hanako TOP] day after tomorrow [what
 o] kau to yakusoku shita no?
 ACC] buy-PRES COMP promise do-PAST NML
 ‘What has Masao promised to buy tomorrow, and what has Hanako
 promised to buy the day after tomorrow?’
- (5)?*[Masao wa] ashita kau to, (soshite) [Hanako wa]
 [Masao TOP] tomorrow buy-PRES COMP (and) [Hanako TOP]
 asatte [nani o] kau to yakusoku shita no?
 day after tomorrow [what ACC] buy-PRES COMP promise do-PAST NML
 ‘(Same as (4))’

In (4), the first conjunct appears to be missing the string *nani o kau to yakusoku shita no* at its right edge. If what is responsible for this apparent ellipsis is gapping rather than right-node raising, it is expected to be possible to interpret sentence (5) as missing the string *yakusoku shita no* at its right edge and the string *nani o* at the location between *ashita* and *kau to*, yielding a structure that would express the same meaning as (4). Such an interpretation, however, is not available for sentence (5), lending support to the view that the kind of apparent ellipsis we are considering here takes place only at the right edge of a conjunct, as predicted by the RNR analysis. While the first consideration above does not rule out the possibility that Japanese syntax has both right-node raising and gapping, this second consideration arguably allows us to draw a stronger conclusion: Japanese has right-node raising, but not gapping.

Likewise, we do not view a sentence like (1) as a case of argument-cluster coordination (a phenomenon in which arguments of a predicate form a constituent and is coordinated with another constituent consisting of arguments of the same predicate (Mouret (2006))), either, because what appears to be the initial conjunct

in a sentence like (1) does not have to be a sequence of arguments of the same predicate, as shown by an example like (6). In (6), what constitutes the apparent initial conjunct *Hanako wa aoi* is made up of a topicalized nominative subject of the verb *eranda* and an adjective that modifies the noun *kusuri*, and are not co-arguments of the same predicate.

- (6) Hanako wa aoi, (soshite) Masao wa akai kusuri o
 Hanako TOP blue-PRES (and) Masao TOP red-PRES pill ACC
 eranda.
 choose-PAST
 ‘Hanako chose a blue pill, and Masao chose a red pill.’

There is one caveat to keep in mind. Strictly speaking, what sentences like (6) show is that a sentence like (1) *can* be analyzed as a case of right-node raising. They do not rule out the possibility that a sentence like (1) might be syntactically ambiguous between a structure involving right-node raising and one involving argument-cluster coordination. Thus, throughout the present paper, we will make an attempt to base our argumentation on example sentences that are not analyzable as instances of argument-cluster coordination.

In what follows, we will consider the following three possible analyses of clausal right-node raising in Japanese. The first possibility we consider is that sentence (1) may involve coordination of two tensed clauses, as shown in (7).

- (7) [[Hanako ga yama e itta], [Masao ga kawa e itta]]
 → Hanako ga yama e, Masao ga kawa e itta

In this analysis, what is right-node-raised in (1) is the tensed verb *itta*.

The second possibility we consider is that the sentence may involve coordination of two tenseless clauses, as shown in (8).

- (8) [[Hanako ga yama e ik-] [Masao ga kawa e ik-]] ta
 → Hanako ga yama e, Masao ga kawa e ik- ta

In this analysis, what is right-node-raised is the verb stem *ik-*. Since the verb stem is a bound morpheme, the pre-RNR structure that is posited in this analysis is not something that can be used as a surface form. The structure becomes a pronounceable sentence only after the verb stem is right-node-raised and the verb stem and the sentence-final tense morpheme *-ta* are combined to yield a phonological word *itta*.

And the third possible analysis we will consider is one in which sentence (1) is derived by applying right-node raising to the sentence in (9), in which the first clause ends with *iki*, the so-called infinitive form of the verb *ik-* ‘to go’.

- (9) Hanako ga yama e iki, Masao ga kawa e itta.
 Hanako NOM mountain to go-INF Masao NOM river to go-PAST
 ‘Hanako went to the mountain and Masao went to the river.’

A clause ending with the infinitive form of a verb is often interpreted as being semantically conjoined with the immediately following clause, while it is not clear whether the first clause in such a structure is syntactically a conjunct or an adjunct. In this analysis, in which (9) is taken to be the pre-RNR form of (1), what is right-node-raised out of the first clause must be the infinitive form *iki*, and what is right-node-raised out of the second clause must be either the verb stem *ik-* or the tensed verb *itta*. We view this third analysis as something conceivable because it has been shown by Shiraishi & Abeillé (2016) that there is a type of right-node raising in which slightly different forms of a verb are right-node-raised as if they were identical to each other.

It will be our contention in this paper that there is evidence that the first and the second analysis are both allowed in the grammar of Japanese whereas there is no evidence that the third analysis is allowed in the grammar. More specifically, we will argue that the sentence in (1) is structurally ambiguous between the first analysis and the second analysis, and that there are sentences that are amenable only to the first type of analysis as well as sentences that are amenable only to the second type of analysis.

