

Crosslinguistic influence in bilingual morphosyntactic processing: Effects of language-common, language-contrasting, and language-specific information

Research Article

Cite this article: Kim, H. (2025). Crosslinguistic influence in bilingual morphosyntactic processing: Effects of language-common, language-contrasting, and language-specific information. *Bilingualism: Language and Cognition* **28**, 172–184. <https://doi.org/10.1017/S1366728924000269>

Received: 4 November 2022
Revised: 3 March 2024
Accepted: 4 March 2024
First published online: 11 April 2024

Keywords:

L2 sentence processing; Korean numeral quantifier; crosslinguistic influence; self-paced reading

Author for correspondence:

Hyunwoo Kim;
Email: hyunwoo2@yonsei.ac.kr

Hyunwoo Kim

Yonsei University, Seoul, Republic of Korea

Abstract

This study investigated how second language (L2) learners process the Korean numeral quantifier construction by using transferable and nontransferable information. For Chinese-speaking learners of Korean (Chinese group), agreement between an honorific numeral quantifier and a noun in Korean constitutes transferable information in the canonical structure and nontransferable L2-specific information in the scrambled structure. For Japanese-speaking learners of Korean (Japanese group), this information gives rise to crosslinguistic conflicts in both structures. The results from a self-paced reading task showed that the Japanese group did not exhibit sensitivity to grammatical errors in both structures, whereas the Chinese group detected the agreement violation in the canonical but not in the scrambled structure. When a context sentence was provided to license scrambling in the test sentence, however, another group of Chinese-speaking learners of Korean showed sensitivity to the violation. These findings suggest varying degrees of crosslinguistic influence in L2 sentence processing.

1. Introduction

Bilingual or second language (L2) processing often diverge from first language (L1) processing due to influence of various linguistic and nonlinguistic constraints. Among various factors affecting L2 processing, the role of learners' L1 has received critical attention in the L2 acquisition and processing literature, most often discussed in terms of crosslinguistic influence. Previous studies have reported mixed findings on crosslinguistic influence in L2 morphosyntactic processing, particularly with regard to whether L2 learners can process target structures that are not transferable from their L1 (e.g., Hopp, 2017; Jackson & Dussias, 2009; Morett & MacWhinney, 2013; Omaki & Schulz, 2011; Tolentino & Tokowicz, 2011; Trenkic et al., 2014; Witzel et al., 2012).

The discrepancies across studies may be due in part to inconsistencies in terms of linguistic materials, learner characteristics, and the types of tasks involved, rendering it difficult to draw a clear conclusion from these study findings regarding the impacts of different degrees of crosslinguistic differences on L2 morphosyntactic processing. To address this issue, Lago and colleagues (Lago et al., 2021) proposed that studies attempting to precisely assess the effect of crosslinguistic influence would need to compare two groups of speakers “whose L1s differ in the behavior of the target phenomenon” (p. 6). Such a between-group comparison allows for rigorous investigation of L1 influence by helping bypass confounding effects specific to a study design (Jarvis, 2010; Lago et al., 2021).

In light of the inconclusive findings regarding the L2 processing of nontransferable information and motivated by the need for a systematic comparison of learner groups with different L1 backgrounds employing the same experimental set-up, the present study sought to investigate the extent to which two groups of L2 learners exploit information that gives rise to different degrees of crosslinguistic interference during L2 morphosyntactic processing. Specifically, we examined the online processing of the honorific numeral quantifier (NQ) construction in Korean by involving two proficiency-matched learner groups, Mandarin Chinese- and Japanese-speaking learners of Korean. As will be discussed below, Mandarin Chinese (henceforth ‘Chinese’) and Korean pattern alike in that both languages have an honorific NQ, yet only Korean, but not Chinese, instantiates an agreement between a noun phrase (NP) and its NQ in a scrambled structure. In contrast, Japanese contrasts with Korean and Chinese in that it only uses a nonhonorific quantifier both in canonical and scrambled structures, showing crosslinguistic conflicts. Given these crosslinguistic differences across the three languages, this study investigated the ability of the two learner groups to detect agreement

© The Author(s), 2024. Published by Cambridge University Press. This is an Open Access article, distributed under the terms of the Creative Commons Attribution licence (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted re-use, distribution and reproduction, provided the original article is properly cited.

violations under different conditions that present language-common, language-contrasting, and language-specific information during sentence processing. By investigating how bilinguals utilize their L1 and L2 knowledge when processing the honorific NQ construction in Korean, this study seeks to advance our understanding of the role of crosslinguistic influence in bilingual processing mechanisms.

2. Theoretical models on crosslinguistic influence on L2 processing

The effects of crosslinguistic influence, which refers to the influence of an L2 learner's knowledge in one language on the processing of another language, have been extensively investigated across various linguistic domains and among a wide range of language populations (e.g., Altarriba, 1992; De Groot & Nas, 1991; Dijkstra & van Heuven, 2002; Spivey & Marian, 1999). This section provides a review of various models and previous studies on crosslinguistic influence, with the particular focus on morphosyntactic processing.

While ample evidence suggests that L2 learners can achieve native-like processing when target structures are analogous across the learners' L1 and L2 (e.g., Tokowicz & MacWhinney, 2005; Tolentino & Tokowicz, 2011), it remains less clear whether L2 learners can effectively utilize nontransferable L2-specific cues for sentence processing (Trenkic et al., 2014). Some theoretical perspectives propose divergent processing behaviors between L1 and L2 speakers. One such perspective is the blocking or overshadowing account (Ellis, 2006; Ellis et al., 2014; Luk & Shirai, 2009), according to which prior knowledge acquired in the L1 can hamper the acquisition of new information in the L2. This account assumes that previously established L1 knowledge may block or overshadow the formulation of novel forms and functions in the L2, creating attentional biases in acquiring L2 information not present in a learner's L1, including both L1-L2 contrasting and L2-unique cues (Ellis, 2006). Applying this notion to L2 processing, it is hypothesized that L2 learners may encounter processing challenges with L2-unique structures due to the deeply entrenched L1-based parsing routines, which may impede the efficient processing of novel information in the L2.

Similarly, the Morphological Congruency Hypothesis (Jiang et al., 2011) and the Feature Reassembly Hypothesis (Lardiere, 2008) predict persistent difficulties with L2-unique cues. The Morphological Congruency Hypothesis posits that the absence of morphological correspondence between the L1 and L2 can impede processing. According to the Feature Reassembly Hypothesis, L2 learners construct their L2 grammar by reassembling linguistic features from their L1 to those corresponding to the L2. In the context of L2 processing, these hypotheses predict difficulties for L2 learners when the target features lack counterparts in the L1. Consistent with these perspectives, several studies have shown that L2 learners experience great difficulties when target information is unique to the L2 and thus cannot be derived from their L1 – for example, when learners from an article-lacking language background process English articles (e.g., Luk & Shirai, 2009), or when English speakers are required to use gender information in articles to predict an upcoming noun during the online processing of Spanish noun phrases (e.g., Grüter et al., 2012).

In contrast, other views argue that L2-unique information does not always result in severe processing difficulty. One prominent perspective is the competition model (MacWhinney, 2008, 2013), which predicts that structures unique to the target are

less problematic for L2 learners in comparison to structures involving L1-L2 contrasting features. According to this model, L2 learners often attempt to map their L1 structures onto L2 counterparts when they perceive similarities between the languages (MacWhinney, 2013). Consequently, when there is a contrast between L1 and L2 forms, L2 learners struggle with L1-L2 mapping due to competition between interpretations, potentially leading to processing difficulties. Simultaneously, the model suggests that unique L2 structures pose fewer processing difficulties because there is no direct conflict or negative transfer from L1 structures (MacWhinney, 1987, 2008, 2013). Key evidence of this model comes from studies demonstrating that English-speaking L2 learners show native-like sensitivity to grammatical violations of L2-specific information – for example, in the processing of gender agreement information in Spanish (e.g., Tokowicz & MacWhinney, 2005) or case marking information in German (e.g., Jackson & Dussias, 2009) – and studies demonstrating that speakers of Chinese, an article-lacking language, can efficiently capitalize on article information to resolve reference in English sentences during online processing (e.g., Trenkic et al., 2014).

In contrast to the models discussed earlier, some perspectives propose that apart from crosslinguistic influence, L2 processing is largely dependent upon processing-related variables, such as proficiency, language experience, working memory capacity, and the complexity of linguistic stimuli (see Cunnings, 2017, for discussion). According to these approaches, L2 learners may converge on native-like processing when the challenges associated with processing-related factors are mitigated. Extending this idea to the current study's context, it is possible to observe native-like performance among L2 learners if the processing-related burdens are reduced. For example, processing less complex structure may allow L2 learners to effectively integrate L2-specific cues, leading to a more native-like pattern of processing (e.g., Wen et al., 2010).

While no consensus has emerged regarding the use of L2-specific information in syntactic processing, an important caveat is raised when interpreting these previous findings. Because previous studies vary in their experimental methods, linguistic materials, and participant characteristics, it is not straightforward to disentangle the effect of crosslinguistic differences from learner- and task-specific variables (Jarvis, 2010; Lago et al., 2021). To address this issue, this study investigates the processing patterns of two proficiency-matched groups of Chinese- and Japanese-L1 learners, using the honorific NQ construction in Korean. Before presenting our findings, we briefly review how distinctively the target phenomena are realized in Korean, Chinese, and Japanese, giving rise to different degrees of crosslinguistic differences.

3. Numeral quantifier constructions in Korean, Chinese, and Japanese

An NQ encodes quantity information of its associate NP, forming an NQ construction (Ko, 2007; Miyagawa, 1989). With differences in the inventories of NQ forms among Korean, Chinese, and Japanese, all these languages require different types of host NQs based on the properties of the modified NP, including factors such as shape, function, and animacy (Gao & Malt, 2009; Lee, 2000).

