

Learning to Express Left-Right & Front-Behind in a Sign versus Spoken Language

Beyza Sümer (beyza.sumer@mpi.nl)

Radboud University, Center for Language Studies, Nijmegen, The Netherlands
International Max Planck Research School for Language Sciences, Nijmegen, The Netherlands

Pamela Perniss (p.perniss@ucl.ac.uk)

Deafness Cognition and Language Research Center, University College London

Inge Zwitserlood (i.zwitserlood@let.ru.nl)

Radboud University, Center for Language Studies, Nijmegen, The Netherlands

Ash Özyürek (asli.ozyurek@mpi.nl)

Radboud University, Center for Language Studies, Nijmegen, The Netherlands
Max Planck Institute for Psycholinguistics, Nijmegen, The Netherlands

Abstract

Developmental studies show that it takes longer for children learning spoken languages to acquire viewpoint-dependent spatial relations (e.g., left-right, front-behind), compared to ones that are not viewpoint-dependent (e.g., in, on, under). The current study investigates how children learn to express viewpoint-dependent relations in a sign language where depicted spatial relations can be communicated in an analogue manner in the space in front of the body or by using body-anchored signs (e.g., tapping the right and left hand/arm to mean left and right). Our results indicate that the visual-spatial modality might have a facilitating effect on learning to express these spatial relations (especially in encoding of left-right) in a sign language (i.e., Turkish Sign Language) compared to a spoken language (i.e., Turkish).

Keywords: Acquisition; sign language; spatial language; left-right; front-behind

Introduction

The visual-spatial modality of sign languages allows spatial relations to be expressed in an analogue manner through the use of locative predicates in the space in front of the body of the signers (Fig. 1a) or by using iconic lexical signs, called relational lexemes (Fig. 1b) (e.g., Emmorey, 2002). Spoken languages, on the other hand, express space categorically, and mainly through abstract labels such as adpositions (see example 1).

Example 1. Kalem kağıd+ın sol+un+da (Turkish)
Pen paper+GEN left+POSS+LOC
“Pen is at the left of the paper”

This raises interesting questions about whether some spatial expressions can be acquired earlier in sign

languages due to the iconic correspondences between form and meaning than in spoken languages. Here we investigate these questions in the domain of viewpoint-dependent spatial relations (i.e., left-right & front-behind).

Acquisition of viewpoint-dependent spatial relations has been reported to appear later than the spatial relations that do not include a viewpoint such as relations of containment (e.g., in) or support (e.g., on) (Piaget & Inhelder, 1971; Johnston, 1988). However, these studies are restricted to spoken languages, and there are few studies on sign languages, which are interesting in this domain due to the modality’s visual-iconic and embodied affordances. The studies with sign language acquiring children have so far focused only on the comprehension of these spatial terms (Martin & Sera, 2006; Morgan, Herman, Barriere, & Woll, 2008), and deaf children acquiring a sign language were found to lag behind children acquiring a spoken language in comprehending these spatial relations (Martin & Sera, 2006). Production studies comparing signing and speaking children in how they learn to express viewpoint-dependent spatial relations in similar tasks are lacking. Thus, the present study compares the acquisition of expressions encoding viewpoint-dependent relations in a sign (Turkish Sign Language; Türk İşaret Dili, TİD) and a spoken language (Turkish), both of which are understudied languages. As such it offers the first comparative study in this domain.

Expressing viewpoint-dependent spatial relations in spoken and sign languages

Encoding certain spatial relations (e.g., pen left of paper) requires interlocutors to impose a viewpoint (either their own or that of their addressees) onto their relational encodings. In spoken languages, speakers’ descriptions usually match how a speaker views a spatial scene (Levelt, 1989), but they may also adopt

the view of their addressee (Schober, 1993). Sign languages are similar to spoken languages in that signers can describe spatial scenes from their own viewpoint, or from the viewpoint of the addressee (Emmorey, 1996).