The findings reported in this paper have implications regarding the basic clause structure of Japanese. There have historically been two schools of thought concerning the syntactic status of the tense morphemes in Japanese. On the one hand, there are authors who argue that a verb stem and the tense morpheme immediately following it always form not just a phonological constituent but a morphosyntactic constituent as well (see Sells (1995) among others). This line of thinking is often referred to as the lexicalist hypothesis in the literature. On the other hand, there are authors who argue that a verb stem and the tense morpheme immediately following it do not necessarily form a morphosyntactic constituent (see Tokieda (1950) and Fukui & Sakai (2003) among others). This view is sometimes referred to as the non-lexicalist view in the literature. The theory that we will advance in this paper, according to which the structure shown in (8) above is possible, entails that, at least in some cases, the tense morphemes in Japanese are syntactically independent and take tenseless clauses as complements. Thus, if the view that we are going to advocate is correct, the lexicalist hypothesis needs to be abandoned.

Before proceeding, we wish to clarify exactly what it means to reject the lexicalist hypothesis in the present context. It is an indisputable fact that a string made up of a tense morpheme and a verb stem immediately preceding it always form a phonological constituent (more specifically, a phonological word) in Japanese. At the same time, there is no easily available evidence that a string of that form is not a morphosyntactic constituent. Our claim in the present paper is that a string that is indisputably a phonological constituent can nevertheless be analyzed by the language learner as a morphosyntactic non-constituent, even when there is no easily available evidence for such an analysis.

2 RNR of mismatched verb forms?

We begin by examining the third type of analysis mentioned above. This analysis appears viable for RNR constructions like (1), which involve conjunction. The analysis, however, encounters a problem when it is applied to examples involving disjunction, such as (10).

- (10) Hanako ga yama e, mata wa Masao ga kawa e itta.
 Hanako NOM mountain to or Masao NOM river to go-PAST
 ‘Hanako went to the mountain, or Masao went to the river.’

The pre-RNR structure posited for sentence (10) in this analysis is shown in (11). The problem is that sentence (11) is considerably unnatural as a sentence expressing simple disjunction of two propositions.

- (11) [Hanako ga yama e iki], mata wa [Masao ga kawa e
 [Hanako NOM mountain to go-INF] or [Masao NOM river to
 itta].
 go-PAST]

The sentence in (11) is acceptable as a sentence expressing something along the lines of “Hanako habitually went to the mountain and Masao habitually went to the river, and on any given day, one of the two types of events (namely either Hanako going to the mountain or Masao going to the river) took place,” but it does not express simple disjunction, which *can* be expressed by (10).

Our assertion that a sentence like (11) cannot express simple disjunction devoid of the implication of habituality is justified by the result of a questionnaire study we conducted using (12) as one of the experimental sentences.

- (12) ??[Seifu-gun ga byôin o kûbaku shi], mata wa
 [government forces NOM hospital ACC air strike do-INF] or
 [hanran-gun ga byôin no sugu chikaku no buki-ko o
 [rebel forces NOM hospital GEN immediate vicinity GEN arsenal ACC
 bakuha shita] rashii.
 explode do-PAST] it appears
 ‘It appears that either the government forces did an air strike on the hospital or the rebel forces exploded the arsenal in the immediate vicinity of the hospital.’
 <1, 4, 6, 4>

The respondents of the questionnaires mentioned in the present paper were all students at the University of Tokyo, and received 500 yen as a compensation for their time. The respondents were asked to judge the acceptability of given sentences on the scale of 1 to 4 described in Table 1. The order of sentences was randomized for each respondent. Each sentence was accompanied by a description of what the

Table 1: The 4-point scale used in the questionnaires

rating	meaning of the rating
1	‘The sentence is perfectly natural under the intended reading.’
2	‘The sentence is slightly unnatural under the intended reading.’
3	‘The sentence is considerably unnatural under the intended reading.’
4	‘The sentence is completely impossible under the intended reading.’

intended reading of that sentence was. The four figures shown after sentence (12) and some other sentences below indicate the number of respondents who chose 1, 2, 3, and 4 respectively for those sentences. A sentence for which the mean acceptability rating was R is shown throughout this paper with no symbol if $1 \leq R < 2$, with ‘?’ if $2 \leq R < 2.5$, with ‘??’ if $2.5 \leq R < 3$, with ‘?*’ if $3 \leq R < 3.5$, and with ‘*’ if $3.5 \leq R \leq 4$.

The questionnaire results reported in this paper come from six different questionnaire studies. The questionnaire for sentence (12) included three experimental sentences and 12 filler sentences, and involved 15 respondents. (The other experimental sentences contained in this questionnaire were structurally and lexically similar to sentence (12) but did not involve right-node raising.) The questionnaire for sentences (18), (23), (24), (25), (26), and (27) included six experimental sentences and nine filler sentences, and involved 10 respondents. The questionnaire for sentences (28), (29), (33), (34), (35), and (37) included six experimental sentences and nine filler sentences, and involved 15 respondents. The questionnaire for sentences (36) and (38) included two experimental sentences and 14 filler sentences, and involved 28 respondents. The questionnaire for sentence (41) included three experimental sentences and 20 filler sentences, and involved 15 respondents. (The other experimental sentences in this questionnaire were structurally and lexically similar to (41) but did not involve coordination.) And the questionnaire for sentences (42) and (43) included three experimental sentences and 20 filler sentences, and involved 11 respondents. (The remaining experimental sentence in this questionnaire was structurally and lexically similar to (42) and (43), but contained only one accusative noun phrase.) What we call filler sentences here are sentences that are irrelevant to the present paper. Some of those sentences were in fact not literally fillers but were included in the questionnaire for some specific purposes.

