

Generalisable patterns of gesture distinguish semantic categories in communication without language

Gerardo Ortega (gerardo.ortega@mpi.nl) &
Ash Özyürek (asli.ozyurek@mpi.nl)

Centre for Language Studies, Radboud University and
Max Planck Institute for Psycholinguistics,
Wundtlaan 1, 6525XD Nijmegen, The Netherlands

Abstract

There is a long-standing assumption that gestural forms are geared by a set of modes of representation (acting, representing, drawing, moulding) with each technique expressing speakers' focus of attention on specific aspects of referents (Müller, 2013). Beyond different taxonomies describing the modes of representation, it remains unclear what factors motivate certain depicting techniques over others. Results from a pantomime generation task show that pantomimes are not entirely idiosyncratic but rather follow generalisable patterns constrained by their semantic category. We show that a) specific modes of representations are preferred for certain objects (acting for manipulable objects and drawing for non-manipulable objects); and b) that use and ordering of deictics and modes of representation operate in tandem to distinguish between semantically related concepts (e.g., "to drink" vs "mug"). This study provides yet more evidence that our ability to communicate through silent gesture reveals systematic ways to describe events and objects around us.

Keywords: pantomime, gesture, action/object distinction, modes of representation, iconicity

Introduction

Speakers have at their disposal several strategies to represent a referent with their gestures. If referring to a glass, for example, a speaker may choose to produce a gesture representing how it should be held, may describe its outline or perhaps would depict its cylindrical volume in a three-dimensional space. Despite most research focusing on the relationship between gesture and speech, the mechanisms responsible for the specific form that iconic gestures adopt remains largely unknown. It is unclear for example whether the physical characteristics of the referent may play a role in gestural production and also whether people develop systematic strategies to make distinctions between different semantic categories (i.e., between actions and objects).

Research investigating the mechanisms responsible for gestural production is paramount to better understand the cognitive system that allows efficient communication through the manual channel. A powerful tool to do so is by investigating communication in the absence of speech. By exploring individual's communicative strategies through silent gestures, we open a window into humans' internal representations as well as into our capacity to convey information in the most effective way. Studies have shown

that after stripping communication from a conventionalised language, speakers of linguistically diverse languages converge in the same strategies to express information through their silent gestures (Goldin-Meadow, So, Özyürek, & Mylander, 2008a; Hall, Mayberry, & Ferreira, 2013; Langus & Nespors, 2010; Özçalışkan, Lucero, & Goldin-Meadow, 2016). It remains an empirical question whether individuals develop systematic strategies to make distinctions across different semantic domains when expressing concepts in silent gesture. It is possible that gestures will be highly idiosyncratic and thus will be executed in different ways. Alternatively, it is possible that individuals' knowledge of the world may interact with the available techniques of gestural representations and as a consequence, their silent gestures will display a high degree of systematicity.

Communication without language and what it reveals about the mind

Studies employing descriptive and empirical methods have found that language is a very important factor that shapes many of our cognitive processes (Carroll, 1956; Flecken, Von Stutterheim, & Carroll, 2014; Majid & Burenhult, 2014; Sapir, 1921). Despite the significant differences observed in the behaviours of speakers of linguistically diverse languages, most effects disappear when we look at their silent gestures. To date there are now several studies showing that, regardless of their language, speakers converge in the strategies used to represent events in this mode of manual communication.

One of the most studied effects is the sequencing of events and the agents that perform them (i.e., word order) during manual communication without speech. Languages vary considerably in the way they order the constituents of a sentence. For instance, English favours a Subject-Verb-Object sequence while Turkish prefers a Subject-Object-Verb. One would expect that when expressing an event through silent gesture, English speakers would favour an S-V-O order while Turkish speakers would follow an S-O-V ordering, as they would in their native language. However, this is not the case as speakers of both languages coincide in an S-O-V sequencing of pantomimes. This effect has been proven to be quite robust as has been replicated in multiple occasions with speakers of very diverse languages (Goldin-Meadow et al., 2008a; Hall et al., 2013; Langus & Nespors, 2010).

