

MACSIM 8

Mid-Atlantic Colloquium of Studies in Meaning



New York University
April 6th, 2019

MACSIM 8

Mid-Atlantic Colloquium of Studies in Meaning

Hosted by New York University,
April 6th, 2019

Schedule of events

Time	Location	Event
9:30 - 10:00	Jurow Hall	<i>Breakfast, registration, & welcome</i>
10:00 - 10:30	Jurow Hall	Haoze Li (NYU): Making wh-phrases dynamic: A case study of Mandarin wh-conditionals
10:30-11:00	Jurow Hall	Yu Cao (Rutgers): Causal and Instrumental <i>How</i> Questions
11:00-11:15	Jurow Hall	<i>Coffee Break</i>
11:15-11:45	Jurow Hall	Nattanun (Pleng) Chanchaochai (co-authored with J��r��my Zehr, both Penn): Ambidirectionality and Thai mid-scale terms: when ‘warm’ means less hot
11:45-12:15	Jurow Hall	Ibtisam Ammouri (CUNY): The prohibition on indefinite subjects in Arabic
12:30-1:40	10 Washington Place	<i>Lunch/Poster session 1 – see next page</i>
1:40-2:50	10 Washington Place	<i>Lunch/Poster session 2 – see next page</i>
3:00-3:30	Jurow Hall	Annemarie van Dooren (UMD): Figuring out epistemic uses of English and Dutch modals: The role of aspect
3:30-4:00	Jurow Hall	Dionysia Saratsli (co-authored with Stefan Bartell and Anna Papafragou, all UDel): Artificial language learning and the learnability of semantic distinctions: the case of evidentiality
4:00-4:30	Jurow Hall	Emory Davis (co-authored with Barbara Landau, both Johns Hopkins): Seeing vs. Seeing That: Interpreting reports of direct perception and inference
4:30-4:45	Jurow Hall	<i>Coffee Break</i>
4:45-5:45	Jurow Hall	Invited Talk: Melissa Fusco (Columbia): Agential Free Choice
5:45-6:00	Jurow Hall	<i>Business meeting</i>
6:30-8:30	Vapiano, 13th and University Pl.	<i>Dinner</i>

MACSIM 8

List of poster presentations

Poster Session 1 (12:30-1:40, 10 Washington Place)

- Omar Agha** (NYU): Event structure and non-culmination in Khoekhoe
Stefan Bartell (UDel): Effect of indefinite form on donkey anaphora interpretation
Karen Clothier (Johns Hopkins): What counts as ‘many’?
Tris Faulkner (Georgetown): An experimental investigation of mood variation in Spanish emotive-factive clauses
Mina Hirzel (UMD) (co-authored with **Valentine Hacquard** (UMD) and **Ailís Cournane** (NYU)): Young children use different modals for different “flavors”
Jinwoo Jo and Bilge Palaz (UDel): Licensing pseudo incorporation in Turkish
Alyssa Kampa (co-authored with **Kaja Jasińska** and **Anna Papafragou**) (UDel): Scalar implicature development in 4- and 5-year-olds is supported by language and executive function networks
Yeonju Lee (CUNY): Wh-in-situ interrogatives through the lens of split wh-NPIs
Maxime Tulling (NYU): Neural correlates of linguistic modality
Yosiane White (Penn): Words take time: Auditory stimuli and strategic processing in semantic priming
Akitaka Yamada (Georgetown): Embedded speech act layers and enhancement effect
Jérémy Zehr (UPenn) (co-authored with **Paul Egré** (IJN-ENS Paris)): Contradictory descriptions with absolute adjectives

Poster Session 2 (1:40-2:50, 10 Washington Place)

- Anouk Dieuleveut** (co-authored with **Valentine Hacquard**) (UMD): Implicative inferences of ability statements with perception verbs
Kajsa Djärv (Penn): Propositional attitude reports: The syntax of presupposition & assertion
Lucia Donatelli (Georgetown): The morphosemantics of Spanish gender: A case study of pseudo-incorporation
Michael Donovan (UDel): Uses of oddball imperatives
Ivana Đurović (CUNY): Neg-raising asymmetry in SerBo-Croatian
Megan Gotowski (co-authored with **Kristen Syrett**) (Rutgers): Probing children’s early comprehension of comparative constructions
Ioana Grosu (NYU): An extended minimal networks theory for backtracking counterfactuals
Yue Ji (co-authored with **Anna Papafragou**) (UDel): Boundedness in event and object cognition
Matthias Lalis (Johns Hopkins): Distributed neural encoding of binding to thematic roles
Jane Lutken (Johns Hopkins): An optimality theory analysis of scope marking at the syntax/semantics interface
Yağmur Sağ (Rutgers): The curious case of measure semantics
Benjamin Shavitz (CUNY): Analyzing the infelicity of tantalizing statements
Sigwan Thivierge (UMD): High shifty operators in Georgian indexical shift

Abstracts

Event structure and non-culmination in Khoekhoe

Omar Agha (NYU)

Introduction Khoekhoe (aka Damara-Nama, in Central Khoisan) has an aspect marker HA in (1), which is systematically ambiguous between a present progressive (ongoing) reading and a perfective/completive meaning with both activity predicates and accomplishment predicates.

(1) *Bare HA, no tense*

tǎntākō-p kē ǀʔũú hǎǎ
Tantako-M.SG DECL eat HA

‘Tantako is eating/has eaten.’

(2) *HA under past tense*

tǎntākō-p kē kò ǀʔũú hǎǎ ʔĩĩ
Tantako-M.SG DECL RCT.PST eat HA COP.PST

‘Tantako had eaten.’

I will show that a uniform semantic analysis of HA is best achieved by manipulating the causal structure of accomplishments, following work by Ramchand (2008) and Tatevosov (2008). Specifically, HA is a modifier of event predicates that may attach at different levels within the verb phrase, producing two readings of HA with accomplishments. This analysis is contrasted with an approach that factors out non-culmination into a separate imperfective operator, such as Altshuler (2014).

Data Accomplishments under bare HA need not have culminated, as shown in (3). HA may also occur under a recent (or remote) past tense marker such as *ko* ‘RCT.PST’, as in (2). In the past tense, HA forms a construction whose semantics resemble the English past perfect. This can be shown by using a *when*-clause to pull apart reference time and event time (full tests are omitted here for space reasons, but included in the paper).

Although (1) (an activity sentence) and (3) (an accomplishment sentence) are both ambiguous, bare *hǎ* contributes different aspectual information depending on the inner aspect (*aktionsart*) of the predicate. The full pattern is summarized in the table below.

Inner Aspect	Non-past HA	Predicates tested
Achievements	perfective, result state still holds	<i>break the lamp</i>
Accomplishments	ambiguous: perfective/continuous	<i>cook one pig, eat the potato</i>
Activities	ambiguous: perfective/continuous	<i>dance</i>
States	inchoative	<i>be happy, be tall, be healthy</i>

Notice that there is no ambiguity with states or achievements.

Analysis: Informal Summary In all its uses HA is a modifier of event predicates that locates some eventuality at the reference time. But in accomplishments, the eventuality that HA modifies may be either a result state or a process. States and achievements, on the other hand, do not provide two separate eventualities for HA to modify, and are therefore not ambiguous. (In what follows, I use *event* to mean *eventuality*, which includes both dynamic events and states.)

Framework There are many possible ways to implement this idea. In the present work we pursue a compositional approach to inner aspect, following a simplified version of Ramchand’s (2008) First Phase Syntax. In the compositional framework, the logical form of an accomplishment verb phrase like *tĩĩ-p-a sǎĩ* ‘cook a *tĩĩp*’ in (3) contains an existentially quantified *process* event and an existentially quantified *result* event (or result state). These distinct event variables are introduced by distinct operators *proc* and *res*, whose scope is shown in (4). (Ramchand’s *init* operator does not play a crucial role here.)

- (3) tǎntākō-p kē tīrī-p-a sǎi hǎi
T-M.SG DECL wild.pig-M.SG-OBL cook HA (4) $\llbracket init \rrbracket(\llbracket proc \rrbracket(\llbracket res \rrbracket(ttirip)))$
‘Tantako has cooked/is busy cooking the *tirip*.’

Entries for *proc* and *res* are given in (5) and (6). (Note that these entries have been substantially simplified from Ramchand (2008).) Let the types *e* and *t* be as usual. The type *v* is the type of eventualities, and the annotation x_σ says that *x* is a variable of type σ .

$$(5) \quad \llbracket proc \rrbracket = \lambda P_{\langle v, t \rangle} \lambda e_v. \text{cooking}(e) \wedge \exists e'_v [\text{cause}(e, e') \wedge P(e')]$$

$$(6) \quad \llbracket res \rrbracket = \lambda x_e \lambda e_v. \text{cooked}(e, x)$$

Entry HA combines with an event predicate *P*, returns the set of *P*-events that hold throughout the reference time *t*, and presupposes that *P* has the subinterval property down to some granularity. The presupposition is not formalized here, but can be expressed via stratified subinterval reference (Champollion, 2017). The variable *t* over time intervals is initially free, but gets bound by λ -abstraction later in the derivation.

$$(7) \quad \llbracket HA \rrbracket_{\langle \langle v, t \rangle, \langle v, t \rangle \rangle} = \lambda P_{\langle v, t \rangle} \lambda e_v. P(e) \wedge \tau(e) = t$$

In an accomplishment sentence like (3), HA has at least two possible scopes, given in (8). The first option in (8a) requires that the result state (the pig being cooked) occurs at reference time. The second option in (8b) requires that the process of cooking occurs at reference time.

- (8) a. $\llbracket proc \rrbracket(\llbracket HA \rrbracket(\llbracket res \rrbracket(ttirip)))$ (result holds at *t*)
 $= \lambda e_v. \text{cooking}(e) \wedge \exists e'_v [\text{cause}(e, e') \wedge \text{cooked}(e', x) \wedge \tau(e') = t]$
b. $\llbracket HA \rrbracket(\llbracket proc \rrbracket(\llbracket res \rrbracket(ttirip)))$ (process ongoing at *t*)
 $= \lambda e_v. \text{cooking}(e) \wedge \tau(e) = t \wedge \exists e'_v [\text{cause}(e, e') \wedge \text{cooked}(e', x)]$

In contrast, the LFs for states and achievements (omitted for space reasons) only contain at most one possible attachment site for HA.

Conclusion A compositional view of inner aspect naturally gives rise to certain ambiguities by allowing certain aspectual operators to modify different events within a complex event description. In the paper, I extend this approach to include the behavior of states, achievements, and activities under bare HA, and to the properties of HA under different tenses.

References

- Altshuler, Daniel. 2014. A typology of partitive aspectual operators. *Natural Language & Linguistic Theory* 32(3):735–775.
Champollion, Lucas. 2017. *Parts of a Whole: Distributivity as a Bridge between Aspect and Measurement*. Oxford Studies in Theoretical Linguistics. Oxford, New York: Oxford University Press. ISBN 978-0-19-875512-8.
Ramchand, Gillian. 2008. *Verb Meaning and the Lexicon: A First-Phase Syntax*. Cambridge Studies in Linguistics ; 116. Cambridge, UK %3B New York: Cambridge University Press. ISBN 978-0-511-38791-3.

Tatevosov, Sergei. 2008. Subevental structure and non-culmination. In O. Bonami and P. Cabredo Hofherr, eds., *Empirical Issues in Syntax and Semantics* 7, 393–422.

The prohibition on indefinite subjects in Arabic
Ibtisam Ammouri

This research sets out from the question: why are indefinites excluded from subject positions in Arabic (as in (1) below)? First, I show that this prohibition can be explained using the Mapping Hypothesis (Diesing, 1992), but only with the added assumption that subjects in Arabic cannot reconstruct to Spec-vP. Second, I provide independent evidence that lack of subject reconstruction is indeed a general phenomenon in the language. Finally, I propose a structural reason for the absence of reconstruction effects.

- (1) *klaab ʕam bilʕabu bi-s-saaha
dogs PROG playing.PL in-the-yard.
Intended: ‘dogs are playing in the yard’.

It is known that Bare Nouns (BNs) cannot appear in subject positions in many languages, such as Romance. Chierchia (1998) proposed that the reason for this restriction is that BNs in these languages are parametrically set to denote properties type $\langle e, t \rangle$, and subject positions require elements whose semantic type is argumental (i.e. GQs type $\langle et, t \rangle$ or referring expressions type e). To fix the type mismatch problem, Romance languages require weak determiners in subject positions, as in (2):

- (2) *(Unos) chicos han entrado. [Chierchia, 1998:342, ex. (4.a)]
a.PL kids have entered.

Despite the fact that Arabic seems to have the same parametric setting as Romance by Chierchia’s definition, weak determiners do not fix the problem, as shown in (3) below. This indicates that type incompatibility is not the source of the ungrammaticality of indefinite subjects in Arabic.

- (3) * talat/kam/ktiir walad faat-u.
Three/some/many kid entered-PL.
Intended ‘some kids entered’.

A different approach to the distribution of BNs was proposed by Diesing (1992) based on data from Germanic languages, where BNs are allowed in subject positions, but their interpretation depends on the kind of predicate they combine with. Subjects of Stage-Level predicates are ambiguous between the kind and the existential readings (4.a), while subjects of Individual-Level predicates only have the kind reading (4.b).

- (4) a. firemen are available. (Characteristically/now at this station)
b. firemen are altruistic. (Characteristically)

Neither of these readings is available in Arabic since indefinites are ungrammatical in subject positions, and reference to kinds is done using definite nouns only. Nonetheless, the framework in which Diesing explains the data in (4) is useful in understanding the ban on indefinite subjects in Arabic. Following Heim (1982), Diesing assumes that weak indefinites do not have quantificational force and thus introduce free variables to the semantic representation. In the absence of overt quantifiers, these variables are bound by default covert operators, with GEN scoping over the restriction of the sentence and $\exists$ over its nuclear scope. Diesing then proposes the Mapping Hypothesis (MH) which states that material inside the VP¹ is mapped to the nuclear scope and material above the VP is mapped to the restriction. With the MH, the ambiguity of (4.a) is explained as a result of the possibility of interpreting the variable in its surface position, where it is bound by GEN and understood generically, or reconstructing the subject to its base position within the vP, where it is bound by $\exists$ and understood existentially.

¹ To restate the MH in Minimalist terms, I will use vP, rather than VP. The crucial point is that the nuclear scope is the maximal projection of the verb and subjects may be reconstructed to their thematic position in it.

The unavailability of the existential reading in (4.b) according to Diesing, stems from the fact that subjects of Individual-Level predicates are born in Spec-IP, where the only available binder is GEN. That is, subjects as in (4.b) cannot reconstruct to the scope of Existential Closure (the vP) because they were never there.

The crucial idea that I would like to adopt from the MH is that the existential reading of subjects is contingent upon reconstruction to Spec-vP. If reconstruction is blocked for some reason, then the facts of Arabic can be explained. Given that the subject variable cannot be bound by GEN (since reference to kinds is not achieved with indefinites), it will remain unbound in cases like (1), represented in (5). Clearly, this is not a truth value, but an $\langle e, t \rangle$ function which is expected to be ungrammatical as a sentence.

(5) $x.[Dogs(x) \text{ and } Playing(x)]$

But is there independent evidence to the unavailability of subject reconstruction in Arabic? The answer is yes, reconstruction effects observed with quantified subjects in other languages are absent in Arabic as well. The scope ambiguity in (6), for example, does not arise in Arabic (7):

(6) we cannot start because all the guests have not arrived yet. [All > NEG, NEG > All]

(7) mniʔdar-ʃ nibda laʔinno kul ʔid-djuuf baʔed-hen ma wisʕluuf. [All > NEG]

The ambiguity in (6) is believed to result precisely from the possibility of reconstruction. If the QP *all the guests* is interpreted in its surface position, the sentence means no guest has arrived, but if it is reconstructed to Spec-vP (below negation), the sentence means not all guests have arrived. In the Arabic translation (7) on the other hand, the QP only takes wide scope with respect to negation, the reading achieved by reconstruction is not available even though it is a more likely scenario. With this independent evidence, the question of this research can be answered as follows: indefinites cannot be subjects because (i) subjects must reconstruct to Spec-vP to get the existential reading and (ii) subjects in Arabic do not undergo reconstruction.

This answer immediately leads to another question and that is why is subject reconstruction blocked in Arabic? As a speculation, I would like to propose a structural account similar to Diesing's treatment of subjects of Individual-Level predicates. But unlike Diesing, I do not assume that subjects in Arabic originate at the specifier of the functional head (IP or TP). Instead, I would like to suggest that the subject is not an argument of the predicate at all, but rather an external phrase binding a silent pronoun which serves as the agent of the predicate and the grammatical subject of the TP domain, as in (8).

(8) [subject] $\lambda_1 [\dots [_{TP} \text{pro}_1 \dots \text{predicate}]]$

Given that Arabic is a Null-Subject Language, this assumption is plausible. The structure is interpretable by Predicate Abstraction (PA), but only if there is a binder to the variable inside the subject phrase (e.g. a strong quantifier). What lends precedence to this idea is the existence of the Multiple Subject Construction, in which a *t*-type sentence is used as a predicate (9), and the fact that such a sentential predicate can be coordinated with what looks like a bare predicate (10):

(9) ram raas-o buʔaʕ-o

Ram head-his hurts-him.

'Ram (is such that) his head hurts'.

(10) Ram raas-o buʔaʕ-o w-mʕasʕsʕeb

Ram head-his hurts-him and-upset.

'Ram (is such that) his head hurts and he is upset'.

The coordination in (10) is possible only if *upset* and *his head hurts him* are of the same semantic type (*t*), suggesting that apparent simple predicates are in fact full sentences with a silent *pro* in the subject position. If the pronoun bound by the matrix subject is not itself a subject (e.g. *his/him*), it cannot be dropped. This analysis captures Aoun et al.'s (2009) observation that subjects in Arabic might seem like Topics binding resumptive pronouns although they do not behave like typical CILDs. It remains to be determined what kind of projection dominates the structure in (8).

References

Aoun, J. E., Benmamoun, E., & Choueiri, L. (2009). *The Syntax of Arabic*. Retrieved from <https://ebookcentral.proquest.com>.

Chierchia, G. (1998). Reference to Kinds across Language. *Natural Language Semantics*, 6(4), 339-405.

Diesing, M. (1992). *Indefinites* (Linguistic inquiry monographs; 20). Cambridge, Mass.: MIT Press.

Heim, I. (1982). The semantics of definite and indefinite noun phrases. *University of Massachusetts at Amherst dissertation*.

Effect of Indefinite Form on Donkey Anaphora Interpretation

Stefan Bartell

Introduction Previous approaches to donkey anaphora have focused on different factors that bias toward existential or universal interpretations (readings). However, they generally do not discuss effect of form of indefinite on readings. D-type theories such as (Elbourne, 2005) that treat all donkey pronouns as definite descriptions don't predict an effect of indefinite on reading. Dynamic accounts such as (Groenendijk and Stokhof, 1991) don't either. Here I present experimental and non-experimental evidence for an effect of indefinite on donkey sentence readings and discuss implications for deciding between competing theories of donkey anaphora such as those mentioned above.

In (Bartell, 2018) I presented evidence for an effect of indefinite on donkey sentence reading in Hungarian based on introspective judgments of eight speakers. Universal readings were preferred (numerically) in the pattern (1). Also, 'one/a' indefinites and bare nouns preferred existential readings, while free choice item phrases preferred universal.

(1) 'one/a' < bare noun < free choice item

Based on (Bartell, 2018), I hypothesized that English would show a similar effect of indefinite form to Hungarian: *some* indefinites would generate the least universal (most existential) readings, followed by *a*, followed by *any*; this pattern is illustrated in (2). The hypothesis is based on the idea that indefinites that can or prefer to take wider scope yield more existential readings, while those that take narrower scope yield more universal.

(2) *some* < *a* < *any*

Method Based on (Geurts, 2002)'s method, participants were presented with donkey sentences and asked to judge whether they were true or false with respect to scenarios; these were only compatible with an existential reading such that a judgment of "True" corresponded to an existential reading and "False" to universal. 18 different donkey sentences and corresponding scenarios were constructed. All were of the same form including a relative clause and *every* NP. Six scenarios and sentences were presented to each subject, two each with *some*, *a*, and *any*, along with 12 fillers. Two sets of 72 subjects were recruited on Amazon's Mechanical Turk website, each for a separate sub-experiment: one with donkey sentences in present tense and one with past tense and temporal modifiers such as 'last year'; stimuli were otherwise identical.

Results Considering only subjects who answered all filler questions correctly, as for past tense, *some* and *a* and *a* and *any* did not differ, but *some* and *any* did ($t(57) = 2.80, p = 0.007$). As for present tense, *some* and *a* did not differ, while *a* and *any* did ($t(49) = 2.91, p = 0.005$), as did *some* and *any* ($t(49) = 4.2, p = 0.0001$). For the combination of present and past tense, in contrast, the predicted pattern (2) did hold. *some* and *a* differed ($t(107) = -2.17, p = 0.03$), as did *a* and *any* ($t(107) = 3.10, p = 0.002$), as did *some* and *any* ($t(107) = 4.96, p < .0001$). In addition, present tense donkey sentences elicited more universal readings than past tense (Welch's $t(299.11) = 3.82, p = 0.0002$).

Table 1: Proportion Universal Readings

tense/indefinite	some	a	any
past	0.21	0.24	0.33
present	0.32	0.44	0.62

Discussion Experimental data on English indefinites support the hypothesized effect of indefinite on donkey sentence readings along the lines of introspective judgments in Hungarian. They seem difficult to reconcile with purely D-type or dynamic approaches but are compatible with a (Chierchia, 1995)-inspired hybrid approach in which pronouns allow different means of anaphora resolution. Furthermore, I propose that pronoun resolution is affected by antecedent semantics. Indefinite determiners such as *some* bias toward dynamic binding-like resolution, while *any* biases toward a D-type strategy. Following (Chierchia, 1995), a dynamic binding-like strategy biases toward existential readings, while a D-type strategy biases toward universal. Also, these results are compatible with an account of donkey readings based in lexical underspecification of (a) indefinites, cf. (Brasoveanu, 2008). Present tense may allow universal quantification over situations (cf. (Elbourne, 2005)), while past tense or a temporal modifier implying a single situation (or event of the nuclear scope predicate) may block it.

(Schwarz, 2009) contrasts weak (uniqueness based) and strong (familiarity) definites in German and pursues D-type and dynamic analyses respectively. (Patel-Grosz and Grosz, 2017) extend (Schwarz, 2009)’s distinction to two types of German pronouns. A natural question is whether different types of pronoun bias toward different means of anaphora resolution, namely D-type vs. dynamic. A second question is whether donkey sentence reading can be influenced by pronoun form (due to different resolution). Results of another survey I conducted on acceptability of donkey sentences in Hungarian crossing indefinite and pronoun forms indicate that antecedent indefinite form may influence means of resolution and that different forms of pronoun are aligned with different types of resolution. 17 Hungarian speakers judged the acceptability of donkey sentences on a scale of 1 to 5. Results are given in Table (2) and summarized in (3).

Table 2: Acceptability of Hungarian Donkey Sentences

pronoun/indefinite	‘one/a’	bare noun	free choice item
null	2.8	3.2	4.1
demonstrative	4.4	4.5	2.4

- (3) null: ‘one/a’ = bare noun < free choice
demonstrative: ‘one/a’ = bare noun > free choice

One explanation for the observed interaction between indefinite and pronoun form is that different antecedent indefinites prefer different resolution, and different pronouns, with different resolution, can accommodate their antecedent better. In particular, free choice item indefinites and null pronouns prefer D-type resolution as opposed to a dynamic binding-like strategy, while ‘one/a’ indefinites and bare nouns and demonstrative pronouns prefer the reverse. If pronoun strength is relative, following (Patel-Grosz and Grosz, 2010), then in Hungarian, null pronouns are weaker than demonstrative. This weak-strong contrast may correspond to more D-type vs. dynamic-like resolution, cf. (Schwarz, 2009). These results provide a parallel to the pattern of reading preference with Hungarian indefinites (1). Those indefinites that preferred a demonstrative pronoun (‘a/one’ and bare nouns) also preferred existential readings. Free choice indefinites preferred a null pronoun and universal readings. In summary, means of anaphora resolution may be influenced by antecedent form as well as pronoun form. With donkey anaphora, different resolution may give rise to different interpretation.

References

- Bartell, S. (2018). Strong donkey indefinites as number neutral. Poster presented at the 2018 U. of Göttingen Workshop on Ambiguity (Ambigo).
- Brasoveanu, A. (2008). Donkey pluralities: plural information states versus non-atomic individuals. *Ling. & Phil.*, 31(2):129–209.
- Chierchia, G. (1995). *Dynamics of Meaning: Anaphora, presupposition, and the theory of grammar*. U. of Chicago Press.
- Elbourne, P. (2005). *Situations and Individuals*. MIT Press.
- Geurts, B. (2002). Donkey business. *Linguistics and philosophy*, 25(2):129–156.
- Groenendijk, J. and Stokhof, M. (1991). Dynamic predicate logic. *Ling. & Phil.*, 14(1):39–100.
- Patel-Grosz, P. and Grosz, P. (2010). On the typology of donkeys: Two types of anaphora resolution. In *Proceedings of Sinn und Bedeutung*, volume 14, pages 339–355.
- Patel-Grosz, P. and Grosz, P. (2017). Revisiting pronominal typology. *Linguistic Inquiry*, 48(2):259–297.
- Schwarz, F. (2009). *Two types of definites in natural language*. PhD thesis, U. Mass Amherst.

Causal and Instrumental *How* Questions

Introduction. Besides the basic use of asking about manners, we have causal *how* questions (HQs) like (1) that target the *cause* by which an *effect* occurs and instrumental HQs like (2) that target the *means* by which a *purpose* is achieved. Often the former can be feasibly answered by a *because* clause but not an agentive *by* gerund, whereas the opposite is true for the latter. (3) shows that HQs with an existential modal may alternate between the two readings depending on the context. With agentivity tests (e.g., Dowty 1979), it can be generalized that *how* has a causal reading if associated with a non-agentive predicate or modal like *sink* or *can* and an instrumental reading if associated with an agentive predicate like *solve* or *drive*.

- (1) A: *How did the Titanic sink?* B: *Because it hit an iceberg.* # *By hitting an iceberg.*
 (2) A: *How did Ben solve the puzzle?* B: *By writing a program.* # *Because he wrote a program.*
 (3) A: *How can you drive so fast?* (a foreigner asks) B: *Because there's no limit in Germany.*
 (a leaner asks) B: *By pressing hard on gas.*

This study attempts to give a first (to the author's knowledge) unified account of both readings that captures this generalization.

Proposal Outline. I keep the lexical entry for manner *hows* (see Stanley 2011 a.o.), a quantifier over event properties as (4); it raises in syntax (see Fig. 1) and leaves a trace (a variable P) that interacts with a Cause/modal/By head in question nuclei to yield causal and instrumental HQs.

$$(4) \llbracket \text{how} \rrbracket = \lambda Q. \bigcup_{P \in \mathcal{D}_{s \rightarrow vt}} Q(P)$$

Distribution of these readings is captured by the fact that only non-agentive predicates undergo causative coercion; modals carry built-in causality; instrumentality presupposes agentivity. Also, causality, modal causality, and instrumentality are uniformly expressed by counterfactual (CF) dependency (Lewis 1973) between propositions: given Kratzer's (1981) ordering source $f(w)$, q CF-depends on p , i.e., $p \Box_{f(w)} q$ if a world u satisfying both p and q is equally or more similar to the situation described by $f(w)$ than any world v satisfying p but not q , i.e., $u \leq_{f(w)} v$.

Causality. In English only non-agentive predicates undergo causativization (see Dowty 1979; Pylkkänen 2008) or causative coercion (see Sæbø, 2016). Implementing the latter with Pylkkänen's (2008) Cause head as adapted in (5), the trace of *how* modifies CauseP by supplying P as a description of the *cause*; see Fig. 1. $\text{CAUSE}_{w,f}(e_1, e_2)$ wraps the CF-dependency of non-occurrence of an e_2 -like event (*effect*) on non-occurrence of an e_1 -like event (*cause*), given a circumstantial ordering source $f(w)$.

- (5) a. $\llbracket \text{Cause}_{w,f} \rrbracket = \lambda P \lambda e_1. \exists e_2 P_w(e_2) \wedge \text{CAUSE}_{w,f}(e_1, e_2)$
 b. $\text{CAUSE}_{w,f}(e_1, e_2) \Leftrightarrow \lambda u. \nexists e'_1 \forall P P_w(e_1) \rightarrow P_u(e'_1) \Box_{f(w)} \lambda v. \nexists e'_2 \forall P P_w(e_2) \rightarrow P_v(e'_2)$

After $\exists$ -binding e_1 and abstracting over w , we derive a question nucleus that says were a P -event not to occur, most likely a P^{VP} -event would not occur either. Applying this to (1), where $P^{\text{VP}} = \lambda w \lambda e. \text{sink}_w(e) \wedge \text{th}_w(e) = \text{Titanic}$, we can derive a set of propositions of the form *a P-event causes Titanic's sinking*.

Modal Causality. Following von Stechow and Iatridou (2005), Kratzer's (1981) theory is cast into syntax as in (6). With the help of an operator T , the trace of *how* modifies a circumstantial modal base $g(w)$ by adding occurrence of a P -event; see Fig. 2.

- (6) a. $\llbracket \text{can}_{w,f} \rrbracket = \lambda g \lambda p. \exists v v \in \perp_{f(w)} \cap g(w) \wedge p(v)$, where $\perp_{f(w)} W = \{u \in W \mid \forall v \in W u \leq_{f(w)} v\}$
 b. $T = \lambda P \lambda g \lambda w. g(w) \cup \{\lambda u. \exists e_1 P_u(e_1)\}$

When no world compatible with $g(w)$ is ranked as closest to a situation given by $f(w)$, admitting a new fact into $g(w)$ might make a difference. So after abstracting over w , we derive a question

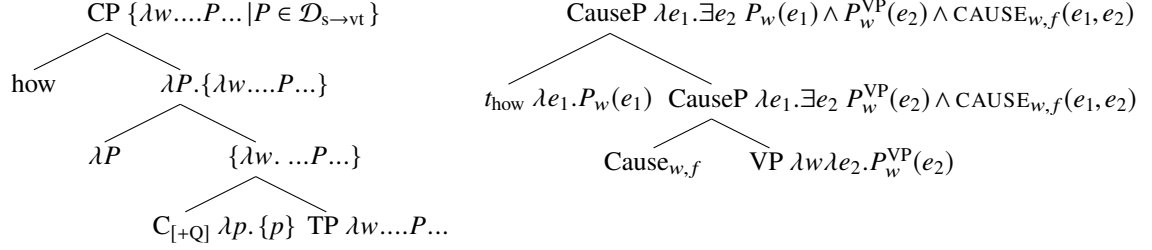


Figure 1: **Left:** Schematized HQ derivation. **Right:** *How* interacting with *Cause*.

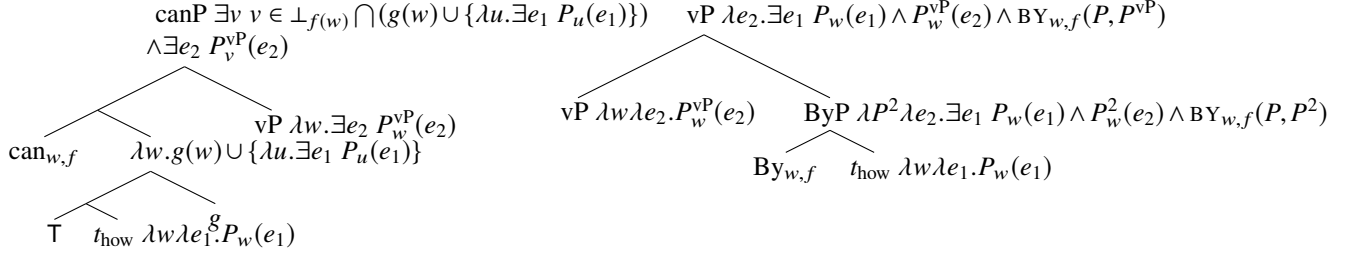


Figure 2: **Left:** *How* interacting with *can*. **Right:** *How* interacting with *By*.

nucleus that says if a P -event occurred, then a P^{VP} -event would be epistemically/deontically possible. Applying this to (3), where $P^{VP} = \lambda w \lambda e. drive\text{-}so\text{-}fast_w(e) \wedge ag_w(e) = \text{you}$, we can derive a set of propositions roughly of the form *if a P -event occurs, you can drive so fast*. Interestingly, the question nucleus we just derived entails that occurrence of a P^{VP} -event CF-depends on $g(w)$ plus an occurrence of a P -event. In this sense, modals carry built-in causality.

