

There is more to gesture than meets the eye: Visual attention to gesture's referents cannot account for its facilitative effects during math instruction

Miriam A. Novack¹ (mnovack1@uchicago.edu)
Elizabeth M. Wakefield¹ (ewakefield@uchicago.edu)
Eliza L. Congdon¹ (econgdon@uchicago.edu)
Steven Franconeri² (franconeri@northwestern.edu)
Susan Goldin-Meadow¹ (sgm@uchicago.edu)

¹Department of Psychology, University of Chicago,
5848 S University Ave, Chicago, IL 60637 USA

²Department of Psychology, Northwestern University
2029 Sheridan Rd, Evanston, IL 53706 USA

Abstract

Teaching a new concept with gestures – hand movements that accompany speech – facilitates learning above-and-beyond instruction through speech alone (e.g., Singer & Goldin-Meadow, 2005). However, the mechanisms underlying this phenomenon are still being explored. Here, we use eye tracking to explore one mechanism – gesture's ability to direct visual attention. We examine how children allocate their visual attention during a mathematical equivalence lesson that either contains gesture or does not. We show that gesture instruction improves posttest performance, and additionally that gesture *does* change how children visually attend to instruction: children look more to the problem being explained, and less to the instructor. However looking patterns alone cannot explain gesture's effect, as posttest performance is not predicted by any of our looking-time measures. These findings suggest that gesture does guide visual attention, but that attention alone cannot account for its facilitative learning effects.

Keywords: Gesture; eye tracking; learning; visual attention

Introduction

Teachers use more than words to explain new ideas; they often accompany their speech with gestures – hand movements that express information through both form and movement patterns. Teachers gesture spontaneously in instructional settings (Alibali et al., 2014) and controlled experimental studies have found that children are more likely to learn novel ideas from instruction that includes speech and gesture, than speech alone (e.g., Ping & Goldin-Meadow, 2008; Singer & Goldin-Meadow, 2005; Valenzano, Alibali, & Klatzy, 2003).

Gesture might improve learning by conveying multiple ideas simultaneously (Singer & Goldin-Meadow, 2005), engaging the motor system (Macedonia, Muller, & Friederici, 2011), and linking abstract ideas to concrete objects in the environment (Valenzano et al., 2003). One understudied potential benefit of gesture is that it engages and directs visual attention. Here we used instruction on the

concept of mathematical equivalence as a case study to test how gesturing towards a novel mathematical equation affects not only children's learning outcomes, but also their allocation of visual attention across the equation.

There are reasons to think that gesture's ability to direct visual attention to relevant objects may underlie its positive effects. Because gesture is a spatial, dynamic social cue, it can focus a listener's visual attention on a specific part of the visual environment. Even young infants will shift their visual attention in response to gesture (Rohlfing, Longo, & Bertenthal, 2012). This could, in turn, increase the likelihood that children would focus on crucial aspects of a problem being taught, and would thus learn more from instruction. Learners likely need to attend to the critical information in an instructional context in order to learn from it. For example, toddlers are more likely to learn pairings between objects and labels if their attention is focused on the object while it is being labeled (Yu & Smith, 2012). If gesture during instruction highlights important features of the problem and causes learners to visually fixate on these features while relevant information is being provided in speech, that increased looking should lead to better learning.

Previous work using eye tracking to understand how people process gesture has focused on visual processing of naturally produced gesture during face-to-face communication, such as when watching a person tell a story. Most of this work has been descriptive, documenting where interlocutors focus their visual attention during communication rather than documenting how patterns of visual attention affect comprehension. Overall, the findings suggest that looking directly toward a speaker's hands is actually quite rare (e.g., Gullberg & Holmqvist, 2006; Gullberg & Kita, 2009). Instead, listeners prefer to look mostly at a speaker's face and spend little time overtly attending to gesture. On the rare occasions when interlocutors do look directly at a gesture, it is typically because the speaker himself is looking towards his own

hands, or is holding a gesture in space for an extended period of time (Gullberg & Kita, 2009).