A similar effect has been observed in the description of motion events. English is a satellite-framed language and as such encodes information of manner (e.g., run) and path (e.g., towards) in a clause. Turkish, in contrast, is a verb-framed language often dropping information about the manner and expressing information about the path only (e.g., enter). When asked to express motion events, the co-speech gestures produced by speakers convey the same information as in their speech (i.e., English speakers express manner and path in their gestures and Turkish speakers express only path) (Kita & Özyürek, 2003; Özyürek, Kita, Allen, Furman, & Brown, 2005). However, the same speakers default to a single strategy when they are expressing the same information in silent gesture. That is, when speakers are in this mode of communication, they no longer align their gestural strategies with the information conveyed in their speech but rather they resort to a strategy that is shared across speakers of different languages (Özçalışkan et al., 2016).

Together these studies demonstrate that while language is an important factor that governs many of our cognitive behaviours, communication through silent gesture overrides any linguistic influence and generates communication with unique properties shared across speakers of different languages. The similarities in patterns observed in pantomimes in different domains (i.e., word order, motion events) have been interpreted as silent gesture being a window onto our internal representations (Özçalışkan et al., 2016) and to our capacity to package information in the most communicatively effective manner (Goldin-Meadow, So, Özyürek, & Mylander, 2008b; Hall et al., 2013).

Another domain that has the potential to reveal a high degree of systematicity in silent gestures is the representation of concepts across different semantic categories. If silent gestures are also prone to a high degree of systematicity across different individuals, it is possible to expect generalisable patterns in the modes of representation used in specific semantic domains. This possibility has not yet been explored and remains an empirical question.

Manual modes of representation

There is general consensus that gestures may adopt at least four modes of representation. *Acting* (or handling) denotes how an object is manipulated (e.g., supination of a closed fist for 'key'). *Representing* (or instrument) uses the hand to recreate the form of an object (e.g., an extended index finger to represent a 'toothbrush'). *Drawing* (or tracing) describes the outline of a referent (e.g., two index fingers tracing a square to represent a 'window'). *Moulding* depicts the three-dimensional characteristics of an object (e.g., cupped hands describing the shape of a 'vase') (Müller, 2013). Beyond different taxonomies describing the modes of representation that gestures can adopt, it remains unclear what factors motivate certain techniques over others.

Müller (2013) proposes that during their narrations, speakers express in both the spoken and manual channel the relevant aspects of a scene or event. Importantly, the form of the gestures will depend primarily on the speakers' focus

of attention or what is considered to be the relevant information. If a speaker, for instance, wants to emphasise the specific way to handle an object he will use a depicting technique that expresses this information. If, in contrast, the main focus of his narrations is the form of an object he will produce a gesture in which the hand configurations represent the shape of the referent. If the focus of his narration is the three-dimensional form of an object he will probably describe the volume of the referent in space. In other words, iconic gestural forms express visual information specifically tailored to describe a unique event.

Recent evidence has shown, however, that people's gestures are not entirely dependent on speakers' focus of attention but rather are constrained by the affordances of the referent. Using a referential paradigm, speakers were asked to describe different objects from a visual prompt. Stimuli were categorised as having high or low affordances (i.e., the degree to which objects allow to be manipulated). The analysis of participants' co-speech gestures showed that objects with high affordances (e.g., wine glass) were often represented through an acting strategy. In contrast, items with low affordances (e.g., sink) were described using a drawing strategy (Masson-Carro, Goudbeek, & Krahmer, 2015).

This study suggests that the form of iconic gestures, at least in co-occurrence of speech, are somewhat constrained by the affordances of the referent. A question that remains unanswered is whether the different representational techniques are also deployed systematically depending of the type of referent (i.e., manipulable and non-manipulable) in silent gesture. Further, it remains an empirical question whether gesturers will resort to a different strategy to make distinctions between actions and objects, when there is no speech to aid marking this differentiation.