The questionnaire result for sentence (12) indicates that the sentence, which has the same structure as (11) but pragmatically disfavors habitual interpretation unlike (11), is considerably unnatural. If we assume (i) that sentence (10) can be derived from sentence (11) through application of a particular type of RNR and (ii) that the type of RNR invoked in generating (10) is meaning-preserving, we predict incorrectly that sentences like (11) and (12) must be able to express simple disjunction, since (10) is capable of expressing simple disjunction. Thus, if assumption (ii) above can be shown to be correct, then we will be able to conclude that assumption (i) must be incorrect. The question, of course, is whether assumption (ii)

can be shown to be correct.

Right-node raising can be meaning-changing under certain circumstances, but there is a reason to believe that the type of right-node raising that is invoked in generating (10) must be meaning-preserving. As noted in Yatabe (2012) and Valmala (2013), when right-node raising is meaning-changing, there has to be a prosodic boundary immediately preceding the right-node-raised expression, so that the right-node-raised expression is pronounced as an independent prosodic constituent (or a sequence of independent prosodic constituents) detached from the phrase (typically a coordinate structure) out of which it has been right-node-raised. In the case at hand, namely sentence (10), the right-node-raised expression is either the verbal expression *itta* ‘go-PAST’ as a whole or the verb stem that is at the left edge of that expression. There is no prosodic boundary immediately preceding the verb stem, and the verbal expression *itta* is pronounced as a normal part of the prosodic constituent that comprises the immediately preceding expression *kawa e* ‘river to’ and the verbal expression. This suggests that, even if the sentence in (10) had been derived from (11) by right-node-raising the verbal expression *itta* or a part of it, the right-node raising involved could not have changed the meaning of the sentence.

We therefore conclude that a sentence like (10) is not derived from a structure like (11).

From a logical point of view, it is possible that a sentence like (1), involving conjunction, can be derived from (9), even if a sentence like (10), involving disjunction, is not derived from (11). We believe, however, that that is a remote possibility. For one thing, it seems crosslinguistically common for there to be parallelism between structures involving conjunction and structures involving disjunction. For another, whatever mechanism derives sentence (10) will derive sentence (1) from a source distinct from (9), thus obviating the need to have a mechanism that derives (1) from (9). Therefore Occam’s razor justifies a certain amount of prejudice against the view that (1) can be derived from (9).

3 RNR out of tensed clauses

Next, we will consider whether there are sentences that must be analyzed as involving RNR of a tensed verb out of coordinated tensed clauses, as depicted in (7). It turns out that there clearly are such sentences. (13) is one such sentence.

- (13) Hanako wa osoraku yama, Masao wa osoraku kawa e,
 Hanako TOP probably mountain Masao TOP probably river to
 (sorezore) itta.
 (individually) go-PAST
 ‘Hanako probably went to the mountain and Masao probably went to the
 river (and the two of them were acting individually).’

Since topic phrases like *Hanako wa* and *Masao wa* cannot appear inside a tenseless

phrase (see Takubo (1987)), this sentence can only be analyzed as involving RNR of the tensed verb *itta* ‘go-PAST’ out of two coordinated tensed clauses.

There are two potential problems with this account that need to be addressed. The first potential problem concerns the grammatical status of the postulated pre-RNR structure. In the account we are advocating here, sentence (13) is derived from a structure like (14).

- (14) Hanako wa osoraku yama e itta, Masao wa osoraku kawa e
 Hanako TOP probably mountain to go-PAST Masao TOP probably river to
 itta.
 go-PAST
 ‘Hanako probably went to the mountain, Masao probably went to the river.’

The problem is that it is not intuitively obvious that this string is allowed as a possible sentence in Japanese; example (14) is an acceptable string in Japanese, but it is conceivable that it is licensed only as a sequence of two independent sentences, rather than as a single grammatical sentence. Our account cannot be correct if a string like (14) is not allowed to be a single grammatical sentence.

This potential problem turns out not to be a real problem for our account, since an example like the following indicates that a juxtaposition of two sentences like (14) can indeed be licensed as a single syntactic constituent in the language.

- (15) Kare wa kekkyoku iwanakatta, [kare-jishin ga iku, kare-jishin
 he TOP ultimately say-NEG-PAST [he himself NOM go-PRES he himself
 ga tatakau to].
 NOM fight-PRES COMP]
 ‘He ultimately did not say that he would go himself and he would fight himself.’

The string *kare-jishin ga iku, kare-jishin ga tatakau* ‘he would go himself and he would fight himself’ in this sentence can only be analyzed as a syntactic constituent consisting of two juxtaposed clauses.

It might seem possible to view sentence (15) as having been derived from (16) by right-node-raising the sentence-final complementizer *to*.

- (16) Kare wa iwanakatta, [kare-jishin ga iku to, kare-jishin
 he TOP say-NEG-PAST [he himself NOM go-PRES COMP he himself
 ga tatakau to].
 NOM fight-PRES COMP]
 ‘He ultimately did not say that he would go himself, that he would fight himself.’

If that is a possible analysis of sentence (15), then the sentence will no longer provide evidence that two juxtaposed tensed clauses can form a syntactic constituent. It is, however, arguably impossible to analyze (15) as a result of such application

of RNR, because (15) is not synonymous with (16). Sentence (15) can mean that the man referred to did not say “I will go myself, and I will fight myself.” On this reading, the sentence can be true even if the man expressed the content of one of the two embedded clauses, as long as he did not express the content of the other embedded clause. On the other hand, sentence (16) cannot express that meaning; it can only mean that the man did not express the content of either of the two embedded clauses.

