

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Acquisition of words and case markers in a novel language: Artificial language learning by Korean and English speakers

Permalink

<https://escholarship.org/uc/item/5hf3h2b7>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Nahm, Susie

Bundgaard-Nielsen, Rikke L.

Wang, Yizhou

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Acquisition of words and case markers in a novel language: Artificial language learning by Korean and English speakers

Susie Nahm¹, Rikke Bundgaard-Nielsen¹, Yizhou Wang¹

¹School of Languages and Linguistics, University of Melbourne

Abstract

The acquisition of second language morphosyntactic features, such as case systems, is particularly challenging for adult learners, and the extent to which these difficulties are shaped by learners' first language remains a core question in theories of language acquisition. The present study explores the acquisition of lexical semantics and morphosyntax of a case-marking artificial language, *Katopu*, in an artificial language learning task with Korean-speaking and English-speaking adults. The results showed that both groups were able to acquire *Katopu* through the exposure to input and feedback, regardless of L1–L2 (dis)similarities and in the absence of explicit instruction, exhibiting comparable learning trajectories in accuracy and response time. Response time data further suggested that the two groups shared a similar processing pattern that semantic and syntactic knowledge were processed using separate mechanisms. The study highlights that sufficient input and feedback can effectively support language acquisition regardless of learners' language background.

Keywords: Artificial Language Learning Experiment; Morphosyntactic Acquisition; Language Acquisition; L1 Transfer; Language Background

Introduction

The world's languages differ in how they encode grammatical functions. Some languages largely depend on morphology to encode case information. For example, in Korean, case relations are primarily expressed through case suffixes, and switching the position of the subject and object noun phrases does not affect the semantic relationship, provided that the case markers correctly indicate the roles. Other languages, such as English, make limited use of case marking morphology; instead, the thematic roles, such as agents and patients, are determined by the language's fixed word order.

Sentence processing studies comparing speakers of case-marking ([+case]) and non-case-marking [–case]) languages show that L1 case systems shape how structural information is processed (Frenck-Mestre et al., 2022; Mitsugi & Macwhinney, 2016). Speakers of [+case] languages automatically exploit case cues for incremental processing, whereas speakers of [–case] languages exhibit limited reliance on case morphology and instead primarily use canonical word order to interpret sentence meaning.

Existing research demonstrates that difficulties in acquiring and processing a second language (L2) case system are often associated with the differences between the first

language (L1) and L2, particularly, the lack of overt case morphology in learners' L1 (Ahn, 2015; DeKeyser, 2005; Ellis, 2022). Several linguistic models have attempted to account for the difficulties in acquiring L2 morphosyntactic features (e.g., gender, case, person).

The Failed Functional Features Hypothesis (FFFH: Hawkins & Chan, 1997) and the Full Transfer/Full Access (FT/FA: Schwartz & Sprouse, 1996) are Universal Grammar (UG) based models that make opposing predictions about the acquisition of grammatical features absent in learner's L1. The FFFH proposes that UG is only partially available to post-critical period L2 learners, such that access is constrained by the extent to which L1 and L2 parameters are mutually interpretable. When comparable or identical structures are not present in the L1, learners are predicted to fail to construct native-like representations of L2 morphosyntactic structures (Hawkins & Franceschina, 2004).

In contrast, the FT/FA model claims that learners' entire L1 grammatical representation is transferred to the initial state of L2 acquisition and further assumes that the UG remains fully accessible to all learners regardless of the age of acquisition. Given full access to the UG and sufficient L2 input, learners can restructure the L1-based parameters to ultimately acquire the target-like L2 structure, even in the absence of structural or conceptual L1–L2 overlap.

While the FFFH and FT/FA models maintain contrasting positions regarding whether the acquisition of L2 features depends on the L1–L2 structural overlap, Sabourin et al. (2006) propose a more nuanced account that, despite the different surface configurations, facilitative L1 transfer effects may occur if there is conceptual overlap. Here, a distinction between Surface Transfer (ST) and Deep Transfer (DT) is made. ST occurs if the target structure is present in both L1 and L2 and shares a similar syntactic configuration. DT is driven by the existence of shared abstract grammatical categories, which are expressed differently in morphology, and its facilitative effects may be weaker than those of ST (Sabourin & Haverkort, 2003; Sabourin & Stowe, 2008).

