

Conceptual Event Units of Putting and Taking in Two Unrelated Languages

Rebecca Defina (rebecca.defina@mpi.nl)

Max Planck Institute for Psycholinguistics &
International Max Planck Research School for Language Sciences, 6500AH Nijmegen, The Netherlands

Asifa Majid (asifa.majid@mpi.nl)

Max Planck Institute for Psycholinguistics, 6500AH Nijmegen, The Netherlands
Donders Institute for Brain, Cognition and Behaviour, Radboud University, Nijmegen, The Netherlands

Abstract

People automatically chunk ongoing dynamic events into discrete units. This paper investigates whether linguistic structure is a factor in this process. We test the claim that describing an event with a serial verb construction will influence a speaker's conceptual event structure. The grammar of Avatime (a Kwa language spoken in Ghana) requires its speakers to describe some, but not all, placement events using a serial verb construction which also encodes the preceding taking event. We tested Avatime and English speakers' recognition memory for putting and taking events. Avatime speakers were more likely to falsely recognize putting and taking events from episodes associated with take-put serial verb constructions than from episodes associated with other constructions. English speakers showed no difference in false recognitions between episode types. This demonstrates that memory for episodes is related to the type of language used; and, moreover, across languages different conceptual representations are formed for the same physical episode, paralleling habitual linguistic practices.

Keywords: Conceptual event units; event segmentation; serial verb constructions; linguistic relativity.

Introduction

Events occur in a continuous stream with no clear boundaries between them. Despite this continuity, we think and talk about events in terms of discrete and divisible units. Previous research has largely focused on the factors influencing the segmentation of events. This paper examines the question from a complementary perspective: what factors might lead event elements to be grouped together into a single conceptual event unit.

When we perceive ongoing activity, we segment it automatically and unconsciously (Kurby & Zacks, 2008; Zacks et al., 2001a). The conceptual event units thus created are structured hierarchically. Each event unit is made up of smaller units, which in turn combine to form larger units (Zacks, Tversky, & Iyer, 2001b). So, what counts as a single conceptual event unit depends to some extent on which level of granularity we are talking about. The choice of granularity level appears to be made at the point of reporting. Prior to that, people segment events at multiple levels of granularity simultaneously (Zacks, Speer, Swallow, Braver, & Reynolds, 2007).

Previous research shows that event units are determined by at least three main factors. First, the inherent properties

of events, such as points of greater motion, have a large effect on where event boundaries are placed (Newton, Engquist, & Bois, 1977; Zacks, 2004). Second, repeated co-occurrence, particularly in different contexts, encourages event elements to be grouped together, regardless of their inherent properties (Avrahami & Karev, 1994). Finally, the particular event schema that the person engages for an event affects the way they segment it (Zacks et al., 2007); for instance, whether or not a person understood the actor's goal influences the way a participant segments the actor's behavior (Zacks, 2004). The fact that event schemas influence conceptual event structure suggests that language may also play a role here. This paper explores this possibility.

Previous cross-linguistic research on the role of language in event cognition has largely focused on differences in the encoding of manner and path in motion events. The results have been mixed: Some studies have found language effects (e.g., Filipović, 2011; Finkbeiner, Nicol, Greth, & Nakamura, 2002; Kersten et al., 2010), but others have not (e.g., Gennari, Sloman, Malt, & Fitch, 2002; Loucks & Pederson, 2011; Papafragou, Massey, & Gleitman, 2002). More recently, scholars have begun to explore other aspects of language and how they might influence event cognition, particularly with respect to causal actions (e.g., Fausey & Boroditsky, 2011; Wolff, Jeon, & Li, 2009). For example, Wolff et al. (2009) found that the semantic property of whether or not a language allowed an intermediary actor to function as an agent affected both the syntactic and non-linguistic partitioning of events, consistent with the proposal that language may play a role in event segmentation.

In Wolff et al.'s (2009) study both the semantic and the syntactic differences are potential instigators of the non-linguistic event segmentation patterns. The current study narrows in on the potential link between syntactic encoding in particular and the concomitant non-linguistic partitioning of events.