Action-object distinctions in the visual modality

Most of the investigations attempting to understand how the manual channel makes distinctions between actions and objects come from sign language research. The first studies exploring this issue found that in American Sign Language (ASL) pairs like HAMMER and TO-HAMMER are formally marked in the movement of the sign. While actions have a continuous movement (e.g., TO-HAMMER), objects have a restrained, repeated movement (e.g., HAMMER) (Supalla & Newport, 1986). It has also been reported that signs for actions tend to use a larger signing space and are less marked through non-manual features (i.e., mouthings) than signs for objects. These characteristics are not universal given that different sign languages use different formal features to mark these distinctions as has been documented for Australian Sign Language (Auslan) (Johnston, 2001) and Russian Sign Language (RSL) (Kimmelman, 2009). Interestingly, emerging sign languages do not seem to exhibit a clear mechanism to make such distinctions. Al-Sayyid Bedouin Sign Language (ABSL) is an emerging sign language that is gradually developing mechanisms to mark these distinctions more overtly.

More recently, studies have shown that an effective mechanism to make distinctions between actions and objects in the absence of speech is through the representation of the referent with different depicting strategies (Padden et al., 2013). When users of different sign languages were asked to represent vignettes of objects and agents manipulating objects (actions), one can observe that there is systematicity in their patterning of use. ASL and ABSL signers tend to depict actions through acting depictions and objects through representing depictions. This distinction is language-specific because users of an unrelated sign language (New Zealand Sign Language) favour the opposite patterns (i.e., acting for nouns and representing for verbs). Interestingly, this study also revealed that when asked to perform the same task, hearing people always converge in the same strategy to represent the referent. That is, silent gesturers predominantly favour an acting strategy for all their gestural depictions (actions and objects alike) with few instances of representing depiction (Padden et al., 2013; Padden, Hwang, Lepic, & Seegers, 2015). The notion of *patterned iconicity* postulates that sign languages may differ in the strategy used to make action-object distinctions, but they systematically exploit the available depicting possibilities (modes of representation) to make such differentiations. In contrast, gesturers default to the same strategy (acting) for both semantic categories.

These studies show that sign languages alter the phonological structure of the sign or their mode of representation to distinguish actions from objects while pantomime overall defaults to the acting technique. A shortcoming of these studies is that they have limited their observations to the techniques of depiction only within two semantic domains (tools and actions with tools). Given that gestures are holistic units without sub-lexical components (McNeill, 1992) individuals are unlikely to modify their gestures' kinematics to make semantic distinctions in a similar way as signs. It is possible, however, that they may deploy additional strategies and bodily cues such as pointing, showing, eye-gaze, and sequences of gestures to mark such distinctions.

The Present Study

In the present study we turn to the production of silent gesture to investigate whether actions vs. objects and their affordances (i.e., manipulable vs. non-manipulable) modulate the strategy used by speakers to represent a referent manually. More specifically, we ask 1) do individuals use a specific depicting strategy for each semantic category; and 2) what are the additional strategies implemented to make semantic distinctions. To that end, we implemented a pantomime generation task to a group of Dutch speakers and described the gestures produced for a list of words, the strategy they used for each semantic category, and the strategies they implemented to differentiate actions from objects.

Methodology

Participants

Twenty native speakers of Dutch (10 females, age range: 21-46, mean: 27 years) living in the area of Nijmegen, the Netherlands took part in the study.

Procedure

Participants were tested individually in a quiet room with two cameras from different angles recording their gestures. They were told that the task consisted of generating a sign or gesture that conveyed exactly the same meaning as the word on the screen. They were explicitly told two rules: they were not allowed to speak or say the target word; and they could not point at any object present in the room (e.g., pointing at the table or at a wall). They were also told that their videos were going to be shown to another participant who would have to guess the meaning of their gesture.

The stimuli consisted of a total of thirty words from three semantic categories: 10 actions with an object (e.g., to phone, to smoke), 10 manipulable objects (e.g., telephone, lighter), and 10 non-manipulable objects (e.g., pyramid, floor). Words were presented in black font on a white background in a different randomised list for each participant. We decided against presenting the stimulus materials with a visual cue so as to avoid prompting participants. Each trial started with a fixation cross in the middle of the screen for 500 ms and this was followed by the word participants had to represent with their gestures. The target word remained on the screen for 4000 ms during which participants had to come up with their gestural depictions. We limited the allowable time for gestural production so as to force participants to produce their most intuitive responses. Participants' renditions were video recorded and later annotated using the software ELAN (Lausberg & Sloetjes, 2009).