Recently, a series of artificial language learning studies (Monaghan et al., 2021; Rebuschat et al., 2021; Walker et al., 2020) have provided further evidence that the absence of L1 case marking affects subsequent acquisition: English-speaking adults could learn novel words and word order but failed to acquire case markers through exposure to target language input, as predicted by the FFFH. These studies

tested whether adults could use statistical knowledge of the co-occurrence between the form and meaning in the input to acquire both vocabulary and grammar without instruction or feedback. Here, participants with no prior knowledge of the target language were exposed to a [+case] artificial language with SOV or OSV word order and case suffixes marking object and subject, like Korean.

In Monaghan et al. (2021), Rebuschat et al. (2021), and Walker et al. (2020), participants viewed one or two scenes depicting transitive events while hearing a sentence that matched only one scene. The results showed that while exposure to naturalistic input alone supported learning of lexical items and word order, acquisition of case markers was only possible when learners were given additional support through immediate *feedback* (Monaghan et al., 2021). These results show that input alone may be insufficient for acquiring L2 morphosyntactic features that lack structural and conceptual overlap with the L1 features, while highlighting the crucial role of feedback in facilitating the acquisition of such features, a factor often neglected in UG-based theories of language acquisition.

However, it remains unclear how learning differs for learners whose L1 shares *conceptual overlap* with the L2 features. If DT is robust, speakers of [+case] languages should acquire a novel [+case] language to a greater extent than [-case] language speakers, regardless of surface dissimilarity. The present study examines the effects of L1–L2 conceptual overlap on the naturalistic acquisition of lexical and morphosyntactic features using an artificial language learning paradigm with two language groups (L1 Korean; L1 English), which differ in case-marking strategies.

Moreover, we aim to identify the cognitive mechanisms that support *efficient* language processing across all learners, such as those proposed by the Dual-System Model Hypothesis (DSMH; Ullman, 2001, 2020; Ullman et al., 1997), which posits two distinct operations: procedural system supporting rule-governed grammar and regular morphology, and declarative system for memorised lexical items and irregularly inflected forms. As memory retrieval is assumed to be more effortful and conscious than grammatical computation, semantic information is predicted to require greater processing costs. If all learners are guided by universally efficient cognitive operations despite L1-specific strategies, the two groups should exhibit similar processing patterns. The present study tests hypotheses presented below:

- H1. The FFFH predicts that English speakers, but not Korean speakers, may struggle to acquire novel case markers.
- H2. In contrast to the FFFH, the FT/FA predicts that Korean and English speakers can acquire *Katopu* vocabularies and morphosyntax through input (and feedback).
- H3. The DT model predicts that Korean speakers may acquire the *Katopu* case system more quickly and to a greater extent than English speakers.
- H4. Based on the DSMH, we expect semantic features to take longer to process in both Korean and English groups.

Methods

Participants

Twenty-four participants were recruited: twelve native speakers of Australian English (Age $M = 23.4$, $SD = 5.5$; 6 females, 6 males) and twelve native speakers of Korean (Age $M = 29.6$, $SD = 9.3$; 8 females, 4 males). These language groups were selected based on their typological difference in case marking (Korean uses case suffixes allowing for flexible word order; English relies on fixed word order). The English participant group was largely monolingual, and no one had any proficiency in case-marking languages. All Korean participants had high L2 English proficiency.

Materials

A case-marking artificial language, *Katopu*, was created for the present study. Using an artificial language allows control over all relevant linguistic variables, including language complexity, morphosyntactic properties, word order configuration, and orthography. It enables control over phonological, phonotactic and prosodic similarity, thereby minimising extraneous linguistic effects on learning, and permits the introduction of novel vocabulary and syntax that are unfamiliar to all participants. In sum, this design enables investigation of the earliest stages of language acquisition as it ensures that all participants are naïve to the language.

Phonology and Prosody

The *Katopu* phonological inventory consists of vowels and consonants shared by Korean and English: consonants /m, n, p, t, k, b, d, g/ and vowels /a, i, o, u/. *Katopu* morphemes conform to the phonotactics and syllable distributional patterns of the two languages—all syllables have a /CV/ structure—limiting segmental cross-language difficulties and differences. Content words were disyllabic. Case markers were monosyllabic, reflecting that function words universally tend to have fewer syllables than content words (Cutler, 1996).