One type of syntactic structure that is particularly interesting for event cognition is serial verb constructions (SVCs). These constructions allow multiple verbs to be placed within a single clause without coordination or subordination (Aikhenvald, 2006; Durie, 1997). The particular syntactic features vary across languages, though there is a shared set of core, prototypical features

(Aikhenvald, 2006; Foley, 2010). Though generally absent in European languages, SVCs are common cross-linguistically. Some languages have a particularly high rate of SVC use and these are called serializing languages. Languages with no SVCs, such as English, are called non-serializing languages.

It has been claimed that SVCs always refer to conceptualizations of a single event (Aikhenvald, 2006; Comrie, 1995). Take the examples below, from Avatime, a Ghana-Togo Mountain language from the Kwa branch of the Niger-Congo language family. The SVC in example (1a) describes what appears to be a single event: a man cutting firewood with the axe he picked up for that purpose. In contrast, the two Avatime simple, single verb clauses in example (1b) describe a less integrated scene of a man picking up an axe (maybe not with the immediate or sole purpose of cutting firewood) and then cutting firewood (not even necessarily with the axe just mentioned).

1. (a) *A-kò kàwɛ-à tsǎ ònyì-nè.*
3S¹-take axe-DEF cut firewood-DEF
'He cut the firewood with the axe.'
- (b) *A-kò kàwɛ-à. A-tsǎ ònyì-nè.*
3S-take axe-DEF 3S- cut firewood-DEF
'He took the axe. He cut the firewood.'

While there is a strong feeling among linguists that SVCs should – and do – refer to single conceptual event units (Aikhenvald, 2006; Comrie, 1995; Durie, 1997), the relationship has not been directly tested.

The best evidence for this relationship, to date, comes from a study conducted by Givón (1990, 1991). He tested conceptual event units by investigating the production process. Speakers pause when they are encoding the next unit of speech (Goldman-Eisler, 1968). Givón thus took pauses in speech as an indication of conceptual cohesion: speech that was encoded together, and so between pauses, was taken to refer to single event units. The frequencies of speech internal pauses in different clause types were compared across three languages of Papua New Guinea (Kalam, Tairora and Tok Pisin), which use verb serialization to different degrees. Givón found that pauses were no more frequent within SVCs than they were within simple clauses with a single verb. From this, he concluded that SVCs and simple clauses both refer to single conceptual event units (contra Pawley, 1987). Note that this study only tests chunking at the linguistic level. It does not provide evidence about cognitive event segmentation. To do that, an independent test of conceptual event structure is required.

The present study aims to conduct just such an independent test. It focuses on placement events in the

serializing language Avatime, to test the following two hypotheses: 1) that SVCs correspond to single conceptual events and 2) that differences in linguistic descriptions of events correlate with differences in conceptual event units.

In Avatime, most placement actions, like putting a cup on a table, or a banana in a basket, must be described using both a take verb and a put verb in an SVC², as in example (2). The grammar of the language requires speakers to encode the taking part as well as the placing part of the event, even if the person only saw the placing. Note that it is logically necessary for an object that is being placed to have been taken at some earlier point in time. So, it is not as strange as it may at first appear for a language to require the preceding taking event to also be encoded. The same construction is also used to describe cases when both the taking and placing events are seen. So an alternative interpretation of (2) is: *S/he took the banana and put it in the basket.*

2. *A-kò kòranti-ɛ kpe ní kàsɔ-yà mè*
3S-take banana-DEF put LOC basket-DEF inside
'S/he put the banana into the basket.'

There is a small set of placement events that are described without a take-put SVC. These exceptional events include putting an article of clothing or jewelry on a body part (in its canonical location), and pouring liquids. These are described using either a put verb in a simple clause (3) or a put verb combined with a pouring manner verb in an SVC (4). It is strongly dispreferred to describe such actions using an SVC with a take verb.