Instrumentality. In analogy to a *Cause* head, I propose a *By* head in (7) to introduce instrumentality. The trace of *how* provides a event property P that describes a *means* (see Fig. 2), which a *By* head relates to an event property that describes a *purpose*. The latter is presupposed to be agentive as a *purpose* necessarily belongs to an agent. Following Rissman (2011) I analyze instrumentality as teleological modality discussed by von Fintel and Iatridou (2005), where a teleological ordering source $f(w)$ ranks worlds according to how well they satisfy an agent's ideals. $BY_{w,f}(P^1, P^2)$ renders von Fintel and Iatridou's theory in terms of CF dependency: given a teleological ordering source, an occurrence of a P^1 -event (*means*) CF-depends on an occurrence of a P^2 -event (*purpose*); in plain, some world where the purpose co-occurs with the means is equally or more ideal than any world the purpose occurs without that means.

- (7) a. $\llbracket By_{w,f} \rrbracket = \lambda P^1 \lambda P^2 : AGENTIVE(P) \lambda e_2. \exists e_1 P_w^1(e_1) \wedge P_w^2(e_2) \wedge BY_{w,f}(P^1, P^2)$
 b. $AGENTIVE(P) \Leftrightarrow \forall w, e P_w(e) \rightarrow \exists x ag_w(e) = x$
 c. $BY_{w,f}(P^1, P^2) \Leftrightarrow \lambda u. \exists e_2 P_u^2(e_2) \Box_{\rightarrow f(w)} \lambda u. \exists e_1 P_u^1(e_1)$

After $\exists$ -binding e_2 and abstracting over w , we derive a question nucleus that says a P^{VP} -event is done by doing a P -event. Applying this to (2), where $P^{VP} = \lambda w \lambda e. solve_w(e) \wedge ag_w(e) = \text{Ben} \wedge th_w(e) = \text{the-problem}$, we can derive a set of propositions of the form *Ben solved the problem by doing a P -event*. In (3), instead of modifying the modal base assignment, *how* could have started combining with *By*. Hence the instrumental reading besides the causal reading derived before.

Conclusion. Sticking to a single lexical entry of *how*, I provide a unified account of causal and instrumental HQs under CF dependency. Distribution of different HQs is captured by the connection between non-agentivity and (modal) causality, and agentivity and instrumentality.

Selected References

- Dowty, David. 1979. *Word meaning and Montague grammar*. D. Reidel Publishing Company.
- von Fintel, Kai, and Sabine Iatridou. 2005. What to do if you want to go to Harlem: Anankastic conditionals and related matters. Ms.
- Kratzer, Angelika. 1981. The notional category of modality. In *Words, worlds, and contexts. New approaches in word semantics*, ed. H.-J. Eikmeyer and H. Rieser, 38–74. Berlin: Mouton de Gruyter.
- Lewis, David. 1973. Causation. *The Journal of Philosophy* 70:556–567.
- Pylkkänen, Liina. 2008. *Introducing arguments*. MIT Press.
- Rissman, Lilia. 2011. Instrumental with and use: modality and implicature. In *Proceedings of SALT 21*, 532–551. LSA.
- Sæbø, Kjell Johan. 2016. “How” questions and the manner–method distinction. *Synthese* 193:3169–3194.
- Stanley, Josan. 2011. *Know how*. Oxford University Press.

Ambidirectionality and Thai mid-scale terms: when ‘warm’ means less hot

Nattanun Chanchaochai and Jeremy Zehr - University of Pennsylvania

Empirical observations It may seem an uncontroversial thing to say that *to get warmer* means to undergo an increase rather than a decrease in temperature. This, however, may not appear intuitive to Thai speakers, for the Thai translation of ‘get warmer,’ *ʔùn k^hûm* (literally ‘warm ascend’) can describe not only increases in temperature, but also decreases from hot to moderately warm [1]. The same observation holds for *salû:a k^hûm* (‘dim ascend’) and *c^hû:n k^hûm* (‘damp ascend’) which can respectively describe not only increases in darkness or wetness, but also changes from highly to moderately dark or wet. Such ambidirectional interpretations are unavailable for more extreme scalemates (*hot/cold*, *dark/bright*, *wet/dry*). After rejecting two other possible analyses, we propose that the Thai mid-scale predicates are semantic equivalents of English *warm*, *dim* and *damp* and we give a semantics for *k^hûm* (‘ascend’) that accounts for cases of ambidirectionality.

To be rejected 1: mild rather than warm One might consider translating *ʔùn* as *mild*, and *ʔùn k^hûm* as *get milder*, which also exhibits ambidirectionality [2]. Such an analysis has two weaknesses: first, it would require new, parallel translations for *salû:a* (‘dim’) and *c^hû:n* (‘damp’), and, second, it predicts *too mild* to be a good translation for excess-constructions built with *ʔùn*. This prediction is not borne out: while *too mild* roughly means *too moderate* [3], the Thai sentence [4] unidirectionally denotes excessively *high* temperatures in much the same way as *too warm*.

To be rejected 2: turn A rather than get A-er In another plausible type of account, *salû:a k^hûm* (‘dim ascend’) would receive a non-scalar interpretation along the lines of *turn dim*. Such an analysis would be compatible with ambidirectionality [5], but further empirical observations lead us to discard it: regardless of the direction of the change, *salû:a k^hûm* (‘dim ascend’) can be modified by a measure phrase referring to the *difference* in illumination [6], whereas *turn 50 lumens dim* describes a *final* illumination of 50 lumens.

Our proposal: k^hûm as moving away from alternative We propose that *k^hûm* describes changes whose *initial* state satisfies a salient alternative of the scalar predicate, and whose *final* state satisfies the scalar predicate itself *instead* [7]. We make two additional assumptions: (i) {*cold*, *warm*, *hot*}, {*bright*, *dim*, *dark*} and {*dry*, *damp*, *wet*} represent salient alternative sets, and (ii) *hot*, *dark* and *wet* respectively entail *warm*, *dim* and *damp* at the literal level (i.e. Thai and English are alike). Since we have assumed two alternatives for each predicate, composition with *k^hûm* can always follow two different paths. When composing with *warm*, one path (*cold* as the alternative) results in what could be paraphrased as *warm but no longer cold*, describing an increase in temperatures; the result of the other path (*hot* as the alternative) could be paraphrased as *warm but no longer hot*, describing a decrease in temperatures. When composing with *hot*, choosing *cold* as the alternative results in the expected change, *hot and no longer cold*; choosing the *warm* alternative, however, results in a literal contradiction, paraphrasable as # *hot but no longer warm*. This result is general, given our assumptions: since it is impossible to literally satisfy an *entailing* predicate without at the same time satisfying an *entailed* one, only one path is left for *hot*, *dark* and *wet*, which therefore always yield unidirectional interpretations. As for *cold*, *bright* and *dry*, the change can only go one way, since each has two alternatives that share the same orientation (e.g. *cold* and *not warm/hot*). The semantics we propose needs two refinements. For one, we need a semantic value that can combine with a measure phrase [6]. Second, native speakers’ judgments suggest that the change need not *complete* a move away from the alternative nor up to satisfying the predicate itself [1]. We give our final proposal in [8] where we (i) change the type of the semantic value so that it denotes a degree corresponding to the difference between the degrees at the initial and at the final states, and (ii) quantify over consistent *standard* functions (a method reminiscent of delineation semantics, e.g. Klein 1980) as well as (iii) over *expansions* of the change.

Discussion Our observations on Thai show that scalar expressions give rise to semantic effects that go beyond what is attested in English. We proposed that Thai has an expression, *k^hûm*, that quantifies over its scalar complement’s alternatives. We make two final remarks. First, anecdotal evidence of *English*-speaking children using “warmer” to mean *less hot* suggests a similar semantic analysis of mid-scale comparatives, which invites further investigation. Second, our semantics gives a central role to alternatives. Since the Thai counterpart of *cool* is typically not used in the same types of context as *cold*, it is not an alternative to *cold* and does not normally exhibit ambidirectionality. Remarkably, native speakers’ judgments suggest that ambidirectionality becomes conceivable (if not entirely natural) for *cool* in the rare contexts that license both *cool* and *cold*. That is, if one were to manipulate the context so as to make any two unrelated scalar terms salient asymmetrically entailing alternatives (an *ad-hoc* scale) one would expect the same kind of ambidirectionality. Conversely, if a scalar predicate lacks any salient alternative, one predicts composition with *k^hûm* to be infelicitous, to the extent that the existential quantification over alternatives would yield trivial falsity. We leave this prediction open to further empirical study. Finally, while our account assumes the existence of a set of salient alternatives, it gives no indication as to how that set is determined, or what the appropriate notion of salience is. Researchers have started to tackle such issues from an experimental perspective (see Doran et al. 2012, van Tiel et al. 2012, Schwarz et al. 2016, McNally 2017) but the question remains a matter of empirical debate at the moment.

- [1] tɔ:n-ní: man kô: jaŋ **jen/rɔ:n** jù: ná? t^hũŋ man cà? **?ùn k^hûm** nít-nuŋ kô:t^hŷ?
now it EMP still **cool/hot** ASP FP although it AUX **warm ascend** a little despite
‘It is still **cool/hot**, although it got slightly closer to a moderate temperature.’
- [2] The weather is too **warm/cold**. I’ll wait until it **gets milder**.
- [3] The weather is **too mild** to have an outdoor ice rink, *or* an outdoor swimming pool.
- [4] ná:m kê:w ní: man **?ùn k^hûm** ná?
water CLS-glass this it **warm too** FP
‘This glass of water is too warm’ / # ‘This glass of water is not hot enough.’
- [5] The experiment room was very {**dark / bright**} at first, but then the light turned **dim**.
- [6] mú:a-kí: man **mú:t** mâ:k lɔ:j tɔ:n-ní: **salũ:a k^hûm** ma: hâ:sip lumên lé:w
just now it **dark** very EMP now **dim ascend** DEI 50 lumens already
‘It was very dark before. Now it has become 50 lumens brighter.’
- [7] $\lambda A. \lambda x. \lambda e. \exists B \in \text{Alt}(A) [B(x, e_{\text{start}}) > \text{std}(B) > B(x, e_{\text{end}})] \wedge A(x, e_{\text{end}}) > \text{std}(A).$
 $\approx x$ now meets *A*’s standard but no longer meets its **alternative**’s
- [8] i. $\lambda A. \lambda x. \lambda e. \lambda d. d = \text{diff}(A(x, e_{\text{end}}), A(x, e_{\text{start}})) \wedge$
ii. $\exists s' \sim \text{std}, B \in \text{Alt}(A) [B(x, e_{\text{start}}) > s'(B) > B(x, e_{\text{end}}) \wedge A(x, e_{\text{end}}) > s'(A)] \wedge$
iii. $\exists d', h_{\text{horizon}} [d' = \text{diff}(h_{\text{horizon}}, A(x, e_{\text{start}})) \wedge d' \geq d \wedge h_{\text{horizon}} > \text{std}(A)].$
 $\approx$ degrees representing the amplitude of the change such that *x*:
- no longer meets *B*’s standard but still meets *A*’s for some **consistent** shift of standards
- meets *A*’s actual standard after a change at least as big as the present one

References. Doran, R., Ward, G., Larson, M., and McNabb, Y. (2012). A novel experimental paradigm for distinguishing between what is said and what is implicated. *Language*, 88:124–154. Klein, E. (1980). A semantics for positive and comparative adjectives. *L&P*, 4:1–45. McNally, L. (2017). Scalar alternatives and scalar inference involving adjectives. In Ostrove et al. editors, *Asking the Right Questions: Essays in Honor of Sandra Chung*, pages 17–27. Schwarz, F., Zehr, J., Grodner, D., and Bacovcin, H. A. (2016). Subliminal priming of alternatives does not increase implicature responses. Poster presented the *Logic and Language in Conversation Workshop in Utrecht*. van Tiel, B., van Miltenburg, E., Zevakhina, N., and Geurts, B. (2016). Scalar diversity. *JoS*, 33:137–175.

Introduction. At least three different interpretations of *many* have been repeatedly described in the literature ([2]; [3]; [4]; c.f. [5] for an overview): the cardinal; the proportional; and the reverse proportional interpretations. In [4]’s original characterization of the reverse proportional reading, the arguments of *many* were saturated in the reverse order to that which they appear in the surface syntax. Subsequent work ([2]; [6]; [5]) instead posited that the reverse proportional reading can be derived from the information structure of the D(iscourse)-tree ([7]). Assuming structured meanings for the questions in these D-trees, this proposal aims derive the regular, reverse proportional and cardinal meanings of *many* pragmatically.

Data. Following [2], [4], [5] and [6], the sentence in (1) can be interpreted in at least two ways, illustrated in (1a) and (1b), depending on the pitch contour of the utterance, whether it be taken as Focus ([6]) or Contrastive Topic ([2]).

(1) Many Scandinavians have won the Nobel Prize in Literature.

a. Many Scandinavians have won the Nobel Prize_{F/CT}. *regular proportional*

Of all the things that Scandinavians do, the proportion that have won the Nobel Prize in Literature is larger than the proportion of them that have done other things.

b. Many Scandinavians_{F/CT} have won the Nobel Prize. *reverse proportional*

Of all the people that have won the Nobel Prize in Literature, the proportion of them that have been Scandinavians is larger than the proportion of winners from other countries.

[2] and [5] propose that both interpretations in (1) can be derived from different pitch contours that are typically equated with Focus(F) or Contrastive Topic (CT) marking ([7]).

The analysis. This proposal follows [5] in assuming a degree-based account of *many*, where *many* is decomposed into a degree morpheme, $\llbracket POS \rrbracket$ ((2)) and a generalized quantifier typed *many* with a degree parameter ((3)):

$$(2) \llbracket POS \rrbracket = \lambda Q_{\langle \langle d,t \rangle, t \rangle} . \lambda P_{\langle d,t \rangle} . P \in Q . L_{\langle \langle d,t \rangle, \langle d,t \rangle \rangle} (Q) \subseteq P$$

$$(3) \llbracket many \rrbracket = \lambda d . \lambda P_{\langle e,t \rangle} . \lambda Q_{\langle e,t \rangle} . (|P \cap Q| : |P|) \geq d$$

In this account, Q (the comparison class) is a set of sets of degrees derived from the discourse context and P (the comparison term) is a set of degrees. In the general case, a *many*-utterance asserts that a so-called neutral segment (a measure of central tendency) derived from Q , is a proper subset of P . Critically, a context sensitive F/CT operator ($\sim$) can either associate internally with the sister of *many*, or externally; the reverse reading arises in the former case, the regular reading in the latter.

To account for how the particular elements of Q arise, and the conditions which license one or the other Focus association, the current proposal adopts D-trees as a formal representation of information structure in a discourse ([7]). The discourse participants jointly building these D-trees have internal mental states, which represent their world knowledge, goals, etc., that specify the level of precision required for an answer to be relevant and resolving (c.f. [8]; [9]). Within these D-trees the $\sim$ operator in *many*-utterances is taken to be licensed by – and anaphoric to – dominating questions of the form “How many...”. Thus *many*-utterances are relevant ([10]), resolving answers to “How many...” questions. These questions are denoted as structured meanings following [11], so that for example (4) is represented as (5):

(4) How many Scandinavians have won different honors?

$$(5) \langle \lambda d . (\{ |x : \text{won}(x, \text{Scandinavians})| \} \geq d), \text{DEGREES} \rangle$$

The participants’ world knowledge and goals define or constrain the elements of the restrictor set (*DEGREES*), such that applying the background function $(\lambda d . (\{ |x : \text{win}(NP', x)| \} \geq d))$ to the restrictor set, derives relevant, resolving answers to the question. To account for the effects of F/CT in *many*-utterances, this proposal adopts [12]’s approach in positing that F/CT pitch contour correspond to an operator at LF that anaphorically binds the F/CT-marked element in an utterance to the set of sets derived from the structured meaning of a licensing question, as in (6b) below, bringing these values into the composition of the utterance.

- (6) a. Many Scandinavians have won the Nobel Prize_{F/CT}.
 b. LF: [[POS C]] [1[t₁-many Scandinavians_{F/CT}] have won NP]] ~C]
 $\llbracket C \rrbracket \subseteq$
 $\{\lambda d'. (\{x: \text{Scandinavians}(x)\} \cap \{x: \text{win}(NP, x)\} : |\{x: \text{Scandinavians}(x)\}|) \geq d,$
 $\lambda d'. (\{x: \text{Andorrans}(x)\} \cap \{x: \text{win}(NP, x)\} : |\{x: \text{Andorrans}(x)\}|) \geq d \dots \}$
 c. $L(\llbracket C \rrbracket) \subseteq$

$\lambda d'. (\{x: \text{Scandinavians}(x)\} \cap \{x: \text{win}(NP, x)\} : |\{x: \text{Scandinavians}(x)\}|) \geq d$

Thus, using this approach, the distinct truth conditions illustrated in (1) above are clearly derivable. The licensing question to (1a) above would be as in (4) - (5), the licensing question to (1b) would be as in (7) - (8), with the F/CT alternatives as in (9b):

(7) How many people of each nationality have won the Nobel Prize in Literature?

(8) $\langle \lambda d. (\{x: \text{win}(NP', x)\} \geq d), \text{DEGREES} \rangle$

(9) a. Many Scandinavians_{F/CT} have won the Nobel Prize.

b. LF: [[POS C]] [1[t₁-many Scandinavians_{F/CT}] have won NP]] ~C]

$\llbracket C \rrbracket \subseteq$

$\{\lambda d'. (\{x: \text{Scandinavians}(x)\} \cap \{x: \text{win}(NP, x)\} : |\{x: \text{Scandinavians}(x)\}|) \geq d,$
 $\lambda d'. (\{x: \text{Scandinavians}(x)\} \cap \{x: \text{win}(\text{goldmedal}, x)\} : |\{x: \text{Scandinavians}(x)\}|) \geq d \dots \}$

c. $L(\llbracket C \rrbracket) \subseteq$

$\lambda d'. (\{x: \text{Scandinavians}(x)\} \cap \{x: \text{win}(NP, x)\} : |\{x: \text{Scandinavians}(x)\}|) \geq d$

Furthermore, an account that considers participants' world knowledge and goals can further shed light on the previously observed distinction between cardinal and proportional *many* ([2], [3], [5]) without needing to posit a distinct denotation. For example, consider a situation where the hearer does not know how many Nobel Prizes have been awarded, that is, the set $\{x: \text{won}(\text{np}, x)\}$ would remain well defined, but with no known members. This means that the degree value in the composition, i.e. when intersected with the set $\{x: \text{Scandinavians}(x)\}$ just results in the degrees of the latter set being calculated in the truth conditions. Critically, $\{x: \text{won}(\text{np}, x)\} \neq \emptyset$, meaning that it is assumed to have at least one defined member, that member is simply not known to the hearer. Thus, the calculation of the truth conditions amounts to just the numerator, which is essentially the cardinal denotation proposed by [3], [5] and others, but falls out from the proportional denotation, in a context with a participant with imperfect knowledge.

Conclusion. This proposal is designed specifically to capture the pragmatic characteristics of the proportional interpretation of *many* by integrating a proportional denotation of *many* with a focus-sensitive account using structured meanings for questions. Crucially, adopting a structured meaning account of questions and a decomposed degree-style denotation of *many*, provides the antecedents for the F/CT operator in a format amenable to composition – without recourse to, e.g. possible world semantics – and a format which intrinsically represents the point-wise alternatives relevant for computing the truth conditions, of a proportional and a cardinal *many*, avoiding the need to posit multiple denotations.

References. [1] Barwise, J., & Cooper, R. (1981). Generalized quantifiers and natural language. In *Philosophy, Language, and Artificial Intelligence*. [2] Cohen, A., (2001). Relative readings of *many*, *often*, and generics. *NLS*. [3] Partee, B. (1988). Many quantifiers. *ESCOL* 5. [4] Westerståhl, D. (1985). Logical constants in quantifier languages. *Linguistics and Philosophy*. [5] Romero, M., (2015). POS and the many readings of *many*. *NELS* 46. [6] Herburger, E. (2000). *What counts: Focus and quantification*. [7] Büring, D. (2003). On D-trees, beans, and B-accent. *L&P*. [8] van Rooy, R. (2003). Questioning to resolve decision problems. *L&P*. [9] Ginzburg, J., (1997). Interrogatives: Questions, Facts, and Dialogue. In *The Handbook of Contemporary Semantic Theory*. [10] Roberts, C., (2012). Information structure in discourse: Towards an integrated theory of pragmatics. *S&P*. [11] Krifka, M. (2001). For a structured meaning account of questions and answers. *Audiatur vox sapientia*. 52. [12] von Stechow, K. (1994). *Restrictions on quantifier domains* (Doctoral dissertation). [13] Schwarzschild, R. (1999). GIVENness, AvoidF and other constraints on the placement of accent. *NLS*. [14] Feigenson, L., et al.. (2004). Core systems of number. *TICS*. [15] Solt, S. (2009). Notes on the comparison class. In *International workshop on vagueness in communication*. [16] Solt, S. (2015). Q-adjectives and the semantics of quantity. *J of S*. [17] Rett, J. (2014). *The semantics of evaluativity*.

Seeing vs. Seeing That: Interpreting reports of direct perception and inference

Emory Davis & Barbara Landau, Johns Hopkins University

There is evidence that young children can reason about and differentiate direct perception and inference, but do not fully master comprehension of the particular linguistic forms that can distinctly encode the two different knowledge sources until later in development (Ünal & Papafragou, 2018; Ünal & Papafragou, 2016; Winans *et al.*, 2015). In English, this difference can be marked lexically (e.g. *see* vs. *think* or *guess*) or syntactically. For example, perception verbs with small clause complements (“I saw something happen”) report direct perception of an event, while perception verbs with sentential complements (“I saw that something happened”) can report either direct perception or inference about an event.

We sought to determine whether and when young English-speaking children have mapped the conceptual distinction between direct perception and inference to different syntactic frames expressing this distinction, as well as what pragmatic or other factors may support the inference interpretation for both adults and children. In a series of three experiments, we presented adult and child participants with eight illustrated stories in which one character directly perceives an event and/or a second character encounters only visual evidence that could lead to an inference that the event had occurred. We then asked participants to make judgments about either Direct Perception sentences containing *see* with a small clause complement (e.g. “I saw Fido eating the cookies”), or Inference sentences containing *see* with a sentential complement (e.g. “I saw that Fido had eaten the cookies”).

Across the three experiments, we found that children under 7 years old did not consistently differentiate the Direct Perception and Inference sentences, and showed a preference for interpreting both sentence types as reporting direct perception. Adults and children 7 and older did interpret these two sentence types differently, but only when pragmatics or context supported inference readings for *see* with a sentential complement; specifically, by presenting the two sentence types together – thus contrasting the frames – or by presenting the Inference sentences in cases where there was only an instance of inferring about, and not directly perceiving, an event. These findings indicate that while direct visual perception is the preferred or default interpretation for utterances containing *see* for both children and adults, younger children cannot make use of the same syntactic or pragmatic cues as older children and adults in comprehension. This suggests that English-speaking children under 7 may not understand that *see* can be used to report visually-based inference, and are still in the process of learning the syntax and semantics of perception verbs and integrating their semantic knowledge with pragmatic and contextual information.

References

- Ünal, E. & Papafragou, A. (2016). Production-comprehension asymmetries and the acquisition of evidential morphology. *Journal of Memory and Language*, 89, 179-199.
- Ünal, E. & Papafragou, A. (2018). Relations Between Language and Cognition: Evidentiality and Sources of Knowledge. *Topics in Cognitive Science*, 1-21.
- Winans, L., Hyams, N., Rett, J., & Kalin, L. (2015). Children’s comprehension of syntactically-encoded evidentiality. *Proceedings of NELS 45*, 189–202.

Implicative inferences of ability statements with perception verbs

Anouk Dieuleveut and Valentine Hacquard

This paper discusses the implicative inference that arises with ability statements involving perception verbs (e.g. *see*, *hear*, *smell*; arguably *tell*, *remember*). We show that when *can* combines with a perception verb, it either expresses a ‘general’ ability (1b) or an ‘actualized’ one (1a). In this, perception verbs differ from regular eventives like *swim*, which give rise exclusively to general ability readings (1c), and to statives like *know the answer*, which lead to infelicity (1d):

- (1) a. Ann can see the ghost of her mother right now, #but she doesn’t see it.
 b. Ann can see the ghost of her mother in general, but she doesn’t see it now.
 c. Ann can swim {right now/in general}, but she {doesn’t/isn’t swimming now}.
 d. #Ann can know the answer.

We argue that the paradigm in (1) follows from the way *can* interacts both with (i) grammatical aspect—which is responsible for the implicative/general ability contrast, as has been argued for “actuality entailments” with past ability modals, and (ii) lexical aspect—which is responsible for allowing perception verbs, but not regular eventives and statives to have the extra implicative reading.

Implicative inference & grammatical aspect – Past ability statements give rise to both general ability and implicative readings for all predicates, regardless of their lexical aspect. This contrast is tied to grammatical aspect, as evidenced by languages like French that distinguish perfective and imperfective aspects overtly in the past: in (3), the general ability reading arises with imperfective (3a), the implicative one, with perfective (3b) (Bhatt’s 1999 “actuality entailment”).

- (2) {In her twenties/#Yesterday} Ann was able to swim, but she didn’t.
 (3) a. Anne pouvait nager, mais elle n’a pas nagé. French
 Anne can-past-**impf** swim, but she did not swim
 b. Anne a pu nager, #mais elle n’a pas nagé.
 Anne can-past-**perf** swim, #but she did not swim

Bhatt (1999) proposes that the basic meaning of past ability statements is implicative, and derives the general ability reading from an additional modal layer, a genericity operator (GEN), associated with the imperfective. Here we follow Hacquard’s (2006) implementation, sketched in (4): with root modals, including ability *can*, aspect scopes over the modal but quantifies over the VP-event. Because it outscopes the modal, perfective anchors the VP-event in the world of evaluation, yielding an actual event (4b); imperfective, however, introduces both event quantification *and* quantification over worlds (4a), anchoring the events in the ‘generic’ worlds, which need not include the actual world:

- (4) a. [_{TP} Past [_{AspP} Gen [_{ModP} can [_{VP} Ann swims]]]]
 = *in all generic worlds w accessible from w*, for all relevant e in w, there was a w’ which is a swim by Ann*
 b. [_{TP} Past [_{AspP} Perf [_{ModP} can [_{VP} Ann swims]]]]
 = *There was an e in w*, which in some w’ is a swim by Ann*

We propose that the ‘general ability’ readings in (1b) and (1c) are similarly due to the presence of GEN, which is also associated with the simple present. The implicative reading in (1a) arises in the absence of GEN, when aspect simply anchors the seeing event in the actual world.

Implicative inference & lexical aspect – Why do only perception verbs lead to an implicative inference with present ability statements? We propose that this is due to their hybrid stative/eventive behavior (Dowty 1979). With regular **eventives**, the simple present forces a generic (habitual) reading; an ‘ongoing’ reading requires progressive aspect, as shown in (5). Perception verbs differ from regular eventives: with the simple present, they are ambiguous between a generic/habitual interpretation, responsible for the general ability reading in (1b), and an ongoing interpretation, responsible for the implicative reading in (1a).

- (5) a. Ann swims {every day/*right now}.
b. Ann sees the ghost {every day/right now}.

But perception verbs also differ from **statives** like *know the answer*, which allow an ongoing interpretation with the simple present, but are infelicitous with ability *can* (1d). According to Hackl (2001), this ban against statives is due to ability *can* requiring a prejacent that expresses a change of state. This requirement may follow more generally from a general constraint against the vacuous use of modals, Condoravdi’s (2001) diversity condition, which requires that the eventuality described by the prejacent not be *settled* throughout the worlds of the modal base. We propose that the reason why **perception verbs** do not lead to infelicity with ability *can* is that, unlike statives, they still express a change of state.

Finally, when ability *can* combines with an eventive in the progressive, it leads to a similar infelicity as with statives. Here again, we take this infelicity to follow from the fact that with progressive aspect, the eventive no longer expresses a change of state.

- (6) #Ann can be swimming.

Conclusion – We argue that the peculiar behavior of perception verbs with ability *can* follows from three facts: (i) ability modals scope below aspect (Hacquard 2006) and are thus susceptible to actuality entailments; (ii) ability modals require a change of state prejacent (Hackl 2001); (iii) perception verbs have a hybrid eventive/stative nature (Dowty 1979): like statives, and unlike regular eventives, they allow an on-going reading with the present tense. This on-going reading is responsible for the implicative inference. The absence of implicativity with general ability readings is due to the intervention of a genericity operator, associated both with the simple present and the imperfective, which introduces an additional layer of modality. Perception verbs however differ from statives, in that they still express a change of state, and thus can appear with ability *can*.

References

- Bhatt, R. (1999). Ability modals and their actuality entailments. *Proceedings of WCCFL 17*.
- Condoravdi, C. (2001). Temporal interpretation of modals-modals for the present and for the past. *The construction of meaning*.
- Dowty, D. R. (1979). Word meaning and Montague grammar: the semantics of verbs and times in generative semantics and in Montague's PTQ. Dordrecht, Holland: D. Dordrecht: Reidel.
- Hackl, M. (2001). Ability ascriptions, ms, MIT.
- Hacquard, V. (2006). *Aspects of modality*. PhD thesis, MIT.
- Hacquard, V. (2009). On the interaction of aspect and modal auxiliaries. *Linguistics and Philosophy* 32(3).
- Kratzer, A. (1991). Modality. *Semantics*.

Propositional Attitude Reports: the Syntax of Presupposition & Assertion

Kajsa Djärv, University of Pennsylvania

Introduction. Propositional attitude verbs (e.g. *say*, *believe*, *know*) are known to be selective about the types of constructions that may occur in their complements. Following Emonds (1970), Hooper and Thompson (1973) identified a set of constructions that, while typically confined to matrix clauses, are also possible under a restricted set of verbs, e.g. (1). Other so-called “Main Clause Phenomena” [MCS] include speaker-oriented adverbs, V-to-C movement [C-V2], scene-setting adverbs, and VP-preposing. The study of MCS has been centered around two problems: (a) identifying the types of lexical/semantic-pragmatic *contexts* that license MCS; and (b) properly characterizing the syntactic and interpretive properties associated with the MCS *themselves*.

Theoretical background. The received view, since H&T, is that the availability of MCS is positively correlated with *assertion*, and negatively correlated with *presupposition*. Broadly, there are two schools of thought: On **positive accounts** (Wechsler 1991; Truckenbrodt 2006; Wiklund et al. 2009; Wiklund 2010; Jensen and Christensen 2013; Julien 2009, 2015; Woods 2016a,b, a.o.), “assertive” verbs such as *say* and *believe* are taken to select or license clauses with an extended C-domain, endowed with features pertaining to Common Ground [CG] management (Bianchi and Frascarelli, 2009), such as Topic, Focus, and Illocutionary Force (à la Rizzi 1997; Speas and Tenny 2003). Topicalization, C-V2 etc. are then *triggered* by features in the C-domain. On **negative accounts** however, “presuppositional” verbs such as *doubt*, *accept*, *regret*, and *know* select clauses headed by some definite or nominal element (à la Kiparsky and Kiparsky 1970). The nominal/D-layer in the embedded clause then effectively *blocks* the derivation of different MCS (e.g. Haegeman and Ürögdi 2010; De Cuba and Ürögdi 2010; Haegeman 2014; Kastner 2015). Further theoretical consensus however, has been hard to reach. We identify three key reasons for this.