In sign languages, encoding spatial relations is mainly realized through polymorphic *classifier predicates* where the signers' hands represent the objects (e.g., a smaller, foregrounded Figure, and a larger, backgrounded Ground) in the spatial configuration, and their relative locations are mapped onto the signing space in an analogue way to the real space depicted. In Fig. 1a, the signer positions her hands relative to her body and to each other to encode the spatial relation between the pen and the paper from her own viewpoint (Emmorey, 1996).



Figure 1: TİD signers' descriptions of the spatial relation of the pen with respect to the paper using (a) classifier predicates and (b) a relational lexeme for left. In the whole utterance, these signs are typically preceded by lexical signs of PAPER (Ground) first, and then PEN (Figure) (not shown here)

To describe viewpoint-dependent spatial relations, signers can also use body-anchored categorical lexical signs (i.e., relational lexemes), instead of, or in addition to classifier predicates in the same utterance. In (Fig. 1b), the signer uses a relational lexeme meaning left to describe the location of the pen in the picture above in relation to the paper. As these examples show, relational lexemes are more categorical than analogue representations conveyed by classifier predicates. Their visual forms are directly anchored to the coordinates of the signers' body (see Fig. 2a, b, c for other relational lexemes in TİD).



Figure 2: TİD signs for (a) right, (b) front, and (c) behind.

Learning to express viewpoint-dependent spatial relations in spoken and sign languages

Studies about the acquisition of viewpoint-dependent spatial relations in spoken languages show that children initially use these terms to refer to their own left-right or front-back. At later stages, they start using these terms to refer to left-right and front-behind of other people, objects, and relative locations (Piaget & Inhelder, 1971; Harris, 1972; Kuczaj & Maratsos, 1975; Conner & Chapman, 1985; Roberts & Aman, 1993).

The studies mentioned so far mainly compared the learning of spatial relations that require a viewpoint with ones that do not. Within viewpoint-dependent spatial relations, left-right distinctions are reported to be more difficult to acquire due to the bilateral symmetry of many objects when compared to front-behind distinctions, which are usually identifiable by distinct perceptual features of objects (Shepard & Hurwitz, 1984; Harris, 1972).

In sign languages, there are only two studies that have investigated the acquisition of viewpoint-dependent spatial relations, and they focus only on comprehension. The findings of these studies show that sign language acquiring children learn the constructions whose comprehension requires mental rotation (i.e., transposition of the signer's left-right, front-behind to their own) later than the ones that do not (i.e., above-below) (Martin & Sera, 2006; Morgan, et al., 2008). Moreover, comparing deaf children acquiring American Sign Language and hearing children acquiring English, Martin & Sera (2006) also observed that deaf children lagged behind age-matched hearing children in the comprehension of these spatial relations.

As mentioned earlier, production studies with sign language acquiring children in this domain are lacking. In general, using terms of viewpoint-dependent spatial relations seems to be challenging for children acquiring a spoken language (Piaget & Inhelder, 1971; Johnston, 1988). Whether this challenge can be overcome by sign-language-acquiring children via exploiting the visual-spatial modality in the expression of these spatial relations is a question that has not been previously investigated.

The Present Study

We suggest three hypotheses about the acquisition of viewpoint-dependent spatial relations in Turkish and TİD: (a) *Similar developmental patterns*: In line with the literature that suggests a universal pattern for the late emergence of viewpoint-dependent spatial relations, we might assume a similar developmental pattern for sign and spoken language acquisition. If this is the case, TİD acquiring children will learn to express these spatial relations late and at similar ages with hearing children. Thus, there will be no effect of the modality of the language being acquired, but general cognitive developmental principles will apply to acquisition of both sign and spoken languages. (b) *Facilitation through visual modality*: One might assume that the iconic properties of the sign language

constructions through which spatial relations are produced (i.e., classifier predicates, Fig. 1a), and especially the use of lexical signs that are directly executed on the signer body (Fig. 1b) may facilitate learning to express these spatial relations in a sign language. On the other hand, learning arbitrary mappings between linguistic labels and spatial relations in spoken languages may present challenges for children. In this case, TİD acquiring children will learn to express these terms earlier than children acquiring Turkish. (c) *Delay due to late comprehension in TİD*: Finally considering previous studies that show that sign-language-acquiring children lagged behind spoken-language-acquiring children in comprehending viewpoint-dependent spatial relations, the production of these spatial relations in a sign language (i.e., TİD) will also appear later than in a spoken language (i.e., Turkish).