3. *A-kpe likùto-lè*
3S-put hat-DEF
'S/he put the hat on.'
4. *E-nyi kùni-ò kpe ní kèzi-à mè*
3S-pour water-DEF put LOC bowl-DEF inside
'S/he poured the water into the bowl.'

The patterns of placement event descriptions in Avatime, and the claim that SVCs describe single conceptual events, lend themselves to experimental testing. Previous research has shown that people mentally fill in parts of event units that they have not actually seen (Strickland & Keil, 2011). We can build on this finding to test whether Avatime speakers treat take-put episodes as single event units. Specifically, if Avatime speakers see a videoclip showing a

¹ Abbreviations used: 3 '3rd person', DEF 'definite', LOC 'locative', S 'singular', ` 'low tone', ^ 'high tone', mid tone is unmarked.

² As with many languages with this type of construction, the take verb acts like an object marker and allows the two objects (thing placed, and location where placed) to be expressed (Lord, 1993). However, unlike some languages, in Avatime the take verb still maintains much of its original lexical semantics in these cases. Different take verbs will even be used to mark differences in the type of taking done.

general placement action, which they would describe using a take-put SVC, they should be more likely to falsely recognize a corresponding taking action. In contrast, if Avatime speakers see a videoclip showing a placement action, which they would not describe with a take-put SVC, such as putting on clothing or pouring a liquid, they should not falsely recognize a corresponding taking action.

To control for the possibility that putting events and their corresponding taking events are generally more cohesive than the donning of clothing or pouring of liquids and their corresponding taking events, we tested a control group of English speakers. English speakers describe general placement events with a single put verb which takes the thing moved as the object and the location as a prepositional phrase. For instance, *She put the book on the table*. The pouring of liquids is described using the same structure as general placements, but the verb is specific to pouring. For instance, *He poured water into the glass*. The putting on of clothing and jewelry is described using essentially the same structure but the location is often not expressed. For instance, *She put the necklace on*. There are no cases where the grammar of English requires the corresponding taking event to also be encoded. Hence, English speakers are not predicted to have differences in false recognition rates to these take events.

Methods

Participants

Thirty-four native speakers of Avatime, aged 11-16 (mean 14.1 years), were recruited at Vane Junior High School, Ghana. Four Avatime speakers were tested but excluded due to technical difficulties or for consistently answering either yes or no for all items. Thirty-three native speakers of English, aged 11-17 (mean 14.2 years), were recruited in the Blue Mountains and Sydney, NSW, Australia.

All Avatime speakers were fluent in Ewe and English and 11 additionally spoke Twi. One English speaker was also fluent in German, two spoke Spanish, one fluently and the other moderately. Of the remaining English speakers, 9 were completely monolingual and 21 had very limited knowledge of another language (French, German, Italian, Japanese, Korean or Latin).

Materials

80 paired putting and taking events were filmed in a single location inside the Max Planck Institute, Nijmegen. They were acted out by two Dutch university students, one male and one female. Each videoclip lasted 3-4 seconds.

A paired putting and taking episode showed the same actor removing an object from one location and placing it in another. For instance, in Figure 1(a) a man takes a banana from the shelf and places it on a plate, in Figure 1(b) a woman takes a necklace from a bag and places it on her

neck. Across episodes, the camera angle and position of the actor in the room were kept constant.



Figure 1: Sample frames from the two videoclips (a) ‘man takes banana from shelf’ and ‘man puts banana on plate’; (b) ‘woman takes necklace from bag’, and ‘woman puts necklace on.’

Objects and locations were selected so as to be familiar to both Avatime and English speakers. The source location of the taking event was always different from the goal location of the putting event. Across episodes, the object, locations, position of the actors, and camera angle varied.

Of the 40 episodes, half had general placement events of the type described using take-put SVCs in Avatime, while the other half did not (the donning of clothing and pouring). Descriptions of the items by Avatime participants at the end of the experiment confirmed this distinction: The placement events in the SVC category were described using take-put serial verb constructions 96.2% of the time ($SD = 1.8$). The placement events in the Non-SVC category were described using take-put serial verb constructions 6.5% of the time ($SD = 1.7$). For ease of reference, both the putting and taking events in an episode will be referred to as either SVC or Non-SVC according to the type of putting event.