Problem 1. Assertion and presupposition are themselves complex and multifaceted concepts (e.g. Stalnaker 1974). What aspects of these notions are relevant to the syntax? While some authors take the relevant dimension to be speaker/attitude holder commitment to the embedded proposition (p), others point to p being discourse new information. Yet others take factivity to be relevant.

Problem 2. The empirical and theoretical status of (doxastic) factives: do they *in fact* permit MCS, and are they predicted to do so, given the semantic underpinning of the syntactic theory (e.g. Simons 2007)? Negative accounts claim that *all* factives disallow MCS, while positive accounts take at least the doxastic factives (e.g. *discover*) to allow MCS.

Problem 3. Evaluating apparent disagreements about *specific* MCS. For instance, Bianchi and Frascarelli (2009) give (2) to show that English topicalization is licensed under emotive factives, in direct contrast to (1b). However, these judgments are subtle and potentially context-sensitive. Apparently conflicting empirical claims of this type may simply be due to a failure to control properly for potential pragmatic confounds. Moreover, theories about the interpretive constraints on MCS are typically based on acceptability judgments/distributional data for MCS under a small set of verbs, taken to represent larger *semantic classes* (see Problem 1). However, it is far from clear what the reality of these classes are, and which verbs actually belong to which class. Are (2) and (1b) *in fact* contradictory judgments, or do they represent some (unknown) dimension of variation?

Summary, problems. Without *comparable* data from different MCS across different languages, which controls for contextual and lexical properties of the relevant sentences, it is difficult to falsify and evaluate competing theoretical accounts. For instance, the current state of the literature is compatible with negative accounts being correct, in theory, about MCS being blocked in “presuppositional contexts”, but mistaken in their empirical assumptions about the doxastic factives. However, it may equally be true that negative accounts are right, about English topicalization, while positive accounts are right, about German C-V2.

Current Study. This talk presents results from a large-scale cross-linguistic experimental study,

investigating the specific lexical and semantic-pragmatic constraints on four different MCS, across three languages. We collected judgments of acceptability and judgments of interpretation, for the same exact same 40 sentences. Each of the 40 critical items (and the 32 fillers and controls) consisted of a unique verb+lexical content combination, set in exactly the same discourse context (Tab. 1). The study manipulated the following independent variables: verb and verb-class, matrix negation, type of MCS [C-V2; Topicalization; Scene-setting Adverbs; Speech Act Adverbs; Unmarked controls], and language [English; Swedish; German] (Tab. 2). Each subject thus saw the same 40 critical sentences involving 20 positive and negative verbs from five purported lexical classes, argued (along with negation) to differ with respect to the licensing of MCS. For an objective measure of the pragmatic dimensions of interest, the 40 critical items were independently tested in the unmarked control version for: speaker commitment to p; attitude holder [AH] commitment to p; likelihood that p is discourse novel. All judgments were given on a 9-point Likert Scale with the end-points marked. 1,272 participants took part in the study. The z-scored data was analyzed using linear mixed-effects models, predicting the acceptability of the different MCS-variants from verb identity and class, plus the three pragmatic factors.

Summary of Main Results.

- Robust association across languages of each of the three discourse properties with different verb class/polarity conditions (Figure 1).
- Results from the two discourse conditions in English shows these are not very sensitive to the discourse context (Figure 1).
- Speaker belief is only *lexically* associated with factive verbs (as expected on a lexical view of factivity) (Figure 1).
- Predictions by attitude holder belief and discourse novelty hypotheses come apart for response predicates and emotive factives (Figure 1).
- E-V2 is predicted by discourse novelty status of p (its distribution looks like that posited for MCS by Hooper and Thompson 1973) (Figure 2).
- Other MCS are not distinguished in terms of any of the dimensions tested here (pragmatic or lexical factors); whatever their licensing conditions are, they do not distribute like MCS in the sense of H&T) (Figure 2).

Discussion. These results have important implications for the separate question of what a theory of MCS should look like. First, we find that MCS is a much more heterogeneous class than previously thought, both in terms of distribution and pragmatic licensing conditions. Second, these results allow us to falsify a number of popular theoretical claims, while strongly supporting others. They support the view that C-V2 is licensed by Discourse Novelty (as argued by Caplan and Djärv 2017 based on Swedish corpus data), but not related to the presence of a belief(p) or commitment-to-p context (à la Truckenbrodt 2006; Wiklund 2010; Julien 2015; Woods 2016b,a,b). Notably, the robust interaction of matrix negation and predicate type is evidence against the hypothesis that the availability of MCS is due to local lexical selection for a particular type of clause (contra e.g. Wiklund et al. 2009; Kastner 2015). Finally, while the results support the view that C-V2 is ruled out in contexts where p is discourse old, we find strong evidence that this does not track Factivity (contra Kastner 2015, and Haegeman and colleagues). (These results are still compatible with there being a common denominator for of each type of MCS investigated, such as “Common Ground management” (Bianchi and Frascarelli, 2009)). In the presentation, we also discuss the implications from the current study for the question of whether certain predicates select DP complement.

Selected References. Bianchi, Valentina, and Mara Frascarelli. 2009. Is topic a root phenomenon? Haegeman, Liliane. 2014. Locality and the distribution of Main Clause Phenomena. Haegeman, Liliane, and Barbara Ürögdi. 2010. Referential CPs and DP: An operator movement account. Hooper, Joan, and Sandra Thompson. 1973. On the applicability of root transformations. Julien, Marit. 2015. The force of V2 revisited. Kastner, Itamar. 2015. Factivity mirrors interpretation: The selectional requirements of presuppositional verbs. Kiparsky, Paul, and Carol Kiparsky. 1970. Fact. Truckenbrodt, Hubert. 2006. On the semantic motivation of syntactic verb movement to C in German. Wechsler, Stephen. 1991. Verb second and illocutionary force. V2 in Scandinavian that-clauses.

- (1) a. [This book]_i, Mary read. (English Topicalization)
 b. John {**thinks**/***regrets**} that [this book]_i, Mary read t_i.
 (Maki et al., 1999; Haegeman and Ürögdi, 2010; Haegeman, 2012; De Cuba, 2017; De Cuba and Ürögdi, 2010; Kastner, 2015)
- (2) I **am glad** that [this unrewarding job]_i, she has finally decided to give up t_i.

Tab. 1. Structure of experimental items:

- *Background*: Two friends, Jane and Sarah, run into each other. Jane says:
- *Target Sentence*: Guess what! I just talked to Mary, and **she said that Lisa lost her job!**
- *Questions to measure acceptability and interpretation of the embedded proposition*:
 - *Acceptability*: To me, this sentence sounds: Completely unnatural – Completely natural
 - *Discourse New*: It is likely – not likely that Jane and Sarah have previously talked about Lisa losing her job.
 - *Speaker Belief*: As far as Jane is concerned, Lisa lost her job. [No – Maybe – Yes]
 - *AH belief*: As far as Mary is concerned, Lisa lost her job. [No – Maybe – Yes]

Tab. 2. Independent variables:

- Verb Identity and Class (between-item)
 1. Speech Act: *say, mention, tell me, claim*
 2. Doxastic Non-factive: *believe, assume, reckon, guess*
 3. Response Stance (*accept, admit, doubt, deny*)
 4. Emotive Factive: *appreciate, resent, love, hate*
 5. Doxastic Factive: *discover, find out, notice, hear*
- Matrix Negation: Verb, ¬Verb (within-item)
- MCS and language (between-subject)
 1. C-V2 (Sw, Ger)
 2. Topicalization (Eng)
 3. Scene setting Adv (Sw, Eng, Ger)
 4. Speech Act Adv (Sw, Eng, Ger)
 5. Unmarked controls (Sw, Eng, Ger)
- Interpretation (between-subject)
 1. Speaker belief that p
 2. Attitude holder belief that p
 3. p as discourse new

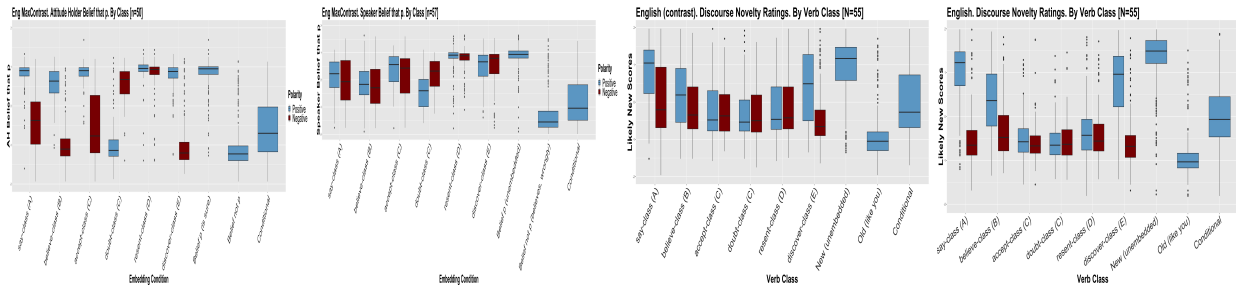


Figure 1: From left to right: Attitude Holder belief that p; Speaker belief that p; Likelihood that p is discourse new (MAXCONTRAST); Likelihood that p is discourse new (MAXNEW). Blue = Positive; Red = Negated.

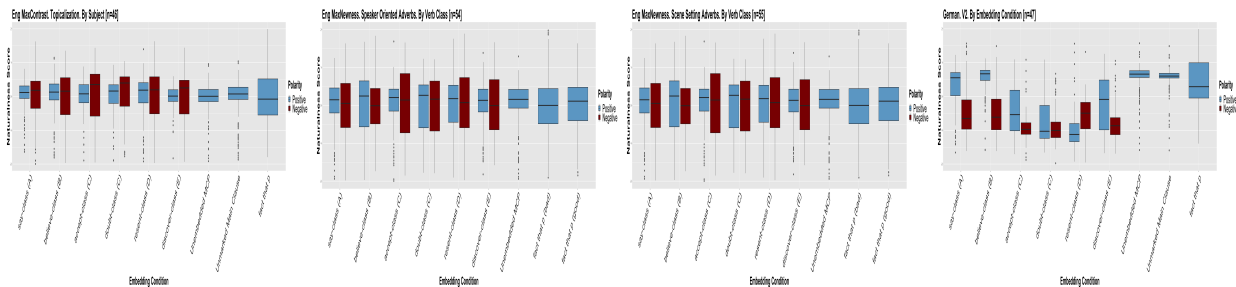


Figure 2: Four types of embedded MCS. From left to right: Topicalization; Speech Act Adverbs, Scene Setting Adverbs, Verb Second.

The Morphosemantics of Spanish Gender: A Case Study of Pseudo-Incorporation

Lucia Donatelli, Georgetown University
led66@georgetown.edu

Introduction. Spanish bare nouns (BNs) exhibit syntactic and semantic behaviors that parallel patterns noted across the literature for pseudo-incorporated nominal expressions (Espinal & McNally 2011; Dayal 2011). I present an analysis of Spanish BNs that unifies verbal and prepositional incorporation and that sheds light on gender assignment and interpretation processes in the language.

Data. BNs are the most syntactically restricted set of nominals in Spanish, appearing in some predicative structures (1a) and as the object of a reduced number of verbs (1b-e). In all cases, BNs must stay close to their verbal host and are quite restricted in how they can be modified. Semantically, Spanish BNs (i) have reduced discourse transparency and provide bad support for pronominal anaphora (1c); (ii) exhibit number neutrality (1d); (iii) have narrow scope; and (iv) occur in constructions that are institutionalized in some manner, often referred to as the “establishedness effect” (2).

Data from BNs suggests that not all gender is equal in its semantic weight. In Spanish, two types of gender are typically recognized: *grammatical gender*, a gender without semantically interpreted gender inferences, and *natural gender*, a gender of the animate entity in question with semantic inferences that more often than not correlates with the gender specification visible on the noun (Kramer 2015). All nouns trigger gender agreement with items such as determiners and adjectives.

Incorporation-like behaviors seem to disappear when the BN in question is morphologically complex, as can be seen with a natural/interpretable gender feature or a diminutive suffix (2). Nevertheless, natural gender features are acceptable in “institutionalized” (3b) or predicative (1a) constructions. The contrast in interpretation with apparent “interpretable” gender features additionally depends on the noun in question. In (3a), while the sentence is felicitous in the context of Inés finding a male or female secretary or nurse, (3b) is only felicitous in the context of a female spouse or partner.

Analysis: Pseudo-Incorporation. I propose that V-N incorporation in Spanish is primarily syntactic, and a THEME argument that denotes a property restricts, rather than saturates, the predicate (Chung & Ladusaw 2004). Verbs that participate in V-N incorporation possess a HAVE/possession subcomponent (Harley 2004), which I analyze as a prepositional element that raises and adjoins to V and allows N to incorporate in contextually appropriate situations.

This analysis is supported by data that Spanish exhibits incorporation-like behavior with prepositional complements of verbs (4), found in a variety of Río de la Plata Spanish. These constructions display similar patterns as for V-N constructions: they are institutionalized (4a); morphologically restricted to the type of preposition (4b); unable to be referred to pronominally; and cannot be modified with gradable adjectives or telic predicates.

I propose that prepositional incorporation like (4) occurs with unergative verbs and functions as an event-kind classifier. This analysis draws on work by Myler (2013) on preposition incor-

poration for GOAL arguments; as well as work on non-canonical objects in Chinese (Barrie & Li, 2012; Zhang, 2018) that analyzes non-THEME incorporation structures as event kind-classifying elements. Spanish as a language thus displays properties typical of both pseudo-noun incorporation, when V-P-N incorporation is possible with verbal HAVE subcomponents and THEME arguments; and of non-canonical objects, when only partial incorporation occurs in a V-P-N structure due to the object being an oblique argument.

Analysis: Gender. Though incorporation in Spanish at first glance seems only to occur with grammatical/uninterpretable gendered nouns that are morphologically and semantically less complex, (3) refutes this. The data suggests that a more refined analysis of the gender assignment processes and the semantic contribution of natural gender in Spanish is necessary. For constructions like (3a), I propose that feminine $u[+FEM]$ gender is uninterpretable, contrary to the animate status of the noun itself. For constructions like (3b), feminine interpretable gender $i[+FEM]$ is permitted as a result of the noun's root denotation, which denotes a counterpart to a set. For constructions like (1a), the syntactic structure allows the subject to value the gender features of the predicate, which originate as uninterpretable. This analysis is further supported by data from ellipsis, in which nouns with uninterpretable gender features license ellipsis constructions (5a-c), while those with interpretable gender features do not (5d).

To explain the restriction on the incorporation of nouns with interpretable gender features, I propose that a feature LATT, which denotes semantic number, is present on these nouns (Heycock & Zamparelli 2005). Natural feminine $i[+FEM]$ nouns possess a $[-LATT]$ feature, which makes them semantically singular and unable to incorporate as a property-denoting noun, unless licensed in a unique configuration such as (3b) where a noun's root denotation necessitates an interpretable counterpart. This stands in contrast to bare, uninterpretable nouns that do not possess such a feature, enabling them to modify the verbs they incorporate to without denoting a specific entity.

References.

- (1) Barrie, M., & Li, Y. H. A. 2012. Noun incorporation and non-canonical objects. In *West Coast Conference on Formal Linguistics* (Vol. 30).
- (2) Chung, S., & Ladusaw, W. A. 2003. *Restriction and saturation* (Vol. 42). MIT press.
- (3) Dayal, V. 2011. Hindi pseudo-incorporation. *Natural Language & Linguistic Theory*, 29(1), 123-167.
- (4) Espinal, M. T., & McNally, L. 2011. Bare nominals and incorporating verbs in Spanish and Catalan. *Journal of Linguistics*, 47(1), 87-128.
- (5) Harley, H. (2004). Wanting, having, and getting: A note on Fodor and Lepore 1998. *Linguistic Inquiry*, 35(2), 255-267.
- (6) Kuguel, I., & Oggiani, C. 2016. La interpretación de sintagmas preposicionales escuetos introducidos por la preposición en. *Cuadernos de Lingüística de El Colegio de México*, 3(2), 5-34.
- (7) Kramer, R. 2015. *The morphosyntax of gender* (Vol. 58). Oxford University Press.
- (8) Myler, N. 2013. On coming the pub in the North West of England: Accusative unaccusatives, dependent case, and preposition incorporation. *The Journal of Comparative Germanic Linguistics*, 16(2-3), 189-207.
- (9) Sudo, Y., & Spathas, G. 2016. Natural gender and interpretation in Greek. Ms., UCL and Universität Stuttgart/Humboldt Universität zu Berlin.
- (10) Zhang, N. N. 2018. Non-canonical objects as event kind-classifying elements. *Natural Language & Linguistic Theory*, 36(4), 1395-1437.

- (1) a. Amalia es [médica / abogada / profesora / bombera].
Amalia is [doctor.F.SG / lawyer.F.SG / professor.F.SG / firefighter.F.SG]
'Amalia is (a) doctor / lawyer / professor / nurse / firefighter.'
- b. María tiene [coche/ casa en la playa/ tarjeta de crédito/ etc.].
María has [car.M.SG/ house.F.SG at the beach/ card.F.SG of credit/ etc.]
'María has (a) [car/ house at the beach/ card of credit/ etc.]'
- c. Hoy lleva falda. Se #la regalamos el año pasado.
today wear skirt.SG.F. her it.SG.F gave the year last
'Today she is wearing (a) skirt. We gave it to her last year.'
- d. Busco piso. [Uno en Barcelona./ Uno en Barcelona y uno en Girona.]
I-look-for flat.SG.M. [one.SG.M in Barcelona./ one.SG.M in Barcelona and one.SG.M in Girona.]
'I'm looking for a flat. [One in Barcelona./One in Barcelona and one in Girona.]'
- e. No encuentro película que me guste.
Not find film.SG.F that me please
'I cannot find a film to my taste.'
- (2) Elena tiene [perro / *perra / *perrito]
Elena has dog.M.SG / dog.F.SG / dog.DIM.M.SG
'Elena has (a) dog / *(female) dog / *little dog.'
- (3) a. Inés busca [secretaria / enfermera]
Inés looks-for secretary.F.SG / nurse.F.SG
'Inés is looking for (a) secretary / nurse.'
- b. Carlos busca [esposa / novia]
Carlos looks-for wife.F.SG / girlfriend.F.SG
'Carlos is looking for (a) wife / girlfriend.'
- (4) a. Los alumnos suelen estudiar en biblioteca.
the student tend-to study in library.F.SG
'The students tend to study in (the) library.'
- b. Siempre lleva la ropa fina a/ de/ *hacia/ *desde tintorería.
always wear the clothing nice to/ from/ towards/ from dry.cleaner.F.SG
'[(S)he] always brings nice clothing to/ from/ towards/ from the dry cleaner's.'
- (5) a. Pablo es doctor y Marta también.
Pablo is doctor.M.SG and Marta also
'Pablo is a doctor, and Marta is too.'
- b. (?)Marta es doctora y Pablo también.
Marta is doctor.F.SG and Marta also
'Marta is a doctor, and Pablo is too.'
- c. Pablo es actor y Marta también.
Pablo is actor.M.SG and Marta also
'Pablo is an actor, and Marta is too.'
- d. *Marta es actriz y Pablo también.
Marta is actress.F.SG and Pablo also
'Marta is an actress, and Pablo is too.'

Uses of Oddball Imperatives

Imperative clause type is distinct from declarative and interrogative clause type in many languages (Aikhenvald 2010). A well-known feature of imperatives is that they have a range of uses available to them (Portner 2007, Condoravdi and Lauer 2012, Kaufmann 2012). In addition to their canonical usage as commands, imperatives can typically participate additionally as wishes, permission, and advice.

- (1) Get out of here! (Command)
- (2) Please don't rain tomorrow! (Wish)
- (3) Feel free to take a cookie. (Permission)
- (4) Take the interstate north for 2 hours. (Advice)

Dubbed the *problem of quantificational homogeneity* by Kaufmann 2012, this range of meanings is generally taken to be a central clue in determining the proper denotation for imperatives. This problem gains an additional complication when looking at several types of *oddball* imperatives. There are certain constructions that can be used for commanding, and sometimes the range of uses available to imperatives, without obviously being part of imperative clause type. For example, *general prohibitives* (Donovan 2018) have the full range of uses available to them that imperatives do.

- (5) No smoking in here! (Command)
- (6) Please, no raining tomorrow! (Wish)
- (7) Fine, no cookies for you then. (Concession)
- (8) No rubbing the infected area. (Advice)

A reasonable conclusion based on this is that both imperative and general prohibitive derive their semantics from the same modal. If the semantics of that modal provides the range of meanings available to imperatives, then it is expected that both expressions will have the same range of meanings available to them.

On the hand, Goal-Oriented Location Commands (GOLCs), seem to be limited in their usage relative to imperative and general prohibitives. GOLCs are limited to directive uses only.

- (9) Feet on the floor! = You must put your feet on the floor. (Command)
- (10) #Package in the mailbox! = I hope the package is in the mailbox. (Wish)
- (11) #Coats in the cabinet! = You are permitted to put your coats in the cabinet (Permission)

Why would different imperative-like constructions have a different range of meanings available to them? Following the same logic used for imperative and general prohibitives, the lack of available readings for GOLCs seems to indicate that they are derived from a different source than imperatives. Given that they are restricted to command usage, it seems straightforward that their meaning is derived directly from necessity. However, it has been suggested that in the case of morphological imperatives, their meaning is also derived from deontic necessity (Kaufmann 2012). We are therefore left with a puzzle. If the meanings of standard imperatives and non-canonical imperatives are both derived from the same source, they should have the same range of meanings available to them. If the meanings of standard imperatives and non-canonical imperatives are derived from different sources, why do they pattern similarly to imperatives in so many respects (for example, both constructions are addressee-oriented)?

Selected References: • Aikhenvald, Alexandra Y. 2010. Imperatives and commands. Oxford: Oxford University Press. • Condoravdi, C. and S. Lauer. 2012. Imperatives: Meaning and illocutionary force. *Empirical issues in syntax and semantics* 9, 37–58. • Donovan, M. 2018. General Prohibition: The Deletion of Allow-Predicates by an Imperative Modal in English. *PLC Working Papers* 41. • Kaufmann, Magdalena. 2012. Interpreting Imperatives. Dordrecht: Springer. • Portner, P. (2007). Imperatives and modals. *Natural Language Semantics*, 15(4), 351–383

Puzzle. Bošković and Gajewski (2011) claim that SerBo-Croatian (SC) does not have neg-raising (NR), and provide an example where a strong NPI ‘at least two years’ is not licensed under a negated instance of the NR verb *mislim* (‘think’). In this paper I show that, although the verb ‘think’ blocks long distance licensing of strong NPIs (1), ‘want’ does not (2). I propose that the cause of this asymmetry lies in the differences in non-truth-conditional meaning of the attitude verbs ‘think’ and ‘want’ in SC.

Theoretical Background. I adopt the approach to NR that was assumed by Gajewski (2007) (originally from Bartsch, 1973). Gajewski accounts for the fact that sentence (3a) has the reading given in (3b) through an Excluded Middle presupposition (EM). Namely, ‘think’ in sentence (3a) triggers the presupposition that the attitude holder is opinionated (OPN) with respect to its complement. A formal representation is given in (4) and (5).

I follow the standard assumption that NPIs require a downward-entailing (DE) environment (Ladusaw 1979 and many since). For strong NPIs, I take Gajewski's (2011) position that the licensing environment must be DE when both its at-issue and non-at-issue inferences are taken into account. (In this respect I depart from Zwarts' 1998 generalization that strong NPIs require anti-additive environment in order to be licensed.)

NR verbs, then, have the Excluded Middle inference as a non-at-issue meaning. In (6), I show that when they are negated, the total of their at-issue and non-at-issue inferences make their complement DE, thus making that complement a licensing environment for strong NPIs.

Proposal. Gajewski's (2011) account, however, is insufficient to distinguish between the behavior of ‘think’ and ‘want’ in SC, as is desired given the data. The question I will try to answer to is what could be the reason for this difference in the behavior of the two attitude verbs in SC. To do so I suggest two additions to Gajewski's (2011) account: sensitivity to anti-presuppositions and a weaker version of Condoravdi's (2002) diversity condition as a presupposition of ‘want’.

The argument of the verb ‘want’ can be a proposition that is false as in (8b), but it is odd if its argument is true, as in (8a). This illustrates that Condoravdi's (2002) diversity condition is not completely symmetric in this respect. I propose that $\llbracket \text{want} \rrbracket$ presupposes that its argument is not necessary in the attitude holder's belief state. Together with the excluded middle presupposition, the total inferences of $\neg \text{want } p$ will be downward entailing. This preserves the licensing of strong NPIs under negated ‘want’. In (9) we see how this approach works with ‘want’.

I consider the possibility that the anti-presupposition of the verb ‘think’ is responsible for the weakening of the effect of the strong NPI licensing. Gajewski's (2011) shows that strong NPIs are sensitive to implicatures and presuppositions of the attitude verb. To account for the data from SC, I add that strong NPIs are also sensitive to anti-presuppositions of ‘think’, as well. Namely, ‘think’ anti-presupposes that the speaker deems the argument of ‘think’ possible.

The question is whether the conjunction of the assertion and the anti presupposition entails the conjunction of EM-strengthened meaning of the stronger version of the assertion and its anti-presupposition and is therefore DE (9). If so, Gajewski's condition holds. The answer is no.

There is no entailment and the environment is not DE, and therefore the strong NPI is not predicted to be licensed in this environment.

(1) *Ne mislim da je izašla iz zemlje **najmanje dve godine**.

NEG **think**.1st.PRES that is leave from country **at least two years**

I don't think she has left the country in at least two years.

(2) Ne želim da izade iz zemlje **najmanje dve godine**.

NEG **want**.1st.PRES that leave from country **at least two years**

I don't want her to leave the country in at least two years.

(3)

a. Bill doesn't think that Mary is here.

b. Bill thinks that Mary is not here.

(Gajewski 2007)

(4)

a. $\llbracket \text{believe} \rrbracket = [\lambda p . \lambda x : \text{OPN}(p)(x) . \text{BEL}_x \subseteq p]$

b. $\text{OPN}(p)(x) = 1$ iff $\text{BEL}_x \subseteq p \vee \text{BEL}_x \subseteq \neg p$

(5)

a. $\llbracket \text{Bill doesn't think that Mary is here} \rrbracket$ is defined only if $\text{OPN}(\llbracket \text{Mary is here} \rrbracket)(\llbracket \text{Bill} \rrbracket)$

b. $\llbracket \text{Bill doesn't think that Mary is here} \rrbracket = 1$ iff $\text{BEL}_{\text{Bill}} \not\subseteq \llbracket \text{Mary is here} \rrbracket$

c. $(5a)+(5b) \rightarrow \llbracket \text{Bill doesn't think that Mary is here} \rrbracket = 1$ only if $\text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is here} \rrbracket$

(6)

a. Bill doesn't think that Mary is a student.

b. Bill doesn't think that Mary is a linguistics student.

(7)

a. $\llbracket \text{Bill doesn't think that Mary is a student} \rrbracket = 1$ only if $\text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket$

b. $\llbracket \text{Bill doesn't think that Mary is a linguistics student} \rrbracket = 1$ only if $\text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

c. $\text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket \Rightarrow \text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

d. Bill doesn't think that Mary has been a student in months.

(8)

a. #Bobby wants this year to be 2019.

b. Bobby wants this year to be 2020.

(9)

a. $\llbracket \text{Bill doesn't want Mary to be a student} \rrbracket$ presupposes & asserts:

$\diamond \neg \llbracket \text{Mary is a student} \rrbracket \& \text{BOUL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket$

b. $\llbracket \text{Bill doesn't want Mary to be a linguistics student} \rrbracket$ presupposes & asserts:

$\diamond \neg \llbracket \text{Mary is a linguistics student} \rrbracket \& \text{BOUL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

c. $\diamond \neg \llbracket \text{Mary is a student} \rrbracket \& \text{BOUL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket \Rightarrow$

$\diamond \neg \llbracket \text{Mary is a linguistics student} \rrbracket \& \text{BOUL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

(10)

a. $\llbracket \text{Bill doesn't think that Mary is a student} \rrbracket$ presupposes & asserts:

$\diamond \llbracket \text{Mary is a student} \rrbracket \& \text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket$

b. $\llbracket \text{Bill doesn't think that Mary is a linguistics student} \rrbracket$ presupposes & asserts:

$\diamond \llbracket \text{Mary is a linguistics student} \rrbracket \& \text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

c. $\diamond \llbracket \text{Mary is a student} \rrbracket \& \text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a student} \rrbracket \Rightarrow$

$\diamond \llbracket \text{Mary is a linguistics student} \rrbracket \& \text{BEL}_{\text{Bill}} \subseteq \neg \llbracket \text{Mary is a linguistics student} \rrbracket$

References

- Bartsch, R. (1973). “Negative Transportation” gibt es nicht. *Linguistische Berichte*, 27.
- Bošković, Željko, and Gajewski, Jon. (2011). ‘Semantic correlates of the NP/DP parameter’. *Proceedings of the North East Linguistic Society 39*. GLSA, University of Massachusetts, Amherst.
- Condoravdi, Cleo. (2002). Temporal interpretation of modals: Modals for the present and the past. *The construction of meaning*, ed. by David Beaver, Luis Casillas Martinez, Brady Clark and Stefan Kaufmann, 59-88. Stanford, CA: CSLI Publications.
- Gajewski, Jon. (2007). ‘Neg-Raising and Polarity’. *Linguistics and Philosophy* 30(3)
- Gajewski, Jon. (2011). ‘Licensing Strong NPIs’. *Natural Language Semantics*. 19(2)
- van der Wouden, Ton. (1995). A problem with the semantics of neg-raising. Ms.
- Zwarts, Frans. (1998). Three types of polarity. In F. Hamm, & E. Hinrichs (Eds.), *Plural quantification* (pp. 177–238). Dordrecht, Kluwer.

An Experimental Investigation of Mood Variation in Spanish Emotive-factive Clauses

Introduction:

In this talk I will be discussing the pragmatic uses of mood variation in Spanish emotive-factive complement clauses. I will show both naturally occurring data as well as judgements obtained from an experiment, that indicate that informational quality (if the information is ‘new’ or ‘old’) influences a speaker’s choice of mood.

Background:

Factive predicates, such as *know* and *regret*, presuppose the veridicality of their embedded clauses. There are two kinds: those that are neutral (e.g. *know*; *remember*), and those which are evaluative (e.g. *sad (that)*, *happy (that)*, *interesting (that)* etc.) (Portner, 2018). Whereas neutral factives are generally considered ‘indicative governors’ (since their default selection tends to be the indicative), emotive-factives are far less uniform in their preferences for mood; they may take the indicative (Greek, Romanian, Bulgarian, Hungarian), subjunctive (French, Italian, Spanish), or both moods (Catalan, Brazilian Portuguese, Turkish) (Giannakidou, 2015; Portner, 2018; Quer, 1998, 2009).

Mood Alternation in Spanish:

It is generally thought that the subjunctive is the required mood in the complements of Spanish emotive-factives (Giannakidou, 2015; Gili Gaya, 1960; Manteca Alonso-Cortés, 1981; Terrell and Hooper, 1974). It has, however, been noted by some that it can also be accepting of the indicative (Blake, 1982; Crespo del Río, 2014; Farkas, 1992a; García and Terrell, 1977; Gregory and Lunn, 2012; Lipski, 1978; Quer, 1998, 2009; Silva-Corvalán, 1994; Terrell and Hooper, 1974). It has furthermore been observed that the Spanish mood alternation has an effect on interpretation. In particular, it relates to a difference between assertion (indicative) and non-assertion (subjunctive) (Borrego et al., 1986; Bosque, 1990; Collentine, 2010; Gregory and Lunn, 2012; Lavandera, 1983; Majías-Bikandi, 1994; Portner, 2018; Quer, 2009; Sessarego, 2016; Terrell and Hooper, 1974), or new information (indicative) vs. familiar information (subjunctive) (Lunn, 1989; Gregory and Lunn, 2012). However, the only examples that have been provided of indicative complements under emotive-factive verbs are the following constructed minimal pairs from Terrell and Hooper (1974):

- (1) a. *Es bueno que Ud. **llega**-PRESENT-INDIC a tiempo.*
b. *Es bueno que Ud. **llegue**-PRESENT-SUBJ a tiempo.*
‘It’s good that you arrive on time.’
- (2) a. *Me sorprendió que **vino**-PAST-INDIC.*
b. *Me sorprendió que **viniera**-PAST-SUBJ.*
‘It surprised me that you came’.