To test crosslinguistic influence in L2 processing of Korean sentences, this study focuses on honorific agreement between an NQ and an NP. The phenomenon of honorific NQ-NP

agreement in Korean encompasses a diverse range of properties, including morphosyntactic, semantic, and pragmatic features (J.-B. Kim & Sells, 2007; Kim-Renaud, 2001; Sohn, 2001). The morphosyntactic and semantic aspects of this phenomenon are best captured by the use of distinct forms of NQs to modify an NP depending on its meaning and social status. Consider (1), for example.

(1) Korean

sensayngnim / *ai sey-pwun
teacher / child three-CL_{HON}¹
'three teachers / three children'

In typical circumstances, the honorific NQ, *pwun*, modifies the NP with an honorific status, *sensayngnim* 'teacher'. However, when it co-occurs with the NP *ai* 'child', which lacks honorific features, it leads to an agreement violation. This nonhonorifiable NP should be modified instead by the NQ, *myeong*, which denotes nonhonorific human entities.

While the consistent marking of an honorific NQ for an NP possessing an honorific status suggests the operation to be primarily morphosyntactic (Kim-Renaud, 2001; Sohn, 2001), it is important to note that the morphosyntactic analysis of this phenomenon can sometimes be overridden by pragmatic conditions. For example, in particular social contexts such as customer service in department stores or restaurants, an NP like *ai* 'child' can be paired with the honorific NQ, *pwun*, when the child is treated as a customer (Nam, 2005). This pragmatic constraint equally applies to the general honorific systems in Japanese and Chinese. However, for the purpose of the current research objectives, this study opts to disregard pragmatic factors and instead concentrates on the morphosyntactic and semantic properties of honorific agreement, assuming that an honorific NQ is predominantly associated with an NP carrying an honorific status (e.g., teacher), rather than an NP lacking such status (e.g., child). In these morphosyntactic and semantic analyses, we adopt the approach established by previous studies, characterizing the honorification of an NQ and an NP using the [$\pm$ HON] features (Choi, 2010; Kwon & Sturt, 2016; Martin, 1992; see also Kim & Sells, 2007, for a finer-grained analysis of Korean honorification).

A similar approach can be applied to the Chinese language. Chinese utilizes three main classifiers that modify human nouns: *ge*, *ming*, and *wei* (Gao & Malt, 2009). While *ge* can be used broadly with various types of nouns, *ming* is typically used in formal contexts, specifically denoting individuals of diverse social statuses. In contrast, *wei* is used to modify a person with a respectful meaning, indicating a more polite or honorific way of referring to an individual. Therefore, without taking specific discourse conditions into account, the noun *haizi* 'child' can be modified by the classifiers *ge* or *ming* but not by *wei*, as in (2). This distinction in the use of classifiers in Chinese aligns with the functional and semantic nuances associated with honorific agreement in Korean.

(2) Chinese

san wei laoshu / *haizi
three CL_{HON} teacher / child
'three teachers / three children'

On the contrary, in Japanese, the nonhonorific NQ *nin* is consistently used for both honorifiable (e.g., *sensei* 'teacher') and nonhonorifiable NPs (e.g., *kodomo* 'child'), as shown in (3).

(3) Japanese

sensei / kodomo san-nin
teacher / child three-CL

Therefore, Korean and Chinese converge in honorific agreement realization when an NP and an NQ are linearly adjacent, which mandates them to concord in semantic and functional features: for example, [+Human, +Honorific]. In this case, no crosslinguistic conflict is expected for Chinese-speaking learners when they compute agreement between NPs and the honorific NQ in Korean. In contrast, Japanese-speaking learners are expected to experience crosslinguistic competition when processing honorific NQ agreement in Korean because their L1 uses the same NQ form modifying all types of human referents whereas the L2 requires the use of different NQ forms depending on the honorific status of the associate NP.

Notably, the NQ construction in Korean allows for an NP to be separated from the modifying NQ as long as they form structural sisters of the same constituent (Kang, 2002; Ko, 2007; M. Park & Sohn, 1993), as shown in (4).

- (4) a. Sensayngnim-i cip-ey [haksayng-ul sey-myeng /
teacher-NOM home-LOC student-ACC three-CL /
*sey-pwun]
three-CL_{HON}
chotayhayssta.
invited
'The teacher(s) invited three students to their home.'
- b. Haksayng-ul_i sensayngnim-i cip-ey [t_i sey-myeng /
student-ACC teacher-NOM home-LOC three-CL /
*sey-pwun]
three-CL_{HON}
chotayhayssta.
invited
'The teacher(s) invited three students to their home.'

The sisterhood relationship between an NP and its NQ is captured by a syntactic dependency (Ko, 2007). In (4a), for example, the NQ *two-myeng* 'two-CL' modifies the object NP *haksayng* 'student' in the adjacent position. In (4b), the object NP moves to the sentence-initial position, forming a long-distant dependency with the stranded NQ (Sportiche, 1988). Therefore, the NQ should be associated only with the scrambled object NP *haksayng* 'student' but not with the subject NP *sensayngnim* 'teacher'.

The displacement of an NP from the base position in the NQ construction, or so-called NQ floating (Miyagawa, 1989), is well-attested in case-marking languages like Korean and Japanese. NQ floating in these languages is known to be triggered by several factors, including syntactic, semantic, and pragmatic conditions (e.g., Hamano, 1997; Ko, 2007). To compute nonlocal dependencies in sentences like (4b), comprehenders need to integrate multiple types of information, such as case-marking information for the arguments, their semantic features related to honorific status, and the pragmatic condition that licenses the movement of an NP – namely, the topicalization of the moved NP (Ko, 2007; M. Park & Sohn, 1993; Sohn, 2001). In contrast, Chinese lacks both scrambling and an overt case-marking system, and thus structures involving NQ floating in scrambled sentences like (4b) are not allowed in this language (Kobuchi-Philip, 2007). Instead, the language permits NP fronting through topicalization, resulting in a long-distance dependency between an NP and an NQ, as illustrated in (5).

However, due to the word order differences between Chinese and Korean, Chinese learners of Korean need to acquire the formulation of NP–NQ agreement in the new sentence configurations when they are separated. Specifically, in contrast to the integration of a moved NP with its associate NQ after the main verb in their L1, Chinese-speaking learners should integrate this information in the preverbal complement position in L2 Korean.

- (5) Naxie xuesheng, laoshi qing-le san-ge
those student teacher invited three-CL
‘The teacher(s) invited three students.’

In contrast, Japanese allows for NQ floating, similar to Korean, but without honorific agreement between an NP and an NQ, as illustrated in (6).

- (6) Gakusei-o sensei-ga ie-ni san-nin
Student-ACC teacher-NOM home-LOC 3-CL
shotaishita.
invited
‘The teacher(s) invited three students to their home.’

Based on these crosslinguistic similarities and differences across Chinese, Japanese, and Korean, the current study probed whether L2 learners can detect NP–NQ violations during L2 processing when the target structures include information giving rise to different degrees of crosslinguistic conflicts. In a self-paced reading experiment, we presented participants with Korean sentences in canonical and scrambled word orders where a nonhonorific- or honorific-marked NQ modifies a nonhonorifiable NP in different positions, as in (7). In both the canonical and scrambled conditions, the use of an honorific-marked NQ (e.g., *sey-pwun*) is infelicitous because it consistently modifies an NP lacking honorific status (e.g., *haksayng* ‘student’). Previous studies on L1 processing have demonstrated that Korean and Japanese native speakers exhibit sensitivity to NQ–NP agreement in canonical and scrambled word orders during online processing (e.g., H. Kim, 2018; Suzuki & Yoshinaga, 2013). However, little is known about how different types of crosslinguistic influence, as triggered by the target construction, affect L2 learners’ sensitivity to honorific NQ–NP agreement.

- (7) a. match/mismatch conditions in the canonical word order
Sensayngnim-i cip-ey [haksayng_i-ul sey-myeng_i /
teacher-NOM home-LOC student-ACC three-CL /
*sey-pwun_i]
three-CL_{HON}
chotayhay-se cenyek-ul mekesseyo.
invite-and dinner-ACC ate
‘The teacher(s) invited the three students to home and had dinner.’
- b. match/mismatch conditions in the scrambled word order
Haksayng_i-ul sensayngnim-i cip-ey [t_i sey-myeng_i /
student-ACC teacher-NOM home-LOC three-CL /
*sey-pwun_i]
three-CL_{HON}
chotayhay-se cenyek-ul mekesseyo.
invite-and dinner-ACC ate
‘The teacher(s) invited the three students to home and had dinner.’

The honorific agreement (7a) constitutes transferable information for Chinese speakers whose L1 also instantiates honorific NP–NQ agreement. For sentences involving an NQ floating (7b), Chinese speakers face the challenge of computing honorific agreement in a non-local configuration without relying on their L1 knowledge because Chinese lacks overt case marking and does not permit non-local NP–NQ agreement in the preverbal complement position. In this case, no crosslinguistic competition is expected for Chinese-speaking learners as the long-distance dependency of honorific NQ agreement is unique to Korean. Unlike Chinese speakers, Japanese speakers will experience crosslinguistic competition in both canonical and scrambled sentences due to Japanese employing the same NQ for human NPs irrespective of their honorific status. Therefore, when processing honorific NP–NQ agreement in Korean sentences such as (7), Japanese speakers must deactivate NQ information in their L1, which includes [+Human, –Honorific] feature, while adjusting the NQ feature to align with the L2.

Additionally, we may observe that both Chinese and Japanese speakers exhibit greater difficulty with NP–NQ agreement in scrambled structures compared to canonical structures. The phenomenon of NQ floating in Korean not only entails a long-distance dependency but also increases structural complexity through NP movement within a noncanonical structure. As a result of this increased complexity, L2 learners may encounter challenges when computing agreement between an NP and an NQ in scrambled structures. Table 1 outlines the characteristics of Korean, Chinese, and Japanese regarding honorific NQ–NP agreement, while Table 2 summarizes the predictions of each theoretical model concerning the L2 processing of target structures.