Coding and data analysis

For each target word, participants were observed to produce one gesture or sequences of gestures to depict the referent. Following a strict coding criteria, all gestures produced for each item were annotated. Each gesture or sequences of gestures would consist minimally of a preparation phase, a stroke and a (partial/full) retraction. Once all the gestures were isolated, we classified them according to their mode of representation. Adapting the taxonomy developed by Müller (2013), we categorised each gesture as follows: *Acting* if the gesture represented how the referent is manipulated; *representing* if the hands were used to recreate the form of an object; and *drawing* if participants used their hands to describe the outline or the three-dimensional characteristics of an object (note that we collapsed Müller's drawing and moulding categories into one). Aside from these modes of representation, we also included the category *deictic* which consisted of pointing, showing and/or ostensive eye-gaze to elements of the gesture (see Figure 1 for examples). This is

not a mode of representation *per se* but we decided to include this category in the analysis given the high prevalence of this strategy to make semantic distinctions.

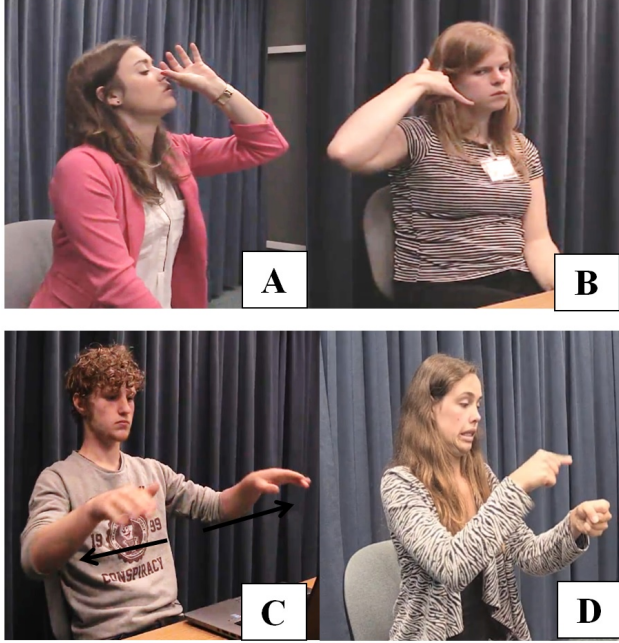


Figure 1: Examples of participants using different modes of representation (Müller, 2013) and deictics. A) The *acting* technique represents how an object is manipulated (e.g., ‘to drink’); B) in *representing* the hands recreate the shape and form of an object (e.g., ‘telephone’); C) *drawing* traces the outline or three dimensional features of an object (e.g., ‘table’), and D) deictics like points, showing or eye-gaze are used to highlight features of a gesture (e.g., pointing at a fist holding an imaginary mug).

After the whole dataset was annotated and categorised according to the gestures’ mode of representation, we calculated the number of gestures produced per item per participant across the three semantic categories (actions with an object, manipulable objects, and non-manipulable objects). Then, we calculated the proportion of the different modes of representation per semantic category; and finally, we calculated the proportion of deictics used in the three different categories.

Results

We calculated the number of gestures produced per item per participant across the three semantic domains. We found that actions with objects elicited the least number of gestures (range: 1-2 gestures; mean: 1.1 gestures, $SD = 0.12$), followed by non-manipulable objects (range 1-4 gestures; mean: 1.49 gestures, $SD = 0.33$), and manipulable objects elicited the highest number of gestures (range: 1-4 gestures; mean: 1.77 gestures, $SD = 0.39$). On the arcsine

transformed values, a one-way ANOVA revealed a significant difference in the number of gestures produced for each semantic category $F(2,38) = 40.14$, $p < 0.0001$, $\eta^2 = 0.679$. Pairwise comparisons after Bonferroni corrections revealed that the number of gestures produced for each category is significantly different from one another. Actions with objects was significantly lower than manipulable objects [$t(19) = 8.39$, $p < 0.0001$] and non-manipulable objects [$t(19) = 5.24$, $p < 0.0001$]. Non-manipulable objects elicited significantly fewer gestures than manipulable objects [$t(19) = 3.98$, $p < 0.001$].