While this descriptive work on visual attention to gesture during spontaneous discourse is informative, we cannot assume the findings will be consistent in instructional settings. First, unlike the discourse studies described above, classroom teachers often gesture towards or near objects (Alibali & Nathan, 2012). In fact, most of the behavioral work that investigates the utility of teachers' gestures has been in situations where gestures are performed in reference to objects (e.g., Ping & Goldin-Meadow, 2008; Singer & Goldin-Meadow, 2005; Valenzeno, et al., 2003). For example, children learn more when a teacher gestures toward a math problem that is written on a chalkboard (Singer & Goldin-Meadow, 2005). This means that in most formal instructional settings, learners have three demands on their visual attention capacities – the instructor who is speaking, the gestures she produces, and the objects she is gesturing toward. Thus, the way in which gesture affects allocation of visual attention in these situations may differ drastically from other kinds of conversational settings.

Second, and more importantly, the way gesture captures or directs visual attention during instruction may have different cognitive implications than how gesture functions in discourse. Specifically, learning involves more than just comprehension of the content of a message; it requires that learners integrate the presented information with their existing knowledge to arrive at a novel conceptual state. This is a non-trivial difference between comprehension and learning, and it may mean that gesture necessarily serves a different function in an instructional context than it does during casual conversation. If learners are sensitive to this, then we might expect that the way gesture affects visual attention during instruction will meaningfully map onto learning outcomes.

In the current study, we ask how gesture directs visual attention for 8-10 year-old children who are learning how to solve missing addend equivalence problems (e.g., $2+5+8 = __+8$). We use eye tracking to compare children's visual attention to instructional videos with either speech alone, or speech with accompanying gesture. Previous work using a similar paradigm has found that giving children relatively brief instruction, using example problems, and allowing children to solve additional problems themselves results in an increased understanding of mathematical equivalence. Importantly, incorporating gesture into instruction boosts this understanding (e.g., Singer & Goldin-Meadow, 2005) relative to instruction with speech alone. In the present study, we use a *grouping* gesture during instruction. This gesture involves producing a V-point to the first two numbers in a missing addend equivalence problem followed by a point to the blank space. This V-point gesture represents the idea that one can solve the equation by adding, or grouping, the first two addends and putting that total in the blank. This V-point gesture is one produced spontaneously by children who already understand how to solve these sorts of problems (e.g., Perry, Church, &

Goldin-Meadow, 1988) and has also been shown to lead to learning when taught to children (Goldin-Meadow, Cook & Mitchell, 2009). Furthermore, this particular gesture is of interest because it contains both *deictic* properties (pointing to specific numbers) and *iconic* properties (representing the idea of grouping through its form). Therefore, the benefits of learning from this type of gesture could arise from looking to the gesture itself, from looking to the numbers that the gesture is referencing, or from some combination therein.

Methods

Participants

Data from 50 participants were analyzed for the present study. Children between the age of 8 and 10 (*mean age* = 8.8 years) were recruited through a database maintained by the University of Chicago Psychology Department and tested in the laboratory. The sample includes 26 children in the Speech+Gesture Condition (14 females) and 24 children in the Speech Alone Condition (14 females). All children in the current sample scored a 0/6 on a pretest, indicating that they did not know how to correctly solve mathematical equivalence problems at the start of the study. Prior to the study, parents provided consent and children gave assent. Children received a small prize, and \$10 compensation for their participation.

Materials

Pretest/Posttest. The pretest and posttest each contained 6 missing addend equivalence problems, presented in one of two formats. In Form A, the last addend on the left side of the equals sign was repeated on the right side (e.g., $a+b+c = __+c$) and in Form B, the first addend on the left side of the equals side was repeated on the right side (e.g., $p+q+r = p'+__$). Both pretest and posttest consisted of 3 of each problem type.

Eye Tracker. Eye tracking data were collected via corneal reflection using a Tobii 1750 eye tracker with a 17 inch monitor. Tobii software was used to perform a 5-point calibration procedure using standard animation blue dots. This was followed by the collection and integration of gaze data with the presented videos using Tobii Studio (Tobii Technology, Sweden).