As we noted in the previous section as well, right-node raising can be meaning-changing under certain circumstances, but the difference in meaning between (15) and (16) cannot be ascribed to right-node raising, if we are correct in assuming that the meaning-changing kind of right-node raising always creates a prosodic boundary immediately before the right-node-raised expression; the sentence-final complementizer *to* in sentence (15), which is the right-node-raised expression in the hypothetical scenario under discussion, does not have to be preceded by an intonational break, and can be pronounced as part of a phonological word that comprises the immediately preceding verbal expression *tatakau* ‘fight-PRES’ and the complementizer. Thus, sentence (15) must be generated without application of RNR at least when the sentence-final complementizer is not immediately preceded by an intonational break, and we can therefore conclude that a juxtaposition of two tensed clauses is allowed to form a syntactic constituent.

The second potential problem with the proposed account of sentence (13) is that the postulated source for it, namely (14), cannot be used in all contexts in which (13) can be used. A case in point is the contrast between (17) and (18).

- (17) Daijōbu sa, Hanako wa osoraku yama, Masao wa osoraku
 OK I assure you Hanako TOP probably mountain Masao TOP probably
 kawa e itta kara.
 river to go-PAST because
 ‘It’s going to be OK, I assure you, because Hanako probably went to the mountain and Masao probably went to the river.’
- (18) ?Daijōbu sa, Hanako wa osoraku yama e itta, Masao
 OK I assure you Hanako TOP probably mountain to go-PAST Masao
 wa osoraku kawa e itta kara.
 TOP probably river to go-PAST because
 ‘It’s going to be OK, I assure you, because Hanako probably went to the mountain and Masao probably went to the river.’
 <2, 5, 0, 3>

In the proposed account, (17) is derived from (18) by right-node-raising the string *e itta* ‘to go-PAST’ out of the two embedded clauses. Thus, the fact that (18) is slightly awkward unlike (17) appears problematic.

In our view, this is also not a real problem for the proposed account. The reason sentence (18) is awkward most probably has to do with the fact that the sentence-final morpheme *kara* ‘because’ is an enclitic, i.e. an expression that needs

to be phonologically dependent on an expression that immediately precedes it. This view receives support from the fact that the syntactic structure exemplified by (18) is perfectly acceptable when the sentence-final morpheme is clearly not an enclitic, as in (15) above. The complementizer *to*, which immediately follows the juxtaposed tensed clauses in (15), can be pronounced as an independent phonological word, separated from the preceding expressions by an intonational break, as in (19), where the use of a comma before *to* is meant to indicate the presence of an intonational break there.

- (19) Kare wa kekkyoku iwanakatta, [kare-jishin ga iku, kare-jishin
he TOP ultimately say-NEG-PAST [he himself NOM go-PRES he himself
ga tatakau, to].
NOM fight-PRES COMP]
'(Same as (15))'

In contrast, the postposition *kara*, which immediately follows the juxtaposed tensed clauses in (18), cannot be pronounced as an independent phonological word; there cannot be an intonational break immediately before that postposition. These observations justify our hypothesis that *kara* is an enclitic whereas *to* is not. Thus, we can capture both the awkwardness of (18) and the well-formedness of (15) by postulating a constraint like (20).

- (20) An enclitic like *kara* must not immediately follow a coordinate structure, when it is not possible for the enclitic to become phonologically dependent on a host that is part of each of the conjuncts (such as an expression that has been right-node-raised out of each of the conjuncts).

This constraint is consistent with the overall theory that we are arguing for in this paper. Moreover, the postulated constraint would not be an unreasonable one; it is arguably a mirror image of the constraint that blocks expressions like (21) and (22) in French (see Bonami & Tseng (2010) for a recent discussion of phenomena of this type).

- (21) *de le père et la mère
of the father and the mother
(22) *du père et la mère
of the father and the mother

Suppose that the preposition *de* is a proclitic (i.e. an expression that needs to be phonologically dependent on an expression that immediately follows it) when its complement is either a non-coordinate structure that starts with the determiner *le* or a coordinate structure one of whose conjuncts starts with *le*. Suppose also that French has a constraint that prohibits a proclitic like *de* from preceding a coordinate structure when it is not possible for the proclitic to become phonologically dependent on a host that is part of each of the conjuncts. Then (21) and (22) will

both be correctly ruled out because in (21) *de* is not phonologically dependent on any host and in (22) *de* is phonologically dependent on a host that is part of the first conjunct alone.

At first blush, the analysis that we have proposed seems to be contradicted by the following observation: a sentence like (18) improves when the word *soshite* ‘and’ is added between the two juxtaposed embedded clauses, as in (23).

- (23) Daijôbu sa, [Hanako wa osoraku yama e itta, soshite
 OK I assure you [Hanako TOP probably mountain to go-PAST and
 Masao wa osoraku kawa e itta kara].
 Masao TOP probably river to go-PAST because]
 ‘It’s going to be OK, I assure you, because Hanako probably went to the
 mountain and Masao probably went to the river.’
 <6, 4, 0, 0>

If addition of the word *soshite* does not alter the syntactic structure involved, sentence (23) is expected to be as awkward as sentence (18), but that expectation is not fulfilled. The Wilcoxon signed-rank test showed that (23) was rated as significantly more acceptable than (18) ($Z = 2.21$, $p = 0.03$).

We submit that addition of *soshite* in this case does alter the syntactic structure involved. More specifically, we hypothesize that what looks like two juxtaposed clauses in a sentence like (23) is in fact not a coordinate structure but a non-coordinate headed structure such that what looks like the second conjunct in it (that is, the clause that starts with the word *soshite*) is its sole head and what looks like the first conjunct in it is an adjunct. If this hypothesis is correct, the enclitic *kara* in (23) can become phonologically dependent on the immediately preceding verbal expression *itta* without violating the constraint in (20).