Design

The 40 put and take episodes, resulted in 80 individual items, each consisting of a sole put event or take event. The 80 items were divided into two sets: only one part of a put-take episode featured in each set. Pilot testing with Avatime speakers showed that remembering all 40 learning items in one go was too difficult, so testing was divided into two blocks. In each block of a given set, there were 5 SVC put events, 5 Non-SVC put events, 5 SVC take events, and 5 Non-SVC take events. Blocks were counterbalanced across participants. Within each block, items appeared in one of four random orders.

Procedure

Participants were asked to watch a series of videoclips and to remember them as best they could. They were told that

they would later be shown more videos, some exactly the same as the ones they had seen and some different, and that their task was to tell the experimenter which videoclips were the same and which were not.

Participants watched videoclips one at a time. The videoclips were separated by a black screen lasting 1 second. After the learning phase, there was a 5 minute distraction task unrelated to the experiment. Participants were then tested for their memory of the 20 videoclips they had just seen, plus their 20 unseen counterparts. So, if a participant saw a girl put on a necklace in the learning phase, they now, in the testing phase, also saw the girl taking the necklace out of the bag. Participants indicated whether each event was the same or different to the events they had seen previously. After finishing testing for the first block, participants saw the second block of 20 items and were tested for memory of those as described above.

After completing the memory experiment, participants viewed all the videoclips again and were asked to describe "what the person did".

Avatime instructions were translated by a native Avatime speaker fluent in English in consultation with the experimenter. Instructions and responses were given verbally in the participant's native language. Responses were recorded using an Olympus LS-10 flash recorder with a headset microphone.

Participants were tested individually and the same procedure was used for English and Avatime participants. The whole experiment lasted approximately 45 minutes.

Results

Responses to seen and new items were analyzed separately using 2 construction-type (SVC or Non-SVC) x 2 event-type (put or take) x 2 language (Avatime or English) x 2 block order (AB or BA) mixed ANOVAs, with construction and event type being within-participant factors, and language and block order between-participant factors. The dependent variable was the number of reported recognitions. Block order was not significant for seen ($F(1,59) = 0.62, p = 0.43, \eta_p^2 = 0.01$) or new items ($F(1,59) < 0.01, p = 0.94, \eta_p^2 = 0.01$), so we collapsed over this factor.

We first tested whether participants were able to correctly recognize the items they had seen. The overall accuracy was 80.7% for Avatime speakers and 83.6% for English speakers. The difference between language groups was not significant, $F(1, 61) = 0.92, p = 0.34, \eta_p^2 = 0.02$. There was a main effect of event-type, $F(1,61) = 9.20, p < 0.01, \eta_p^2 = 0.13$. Putting events were remembered more accurately ($M = 8.50$) than taking events ($M = 7.92$). There was also a main effect of construction, $F(1,61) = 9.81, p < 0.01, \eta_p^2 = 0.14$, and a just significant interaction between construction-type and language $F(1,61) = 3.94, p = 0.05, \eta_p^2 = 0.06$. English speakers remembered Non-SVC events more accurately ($M = 8.73$) than SVC events ($M = 7.99$). Avatime speakers showed no difference in recognition between

previously seen SVC events ($M = 7.98$) and Non-SVC events ($M = 8.15$). There were no other interactions.