Instructional videos. Two sets of 6 instructional videos were created to teach children how to solve Form A missing addend math problems (e.g., $5+6+3 = __+3$) – one set for children in the Speech Alone condition and one set for children in the Speech+Gesture condition. All videos showed a woman standing next to a Form A missing addend math problem, written in black marker on a white board. At the beginning of each video, the woman said, "Pay attention to how I solve this problem", and then proceeded to write the correct answer in the blank (e.g., writing 11 in the previous example). She then described how to solve the

problem, explaining the idea of *equivalence*: “*I want to make one side equal to the other side. 5 plus 6 plus 3 equals 14, and 11 plus 3 is 14, so one side is equal to the other side.*” During this spoken instruction, the woman kept her gaze on the problem. In the Speech+Gesture videos, the woman accompanied her speech with a gesture strategy. When she said “*I want to make one side...*”, she simultaneously produced a V-point with her index and middle finger to the first two addends, then, as she said “...the *other side*” she moved her hand across the problem, bringing her fingers together to point to the answer with her index finger. She produced no gestures in the Speech Alone videos. To ensure that the speech was identical across the two training conditions, the actress recorded a single audio track for each problem, prior to filming. Each of the twelve videos was approximately 25 seconds long.

Procedure

Children first completed a written pretest containing 6 missing addend math problems. All children in the current sample scored 0/6¹. The experimenter then wrote children’s (incorrect) answers on a white board and they were asked to explain their solutions.

Next, children sat in front of the eye tracking monitor, approximately 18 inches from the screen, and were told they would watch instructional videos that would help them understand the type of math problems they had just solved. Their position was calibrated and adjusted if necessary, then they began watching the first of the 6 instructional videos (either Speech Alone, or Speech+Gesture, depending on the assigned training condition). At the conclusion of each video, children were asked to solve a new missing addend problem on a small, hand-held whiteboard, and were given feedback on whether or not their answer was correct (e.g., “that’s right, 10 is the correct answer” or “no, actually 10 is the correct answer”). All problems shown in the instructional videos were Form A, and all problems that children had the opportunity to solve were Form A.

After watching all 6 instructional videos and having 6 chances to solve their own problems during training, children completed a new, 6-question paper-and pencil posttest. The posttest, like the pretest, included 3 Form A problems and 3 Form B problems. As children saw only Form A problems during training, we refer to these as “Trained” problems and Form B as “Transfer” problems.

Results

Behavioral Results

Training. Figure 1 shows the proportion of participants in each condition who answered problems correctly during training. A mixed-effects logistic regression predicting the log-odds of success on a given training problem with problem number (1-6) and condition (Speech Alone,

Speech+Gesture) as fixed factors and subject as a random factor revealed a positive effect of training problem ($\beta=0.91$, $SE=0.15$, $z=6.21$, $p<.001$), indicating that children became more likely to correctly answer problems as training progressed. There was, however, no effect of condition during training ($\beta=0.03$, $SE=0.72$, $z=0.04$, $p=.96$, indicating that learning rates during training did not differ by condition. By the final training problem, over 90% of participants in both groups were answering the problems correctly, which suggests that both types of instruction were equally comprehensible.

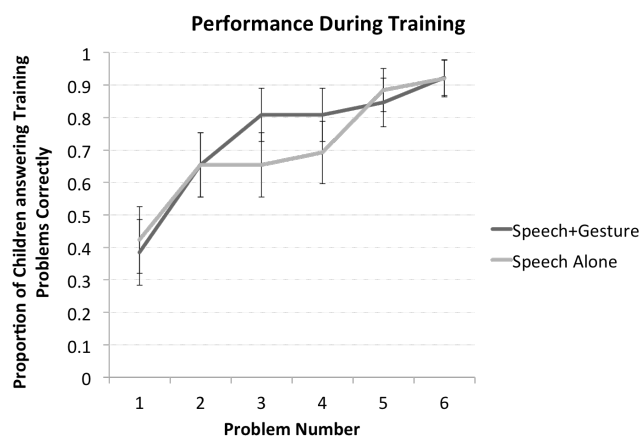


Figure 1. Performance during training on practice problems. Learning increased across the 6 problems, but was not different across the two training conditions.