Table 1 shows the proportion of instances in which participants produced a single vs. multiple gestures to describe a referent across conditions. After removing passes and wrong targets (e.g., the target word ‘to sieve’ *zeven* often elicited the gesture ‘seven’ *zeven*) we can see that the vast majority of actions with tools elicited a single gesture, which often depicted how the action is executed (see Figure 2 for the gesture ‘to-drink’). In contrast, manipulable objects were predominantly depicted with more than one gesture (see Figure 2 for the gesture ‘lighter’). Non-manipulable objects have a split with an almost equal proportion of items being depicted with a single or multiple gestures.

Table 1: Proportion of concepts depicted with a single or multiple gestures across conditions (N=200)

	Action with objects	Manipulable objects	Non-manipulable object
Single	0.91	0.35	0.45
Multiple	0.09	0.61	0.46
Wrong	0.01	0.05	0.10
	1.00	1.00	1.00

In order to explore whether the gestural forms are restricted by the affordances of the referent, we looked at the mode of representation used across semantic categories (acting, representing or drawing). We focused on the instances in which a single gesture had been elicited to represent a concept. Table 2 shows that actions with objects and manipulable objects use predominantly the acting strategy (i.e., how an object is used). In contrast, non-manipulable objects resort more often to a drawing technique.

Table 2: Proportion of concepts depicted with different modes of representations across conditions (one-gesture depictions only)

	Action w/object (N=181)	Manipulable object (N=70)	Non-manipulable object (N=89)
Acting:	0.86	0.81	0.20
Drawing	0.00	0.01	0.57
Representing	0.14	0.17	0.22
	1.0	1.0	1.0

Finally we looked at the instances in which participants included a deictic (i.e., pointing, showing, ostensive eye-gaze) to refer to a specific feature of their gesture. We observed that out of the whole data set, actions with objects (N=199) and non-manipulable objects (N=180) elicited a small proportion of deictics (0.04 and 0.09 respectively). In contrast, manipulable objects (N=191) elicited significantly more deictics (0.25). For instance, to represent ‘lighter’ participants would perform the action of lighting a cigarette and then they would point at an imaginary lighter (See Figure 2). Similarly, for ‘toothbrush’ they would pretend to be brushing their teeth with a handling handshape (i.e., closed fist) and then they would raise it and show it to the camera.

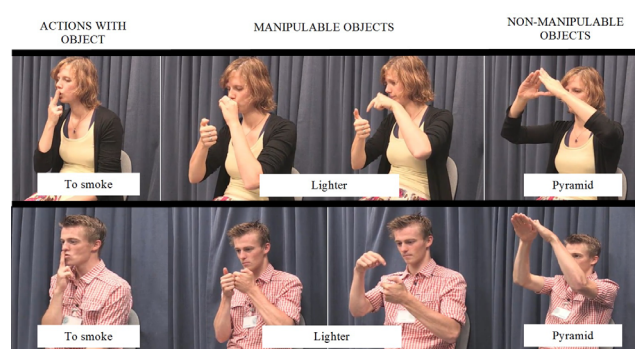


Figure 2: Examples of the modes of representation in pantomimes across different semantic categories. Actions with object were depicted with a single gesture representing how to manipulate an object (e.g., ‘to smoke’); manipulable objects were represented with the gesture of an action followed by a deictic (e.g., lighter was depicted with a pantomime of lighting a cigarette and then pointing at an imaginary lighter). Non-manipulable objects were more frequently depicted with drawing depictions (e.g., ‘pyramid’).