Our hypothesis concerned false recognitions to previously unseen or new items. It was predicted that there would be a three-way interaction between construction-type, event-type and language. Avatime speakers would have more false recognitions for taking events if the corresponding put event was one that they would describe using a take-put serial verb construction. English speakers should have the same rates of false recognition for SVC and Non-SVC type events. There was no statistically significant 3-way interaction, $F(1,61) = 0.01, p = 0.92, \eta_p^2 < 0.01$. However, there was a main effect of language, $F(1,61) = 14.34, p < 0.01, \eta_p^2 = 0.19$. Avatime speakers, in general, had more false recognitions ($M = 2.83$) than English speakers ($M = 1.58$). There was also a main effect of construction-type $F(1,61) = 4.36, p = 0.04, \eta_p^2 = 0.07$. SVC events had more false recognitions ($M = 2.37$) than Non-SVC events ($M = 2.04$). More interestingly, there was a significant interaction between language and construction-type, $F(1,61) = 4.36, p = 0.04, \eta_p^2 = 0.07$, see Figure 2. Avatime speakers had more false recognitions for SVC type events in general ($M = 3.17$) than for Non-SVC type events ($M = 2.50$). English speakers, on the other hand, had the same false recognition rates for SVC ($M = 1.58$) and Non-SVC events ($M = 1.58$). This suggests that Avatime speakers remember events described with SVCs differently to those which are not; and that this effect is not due to properties of the events themselves, since English speakers fail to show a difference across these event types.

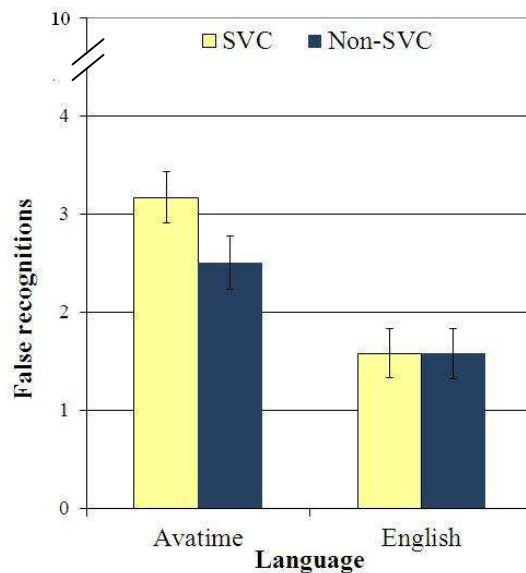


Figure 2: Average false recognitions as a function of language and construction type.

Finally, there was an unpredicted interaction between language and event type, $F(1,61) = 4.51, p = 0.04, \eta_p^2 =$

0.07). Avatime speakers showed more false recognitions for put events ($M = 3.08$) than take events ($M = 2.58$). In contrast, English speakers had slightly more false recognitions for take events ($M = 1.71$) than for put events ($M = 1.44$). There were no other significant interactions.

Discussion

Avatime speakers, but not English speakers, displayed more false recognitions for put and take events from SVC episodes than from the equivalent events in non-SVC episodes. This is consistent with the suggestion that language may play a role in conceptual event structure. Avatime speakers appear to construct a single conceptual event unit that includes the taking and putting event segments precisely when the putting event is one that they would describe with a take-put SVC.

Our initial prediction was that false recognitions would occur only with the take part of the episode. This was because it is the put event that determines whether or not an SVC is used, and it is this use of the SVC which is predicted to determine whether or not the take action is included in the event unit. For example, picking up a necklace should only be combined with its corresponding putting event if the putting event is something like putting the necklace in a bag, not putting it around your neck. Our results show, however, that Avatime speakers falsely recognize both the take and put parts of the SVC episodes regardless of which part they saw first. This indicates that as soon as both parts have been seen and understood to form an SVC type episode, Avatime speakers join the taking and putting actions together into a single conceptual event unit. Familiarity with either part then spreads to the unit as a whole, resulting in false recognition of the unseen part, be it a putting or a taking action.

These results show a correlation between conceptual event units and linguistic structure, but from these results alone we cannot say whether language influences conceptual structure, conceptual structure influences language, or whether some third factor is involved. There is some ancillary evidence that language may play a causal role here. For example, Trueswell and Papafragou (2010) found that people under high cognitive load directed attention to event elements considered important in their language; while Zacks et al. (2001b) found that there was greater alignment between fine- and coarse-level segmentation when speakers described events while they segmented them, rather than describing them later. To determine whether linguistic encoding is critically involved in this experiment further study would be needed.