Given that language-specific intonation and stress patterns can be used to discern semantic and syntactic information (Tyler & Cutler, 2009), and that Korean and English differ in these systems, it was important to control the intonation and stress patterns implemented in the auditory *Katopu* input. For this reason, the audio input presented alongside the image material was generated using Google's *Wavenet* Text-to-Speech (TTS) in the following way: All *Katopu* /CV/ syllables were separately generated and then concatenated, using an artificial female voice. This ensured that all words were durationally similar (210 ms per syllable; no pause between syllables), each syllable and word had similar pitch contours and intensities, and importantly, that each word was stress-less. It also ensured that all sentences were devoid of phrasal prosody that might have influenced participants' perception of word boundaries. There was no prosodic break between any syllables in the audio.

Vocabulary and Syntax

The *Katopu* vocabulary comprises four disyllabic nouns (*bapi* ‘boy’, *kubu* ‘bear’, *nino* ‘girl’, *topa* ‘rabbit’); two disyllabic verbs (*gado* ‘kick’, *numa* ‘hit’); and two monosyllabic case markers (*bi-* ‘OBJ’, *mo-* ‘SUBJ’). None of the pseudowords are shared with Korean or English. The verbs were chosen as these words clearly depict the agent-patient relationship.

Katopu syntax differs from both English and Korean. It has fixed verb-subject-object (VSO) word order and prenominal case markers. A hyphen was inserted between the marker and the noun in the written sentences as an additional cue, but no pause was present in the audio stimuli. Verbs maintaining their position is characteristic of both Korean and English canonical transitive sentences (Lee, 2007). Accordingly, *Katopu* verbs consistently appear at the first position in grammatical and ungrammatical items.

Stimuli

Using the *Katopu* vocabulary, a total of 72 sentences were generated: 24 grammatical sentences (see sentence below):

Numa mo-kubu bi-bapi
Hit SUBJ-bear OBJ-boy
‘The bear hits the boy’

Forty-eight *ungrammatical* sentence items of two types were also generated: 24 with semantic errors (12 subject and 12 object substitution) and 24 with syntactic errors (12 double-marking and 12 flipped-marking). Semantic errors indicate limitations in lexical acquisition. Successful detection of double-marking errors (e.g., **numa mo-kubu mo-bapi*) suggests sensitivity to the syntactic distribution of case markers, whereas detection in flipped errors indicates acquisition of semantic roles and correct use of case markers.

Procedure

The present experiment employed an artificial language learning task in which participants were exposed to auditory and written sentences in *Katopu* co-presented with visual scenes that either matched or mismatched the information encoded in the sentences. Participants indicated via button press (“F” or “J”) whether the linguistic material matched the scene and conformed to the *Katopu* grammar. This task is similar to three recent studies (Monaghan et al., 2021; Rebuschat et al., 2021; Walker et al., 2020), which assume that learners exploit statistical co-occurrences between linguistic and non-linguistic input to acquire novel features; however, the procedure was modified to test the specific hypotheses of the present study.

The present experiment was delivered via the *Psytoolkit* server (Stoet, 2010, 2017), and participants could complete the task on their own laptop using their own headphones. Participants were informed that they would hear sentences in artificial language and may figure out its words and rules through the experiment. They were instructed to respond as accurately and quickly as possible. The experiment consisted

of three parts: a passive *learning block* with 24 trials, four *training blocks* (Block 1–4; 48 trials each), and finally a *generalisation block* (Block 5; 48 trials), introducing a novel character (*nino*, ‘girl’) in *Katopu*—resulting in 24 passive learning trials and 240 active judgement trials (total duration approximately 30 minutes). The audio and images presented for each trial in the experiment were identical for the English and Korean participants, but the orthography of *Katopu* was presented in the participants’ respective L1s (see Figure 1).

In both the learning and training blocks, participants were presented with randomised sentences constructed from three nouns, two verbs, and two case markers, resulting in 12 unique grammatical sentences. In the 24-trial learning block, the 12 sentences appeared twice: once paired with the correct image facing right and once with the same image flipped to face left, while preserving the agent-patient roles, thereby controlling for potential saliency or positional cues associated with fixed thematic role placement. Each image and written sentence was displayed for 10 seconds, with the auditory stimuli presented synchronously.

The four training blocks differed from the learning block in that participants judged whether each sentence matched the displayed image and received explicit correctness feedback (e.g., “correct answer” or “incorrect answer”). However, no information about the correct form was provided following incorrect responses. To reduce task complexity, verbs were blocked such that only one verb appeared in each training block. See figure 1 below for the example training block trial.