One piece of evidence for this hypothesis comes from observations like the following.

- (24) ?Kimi wa, [[sono biru ni kaminari ga ochita], soshite
 you TOP [[that building LOC lightning NOM fall-PAST] and
 [kekka-teki ni nani ga okita] kara] komatta no?
 [as a result what NOM happen-PAST] because] be troubled-PAST NML
 ‘What is the thing *x* such that you got into trouble because a lightning hit
 that building and *x* happened as a result?’
 <3, 3, 3, 1>
- (25)?*Kimi wa, [[sono biru ni nani ga ochita], soshite [kekka-teki ni
 you TOP [[that building LOC what NOM fall-PAST] and [as a result
 kaji ga okita] kara] komatta no?
 fire NOM happen-PAST] because] be troubled-PAST NML
 ‘What is the thing *x* such that you got into trouble because *x* hit that build-
 ing and a fire broke out as a result?’
 <1, 1, 3, 5>

In both these sentences, the word *kara* ‘because’ takes as the complement a sequence of juxtaposed clauses joined by *soshite*, and one of the clauses contains the *wh* expression *nani* ‘what’. The *wh* word is contained in the second of the juxtaposed clauses in (24), and it is contained in the first of the juxtaposed clauses in (25). If the juxtaposed clauses constitute a normal coordinate structure, these sentences are expected to have the same level of acceptability, but the Wilcoxon signed-rank test revealed that sentence (24) was rated as significantly more acceptable than sentence (25) ($Z = 2.73, p < 0.01$). We take this to be a piece of evidence that the juxtaposed clauses in sentences like (23), (24), and (25) do not constitute coordinate structures.

The complementizer *to*, which is not an enclitic, contrasts with *kara* in this regard as well, as shown by the following examples.

- (26) ?Kimi wa, [[sono biru ni kaminari ga ochita], soshite
you TOP [[that building LOC lightning NOM fall-PAST] and
[kekka-teki ni nani ga okita], to] omotteru no?
[as a result what NOM happen-PAST] that] think NML
‘What is the thing *x* such that you think that a lightning hit that building
and *x* happened as a result?’
<1, 6, 3, 0>
- (27) ?Kimi wa, [[sono biru ni nani ga ochita], soshite [kekka-teki ni
you TOP [[that building LOC what NOM fall-PAST] and [as a result
kaji ga okita], to] omotteru no?
fire NOM happen-PAST] that] think NML
‘What is the thing *x* such that you think that *x* hit that building and a fire
broke out as a result?’
<0, 7, 3, 0>

In both these sentences, *to* takes as the complement a sequence of two juxtaposed clauses, with the word *soshite* in between. In (26), the second of those juxtaposed clauses contains a *wh* word *nani*, and in (27), the first of the juxtaposed clauses contains that word. There is no discernible difference in acceptability between the two examples. This observation makes sense if we assume that a sequence of juxtaposed clauses with *soshite* in between is structurally ambiguous and can be analyzed not only as a non-coordinate headed structure but also as a normal coordinate structure. When such a sequence of clauses is followed by *kara*, it has to be analyzed as a non-coordinate structure because of the constraint stated in (20). On the other hand, when such a sequence is followed by *to*, it can be analyzed as a normal coordinate structure, with the result that a *wh* word is allowed to occur in any of the juxtaposed clauses, albeit somewhat marginally.

Thus, there does not appear to be any fundamental problem with the hypothesis that a tensed verb can be right-node-raised out of juxtaposed tensed clauses in Japanese.

4 RNR out of tenseless phrases

In this section, it will be argued that there are sentences that are amenable only to the analysis depicted in (8), which is incompatible with the lexicalist hypothesis. Our argument here is based on sentence (28).

- (28) Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
 every morning regularly [about 15 minutes jogging ACC do-PRES or]
 chôshoku mae ni udetatefuse, chôshoku go ni sukuwatto o shita.
 before breakfast pushup after breakfast squat ACC do-PAST
 ‘Every morning, I regularly either jogged for about 15 minutes or did
 pushups before breakfast and squats after breakfast.’
 <6, 6, 0, 3>

In the latter half of this sentence, the string *chôshoku mae ni udetatefuse* ‘pushups before breakfast’ and the string *chôshoku go ni sukuwatto* ‘squats after breakfast’ are juxtaposed with each other. Since neither of the juxtaposed strings consists of co-arguments of the same predicate, this portion of the sentence cannot be regarded as an instance of argument-cluster coordination. Thus, if we are to adhere to the lexicalist hypothesis, it has to be assumed that this sentence is derived from sentence (29) by right-node-raising the accusative case marker *o* and the tensed verb *shita*.

- (29)?*Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
 every morning regularly [about 15 minutes jogging ACC do-PRES or]
 chôshoku mae ni [udetatefuse o] shita, chôshoku go ni [sukuwatto
 before breakfast [pushup ACC] do-PAST after breakfast [squat
 o] shita.
 ACC] do-PAST
 ‘(Same as (28))’
 <0, 3, 5, 7>

This assumption, however, is problematic. As shown by the questionnaire result, sentence (29) is considerably unnatural under the intended interpretation. The Wilcoxon signed-rank test showed that (29) was rated as significantly less acceptable than (28) ($Z = 2.94$, $p < 0.01$). The only meaning that sentence (29) can express appears to be something along the lines of “Every morning, I regularly either jogged for about 15 minutes or did pushups before breakfast, and I did squats after breakfast.” In other words, whereas the structure of the verb phrase in (28) is (30), the structure of the verb phrase in (29) seems to be (31).