Discussion

In this study we investigated whether certain modes of representations were typically bound to a specific semantic domain and whether there were generalisable patterns observed across different participants. We also looked into the different strategies deployed to distinguish different word types such as actions and objects. The results from a pantomime elicitation task revealed that pantomimes show systematic patterns when speakers are asked to represent a referent in silent gesture. Actions with objects tend to be expressed with a single gesture using an acting mode of representation (i.e., representing how an object is manipulated). Non-manipulable objects in contrast tend to be represented with a drawing strategy and with more than one gesture. Interestingly, manipulable objects elicited significantly more gestures than the other two categories and the most common strategy used was also acting.

Crucially, we observed that manipulable objects elicited significantly more deictics than the other two categories. That is, participants would pantomime an event (e.g., eating soup) and then highlight part of this gesture with a deictic (e.g., pointing at an imaginary spoon).

The present data replicates earlier findings that speakers tend to rely on the acting mode of representation in their gestures (Padden et al., 2013, 2015; van Nispen, van de Sandt-Koenderman, Mol, & Krahmer, 2014). However, we also find that this mode of representation falls out of favour when the referent does not allow an effective way of depiction. Gesturers seem to switch from one strategy to the other depending on whether the referent allows for certain modes of representation (i.e., drawing is favoured in the depiction of non-manipulable objects). These findings go in line with recent research showing that the shape of an object and the possible ways to interact with it modulate the form of co-speech gestures (Masson-Carro et al., 2015). These findings also resonate work showing that gesture production relates to simulation of actions (Cook & Tanenhaus, 2009). It is possible that speakers default to an acting strategy in their gestures because they are simulations of their experiences with objects. However, when an object does not lend itself to a clear form of manipulation or the affordances of the object does not permit the use an acting strategy, individuals will turn to an alternative strategy to represent it.

When we look at the different strategies adopted to make distinctions between actions with objects and manipulable objects we see that gesturers do not align different modes of representation to a specific category, as has been shown for established or emerging sign languages (Johnston, 2001; Kimmelman, 2009; Supalla & Newport, 1986). Instead, we see that gesturers complement their acting strategies with deictics to highlight the focus of their gesture. That is, gesturers feel the communicative need to inform the addressee that the intended referent is not the action they are depicting, but the object at hand.

This study adds to our current understanding of gesture production, and the factors that drive their form. However, we should be cautious about the generalisation of these results given that most research in this domain has focused on the gestures produced by Dutch speakers. Future work should investigate whether speakers of different languages adopt the same strategies in silent gesture regardless of their native language. By looking at communication in the absence of speech we see gesturers devise strategies to express complex notions such as action-object distinctions. These strategies operate in tandem with individuals’ knowledge of the world and the available strategies to represent a referent. These strategies may be the raw materials of emerging sign languages and the foundations of a conventionalised manual linguistic system.

Acknowledgments

This work was supported by a Veni grant by the Netherlands Organisation for Scientific Research (NWO) awarded to the first author. We would like to thank Deniz

Cetin and Renske Schilte for their collaboration in the annotation of the data.