To test these hypotheses, deaf children and adults who have learned TİD natively (i.e., from deaf parents) and age-matched Turkish speakers were given picture description tasks where pictures showed two objects configured in left-right (i.e., lateral axis) and front-behind (i.e., sagittal axis) relations.

Participants

Data were elicited from deaf children and hearing children in two age groups (younger children, mean age: 5;2 years & older children, mean age: 8;1 years; N=10 in each group). Their data were compared to those of adults (N=10 for each language). All deaf adults are also native signers of TİD, and all participants reside in İstanbul, Turkey.

Stimuli and Procedure

Signers/speakers were asked to sit opposite an interlocutor (a hearing or deaf confederate depending on the language condition). Stimulus pictures, presented on a computer screen, showed two objects located on either the lateral axis (left-right; e.g., pen to left of paper) (N=6) or the sagittal axis (front-behind; e.g., cup behind a box) (N=6) (see Fig. 3). The Ground objects in the pictures did not have intrinsic fronts or backs. Target pictures were presented with 3 other pictures that showed the same/different objects in different spatial configurations, and remained visible on a computer screen during the description to avoid memory effects. The addressee, who did not see the screen, was given the same 4 pictures on a separate paper, and was asked to find the described one.

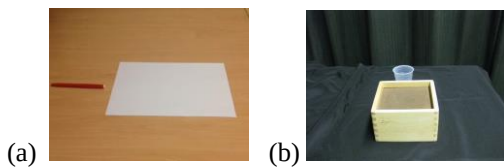


Figure 3: Examples of stimuli pictures used to elicit encodings for (a) left-right and (b) front-behind.

Results

Mean proportions of descriptions encoding a spatial relation and different strategies to encode these relations were calculated for each person and group. Arcsine transformation was applied to all the data, but the mean proportions and standard errors in the graphs are reported from the untransformed data.

First, we investigated how frequently the spatial relation between the objects was encoded by different age groups in each language. A 3-way mixed ANOVA with spatial type (within subjects: left-right and front-behind), age (between subjects: adult, older children, younger children), and language (within subjects: TİD, Turkish) as factors revealed a main effect of age, $F(2,54)=8.83$, $p<.001$, $\eta^2_p=.25$, but not for spatial type, $F(1,54)=.85$, $p=.36$, $\eta^2_p=.02$, or language, $F(1,54)=.69$, $p=.41$, $\eta^2_p=.01$. There was no interaction between spatial type and language, $F(1,54)=1.23$, $p=.27$, $\eta^2_p=.02$; spatial type and age, $F(2,54)=1.09$, $p=.34$, $\eta^2_p=.04$; and between language and age, $F(2,54)=.11$, $p=.89$, $\eta^2_p=.004$. There was no 3-way interaction among the variables, $F(2,54)=.31$, $p=.74$, $\eta^2_p=.01$. Post-hoc comparisons (Bonferonni) for the main effect of age indicated that older children expressed the relational encoding between the Figure and the Ground as frequently as adults ($p=.63$), while younger children expressed them significantly less frequently than adults ($p<.001$) and older children ($p=.02$). This pattern was the same for both deaf and hearing children.

Table 1: Mean proportions and (SEs) of frequency of encoding of a spatial relation by the different age groups in TİD and Turkish.

Participants	TİD	Turkish
Adults	.98 (.02)	1.00(.00)
Older Children	.94(.04)	.96(.02)
Younger Children	.81(.08)	.75(.12)

As the next step, we investigated what types of relational encoding linguistic strategies were preferred by adults and children, and how adult-like children in each age group were. Since it is hard to equate language strategies to encode a spatial relation in these languages, the analyses were conducted separately for TİD and Turkish.

Linguistic strategies used to encode a locative relation in TİD

We categorized TİD strategies into three groups: classifier predicates (Fig. 1a), relational lexemes (Fig.