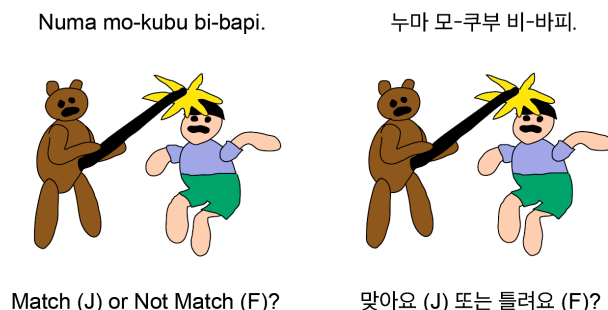


Figure 1: Example of a training block trial.
English (left) and Korean (right).

The 48-trial generalisation block differed from the training by introducing a novel character—*Nino* ‘girl’ with explicit image-referent mapping. Consequently, this block included four nouns, two verbs, and the two case markers. The distribution of grammatical, ungrammatical trials and error types were identical to the training blocks.

An inter-trial interval was imposed in block 1–5 to allow participants time to verify their hypothesis: 3-s delay followed correct responses, and 5-s delay followed incorrect responses. The longer delay was intended to promote task engagement and discourage non-deliberate responding. To minimise fatigue-related confounds, short breaks provided inserted between blocks, with participants able to determine the duration of the pause (0–10 minutes).

Results and discussion

General Learning Trajectories

To assess learning trajectories, accuracy and RT were analysed using General Linear Mixed Model (GLMM) and Linear Mixed Model (LMM) with the lmerTest package version 3.1-3 (Kuznetsova et al., 2017) in R version 4.4.2 (R Core Team, 2024). Accuracy was modelled using a binomial link directly on trial-by-trial data with the following formula: $Status \sim Language * Block + Language * Generalisation + (1 + Block | Participant) + (1 | Sentence Item)$. Language was dummy coded with English as the reference. This allowed testing the main hypothesis of whether the two language groups have different learning patterns across five blocks, as FFFH but not FT/FA would predict (full model in Tables 1, 2; overall learning curves of both groups in Figure 2).

The GLMM showed no initial accuracy difference between the two groups ($\beta = -0.099$, $p = .862$), no significant interactions with block ($\beta = 0.021$, $p = .855$) or generalisation

($\beta = -0.255$, $p = .571$), indicating similar accuracy across all blocks. Both groups performed above chance (50%) in block 1 ($\beta = 1.056$, $p = .011$; equalling 74%), suggesting some learning during exposure or early trials with feedback.

RTs for correct grammatical trials were log-transformed to address skewness and modelled using an LMM: $Log-RT \sim Language * Block + Language * Generalisation + (1 + Block | Participant) + (1 | Sentence Item)$. The LMM revealed differing RT learning curves between groups (Table 1), though both showed decreasing RTs over blocks, indicating increased response automaticity. Block was significant ($\beta = -0.053$, $p = .015$), and the interaction between block and Korean group was significant ($\beta = -0.103$, $p = .001$), suggesting that Korean participants' RTs declined more steeply, likely due to initially longer RTs (attributable to properties of Hangul (He et al., 2018) rather than differences in grammatical learning). Post hoc tests on estimated marginal means confirmed significant group differences in blocks 1 ($M = 3864$ and 5127 ms, $p = .003$), block 2 ($M = 3400$ and 3987 ms, $p = .015$), but not in blocks 3–5 ($p > .05$).