Naturally Occurring Data:

I obtained the following samples of data from *El Corpus del Español* (The Corpus of Spanish). They demonstrate that use of the indicative is influenced by how new the information is to the hearer/reader. Example (3) discusses a new brand and the products it advertises to blog readers who had no knowledge of the brand's inception. Example (4) details a mother sharing her concerns about her baby's eating habits to readers who were neither acquainted with her nor with her situation.

- (3) *¡Chiquillas! ¡Estoy demasiado contenta de poder presentarles por primera vez en el blog a la marca COE! ¡Me encanta que todos los productos de la marca **vienen** con un sticker que indica el olor y el estado de ánimo que genera!*

Girls! I am too happy to be able to introduce you to a brand called 'COE' for the first time in this blog. I am happy that all of the brand's products **come-INDIC** with a sticker that indicates the smell and mood that it generates.

- (4) *Hola, mi bebe tiene 7 meses, está bien en el peso y el tamaño para su edad, pero me preocupa que no le **agrada** mucho la comida, todavía toma leche materna.*

Hello, my baby is 7 months old and his weight and size are good for his age, but I'm worried that food does not **please-INDIC** him very much, he's still drinking breast milk.

Experiment:

Nineteen native speakers (NSs) of Spanish from Latin America and Spain performed 2 Acceptability Judgment Tasks (AJT). The first was contextualized, while the other contained stand-alone evaluative sentences. The contextualized AJT contained various evaluative expressions, obtained and adapted from *El Corpus del Español*. Instances of both the indicative and subjunctive were included, with each mood preceded by contexts that were created to signal if the evaluated information would have been 'new' or 'old' to the recipient in question. The context-free AJT also included both moods, but with no information about the context in which the speaker would have said them.

Examples:

Contextualized AJT: English Translation of Sample Item Showing the Indicative as Used With 'Old' Information

Comment extracted from a newspaper article titled: "All that you need to know about the iPod nano 5. This comment is for readers of a newspaper who are familiar with the changes that *Apple* is planning to implement:

Original Sentence: *Como hemos comentado al principio del artículo, es asombroso que Apple **ha conseguido** meter aún más tecnología y prestaciones en el iPod nano sin cambiar sus dimensiones.*

‘As we mentioned at the beginning of the article, it’s amazing that Apple has-INDIC decided to add even more technology and features to the iPod nano without having to change its dimensions.’

‘This sounds: a. very good b. acceptable c. odd d. unacceptable’

Context-free AJT:

(6) *Es malo que no le conozcan porque es un tipo fenomenal.*

‘It is bad that you all don’t know-SUBJ him because he’s a phenomenal guy.’

‘This sounds: a. very good b. acceptable c. odd d. unacceptable’

Results and Conclusions:

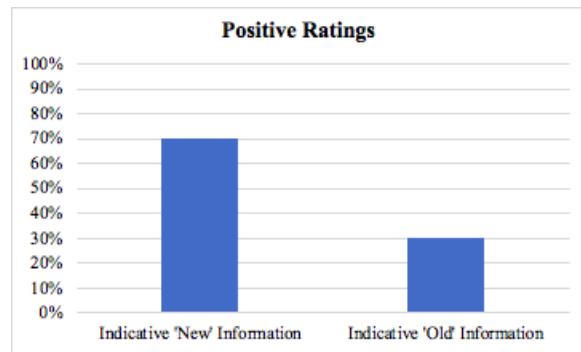


Figure 1

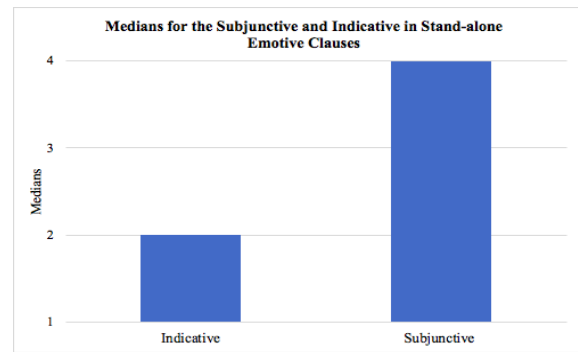


Figure 2

Median Scale for Figure 2:

Very good – 4 Acceptable – 3 Odd – 2 Unacceptable – 1

Mood variation in Spanish emotive-factive clauses is meaningful although this is not recognized in descriptive or pedagogical grammar. Since emotive-factives tend to relay information that is already known to a listener, the default subjunctive mood (the ‘un-assertive’ mood) is regularly used since the evaluated content does not need to be highlighted. When the information is new, the acceptability of the indicative increases, since it, as the more assertive mood, is used to call the reader/hearer’s attention to the evaluated content. As seen in the figures above, the acceptability of the indicative decreases when the information is old, or when the reader/hearer had no preceding context to prompt its usage.

Agential Free Choice

Melissa Fusco

MACSIM8 @ NYU Invited Speaker

The Free Choice effect—whereby $\Diamond(p \text{ or } q)$ seems to entail both $\Diamond p$ and $\Diamond q$ —has long been described as a phenomenon affecting the deontic modal “may”. This talk presents an extension of the semantic account of deontic free choice defended in Fusco 2015 to the agentive modal “can”, the “can” which, intuitively, describes an agent’s powers.

I begin by sketching a model of inexact ability, which grounds a modal approach to agency (Belnap & Perloff, 1998; Belnap et al., 2001) in a Williamson (1992, 2014)-style margin of error. A classical propositional semantics combined with this framework can reflect the intuitions highlighted by Kenny (1976)’s much-discussed dartboard cases, as well as the counterexamples to simple conditional views recently discussed by Mandelkern et al. (2017). I substitute for classical disjunction an independently motivated generalization of Boolean join—one which makes the two diagonally, but not generally, equivalent—and show how it extends free choice inferences into a simple object language.

References

- Belnap, Nuel, and Michael Perloff (1998). “Seeing to it that: a canonical form for agentives.” *Theoria*, 54: pp. 175–199.
- Belnap, Nuel, Michael Perloff, and Ming Xu (2001). *Facing the Future: Agents and Choices in Our Indeterminist World*. Oxford University Press.
- Fusco, Melissa (2015). “Deontic Modality and the Semantics of Choice.” *Philosophers’ Imprint*, 15(28): pp. 1–27.
- Kenny, Anthony (1976). “Human Abilities and Dynamic Modalities.” In Manninen, and Tuomela (eds.) *Essays on Explanation and Understanding*, Dordrecht Reidel.
- Mandelkern, Matthew, Ginger Schultheis, and David Boylan (2017). “Agentive Modals.” *Philosophical Review*, 126(3): pp. 301–343.
- Williamson, Timothy (1992). “Inexact Knowledge.” *Mind*, 101(402): pp. 217–242. Williamson, Timothy (2014). “Very Improbable Knowing.” *Erkenntnis*, 79: pp. 971–999.

Probing children's early comprehension of comparative constructions

Megan Gotowski and Kristen Syrett

INTRODUCTION: Previous studies of children's comprehension of comparatives (*Alex is taller than Joe (is)*) have reported variable interpretation patterns (see e.g., Townsend 1974; Donaldson & Wales, 1970; Bishop & Bourne 1985; Gathercole, 1985; Gor & Syrett 2015; Arie, Syrett, & Goro 2017; a.o.). Up till now, there has not been a coherent account of the source of these non-adult-like interpretations across tasks, since the results hint at different possibilities attributed to either an immature grammar or to immature processing. However, they might also indicate children's appeal to a cross-linguistically licensed interpretation assigned to the English form, if we assume Universal Grammar delimits the space of possibilities available to a young child acquiring these constructions. In this research, we probe children's interpretation of comparative constructions, with a goal of ruling out key interpretational variants as the source of apparent non-adult-like interpretations of the English comparative construction. We argue that while many children actually *do* demonstrate adult-like comprehension, others resort to interpretations arising from the degree constructions within the language they are acquiring.

HYPOTHESIS SPACE: Let us propose that the hypothesis space is constrained by non-English interpretations, an appeal to a similar non-comparative structure in their own language, incremental processing of the comparative. We can therefore generate a set of hypotheses about non-target interpretations, given the English surface structure.

H₁: Children will re-interpret the comparative along the lines of an A-not-A analysis, and create polarity partitions (see Schwarzschild, 2008), as in (1)—a strategy observed in languages like Hixkaryana (2), and languages that are claimed to lack degrees, as in Motu (3).

(1) Alex is tall and Joe is not tall.

(2) Kaw-ohra naha Waraka, kaw naha Kaywerye Stassen (1985)
tall- NOT he.is Waraka, tall he-is Kaywerye
'Kaywerye is taller than Waraka.'

(3) Mary na lata, to Frank na kwadogi Beck et al. (2009); Hohaus et al. (2014)
Mary TOP tall, but Frank TOP short
'Mary is taller than Frank.'

H₂: Children may misanalyse the standard-marker *than* as signaling conjunction and assume a coordination analysis, along the lines of what has been proposed for mis-analysis of relative clauses (Tavakolian 1981) (although see Syrett & Lidz 2009). See (4).

(4) Alex is tall and Joe is tall.

H₃: Either as a result of immature processing or an inability to appeal to an explicitly designated contextual standard, children may 'ignore' the standard phrase, interpreting only the matrix clause with the positive gradable adjective, as in (5).

(5) Alex is tall.

EXPERIMENT 1: 43 children (age 3-5) and a control set of adults participated. Across target trials, participants were presented with a set of objects corresponding to four GAs (*big, long, full, bumpy*). They were first asked to CATEGORIZE them, but placing the 'A ones' on a red felt rectangle to the left and the 'other ones' on a blue felt rectangle to the right. They then heard a series of questions in which they were asked to COMPARE the objects. The prompts had the format, *Is X A-er than Y?* but objects were pre-selected from the 'A' or 'other' (i.e., 'not A')

group, based on four target conditions (multiple trials each across target GAs), which in turn allowed us to make clear predictions about response patterns, as summarized in Table 1 below. Control items involved categorization based on object kind.

Table 1. Predicted response patterns to *Is X A-er than Y?* questions in Exp. 1, given four approaches

	$X = A, Y = A$	$X = \neg A, Y = \neg A$	$X = A, Y = \neg A$	Reverse
H ₀ : 'X is Aer than Y'	Yes	Yes	Yes	No
H ₁ : 'A-not-A'	No	No	Yes	No
H ₂ : 'A and A'	Yes	No	No	No
H ₃ : 'X is A'	Yes	No	Yes	No

Overall % correct responses to each category, based on the top row of Table 1 were 83%, 84%, 77%, and 93%, respectively, although these percentages do not capture response patterns. We thus categorized each individual participant based on their response pattern, as above. Adults patterns just as anticipated. 29 of the 43 children displayed an adult-like 'X is Aer than Y' pattern of interpretation. No child appeared to consistently appeal to either H₂ or H₃, and only one was consistently H₄. The remainder of the 13 children displayed an inconsistent pattern across items (n=8), or else appear to have resorted to other degree constructions in English (e.g., an equative).

EXPERIMENT 2: 20 children (age 4-5) and a control set of adults participated. Participants were shown a series of Powerpoint slides, each displaying two animals, each of which had two sets of objects (see Figure 1). A puppet delivered a target comparative statement about each scene, first in prediction mode without the images, then again once the images were displayed. In Block 1, the comparative featured one subject, while in Block 2, it featured two subjects. Participants either accepted or rejected the statement, occasionally providing justifications.

Figure 1: Sample trial slides and target statements for Exp. 2



The lion has more baseballs than strawberries. (F) The tiger has more cookies than the panda has apples. (F)
 Only 12 of the 20 children provided adult-like responses, and only one consistently adopted H₄. Seven others displayed a response pattern that diverged from the predicted hypothesis space, in which it appears they were comparing within or across object kinds for each animal (e.g., comparing the cardinality of object 1 and object 2 for each animal). Surprisingly, the % correct was higher for the 'different subject' block (70% v. 82%), presumably because children thought they should incorporate the entire scene, and interpreted the elided subject in the standard without an identity constraint (i.e., The lion has more baseballs than the alligator...).

CONCLUSIONS: The results indicate that children do not uniformly adopt a consistent analysis of the English comparative. While many have mastered it by age 4-5, others appeal to a variety of interpretive strategies. Crucially, they are *not* consistently adopting a non-degree-based partition analysis, nor are they appealing to a conjunction analysis that reflects an immature grammar or an inability to incorporate the standard. Instead, their non-adult-like strategies seem to reflect an appeal to other degree constructions in the language they are acquiring or an appeal to aspects of the scene in the resolution of the ellipsis inherent to comparatives.

An Extended Minimal Networks Theory for Backtracking Counterfactuals

Ioana Grosu (New York University)

Background. I propose an account for backtracking in counterfactuals which improves upon predictions from the Minimal Networks Theory given in Hiddleston (2005).

The way in which Minimal Networks Theory is described in Hiddleston (2005) and Rips (2010) only directly accounts for certain cases of backtracking counterfactuals. The simplest cases are the ones in which there is exactly one minimally altered network under consideration. For these cases, the Minimal Networks Theory predicts that the truth of the counterfactual is dependent on two aspects. The first is the truth of the antecedent. The second is the similarity of the counterfactual world to the actual world. Since Minimal Networks Theory calculates similarity based on the sets of causal breaks and intact variables, the truth of the counterfactual is therefore dependent on whether the world given the antecedent and consequent is minimally altered (i.e., has a minimal set of breaks and a maximal set of intact variables), when compared to other worlds in which the antecedent is true. If only one minimally altered network is generated, then the counterfactual which generates that network is judged to be true.

Issues arise for Minimal Networks Theory in cases where there is no unique minimally altered network under consideration. That is, cases where multiple minimally altered networks are being generated, or no obvious minimally altered network is generated. For these, Hiddleston, and consequently Rips and Edwards (2013) (who provides an extension to Minimal Networks Theory based on conditional probabilities), claim that a consequent is true only if it is true in all possible minimally altered networks in which the antecedent is also true. While this works for some cases in which multiple minimally altered networks are generated, it is a very coarse grained analysis that fails in cases where context-dependent factors differentiate between possible minimally altered networks. Furthermore, it provides no account when there is no obvious minimally altered network generated.

Aim. Minimal Networks Theory allows for use of context dependent factors, but does not explicitly define the way in which these context dependent factors are taken into account. My aim is to introduce an extension to Minimal Networks Theory which allows for an evaluation of counterfactuals not only based on causal breaks and intact variables, but also on other factors that influence people's judgments of the status of variables within the system. I crucially rely on the notion of mutability, which is informally defined as the ease with which alternatives to the given fact come to mind. I build off of accounts such as those in Dehghani et al. (2012) and Lucas and Kemp (2015). I extend the definitions of causal breaks and intact variables in order to incorporate the effects of fact mutability.

Proposed Extension. My proposed extension rectifies the fact that Hiddleston does not provide an explicit way to incorporate context-dependence into his work. I consider counterfactual scenarios such as the following (adapted from Kahneman et al. (1982), discussion from Kahneman and Miller (1986)). This example highlights the fact that routines are treated differently from exceptions, and that exceptions are more mutable than routines.

“On the day of the accident, Mr. Jones left his office at the regular time. He sometimes left early to take care of home chores at his wife’s request, but this was not necessary on that day. Mr. Jones did not drive home by his regular route. That day was exceptionally clear and Mr. Jones told his friends at the office that he would drive along the shore to enjoy the view. As commonly happens in such situations, the Jones family and their friends often thought and often said “If only...” during the days that followed the accident. How did they continue that thought? Please write one or more likely completions.”

In this case, where the route is the exceptional variable and the time is routine, participants overwhelmingly chose to alter the route, as opposed to the time. When presented with a similar version of the scenario, in which the time was the exceptional variable instead of the route, participants chose to alter the time. As it is presented in Hiddleston (2005), Minimal Networks Theory cannot account for these findings. The sets of causal breaks are equally minimal for the legal causal networks in both cases, and the sets of intact variables are equally maximal.

Therefore, Minimal Networks Theory predicts that participants would not exhibit a preference between the counterfactual case in which Mr. Jones changes his route, and the counterfactual case in which Mr. Jones changes the time of departure. For both of these cases Minimal Networks Theory predicts that participants would choose to write that *either* a changed time of departure or a changed route could have been the case. This contrasts with reported judgments from Kahneman et al. (1982), which show that (a) is the case.

- a. If Mr. Jones had lived, he would have driven along his regular route.
- b. If Mr. Jones had lived, he would have left work early.

By incorporating context dependent factors, however, the difference in judgments between the two scenarios (i.e., the scenario in which the route is the exception and the scenario in which the time is the exception) can be explained.

I propose the following extension to the definitions of causal breaks and intact variables. All causal breaks that do not correspond to maximally mutable facts are relevant causal breaks and all intact variables that do not correspond to minimally mutable facts are relevant intact variables. I define maximal mutability on the basis of the mutable features applicable to a fact. A maximally mutable fact is a fact described using the maximal set of mutable features, when compared to other facts in a given model. I define a fact as the value of a non-antecedent variable at the actual world. In addition, I propose a relative ranking between relevant breaks and relevant intact variables. These extensions are based on the idea that, when deciding on which minimal network to take into account, people will primarily take into consideration salient changes to the network.

By incorporating these extensions to the theory proposed in Hiddleston (2005), I can now account for cases such as the one in Kahneman et al. (1982), as well as for some other problem cases for Minimal Networks Theory in which more than one minimally altered network is generated.

References

- Dehghani, M., Iliev, R., and Kaufmann, S. (2012). Causal explanation and fact mutability in counterfactual reasoning. *Mind & Language*, 27(1):55–85.
- Hiddleston, E. (2005). A causal theory of counterfactuals. *Noûs*, 39(4):632–657.
- Kahneman, D. and Miller, D. T. (1986). Norm theory: Comparing reality to its alternatives. *Psychological review*, 93(2):136.
- Kahneman, D., Slovic, S. P., Slovic, P., and Tversky, A. (1982). *Judgment under uncertainty: Heuristics and biases*. Cambridge university press.
- Lucas, C. G. and Kemp, C. (2015). An improved probabilistic account of counterfactual reasoning. *Psychological review*, 122(4):700.
- Rips, L. J. (2010). Two causal theories of counterfactual conditionals. *Cognitive science*, 34(2):175–221.
- Rips, L. J. and Edwards, B. J. (2013). Inference and explanation in counterfactual reasoning. *Cognitive Science*, 37(6):1107–1135.

Learning to map modals to meanings: an elicited production study on force and flavor

The same modal words can be used to express possibility and necessity in different “flavors”. This creates a complex one-to-many mapping from word-form to meaning for language learners. For example, in (1), *must* is used to express a goal-oriented (teleological) necessity, while in (2), *must* is used to express a knowledge-based (epistemic) necessity. Furthermore, different modals can express the same flavor in a many-to-one mapping relationship (e.g., *have to* or *should* express the same meaning as *must* in (1)).

- (1) Kat must/has to/should go down the pink path
→ given her goal to get to the bakery
- (2) Nick must be hiding in the pink box
→ given our knowledge that the only other box is empty

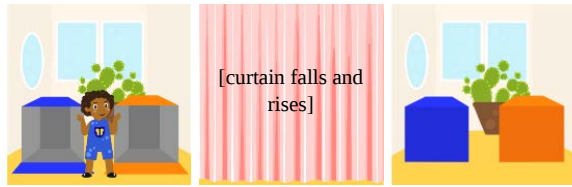
Children acquiring modals must learn to (a) map the same modal to different meanings, and (b) map the same meaning to different modals. How do children work out the complex relations involved in this modal space?

Corpus studies show that children use modals to express non-epistemic “flavors” before epistemic ones (e.g., Wells, 1979). Comprehension studies show that children are likely to accept possibility modals in necessity contexts, and sometimes also necessity modals in possibility contexts, unlike adults (Noveck, 2001; Ozturk & Papafragou, 2014). For example, children accept descriptions like (2) in situations where there is more than one possible hiding location. This behavior could be due to pragmatic or conceptual difficulty, as argued for in the literature, but it could also be due to children not having yet figured out the force of the modals tested, nor the range of flavors it can express. Knowing which modals children prefer to produce in different carefully-controlled contexts will help us better understand which meanings children use a modal to express.

Building on Cournane, (2014), we used a sentence-repair task to elicit modals in teleological and epistemic necessity and possibility contexts. Children are trained to repeat pre-recorded sentences to a shy snail puppet (who cannot hear the stories because he’s hiding in his shell), in which a glitch has replaced some words with white noise, and ‘repair’ the sentences by filling in the missing word. In test sentences, the white noise occurs where an adult would use a modal, which allows children to supply their own modal to fit the controlled context.

We crossed three factors: force (necessity vs. possibility), “flavor” (teleological vs. epistemic) and age (children: (n=46, (24 female)), ages 3;0;13 to 5;4;15, mean=4;1;29 adults: (n=24, (18 female)), ages 18 to 28, mean=21). Each participant saw four trials in each condition, and participants either saw all of the teleological trials first, or all of the epistemic trials first, with force pseudo-randomized throughout the experiment (participants either saw necessity trials first or possibility trials first). Sample contexts are given in (3a-b).

(3) a. Epistemic-possibility context:



This time, there's a blue box and an orange box, so there are two hiding spots...

Where's Nick? One spot is in the orange box...
...look! **Or, Nick *glitch* be hiding in the blue box**

b. Teleological-necessity context:



Now, Kat is going to the balloon store to get balloons! There are two ways to get to the balloon store. One way is to go down the brown path...

...but, uh oh! It's blocked! **So, Kat *glitch* go down the pink path**

Results: While 99% of adult responses were modals, only about 36% of child responses were. The rest of the time, children prefer to repair the sentence by adding tense to *go* or *be*, or (less frequently) by repeating *go* or *be* in their bare form, but without filling in the glitch (4).

(4) **Teleological:** {I think} Kat {goes/went/is going/go} down the yellow path

Epistemic: {I'm pretty sure/I guess/maybe} Nick {is/was/be} hiding in the yellow box
{maybe/I guess}

Figure 1. Which modals do **adults** prefer?

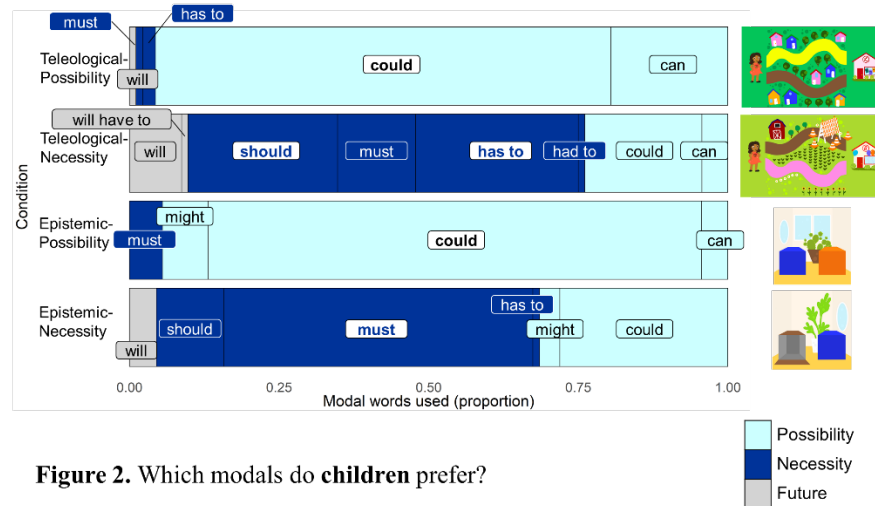
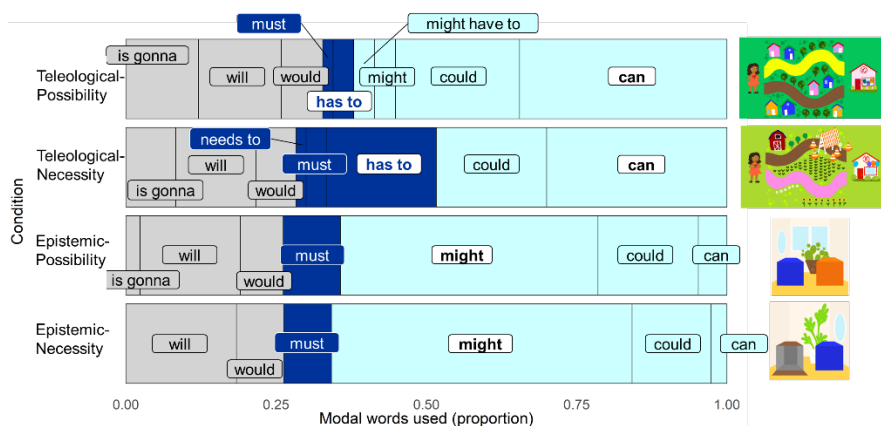


Figure 2. Which modals do **children** prefer?



Of the target modal responses, children prefer to use *can* in both teleological necessity (where adults prefer *have to* and *should*) and possibility contexts (where adults prefer *could*), and *might* in both epistemic necessity (where adults prefer *must*) and possibility contexts (where adults prefer *could*). Unlike adults, children use *must* and *could* the same amount in each “flavor” context. Unlike children, adults very rarely used *might*.

On the force dimension, children do not use different modals in epistemic possibility and necessity contexts, but they do in teleological necessity and possibility contexts: children are more likely to use *have to* to express teleological necessity than they are to express teleological possibility. However, this pattern only surfaces after aggregating child responses. Individually, if children used a necessity modal in necessity contexts, then they also likely used that modal in possibility contexts, echoing their non-discriminatory behavior reported in comprehension studies. Our findings suggest that children ages 3-4 years-old may not yet have robust modal words to describe epistemic and teleological possibility and necessity (in particular, epistemic necessity).

Selected References

- Cournane, A. (2014) In Search of L1 Evidence for Diachronic Reanalysis: Mapping Modal Verbs, *Language Acquisition*, 21:1, 103-117, DOI: 10.1080/10489223.2013.855218
- Cournane, A. (2015). Revisiting the epistemic gap: Evidence for a grammatical source. In *Proceedings of the 39th annual Boston University conference on language development*.
- Cournane, A. (2015). Modal development. PhD. Thesis. UToronto.
- Kuczaj & Maratsos (1975). What children can say before they will. *Merrill-Palmer Quarterly of Behavior and Development* 21, 89–111.
- Noveck, I. A. (2001). When children are more logical than adults: Experimental investigations of scalar implicature. *Cognition*, 78(2), 165-188.
- Ozturk, O., & Papafragou, A. (2015). The acquisition of epistemic modality: From semantic meaning to pragmatic interpretation. *Language Learning and Development*, 11(3), 191-214.
- Wells (1979). Learning and using the auxiliary verb in English. In Lee (ed.), *Language Development*, 250-270. Croom Helm.

Midpoints and Endpoints in Event Cognition

Yue Ji, Anna Papafragou University of Delaware

Events unfold over time, i.e., they have a beginning and an endpoint. Previous studies have illustrated the importance of endpoints for event perception and memory (Lakusta & Landau, 2005; Papafragou, 2010; Strickland & Keil, 2011; Zacks & Swallow, 2007). However, this work has not compared endpoints to other potentially salient points in the internal temporal profile of events (e.g., midpoints) and has only discussed events with a self-evident endpoint. Across a broader range of events, some have a natural, inherent endpoint and are temporally *bounded* (e.g., crack an egg into a bowl, eat a sandwich), while others are unspecified about when they come to an end and are temporally *unbounded* (e.g., stir an egg in a bowl, eat cheerios). This distinction is systematically encoded in language (van Hout, de Swart, & Verkuy, 2005) and can be characterized with the property of homogeneity, i.e., *bounded* events have a non-homogeneous internal structure leading to a “culmination” (Parsons, 1990) while *unbounded* events have a homogeneous internal structure (Krifka, 1989). In the present study, we compared event endpoints with midpoints in both bounded and unbounded events and hypothesized that the differences in internal event structure should affect how viewers process and weigh temporal slices of different events.

We created videos showing bounded and unbounded events and then introduced brief interruptions which took up one-fifth of the total video duration (range: 0.8-2.4s) to block either the temporal midpoints or endpoints of the events. The experiment adopted a variant of the “picky puppet task” (Waxman & Gelman, 1986), where participants were invited to watch a couple of videos and were told that the girl in the videos was very picky: she liked some of her videos but not the others. The task was to figure out what kind of videos the picky girl liked. In the training phase, participants were presented with 8 pairs of videos; in each pair, the two videos showed the same event but differed in the placement of the interruption (see Figure 1). After each video, participants heard either “The girl likes the video” or “The girl doesn’t like the video”, depending on the girl’s preference for either interruptions blocking event middles (i.e., mid-interruptions) or interruptions blocking event endpoints (i.e., end-interruptions). Participants were randomly assigned to either Bounded or Unbounded condition, depending on which type of events they were exposed to throughout the experiment. At test, participants watched 8 new videos with a mid- or an end-interruption and decided whether the picky girl would like them or not. The results revealed a significant interaction between the girl’s preference (Likes mid-interruption vs. Likes end-interruption) and event type (Bounded events vs. Unbounded events): participants who watched videos of bounded events had better performance when the picky girl liked mid-interruptions rather than end-interruptions, but participants exposed to videos of unbounded events performed equally well in identifying preferences for either type of interruptions ($F(1, 116) = 6.26, p = .014$; see Figure 2). This suggests that for bounded events, blocking the endpoint was more disturbing (and hence less acceptable as a “preference”) compared to blocking the midpoint but for unbounded events, interruptions at the two time points were treated largely identically.

Our findings provide new evidence for the importance of event endpoints, extending previous literature that compared endpoints and start points in motion events (e.g., Lakusta & Landau, 2005; Papafragou, 2010; Regier & Zheng, 2007). More importantly, the results demonstrate that the salience of endpoints depends on event boundedness, i.e., endpoints are weighted over midpoints only in bounded events that have a finely differentiated internal structure but not in unbounded events that have a homogeneous structure.



Figure 1. Examples of a training trial for a bounded event (folding up a handkerchief) that includes the two versions of the event: (a) mid-interruption (actor in yellow shirt), (b) end-interruption (actor in blue shirt).

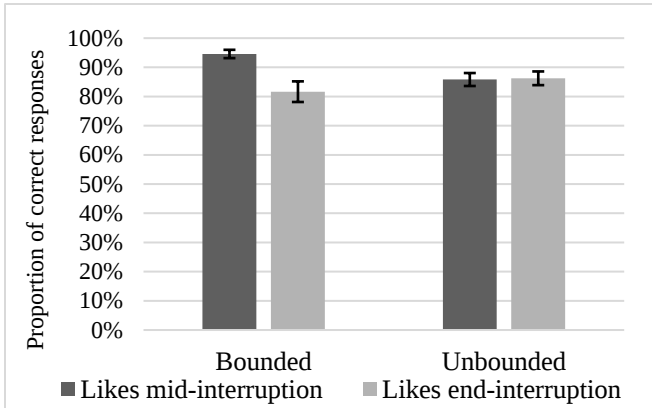


Figure 2. Proportion of correct responses. Error bars represent standard error.

References

Krifka, M. (1989). Nominal reference, temporal constitution and quantification in event semantics. In R. Bartsch, J. van Benthem, & P. van Emde Boas (Eds.), *Semantics and Contextual Expression, Groningen-Amsterdam Studies in Semantics 11*. Dordrecht: Foris Publications.

Lakusta, L., & Landau, B. (2005). Starting at the end: the importance of goals in spatial language. *Cognition*, 96, 1-33.