- (30) [VP1 [VP2 VP3]]
 (31) [[VP1 VP2] VP3]

Thus, (29) cannot be the pre-RNR structure of (28) unless it is assumed that RNR can induce restructuring of the kind that can transform (31) into (30). Such an assumption appears to us to be far-fetched in the first place, and it is made all the more implausible by the fact that there is no intonational break immediately before the right-node-raised expression in (28), a fact that suggests that the right-node raising involved in generating (28) is of the meaning-preserving type.

In contrast, such a problematic assumption is not forced on us if the analysis depicted in (8) is applied to (28). On such an account, sentence (28) can be generated as follows.

- (32) [Mai-asa chanto
 [[jûgo-fun gurai jogingu o suru ka]
 [[chôshoku mae ni udetatefuse o s-]
 [chôshoku go ni sukuwatto o s-]]]-ta
 ↓
 [Mai-asa chanto
 [[jûgo-fun gurai jogingu o suru ka]
 [[chôshoku mae ni udetatefuse]
 [chôshoku go ni sukuwatto o s-]]]-ta]

The bound morpheme *s-* is the verb stem of the verb *suru* ‘to do’, and *-ta* is the past tense morpheme. In this analysis, the complement of the past tense morpheme has a structure like (30), where VP1, which ends with *ka* ‘or’, is headed by the present tense form of a verb (*suru*), whereas VP2 and VP3 are both headed by a verb stem (*s-*). What is right-node-raised is the sequence made up of the accusative case marker *o* and the verb stem *s-*. After the application of RNR, the verb stem and the tense morpheme are combined to become the phonological word *shita*.

The following three examples are variants of sentences (28) and (29), and exhibit the same pattern of acceptability as those sentences.

- (33) Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
 every morning regularly [about 15 minutes jogging ACC do-PRES or]
 chôshoku mae ni udetatefuse, chôshoku go ni sukuwatto o shita?
 before breakfast pushup after breakfast squat ACC do-PAST
 ‘Did you regularly either jog for about 15 minutes or do pushups before
 breakfast and squats after breakfast, every morning?’
 <11, 2, 2, 0>
- (34) *Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
 every morning regularly [about 15 minutes jogging ACC do-PRES or]
 chôshoku mae ni udetatefuse o shita? chôshoku go ni sukuwatto o
 before breakfast pushup ACC do-PAST after breakfast squat ACC
 shita?
 do-PAST

‘(Same as (33))’

<1, 1, 4, 9>

- (35) *Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
every morning regularly [about 15 minutes jogging ACC do-PRES or]
chôshoku mae ni udetatefuse o shita, chôshoku go ni sukuwatto o
before breakfast pushup ACC do-PAST after breakfast squat ACC
shita?

do-PAST

‘(Same as (33))’

<0, 1, 7, 7>

Sentence (33) is an interrogative variant of sentence (28), and is as acceptable as the latter. Sentences (34) and (35) are interrogative variants of (29), and are both even less acceptable than the original, non-interrogative sentence. The only difference between (34) and (35) is that the former contains two question marks whereas the latter contains only one question mark.

The following example shows that sentence (29) does not become acceptable even if the word *soshite* is added between the two juxtaposed tensed clauses.

- (36) *Mai-asa chan to, [jûgo-fun gurai jogingu o suru ka],
every morning regularly [about 15 minutes jogging ACC do-PRES or]
chôshoku mae ni udetatefuse o shita, soshite chôshoku go ni
before breakfast pushup ACC do-PAST and after breakfast
sukuwatto o shita.

squat ACC do-PAST

‘(Same as (28))’

<0, 0, 9, 19>

This observation is consistent with what our hypothesis leads us to expect.

The following example, which is modelled after an example discussed in Kuroda (2003), shows that the process that we have claimed takes place inside the complement of a tense morpheme can take place inside the complement of the causative morpheme (*s*)*ase*. This observation adds to the plausibility of the proposed account.

- (37) Hanako wa Masao ni, [sôji o shite fuyôhin o
Hanako TOP Masao DAT [cleaning ACC do-GER unnecessary items ACC
subete shobun suru ka], heya-dai o kyô jû, chûshajô-dai
all get rid of-PRES or] rent ACC within today parking space fee
o kongetsu chû ni zengaku shiharawaseru koto ni
ACC withing this month DAT the entire amount pay-CAUS-PRES NML DAT
shita.
do-PAST

‘Hanako decided to make Masao do one of two things, where option 1 was to clean the place and get rid of all the unnecessary items, and option 2 was to pay up the rent before the end of the day and the parking space fee before the end of the month.’

<10, 4, 1, 0>

In this sentence, the causative morpheme *ase* (which is embedded in the phonological word *shiharawaseru*) takes a complement whose pre-RNR structure has the form shown in (30), where VP1 is a verb phrase followed by *ka* ‘or’ (i.e. the bracketed expression in (37)), and VP2 and VP3 are both tenseless verb phrases ending in the verb stem *shiharaw-*. The verb stem (together with the dative case marker *ni* and the noun *zengaku* ‘the entire amount’) is right-node-raised out of VP2 and VP3, and fuses with the causative morpheme and the tense morpheme to become the phonological word *shiharawaseru*. There is arguably no other way to analyze the structure of (37).