4. Experiment 1

4.1. Participants

The study involved 48 Korean-L2 speakers comprising 24 L1 (Mandarin) Chinese speakers (Chinese group) and 24 L1 Japanese speakers (Japanese group), who were recruited from international student pools at local universities in South Korea. The Chinese group began learning Korean at the mean age of 22.7 ($SD = 4.6$) and had stayed in Korea for an average of 3.3 years ($SD = 2.1$) at the time of testing. The Japanese group had their first exposure to Korean at the mean age of 20.5 ($SD = 4.6$) with an average of 4.6 years of staying in Korea ($SD = 3.4$). Independent samples *t*-tests showed no significant differences between the two groups in their age ($t(46) = 0.655$, $p = .516$), onset of learning Korean ($t(46) = 1.623$, $p = .111$), and length of residence in Korea ($t(46) = -1.615$, $p = .113$).

To establish that the L2 participants in our study possessed an adequate level of proficiency for processing the target sentences,

Table 1. Properties of honorific agreement in canonical and scrambled structures in Korean, Chinese, and Japanese

Language	Canonical structure	Scrambled structure
Korean	Honorific agreement possible	Honorific agreement possible
Chinese	Honorific agreement possible	Scrambling not allowed; no case marking system present
Japanese	Only non-honorific agreement possible	Only non-honorific agreement possible

Table 2. Predictions of theoretical models on the processing of target structures by Chinese- and Japanese-speaking L2 learners

Model	Chinese-speaking L2 learners	Japanese-speaking L2 learners
blocking or overshadowing account	Difficulty in scrambled structure	Difficulty in both canonical and scrambled structure
Morphological Congruency Hypothesis & Feature Reassembly Hypothesis	Difficulty in scrambled structure	Difficulty in both canonical and scrambled structure
Competition model	Less difficulty in scrambled structure (compared to Japanese-speaking learners)	Difficulty in both canonical and scrambled structure

we attempted to recruit highly advanced learners who had achieved a level within the two top tiers (Level 5 and Level 6) in the Test of Proficiency in Korean (TOPIK), an official Korean proficiency test. The descriptors provided by the test developers indicate that these levels correspond to language proficiency suitable for using Korean in professional or academic settings. Therefore, we assumed that our L2 participants would possess the necessary proficiency to comprehend the linguistic stimuli presented in our study.

We also measured the current proficiency levels of our participants via self-rated proficiency and a modified version of TOPIK. For the self-rated proficiency, participants were asked to provide ratings of their Korean fluency on a ten-point scale (1 indicating the lowest and 10 indicating the highest fluency) in the four domains of reading, writing, listening, and speaking. The mean ratings for the Chinese group were 8.1 ($SD = 1.1$) in reading, 6.6 ($SD = 1.3$) in writing, 7.8 ($SD = 1.5$) in listening, and 7.0 ($SD = 1.3$) in writing. The ratings for the Japanese group were 7.8 ($SD = 1.3$) in reading, 6.9 ($SD = 1.6$) in writing, 8.3 ($SD = 1.0$) in listening, and 7.3 ($SD = 1.6$) in writing. None of the ratings had statistical differences between the two language groups (see Table 3). In the modified version of the TOPIK (Jeong, 2017), participants read 20 sentences with blanks and chose the correct morphemes or words for each blank from four choices. The mean accuracy scores were 62.9% ($SD = 17.7$) for the Chinese group and 66.9% ($SD = 13.6$) for the Japanese group, which were not statistically different. The results of the self-ratings and the TOPIK scores indicate that the two language groups were closely matched in their Korean proficiency. We factored in the individual variance of the proficiency scores by including the TOPIK scores as a continuous variable in the analysis model.

In addition to the L2 groups, 24 Korean speakers (Korean group) served as a control group. The study was approved by the ethics committee of University of Hawaii. All participants received monetary compensation for their participation. Participant information is summarized in Table 3.

4.2. Materials and procedure

During a single visit to the lab, each participant completed a self-paced reading task, followed by a language background questionnaire and then an acceptability judgment task as measurement of their knowledge of the target structures.

Acceptability judgment task

This task was conducted to ensure that participants had knowledge of the NP-NQ agreement in canonical and scrambled sentences in Korean. The task consisted of 20 quadruplets of Korean sentences in which the NP-NQ agreement (match, mismatch) and word order (canonical, scrambled) were manipulated as in (7). The experimental items were identical to those used in the self-paced reading task, with the exception that each sentence was truncated after the first verb (e.g., *teachers-NOM home-LOC students-ACC three-CL/*three-CL_{HON} invited* 'The teachers invited the three students to home'). Since an NQ in the scrambled condition can be associated either with the fronted NP or the subject NP, we ensured that all experimental sentences included an adverbial phrase following the subject NP to avoid this ambiguity. The experimental items were counterbalanced across the four conditions using the Latin square design (5 items per condition), and each participant read only one of the four types of a single item. Experimental items were intermixed with 30 distractor items, including sentences with an honorific-marked NQ modifying the subject NP, canonical and scrambled sentences without NQs, and simple intransitive sentences. Half of the distractor items were grammatical, and half were ungrammatical. The items were pseudorandomized to ensure that experimental items in the same condition did not appear in a row. The experimental sentences are provided in Appendix A as online supplementary materials to this paper.

During the task, participants read sentences on a computer screen and provided judgment ratings for given sentences on a scale from 1 (very unnatural) to 4 (very natural), along with an additional option of "I don't know". Each item was presented on

Table 3. Information of participants in Experiment 1

	Korean group	Chinese group	Japanese group	P value (Chinese vs. Japanese)
Age	26.9 (5.8)	26.0 (3.9)	25.1 (4.9)	.428
Mean years of staying in Korea	-	3.3 (2.1)	4.6 (3.4)	.369
Mean onset of studying Korean	-	22.7 (4.6)	20.5 (4.6)	.261
Mean self-ratings (1-10)	9.3 (0.8)	7.4 (1.0)	7.6 (1.0)	.582
Mean TOPIK score (%)	-	62.9 (17.7)	66.9 (13.6)	.390

Note. The values in the parentheses indicate standard deviations

a single page, and participants were disallowed from returning to previous pages to correct their answers. Prior to the task, participants received oral and written instructions on the task and worked through two practice items. The procedure took approximately 10 minutes for the control group and 15 minutes for the L2 groups.

Self-paced reading task

The materials for the self-paced reading task consisted of 20 sets of experimental items counterbalanced across four conditions, as in (7). We focused on two specific regions as regions of interest: the phrase containing the NQ (e.g., *sey-myeng* / **sey-pwun* '3-CL') and the subsequent verb (e.g., *chotayhay-se* 'invite-and'). The experimental items were intermixed with 30 fillers with various types of structures.

Participants individually completed the task in a quiet lab under the supervision of a research assistant. The task was run using a moving window display (Just et al., 1982) via the Ibex Farm platform (Drummond, 2013). The items were presented randomly by the program. Each trial began with a series of dashes appearing on the screen, marking the positions of words to appear. When a participant pressed the spacebar, the sentence-initial word appeared in the first dash position. The next spacebar press concealed the first word while revealing the next word. This button-press repeated until participants read the last word. Each target sentence was presented word by word in 7 regions (Rs). A true-false question followed each sentence to check participants' understanding of the previous sentence. For example, the question for the sentences in (7) was "Sensayngnim-i haksayng-kwa cemsim-ul mekesseyo? (Did the teachers have lunch with the students?)". Participants responded to the question by clicking on a "yes" or "no" option that appeared below the question. The position of the correct answer was counterbalanced across items. Participants worked through five practice items before the main experiment. The whole task took about 15 minutes for the control group and 20 minutes for the L2 groups.

4.3 Results and discussion

In order to establish that our participants had a solid understanding of honorific NQ-NP agreement, we first present the findings

from the acceptability judgment task, followed by the results from the self-paced reading task.

Acceptability judgment task

Prior to data analysis, we screened for the choice of "I don't know" option (0.4% in the Chinese group, 0.2% in the Japanese group, and 0.3% in the Korean group). Figure 1 presents the judgment ratings for the experimental sentences after removing these responses. All three groups showed similar judgment patterns: They gave higher acceptability ratings for the sentences in the match than the mismatch condition for both canonical and scrambled word order conditions.

To scrutinize the judgment patterns in detail, we conducted an ordinal regression analysis, which is appropriate for analyzing ordinal data such as acceptability ratings (Verissimo, 2021). In this analysis, we generated a cumulative link mixed effects model (Christensen, 2015), using the logit link function and flexible thresholds. The model included fixed effects of *Group* (Chinese group, Japanese group, Korean group), *Agreement* (match, mismatch), and *Word order* (canonical, scrambled), as well as random effects of participants and items. Because our primary objective was to investigate participants' sensitivity to agreement violations in distinct structures by different language groups, we nested the effects of *Agreement* and *Word order* within each group. We initially constructed the maximal random effects structure (Barr et al., 2013). However, due to potential Type I error, we simplified the random effects structure through model comparisons using Akaike Information Criterion (AIC) and likelihood ratio tests (Matuschek et al., 2017). The final model formula for each region was Judgment ratings ~ Group/(Agreement*Word order) + (1 + Agreement*Word order | participant) + (1 + Group | item). When assessing the effect of *Group*, we treated the Korean group as a baseline.² The modelling was conducted using the ordinal package CLMMs (Christensen, 2015) in R version 4.0.3 (R Core Team, 2020).