Papafragou, A. (2010). Source-goal asymmetries in motion representation: Implications for language production and comprehension. *Cognitive Science*, 34, 1064-1092.

Regier, T., & Zheng, M. (2007). Attention to endpoints: A cross-linguistic constraint on spatial meaning. *Cognitive Science*, 31, 705-719.

Strickland, B., & Keil, F. (2011). Event completion: Event based inferences distorts memory in a matter of seconds. *Cognition*, 121, 409-415.

van Hout, A., de Swart, H., & Verkuyl, H. J., (2005). Introducing perspective on aspect. In H. J. Verkuyl, H. de Swart, & A. van Hout (Eds), *Perspectives on aspect* (pp. 1-17). Dordrecht: Springer.

Zacks, J., & Swallow, M. (2007). Event segmentation. *Current Directions in Psychological Science*, 16, 80-84.

Licensing pseudo-incorporation in Turkish

Jinwoo Jo and Bilge Palaz

University of Delaware

Turkish exhibits the phenomenon of so-called pseudo-incorporation (PI), in which a bare nominal shows the semantic properties of incorporated nouns without forming a single morphological unit with a verb. What is particularly interesting about Turkish is that PI appears to take place in a wider range of environments than it does in many other languages, which usually allow theme PI only. Specifically, Öztürk (2004, 2009) reports that in Turkish, not only theme but also agent may be pseudo-incorporated (PI-ed) into a verb as shown in (1) and (2), respectively.

- | | | | |
|-----|-------------------------|-----|-----------------------|
| (1) | Ali kitap oku-du. | (2) | Ali-yi arı sok-tu. |
| | Ali book read-PST | | Ali-ACC bee sting-PST |
| | ‘Ali did book-reading.’ | | ‘Ali got bee-stung.’ |

In addition to theme and agent PI, we report that Turkish also allows goal to be PI-ed into a verb as in (3), and that it even allows more than one nominal to undergo PI as exemplified in (4).

- | | |
|-----|---|
| (3) | Öğretmen hasta öğrenci-yi doktor-a yolla-dı. |
| | teacher sick student-ACC doctor-DAT send-PST.3SG |
| | ‘The teacher did to-doctor-sending the sick student.’ |
| (4) | Maç-ta çık-an olaylar yüzünden taraftar futbolcu döv-dü. |
| | game-loc happen-rel incidents due.to supporter footballer beat-pst |
| | ‘Supporter’s footballer-beating took place due to the incidents in the game.’ |

The nominals in question in (2–4) exhibit the typical properties of PI-ed nominals such as weak referential force, weak interpretation under ellipsis, name-worthiness, number neutrality, and obligatory narrow scope with respect to logical operators, etc., indicating that agent, goal, and multi-nominal PI are in fact possible in Turkish. See Appendix on page 3 for a couple of examples that suggest the existence of each of the non-canonical forms of PI in Turkish.

The PI facts in Turkish pose non-trivial problems for the traditional complementation approach to PI, which claims that PI takes place only between a verb and its complement. Assuming the θ -structure correlation and the X'-theoretic view of phrase structures, the cases of agent, goal, and multi-nominal PI can hardly be accounted for in any straightforward manner under the complementation approach. Moreover, the contrast between (5) and (6) below suggests that the PI-ed theme and agent in fact occupies distinct structural positions: VP-adjoining adverbs may appear before PI-ed theme (5), but they are banned from appearing before PI-ed agent (6).

- | | | | |
|-----|---------------------------------------|-----|-----------------------------------|
| (5) | Ali (güzel) şarkı söyle-di. | (6) | Ali-yi (*kötü) polis döv-dü. |
| | Ali (beautiful) song say-PST | | Ali-ACC (*bad) police beat-PST |
| | ‘Ali did song-singing (beautifully).’ | | ‘Ali got police-beaten (*badly).’ |

Under the complementation analysis, PI-ed nominals must occupy Compl,VP regardless of the θ -role they are associated with; accordingly, it is expected that a VP-adjoining adverb can appear before a PI-ed nominal regardless of the θ -role it is associated with, contrary to fact as illustrated above. The contrast in (5–6) can be given a straightforward account if PI-ed theme and agent do not occupy the same structural position. That is, if *şarkı* ‘song’ in (5) is at Compl,VP but *polis* ‘police’ in (6) at Spec,VoiceP, then the incompatibility of the VP-adjoining adverb in (6) can be easily ascribed to the inappropriate attachment site of the adverb: it is attached to the edge of VoiceP, not VP.

The problems of the complementation analysis reviewed above call for a new analysis of PI, at least for Turkish, in which the target nominals of PI are allowed to occupy the structural positions other than the complement of a verb. In this paper, we offer one such analysis based on the non-saturating mode of semantic composition, *Restrict*, proposed by Chung and Ladusaw (2004).

Specifically, we first assume that the PI interpretation is attained when a property-denoting bare nominal and a predicate are composed via *Restrict* as illustrated below.

- (7) a. $[_{VP} \text{ kitap oku}]$ b. $\text{Restrict}(\text{book}, \lambda x \lambda e[\text{read}(e, x)]) = \lambda x \lambda e[\text{book}(x) \ \& \ \text{read}(e, x)]$
 book read

With this assumption, we propose that *Restrict* is subject to an LF condition formalized in (8).

- (8) *Restrict* may apply between a property-denoting nominal N and a predicate P, only if P does not dominate any predicate Q such that Q is saturated by an entity-denoting argument.

According to (8), PI is expected to be possible, regardless of the θ -role that the target nominal is associated with, as long as the target predicate does not have a history of saturation in the previous steps of semantic composition. Crucially, in Turkish, an ACC-marked theme argument A-moves to Spec, VoiceP (Kelepir 2001). The movement operation creates an environment where the property-denoting nominal at Spec, VoiceP (agent; Kratzer 1996) or Spec, ApplP (goal; Marantz 1993) and its sister predicate, Voice' or Appl', can undergo PI, in that the theme argument is *extracted out of VP*, and thus the target predicate no longer dominates any predicate saturated by an entity-denoting nominal. The assumption needed for this analysis is that an A-trace, as a mere member of A-chain, does not have the ability to saturate a predicate. What saturates a predicate is an entire A-chain, which is in line with the common view that an A-chain as a whole forms a semantic argument of a predicate, not a member of it. Turning to multi-nominal PI, note in (7) that the predicate stays unsaturated when it is composed with a nominal via *Restrict*. What this means under the view of (8) is that *Restrict* may apply more than once in a single derivation, for the previous application of *Restrict* does not saturate its target predicate so that the target predicate of the later application of *Restrict* has no history of saturation. Hence, the possibility of multi-nominal PI.

The current approach not only accounts for the non-canonical forms of PI, but it also correctly rules out ungrammatical cases of them. For instance, although agent PI is possible in the unergative and the transitive, it is not allowed in the ditransitive. Under the current approach, this is because in the ditransitive, the goal argument saturates the predicate (Appl') below VoiceP, and consequently, the target predicate of agent PI (Voice') dominates the predicate Appl' saturated by a goal argument in violation of (8). Note that the impossibility of agent PI in the ditransitive is only when goal is interpreted as an individual argument. If goal also undergoes PI, agent PI becomes possible deriving the agent-goal PI construction exemplified in (9).

- (9) Suçlu-yu vatandaş polis-e ihbar et-ti.
 criminal-ACC citizen police-DAT report do-PST
 'The criminal got citizen-reported-to-police.'

The same holds for multi-nominal PI. Agent-theme PI in the transitive is possible in Turkish, but agent-theme PI to the exclusion of goal in the ditransitive is impossible because the saturation of Appl' by a goal argument bleeds PI of agent under (8).

To summarize, the possibility of non-canonical PI in Turkish is claimed to be due essentially to the existence of a particular syntactic operation in the language (i.e., extraction of theme). The view of the current paper may extend to the analysis of the cross-linguistic variation of PI, where the possibility of non-theme PI is determined according to the way in which the narrow syntax feeds LF in each language.

Appendix

(i) Agent PI

a. *Number neutrality*

Dün Ali-yi tekrar tekrar arı sok-tu.
yesterday Ali-ACC again again bee sting-PST
'Yesterday, Ali kept being bee-stung.' (Different bees could have stung Ali.)

b. *Obligatory narrow scope*

Ali-yi arı sok-ma-dı.
Ali-ACC bee sting-NEG-PST
Possible: 'It was not the case that Ali got stung by {a bee/bees}.' / **Impossible:** 'There {was a bee/were bees} that did not sting Ali.'

(ii) Goal PI

a. *Number neutrality*

Dün öğretmen Ali-yi tekrar tekrar doktor-a yolla-dı.
yesterday teacher Ali-ACC again again doctor-DAT send-PST
'Yesterday, the teacher kept doctor-sending Ali.' (Ali could have been sent to different doctors.)

b. *Obligatory narrow scope*

Öğretmen hasta öğrenciyi doktor-a yolla-ma-dı.
teacher sick student-acc doctor-dat send-neg-pst.3sg
Possible: 'It was not the case that the teacher sent the the sick student to {a doctor/doctors}.' / **Impossible:** 'There {was a doctor/were doctors} to whom the teacher did not send the sick student.'

(iii) Multi-nominal PI

a. *Number neutrality*

Taraftar tekrar tekrar futbolcu döv-dü.
supporter again again footballer beat-pst
'Supporter's footballer-beating took place again and again.' (Different supporters could have kept beating different footballers.)

b. *Obligatory narrow scope*

Taraftar futbolcu döv-me-di.
supporter footballer beat-neg-pst
Possible: 'It was not the case that {a supporter/supporters} beat {a footballer/footballers}.' / **Impossible:** 'There {was a supporter/were supporters} that did not beat {a footballer/footballers}.' , 'There {was a footballer/were footballers} that {was/were} not beaten by {a footballer/footballers}.' , 'There were {a supporter/supporters and a footballer/footballers} that were not involved in the beating event.'

Selected References Chung, Sandra, and William A. Ladusaw. 2004. *Restriction and Saturation*. Cambridge, MA: MIT Press. • Kelepir, Meltem. 2001. Topics in Turkish syntax. Doctoral dissertation, Massachusetts Institute of Technology. • Kratzer, Angelika. 1996. Severing the external argument from its verb. In *Phrase Structure and the Lexicon*, edited by John Rooryck and Laurie Zaring, 109–137. Dordrecht: Kluwer. • Marantz, Alec. 1993. Implications of asymmetries in double object constructions. In *Theoretical Aspects of Bantu Grammar*, edited by Sam A. Mchombo, 113–151. Stanford: CSLI. • Öztürk, Balkız. 2004. Case, referentiality and phrase structure. Doctoral dissertation, Harvard University. • Öztürk, Balkız. 2009. Incorporating agents. *Lingua* 119:334–358.

Scalar implicature development in 4- and 5-year-olds is supported by language and executive function networks

Alyssa Kampa, Kaja K. Jasińska, Anna Papafragou
University of Delaware

Conversational principles lead adults to interpret a sentence like “Some children eat broccoli” to mean not all children do (Grice, 1989). This inference, a “scalar implicature,” (SI) requires pragmatic reasoning about speaker intentions, which relies on complex linguistic and cognitive abilities. Neuroimaging studies show recruitment of cognitive (i.e. executive functions; EF) and social (i.e. Theory of Mind; ToM) networks during pragmatic reasoning in adults (van Ackeren et al., 2012; Champagne-Lavau & Stip, 2010); however, little is known about the neural systems underlying development of pragmatic and social reasoning before age 5 (Shetreet et al., 2014; Gweon et al., 2012). Pragmatic reasoning develops during preschool, with some studies reporting SI success as early as 3.5-years (Stiller et al., 2015), when cognitive (EF) and social (ToM) abilities are concurrently developing. However, other studies have reported SI failures in children as old as 9 (Noveck, 2001). Examining developing brain networks can provide insight into the trajectories of pragmatic development to inform current inconsistencies that behavioral results alone cannot address. We ask how cognitive, linguistic, and social abilities, supported by the maturation of specific neural systems, contribute to pragmatic development. Using functional Near Infrared Spectroscopy (fNIRS) neuroimaging, we investigate how neural systems that support SI development change between ages 4 and 5, and into adulthood.

Participants completed a behavioral battery (PPVT, DCCS, Sally-Anne task) and neuroimaging SI task (adapted from Kampa & Papafragou, 2017). In this referent selection task, participants saw a pair of photos (Fig. 1) which differed only in whether the observer had full visual access to the contents of the box (full-knowledge) or limited visual access (limited-knowledge). Participants were told that the girl would describe what she sees in one of the two boxes. Participants would hear either “I see a spoon and a bowl” (strong) or “I see a spoon” (weak); a pragmatic responder should link the strong statement to the full-knowledge speaker and the weak statement to the limited-knowledge speaker, in accordance with principles of informativeness. Imaging data were analyzed using NIRS-SPM-v4 (Jang et al., 2009).

Results indicate activation of linguistic and cognitive networks during SI derivation for children and adults. Adults (N=26) showed greater activation for the weak vs. strong condition (i.e. implicature derivation) in the left VLPFC, including Inferior Frontal Gyrus (LIFG), a region associated with higher-level linguistic and cognitive processing, and the right DLPFC, associated with EF. 4- and 5-year-old children (initial pilot: N=12) who succeeded on the task displayed similar patterns of activation in both regions, suggesting early recruitment of language and EF networks in pragmatic reasoning. Ongoing work will yield further insight into neural recruitment for SI in children with varying pragmatic reasoning ability.

These data provide new insight into the pragmatic development in 4- and 5-year-olds; by age 4, some children show adult-like recruitment of language and executive function networks during scalar implicature derivation. These findings not only address age inconsistencies in the behavioral literature, but also provide broader insights into the development of linguistic and cognitive neural systems in young children.

Figure 1. An example test trial from the SI task.

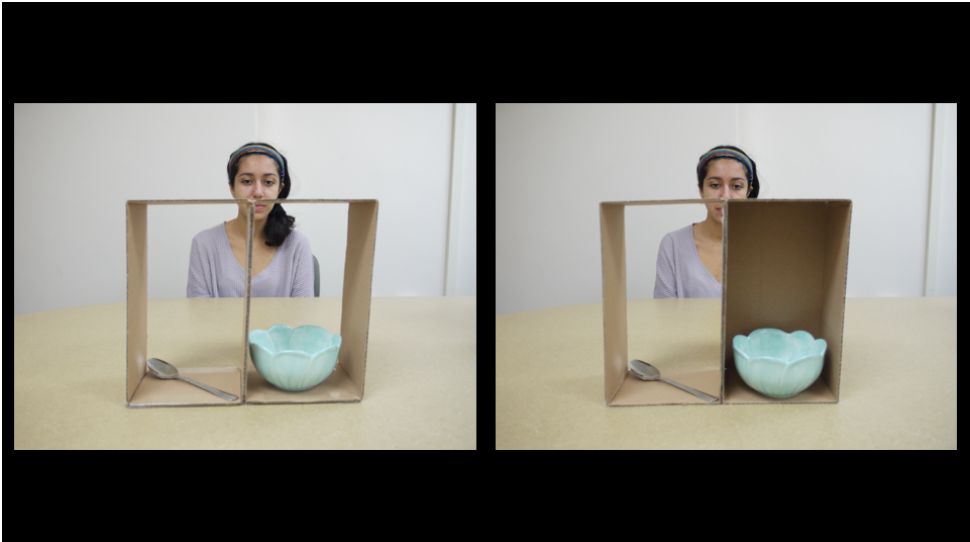
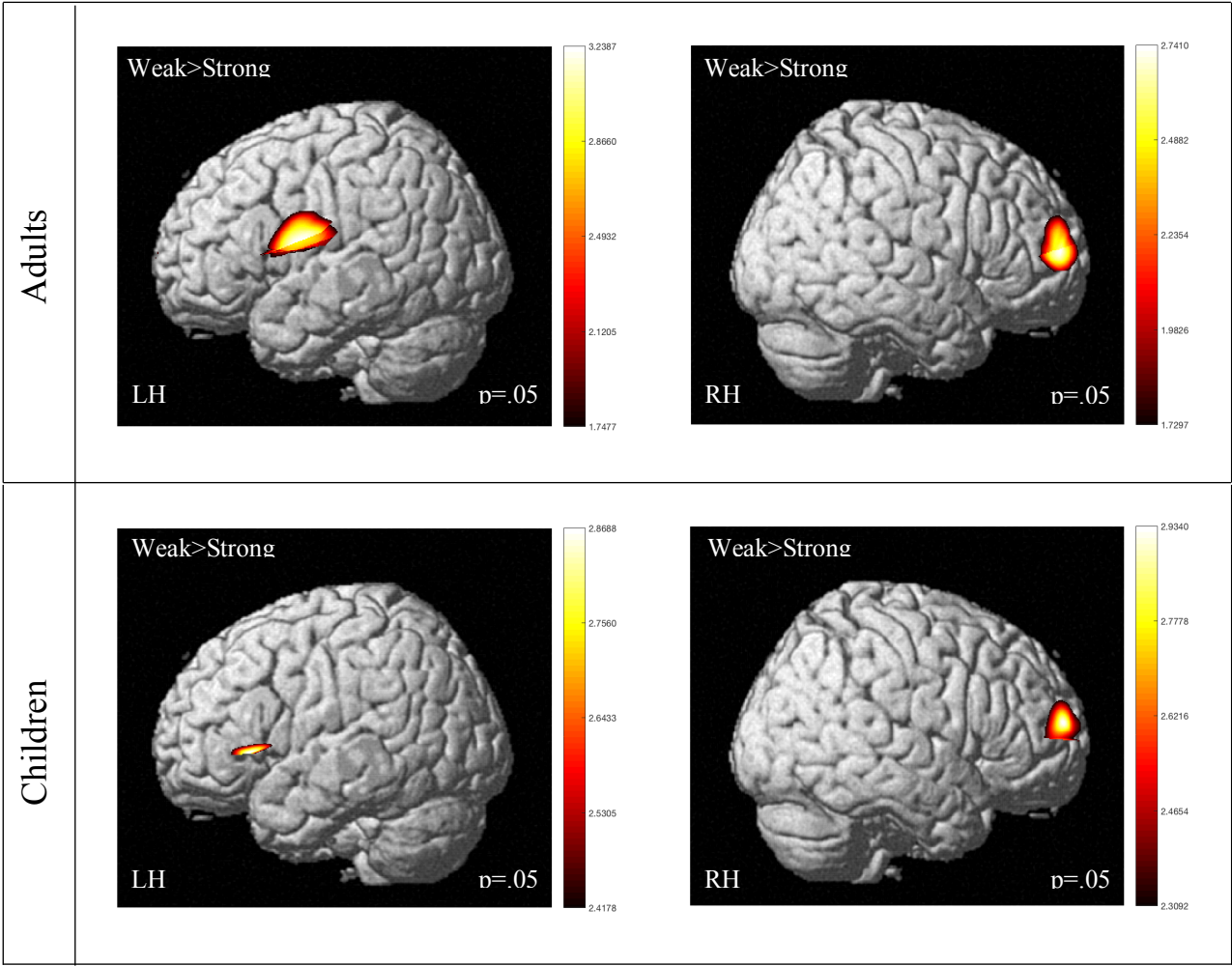


Figure 2. Neuroimaging results from the SI task.



Distributed neural encoding of binding to thematic roles

Matthias Lalisce

Problem Statement Much neurolinguistic research is directed towards validating and arbitrating between detailed linguistic or psycholinguistic theories. Parallel to this endeavor is a research programme, articulated by Poeppel et al. (2012), that aims to align the “parts lists” of linguistics and neuroscience by relating neural observables to the representations and computations proposed in linguistics and psycholinguistics. In this work, we pursue the alignment program by way of the *Structure Encoding Problem*: how is the formal property of structural sensitivity of linguistic representations realized in patterns of neural activation? The present work investigates the human brain’s solution to this problem for the encoding of propositions, where the distinct thematic role assignments of *cat* and *dog* in the propositions expressed by “the cat chased the dog” and “the dog chased the cat” distinguish their meanings. Reanalyzing fMRI neuroimaging data generated by the landmark work of Frankland and Greene (2015)—henceforth F&G—we contrast the localist solution to the Structure Encoding Problem considered by F&G with a distributed hypothesis derived from theoretical work in AI (Smolensky, 1990).

Background F&G localized a pair of regions of interest (ROIs) in the superior temporal sulcus (STS) that selectively encode the identities of agents and patients in a sentence. In ROI-A (the “agent” region), it is possible to decode the identity of an agent above chance, but not the patient, while in ROI-P (the “patient” region), the opposite is true. Since these regions were found to be non-overlapping, F&G liken them to “the data registers of a computer”, which play a crucial role in the neural representation of proposition-level structures by explicitly storing the values of particular semantic roles.

This paper implements a stronger test of this claim than is provided in the original work by explicitly modeling the composition of neural patterns into propositional structures, comparing this to the case when composition is not modeled. Let $\langle X_a, Y_p \rangle$ denote a proposition with X as agent and Y as patient. We assume (1) and set out to arbitrate between (2) and (3):

- (1) **Superposition.** The neural encoding of the proposition $\langle X_a, Y_p \rangle$ is the vector sum (superposition) of the activation patterns encoding X_a and Y_p : $X_a + Y_p$
- (2) **Localist hypothesis.** The activation patterns for agent bindings X_a and patient bindings Y_p reside on disjoint sets of representational units.
- (3) **Distributed hypothesis.** The units supporting the activation patterns for agent bindings X_a and patient bindings Y_p are not disjoint.

(1) says that filler-role bindings are associated with points (vectors) in the neural state space, which are combined via pattern superposition—an explicit proposal for the neural realization of compositionality.

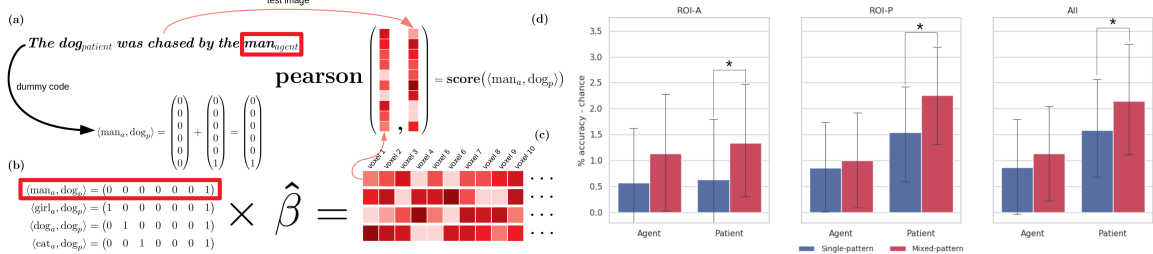


Figure 1: **Right** Visualization of the decoding procedure for the mixed-pattern model. To decode the agent, the patient is fixed to its true value and the predictions for the full proposition are compared to a held-out image. **Left** Graph of the results.

Method We consider the two ROIs localized by F&G, and evaluate their information content by attempting to decode the identity of agents and patients from the voxels

in those regions. The results reported in F&G are compatible with a superposition account of filler-role binding in a rather uninteresting sense, equivalently expressed as a kind of vector concatenation. However, a decoding methodology based on independent, single-role decoding is unable to arbitrate decisively between hypotheses (2) and (3). The data from each trial of the experiment are “mixed”, meaning that a trial with $\langle \text{man}_a, \text{cat}_p \rangle$ contains pattern components from both man_a and cat_p . The intuition behind our approach is the following. In an ROI that contains information only about agents, the patterns cat_p and girl_p will, on average, be identical. If the ROI is, in fact, sensitive to patients, then these patterns will vary systematically—but from the point of view of a model of agents where the patients are ignored, this systematic variation will appear to be noise. This could lead to poor decoding accuracy for, say, patients if the signal associated with agents is somewhat stronger—a false negative.

A stronger test of representational specificity, then, is to explicitly compare models that do and do not model pattern composition. To this end, we evaluate two classes of predictive models. In the **single-pattern** models, patterns for the fillers in the agent and patient roles are estimated in independent linear regressions, and are also independently compared with each held-out image in decoding. Our central manipulation is to fit **mixed-pattern** models that estimate filler patterns for both roles within each region. Then, when decoding experimental conditions from held-out trials, patterns are synthesized from the learned regression coefficients for *both* roles, modeling the entire proposition, rather than just one of its constituents.

Data 25 participants underwent fMRI while reading sentences like “the dog_a chased the cat_p ”—treated identically to “the cat_p was chased by the dog_a ”. Details in F&G’s paper.

Predictions If the mixed-pattern model does not perform better than the single-pattern model, we conclude that a region contains information about the contents of just one role. If it does perform better, we conclude that informative patterns about both roles are superposed within that region. There are four possible fillers, so chance is 25%.

Results In single-pattern decoding, we fail to replicate F&G’s finding that the ROI-A is role-selective, finding no difference between agents and patients. Furthermore, single-role decoding is not significant for either role in that region. On the other hand, single-role decoding is significant for both roles in ROI-P, providing *prima facie* evidence that the representations there are distributed over the same voxels. The key test is the comparison between the single- and mixed-pattern models. Here, we find that in both ROI-A and ROI-P, addition of information about the identity of the agent improves the accuracy of patient decoding, including when using all voxels (“All”).

Conclusions Formal reasoning about the nature of pattern superposition led us to detailed predictions about two possible representational architectures for representing multiconstituent structure neurally. Contrary to F&G’s original findings, our new analysis provides evidence that the Structure Encoding Problem is not solved in the manner of a classical computer, but rather by superposing distributed patterns of activation across overlapping sets of voxels, in line with hypothesis (3).

Frankland, S. M. and Greene, J. D. (2015). An architecture for encoding sentence meaning in left mid-superior temporal cortex. *Proceedings of the National Academy of Sciences*, 112(37):11732–11737.

Poeppel, D., Emmorey, K., Hickok, G., and Pytkkannen, L. (2012). Towards a new neurobiology of language. *The Journal of Neuroscience*, 32:14125–14131.

Smolensky, P. (1990). Tensor product variable binding and the representation of symbolic structures in connectionist networks. *Artificial Intelligence*, 46:159–216.

Wh-in-situ interrogatives through the lens of split wh-NPIs

A major disagreement on Korean (and Japanese) wh-interrogatives is whether movement is involved. On the movement approach (Huang 1982, Hagstrom 1998), the wh-item and the question particle must form a local relation at some point of the derivation, while on the in-situ approach (Shimoyama 2006, Beck 2006), a local relation is not necessary, requiring no movement. This paper argues for the latter by comparing wh-interrogatives with split wh-NPIs. I will show that wh-interrogatives and split wh-NPIs involve identical association mechanisms via ALTERNATIVE SEMANTIC COMPOSITION; however, distributional differences between these wh-constructions will be presented, which argues for the in-situ approach of wh-interrogatives.

Split wh-NPIs: Wh-NPIs in Korean consist of a wh-item and the focus particle *-to* meaning ‘also’ or ‘even’ and are argued to obey a clausemate condition (Hong 1995, Tieu and Kang, 2014). When the clausemate condition is violated (1a), moving *-to* closer to the licenser, creating a split wh-NPI, can improve the sentence (1b). (The wh-item and *-to* must belong to the same phonological phrase to get the NPI reading.)

- (1) a. * Mary-nun [**etten haksayng-to** chayk-ul ilk.ess.ta-ko] malha.ci **anh**-ass-ta.
 M-Top which student-Foc book-Acc read-Comp say Neg-Past-Decl
 ‘Mary didn’t say that any student read books.’
 b. Mary-nun [**etten haksayng-i** chayk-ul ilk.ess.ta-ko-**to**] mal.ci **anh**-ass-ta.
 M-Top which student-Nom book-Acc read-Comp-Foc say Neg-Past-Decl

Uniform LF behavior: It is well known that wh-interrogatives exhibit intervention effects that an intervener (e.g., a focus) between the wh-item and the question particle causes ALTERNATIVE SEMANTIC COMPOSITION failure (2a) (Shimoyama 2006, Beck 2006). The Intervention effects are also observed in split wh-NPIs as in (2b), suggesting that the association between the wh-item and *-to* is also achieved via ALTERNATIVE SEMANTIC COMPOSITION.

- (2) a. * Mary-man **nwukwu**-lul chohaha-**ni**?
 M-only who-Acc like-Q
 ‘Who does only Mary like?’
 b. * Mary-nun [Tom-man **etten chayk**-ul ilk.ess.ta-ko-**to**] malha.ci anh-ass-ta.
 M-Top T-only which book-Acc read-Comp-Foc say Neg-Past-Decl
 ‘Mary didn’t say that only Tom read any books.’

Non-uniform syntactic behavior: The association involved in split wh-NPIs is syntactically more constrained than wh-interrogatives. Split wh-NPIs, unlike wh-interrogatives (4), exhibit island effects (3).

- (3) a. * Mary-nun [sinbal-hako **etten os**]-**to** sa.ci anh-ass-ta.
 M-Top shoes-and which clothes-Foc buy Neg-Past-Decl
 ‘Mary didn’t buy shoes and any clothes.’
 b. * Mary-nun [**nwukwu**-ka ilccik o.ass-ki ttaymwune]-**to** hwana.ci anh-ass-ta.
 M-Top who-Nom early came-Nz because-Foc get.angry Neg-Past-Decl.
 ‘Mary wasn’t angry because anyone came early.’
 (4) a. Mary-nun [sinbal-hako **etten os**]-ul sa.ss-**ni**?
 M-Top shoes-and which clothes-Acc buy bought-Q
 ‘What is the clothing article x such that Mary bought shoes and x?’
 b. Mary-nun [**nwukwu**-ka nuckey oass-ki ttaymwune] hwana.ss-**ni**?
 M-Top who-Nom late came-Nz because-Foc got.angry-Q.
 ‘Who is the person x such that Mary got angry because x came late?’

The presence of island effects in (3) is unexpected given that association via ALTERNATIVE SEMANTIC COMPOSITION is claimed not to be sensitive to syntactic boundaries (Shimoyama 2006). I propose that the contrast in (3-4) is attributed to particle movement of *-to* in split wh-NPIs, and the absence of such movement in wh-interrogatives contra Hagstrom (1998). I argue that the wh-item and *-to* in split wh-NPIs are base-generated together, but the latter moves to the edge of a phase

to interact with the NPI licenser in a different spell-out domain following (Bošković's 2007). This analysis is supported by the fact that *-to* is adjacent to the wh-item when the licenser is in the same clause (5).

- (5) Mary-ka **nwukwu-to** manna.ci **anh-ass-ta**.
 M-Nom who-Foc meet Neg-Past-Decl
 'Mary didn't meet anyone'

Moreover, split wh-NPIs exhibit apparent locality condition between the wh-item and *-to*, unlike in wh-interrogatives, as illustrated in (6). This can be captured if *-to* is subject to a ban on superfluous steps (Chomsky, 1995), which is an economy constraint imposed on movement.

- (6) a. Mary-nun [Tom-i **etten chayk-ul** ilk-nun-ta-ko] malha.y-ss-**ni**?
 M-Top T-Nom which book-Acc read-Pres-Decl-Comp say-Past-Q
 'Which book did Mary say that Tom was reading?'
 b. Mary-nun [Tom-i **etten chayk-ul** ilk-nun-ta-ko] malha.ci-**to anh-ass-ta**.
 M-Top T-Nom which book-Acc read-Pres-Decl-Comp say-Foc Neg-Past-Decl
 ?* 'Mary didn't say that Tom read any books.'

Since the clausemate condition violation (1a) can be ameliorated when *-to* moves to the edge of the embedded clause (7), further movement closer to the NPI licenser as in (6b) will be unnecessary and violate the ban on the superfluous steps. If wh-interrogatives involved the same type of movement as split wh-NPIs, assuming that the question particle move to interact with the interrogative C (Hagstrom 1998), the ban on superfluous steps would have stopped the question particle from appearing at the edge of the sentence, where an interrogative C is situated.