1b), and others (e.g., showing the location of the objects through index finger pointing). The results of a 2 (Within subjects, Spatial type: left-right and front-behind) by 3 (Within subjects, Linguistic strategy: classifier, relational lexeme, other) by 3 (Between subjects, Age: adults, older children, younger children) mixed ANOVA yielded main effects for spatial type, $F(1,27)=4.89$, $p=.04$, $\eta^2_p=.15$; linguistic strategy, $F(1.77,47.89)=56.62$, $p<.001$, $\eta^2_p=.68$; and age $F(2,27)=13.39$, $p<.001$, $\eta^2_p=.50$. Due to an interaction between linguistic strategy and spatial type, $F(1.31, 35.50)=7.08$, $p=.007$, $\eta^2_p=.21$, separate analyses were conducted for each spatial type.

Left-Right Encoding in TİD The results of a 3 (Between subjects, Age: adults, older children, younger children) by 3 (Within subjects, Linguistic strategy: classifier, relational lexeme, other) mixed ANOVA showed no main effect for age, $F(2,27)=1.66$, $p=.21$, $\eta^2_p=.11$, but a main effect for linguistic strategy, $F(1.39,37.61)=60.56$, $p<.001$, $\eta^2_p=.69$, without an interaction between them, $F(4,54)=.10$, $p=.98$, $\eta^2_p=.008$. Tests of within-subject controlled comparisons showed that classifier predicates were preferred more frequently than relational lexemes ($p<.001$), and the "other" strategies ($p<.001$). The frequency of using relational lexemes and the ones in the "other" category were found to be similar to each other ($p=.65$). The lack of a main effect for age indicates that deaf children in both age groups used the linguistic forms in three different categories as frequently as deaf adults. Thus, these findings suggest that TİD acquiring children, even the younger ones, are able to employ the different language strategies as frequently as adults to encode left-right (see Fig. 4a).

Front-Behind Encoding in TİD The results of a 3 (Between subjects, Age: adults, older children, younger children) by 3 (Within subjects, Linguistic strategy: classifier, relational lexeme, other) mixed ANOVA yielded a main effect of age, $F(2,27)=14.17$, $p<.001$, $\eta^2_p=.51$, and main effect of linguistic strategy, $F(1.69,45.73)=29.12$, $p<.001$, $\eta^2_p=.52$, without any interaction between them, $F(4,54)=1.03$, $p=.40$, $\eta^2_p=.07$. The controlled contrasts for the main effect of spatial type indicated that classifier predicates were used more frequently than relational lexemes ($p<.001$) and the "other" strategies ($p<.001$). Relational lexemes were observed to be more frequent than the "other" forms, as well ($p<.001$). Post-hoc comparisons (Bonferroni) for the effect of age showed that older ($p<.001$) and younger ($p=.001$) deaf children differed from adults in how frequently they used these strategies. In other words, both age groups of children used classifier predicates and the relational lexemes less frequently than adults, but it was only the older children who used the forms from the "other" category less frequently than adults, while younger ones preferred

these "other" forms more frequently than deaf adults. There was no such difference between two age groups of deaf children ($p=1.00$) (see Fig. 4b).

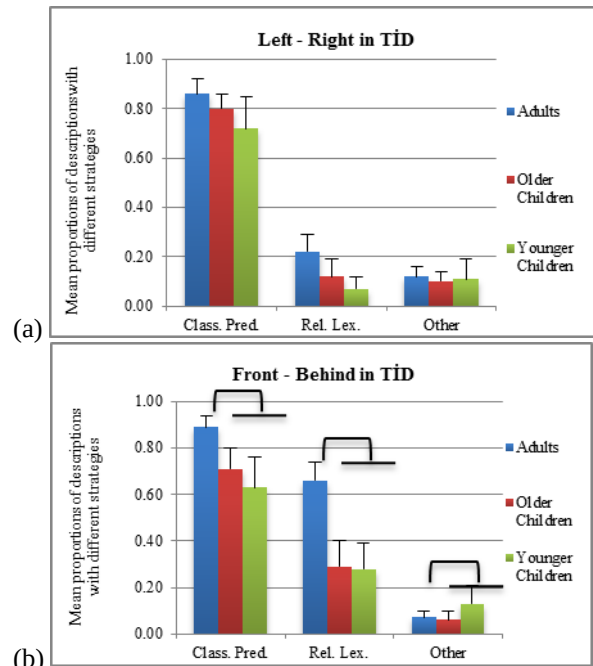


Figure 4: Mean proportions and error bars (representing SE) of descriptions with different strategies available in TİD to encode (a) left-right and (b) front-behind