The model output is summarized in Appendix B as online supplementary materials to this paper. To address our research questions, we focus on the effect of *Agreement* and its interaction with *Word order* for each group. In the Korean group, we found a main effect of *Agreement*, without its interaction with *Word order* ($\beta = -1.433$, $SE = 0.295$, $p < .001$). These results indicate that this

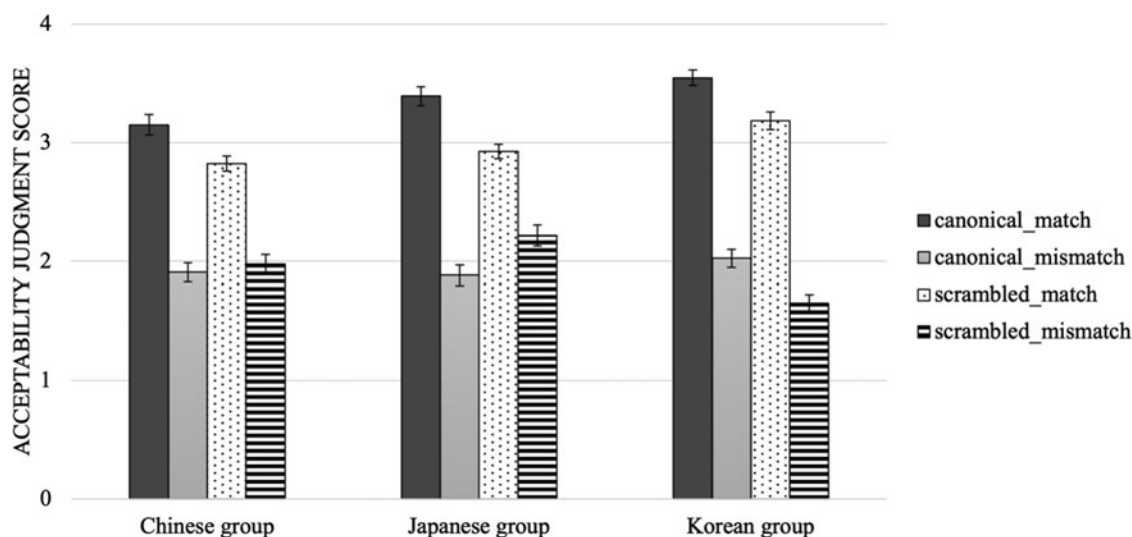


Figure 1. Mean acceptability judgment scores in Experiment 1. Error bars denote 95% confidence intervals.

group accepted the sentences in the match condition more often than those in the mismatch condition for both canonical and scrambled word order. For both Chinese and Japanese groups, there was a robust effect of *Agreement* (Chinese group: $\beta = -2.956$, $SE = 0.350$, $p < .001$; Japanese group: $\beta = -3.347$, $SE = 0.361$, $p < .001$), qualified by its interaction with *Word order* (Chinese group: $\beta = 1.324$, $SE = 0.614$, $p = .031$; Japanese group: $\beta = 2.823$, $SE = 0.624$, $p < .001$). Post-hoc analyses for each word order condition showed similar results across the two groups. For the Chinese group, a significant effect of *Agreement* emerged both in the canonical ($\beta = -3.391$, $SE = 0.555$, $p < .001$) and the scrambled conditions ($\beta = -2.281$, $SE = 0.408$, $p < .001$), with higher acceptance ratings for the sentences in the match than the mismatch condition. Likewise, the model for the Japanese group showed the main effect of *Agreement* both in the canonical ($\beta = -5.219$, $SE = 0.716$, $p < .001$) and the scrambled conditions ($\beta = -1.979$, $SE = 0.426$, $p < .001$), with higher ratings in the match than the mismatch condition. As indicated by the interaction, however, the effect of *Agreement* for these learner groups was more pronounced in the canonical condition compared to the scrambled condition, possibly due to the relatively low acceptance rate for sentences in the scrambled condition.

In sum, both control and L2 groups showed sensitivity to the violation of honorific NQ-NP agreement in canonical and scrambled sentences in the acceptability judgment task. These findings suggest that the L2 participants had native-like knowledge of the syntactic constraints underlying honorific NQ agreement. While the study initially did not consider specific pragmatic conditions, it is important to acknowledge the potential influence of participants' presuppositions about discourse situations while completing the tasks. For example, participants might have associated an honorific NQ with children, considering situations where they are treated as customers. However, the consistent effects of agreement observed across all three groups suggest that any such assumptions did not significantly impact the participants' judgments of the experimental sentences. These findings allow us to establish that any different reading time patterns between the native speaker and L2 groups in the self-paced reading task should not be

ascribed to the L2 learners' insufficient knowledge of the target structures.

Self-paced reading

We first inspected participants' accuracy rates for the truth-value verification questions. The mean accuracy rates were 89.3% ($SD = 4.6$) in the Chinese group, 89.4% ($SD = 3.7$) in the Japanese group, and 90.7% ($SD = 3.6$) in the Korean group. A one-way ANOVA revealed no statistical differences in the accuracy scores across the three groups ($F < .1$). These results indicate that all participants generally paid attention to the sentence meanings during the task. Trials with incorrect responses were removed for further analysis.

Prior to data analysis, the participants' reading time (RT) data were trimmed by eliminating extreme values beyond 2.5 standard deviations from the mean (2.7% of the entire data). The remaining RTs were log-transformed to meet the normal distribution requirement (Ratcliff, 1993). Subsequently, the RTs were residualized to control for individual variations in reading speed and word length across conditions. Following Trueswell et al. (1994), we calculated residual RTs by subtracting the predicted RT, which was obtained based on the linear mixed-effects model including a fixed effect of the number of characters within a region and a random effect of participant, from the observed RT for each participant.

Figures 2, 3, and 4 show residual RT profiles for each group. (Raw RTs are provided in Appendix C.) We focused on the following two regions as the main analysis frames: the NQ as the critical region (R4, e.g., *sey-myeng* / **sey-pwun* 'three-CL / three-CL_{HON}') in which agreement checking is assumed to occur, and the following word as the spillover region (R5, e.g., *chotayhay-se* 'invite-and'). Inspection of the graph for the Korean group suggests that the agreement effect (i.e., longer RTs in the mismatch versus match condition) begins to emerge in the spillover region (R5). For the Chinese group, the agreement effect appears to emerge in R5, but only for the canonical word order conditions. In contrast, the Japanese group shows consistent reading times across the four conditions throughout all regions.

To compare each group's reading time patterns in detail, we created linear mixed-effects regression models for the critical and spillover regions. The models were constructed in the same

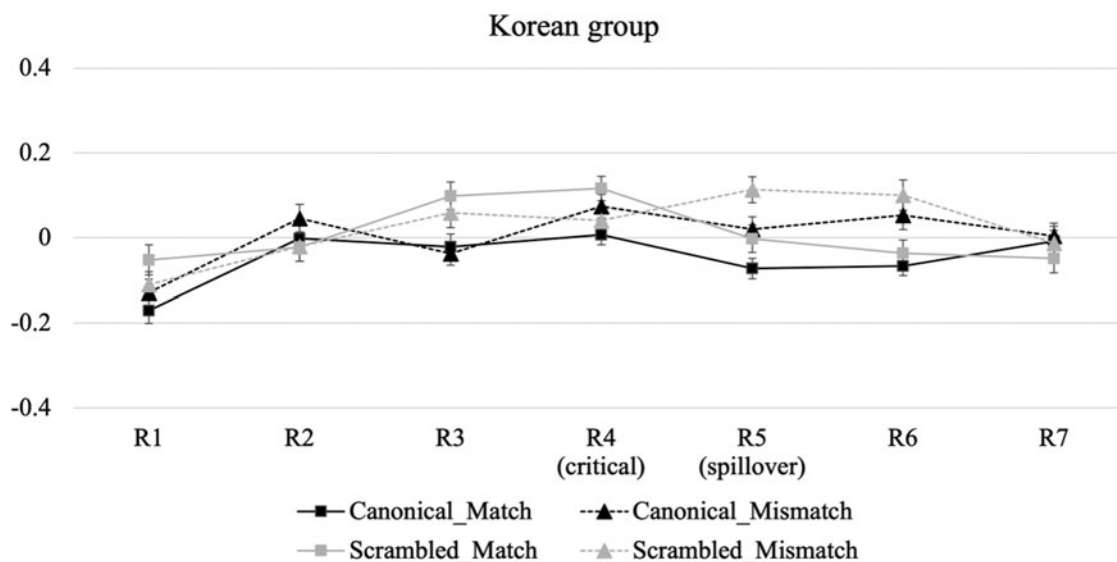


Figure 2. Residual reading time profiles for the native speaker group in Experiment 1. Error bars denote 95% confidence intervals.