Note that sentence (37) itself poses the same problem for the lexicalist hypothesis that sentence (28) does. In order to analyze the sentence in accordance with the lexicalist hypothesis, it is necessary to derive it from (38) by right-node-raising the string *ni zengaku shiharawaseru*, but sentence (38) cannot express the same meaning as (37).

(38) ??Hanako wa Masao ni, [sôji o shite fuyôhin
 Hanako TOP Masao DAT [cleaning ACC do-GER unnecessary items
 o subete shobun suru ka], heya-dai o kyô jû
 ACC all get rid of-PRES or] rent ACC within today
 ni zengaku shiharawaseru, chûshajô-dai o
 DAT the entire amount pay-CAUS-PRES parking space fee ACC
 kongetsu chû ni zengaku shiharawaseru koto ni
 withing this month DAT the entire amount pay-CAUS-PRES NML DAT
 shita.
 do-PAST
 ‘(Same as (37))’
 <2, 7, 11, 8>

To summarize the discussion so far, we have two arguments for the non-lexicalist analysis of (28). Unlike the lexicalist analysis, it does not require us to assume that RNR (more specifically, the type of RNR that does not induce a prosodic boundary immediately before the right-node-raised expression) can induce restructuring of the kind that transforms (31) into (30). Moreover, there is an independent reason to believe that the syntactic structure that it postulates is allowed by the grammar.

In the remainder of this section, we wish to address one apparent problem with our analysis. The account that we have presented above relies on the hypothesis that a structure of the following form can be analyzed as a coordinate structure.

- (39) [[... V PRES] ka [... V]]
 (where V is a verb stem and PRES is a present tense morpheme.)

It is not immediately obvious whether this hypothesis, which we owe to Kuroda (2003), is indeed correct or not.

The first thing to be noted about the structure depicted in (39) is that the first disjunct and the second disjunct belong to different syntactic categories, the former being tensed and the latter being not tensed. This might appear problematic, but it is not; all it means is that this structure is an instance of coordination of unlikes. In general, the conjuncts (including disjuncts) of a coordinate structure do not necessarily have to belong to the same syntactic category, as demonstrated in Bayer (1996) and the literature cited there. We can assume, without a problem, that the grammar of Japanese contains a phrase structure schema that licenses a coordinate structure consisting of one or more *ka*-marked phrases headed by the present-tense morpheme, followed by a VP headed by a tenseless verb stem. One way to deal with coordination of unlikes in general within the HPSG framework is presented in Yatabe (2004).

In our view, the hypothesis that an expression of the form shown in (39) can be a coordinate structure in Japanese is not only unproblematic but empirically justified by the following two considerations.

First, since part of a conjunct cannot be preposed out of the coordinate structure in Japanese (see Yatabe (2003)), the hypothesis in question leads us to expect that a part of the second VP in a structure like (39) cannot be preposed out of the expression that is assumed here to be a coordinate structure, and this expectation is fulfilled, as shown by the contrast between (40a) and (40b).

- (40) a. Mai-asa [jogingu o suru ka] [hon o] yonda.
 every morning [jogging ACC do-PRES or] [book ACC] read-PAST
 ‘Every morning, I either jogged or read a book.’
 b. *Mai-asa [hon o] [jogingu o suru ka] yonda.
 every morning [book ACC] [jogging ACC do-PRES or] read-PAST
 ‘(Same as (40a))’

According to the hypothesis we are pursuing, sentence (40a) contains a coordinate structure in which a VP of the form *jogingu o suru* ‘jogging ACC do-PRES’ and another VP of the form *hon o yom-* ‘book ACC read’ are coordinated by the word *ka* ‘or’. The noun phrase *hon o* is part of the second conjunct in this structure, and hence is expected to be impossible to prepose out of the coordinate structure. This expectation is fulfilled by the unacceptability of sentence (40b). In contrast, if what we took to be the first conjunct of a coordinate structure is instead assumed to be, say, some kind of adjunct, then the unacceptability of (40b) will likely remain mysterious.

Second, we can use the so-called double-*o* constraint to show that the two or more expressions that are joined by the word *ka* in a structure like (39) have

identical syntactic status, as is predicted by the hypothesis that those two or more expressions are conjuncts of a coordinate structure. The double-*o* constraint is a grammatical rule of Japanese whose effect is illustrated by the following example.

- (41)?*[Hizashi ga tsuyoku natte kita node], Satô sensei
[sunlight NOM strong become-GER come-PAST because] teacher Sato
wa [kodomo-tachi o], [seibô o kaburu ka] [hiyake-dome
TOP [children ACC] [school hat ACC put on-PRES or] [sunscreen
o] nuraseru koto ni shita.
ACC] apply-CAUS-PRES NML DAT do-PAST
‘Because the sunlight became strong, Sato, the teacher, decided to make
the children either put on the school hat or apply sunscreen to themselves.’
<1, 3, 5, 6>

In this example, the causee (*kodomotachi* ‘children’) is marked by the accusative case marker *o*, and the complement of the causative morpheme consists of two VPs which both consist of a transitive verb and a grammatical object marked by *o*. This sentence is ruled out by the double-*o* constraint, which prohibits the causee from being accusative when the complement of the causative morpheme is headed by a transitive verb.

Now, compare this sentence with the following two sentences.