Due to different production patterns found for left-right and front-behind in TİD, we investigated if the reason why deaf children lagged behind adults in front-behind encoding could be related to the fact that deaf adults used double strategies, possibly for emphasis (i.e., first a classifier predicate followed by a relational lexeme). So, we examined the frequency of descriptions where deaf participants used double strategies. We observed that deaf adults encoded front-behind by using double strategies more frequently than left-right encodings, and overall did so more frequently than children who preferred either a classifier predicate or a relational lexeme, but not both (see Table 2).

Table 2: Raw numbers and (mean proportions) of descriptions with a "double strategy" in TİD

TİD Signers	Left-Right Encodings	Front-Behind Encodings	Total
Adults	13 (.22)	27 (.47)	40 (.35)
Older Children	5 (.09)	1 (.02)	6 (.05)
Younger Children	1 (.03)	2 (.05)	3 (.04)

Linguistic strategies used to encode a locative relation in Turkish

To encode viewpoint-dependent spatial relations, Turkish-speaking adults either used a general relational term (e.g., *Kalem kağıdın yanında* “pen is at the side of the paper”) or employed a viewpoint-dependent spatial noun (e.g., *Kalem kağıdın solunda* “pen is left of the paper”). They also sometimes encoded a spatial relation on a different axis. For example, while describing a left-right (i.e., lateral axis) configuration (e.g., pen left of paper), they used relational lexemes for front or behind (i.e., sagittal axis). We categorized this as “other” strategy.

In order to see if Turkish acquiring children are similar to adults in how frequently they prefer different spatial strategies to encode viewpoint-dependent relations, we conducted a 2 (Within subjects, Spatial type: left-right and front-behind) by 3 (Within subjects, Linguistic strategy: (left-right/front-behind, side, other) by 3 (Between subjects, Age: adults, older children, younger children) mixed ANOVA. It showed main effects for linguistic strategy, $F(1.23,33.11)=19.17$, $p<.001$, $\eta^2_p=.42$, and spatial type, $F(1,27)=10.03$, $p=.004$, $\eta^2_p=.27$, but not for age, $F(2,27)=2.97$, $p=.07$, $\eta^2_p=.18$. There was an interaction between linguistic strategy and age, $F(4,54)=11.99$, $p<.001$, $\eta^2_p=.47$, and between linguistic strategy and spatial type, $F(1.26,33.99)=14.83$, $p<.001$, $\eta^2_p=.36$, in addition to a 3-way interaction among the variables, $F(4,54)=2.83$, $p=.03$, $\eta^2_p=.17$. So, we conducted separate analyses for encoding left-right and front-behind in Turkish.

Left-Right Encoding in Turkish A 3 (Between subjects, Age: adults, older children, younger children) by 3 (Within subjects, Linguistic strategies: left-right, side, other) mixed ANOVA showed main effects for age, $F(2,27)=3.26$, $p=.05$, $\eta^2_p=.20$, and for linguistic strategy, $F(1.76,47.43)=9.27$, $p=.001$, $\eta^2_p=.26$, with an interaction between them, $F(4,54)=10.71$, $p<.001$, $\eta^2_p=.44$. Due to the interaction, one-way ANOVAs were conducted, and we observed that Turkish acquiring children in both age groups employed spatial nouns for left-right less frequently than adults ($p=.007$ for older children and $p=.006$ for younger children). There was no difference between older and younger hearing children ($p=.80$). Instead, older children preferred the general relational term “side”, and younger ones used a spatial noun for a different axis (front-behind) more frequently than adults ($p=.003$ and $p=.03$, respectively) (see Fig. 5a).