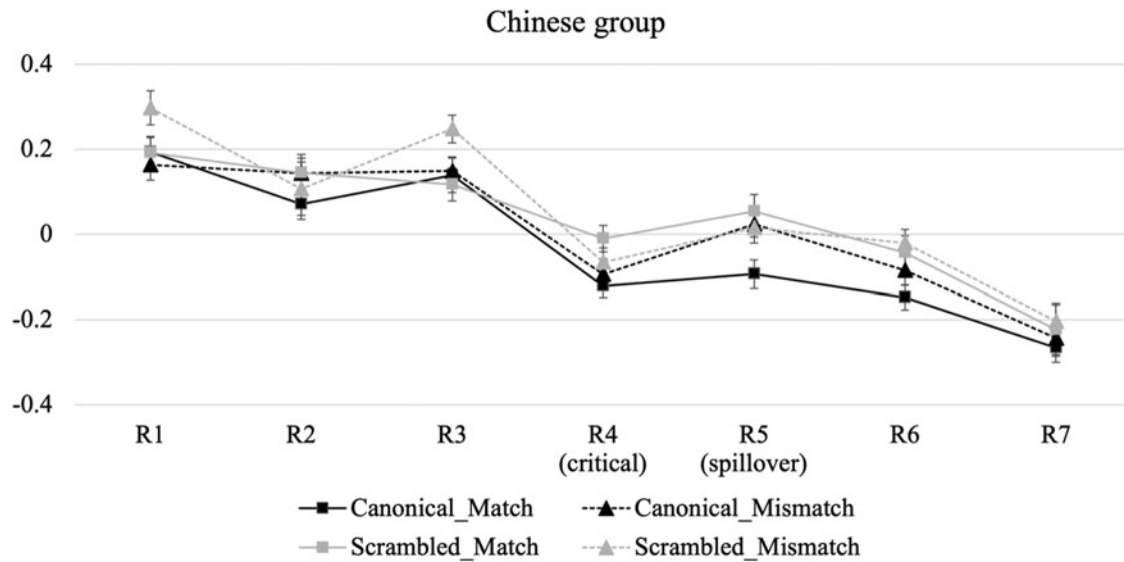


Figure 3. Residual reading time profiles for the Chinese-speaking L2 group in Experiment 1. Error bars denote 95% confidence intervals.

manner as in the acceptability judgment task. The final model formula for each region was $\text{Residual RT} \sim \text{Group}/(\text{Agreement} * \text{Word order}) + (1 + \text{Agreement} * \text{Word order} | \text{participant}) + (1 + \text{Group} | \text{item})$. When assessing the effect of *Group*, we treated the Korean group as a baseline.

The model outcomes for the critical and spillover regions are reported in Appendices D and E, respectively. As in the acceptability judgment task, we focus on effects associated with agreement and its interaction with word order for each group. The model for the critical region showed distinct results across the three groups. For the Korean group, there was a significant interaction between *Agreement* and *Word order* ($\beta = -0.161$, $SE = 0.063$, $p = .012$). Pairwise comparisons within each word order condition revealed no significant RT difference between the match and the mismatch conditions, both for the canonical condition ($\beta = 0.078$, $SE = 0.041$, $p = .065$) and the scrambled condition ($\beta = -0.083$, $SE = 0.044$, $p = .062$). For the Chinese group,

there was a marginal interaction between *Agreement* and *Word order* ($\beta = -0.121$, $SE = 0.067$, $p = .073$). Post-hoc tests revealed no significant effect of agreement both in the canonical ($\beta = 0.051$, $SE = 0.054$, $p = .355$) and in the scrambled word order condition ($\beta = -0.062$, $SE = 0.058$, $p = .290$). The proficiency scores from the TOPIK added to the model did not interact with agreement ($\beta = -0.005$, $SE = 0.008$, $p = .503$). Finally, for the Japanese group, there was no main effect of *Agreement* ($\beta = -0.008$, $SE = 0.033$, $p = .800$) and no interaction with *Word order* ($\beta = -0.093$, $SE = 0.069$, $p = .175$). In addition, participants' proficiency scores did not interact with agreement ($\beta = -0.009$, $SE = 0.008$, $p = .247$). In sum, we found no evidence that the participants in the three groups showed sensitivity to the agreement violation in the critical region.

Turning to the analysis of the spillover region, we found a significant effect of *Agreement* for the Korean group ($\beta = 0.125$, $SE = 0.035$, $p < .001$), induced by longer RTs in the mismatch

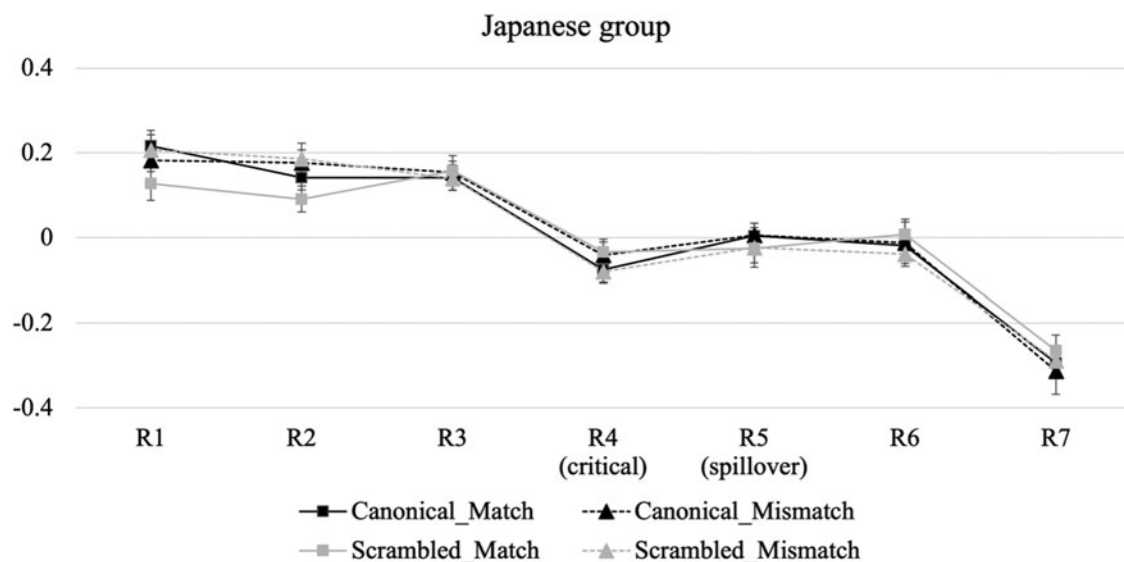


Figure 4. Residual reading time profiles for the Japanese-speaking L2 group in Experiment 1. Error bars denote 95% confidence intervals.

than the match condition. This effect did not significantly interact with *Word order* ($\beta = -0.005$, $SE = 0.075$, $p = .951$), suggesting their sensitivity to the agreement violation both in canonical and scrambled sentences. For the Chinese group, we found a marginal interaction between *Agreement* and *Word order* ($\beta = -0.154$, $SE = 0.080$, $p = .058$). Post-hoc tests revealed a significant effect of *Agreement* in the canonical ($\beta = 0.127$, $SE = 0.056$, $p = .025$), but not in the scrambled word order condition ($\beta = -0.048$, $SE = 0.059$, $p = .418$). The proficiency scores from the TOPIK added to the model did not interact with *Agreement* ($\beta = -0.008$, $SE = 0.008$, $p = .305$). These results indicate that the Chinese group showed sensitivity to the agreement violation only in the canonical word order condition, regardless of L2 proficiency. Finally, for the Japanese group, there was no effect of *Agreement* ($\beta = 0.032$, $SE = 0.039$, $p = .412$) and no interaction with *Word order* ($\beta = -0.034$, $SE = 0.083$, $p = .685$). In addition, participants' proficiency scores did not interact with *Agreement* ($\beta = 0.010$, $SE = 0.010$, $p = .349$). These results suggest the Japanese participants' insensitivity to the violation, irrespective of the conditions and proficiency levels.

In sum, the three groups exhibited different processing patterns in the spillover region. The Korean group showed the agreement effect both in canonical and scrambled sentences, whereas the Japanese group showed no evidence of sensitivity to the agreement violations, either in the canonical or in the scrambled sentences. Given the crosslinguistic conflict between Korean and Japanese in the realization of agreement between an NP with an honorific status and an NQ, the results of the Japanese group suggest that language-contrasting features pose a major processing difficulty for L2 learners. In contrast, the Chinese group showed sensitivity to the agreement effect only in canonical sentences. Since both Korean and Chinese allow for honorific agreement between an NP and an NQ in the canonical structure, the Chinese group's processing pattern in this structure indicates that language-common features do not give rise to a processing challenge for L2 learners. In the scrambled structure, the Chinese group failed to detect the agreement violation, which seems to suggest that language-unique properties impede native-like processing for L2 learners.

Caution should be raised, however, against drawing a firm conclusion from this result. It should be noted that the processing of the NQ-NP agreement in the scrambled structure required the Chinese group not only to use the L2-unique information but also to do so in the noncanonical, scrambled structure. As mentioned in the literature review, floating NQ introduces additional syntactic complexity through NQ movement. Previous research shows that this increased complexity can pose greater challenges for child speakers compared to NQ-NP agreement in canonical sentences (e.g., Suzuki & Yoshinaga, 2013). In particular, the scrambled word order in Korean is known to present considerable difficulties for nonnative speakers (e.g., K. Kim, 2014; Song et al., 1997), which may have been the case for our L2 participants.

However, such difficulty can be alleviated when contextual information pragmatically licenses the use of scrambling (e.g., Kaiser & Trueswell, 2004; M. H. Kim, 2019). In Korean, scrambling typically occurs in contexts where the object NP has been previously mentioned and is thus defocused as old information, while the following subject NP carries the focus as new information. The processing of Korean scrambling in accordance with this information structure has been well-documented in the literature. For example, M. H. Kim (2019) found that both native Korean speakers and Chinese-speaking L2 learners of Korean exhibited

facilitated processing when a scrambled NP in the target sentence was introduced as a topic in the preceding context. These results suggest that the presence of contextual cues can help readers process scrambled sentences by mitigating the difficulties associated with this syntactic structure. Moreover, previous research suggests that L2 learners can approach native-like processing when the processing burdens are alleviated (e.g., Cunnings, 2017; Wen et al., 2010). It thus remains uncertain whether the Chinese speakers' divergent processing pattern in the scrambled structure was due to their inability to use L2-unique information or to the general difficulty of processing noncanonical structures without contextual information that licenses the use of scrambling.

To address this issue, we conducted another self-paced reading study (Experiment 2) in which we added a context sentence preceding the target sentence. In each context sentence, the fronted NP in the scrambled condition was introduced as discourse-new information, thus making object NP scrambling in the following sentence pragmatically felicitous. If the Chinese group's insensitivity to the agreement violation in the scrambled condition in Experiment 1 resulted from their difficulties associated with the processing of scrambled sentences, not from their inability to use L2-specific information, they will show increased sensitivity to the NP-NQ violation in the scrambled condition when the context sentence licenses object-fronting in the ensuing sentence.