- (42) [Hizashi ga tsuyoku natte kita node], Satô sensei
[sunlight NOM strong become-GER come-PAST because] teacher Sato
wa [kodomo-tachi o], [seibô o kaburu ka] [kyôshitsu e]
TOP [children ACC] [school hat ACC put on-PRES or] [classroom to]
modoraseru koto ni shita.
return-CAUS-PRES NML DAT do-PAST
‘Because the sunlight became strong, Sato, the teacher, decided to make
the children either put on the school hat or return to the classroom.’
<3, 6, 2, 0>
- (43) [Hizashi ga tsuyoku natte kita node], Satô sensei
[sunlight NOM strong become-GER come-PAST because] teacher Sato
wa [kodomo-tachi o], [kyôshitsu e modoru ka] [seibô o]
TOP [children ACC] [classroom to return-PRES or] [school hat ACC]
kaburaseru koto ni shita.
put on-CAUS-PRES NML DAT do-PAST
‘Because the sunlight became strong, Sato, the teacher, decided to make
the children either return to the classroom or put on the school hat.’
<4, 5, 2, 0>

In both these sentences, the causee is marked by the accusative case marker, and the complement of the causative morpheme contains two VPs, as in (41) above.

In sentence (42), the first VP in the complement of the causative morpheme is headed by a transitive verb, and the second VP is headed by an intransitive verb. In sentence (43), on the other hand, the first VP is headed by an intransitive verb, and the second VP is headed by a transitive verb. The questionnaire result indicates that these two sentences are equally acceptable.

This arguably means that the two VPs contained in the complement of the causative morpheme in these sentences have identical syntactic status, in conformity with the hypothesis that the structure shown in (39) constitutes a coordinate structure. Given that hypothesis, the status of sentences (41), (42), and (43) can be captured by a constraint like the following.

- (44) The causee argument of a causative morpheme can be accusative only if the complement of that causative morpheme is either a single VP headed by an intransitive verb or a coordinate structure such that one of the conjuncts is headed by an intransitive verb.

In contrast, if the structure shown in (39) were, say, some kind of head-adjunct structure, sentence (42) and sentence (43) would have to differ from each other in acceptability, contrary to what we have seen.

5 Summary

We have examined the fine structure of clausal right-node raising constructions in Japanese, and argued that there are sentences in which a tensed verb is right-node-raised out of coordinated tensed clauses as well as sentences in which a verb stem is right-node-raised out of coordinated tenseless phrases. In the latter case, the tense morpheme has to be assumed to take a tenseless complement clause, and we have noted that the existence of such a structure contradicts the so-called lexicalist hypothesis, according to which a verb stem and the tense morpheme immediately following it always form a morphosyntactic constituent.

References

- Bayer, Samuel. 1996. The coordination of unlike categories. *Language* 72. 579–616.
- Bonami, Olivier & Jesse Tseng. 2010. French portmanteaus as simplex elements. Handout for Jahrestagung der DGfS. Obtained at <http://www.llf.cnrs.fr/sites/llf.cnrs.fr/files/biblio/bonz-dgfs-ho4.pdf> on September 9, 2018.
- Fukui, Naoki & Hiromu Sakai. 2003. The visibility guideline for functional categories: verb raising in Japanese and related issues. *Lingua* 113. 321–375.

- Kuroda, S.-Y. 2003. Complex predicates and predicate raising. *Lingua* 113. 447–480.
- Mouret, François. 2006. A phrase structure approach to argument cluster coordination. In Stefan Müller (ed.), *Proceedings of the 13th international conference on Head-Driven Phrase Structure Grammar*, 247–267. Stanford: CSLI.
- Sells, Peter. 1995. Korean and Japanese morphology from a lexical perspective. *Linguistic Inquiry* 26. 277–325.
- Shiraishi, Aoi & Anne Abeillé. 2016. Peripheral ellipsis and verb mismatch. In Doug Arnold, Miriam Butt, Berthold Crysmann, Tracy Holloway King & Stefan Müller (eds.), *Proceedings of the joint 2016 conference on Head-driven Phrase Structure Grammar and Lexical Functional Grammar*, 662–680. Stanford: CSLI.
- Takubo, Yukinori. 1987. Tôgo-kôzô to bunmyaku-jôhō (Syntactic structure and contextual information). *Nihongogaku* 6(5). 37–48.
- Tokieda, Motoki. 1950. *Nihon Bunpô: Kôgo-hen (Japanese Grammar: Part on the Vernacular)*. Tokyo: Iwanami Shoten.
- Valmala, Vidal. 2013. On right node raising in Catalan and Spanish. *Catalan Journal of Linguistics* 12. 219–251.
- Yatabe, Shûichi. 2003. Does scrambling in Japanese obey the Coordinate Structure Constraint? In *Nihon gengogakkai dai-126-kai taikai yokô-shû (Proceedings of the 126th meeting of the Linguistic Society of Japan)*, 262–267. Kyoto: Linguistic Society of Japan.
- Yatabe, Shûichi. 2004. A comprehensive theory of coordination of unlikes. In Stefan Müller (ed.), *Proceedings of the 11th international conference on Head-Driven Phrase Structure Grammar*, 335–355. Stanford: CSLI.
- Yatabe, Shûichi. 2012. Comparison of the ellipsis-based theory of non-constituent coordination with its alternatives. In Stefan Müller (ed.), *Proceedings of the 19th international conference on Head-Driven Phrase Structure Grammar*, 454–474. Stanford: CSLI.
- Yatabe, Shûichi & Wai Lok Tam. 2017. In defense of an HPSG-based theory of non-constituent coordination: A reply to Kubota and Levine. Ms. Posted at <http://ling.auf.net/lingbuzz/003152>.