The link between SVCs and single conceptual event units has often been suggested as a definitional criterion for SVCs (Aikhenvald, 2006; Comrie, 1995; Durie, 1997). This paper provides the first language external evidence in favor of this often cited relationship. However, SVCs were only compared with sets of separate clauses. To determine

whether or not the link between SVCs and single conceptual event units is useful as a definitional criterion, SVCs should be compared to other types of complex clauses as well. This paper has shown that testing recognition memory is a viable method for investigating the relationship between conceptual event units and syntactic structures.

In addition to the main result discussed above, we found three other effects. 1) Putting events were remembered more accurately than taking events by speakers of both languages. This is in line with predictions based on the asymmetry of sources and goals in motion events (Regier & Zheng, 2007; Papafragou, 2010) and research concerning put and take lexicons (Narasimhan, Kopecka, Bowerman, Gullberg, & Majid, in press; Regier, 1995). 2) Avatime speakers displayed more false recognitions for put events than for take events while English speakers showed the reverse pattern. This is not immediately interpretable and will require further investigation. 3) English speakers remembered both putting and taking events from Non-SVC episodes more accurately than those from SVC episodes. This shows that there may be differences between the episode types which are noticeable by English speakers, even though they do not use SVCs. It seems likely that actions involving clothing as well as pouring actions could be more salient than general taking and putting actions. Although English speakers were sensitive to the distinction between episode types, they nevertheless performed equivalently with respect to memory for new events. So although there may be differences between SVC and Non-SVC episodes these differences alone cannot predict our final results.

Conclusion

This study provides the first evidence for the often claimed connection between serial verb constructions and single conceptual event units. It demonstrates that event elements grouped together in language are grouped together as conceptual event units: Avatime speakers conceptualize a take-put episode as a single event unit exactly when the placement event is one they would describe with a take-put SVC but not if it is from a Non-SVC. English speakers, on the other hand, do not distinguish the two types of events in their syntax, nor do they demonstrate greater event cohesion for the events described by take-put SVCs in Avatime. Thus, speakers' event conceptualisations parallel the linguistic structures used to describe those events.

References

- Aikhenvald, A. (2006). Serial verb constructions in typological perspective. In A. Aikhenvald & R. Dixon (Eds.), *Serial Verb Constructions: A Crosslinguistic Typology*. Oxford: Oxford University Press.
- Avrami, J., & Kareev, Y. (1994). The emergence of events. *Cognition*, 53, 239-261.