5. Experiment 2

5.1. Participants

We recruited another group of Chinese-L1 learners of Korean ($n = 28$), who did not participate in Experiment 1. Participant information is summarized in Table 4. When comparing this group with the two L2 groups in Experiment 1, there were no significant differences among the three groups in terms of the length of residence in Korea, onset age of studying Korean, self-reported Korean proficiency, and TOPIK scores (all $ps > .1$), confirming that the learner characteristics were closely matched across the two experiments.

5.2. Materials and procedure

The participants completed the self-paced reading task, followed by a language background questionnaire and the modified version of the TOPIK during a single lab visit. The materials for the self-paced reading task were identical to those in Experiment 1, except for an inclusion of a context sentence, as illustrated in (8).

(8) Context sentence:

Wuli-hakkyo haksayng-tul-un chincellhayyo.
our-school student-PL-TOP kind
'The students in our school are kind.'

Target sentence:

Haksayng_i-ul sensayngnim-i cip-ey [t_i sey-myeng_i /
student-ACC teacher-NOM home-LOC three-CL /
*sey-pwun_i]
three-CL_{HON}
chotayhay-se cenyek-ul mekesseyo.
invite-and dinner-ACC ate
'The teacher(s) invited the three students to home and had dinner.'

Table 4. Information of participants in Experiment 2

Group	Age	Mean years of staying in Korea	Onset of studying Korean	Averaged self-reported Korean proficiency (1-10)	TOPIK score (%)
Chinese group	27.2 (4.1)	3.1 (1.8)	21.2 (3.1)	7.6 (0.9)	64.4 (15.3)

Note. The values in the parentheses indicate standard deviations

The experimental sentences were intermixed with the same fillers used in Experiment 1 with a context sentence added to each filler sentence. Because the primary objective of Experiment 2 was to probe whether the Chinese group can detect the NP-NQ violation in the scrambled structure in the presence of the context sentence, we only included scrambled structures, which were aligned across the NP-NQ match and mismatch conditions. As in Experiment 1, the NQ (R4, e.g., *sey-myeng* / **sey-pwun*) was analyzed as the critical region and the following word (R5, e.g., *chotayhay-se*) as the spillover region. The experiment proceeded in the same manner as in Experiment 1.

5.3 Results and discussion

Data trimming and analysis procedures were identical to those in Experiment 1. The Chinese group's accuracy rates for the comprehension check-up questions were 88.5% ($SD = 2.5$), which did not statistically differ from those of the Korean group ($p = .384$) or the L2 groups in Experiment 1 (all $ps > .9$), suggesting that these participants equally attended to the task. Trials with incorrect responses were removed for further analysis. We further removed RTs that fell beyond 2.5 standard deviations from the mean (2.5% of the entire data). The remaining RTs were log-transformed and residualized.

Figure 5 shows residual RT profiles for the Chinese group. (Raw RTs are provided in Appendix F.) The reading time patterns for this group indicated a noticeable RT gap between the match and the mismatch conditions in the critical (R4) and the spillover region (R5).

To statistically compare the RTs between the match and the mismatch conditions in the critical and spillover regions, we

created linear mixed-effects models that included the fixed factor of *Agreement* (centered and deviation coded: match: $-.5$, mismatch: $.5$) and the random factors of participants and items. The maximum random effects structure was simplified to address the issue of Type I error. The model formula for each region was Residual RT $\sim$ Agreement + (1 + Agreement | participant) + (1 | item).

The model outcomes for the critical and spillover regions are presented in Appendix G. The models showed a robust main effect of *Agreement* both in the critical ($\beta = 0.078$, $SE = 0.032$, $p = .016$) and spillover regions ($\beta = 0.083$, $SE = 0.032$, $p = .009$), with longer RTs in the mismatch than in the match condition.

We also compared reading time patterns between the two Chinese groups in Experiment 1 and Experiment 2. Linear mixed-effects models were fit to residual RTs for the scrambled conditions in the critical and spillover regions, including *Experiment* (Experiment 1: $-.5$, Experiment 2: $.5$) and *Agreement* (match: $-.5$, mismatch: $.5$) as fixed factors as well as random factors of item and participant. As in the previous models, we simplified the maximal random effects structure. The final model formula was Residual RT $\sim$ Experiment*Agreement + (1 + Agreement | participant) + (1 | item).

The models revealed a significant interaction between *Experiment* and *Agreement*, both in the critical region ($\beta = 0.140$, $SE = 0.059$, $p = .018$) and in the spillover region ($\beta = 0.130$, $SE = 0.063$, $p = .038$). These results indicate that only the Chinese speakers in Experiment 2 showed sensitivity to the agreement violation in the scrambled sentences, suggesting that this group was able to exploit the L2-unique feature (i.e., honorific NP-NQ agreement in the scrambled structure) when scrambling was pragmatically licensed by the context sentence. These

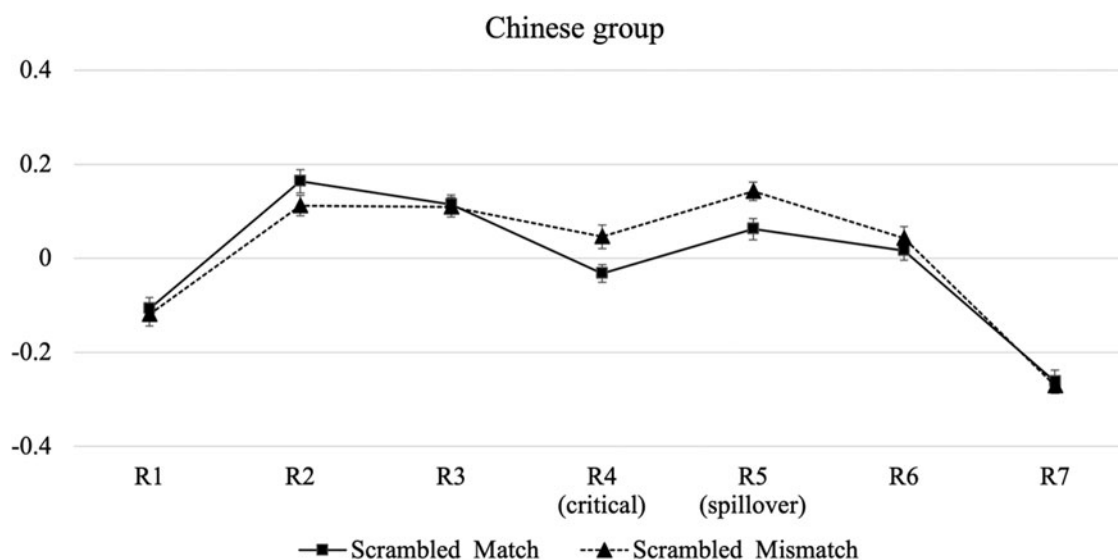


Figure 5. Residual reading time profiles for the Chinese-speaking L2 group in Experiment 2. Error bars denote 95% confidence intervals.

findings confirm that the Chinese speakers' insensitivity to the agreement violation in the scrambled condition in Experiment 1 was mainly due to their cognitive difficulty of processing the scrambled sentences without contextual information, not due to their inability to use the L2-unique information.

6. General discussion

The primary objective of the present study was to test the effect of different degrees of crosslinguistic influence in L2 syntactic processing. In the self-paced reading experiments, we presented two groups of L2 learners, Chinese- and Japanese-speaking nonnative speakers, with Korean sentences requiring agreement between a numeral quantifier (NQ) and a noun phrase (NP) in canonical and scrambled structures. Based on the crosslinguistic similarities and differences among Chinese, Japanese, and Korean in terms of the realization of honorific NP-NQ agreement, we investigated whether Chinese and Japanese speakers would be able to use language-common, language-contrasting, and language-specific information during L2 syntactic processing. In this section, we discuss the results from each L2 group in light of the theoretical approaches to L2 sentence processing reviewed in the literature review section.

In Experiment 1, the Japanese group did not show sensitivity to NP-NQ agreement violations in both canonical and scrambled structures. This finding aligns with our predictions based on the competition model, which suggests that the honorific NQ-NP agreement in Korean presents a crosslinguistic competition for Japanese speakers. Unlike Korean, Japanese uses a nonhonorific NQ to modify both honorifiable and nonhonorifiable human NPs. As a result, Japanese learners might incorrectly activate NQ in their L1, which lacks an honorific feature, leading to competition between the contrasting NQ information types regarding the honorific feature in their L1 and L2. Consequently, they may have experienced increased cognitive demands in inhibiting the influence of their L1 knowledge, making them unable to detect the agreement violation in our study. These findings are consistent with previous research demonstrating L2 learners' persistent challenges in processing target structures that are in conflict with their L1 patterns (e.g., Andersson et al., 2019; Ionin et al., 2021).

Unlike for the Japanese group, the honorific NP-NQ agreement in the canonical structure constitutes a language-common property for the Chinese group, who successfully detected the agreement violation in the canonical condition. The two learner group's divergent patterns reflect the effect of crosslinguistic influence in L2 processing. This result supports the main idea of the competition model that L2 learners attempt to carry over L1 properties that are perceived to match the L2 counterpart, and when there is no crosslinguistic competition, the transferred L1 information serves as a support factor in L2 processing (MacWhinney, 2013).

However, the Chinese group failed to show sensitivity to agreement violations in the scrambled condition, suggesting that L2-unique features may cause a processing difficulty. However, when we added a context sentence that provided pragmatic affordances for object scrambling in the test sentence in Experiment 2, another group of Chinese speakers, whose proficiency was closely matched with that of the L2 groups in Experiment 1, successfully detected the agreement violation in the scrambled structure. These findings suggest that the result from the Chinese group in Experiment 1 was primarily ascribed to their difficulty of processing noncanonical structures without contextual information, not necessarily indicating their inability to use language-specific information.