Posttest. Although the groups did not differ in performance at the end of training, their scores on an immediate posttest reflected an advantage of having learned through Speech+Gesture instruction (see Figure 2). Participants in the gesture condition answered significantly more problems correct at the posttest ($M=4.11$, $SD=2.04$) than participants in the speech condition ($M=2.64$, $SD=2.08$). A mixed-effects logistic regression with problem type (Form A: trained, Form B: transfer) and Condition (Speech+Gesture, Speech Alone) as fixed factors and subject as a random factor showed a significant effect of condition ($\beta = -2.60$, $SE = 0.99$, $z=2.59$, $p<.01$) indicating that posttest performance in the Speech+Gesture Condition was better than performance in the Speech Alone Condition. There was also a significant effect of problem type ($\beta=2.27$, $SE=0.43$, $z=5.31$, $p<.001$), demonstrating that performance on Form A (trained problems) was better than performance on Form B: (transfer problems). There was no significant interaction between Condition and Problem Type ($\beta=0.29$, $SE=0.79$, $z=-0.37$, $p=0.71$).

¹ Children who answered pretest problems correctly ($n=59$) were still run in the study but are excluded from the current analyses.

² There was a gesture space in the Speech Alone video, despite

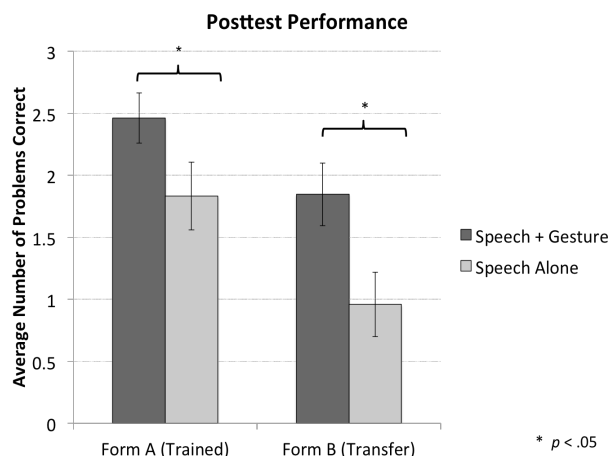


Figure 2. Posttest performance by condition and problem type. Error bars represent ± 1 standard error of the mean.

Eye-Tracking Results

We used a multistep process to analyze the eye tracking data: (1) Areas of interest (AOIs) were generated for the instructor, problem and gesture space² (See Figure 3) using Tobii Studio. Fixations outside of these AOIs were collapsed into “Other”. (2) Data were extracted and processed, such that the AOI a participant fixated in could be determined at 50 msec intervals across the entire length of each problem. (3) Time segments of interest, during which a particular event was happening in the videos (e.g., the instructor stating the equalizer strategy, “*I want to make one side equal to the other side*”) were identified, and total gaze duration during a given time segment in each AOI were computed. (4) We calculated the proportion of time a participant spent in each AOI within each segment collapsed across all six problems. For each participant, eye tracking data were excluded if visual inspection showed that the calibration was off. On average, children in the Gesture Condition contributed data from 4.96 ($SD = 1.34$) trials, and children in the Speech Condition contributed data from 4.90 trials ($SD = 1.34$).

Allocation of visual attention across conditions.

To determine whether patterns of visual attention differed when children were instructed through Speech+Gesture vs. Speech Alone, we considered the proportion of time children spent in each AOI for two time segments of interest. The **strategy** segment encompassed time when the instructor stated the equalizer strategy: *I want to make one side, equal to the other side*. During this segment, spoken instruction was identical across conditions, but children in

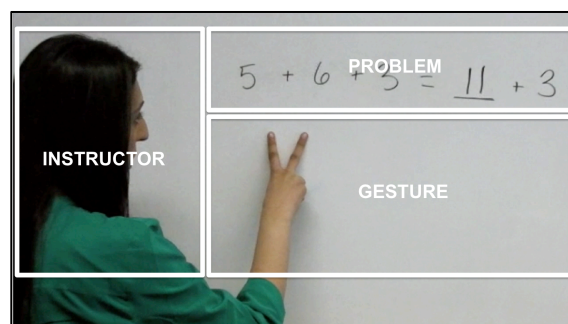


Figure 3. Still shot taking during a gesture segment, with AOIs overlaid.

the Speech+Gesture condition also saw co-speech instructional gestures. As the strategy was explained twice per problem, data from these epochs were combined into one segment of interest. The **explanation** segment encompassed time when the instructor elaborated on the strategy, highlighting the particular addends in the problems (e.g., “*5 plus 6 plus 3 is 14, and 11 plus 3 is 14*”). This segment was visually identical across the experimental groups, allowing us to ask whether the presence of gesture during the preceding **strategy** segment caused children in the Speech+Gesture condition to focus their visual attention in the subsequent **explanation** segment differently than those in Speech Alone instruction.