Front-Behind Encoding in Turkish A 3 (Age: adults, older children, younger children) by 3 (Linguistic strategy: front-behind, side, other) mixed ANOVA yielded no main effect of age, $F(2,27)=.25$, $p=.78$, $\eta^2_p=.02$, but a main effect for linguistic strategy, $F(1.06,28.69)=22.05$, $p<.001$, $\eta^2_p=.45$, with an interaction between them, $F(4,54)=5.67$, $p=.007$, $\eta^2_p=.30$. The results of one-way ANOVAs showed that older children employed a spatial noun for front-behind

as frequently as adults ($p=.26$), but younger ones used them less frequently than adults ($p=.008$). Younger children used the general relational term “side” more frequently than adults ($p=.007$). Children never used the other left-right category to describe front-behind relations (see Fig. 5b). Unlike TİD signers, Turkish-speaking participants did not use double strategies in their relational encodings.

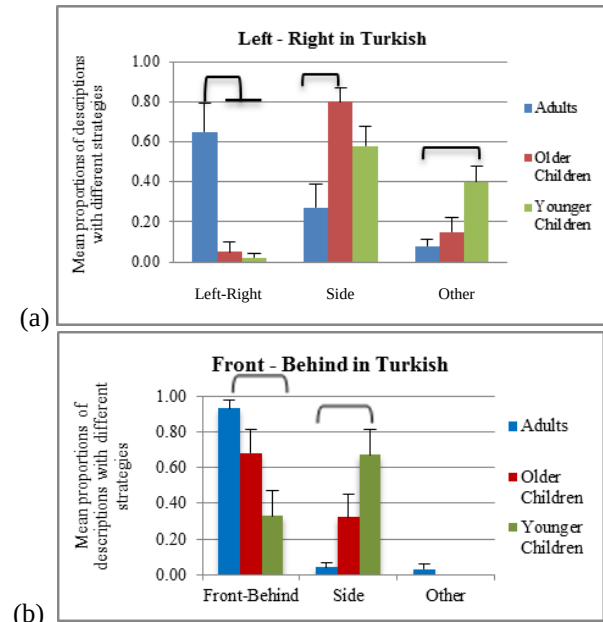


Figure 5: Mean proportions and error bars (representing SE) of descriptions with strategies available in Turkish to encode (a) left-right and (b) front-behind

Discussion and Conclusion

A closer look into the language-specific strategies in a sign (TİD) and a spoken language (Turkish) reveals differences in how children learn to express viewpoint-dependent spatial relations in each language.

To encode left-right, TİD-acquiring deaf children were similar to Turkish deaf adults in how likely they were to use classifier predicates and relational lexemes. One might argue that classifier predicates do not necessarily encode left-right, but rather “next to” relations - as in the case of Turkish “*yanında*” - at the side”. Thus, TİD-acquiring children’s use of classifier predicates as frequently as deaf adults may not necessarily show that they are encoding left-right distinctively. However, these children were also observed to be similar to deaf adults in how frequently they used relational lexemes for left and right, which are more categorical than classifier predicates. Both age groups of Turkish acquiring children, on the other hand, used spatial nouns for left-right much less frequently than Turkish speaking adults. Instead, they mostly preferred to describe the location of the entities as “at the side” of another object. When they attempted to provide spatial encodings more specific than “at the

side", they mostly used "front" or "behind", thus referring to sagittal axis (i.e., front-behind) for the objects located on lateral axis (i.e., left-right). This strategy was not observed among TİD using children.

For front-behind encodings in TİD, deaf children were not adult-like, and used classifier predicates and relational lexemes less frequently than deaf adults. However, deaf adults preferred to encode front-behind by mostly using two different strategies while deaf children almost always used a single strategy in their descriptions. The difference in the frequency of using double strategies by adults but not children might have caused the non-adult-like pattern found for deaf children in encoding front-behind.

For front-behind encodings in Turkish we found that older Turkish speaking children used spatial nouns as frequently as adults. Younger children, on the other hand, still need to learn adult-like use of spatial nouns. They preferred instead to use the general relational term "side".