It should be noted that the successful processing of the Chinese group in Experiment 2 can be interpreted as a result of the interaction of multiple factors, rather than solely attributing it to the existence of L2-unique cues. As mentioned earlier, the use of contextual information can facilitate the processing of scrambled sentences in Korean by providing pragmatic affordance for scrambling (M. H. Kim, 2019). Therefore, it is possible that the Chinese group in our study experienced reduced processing difficulty when they encountered the scrambled sentences preceded by the context sentence, allowing them to better detect the agreement violation. Moreover, the movement of an NP presented as a topic is a common phenomenon in Chinese (Huang et al., 2009), which may have positively influenced the processing of the Chinese speakers in our study. However, it is important to acknowledge that these factors are confounded in the current design, and future research is needed to examine their individual contributions. It appears that these factors interact with one another, jointly helping the Chinese speakers process the L2-unique cues (i.e., nonlocal NP-NQ agreement) in this study.

Combining the results from Experiment 1 and Experiment 2, the Chinese speakers' processing of the Korean NP-NQ agreement in canonical and scrambled structures provides evidence that L2 learners can make online use of nontransferable, L2-specific information, as well as language-common information, consistent with the competition model. This finding is in line with previous findings that L2 learners can process target information not instantiated in their L1 in a native-like way (e.g., Herbay et al., 2018; Jackson & Dussias, 2009; S. H. Park & Kim, 2023; Trenkic et al., 2014). In addition, our findings advance further in demonstrating that L2-specific features can be utilized for L2 sentence processing when the features require integration of multiple sources of information. In order to check the honorific agreement between an NP and an NQ in the scrambled structure in Korean, comprehenders should compute the syntactic dependency between the dislocated items as well as checking the honorific status of the NP. Despite this dual informativity, our Chinese participants showed sensitivity to the agreement violations, suggesting that language-specific properties requiring an integration of multiple cues, including grammatical constraints (i.e., syntactic dependency) and semantic aspects (i.e., honorific status) of an NP-NQ relationship, can still be processed in a target-like way.

The current findings also provide compelling evidence for the robust effect of crosslinguistic influence in L2 syntactic processing. Previous research has characterized bilingual memory representations as the integrated and non-language-selective system that allows bilinguals and L2 learners to access linguistic information from both languages in a parallel manner at every level of representation (Bernolet et al., 2009; Dijkstra & van Heuven, 2002; Hartsuiker et al., 2004). Consistent with this integrated model, the crosslinguistic influence observed in the current study indicate that our L2 participants had constant access to both their L1 and L2 knowledge during their L2 processing. In addition, our findings advance the understanding of the integrated model by demonstrating that the effect of crosslinguistic influence varies depending on the status of the target information shared across languages (i.e., L1-L2 overlapping information, L1-L2 contrasting information, L2-unique information). Moreover, our results highlight the impact of shared linguistic representations in L2 learners on the integration of various types of information, including morphosyntactic, semantic, and pragmatic aspects of honorific NQ-NP agreement.

Despite the evidence of crosslinguistic influence, we also acknowledge the presence of other potential factors that might have affected the results. One such factor is the difference in the inventories of specific classifiers for human entities between Chinese and Korean. As outlined in the literature review, Chinese allows for three main classifiers (*ge*, *ming*, and *wei*) to modify human nouns, while Korean utilizes two classifiers (*myeong*, *pwun*). This discrepancy could have contributed to the Chinese speakers' insensitivity to agreement violations in the scrambled condition in Experiment 1. However, it is important to note that this mapping difference did not significantly affect their processing in the canonical condition, where participants successfully detected the agreement violations.

Second, as a reviewer pointed out, the additional group of native Mandarin speakers in Experiment 2 was presented with scrambled structures alone without the inclusion of canonical structures. This discrepancy in experimental context could have introduced unique experience and strategies compared to processing both canonical and scrambled structures. Future research should address this potential limitation by controlling the experimental contexts provided to participants.

Third, as noted by another reviewer, the current study included a limited number of filler items, which could not have completely diverted participants' attention from the objective of the study. To address this issue, future research should include a larger number of filler items to ensure that participants' attention is directed away from the specific research focus.

Finally, while this study did not account for potential effects of pragmatic conditions that could override the conventional analysis of morphosyntactic agreement between an NQ and an NP, it remains possible that some participants presupposed certain pragmatic contexts, potentially impacting their behavior during the tasks. This issue may be particularly crucial considering that individuals without honorific status, such as a child, can still be honorified in customer service settings, not only in Korean but also in Chinese and Japanese cultures. Therefore, further research is necessary to clarify to what extent such contextual information may influence L2 learners' processing of honorific agreement in Korean.

Despite these limitations, the current study showed that language-common, language-contrasting, and language-specific features have different degrees of impacts on L2 syntactic processing. Overall, these results lend credence to the competition model, which predicts greater processing difficulty for structures giving rise to crosslinguistic conflicts than for structures shared across languages and structures uniquely found in an L2. Further research is required to investigate whether and how the effects of language-contrasting and language-specific information can generalize to L2 learners with various L1 backgrounds across different language phenomena.

Supplementary Material. For supplementary material accompanying this paper, visit <https://doi.org/10.1017/S1366728924000269>

Data availability. The data that support the findings of this study are openly available in the OSF repository, at https://osf.io/q3pk8/?view_only=81a39e3db40548e2ab0b4cc0de657370.

Notes

¹ Abbreviations used in the glosses: ACC = accusative marker; CL = classifier; HON: honorific marker; LOC = locative marker; NOM = nominative marker; PL = plural marker; TOP = Topic marker

² We thank an anonymous reviewer for bringing up this point.

References

- Altarriba, J. (1992). The representation of translation equivalents in bilingual memory. In Harris, R. J. (Ed.), *Cognitive processing in bilinguals* (pp. 157–174). Amsterdam: Elsevier.
- Andersson, A., Sayehli, S., & Gullberg, M. (2019). Language background affects online word order processing in a second language but not offline. *Bilingualism: Language and Cognition*, 22(4), 802–825. <https://doi.org/10.1017/S1366728918000573>
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Bernolet, S., Hartsuiker, R. J., & Pickering, M. J. (2009). Persistence of emphasis in language production: A cross-linguistic approach. *Cognition*, 112(2), 300–317. <https://doi.org/10.1016/j.cognition.2009.05.013>
- Choi, K. (2010). Subject honorification in Korean: In defense of Agr and Head-Spec agreement. *Language Research*, 46(1), 59–82.
- Christensen, R. H. B. (2015). *Ordinal: Regression models for ordinal data*. R package version, 28.
- Cummings, I. (2017). Parsing and Working Memory in Bilingual Sentence Processing. *Bilingualism: Language and Cognition*, 20(4), 659–678. <https://doi.org/10.1017/S1366728916000675>
- De Groot, A. M. B., & Nas, G. (1991). Lexical representation of cognates and noncognates in compound bilinguals. *Journal of Memory and Language*, 30(1), 90–123. [https://doi.org/10.1016/0749-596X\(91\)90012-9](https://doi.org/10.1016/0749-596X(91)90012-9)
- Dijkstra, T., & van Heuven, W. J. B. (2002). The architecture of the bilingual word recognition system: From identification to decision. *Bilingualism: Language and Cognition*, 5(3), 175–197. <https://doi.org/10.1017/S1366728902003012>
- Drummond, A. (2013). *Ibex Farm*. Retrieved from <http://spellout.net/ibexfarm/>
- Ellis, N. (2006). Selective attention and transfer phenomena in L2 acquisition: Contingency, cue competition, salience, interference, overshadowing, blocking, and perceptual learning. *Applied Linguistics*, 27(2), 164–194. <https://doi.org/10.1093/applin/aml015>
- Ellis, N., Hafeez, K., Martin, K. I., Chen, L., Boland, J., & Sagarra, N. (2014). An eye-tracking study of learned attention in second language acquisition. *Applied Psycholinguistics*, 35(3), 547–579. <https://doi.org/10.1080/17470218.2014.961934>
- Gao, M. Y., & Malt, B. C. (2009). Mental representation and cognitive consequences of Chinese individual classifiers. *Language and Cognitive Processes*, 24(7–8), 1124–1179. <https://doi.org/10.1080/01690960802018323>
- Grüter, T., Lew-Williams, C., & Fernald, A. (2012). Grammatical gender in L2: A production or a real-time processing problem? *Second Language Research*, 28(2), 191–215. <https://doi.org/10.1177/02676583124379>
- Hamano, S. (1997). On Japanese quantifier floating. In A. Kamio (Ed.), *Directions in functional linguistics* (pp. 173–97). Amsterdam: John Benjamins.
- Hartsuiker, R. J., Pickering, M. J., & Veltkamp, E. (2004). Is syntax separate or shared between languages? Cross-linguistic syntactic priming in Spanish-English bilinguals. *Psychological Science*, 15(6), 409–414. <https://doi.org/10.1111/j.0956-7976.2004.0069>
- Herbay, A. C., Gonnerman, L. M., & Baum, S. R. (2018). How do French-English bilinguals pull verb particle constructions off? Factors influencing second language processing of unfamiliar structures at the syntax-semantics interface. *Frontiers in Psychology*, 9, 1885. <https://doi.org/10.3389/fpsyg.2018.01885>
- Hopp, H. (2017). The processing of English which-questions in adult L2 learners: Effects of L1 transfer and proficiency. *Zeitschrift für Sprachwissenschaft*, 36(1), 107–134. <https://doi.org/10.1515/zfs-2017-0006>
- Huang, C. T. J., Li, Y. H. A., & Li, Y. (2009). *The syntax of Chinese*. Cambridge: Cambridge University Press.
- Ionin, T., Choi, S. H., & Liu, Q. (2021). Knowledge of indefinite articles in L2-English: Online vs. offline performance. *Second Language Research*, 37(1), 121–160. <https://doi.org/10.1177/026765831985746>
- Jackson, C. N., & Dussias, P. E. (2009). Cross-linguistic differences and their impact on L2 sentence processing. *Bilingualism: Language and Cognition*, 12(1), 65–82. <https://doi.org/10.1017/S1366728908003908>