Discussion: The representative movement analyses of wh-interrogatives are 1) covert wh-movement analysis (Huang 1982) and 2) particle movement analysis (Hagstrom 1998). The covert wh-movement analysis is not compatible with the existence of intervention effects (Pesetsky 2000), and the question particle movement analysis can be ruled out on the basis of comparison with split wh-NPIs. If wh-interrogatives involved particle movement just like split wh-NPIs, we will expect to see the same syntactic restrictions (e.g., sensitivity to islands and a ban on superfluous steps). *The distinct syntactic behaviors can easily be captured, if the wh-item and the question particle are associated non-locally without appeal to movement, as suggested by Shimoyama (2006), while in split wh-NPIs, an extra syntactic operation is involved: -to moves to interact with the NPI licenser.* A possible objection to comparing wh-interrogatives with wh-NPIs, as is done in this abstract, is that intervention effects in the former can be ameliorated by scrambling via short scrambling, while in the latter it cannot as in (7). Scrambled wh-items in wh-NPI constructions, however, always undergo reconstruction unlike in wh-interrogatives. Evidence for this comes from (??), where bleeding of Condition C averts the NPI reading.

- (7) * Mary-nun [**etten chayk-ul**_i Tom-**man** t_i ilk.ess.ta-ko-**to**] malha.ci anh-ass-ta.
 M-Top which book-Acc T-only read-Comp-Foc say Neg-Past-Decl
 'Mary didn't say that Tom read any books.' (No amelioration effect)
 (8) [Mary_i-ui **etten chinkwu-lul**]_j kunye_{i/j}-ka t_j cohaha.ci-to anh-nun-ta.
 M-Gen which friend-Acc she-Nom like-Foc Neg-Pres-Decl
 (i) 'She*_{i/j} does not like Mary_i's any friends.'
 (ii) 'As for Mary_i's some friend, She_i does not even like her.'

While short scrambling in Korean wh-interrogatives allows both bleeding effect of Condition C and the wh-interrogative reading, when (8) has the NPI reading and the bleeding effect cannot co-occur. Under the NPI reading, the c-commanded pronoun *kunye* cannot be coindexed with the R-expression *Mary* (8i). With the bleeding effect, the NPI reading is not available (8ii). The wh-item is interpreted as an indefinite and *-to* as 'even'. If reconstruction feeds Condition C, as Fox (1999) argues, Condition C effect in (8b) with the NPI reading suggests that the scrambled wh-item of split wh-NPIs requires reconstruction to have the NPI reading. Hence, the lack of the amelioration effect in (7) does not undermine the diagnosis of (7) as an intervention effect. If the analysis provided here is on the right track, it also provides new insights on the clausemate condition of wh-NPIs and supports the ALTERNATIVE SEMANTIC approach to the construction.

Making *wh*-phrases dynamic: A case study of Mandarin *wh*-conditionals

Introduction: This paper is a modest attempt to bring together two lines of research on *wh*-questions (*wh*-Qs) to shed light on Mandarin *wh*-conditionals. On one hand, many studies argue that short answers to *wh*-Qs, such (1), are not reducible to ellipsis and hence must be semantically represented (Groenendijk & Stokhof 1989; Jacobson 2016; Xiang 2016). On the other hand, Honcoop (1998) and Haida (2007) suggest that *wh*-phrases have dynamic discourse contributions in the sense of introducing discourse referents (drefs), as evidenced by cross-sentential binding (2). In this paper, I propose that the drefs introduced by a *wh*-phrase can be used to model the short answer to the corresponding *wh*-question. I then discuss how this proposal provides a novel analysis for Mandarin *wh*-conditionals (3), which are conditionals with **co-referring** *wh*-phrases showing up in the antecedent clause and the consequent clause (*jiu* is a conditional marker).

- (1) A: Who enters? (3) **Shéi** xiān jìnlái, **shéi** jiù xiān chī.
 B: Ahn. who first enter who then first eat
 (2) Who₁ won the game? What's his₁ score? 'Whoever enters first eats first.'

Non-interrogative uses of *wh*-phrases are generally taken to be indefinites. The obligatory co-reference of *who*'s in (3) is puzzling and violates the novelty condition of indefinites (Heim 1982).

Update with centering: Following Bittner (2014) and Murray (2010), I assume that a context *c* is a set of structured sequences *s* of drefs (cf. Dekker 1994). Specifically, $s := \langle \top, \perp \rangle$, in which $\top$ is the top sequence representing drefs in the center of attention, while $\perp$ is the bottom sequence representing drefs in the periphery of attention. Sentences denote context change potentials, i.e., functions from context to context. The table below lists some sample lexical items. Proper names can add drefs to $\top$ (when notated with $\uparrow$) or $\perp$. $\top_s + a$ is a shorthand for $\langle \top + a, \perp \rangle$ and $\perp_s + b$ for $\langle \top, \perp + b \rangle$, where $+$ is sequence extension. Proper names are modeled as generalized quantifiers (GQ). The denotation of *Ahn invites Bill* is composed as in (4).

items	denotation	(4)
Ahn [↑]	$\lambda P \lambda c. P(a)(\{\top_s + a \mid s \in c\})$	$\llbracket \text{Ahn invites Bill} \rrbracket =$
Bill	$\lambda P \lambda c. P(b)(\{\perp_s + b \mid s \in c\})$	Ahn [↑] $\lambda x. (\text{Bill } \lambda y. \text{invite}(y)(x)) =$
invite	$\lambda x \lambda y \lambda c. \{s \in c \mid \text{invite}(y)(x)\}$	$\lambda c. \{\langle \top + a, \perp + b \rangle \mid \langle \top, \perp \rangle \in c, \text{invite}(b)(a)\}$

Questions: We follow the spirit of Karttunen's (1977) semantics of *wh*-Qs and propose that *wh*-phrases denote GQs quantifying over proper names, i.e., dynamic GQs, as in (5).

- (5) **who**[↑] := $\lambda f. \bigcup \left\{ f(\mathcal{P}) \mid \mathcal{P} \in \{\text{Ahn}^\uparrow, \text{Bill}^\uparrow\} \right\}$

We assume that in *wh*-questions **only** *wh*-phrases introduce drefs to $\top$ (cf. Murray 2010), since they provide the foreground information and establish sets of alternatives that people restrict their attention to (von Stechow & Zimmermann 1984; Krifka 2001; a.o.). The denotation of *who enters* is a set of context change potentials, i.e., possible sentential answers, as in (6) and Figure 1.

- (6) $\llbracket \text{who enters} \rrbracket = \text{who}^\uparrow \lambda \mathcal{P}. \mathbf{C}(\mathcal{P} \lambda x. (\text{enter}(x))) = \begin{cases} \lambda c. \{\langle \top + a, \perp \rangle \mid \langle \top, \perp \rangle \in c, \text{enter}(a)\} \\ \lambda c. \{\langle \top + b, \perp \rangle \mid \langle \top, \perp \rangle \in c, \text{enter}(b)\} \end{cases}$
 $= \{\llbracket \text{Ahn enters} \rrbracket, \llbracket \text{Bill enters} \rrbracket\}$

Short answers: We can extract possible short answers to a *wh*-Q from the set of possible sentential answers to it by using an operator Λ that takes a question *Q* and returns a dynamic property of sequences *i*. $\top_{s'} - \top_s$ delivers the sequence that is part of $\top_{s'}$ but not $\top_s$. Any sequence *i* that has the property consists of drefs introduced by a possible sentential answer *p* in *Q* (see Figure 2).

- (7) $\Lambda(Q) := \lambda i \lambda c. \bigcup_{p \in Q} \{s' \mid s' \in p(c), \exists s \in c. s \leq s' \ \& \ \top_{s'} - \top_s = i\}$

Quantification over short answers: The present proposal accounts for many phenomena that call for the use of short answers to *wh*-Qs—*wh*-conditionals being one of them. Concretely, I propose

that the two *wh*-clauses in (3) are questions, (see also Liu 2016), denoting the set Q_1 and Q_2 respectively, and each of them is operated on by Λ . The conditional introduced by *jìu* expresses adverbial quantification: a covert adverbial akin to *always* ($\mathbb{A}$) takes the antecedent clause as restriction and the consequent clause as scope (Kratzer 1981; Cheng & Huang 1996; Chierchia 2000). (3), translated as (8), involves a dynamic universal quantification over sequences. In prose, (8) says: all the sequences that are possible short answers to Q_1 are possible short answers to Q_2 .

$$(8) \quad \mathbb{A}_i \left(\underbrace{\Lambda(Q_1)(i)}_{\text{restriction}} \right) \left(\underbrace{\Lambda(Q_2)(i)}_{\text{scope}} \right) = \lambda c. \left\{ s \in c \mid \forall i. \Lambda(Q_1)(i)(c) \neq \emptyset \rightarrow \Lambda(Q_2)(i)(\Lambda(Q_1)(i)(c)) \neq \emptyset \right\}$$

As a result, if *Ahn* is the short answer to *who enters first*, then it is also the short answer to *who eats first* (see Figure 3). This is the underlying reason for why the two *who*'s seem to co-refer.

Pair-list readings: In multiple *wh*-conditionals, the *wh*-phrases in the antecedent clause establish a list of pairs, and the *wh*-phrases in the consequent clause give rise to the same list.

(9) **Shéi** ná-le **nǎ** **dào cài**, **shèi** *jìu* yào bǎ **nǎ** **dào cài** chī-wán.

who take-Asp which CL dish who then must BA which CL dish eat-up
'Everyone who took a dish must finish it.'

(If Ahn took bread and Dufu corn, Ahn must finish beef and Dufu corn; and if Ahn took corn and Dufu bread, Ahn must finish corn and Dufu bread)

Our proposal is compatible with the quantifying-into-question approach in which a multiple *wh*-question can be understood as a conjunction of two questions. For example, the denotation of *who took which dish* is derived in (10). $\sqcap$ is to pointwisely apply dynamic conjunction $\wedge$ to two sets. Finally, different pair lists correspond to different sequences (cf. Bumford 2015).

$$(10) \quad \llbracket \text{who took which dish} \rrbracket = \llbracket \text{Ahn took which dish} \rrbracket \sqcap \llbracket \text{Dufu took which dish} \rrbracket =$$

$$\left\{ \begin{array}{l} \llbracket \text{A took beef} \rrbracket \wedge \llbracket \text{D took corn} \rrbracket \\ \llbracket \text{A took corn} \rrbracket \wedge \llbracket \text{D took beef} \rrbracket \\ \llbracket \text{A took beef} \rrbracket \wedge \llbracket \text{D took beef} \rrbracket \\ \llbracket \text{A took corn} \rrbracket \wedge \llbracket \text{D took corn} \rrbracket \end{array} \right\} = \left\{ \begin{array}{l} \lambda c. \{ \top_s + \mathbf{b} + \mathbf{a} + \mathbf{c} + \mathbf{d} \mid s \in c, \text{take}(\mathbf{b})(\mathbf{a}), \text{take}(\mathbf{c})(\mathbf{d}) \} \\ \lambda c. \{ \top_s + \mathbf{c} + \mathbf{a} + \mathbf{b} + \mathbf{d} \mid s \in c, \text{take}(\mathbf{c})(\mathbf{a}), \text{take}(\mathbf{b})(\mathbf{d}) \} \\ \lambda c. \{ \top_s + \mathbf{b} + \mathbf{a} + \mathbf{b} + \mathbf{d} \mid s \in c, \text{take}(\mathbf{b})(\mathbf{a}), \text{take}(\mathbf{b})(\mathbf{d}) \} \\ \lambda c. \{ \top_s + \mathbf{c} + \mathbf{a} + \mathbf{c} + \mathbf{d} \mid s \in c, \text{take}(\mathbf{c})(\mathbf{a}), \text{take}(\mathbf{c})(\mathbf{d}) \} \end{array} \right\}$$

The *wh*-conditional in (9) expresses: for any sequence i that is a possible short answer to *who took which dish*, i is also a possible short answer to *who must finish which dish*. Given (10), if $i = \mathbf{b} + \mathbf{a} + \mathbf{c} + \mathbf{d}$ is a short answer to the first question, then it is a short answer to the second question, i.e. Ahn must finish beef and Dufu must finish corn.

Coordination: It is well known that the categorial approach (Hausser & Zaefferer 1979) represents the meaning of a *wh*-Q as a set of short answers. However, it cannot represent coordination of *wh*-Qs as sets of short answers (Groenendijk & Stokhof 1989; Xiang 2016). For this reason, it fails to predict the well-formedness of *wh*-conditionals with coordinated *wh*-phrases.

(11) **Nǐ** chī **shěnmē**, **hē** **shěnmē**, **wǒ** *jìu* yào chī **shěnmē**, **hē** **shěnmē**.

you eat what drink what I then must eat what drink what

'No matter what you eat and what you drink, I must eat and drink the same things.'

My proposal can easily capture (11). In the antecedent clause, *you eat what* is conjoined with *you drink what* via $\sqcap$. The short answer is a sequence consisting of a food and a drink. The same mechanism is applied to the consequent clause.

Conclusion: I have proposed a novel way to derive short answers to *wh*-Qs from propositional answers using dynamic semantics. The proposal not only offers an adequate analysis for Mandarin *wh*-conditionals, but can also be extended to English free relatives and quantificational variability effects of *wh*-Qs, which Xiang (2016) has used to motivate the semantic necessity of short answers.

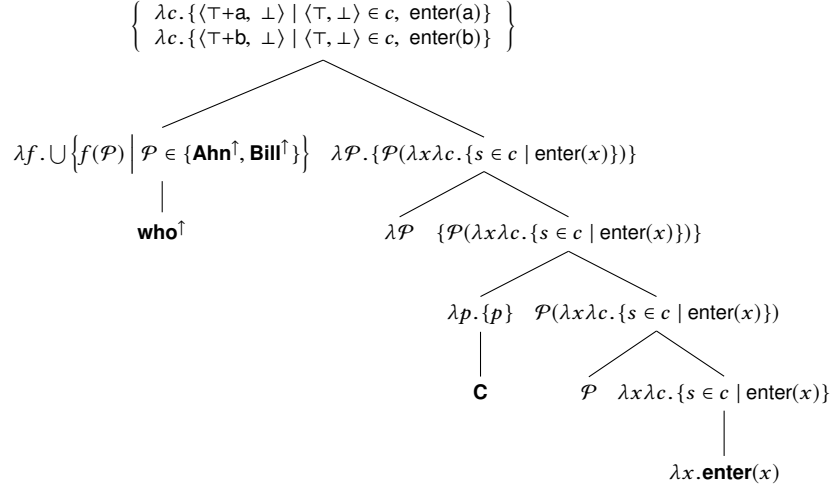


Figure 1: $\text{who}^\uparrow$ undergoes Quantifier Raising, leaving a ‘trace’ $\mathcal{P}$ which is itself typed a dynamic GQ and normally takes scope. In this sense, $\text{who}^\uparrow$ is a higher order dynamic GQ. **C** is the complementizer in the sense of Karttunen (1977), mapping a proposition to a singleton set of the proposition.

$$\{\langle \tau, \perp \rangle\} \xrightarrow{\Lambda(\llbracket \text{who enters} \rrbracket)(a)} \left(\begin{array}{l} \xrightarrow{\llbracket \text{Ahn enters} \rrbracket} \{\langle \tau+a, \perp \rangle\} \xrightarrow{(\tau+a)-\tau=a} \{\langle \tau+a, \perp \rangle\} \\ \xrightarrow{\llbracket \text{Bill enters} \rrbracket} \{\langle \tau+b, \perp \rangle\} \xrightarrow{(\tau+b)-\tau=b} \emptyset \end{array} \right) \bigcup \Rightarrow \{\langle \tau+a, \perp \rangle\}$$

(a) Suppose the sequence i is a that consist of only *Ahn*.

$$\{\langle \tau, \perp \rangle\} \xrightarrow{\Lambda(\llbracket \text{who enters} \rrbracket)(b)} \left(\begin{array}{l} \xrightarrow{\llbracket \text{Ahn enters} \rrbracket} \{\langle \tau+a, \perp \rangle\} \xrightarrow{(\tau+a)-\tau=a} \emptyset \\ \xrightarrow{\llbracket \text{Bill enters} \rrbracket} \{\langle \tau+b, \perp \rangle\} \xrightarrow{(\tau+b)-\tau=b} \{\langle \tau+b, \perp \rangle\} \end{array} \right) \bigcup \Rightarrow \{\langle \tau+b, \perp \rangle\}$$

(b) Suppose the sequence i is b that consist of only *Bill*.

Figure 2: Consider (6). The sequences a and b can make $\Lambda(\llbracket \text{who enters} \rrbracket)$ ‘true’ ($\neq \emptyset$) relative to the input context.

$$\{\langle \tau, \perp \rangle\} \xrightarrow{\Lambda(\llbracket \text{who enters first} \rrbracket)(a)} \{\langle \tau+a, \perp \rangle\} \xrightarrow{\Lambda(\llbracket \text{who eats first} \rrbracket)(a)} \left(\begin{array}{l} \xrightarrow{\llbracket \text{Ahn eats first} \rrbracket} \{\langle \tau+a+a, \perp \rangle\} \xrightarrow{(\tau+a+a)-(\tau+a)=a} \{\langle \tau+a+a, \perp \rangle\} \\ \xrightarrow{\llbracket \text{Bill eats first} \rrbracket} \{\langle \tau+a+b, \perp \rangle\} \xrightarrow{(\tau+a+b)-(\tau+a)=b} \emptyset \end{array} \right) \bigcup$$

Figure 3: The sequence a (only involving *Ahn*) is a possible short answer to *who enters first* and is also a possible short answer to *who eats first*. $\{\langle \tau+a+a, \perp \rangle\}$ indicates *Ahn* enters first and eats first.

Selected references Bitter. 2014. *Temporality*. Oxford. Bumford. 2015. Incremental quantification and the dynamics of pair-list phenomena. *Semantics & Pragmatics*. Cheng & Huang. 1996. Two types of donkey sentences. *Natural Language Semantics*. Chierchia. 2000. Chinese conditionals and the theory of conditionals. *Journal of East Asian Linguistics*. Haida. 2007. *The indefiniteness and focusing of wh-words*. Phd thesis. Murray. 2010. *Evidentiality and the structure of speech acts*. Phd thesis. Liu. 2016. *Varieties of alternatives*. Phd thesis. von Stechow & Zimmermann. 1984. Term answers and contextual change. *Linguistics*. Xiang. 2016. *Interpreting questions with non-exhaustive answers*. Phd thesis.

An Optimality Theory Analysis of Scope Marking at the Syntax/Semantics Interface

Jane Lutken

Background: The phenomenon known as “Scope Marking” (SM) has been described in a variety of languages including German [1], Hindi [2], and Hungarian [3] among others. SM is characterized by the use of a wh-phrase in each clause of a question, as seen in (1) from German.

(1) Was glaubst du, **mit wem** Maria gesprochen hat?

What think you, **with whom** Maria spoken has?

With whom do you think Maria spoke?

Ex. From ([1], 1b)

Analyses of SM fall into two major categories: The Direct Dependency Approach (DDA) and the Indirect Dependency Approach (IDA). DDA analyses typically follow [1] and [4] and analyze the first wh-phrase (*was* in (1)) as an expletive wh-phrase which marks the scope of the true (medial)wh-phrase. Because the scope marker is analyzed as an expletive, the semantic analysis proposed is equivalent to a long-distance (LD) question, formally represented in (2), from [5]. (2), however, would not be the meaning expressed by the SM construction in Hindi according to [2]. Rather, the first wh-phrase is a contentful question over propositions and the ‘second’ question limits the set of possible answers to the ‘first’ question (IDA). In an example such as (1), *with whom Maria has spoken* limits the set of possible answers to *What think you* to only include thoughts about who Maria’s interlocutor was. This analysis is formally represented in (3).

(2) $\lambda p \exists x [\text{person}(x) \wedge p = \text{You think Maria spoke to } x]$

(3) $\lambda p \exists q [\exists x [q = \text{'has spoken' (m,x)}] \wedge p = \text{'think' (j,q)}]$

Puzzle: Both types of analysis have aspects which work well for some SM languages. However, to date there has been no single analysis of SM which satisfactorily accounts for the cross linguistic variation seen. While DDA analyses account for most data in languages like German, they fail to account for Hindi. In contrast, IDA analyses account for Hindi, but do not satisfactorily explain German. Neither account fully explains the pattern of Hungarian.

Proposal: We offer a unified analysis which accounts for both the syntactic and the semantic cross-linguistic variation in SM. To this end, we employ Optimality Theory (OT) [6],[7] which utilizes universal, violable constraints to formalize an input-output (semantic-syntactic) relationship. OT is an ideal tool for this analysis because the puzzle which arises is that a similar syntactic structure in SM-languages does not correspond to similar meanings in these languages. In OT, cross-linguistic variation results from differences in relative rankings of these universal, violable constraints across languages; thus, variation in the input-output pairings are an expected rather than problematic finding.

Establishing the input-output pairings necessitates a detailed analysis of the semantic and pragmatic input. We have identified three pragmatic variables which play crucial roles in determining what type of syntactic output a speaker uses. These include: *Question Under Discussion* (QUD), as described by [8], *Contrastive topic* (CT) [9], and the relevant scope of the matrix and embedded verbs. Different configurations of these pragmatic variables result in a set of possible syntactic outputs which include: Long-distance wh-movement (LD), syntactic Scope Marking (synSM), and sequential questions (seqQ).

We assume that seqQs will always be the optimal output when the matrix verb does not scope over the embedded verb. However, to establish the optimal outputs when the matrix verb does scope over the embedded verb, we created scenarios which manipulated whether or not the subject of the question was a CT and whether or not the question raised by the embedded clause (Q2) was resolved. We conducted a survey asking native speakers of English, German, Hindi, and Hungarian to give acceptability judgments for each syntactic output in each scenario. These languages were chosen in order to show a range of syntactic strategies employed for various semantic inputs. Table 1 summarizes the structures native speakers used for each scenario.

Table 1. Results of acceptability judgments survey

	Condition				
		Q2 unresolved, Subj is CT	Q2 unresolved, Subj not CT	Q2 resolved, Subj is CT	Q2 resolved, Subj not CT
Language	Adult dir. English	LD	LD	LD	LD
	Child dir. English	SeqQ	SeqQ	LD	LD
	Hungarian	SeqQ	SeqQ	synSM	synSM
	German	synSM	SeqQ	synSM	LD
	Hindi	synSM	synSM	synSM	synSM

Conclusion: The results of our survey confirm that there are systematic differences between languages in what type of input results in synSM. For example, in German, synSM is the optimal output when the subject of the sentence is a CT, regardless of whether Q2 is resolved or unresolved. In Hindi, this distinction plays no role and synSM is the optimal output in all cases. We formalized these differences as constraints which vary in ranking across languages. Our OT analysis shows that the cross-linguistic variation seen is attributable to different rankings of a small set of universal, violable constraints. The success of the formalization is not only a testament to OT, but is, to our knowledge the first unified account of SM, addressing both syntactic and semantic analyses as well as pragmatic input and cross-linguistic variation.

References:

- [1] McDaniel, Dana. 1989. Partial and Multiple Wh-Movement. *Natural Language and Linguistic Theory* [2] Dayal, Veneeta Srivastav. 1994. Scope Marking as Indirect Dependency. *NLS*. [3] Horvath, Julia. 1995. Partial Wh-Movement and Wh- “Scope-Markers”. In István Kenesei (ed.), *Levels and Structures*. [4] Riemsdijk, Henk van. 1983. Correspondence Effects and the Empty Category Principle. In Y. Otsu et al. (eds.), *Studies in Generative Grammar and Language Acquisition*. [5] Stechow, Arnim von. 2000. Partial Wh-Movement, Scope-Marking, and Transparent Logical Form. In Uli Lutz, Gereon Müller, & Arnim von Stechow (eds.), *Wh-Scope Marking*. [6] Prince, A. & Smolensky, P. (2004). *Optimality Theory: Constraint Interaction in Generative Grammar*. Oxford: Blackwell. [7] Legendre, Géraldine, Jane Grimshaw, and Sten Vikner (eds.). 2001. *Optimality-Theoretic Syntax*. [8] Roberts, C.1996. Information structure in discourse: Towards an integrated formal theory of pragmatics. In J.H. Yoon & A. Kathol (Eds.), *OSU working papers in linguistics: Papers in semantics*. Columbus: Ohio State University. [9] Büring, D.2003. On D-trees, beans, and B-accent. *Linguistics and Philosophy*.

The Curious Case of Measure Semantics

Yağmur Sağ- Rutgers University

Problem This paper explores *measure constructions* (MCs) in English (*two kilos of apples*) and Turkish (the equivalent of *two kilo apple*) that are composed of a numeral, a classifier (a container noun, e.g. *glass*, or a measure term, e.g. *kilo*), and a substance noun. MCs have two interpretations: the *individuating* and *measure readings*. The former is realized as either the *container* or *portion readings* (Rothstein 2011, Partee & Borschev 2012, Scontras 2014, Khrizman et al 2015).

- (1) Mary brought two glasses/liters of water on the tray. They were blue. CONTAINER READING
- (2) Mary drank two glasses/liters of water, one in the morning, one in the evening. PORTION READING
- (3) Mary added two glasses/liters of water to the soup. MEASURE READING

The character of the measure reading shows variation between English and Turkish. This disparity, to be outlined below, is the main focus here and argued to stem from a two-part semantics that MCs have and two different ways in which these parts are structurally composed.

1. Distributive elements such as reciprocals and *each* are only compatible with the individuating reading of MCs (Rothstein 2011). (4a) is true in a situation where three boxes are put next to each other in the closet, which can be identified as either the container (referring to the boxes) or portion reading (referring to the groups of books coming in boxes). However, it does not describe a situation where the individual books are put next to each other in the closet, referring to *three boxes* as a way of measuring the total amount of books. In contrast, this reading is available in Turkish (4b).

- (4) a. We put the three boxes of books next to each other in the closet.
- b. Üç kutu kitab-ı dolap-ta yan yan-a koy-du-k.
three box book-ACC closet-LOC next next-DAT put-PAST-1PL

2. English MCs can be embedded under mass quantifiers in their measure reading and under count quantifiers in their individuating reading (5) (Rothstein 2011).

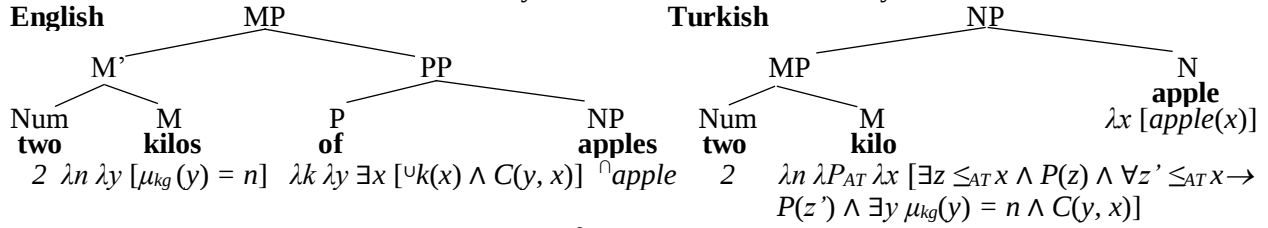
- (5) We gave a little/a few of the twenty kilos of apples to the child we saw on the street.

In the latter case, since the individuated units are kilo-packs of apples, the quantification is over these units, not individual apples. In other words, *a few of the twenty kilos of apples* means a few kilos of the twenty kilos of apples, not a few apples from the given set. The fact that this latter reading is also not available through the measure reading of MCs together with their compatibility with mass quantifiers and incompatibility with distributivity makes it reasonable to assume that English MCs are mass expressions in the measure reading as Rothstein 2011 claims. Conversely, quantification over individual apples is available for Turkish MCs. *Yirmi kilo elmanın bir kaçı* ‘a few of the twenty kilos of apples’ can mean a few apples from the given set. Combining with the distributivity facts, this shows that when the substance noun is count, MCs in Turkish have a count denotation in the measure reading in contrast to English MCs.

Previous Accounts Generally, depending on the type of the substance noun, MCs are taken to denote sets of plural or mass individuals that measure the appropriate amount along a dimension. For e.g., for Scontras 2014 *two kilos of apples* equals to $\lambda x [\cup^{\cap} \text{apple}(x) \wedge \mu_{kg}(x) = 2]$ (cf. Krifka 1989, Champollion 2010). However, under this theory, MCs of English with a count noun have a count denotation, contrasting with the conclusion reached above. Alternatively, Rothstein 2011 claims that when the substance noun is count, it must shift from the count type to the mass type since measurement operates at the mass domain only. So, under her theory *two kilos of apples* is $\lambda x [\text{apples}_{\text{mass}}(x) \wedge \mu_{kg}(x) = 2]$. First, this theory does not account for Turkish MCs. Second, although I follow the idea that measurement occurs at the mass domain, the motivation behind the shift of the count nouns to the mass type remains vague.

Proposal Instead, I argue that measurement universally operates at the domain of portions of matter which is connected to a substance noun by a *Constitution relation* (C) inside the derivation. That is, MCs with the measure reading are composed of two parts, the part with the count or mass substance noun and the part with the measured amount which is always mass. The notion of *portions of matter* and the C relation goes back to Link 1983. The famous example is a ring recently made up from some old gold. Their distinctive properties reveal that even if the ring and the gold in the ring share the portion of matter they are made of, they are not the same entities. They are connected by a C relation, denoted by the materialization function

h , which maps every individual to its corresponding portion of matter, i.e. $C(a, b)$ is true iff $a = h(b)$. If a, b are mass the semantic fact follows trivially because h denotes the identity function on mass individuals.



E: $[[two \ kilos \ of \ apples]] = \lambda y \ [\mu_{kg}(y) = 2 \wedge \exists x \ \cup^{apple}(x) \wedge C(y, x)]$

a set of portions of matter that amount to 2 kilos in weight and constitute a plurality of apples

T: $[[two \ kilo \ apple]] = \lambda x \ [\exists z \leq_{AT} x \wedge apple(z) \wedge \forall z' \leq_{AT} x \rightarrow apple(z') \wedge \exists y \ \mu_{kg}(y) = 2 \wedge C(y, x)]$

a set of pluralities of apples constituted by a portion of matter that amount to 2 kilos in weight

English MCs take the portion of matter, hence the measured amount, as the basis of reference. This generates a mass denotation, which makes MCs compatible with mass quantifiers even if the substance noun is count. Since the set denoted by the substance noun is existentially closed, it is not accessible for reciprocals or count quantifiers. In Turkish MCs, the basis of reference is the substance noun, hence when it is count, the measure expression is also count and available for reciprocals and count quantifiers. This reversal lines up with the existence/absence of *of* which subsequently generates different syntactic structures for the MCs of the two languages ([Schwarzschild 2006](#)). English ones are headed by the classifier which introduces the amount measured, and the noun is the complement of *of* which introduces the C relation. In Turkish - a strict head-final language - due to the absence of *of* they are headed by the noun and the C relation is wired into the denotation of the classifier. This account not only addresses the two-way denotations of MCs but also contributes to the ongoing debates on the semantic and syntactic status of *of*.

Substance nouns of Turkish are singular or mass, contrasting with English nouns which are plural or mass. I follow [Scontras 2014](#) in that the latter are kind terms and get instantiated inside the derivation. Based on [Sag's 2018](#) claim that singular nouns in Turkish are ambiguous between atomic properties and singular kind terms which cannot be instantiated ([Dayal 2004](#)), I propose that they are the simplest form of a predicate in Turkish, atomic if count, mass otherwise. Alternatively, they could be treated similar to singular substance nouns of Brazilian Portuguese. Building on [Pires de Oliveira and Rothstein 2011](#), [Rothstein 2017](#) claims that they are furniture-type mass nouns that are compatible with distributivity unlike in English. However, this account cannot be adopted for Turkish since it patterns with English in that sense.