Strategy segment. Figure 4 shows the proportion of time children spent looking in each of the AOIs during the strategy segment in each condition. On average, children in the Speech+Gesture condition spent a greater proportion of time looking to the problem itself compared to children in the Speech Alone condition (60% versus 48%) ($\beta=0.11$, $SE=0.05$ $t=2.39$, $p<0.05$). In contrast, children in the Speech Alone condition allocated more visual attention to the instructor, compared to children in the Gesture condition (47% vs. 18%) ($\beta=-0.29$, $SE=0.04$ $t=-6.19$, $p<0.01$). Finally, children in the Speech+Gesture condition spent 19% of the time looking to the Gesture space. Unsurprisingly, children in the Speech Alone condition spent significantly less time (3%) in this AOI ($\beta=.16$, $SE=0.02$ $t=5.63$, $p<0.01$) as there was nothing there to draw their attention. Together, these results suggest that gesture does affect visual attention in an instructional context, leading participants to look more to the objects being referenced, and less to the instructor herself.

Explanation segment. Figure 4 also shows the proportion of time spent in each AOI during the explanation segment, with children across both conditions splitting their time evenly between the instructor and the problem. Analyses indicated that there were no differences in looking times to the AOIs by Condition during the explanation segments.

² There was a gesture space in the Speech Alone video, despite the fact that there was never any gesture produced in those videos.

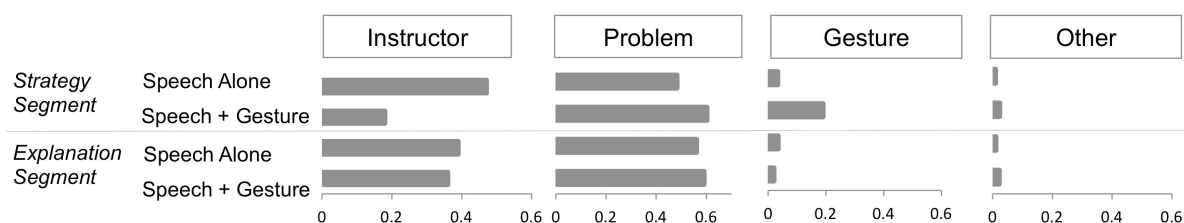


Figure 4. Average proportion of gaze duration across all 6 problems during strategy and explanation segments.

Relation between visual attention and learning.

Given the condition differences between the allocation of visual attention during the strategy segment of instruction, we were interested in whether the focus of attention elicited by the presence of gesture predicted learning outcomes. To explore this we conducted a regression to determine whether looking towards the problem itself (which children did more in the Speech+Gesture condition) predicted posttest performance. Proportion of time looking to the problem did not predict performance on the posttest ($\beta=2.53$, $SE=1.89$, $t=1.34$ $p=0.18$). In other words, the presence of gesture did lead children to look more to objects referenced by gesture but that increase in looking was not responsible for the increase in learning outcomes. Focusing just on the gesture condition, we see that children spend relatively little time looking directly at the gesture (only about 19%), and the amount of looking to the gesture itself, while it is being produced, has no relation to learning outcomes within the gesture condition ($\beta=1.96$, $SE=4.47$, $t=0.44$ $p=0.66$).