In sum findings clearly indicate that Turkish speaking children produce spatial nouns for front and behind earlier than spatial nouns for left and right, confirming previous research in spoken languages. TİD signing children, on the other hand, were adult-like in using classifier predicates as well as categorical relational lexemes especially for left-right earlier than for front-behind, and earlier than their age-matched hearing peers.

These results imply that expressing viewpoint-dependent spatial relations, especially left-right, might be facilitated through directly mapping these terms onto the coordinates of the body. However, this advantage seems to manifest itself in production rather than in comprehension (Martin & Sera, 2006). Thus, the availability of iconic forms in sign language can be an advantage for production (i.e., earlier production) and a disadvantage for comprehension (i.e., not providing enough cognitive challenge for abstraction). Since this study is about the production of these terms, the results reflect the maximized possibility of language modality, but not necessarily the level of cognitive abstraction. Signing children might show adult-like proficiency earlier either because the iconic forms help them to develop these concepts earlier or because the adult-form does not require mastering abstract forms due to the iconicity between form and meaning. Further understanding of the effects of modality in this domain necessitates the analysis of viewpoint (signer vs. addressee) as well as the use of co-speech gestures in such spatial descriptions by speakers – especially for the cases in Turkish where no viewpoint was encoded in speech (Sümer, Perniss, Zwitterlood, Özyürek, forthcoming).

Acknowledgments

This work has been supported by a European Research Council (ERC) Starting Grant awarded to the fourth author. We thank our deaf assistants Sevinç Akın and Şule Kibar for their help with data collecting and coding.

References

- Conner, P. & Chapman, R. (1985). The development of locative comprehension in Spanish. *Journal of Child Language*, 12, 109-123.
- Emmorey, K. (1996). The confluence of space and language in signed language. In P. Bloom, M.A. Peterson, L. Nadel, & M. Garrett (Eds.), *Language and Space*, (pp. 171-210). Cambridge, MA: MIT Press.
- Emmorey, K. (2002). *Language, cognition, and the brain: Insights from sign language research*. Mahwah, NJ: Lawrence Erlbaum Associates.
- Harris, L. (1972). Discrimination of left and right, and development of the logic relations. *Merrill-Palmer Quarterly*, 18, 307-320.
- Johnston, J. R. (1988). Children's verbal representation of spatial location. In J. Stiles-Davis, M. Kritchevsky, & U. Bellugi (Eds.), *Spatial Cognition: Brain bases and development*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Kuczaj, S. A. & Maratsos, M. P. (1975). On the acquisition of "Front, Back", and "Side". *Child Development*, 46, 202-210.
- Martin, A. J. & Sera, M. D. (2006). The acquisition of spatial constructions in American Sign Language and English. *Journal of Deaf Studies and Deaf Education*, 11 (4), pp. 391-402.
- Levelt, W. J. M. (1989). *Speaking: From intention to articulation*. Cambridge, MA: The MIT Press.
- Morgan, G., Herman, R., Barriere, I., & Woll, B. (2008). The onset and the mastery of spatial language in children acquiring British Sign Language. *Cognitive Development*, 23, 1-19.
- Piaget, J. & Inhelder, B. (1971). *A Child's Conception of Space* (F. J. Langdon & J. L. Lunzer, Trans.). New York: Norton (Original work published 1948).
- Roberts, R. J. & Aman, C. J. (1993). Developmental differences in giving directions: Spatial frames of reference and mental rotation. *Child Development*, 64, 1258-1270.
- Schober, M. (1993). Spatial perspective taking in conversation. *Cognition*, 47, 1-24.
- Shepard, R. N. & Hurwitz, S. (1984). Upward direction, mental rotation, and discrimination of left and right turns in maps. *Cognition*, 18, 161-193.
- Sümer, B., Perniss, P., Zwitterlood, I., & Özyürek, A. (forthcoming). The Role of Visual Modality in Development of Encoding Spatial Relations: A Comparison of Sign, Speech, and Co-speech Gesture