Further Implications For some speakers of English, MCs in the measure reading allow distributivity and count quantification as in Turkish. I call this Grammar B of English, for which I assume that it is possible to treat *of* as a PF-inserted element and shift the head from the classifier to the substance noun. Although such a strategy is restricted to a dialectal variation in the measure reading, it is available to all speakers in the individuating reading, yielding the ambiguity between the container (headed by the classifier) and portion readings (headed by the noun) ([Partee & Borschev 2012](#), [Scontras 2014](#)). Turkish MCs differ from English ones in not having the container reading. While (1) can refer to the containers and be followed as 'They were blue.', it can only refer to the water in Turkish, so such a follow-up is infelicitous. I argue that this is because Turkish MCs are always headed by the noun, given the absence of *of*. I believe that it must be harder to reanalyze a structure inserting an element that is not there than reanalyzing a structure by deleting an existing one. Instead, complying with the strict head-final property of the language, for the container reading, a different structure is formed with the reversed order (e.g. *iki su-dolu bardak* 'two water-full glass'). As a final remark, while *of* is an indicator of the structural difference between English and Turkish MCs, its absence does not always implicate a Turkish-like behavior. What is actually at stake is which element the MC is/can be headed by. Namely, depending on the headedness properties of the language in question it is possible for its MCs to lack *of* but be headed by the classifier. E.g., German and Dutch MCs lack *of*, yet still pattern with English both in both readings, and this is expected if they have an English-like structural alignment. Confirming this, [Grestenberger 2015](#) and [Ruys 2017](#) show that German and Dutch MCs are headed by the classifier in the measure and container readings as in English.

Artificial language learning and the learnability of semantic distinctions: the case of evidentiality

Dionysia Saratsli¹

Stefan Bartell¹

Anna Papafragou²

¹*Department of Linguistics and Cognitive Science, University of Delaware*

²*Department of Psychological and Brain Sciences, University of Delaware*

It is often assumed that cross-linguistically more prevalent distinctions are easier to learn (*Typological Prevalence Hypothesis* - TPH).¹ Prior work supports this hypothesis in phonology, morphology and syntax^{2,3,4} but has not addressed semantics. Furthermore, tests of the TPH with children are complicated (e.g., because of the potential role of cognitive development). Here we ask whether the TPH predicts the relative learnability of semantic distinctions in a domain that is not grammaticalized in English and can be taught to adults without native language interference: *evidentiality* (the encoding of information source).

Cross-linguistically, there are three common types of evidential morpheme: Direct (firsthand/perceptual evidence), Inferential (inference based on evidence), and Reportative (hearsay).⁵ In general, evidential systems mark Reportative or Inferential access (systems that only mark Direct access are rare)⁶; the most widespread evidential system involves only Reportative morphemes.⁵ According to TPH, Reportative-marking systems should be the most learnable, while Direct-marking systems the least learnable. We test this prediction using Artificial Language Learning (ALL). A previous learnability study⁷ on evidentiality offered preliminary support for TPH. However, that study used static pictures where Reportative access alone was marked with a salient visual cue that could have boosted system learnability. The current study uses dynamic videos whose visual characteristics are consistent across systems.

English speakers (n=101) were exposed to an “alien” language that was similar to English but had a novel verb-final morpheme, *ga*, and had to figure out what *ga* meant. They were shown 21 videos in which a girl gained access to an event through observation of someone’s action (Direct), inference from visual clues (Inferential), or report (Reportative; 7 videos per access type). For each video, the girl’s access to the event was controlled by a third character (Fig.1). At the end of each video, the girl produced a sentence with or without *ga*. There were three between-subject conditions depending on system (whether *ga* marked Direct, Inferential or Reportative access). Participants later completed a Production task: they watched 12 new videos (4 per access type) and had to complete the girl’s sentences with *ga* if appropriate. They also completed a Comprehension task: they watched 36 videos (12 per access type), on half of which the girl made errors in the use of *ga* (50% misses, 50% incorrect inclusions), and had to say whether *ga* was used correctly or not. A one-way ANOVA conducted on composite Production and Comprehension scores revealed an effect of System ($F(2,98)=6.535$, $p<0.01$; Fig.2). Pairwise comparisons (Bonferroni) revealed a significant advantage of the Reportative system over the Direct ($p=.004$) and the Inferential system ($p=0.01$).

Our data support the TPH, since the typologically prevalent Reportative evidential system was learned best and the rare Direct worst. Furthermore, our data support the conjecture that, cross-linguistically, indirect sources seem to be marked preferentially (and acquired more easily) compared to direct sources. We discuss this pattern in terms of the pragmatic need to mark indirect (and potentially more unreliable) over direct sources of information.

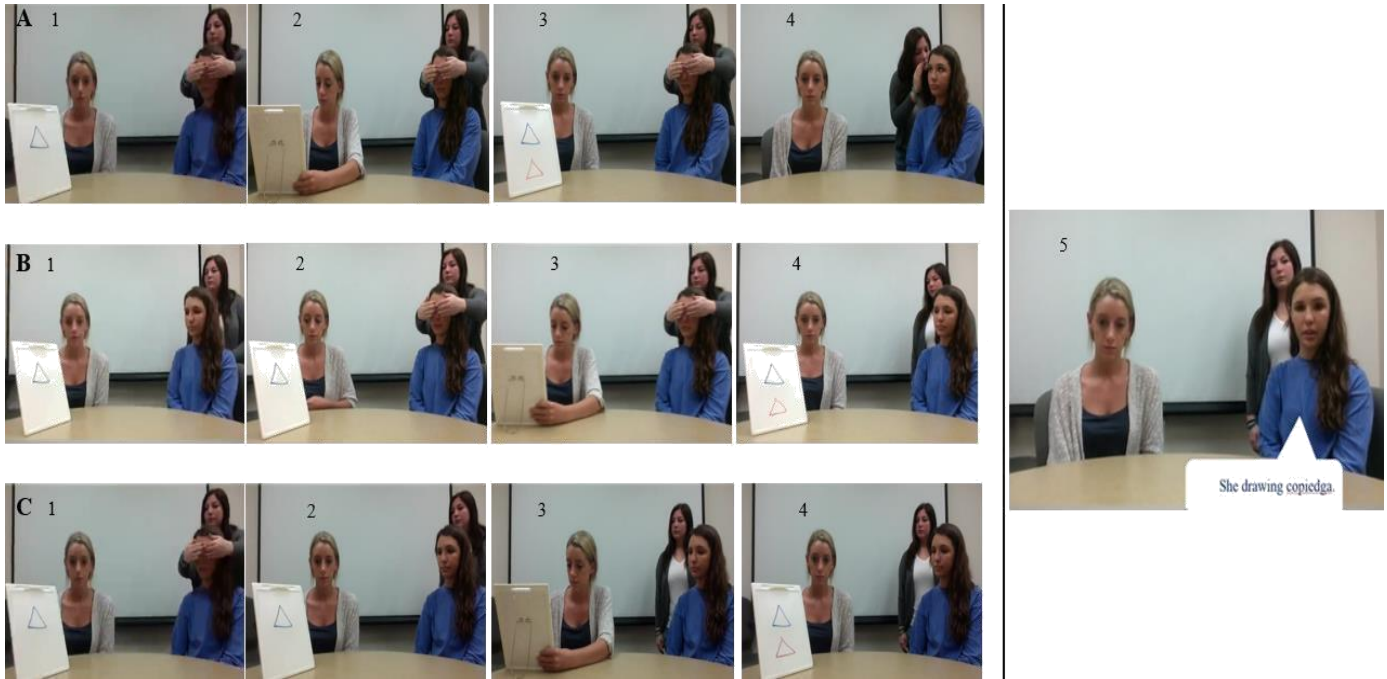


Figure 1. Sample screenshots from one video for each Access type: (A) Reportative, (B) Inferential, (C) Direct. Videos across systems have the same ending (Panel 5). In that panel, the girl in blue utters an evidential sentence ("She drawing copiedga".)

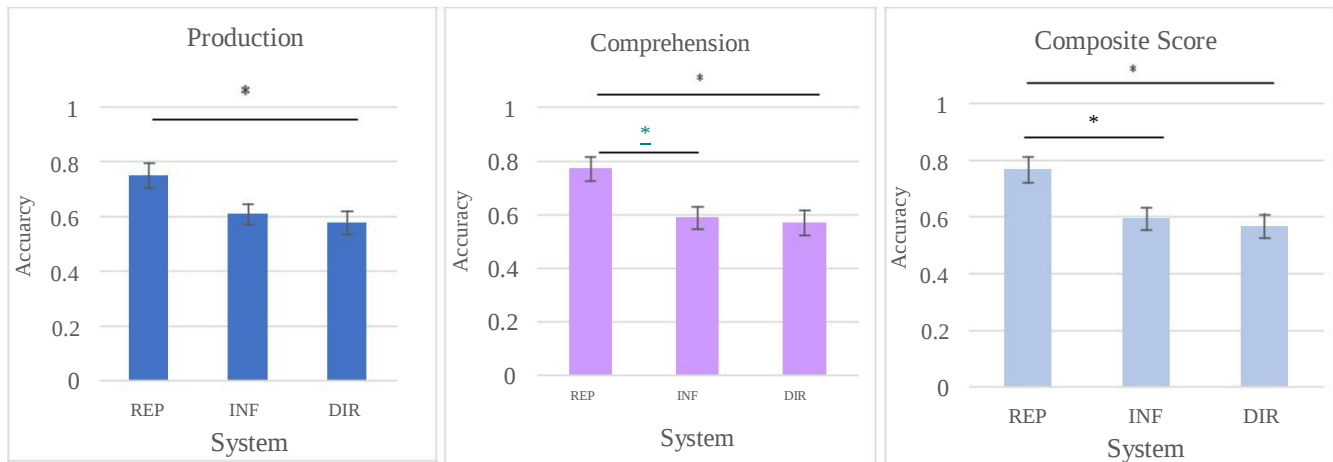


Figure 2. Accuracy Means Across Systems. The composite score represents a combined Production/Comprehension score. Error bars represent ± 1 S.E.

References

- Gentner, D., and Bowerman, M. (2009). Why some spatial semantic categories are harder to learn than others: The typological prevalence hypothesis. In Guo, J. et al., (Eds.), *Crosslinguistic approaches to the psychology of language: Research in the tradition of Dan Isaac Slobin* (pp. 465–480). Hillsdale, NJ: Erlbaum.
- Culbertson, J., and Newport, E. (2015). Harmonic biases in child learners: In support of language universals. *Cognition*, 139: 71–82.
- Culbertson, J., Smolensky, P., and Legendre, G. (2012). Learning biases predict a word order universal. *Cognition*, 122(3): 306–329.
- Finley, S., and Badecker, W. (2009). Artificial language learning and feature-based generalization. *Journal of Memory and Language*, 61(3): 423–437.
- Aikhenvald, A. (2018). Evidentiality: The Framework. In A.Y. Aikhenvald (Ed.), *The Oxford Handbook of Evidentiality* (pp.1-43). Oxford: Oxford University Press.
- De Haan, F. (2013). Semantic distinctions of evidentiality. In Dryer, M. and Haspelmath, M. (Eds.), *The World Atlas of Language Structures Online*. Max Planck Institute for Evolutionary Anthropology, Leipzig.
- [NAMES WITHHELD] (2015). Cross-linguistic variation and the learnability of semantic systems. Paper presented at the 40th Annual Meeting of the Boston University Conference on Language Development. Boston, MA.

Analyzing the Infelicity of Tantalizing Statements

Benjamin C. Shavitz, The CUNY Graduate Center

According to Uegaki & Sudo 2017 and Elliott, Klinedinst, Sudo, & Uegaki 2017, there are three types of clause-embedding predicates: Responsive predicates can embed both declarative and interrogative complements (e.g., *know*), rogative predicates can only embed interrogative complements (e.g., *wonder*), and anti-rogative predicates can only embed declarative complements (e.g., *believe*) (Uegaki & Sudo 2017). Based on this model, all clause-embedding predicates that accept interrogative complements (responsive and rogative predicates) do so without any concern for the nature of the Wh-item upon which a complement is built.

I explore observed data that appear to constitute an exception to this assumed Wh-indiscrimination. The responsive predicate “told me” seems to pose a problem for the assumption that clause-embedding predicates accept interrogative complements without regard for the complements’ Wh-items/complementizers. “Told me” is a responsive predicate, accepting both declarative and interrogative complements:

- (1) Sean told me that Moira went to the store. (*Declarative Complement*)
- (2) Sean told me what Moira did. (*Object Interrogative Complement*)
- (3) Sean told me who went to the store. (*Subject Interrogative Complement*)

But, even though “told me” accepts interrogative complements, it faces problems with the Wh-item “whether” (along with the nearly synonymous complementizer “if”). The following sentence (as well as other sentences with the same structure) has been judged infelicitous:

- (4) #Sean told me whether/if Moira went to the store. (*Binary Interrogative Complement*)

I attempt to account for this seeming exception without violating the integrity of the Wh-indiscrimination principle by appealing to Grice’s conversational maxims and the semantics-pragmatics interface, instead of the semantics-syntax interface. Following Gamut 1991’s interpretations of the Gricean maxims, I appeal to the considerations of relevance and relative logical strength, as well as my own proposed Tantalization Uncooperativeness Condition (which assesses the perceived relevance of something almost stated) in order to provide an explanation for the seemingly problematic “told me whether/if” sentences:

In the case of “told me [interrogative complement]” sentences, there are always logically stronger sentences available to the speaker because, by virtue of having already been told the declarative substance of the interrogative complement, the speaker necessarily has the ability to utter a sentence with a declarative complement that not only introduces the information provided in the interrogative version – namely that the speaker was told something of a certain nature – but that also introduces the additional information of what exactly the speaker was told. Grice asserts that a cooperative speaker will always utter the logically strongest relevant statement available. Since a logically stronger sentence is always available to the speaker of a “told me [interrogative complement]” sentence, but only “whether/if” “told me [interrogative complement]” sentences are infelicitous, the logically stronger sentence for a “told me whether/if” statement must be considered relevant by the listener, while those of other “told me [interrogative complement]” sentences must not. My proposed Tantalization Uncooperativeness Condition accounts for this difference.

The Tantalization Uncooperativeness Condition is defined as follows:

(5) Tantalization Uncooperativeness Condition: If the speaker of an utterance has almost said something but not quite, that something is perceived as relevant by the listener, unless its relevance is overtly dismissed,

with the definition of “almost said something” being:

(6) A speaker has almost said something if the statement she has uttered leaves exactly two belief states available to the listener with respect to the matter discussed in the utterance.

I also attempt to explain certain syntactically unrelated data by extending the analysis developed for “told me whether/if” sentences, in order to provide additional evidence of the analysis’s utility and to further motivate my Tantalization Uncooperativeness Condition: the analysis is used to explain the infelicity of perfectly ambiguous statements devoid of context, such as:

(7) #All the students did not fail the exam.

References

- Elliott, Patrick D., Nathan Klinedinst, Yasutada Sudo, and Wataru Uegaki. 2017. Predicates of Relevance and Theories of Question Embedding. *Journal of Semantics*, 34, 547-554.
- Gamut. L.T.F. 1991. *Logic, Language, and Meaning*, Volume 1: *Introduction to Logic*. Chicago: University of Chicago Press.
- Uegaki, Wataru and Yasutada Sudo. 2017. The Anti-Rogativity of Non-Veridical Preferential Predicates. In Cremers, A., T. van Gessel, and F. Roelofsen (eds.), *Proceedings of the 21st Amsterdam Colloquium*, 492-501.

HIGH SHIFTY OPERATORS IN GEORGIAN INDEXICAL SHIFT

Sigwan Thivierge, University of Maryland

OVERVIEW. The phenomenon known as *indexical shift* induces embedded indexicals to pick up their reference from an attitude event, and not the actual utterance context. Various proposals attribute indexical shift to (i) the ‘shiftability’ of the indexicals (Schlenker, 1999, 2003, et seq.), (ii) a ‘monster’ operator that rewrites the context parameters for the embedded clause (Anand & Nevins, 2004; Anand, 2006; Shklovsky & Sudo, 2014), or (iii) a hybrid of the two (Sundaresan 2018). Given that, in languages with indexical shift, multiple indexicals in the same embedded clause must shift together, this *Shift Together* constraint (Anand & Nevins 2004; Anand 2006) provides strong evidence in favour of the monster operator approach.

Indexical shift has been reported for a wide range of languages, such as Nez Perce (Deal 2014), Uyghur (Sudo 2012), Slave (Rice 1986), Zazaki (Anand & Nevins 2004, Anand 2006), among many others (see also Deal 2018; Sundaresan 2018). Here, I discuss indexical shift in Georgian, which has been underdescribed in the literature. I show that it occurs under both speech and attitude verbs, but also in matrix clauses used to report the speech or thought of others. This suggests that indexical shift is induced by an operator separate from the verb. The operator is exponed by a phrase-final *-o*, which, in multiple embeddings, can appear in each clause.

DATA & ANALYSIS. The two following sentences demonstrate Georgian indexical shift. In (1), a description in the embedded clause can be read *de re*. In this context, Dato knows Bryan Adams for his activism work, but does not know that he’s also a singer. If Dato says to me, “I saw Bryan Adams in Tbilisi,” then, at a concert at which Bryan Adams is singing, I can say to you:

- (1) Dato-m tkv-a v-nax-e es momxreral-i Tbilifi-**o**
 Dato-ERG say-3SG.AOR 1-see-PART.AOR DEM.PROX singer-NOM Tbilisi-O
 ‘Dato_i said I_i saw this singer in Tbilisi.’

Since Dato does not know that Bryan Adams is a singer, the embedded clause cannot be a quotation. Furthermore, a shifted 1st person in the embedded clause must be read *de se*, as in (2) (see Schlenker 1999, Messick 2016).

- (2) Dato-m tkv-a (rom) avad v-ar-**o**
 Dato-ERG say-3SG.AOR C sick 1-be.PRES-O
 ‘Dato_i said I_i am sick.’

✓Earlier today, Dato told me he (Dato) is sick.

Dato, at the hospital for a checkup, happens to glance at the chart of a patient’s blood work. Dato, a doctor himself, sees that the patient is clearly sick, but the name is hard to read. He says to the nurse when she comes in, “This guy is really sick.”

As shown in (3), Georgian indexical shift obeys the *Shift Together* constraint: embedded 1st and 2nd person indexicals must both shift, if they shift at all. I take this restriction to be indicative of an operator that rewrites the context parameters of indexicals in its scope.

- (3) Nino-m u-txr-a Dato-s (rom) da-g-i-nax-e-**o**
 Nino-ERG APPL-tell-3SG.AOR Dato-DAT C PRV-2-APPL-see-AOR.PART-O
 ✓ ‘Nino told Dato that I saw you.’ ✗ ‘Nino told Dato_j that I saw you_j.’
 ✓ ‘Nino_i told Dato_j that I_i saw you_j.’ ✗ ‘Nino_i told Dato that I_i saw you.’

In (4), an unshifted 1st person pronoun can appear in the matrix clause. This behaviour is unsurprising if the shifty operator is introduced by the matrix verb—it can only induce indexical shift

for the embedded clause, not the matrix. Thus, (4) is possible in a context where Nino and Dato have been dating for a significant period of time, and Nino tells me she loves him. Given this information, I can then tell you:

- (4) Nino-m m-i-txr-a (rom) Dato m-i-χvar-s-**o**
 Nino-ERG 1-APPL-say-3SG.AOR C Dato.NOM 1-APPL-love-3SG.PRES-O
 ‘Nino_i told me that I_i love Dato.’

Notably, *-o* may also appear on the *matrix* verb, shown in (5), inducing shift for the matrix 1st person pronoun. The context where this sentence is felicitous is very specific—it must be one where Nino and Dato have been dating for a significant period of time, and Nino tells Gio she loves Dato. If Gio later tells me about this, then I can tell you:

- (5) Nino-m m-i-txr-a-**o** (rom) Dato m-i-χvar-s-**o**
 Nino-ERG 1-APPL-say-3SG.AOR-O C Dato.NOM 1-APPL-love-3SG.PRES-O
 ‘Nino_i told me_k that I_i love Dato.’
 (Where Gio and the matrix 1st person pronoun are co-referent)

In the embedded clause, the matrix attitude verb introduces a monster operator to rewrite the context parameters such that the 1st person pronoun is shifted to refer to Nino, the attitude holder of that matrix verb. In the matrix clause, however, the 1st person pronoun shifts to refer to Gio, which is possible given that Gio is a speaker who is topical in the current conversation.

Notably, there is no higher verb to introduce an operator to shift the matrix 1st person pronoun. It must thus be the case that the operator can merge into the matrix CP and scope over indexicals in the matrix clause, a phenomenon that has not been attested for other languages with indexical shift. This is schematized in English below: the matrix 1st person indexical gets its reference from the shifty operator in the matrix CP—namely, of ‘Gio’—and the embedded shifty operator in the embedded CP layer shifts the reference of the 1st person indexical to ‘Nino’.

- (6) [CP OP [TP Nino_i told me_k [CP OP (that) [TP I_i love Dato]]]]

This data point thus provides novel evidence for independent shifty operators. One of the standard views of monster operators is that they are introduced by speech or attitude verbs; in Georgian, however, there need not be a speech or attitude verb if *-o* is used to report the speech or attitudes of a speaker topical in the conversation. Furthermore, the sentence in (5) shows that shifty operators are not limited to embedded clauses—they can appear in matrix CPs as well.

CONCLUSION & IMPLICATIONS. I have shown that, in Georgian indexical shift, a shifty operator is distinct from attitude verbs and may merge high in the matrix CP structure. These are properties that have not been featured in the literature previously, and thus bear on theories of indexical shift that derive shift via operators which are themselves introduced by a speech or attitude verb. Further, the behaviour of Georgian *-o* may be related to free indirect discourse—that is, matrix *-o* may serve as an indicator of free indirect speech.

SELECT REFERENCES. Anand, P. & Nevins, A. 2004. Shifty operators in changing contexts. *Proceedings of SALT XIV*. Anand, P. 2006. De de se. *Doctoral Dissertation, MIT*. Schlenker, P. 1999. Propositional attitudes and indexicality: a cross-categorical approach. *Doctoral Dissertation, MIT*. Deal, A.R. 2018. Shifty asymmetries: Universals and variation in shifty indexicality. *Ms.* Sundaresan, S. 2018. An alternative model of indexical shift: Variation and selection without context-overwriting. *Ms.*

Title: Neural Correlates of Linguistic Modality

Introduction. The ability to communicate about things outside the here-and-now is a core trait of human language¹, yet its neural underpinnings are understudied. This study investigated the contribution of modal verbs, words that refer to possible states of affairs that are not actual or known. Modal expressions such as *may* and *must* have received a lot of attention in semantics^{2,3}, philosophy⁴ and language acquisition^{5,6}, but we know very little about the online processing of linguistic modality and its neural mechanisms. The extensive literature on the semantic properties of linguistic modality, allow us to postulate clear hypotheses about the operations involved in modal processing. In this study we investigated the neural mechanisms underlying the processing of assertive verbs (like *do*) and modal verbs (*may* and *must*), provided in both likelihood and obligation contexts, in order to gain insight into the basic brain mechanisms involved in modal processing. We focused on the following three properties of modality:

- 1) *general contribution of modality*: Modal expressions differ from factual expressions in their contribution to the discourse. Factual assertions (e.g. 'John loves Mary') are claims about the world under discussion which can be either accepted or rejected, allowing the addressee to **update their beliefs** about this world accordingly⁷. In contrast, modal statements (e.g. 'John must love Mary') do not make any claims about the actual world directly, rather they **postulate possible scenarios** that are compatible with the world under discussion^{2,3}.

Q1: Is there an overall effect for modal processing (versus processing of factual assertions)?

- 2) *modal force variation*: Modals come with different *forces*: e.g. *may* indicates possibility while *must* indicates necessity. The formal representation of the force of a modal verb is often expressed as a quantifier ranging over possible worlds^{2,3}. Necessity verbs like *must* are associated with a **universal quantifier** '∀' (for every x) while possibility verbs like *may* are associated with an **existential quantifier** '∃' (for at least some x).

Q2: Are there differences in modal processing of different forces (possibility versus necessity)?

- 3) *modal flavor variation*: Modals come in different flavors, e.g. in "It must be raining" *must* has a likelihood (epistemic) reading, while in "You must eat vegetables" it has a (deontic) reading of an obligation. The modal flavors differ in the modal base on which they postulate possible worlds. **Likelihood** modals pick out worlds compatible with what is known in the world of evaluation. **Obligation** modals pick out worlds compatible with certain rules and norms.

Q3: Are there differences in modal processing of different flavors (epistemic versus deontic)?

Insight into online processing of modal verbs across these three dimensions could help us gain insight into how modal meaning is computed by the brain, whether it shares machinery with related phenomena, and in which order operations involved in computing modal meaning occur.

Methods. A magnetoencephalography (MEG) study (N=25) compared visually presented sentences (rapid serial visual presentation) containing the ambiguous modals 'may' and 'must' against sentences containing the non-modal verb 'do'. In order to have *do* naturally appear in the same

position as *may* and *must*, our sentences contained VP ellipsis (... and the squires do/may/must too), controlled for elided-VP length and complexity (Fig 1). The interpretation of the ambiguous modals was dependent on prior (pre-normed) contexts biasing towards either an inferential or obligation reading (Fig 2). Target sentences (N=240) were followed by a task sentence, where participants indicated whether these were natural continuations of the story or not.

Results. We defined regions of interest (ROIs) based on previous neurolinguistic literature looking at related phenomena. For the force manipulation we looked at areas [IPS, PCC, IFG] involved in processing semantic elements that are theoretically related to universal and existential quantification: logical quantifiers (*some/all*)^{8,9} and the connectives (*and/or*)¹⁰. For the flavor manipulation we looked at areas [TPJ, STS, mPFC, rACC] involved in theory of mind processing and social cognition^{11,12}. Within these ROIs we did not find any effect for FORCE in the anticipated ROIs. We did find increased activity for the necessity (must) modals in the right Anterior Cingulate Cortex (rrACC), between 400-450 ms (Fig 4A). For the FLAVOR manipulation we observed increased activation for the obligation condition over the likelihood condition between 695-720 ms in the right Superior Temporal Sulcus (rSTS), associated with goal-directed action and intention/desires¹¹ (Fig 4B). These effects did not survive multiple corrections. A full-brain analysis in the time window 100-900ms after target word onset revealed a significant spatio-temporal cluster (Fig 2) reflecting a robust increase for the non-modal conditions over modal ones, at 210-350ms starting around the right Temporoparietal Junction (rTPJ) and spreading up to the right Inferior Parietal Sulcus (rIPS) and right medial surfaces (cuneus - posterior cingulate cortex) (Fig 5).

Discussion. We hypothesize that this increased activation for the non-modal condition may reflect computations involved with evaluation and integration of claims made about the world of evaluation, a process absent from the modal condition as those sentences only assert possible compatibilities with the evaluated world. This belief-updating function is in line with suggestions that the rTPJ plays a role in theory revision and conceptual change¹³ and supports that the right hemisphere is involved in pragmatic processing and contextual coherence^{14,15}. The late effect of obligation>likelihood in the rSTS suggests that flavor information becomes available at a later stage in processing, and is computed after FORCE information is being processed.

References. 1) Hockett, C., F. (1959). Animal “Languages” and Human Language. *Human Biology*, 31(1), 32–39. 2) Kratzer, A. (2012). Modals and conditionals: New and revised perspectives (Vol. 36). Oxford University Press. 3) Hacquard, V. (2006). Aspects of modality. Ph.D. Thesis, MIT. 4) Fintel, K. von, & Gillies, A. S. (2010). Must . . . stay . . . strong! *Natural Language Semantics*, 18(4), 351–383. 5) Cournane, A. 2015. *Modal Development: Input-Divergent L1 Acquisition in the Direction of Diachronic Reanalysis*. PhD Dissertation. University of Toronto. 6) Papafragou, A. (1998). The acquisition of modality: Implications for theories of semantic representation. *Mind & language*, 13(3), 370-399. 7) Stalnaker, R. C. (1978). Assertion. Wiley Online Library. 8) Troiani, V., Peelle, J. E., Clark, R., & Grossman, M. (2009). Is it logical to count on quantifiers? Dissociable neural networks underlying numerical and logical quantifiers. *Neuropsychologia*, 47(1), 104–111. 9) Olm, C. A., McMillan, C. T., Spotorno, N., Clark, R., & Grossman, M. (2014). The relative contributions of frontal and parietal cortex for generalized quantifier comprehension. *Frontiers in Human Neuroscience*, 8. 10) Baggio, G., Cherubini, P., Pischredda, D., Blumenthal, A., Haynes, J.-D., & Reverberi, C. (2016). Multiple neural representations of elementary logical connectives. *NeuroImage*, 135, 300–310.

11) Saxe, R., Carey, S., & Kanwisher, N. (2004). Understanding Other Minds: Linking Developmental Psychology and Functional Neuroimaging. *Annual Review of Psychology*, 55(1), 87–124. **12)** Saxe, R., & Powell, L. J. (2006). It's the Thought That Counts: Specific Brain Regions for One Component of Theory of Mind. *Psychological Science*, 17(8), 692–699. **13)** Corbetta, M., Patel, G., & Shulman, G. L. (2008). The Reorienting System of the Human Brain: From Environment to Theory of Mind. *Neuron*, 58(3), 306–324. **14)** Martin, I., & McDonald, S. (2003). Weak coherence, no theory of mind, or executive dysfunction? Solving the puzzle of pragmatic language disorders. *Brain and language*, 85(3), 451-466. **15)** Kuperberg, G. R., Lakshmanan, B. M., Caplan, D. N., & Holcomb, P. J. (2006). Making sense of discourse: An fMRI study of causal inferencing across sentences. *Neuroimage*, 33(1), 343-361.

Figure 1. A. Example Stimulus Set

Obligation (deontic)	
*(word-by-word, 300ms on/150ms off)	
Context	Normally, only knights sit at the round table.
Target*	
factual	But the king says that their squires do too.
possibility	But the king states that their squires may too.
necessity	But the king learns that their squires must too.
Continuation	They form a circle. (congruent)
Likelihood (epistemic)	
Context	Apparently, knights overhear a lot of secrets in the castle.
Target	
factual	But the servant thinks that their squires do too.
possibility	But the servant concludes that their squires may too.
necessity	But the servant realizes that their squires must too.
Continuation	All the castle's secrets are safe. (incongruent)

Figure 2. Results pre-norming flavor bias

Example: Apparently knights overhear a lot of secrets in the castle. But the servant thinks that their squires are ____ to as well.