Discussion

Although decades of work have found that gesture supports learning when added to instructional contexts, this was the first study to ask how gesture during instruction guides visual attention and facilitates learning through an attentional mechanism. Our behavioral results replicate previous work (e.g., Singer & Goldin-Meadow, 2005). We show that children who learn from watching speech+gesture instruction have more robust learning than children who learn from speech alone, as demonstrated by higher performance on a posttest. Importantly, and surprisingly, we also add a novel finding to the behavioral literature. Whereas most researchers consider posttest performance alone as a measure of learning, we asked how children's performance changed *during* instruction. We show that learning rates *during* instruction did not differ across the two groups, but only emerged after a change in context (i.e., moving from sitting in front of the eye tracker to a desk), and when intermittent reminders of the strategy were not present. This suggests that our learning paradigm may only produce fragile, temporary learning outcomes, but that the addition of gesture to the instruction can help solidify that knowledge. This short-term retention effect corroborates previous work showing that the effects of gesture are particularly good at promoting long-lasting learning (e.g., Cook, Mitchell, & Goldin-Meadow, 2008).

Our eye tracking results demonstrated that at a global level, gesture directs visual attention towards spoken referents in a formal, instructional context, and that children are more likely to focus on referents of gesture than gesture itself. This is interesting, given that the Speech+Gesture videos contained more *items* (i.e., moving hands) for children to look at than the Speech Alone videos, and yet, children in this condition focused the majority of their attention on the problem. Relatedly, it is notable that there was relatively little overt focus on the gesture form, even though previous work suggests that the form of the gesture itself is important for learning in this task (Goldin-Meadow et al., 2009). Finally, in terms of general looking patterns, we found that although gesture affects visual attention when it is being produced, it does not affect visual attention of subsequent speech-only instruction, as seen from our analysis of the explanation segment of instruction.

Our looking time findings suggest similarities between natural communicative gesture, and purposeful, instructional gesture. Like work on communicative gesture, we find that looking directly at gesture is relatively uncommon. However, our results may suggest a difference between natural and instructional gesture contexts: even though fixation on gesture is relatively rare, gesture in instructional contexts may draw more attention than gesture in natural communication. When gesture was present in the current study, all children in the sample looked directly at it, at least for some amount of time. In a study of gesture in discourse, only 9% of gestures were ever fixated (Gullberg & Holmqvist, 2006). This difference may be attributable to the way gestures were used in our instruction that differ from their use in discourse. In our videos, gestures were front-and-center – they were in the middle of the screen, while the instructor was faced away from the child, providing a cue to their importance. In contrast, in previous studies of communicative gesture in discourse, participants see face-to-face communication, where the face may take center stage. Further work examining more types of instructional gesture (and perhaps less salient instructional gestures) may reveal what is driving this difference.

In our final analysis, we asked whether attention to the problem during the strategy segment of instruction led to better posttest performance, with the rationale that finding this link would suggest that at least part of the facilitative effects of gesture in previous studies is driven by its ability to guide attention. Although we did not find evidence that gesture enhances learning by highlighting important features

of a problem, and increasing fixation to those features, it may be possible for gesture to highlight important *relational* aspects of a problem, which will be examined in future work. For example, adults solving these same kinds of math problems are less likely to make errors if they traversed the equal sign, a gaze pattern that may be highlighting the relational structure of the equation (Chesney et al., 2013). Thus it is possible that gesture could lead to useful eye-movement patterns not captured by the current analysis, which could in turn support learning outcomes.

It also remains possible that the effect of gesture on visual attention is not the main mechanism through which gesture facilitates learning. For example, Ping & Goldin-Meadow (2008) found that 5-6 year olds were just as likely to improve their understanding of Piagetian conservation after a lesson that included gesture, irrespective of whether or not the objects to which the gestures referred (i.e., glasses that contained water) were present. In another study, Goldin-Meadow, Cook, & Mitchell (2009) taught children how to solve missing addend equivalence problems by producing a grouping gesture either to the correct, or incorrect addends to be grouped. Remarkably, children learned even if they had produced the V-point to the wrong addends, suggesting that directing visual attention to the wrong place does not disrupt gesture's positive effects on learning. Still, it seems likely that visual attention is part of the story. In fact, in the example given above, Goldin-Meadow et al. (2009) found that although children could learn from an 'incorrect' gesture, they benefitted more from the same gesture, used to highlight grouping of the correct addends, and, presumably, draw visual attention to these addends.