- Jarvis, S. (2010). Comparison-based and detection-based approaches to transfer research. In Roberts, L., Howard, M., Laoire, MÓ, & D. Singleton (Eds.), *EUROSLA Yearbook 10* (pp. 169–192). Amsterdam: John Benjamins.
- Jeong, J. (2017). *A Study on the L2 Korean acquisition of tense, aspect, and modality pre-final endings by L1 Chinese learners*. [Doctoral dissertation, Ewha Womans University].
- Jiang, N., Novokshanova, E., Masuda, K., & Wang, X. (2011). Morphological congruency and the acquisition of L2 morphemes. *Language Learning*, 61(3), 940–967.
- Just, M. A., Carpenter, P. A., & Woolley, J. D. (1982). Paradigms and processes in reading comprehension. *Journal of Experimental Psychology: General*, 111(2), 228–238. <https://doi.org/10.1037/0096-3445.111.2.228>
- Kaiser, E., & Trueswell, J. C. (2004). The role of discourse context in the processing of a flexible word-order language. *Cognition*, 94(2), 113–147. <https://doi.org/10.1016/j.cognition.2004.01.002>
- Kang, B. M. (2002). Categories and meanings of Korean floating quantifiers—with some reference to Japanese. *Journal of East Asian Linguistics*, 11(4), 375–398. <https://doi.org/10.1023/A:1019967311110>
- Kim, H. (2018). Second language processing of Korean floating numeral quantifiers. *Journal of Psycholinguistic Research*, 47(5), 1101–1119. <https://doi.org/10.1007/s10936-018-9581-8>
- Kim, J.-B., & Sells, P. (2007). Korean honorification: A kind of expressive meaning. *Journal of East Asian Linguistics*, 16, 303–336. <https://doi.org/10.1007/s10831-007-9014-4>
- Kim, K. (2014). *Unveiling linguistic competence by facilitating performance*. [Doctoral dissertation, University of Hawaii at Manoa].
- Kim, M. H. (2019). *Processing of canonical and scrambled word orders in native and non-native Korean*. [Doctoral dissertation, University of Illinois at Urbana-Champaign].
- Kim-Renaud, Y.-K. (2001). Change in Korean honorifics reflecting social change. In T. E. McAuley (Ed.), *Language change in East Asia* (pp. 27–46). London: Curzon.
- Ko, H. (2007). Asymmetries in scrambling and cyclic linearization. *Linguistic Inquiry*, 38(1), 49–83. <https://doi.org/10.1162/ling.2007.38.1.49>
- Kobuchi-Philip, M. (2007). Floating numerals and floating quantifiers. *Lingua*, 117(5), 814–831. <https://doi.org/10.1016/j.lingua.2006.03.008>
- Kwon, N., & Sturt, P. (2016). Attraction effects in honorific agreement in Korean. *Frontiers in Psychology*, 7, 1302–1316. <https://doi.org/10.3389/fpsyg.2016.01302>
- Lago, S., Mosca, M., & Stutter Garcia, A. (2021). The role of crosslinguistic influence in multilingual processing: Lexicon versus syntax. *Language Learning*, 71(S1), 163–192. <https://doi.org/10.1111/lang.12412>
- Lardiere, D. (2008). Feature-assembly in second language acquisition. In J. Liceras, H. Zobl and H. Goodluck (eds.), *The role of features in second language acquisition*. Mahwah, NJ: Lawrence Erlbaum Associates, 106–40.
- Lee, C. (2000). Numeral classifiers, (in-)definites and incremental theme in Korean. In Lee, C., & J. Whitman (Eds.), *Korean Syntax and Semantics: LSA Institute Workshop*. Santa Cruz, '91. Thaeaksa: Seoul.
- Luk, Z. P. S., & Shirai, Y. (2009). Is the acquisition order of grammatical morphemes impervious to L1 knowledge? Evidence from the acquisition of plural-s, articles, and possessive's. *Language Learning*, 59(4), 721–754. <https://doi.org/10.1111/j.1467-9922.2009.00524.x>
- MacWhinney, B. (1987). The competition model. In B. MacWhinney (Ed.), *Mechanisms of language acquisition* (pp. 249–308). Mahwah, NJ: Lawrence Erlbaum.
- MacWhinney, B. (2008). A unified model. In P. Robinson & N. C., Ellis (Eds.), *Handbook of Cognitive Linguistics and Second Language Acquisition* (pp. 341–371). New York: Routledge.
- MacWhinney, B. (2013). The logic of the unified model. In Gass SM, & A. Mackey (Eds.), *The Routledge Handbook of Second Language Acquisition* (pp. 211–227). New York: Routledge.
- Martin, S. (1992). *A Reference Grammar of Korean*, Charles E. Tuttle Company, Rutland, Vermont, and Tokyo, Japan.
- Matuschek, H., Kliegl, R., Vasishth, S., Baayen, H., & Bates, D. (2017). Balancing Type I error and power in linear mixed models. *Journal of Memory and Language*, 94, 305–315. <https://doi.org/10.1016/j.jml.2017.01.001>
- Miyagawa, S. (1989). *Structure and case marking in Japanese*. San Diego: Academic Press.
- Morett, L. M., & MacWhinney, B. (2013). Syntactic transfer in English-speaking Spanish learners. *Bilingualism: Language and Cognition*, 16(1), 132–151. <https://doi.org/10.1017/S1366728912000107>
- Nam, Y. (2005). *4 cwukan uy kwuke yehayng [A 4-week journey through Korean]*. Seoul: Sengantang.
- Omaki, A., & Schulz, B. (2011). Filler-gap dependencies and island constraints in second-language sentence processing. *Studies in Second Language Acquisition*, 33(4), 563–588. <https://doi.org/10.1017/S0272263111000313>
- Park, M., & Sohn, K. (1993). Floating quantifiers, scrambling, and the ECP. In S. Choi (Ed.), *Japanese/Korean linguistics 3* (pp. 187–203). Stanford, CA: CSLI Publications.
- Park, S. H., & Kim, H. (2023). Bilingual processing of case-marking and locative verbs in Korean: Use of L2-unique information. *International Journal of Bilingualism*, 27(5), 586–602. <https://doi.org/10.1177/13670069221110385>
- Ratcliff, R. (1993). Methods for dealing with reaction time outliers. *Psychological Bulletin*, 114(3), 510–532. <https://doi.org/10.1037/0033-2909.114.3.510>
- R Core Team. (2020). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing, Vienna, Austria. URL <https://www.R-project.org/>.
- Sohn, H. (2001). *The Korean language*. Cambridge, MA: Cambridge University Press.
- Song, M., O'Grady, W., Cho, S., & Lee, M. (1997). The learning and teaching of Korean in community schools. In Y. H. Kim (Ed.), *Korean language in America 2* (pp. 111–127). Los Angeles, CA: American Association of Teachers of Korean.
- Spivey, M. J., & Marian, V. (1999). Cross talk between native and second languages: Partial activation of an irrelevant lexicon. *Psychological Science*, 10(3), 281–284. <https://doi.org/10.1111/1467-9280.00151>
- Sportiche, D. (1988). A theory of floating quantifiers and its corollaries for constituent structure. *Linguistic Inquiry*, 19(3), 425–449.
- Suzuki, T., & Yoshinaga, N. (2013). Children's knowledge of hierarchical phrase structure: quantifier floating in Japanese. *Journal of Child Language*, 40(3), 628–655. <https://doi.org/10.1017/S0305000912000190>
- Tokowicz, N., & MacWhinney, B. (2005). Implicit and explicit measures of sensitivity to violations in second language grammar: An event-related potential investigation. *Studies in Second Language Acquisition*, 27(2), 173–204. <https://doi.org/10.1017/S0272263105050102>
- Tolentino, L. C., & Tokowicz, N. (2011). Across languages, space, and time: A review of the role of cross-language similarity in L2 (morpho) syntactic processing as revealed by fMRI and ERP methods. *Studies in Second Language Acquisition*, 33(1), 91–125. <https://doi.org/10.1017/S0272263110000549>
- Trenkic, D., Mirkovic, J., & Altmann, G. T. (2014). Real-time grammar processing by native and non-native speakers: Constructions unique to the second language. *Bilingualism: Language and Cognition*, 17(2), 237–257. <https://doi.org/10.1017/S1366728913000321>
- Trueswell, J. C., Tanenhaus, M. K., & Garnsey, S. M. (1994). Semantic influences on parsing: Use of thematic role information in syntactic ambiguity resolution. *Journal of Memory and Language*, 33(3), 285–318. <https://doi.org/10.1006/jmla.1994.101>
- Verissimo, J. (2021). Analysis of rating scales: A pervasive problem in bilingualism research and a solution with Bayesian ordinal models. *Bilingualism: Language and Cognition*, 24(5), 842–848. <https://doi.org/10.1017/S1366728921000316>
- Wen, Z., Miyao, M., Takeda, A., Chu, W., & Schwartz, B. D. (2010). Proficiency effects and distance effects in non-native processing of English number agreement. In K. Franchini, K. Iserman & L. Keil (Eds.), *Proceedings of the 34th annual Boston University Conference on Language Development* (445–456). Somerville, MA: Cascadia.
- Witzel, J., Witzel, N., & Nicol, J. (2012). Deeper than shallow: Evidence for structure-based parsing biases in second-language sentence processing. *Applied Psycholinguistics*, 33(2), 419–456. <https://doi.org/10.1017/S0142716411000427>