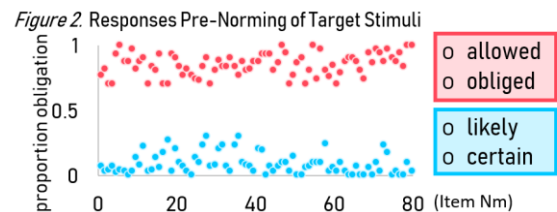


Figure 3. Effects Modal Force and Flavor ROI analysis; uncorrected for multiple comparisons

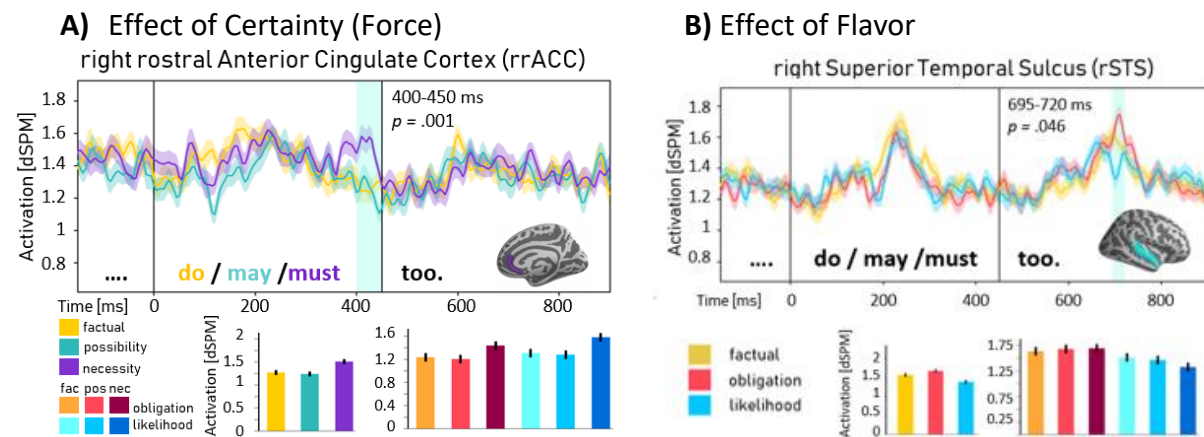
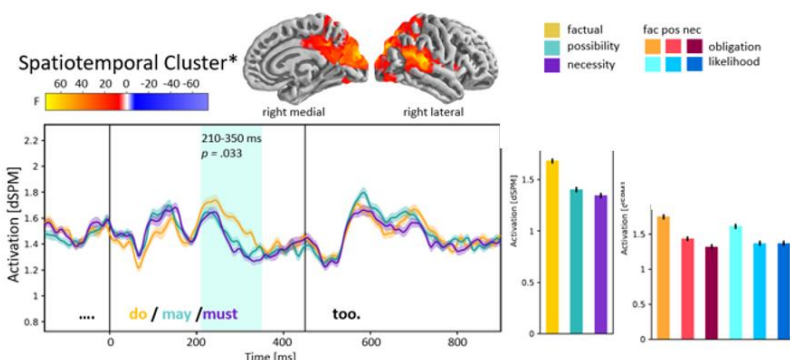


Figure 4. Effect of Factuality



Figuring out epistemic uses of English and Dutch modals: The role of aspect
Annemarie van Dooren, University of Maryland (avdooren@umd.edu)

Introduction. This paper investigates how children figure out that functional modals like English *must* and Dutch *moeten* 'must' can be used to express different 'flavors' of modality: epistemic, deontic, bouletic, and so on (Ex.1). The existing acquisition literature^{[1],[2]} shows that children produce functional modals with epistemic meanings up to a year later than with root (non-epistemic) meanings, suggesting that they may initially fail to realize that these modals can express epistemic meanings in addition to root. We conducted a corpus study on English and Dutch child-directed speech to examine how modality is expressed in speech to young children, to investigate the ways in which the input may help or hinder learners uncover the multiple flavors functional modals can express. Our results^{[3],[4]} suggest that the way adults use functional modals may obscure their polysemy: modals are mainly used either with a root or an epistemic flavor, with an overall strong bias towards root uses. This is even more so in Dutch than in English. Yet, children eventually figure out modal polysemy. To investigate how the linguistic input might help, we explore a distributional difference between roots and epistemics that could give away epistemic flavor, concerning modals' temporal properties, which we track by the aspectual properties of the modal's prejacent. We show that aspect is differently distributed across root and epistemic flavors of functional modals in both English and Dutch - despite the slight differences between modals in these languages. However, the strong usage bias towards root meanings may lead to a weakened signal, suggesting the cue will be useful only if learners expect flavor-based constraints, and use them in combination with cues stemming from the situational context.

Linking hypothesis. The *temporal orientation* (TO) of modals could potentially cue in learners into epistemic flavor. Root and epistemic modals have been claimed to differ in their TO, the time of the event expressed by the modal's prejacent relative to the evaluation time of the modal: root modals are future-oriented, i.e., the time of the event described by the prejacent needs to follow the time of evaluation of the modal, while epistemic modals show no such restriction^[5]. Consider (A)-(C): In (A), both future and present TO are possible, and both root and epistemic interpretations are available. (B) and (C), which respectively trigger a past and a present temporal orientation, seem to only express epistemic possibility.

(A)	John may run.	Future/Present TO	root, epistemic
(B)	John may have run.	Past TO	*root, epistemic
(C)	John may be running.	Present TO	*root, epistemic

Why should root modals be restricted to future TO? Root modality expresses possibilities given a set of circumstances and priorities (desires, orders, goals...). Intuitively, such possibilities are made trivial when the circumstances are already *settled* with respect to the event or state expressed by the prejacent. Epistemic modality, on the other hand, expresses possibilities given a body of knowledge and such possibilities are not *settled* even with respect to a past or present event or state - what we know about the past or the present may be partial. More formally, the restriction on modals' temporal orientation has been argued to follow from a general constraint against the vacuous use of modals, called the *Diversity Condition* (DC)^[5]. The DC requires that the proposition expressed by the modal's prejacent does not hold (or fail to hold) throughout the worlds of the Modal Base, i.e., the set of worlds that the modal quantifies over.

How could a language learner make use of the constraint in (A)-(C) to figure out epistemic flavor? Modals with a past or present orientation could alert learners that the modal expresses a non-root meaning, provided that (i) they expect modal meanings to be governed by something like the DC, and thus expect root meanings to be restricted to future TO, and (ii) the constraint is clearly manifest in the modals produced in the input. TO might sometimes be difficult to determine in context for the learner, however, which is why we study aspectual correlates that largely track TO. In the absence of overt grammatical aspect markers (B)-(C), lexical aspect constrains TO and as such might cue in the learner to epistemic flavor: while bare eventives can be future- or present-oriented (A), bare statives tend to be present-oriented (D).

(D) John may be home. Present TO *root, epistemic

Methods. We examined 12 English (ENG) and 7 Dutch (DU) child-input corpora (Manchester^[6], age range=1;09-3;00, Groningen^[7], age range=1;05-3;06). We extracted adult utterances with modals (ENG: 81,854/564,625 (21.7%); DU: 40,486/18,1003 (22.4%)). We coded modals for SYNTACTIC_CATEGORY [lexical, functional] and FLAVOR [root, epistemic, future] (Ex.2). To determine FLAVOR for the polysemous items, we hand-examined the transcripts. We coded all functional modals for GRAMMATICAL_ASPECT (perfect, progressive) and LEXICAL_ASPECT (stative, eventive, perception verb).

Results. English and Dutch children hear epistemic vocabulary using lexical verbs (e.g. *think, know*) and adverbs (e.g. *maybe*) quite frequently (~5% of total utterances; Ex.3). Epistemics are however rarely expressed by functional modals (e.g. *must*) in both English (9% (n=1,779)) and Dutch (2% (n=235)) (Ex.4). Both grammatical and lexical aspect are distributed differently across root and epistemic flavor: In English, 9.8% (n=171) of the epistemics with a verbal complement take a complement with a perfect or a progressive, compared to less than 1% of the roots (n=167) (Dutch: 4.7% (n=11) vs. 0% (n=1)). 67.6% (n=950) of the English epistemics without grammatical aspect in its complement furthermore takes a stative predicate, compared to 14.3% (n=2,193) of the roots (Dutch: 64.2% (n=97) vs. 3.7% (n=354)) (Ex.5). A generalized linear mixed-effects model supports that a stative complement (containing a progressive, perfect, or bare stative) is a significant predictor of flavor in English and Dutch (Ex.6).

Discussion. Our corpus results show that the root/epistemic asymmetry seen in child production data may be an input effect, as the way adults use functional modals makes it difficult to see that they express epistemic meanings. Yet, children eventually pick up on epistemic meanings. How? The linguistic context provides distributional cues that could help learners. We investigated the TO of modals argued to restrict the distribution of root meanings. Our results show that aspect, which largely tracks TO, distinguishes roots and epistemics in both English and Dutch: roots mostly combine with eventives, epistemics mostly combine with statives. However, given the high proportion of root uses, the number of root modals with a stative preadjacent is actually higher than the number of epistemics with statives (Ex. 5). Does this threaten to make the linking hypothesis useless? Crucially, the majority of roots with statives is future-oriented (counterfactuals and coerced eventives (Ex. 7)) and is as such in line with the constraint on TO. We think that in these particular instances, the context should make the future-orientation salient. In sum, while experiments^[2] have to determine whether children actually figure out epistemic flavor using aspect, we show that the ingredients are available in English and Dutch. A discussion on the differences found between the two languages follows.

- (1) John must eat meat.
 i. 'John is likely a meat eater.' epistemic
 ii. 'John is obliged to/wants to/needs to eat meat.' root (deontic, teleological...)
- (2) Modals lemmas by syntactic category
Functional Aux = can, could, may, must, should, might, shall, will, would
 Quasi-Auxiliaries (QA) = have to, got to, ought to, supposed to, going to
Lexical V = *epis*: know, think, seem...; *root*: want, order, let's...
 Adv = *epis*: maybe, perhaps, probably...
 Adj = *epis*: sure, certain... *root*: able, capable... *epis/root*: possible...

- (3) *Table 1*: Aggregate raw counts of modal utterances in the input by flavor and syntactic category (lexical & functional), English and Dutch (% of total utterances)

	Lexical modality			Functional modality	
	epistemic	root	epi/root	epi/root	future
ENG	15,750 (4.6%)	12,433 (3.7%)	2,434 (0.7%)	20,528 (6%)	22,661 (6.7%)
DU	9,402 (5.2%)	582 (0.3%)	11 (0.01%)	20,765 (11.5%)	463 (0.3%)

- (4) *Table 2*: Aggregate raw counts of input functional modals (% of total utterances)

	epistemic	root	total
ENG	1,779 (9.3%)	17,423 (90.7%)	19,202
DU	235 (1.6%)	14,950 (98.4%)	15,185

- (6) *Table 4*: Results of the model tests the effect of aspect on usage flavor.

(glmer, Flavor~Stative+(1|corpus)
 family=binominal). FLAVOR:
 $\beta = 2.53, <0.0001^{***}$ (ENG),
 $\beta = 3.31, <0.0001^{***}$ (DU)

- (5) *Table 3*: Aggregate raw counts of modal utterances in input by flavor and aspect (% of total utterances)

		epistemic	root
ENG	stative	1,121 (71.1% of 1,577 epi) (grammatical: 171; lexical: 950)	2,360 (15.3% of 15,473 root) (grammatical: 167; lexical: 2,193)
	eventive	448 (28.4%)	11,111 (71.8%)
	perception	8 (0.5%)	2,002 (12.9%)
DU	stative	108 (67.7% of 162 epi) (grammatical: 11; lexical: 97)	355 (3.7% of 9,680 roots) (grammatical: 1; lexical: 354)
	eventive	54 (35.8%)	9,038 (93.4%)
	perception	0 (0%)	287 (3.0%)

- (7) a. You **could** have said hello. counterfactual (Mother, Carl 2;04)
 b. Well they **can** have a tray each if they want. coerced eventive (Mother, Ruth 2;07)

References. [1] Kuczaj & Maratsos (1975). What children can say before they will. *Merrill-Palmer Quarterly of Behavior and Development* 21, 89-111. [2] Cournane (2015). *Modal development*. PhD. Thesis, UToronto. [3] van Dooren, Dieuleveut, Cournane & Hacquard (2017). Learning what *must* and *can* must and can mean. *Proceedings of the Amsterdam Colloquium*. [4] van Dooren, Tulling, Cournane & Hacquard (to appear). Discovering Modal Polysemy. *Proceedings of BUCLD 43*. [5] Condoravdi (2002). Temporal interpretation of modals. In Beaver, Casillas Martinez, Clark & Kaufmann (eds.) *The construction of meaning*. [6] Theakston, Lieven, Pine, & Rowland (2001), The role of performance limitations in the acquisition of verb-argument structure. *Journal of Child Language* 28, 127-152. [7] Wijnen & Verrips (1998). The Acquisition of Dutch Syntax, In Gillis & De Houwer (eds.), *The Acquisition of Dutch*.

Words take time: Auditory stimuli and strategic processing in semantic priming

Yosiane White (University of Pennsylvania)

yosiane@sas.upenn.edu

This study examines participants' task-related strategy use in auditory semantic priming experiments. Semantic priming (SP) occurs when lexical access to a target word is facilitated by a preceding semantically related word (e.g. *dog* – CAT versus *table* – CAT). Visual SP is frequently used to study access to the semantic representation of words (Neely, 1991). More recently, auditory stimuli have been used in SP paradigms for similar purposes. This is despite the fact that visual word presentation is holistic while auditory word presentation is incremental (Cutler, 2002; 2012). Visual SP has proven highly susceptible to strategy use (Neely, 1991). The main parameters that induce these strategic effects are, (1) a large inter-stimulus interval (ISI) between prime and target, giving a participant time to predict the upcoming target (den Heyer et al., 1983), and (2) a high proportion of related pairs, priming participants to expect semantically related targets (McNamara, 2005). The susceptibility of auditory SP to task-related strategy use has not yet been systematically studied. In this study, three experiments suggest that participants do not use target-prediction strategies in auditory SP when a minority of pairs are related, regardless of the length of the ISI. This differs markedly from visual SP, and highlights an advantage of using auditory SP for studying access to semantic representations.

Experiment 1: Exp1 asks whether varying ISI in auditory SP has the same effect as in visual SP. 117 native English undergraduates completed a paired lexical decision task in which they heard 318 primes (randomised across 4 counterbalanced lists), followed by a 200ms or 800ms ISI, and then a semantically related ($\frac{1}{3}$ of the items), unrelated ($\frac{1}{3}$), or nonword ($\frac{1}{3}$) target. RTs (in ms) to the target were measured from the onset of the target sound file. Minimal a-priori trimming and model criticism (Baayen & Milin, 2010) were done before fitting a linear mixed effects model in R. As expected, participants responded significantly faster to related targets than unrelated targets at both the 200ms ($t = -25.9$) and 800ms ISI ($t = -22.1$) (1)(3). Interestingly, the 52ms priming effect in the 800ms ISI condition is significantly smaller than the 64ms effect in the 200ms condition ($t = 2.768$).

This result suggests that either participants are not strategically predicting the targets in this experiment (despite > 90% reported awareness of related pairs in a post-test questionnaire), or participants are using strategies which boosts the priming effect in both ISI conditions, but rapid decay of auditory SP reduces the effect at the long ISI.

Experiment 2: A possible explanation for the results in Exp1 is that the short ISI is not short enough to hinder strategy use. Exp2 uses a between subjects design for ISI with a 200ms ISI and a 0ms ISI to test this. 55 undergraduates participated in Exp2. We replicate the SP effect found with a 200ms ISI in Exp1 ($t = -19.6$), and find significant priming at the 0ms ISI ($t = -19.46$). Further, we find no difference in priming effects between the ISI conditions ($t = -0.51$).

Experiment 3: An alternative explanation for Exp1 is that $\frac{1}{3}$ related pairs is not a low enough proportion to thwart strategy use. Exp3 attempts to reduce the utility of strategic prediction by reducing this proportion to $\frac{1}{6}$ of the items. 110 native English speakers took part. Interestingly, we find very similar priming effects to Exp1 (2). Both the 200ms ($t = -2.17$) and 800ms ($t = -2.75$) ISI conditions yield significant priming, although now we find no difference across the ISIs.

So far, neither a reduction of the ISI from 800ms to 0ms, nor a reduction of the ratio of related pairs from $\frac{1}{3}$ to $\frac{1}{6}$ reduced auditory SP magnitude. This contrasts with the visual SP literature that finds strategy use at ISIs over 200ms and $\frac{1}{3}$ related pairs. The current results support a theory that no target-prediction strategies are being used in auditory SP. A planned Experiment 4 will push this hypothesis by increasing the related pairs to $\frac{1}{2}$ to increase the utility of strategy use.

(1) Experiment 1 reaction times
in ms (with SD) and priming effects by ISI

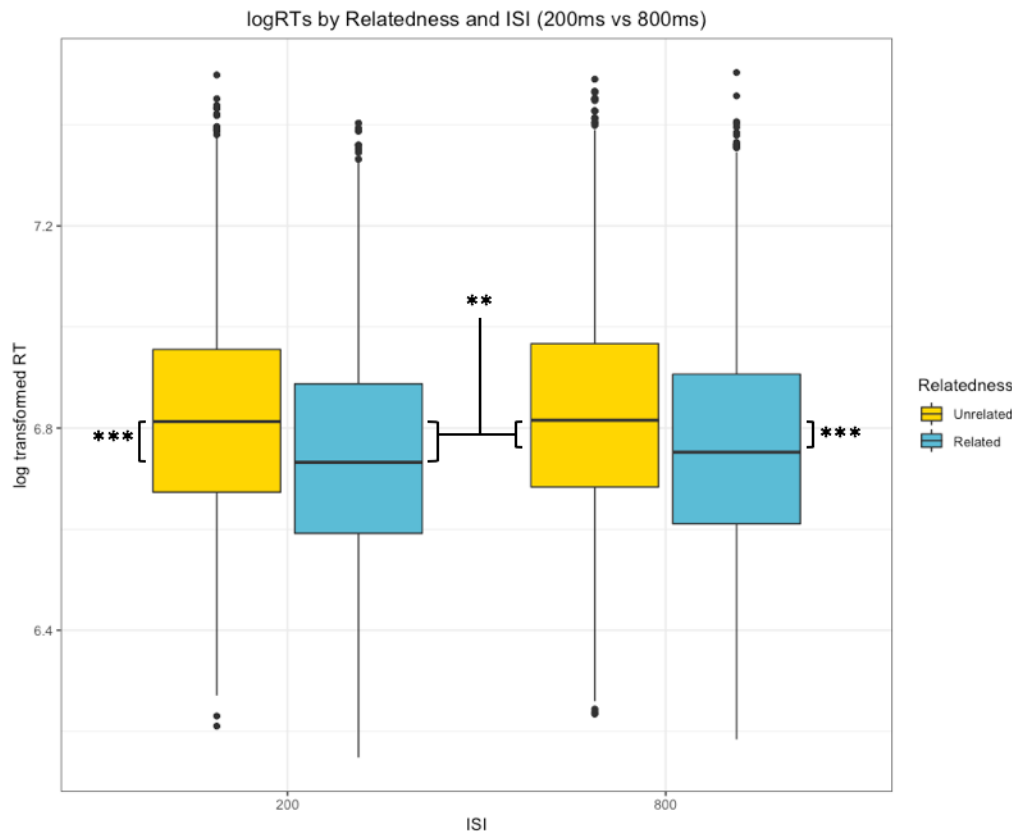
	200ms ISI	800ms ISI
Unrelated	949 (183)	955 (182)
Related	885 (179)	903 (183)
Priming effect	64***	52***

(2) Experiment 3 RTs in ms (with SD)
& priming effects at a 200ms and 800ms ISI

	200ms ISI	800ms ISI
Unrelated	973 (183)	986 (181)
Related	922 (169)	923 (164)
Priming effect	51*	63**

(3) Log RTs by ISI for Experiment 1

(*** = $p < .000$, ** = $p < .01$, * = $p < .05$, ns = not significant)



References

Baayen, R. H., & Milin, P. (2010). Analyzing reaction times. *International Journal of Psychological Research*, 3(2), 12-28. **Cutler, A.** (2002). Lexical access. In *Encyclopedia of cognitive science* (pp. 858-864). Nature Publishing Group. **Cutler, A.** (2012). Native listening: Language experience and the recognition of spoken words. Mit Press. **den Heyer, K., Briand, K., & Dannenbring, G. L.** (1983). Strategic factors in a lexical-decision task: Evidence for automatic and attention-driven processes. *Memory & Cognition*, 11(4), 374-381. **McNamara, T. P.** (2005). *Semantic priming: Perspectives from memory and word recognition*. Psychology Press. **Neely, J. H.** (1991). Semantic priming effects in visual word recognition: A selective review of current findings and theories. In D. Besner & G. W. Humphreys (Eds.), *In Basic processes in reading: Visual word recognition*, Hillsdale, NJ: Erlbaum, 264-336.

Embedded Speech Act Layers and Enhancement Effect

Georgetown University

Akitaka Yamada

Introduction. Recently, researchers have proposed two-tiered model for the syntactic structure above/around CP (Miyagawa 2012, 2017; Kim 2018; Zu 2018; Portner et al. to appear) and the upper layer, aka., the speech act layer, is understood as the structure only available in the root-clause, whereas the lower layer can appear in embedded environments. By examining addressee-honorific markers in Japanese, however, I would argue that speech act layers are also embeddable, contra these previous studies.

Data. Addressee-honorific markers are verbal suffixes that encode the speaker's respect to the addressee. Though Korean and Japanese are well-known for such an honorific system, recent studies have revealed that genealogically unrelated languages exhibit a similar grammatical encodings --- Basque (Oyherçabal 1993; Miyagawa 2012, 2017; Haddican 2015; Antonov 2015; 2016; Zu 2015, 2018), Burmese (Okell 1969: 375; Wheatley 1982; Myint 1999; Kato 2018), Thai (McCready 2014, to appear), Punjabi (Kaur 2018; Kaur and Yamada in prep), Tamil (McFadden ms), and Magahi (Alok and Baker ms; Baker and Alok 2019). In Japanese, addressee-honorific markers are observed under some embedding predicates as shown in (1) (Tagashira 1973; Harada 1976; Nonaka and Yamamoto 1985; Kaur and Yamada in prep; Yamada to appear).

(1) Embedded addressee-honorific marker

[*Watasi-no musuko-ga kabin-o kowasi-te*
I-GEN son-NOM vase-ACC break-CV
simai-masi-ta-koto]-o o-wabi itasi-masi-ta.
MAL-HON_A-PST-C-ACC HON-apologizing do.HON_U-HON_A-PRS
'(i) I am (hereby) apologizing (to you) for my son's having broken
the vase;

With such an embedded addressee-honorific marker, the politeness level of the sentence is enhanced (the ENHANCEMENT EFFECT). First, how is such an embedded addressee-honorific marker licensed? Second, how does the embedded addressee-honorific marker strengthen the politeness level?

Analysis. Assuming that addressee-honorific markers are involved with agreement

(Miyagawa 2012, 2017), I assume that there is a speech act layer in the embedded clause. Following my earlier work (Yamada 2019), I also assume that the denotation of the relevant honorific feature is the tuple of three elements <the speaker, 1, the addressee>. This non-at-issue meaning is shipped to the storage of respect even time we have the speech act layer. So, when the embedded clause is interpreted, the respect object is shipped to the storage as in (2). Then, when the main clause is interpreted, another respect object is shipped to the storage, resulting in the storage expansion in (3).

(2) respect: {<akitaka, 1, satoshi>}

(3) respect: {<akitaka, 1, satoshi>, <akitaka, 1, satoshi>}

In order to summarize what kind of respect is expressed by the sentence, we need to collapse the triples into one representative triple by summing up the politeness intensity expressed; i.e., <akitaka, 2, satoshi>.

Alok, D. & Baker, M. (2018). On the mechanics (syntax) of indexical shift: Evidence from allocutive agreement in Magahi. Ms., Rutgers University. **Antonov, A.** (2015). Verbal allocutivity in a crosslinguistic perspective. **Haddican, B.** (2015). A note on Basque vocative clitics. **Haddican, B.** (to appear). The syntax of Basque allocutive clitics. *Glossa*. **Harada, S.-I.** (1976). Honorifics. **Kato, A.** (2018). Burmese. **Kaur, G.** (2018). Addressee agreement as the locus of imperative syntax. **McCready, E.** (2014). A semantics for honorifics with reference to Thai. **McCready, E.** (to appear). *Honorification and social meaning*. **McFadden, T.** (2017). The morphosyntax of allocutive agreement in Tamil. Ms. LeibnizZentrum Allgemeine Sprachwissenschaft. **Oyharçabal, B.** (1993). Verb agreement with non-arguments: on allocutive agreement. In *Generative studies in Basque linguistics*. **Portner, P.** (forthcoming). Commitment to Priorities. In Fogel, D., Harris, D. & Moss, M. (eds.) *New Work on Speech Acts*. Oxford: Oxford University Press. **Portner, P., Pak, M., Zanuttini, R.** (to appear). The addressee at the syntax-semantics interface: Evidence from politeness and speech style. *Language*. **Zu, V.** (2015). A two-tiered theory of the discourse. **Zu, V.** (2018). *Discourse participants and the structural representation of the context*. **Yamada, A.** (2019). Expressiveness from a Bayesian Perspective. *JELS* 36.

Contradictory Descriptions with Absolute Adjectives

Jeremy Zehr (UPenn) and Paul Egré (IJN-ENS Paris)

1. Borderline contradictions. Several experiments over the last decade indicate that borderline cases of vague predicates license contradictory descriptions such as “x is neither tall nor not tall”, or “x is tall and not tall” [1,2,4]. These so-called borderline contradictions have been regimented in paraconsistent-friendly accounts of vagueness [1,3,4,5], in which “and” and “neither... nor...” descriptions are treated symmetrically. In [6], however, a marked preference was evidenced for descriptions of the form “neither *P* nor not *P*” over “*P* and not *P*” when *P* is a *relative* gradable adjective. [6] left open whether this pattern would extend to *absolute* gradable adjectives. In this paper, we report on an experiment that replicates the findings of [6] for relative adjectives, but shows no such asymmetry for absolute adjectives. This difference invites a revision of the account laid out in [6], by integrating data concerning the treatment of lexical antonyms.

2. Study. Drawing on [6] we presented participants in an online experiment with short vignettes describing target borderline cases for 8 adjectives, asking whether the contradictory descriptions were true, along with two unproblematic true and false control descriptions. Each participant judged either relative or absolute adjectives, either with their syntactic negation (*not tall/not flat*) or their lexical antonym (*short/bumpy*). For absolute adjectives, borderline cases were designed to be cases located very near the closed bound of the scale (see [7] and Examples below).

Figure 1 reports a bar-graph of our results. We fitted logistic regression models predicting the *Yes* answers of participants with over 50% accuracy on both controls (N=138/167). Our factors were *Negation* (syntactic vs. lexical), *Category* (relative vs. absolute) and *Description* (“and” vs. “neither” vs. “ctl-true” vs. “ctl-false”). Random effect variables were included to reflect by-adjective and by-participant variation. For *relative* adjectives, “neither” and “ctl-true” descriptions did not significantly differ (regardless of negation, no interaction) whereas only *syntactic*-“and” descriptions significantly differed from “ctl-false.” For *absolute* adjectives, no significant contrast was found between “neither” and “and” descriptions (regardless of negation, no interaction). All other simple effects were significant. Acceptance of *syntactic*-“and” descriptions was significantly greater for *relative* than for *absolute* adjectives; we found no significant interaction of *Category* × “ctl-false” vs. “and” in the *syntactic* groups.

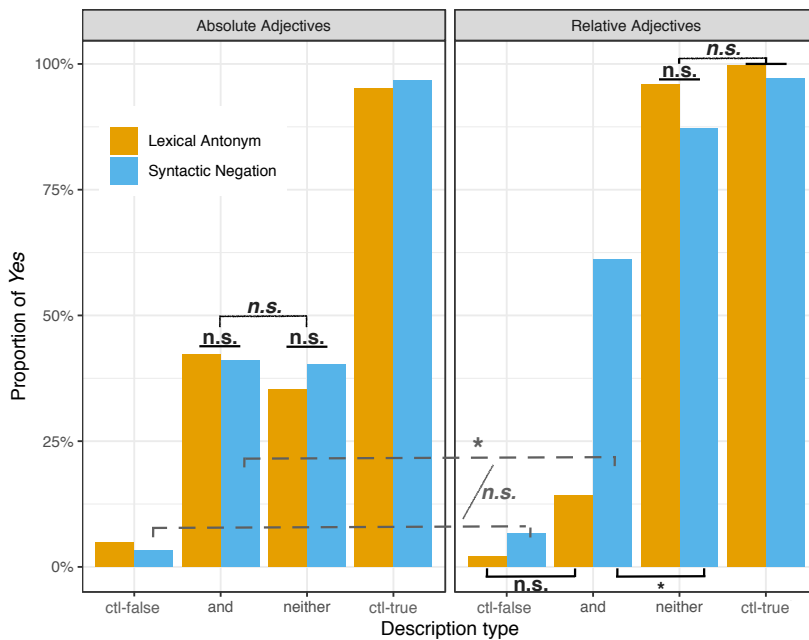


Figure 1.

3. Interpretation. Two main observations can be made about the data. Firstly, syntactic negations and lexical antonyms yield similar acceptance rates with absolute adjectives, but not so with relative adjectives (see the “and” contrast). Secondly, the preference for “neither” descriptions over “and” descriptions in borderline cases is only evidenced for relative adjectives.

To account for both effects, we adopt the strict-tolerant framework of [3], also applied in [6] and [8], in which both relative and absolute adjectives admit *strict* and *tolerant* readings, respectively narrowing and widening their extension, thereby creating gaps and gluts with their negative counterparts. In [6], the preference for “neither” descriptions in relative adjectives is explained by postulating a precedence of strict readings over tolerant readings (see [5]) and by assuming a “strictly” operator to be inserted above predicate negation (interpreting “neither tall or not tall” as “neither strictly tall, nor strictly not tall”). The account is at best incomplete, for it remains silent on lexical antonyms as well as absolute adjectives.

In light of the data, we postulate that (i) *for relative adjectives, lexical antonyms are semantic contraries, rather than contradictories*, leaving a gap between them; (ii) *for absolute adjectives, lexical antonyms are contradictories that leave no gap*, in much the same way as syntactic negations. In (i) we depart from [9]’s account, which treats every antonym $\bar{P}$ as the semantic complement of P . Under assumption (ii), the strict-tolerant account directly explains the symmetric acceptance of contradictory descriptions for absolute adjectives. To illustrate, if *dry* literally denotes a 0% amount of water, by (ii) *not dry* and *wet* denote any amount in the complement region ($> 0\%$) but an amount of 1% can still count as *dry* under a tolerant reading (creating a glut $\rightsquigarrow$ “and”) whereas the same 1% amount can fail to count as *not dry* or *wet* under a strict reading (creating a gap $\rightsquigarrow$ “neither”). For *short* and *tall*, assumption (i) directly predicts the applicability of “neither short nor tall” in the gap region. On the other hand, tolerance may fail to fill the pre-existing gap so as to make *short* overlap with *tall*, thus explaining the massive rejection of “and” descriptions with relative antonyms.

We need one additional assumption, namely (iii) *syntactic negations of relative adjectives can be locally strengthened to their lexical antonym* (see [10]). By (iii), speakers may reinterpret “neither tall nor not tall” as “neither tall nor short.” This explains the near-ceiling acceptance of lexical-“neither”-relative descriptions and, at the same time, does not make “and” descriptions more acceptable, thereby accounting for the lack of a significant interaction.

Overall, the present account is both more general and simpler than the one proposed in [6]: it assumes a local strengthening operation independent of the strict-tolerant machinery.

Examples.

1. A survey on heights has been conducted in your country. In the population there are people of a very high height, and people of a very low height. Then there are people who lie in the middle between these two areas. Imagine that Sam is one of the people in the middle range. Comparing Sam to other people in the population, is it true to say the following? *Sam is neither tall nor short* ☐Yes ☐No *Sam is tall and short* ☐Yes ☐No

Sam is in the middle range ☐Yes ☐No *Sam’s height is very high* ☐Yes ☐No

2. Sam is a blacksmith working in a traditional workshop where they produce swords. In the workshop, there are blades that have no bulges and there are blades that have many small bulges. Then there are blades with exactly one small bulge. Imagine that the blade that Sam is looking at has exactly one little bulge. Comparing the blade that Sam is looking at to the other blades, is it true to say the following?

The blade is neither flat nor bumpy ☐Yes ☐No *The blade is flat and bumpy* ☐Yes ☐No

The blade has exactly one bulge ☐Yes ☐No *The blade has many bulges* ☐Yes ☐No

References. [1] Alxatib, S. & Pelletier, F.J. 2011. The Psychology of Vagueness. *M&L* 26. [2] Serchuk, P. et al. 2011. Vagueness, logic and use. *M&L* 26(5). [3] Cobreros, P. et al. 2012. Tolerant, Classical, Strict *JPL* 41. [4] Ripley, D. 2011. Contradictions at the borders. *Vagueness in Communication*. [5] Cobreros, P. et al. 2015. Pragmatic Interpretations of Vague Expressions. *JPL* 44(4). [6] Egré, P. & Zehr, J. 2018. Are Gaps Preferred to Gluts? *The Semantics of Gradability, Vagueness, and Scale Structure*. [7] Kennedy, C. & McNally, L. 2005. Scale structure, degree modification, and the semantics of gradable predicates. *Language* 81 (2). [8] Burnett, H. 2014. A Delineation Solution to the Puzzles of Absolute Adjectives. *L&P* 37. [9] Krifka, M. 2007. Negated Antonyms. *Presupposition and Implicature in Compositional Semantics*. [10] Ruytenbeek, N. et al. 2017. Asymmetric inference towards the antonym. *Glossa* 2.