References

- Carrol, J. B. (1956). *Language, Thought and Reality: Selected Writings of Benjamin Lee Whorf*. Cambridge, MA: MIT Press.
- Cook, S. W., & Tanenhaus, M. K. (2009). Embodied communication: speakers' gestures affect listeners' actions. *Cognition*, 113(1), 98–104. doi:10.1016/j.cognition.2009.06.006
- Flecken, M., Von Stutterheim, C., & Carroll, M. (2014). Grammatical aspect influences motion event perception: findings from a cross-linguistic non-verbal recognition task. *Language and Cognition*, 6(01), 45–78. doi:10.1017/langcog.2013.2
- Goldin-Meadow, S., So, W. C., Özyürek, A., & Mylander, C. (2008a). The natural order of events: how speakers of different languages represent events nonverbally. *Proceedings of the National Academy of Sciences of the United States of America*, 105(27), 9163–8. doi:10.1073/pnas.0710060105
- Goldin-Meadow, S., So, W. C., Özyürek, A., & Mylander, C. (2008b). The natural order of events: how speakers of different languages represent events nonverbally. *Proceedings of the National Academy of Sciences of the United States of America*, 105(27), 9163–9168. doi:10.1073/pnas.0710060105
- Hall, M. L., Mayberry, R. I., & Ferreira, V. S. (2013). Cognitive constraints on constituent order: evidence from elicited pantomime. *Cognition*, 129(1), 1–17. doi:10.1016/j.cognition.2013.05.004
- Johnston, T. (2001). Nouns and verbs in Australian sign language: An open and shut case? *Journal of Deaf Studies and Deaf Education*, 6(4), 235–257. doi:10.1093/deafed/6.4.235
- Kimmelman, V. (2009). Parts of speech in Russian Sign Language: The role of iconicity and economy. *Sign Language and Linguistics*, 12(2), 161–186. doi:10.1075/sl
- Kita, S., & Özyürek, A. (2003). What does cross-linguistic variation in semantic coordination of speech and gesture reveal?: Evidence for an interface representation of spatial thinking and speaking. *Journal of Memory and Language*, 48(1), 16–32. doi:10.1016/S0749-596X(02)00505-3
- Langus, A., & Nespors, M. (2010). Cognitive systems struggling for word order. *Cognitive Psychology*, 60(4), 291–318. doi:10.1016/j.cogpsych.2010.01.004
- Lausberg, H., & Sloetjes, H. (2009). Coding gestural behavior with the NEUROGES-ELAN system. *Behavior Research Methods*, 41(3), 841–9. doi:dx.doi.org/10.3758/brm.41.3.841
- Majid, A., & Burenhult, N. (2014). Odors are expressible in language, as long as you speak the right language. *Cognition*, 130(2), 266–270. doi:10.1016/j.cognition.2013.11.004
- Masson-Carro, I., Goudbeek, M., & Krahmer, E. (2015). Can you handle this? The impact of object affordances on how co-speech gestures are produced. *Language, Cognition and Neuroscience*, 1–11. doi:10.1080/23273798.2015.1108448
- McNeill, D. (1992). *Hand and mind: What gestures reveal about thought*. Chicago: University of Chicago Press.
- Müller, C. (2013). Gestural modes of representation as techniques of depiction. In C. Müller, A. Cienki, S. Ladewig, D. McNeill, & J. Bressem (Eds.), *Body - Language - Communication: An International Handbook on Multimodality in Human Interaction* (pp. 1687–1701). Berlin: De Gruyter Mouton.
- Özçalışkan, Ş., Lucero, C., & Goldin-Meadow, S. (2016). Does language shape silent gesture? *Cognition*, 148, 10–18. doi:10.1016/j.cognition.2015.12.001
- Özyürek, A., Kita, S., Allen, S. E. M., Furman, R., & Brown, A. (2005). How does linguistic framing of events influence co-speech gestures?: Insights from crosslinguistic variations and similarities. *Gesture*, 5(1), 219–240. doi:10.1075/gest.5.1.15ozy
- Padden, C., Hwang, S.-O., Lepic, R., & Seegers, S. (2015). Tools for Language: Patterned Iconicity in Sign Language Nouns and Verbs. *Topics in Cognitive Science*, 7(1), 81–94. doi:10.1111/tops.12121
- Padden, C., Meir, I., Hwang, S.-O., Lepic, R., Seegers, S., & Sampson, T. (2013). Patterned iconicity in sign language lexicons. *Gesture*, 13(3), 287–305.
- Sapir, E. (1921). *Language*. New York: Harcourt, Brace & World.
- Supalla, T., & Newport, E. L. (1986). How many sits in a chair? The derivations of nouns and verbs in American Sign Language. In P. Siple (Ed.), *Understanding language through sign language research* (pp. 91–132). New York: Academic Press.
- van Nispen, K., van de Sandt-Koenderman, M., Mol, L., & Krahmer, E. (2014). Pantomime Strategies: On Regularities in How People Translate Mental Representations into the Gesture Modality. In *Proceedings of the 36th Annual Conference of the Cognitive Science Society (CogSci 2014)* (pp. 3020–3026). Austin, TX: Cognitive Science Society, Inc.