- Comrie, B. (1995). Serial verbs in Haruai (Papua New Guinea) and their theoretical implications. In J. Bouscaren, J. Franckel, & S. Robert (Eds.), *Langues et langage: Problèmes et raisonnement en linguistique, mélanges offerts à Antoine Culioli*. Paris: University Presses of France.
- Durie, M. (1997). Grammatical Structures in Verb Serialization. In A. Alsina, J. Bresnan, & P. Sells (Eds.), *Complex Predicates*. Stanford, CA: CSLI Publications.
- Fausey, C. M., & Boroditsky, L. (2011). Who dunnit? Cross-linguistic differences in eye-witness memory. *Psychonomic Bulletin & Review*, 18(1), 150–157.
- Filipović, L. (2011). Speaking and remembering in one or two languages: bilingual vs. monolingual lexicalisation and memory for motion events. *International Journal of Bilingualism*, 15(4), 466–485.
- Finkbeiner, M., Nicol, J., Greth, D., & Nakamura, K. (2002). The role of language in memory for actions. *Journal of Psycholinguistic Research*, 31(5), 447–457.
- Foley, W. A. (2010). Events and serial verb constructions. In B. Baker & M. Harvey (Eds.), *Complex Predicates: Cross-Linguistic Perspectives on Event Structure*. Cambridge: Cambridge University Press.
- Gennari, S. P., Sloman, S., Malt, B., & Fitch, T. (2002). Motion events in language and cognition, *Cognition*, 83, 49–79.
- Givón, T. (1990). ‘Verb serialization in Tok Pisin and Kalam: a comparative study of temporal packaging.’ In J. Verhaar (Ed.) *Melanesian Pidgin and Tok Pisin*. Amsterdam: Benjamins.
- Givón, T. (1991). Serial verbs and event cognition in Kalam: an empirical study of cultural relativity. In C. Lefebvre (Ed.), *Serial verbs: grammatical, comparative and universal grammar*. Amsterdam: John Benjamins.
- Goldman-Eisler, Frieda. 1968. *Psycholinguistics: experiments in spontaneous speech*. New York: Academic Press.
- Kersten, A. W., Meissner, C. A., Lechuga, J., Schwartz, B. L., Albrechtsen, J. S., & Iglesias, A. (2010). English Speakers Attend More Strongly Than Spanish Speakers to Manner of Motion When Classifying Novel Objects and Events. *Journal of Experimental Psychology: General*, 139(4), 638–653.
- Kurby, C. A., & Zacks, J. M. (2008). Segmentation in the perception and memory of events. *Trends in Cognitive Sciences*, 12(2), 72–79.
- Lord, C. (1993). *Historical change in serial verb constructions*. Typological studies in language. Amsterdam: Benjamins.
- Loucks, J., & Pederson, E. (2008). Linguistic and non-linguistic categorization of complex motion events. In J. Bohnemeyer & E. Pederson (Eds.), *Event representation in language and cognition*, Language Culture and Cognition. Cambridge: Cambridge University Press.
- Narasimhan, B., Kopecka, A., Bowerman, M., Gullberg, M. & Majid, A. (in press). Putting and taking events: a cross-linguistic perspective. In A. Kopecka & B. Narasimhan (Eds.), *Events of ‘putting’ and ‘taking’: a cross-linguistic perspective*. Amsterdam: John Benjamins.
- Newton, D., Engquist, G., & Bois, J. (1977). The objective basis of behaviour units. *Journal of Personality and Social Psychology*, 35(12), 847–862.
- Papafragou, A. (2010). Source-goal asymmetries in motion representation: Implications for language production and comprehension. *Cognitive Science*, 34, 1064–1092.
- Papafragou, A., Massey, C., & Gleitman, L. (2002). Shake, rattle, “n” roll: the representation of motion in language and cognition. *Cognition*, 84, 189–219.
- Pawley, A. (1987). Encoding events in Kalam and English: different logics for reporting experience. In R. Tomlin (Ed.), *Coherence and grounding in discourse*, Typological Studies in Language. Amsterdam/Philadelphia: John Benjamins.
- Regier, T. (1995). A model of the Human Capacity for Categorizing Spatial Relations. *Cognitive Linguistics*, 6, 63–88.
- 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 distort memory in a matter of seconds. *Cognition*, 121, 409–415.
- Trueswell, J., & Papafragou, A. (2010). Perceiving and remembering events cross-linguistically: Evidence from dual-task paradigms. *Journal of Memory and Language*, 63, 64–82.
- Wolff, P., Jeon, G.-H., & Li, Y. (2009). Causers in English, Korean, and Chinese and the individuation of events. *Language and Cognition*, 1(2), 167–196.
- Zacks, J. M. (2004). Using movement and intentions to understand simple events. *Cognitive Science*, 28, 979–1008.
- Zacks, J. M., Braver, T. S., Sheridan, M. A., Donaldson, D. I., Snyder, A. Z., Ollinger, J. M., Buckner, R. L. & Raichle, M. E. (2001a). Human brain activity time-locked to perceptual event boundaries. *Nature Neuroscience*, 4, 651–655.
- Zacks, J. M., Tversky, B., & Iyer, G. (2001b). Perceiving, remembering and communicating structure in events. *Journal of Experimental Psychology: General*, 130(1), 29–58.
- Zacks, J. M., Speer, N. K., Swallow, K. M., Braver, T. S., Reynolds, J. R. (2007). Event perception: a mind-brain perspective. *Psychological Bulletin*, 133(2), 273–293.