In the present study, we have established that instructional gesture *does* drive children to look at a novel equation differently, and children show increased learning after this type of instruction; we have just also shown that this shift in global looking pattern does not provide a simple causal explanation for this cognitive effect. Future work will consider how more nuanced aspects of visual attention, such as whether it helps children synchronize their looking with spoken instruction, as well as ways in which the ability to guide visual attention may combine with other features of gesture to support learning.

Acknowledgments

Funding for this study was provided by NICHD (R01-HD47450, to Goldin-Meadow), NSF BCS 1056730, and the Spatial Intelligence and Learning Center (SBE 0541957, Goldin-Meadow is a co-PI) through the National Science Foundation. We also thank Kristin Plath, William Loftus, and Aileen Campanaro for their help with data collection, and Amanda Woodward for the use of the eye tracker.

References

Alibali, M. W., & Nathan, M. J. (2012). Embodiment in Mathematics Teaching and Learning: Evidence From Learners' and Teachers' Gestures. *Journal of the*

- Learning Sciences*, 21, 247–286.
doi:10.1080/10508406.2011.611446
- Alibali, M. W., Nathan, M. J., Wolfgram M. S., Church, R. B., Jacobs, S.A., Johnson Martinez, C. & Knuth, E. J. (2014) How Teachers Link Ideas in Mathematics Instruction Using Speech and Gesture: A Corpus Analysis, *Cognition and Instruction*, 32, 65-100. doi: 10.1080/07370008.2013.858161
- Chesney, D. L., McNeil, N. M., Brockmole, J. R., & Kelley, K. (2013). An eye for relations: eye-tracking indicates long-term negative effects of operational thinking on understanding of math equivalence. *Memory & Cognition*, 41, 1079-1095.
- Cook, S. W., Mitchell, Z., & Goldin-Meadow, S. (2008). Gesturing makes learning last. *Cognition*, 106, 1047-1058. doi: 10.1016/j.cognition.2007.04.010
- Goldin-Meadow, S., Cook, S. W., & Mitchell, Z. A. (2009). Gesturing gives children new ideas about math. *Psychological Science*, 20, 267–272. doi:10.1111/j.1467-9280.2009.02297.x
- Gullberg, M., & Kita, S. (2009). Attention to speech-accompanying gestures: Eye movements and information uptake. *Journal of nonverbal behavior*, 33, 251-277. doi: 10.1007/s10919-009-0073-2
- Gullberg, M., & Holmqvist, K. (2006). What speakers do and what listeners look at. Visual attention to gestures in human interaction live and on video. *Pragmatics and Cognition*, 14, 53–82. doi: http://dx.doi.org/10.1075/pc.14.1.05gul
- Macedonia, M., Muller, K., & Friederici, A. D. (2011). The impact of iconic gestures on foreign language word learning and its neural substrate. *Human Brain Mapping*, 32, 982-998. doi: 10.1002/hbm.21084
- Rohlfing, K. J., Longo, M. R., & Bertenthal, B.I. (2012). Dynamic pointing triggers shifts of visual attention in young infants. *Developmental Science*, 15, 426-435. doi:10.1111/j.1467-7687.2012.01139.x
- Singer, M. a, & Goldin-Meadow, S. (2005). Children learn when their teacher's gestures and speech differ. *Psychological Science*, 16, 85–89. doi:10.1111/j.0956-7976.2005.00786.x
- Perry, M., Church, R. B., & Goldin-Meadow, S. (1988). Transitional knowledge in the acquisition of concepts. *Cognitive Development*, 3, 359–400. doi:10.1016/0885-2014(88)90021-4
- Ping, R. M., & Goldin-Meadow, S. (2008). Hands in the air: using ungrounded iconic gestures to teach children conservation of quantity. *Developmental Psychology*, 44, 1277–87. doi: 10.1037/0012-1649.44.5.1277
- Valenzeno, L., Alibali, M. W., & Klatzky, R. (2003). Teachers' gestures facilitate students' learning: A lesson in symmetry. *Contemporary Educational Psychology*, 28, 187–204. doi: 10.1016/S0361-476X(02)00007-3
- Yu, C., & Smith, L. B. (2012). Embodied attention and word learning by toddlers. *Cognition*, 125, 244-262. doi:10.1016/j.cognition.2012.06.016