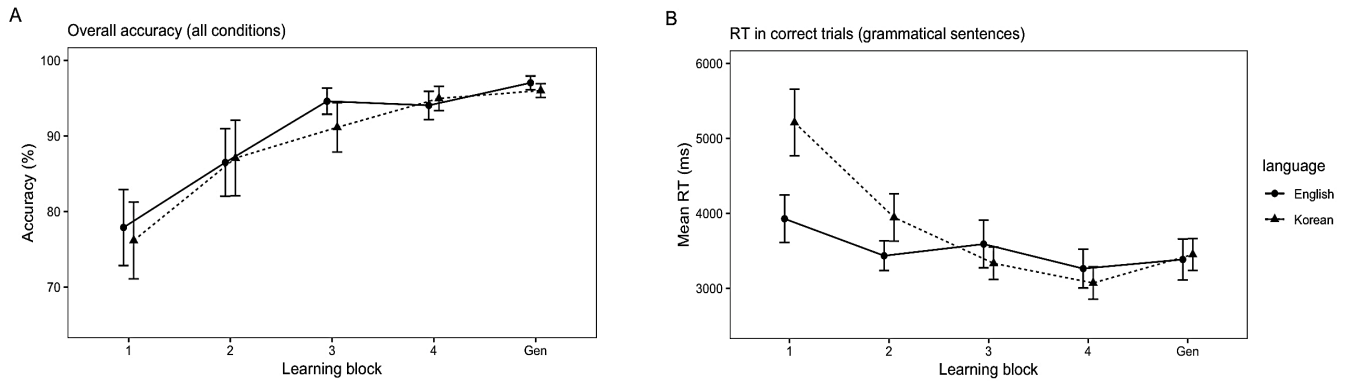


Figure 2: Overall accuracy (left) and mean RT (right).

Table 1: Summary of GLMM demonstrating accuracy.

Fixed effects	Estimate (β)	SE	z	p
Intercepts (Ref=English)	1.056	0.418	2.527	.012*
Language=Korean	-0.010	0.577	-0.173	.863
Block	0.608	0.084	7.218	<.001*
Generalization (Gen)	-0.336	0.364	-0.924	.356
Language=Korean*Block	0.021	0.115	0.182	.855
Language=Korean*Gen	-0.255	0.452	-0.565	.572

Table 2: Summary of LMM demonstrating RT in correct trials.

Fixed effects	Estimate (β)	SE	df	t	p
Intercepts (Eng Block 1)	8.184	0.086	24.0	95.467	<.001*
Language=Kor	0.344	0.121	23.5	2.854	<.001*
Block	-0.053	0.021	31.5	-2.586	.015*
Generalization (Gen)	0.106	0.041	759.4	2.568	.010*
Language=Korean*Block	-0.103	0.029	30.7	-3.539	.001*
Language=Korean*Gen	0.211	0.055	2590.4	3.820	<.001*

* Indicates statistical significance ($p < .05$)

Semantic and Syntactic Learning

We examined the learning trajectories of the two groups (English versus Korean) based on error detection in ungrammatical trials, distinguishing semantic and syntactic anomalies (Figure 3). As with overall accuracy, error detection accuracy was analysed using a GLMM and correct-rejection RTs using an LMM with the formula: *Status/Log-RT ~ Error Type * Language * Block + Error Type * Language * Generalisation + (1 + Block | Participant) + (1 | Sentence Item)*. These models tested whether L1 influenced learning patterns for lexical semantics versus syntactic/case-marking features.

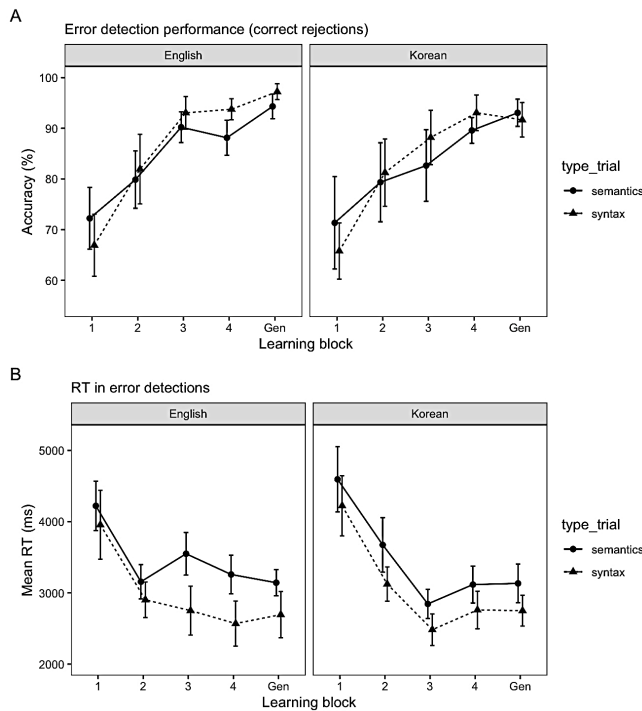


Figure 3: Error detection accuracy and mean RT in correct rejections by L1 English and L1 Korean speakers across learning blocks. Error bars indicate SE of the mean.

The GLMM revealed no overall group difference in error detection accuracy ($\beta = -0.057$, $p = .929$) and no significant interactions with block ($\beta = -0.004$, $p = .980$) or generalisation ($\beta = -0.048$, $p = .944$), consistent with overall accuracy results (Table 1) and indicating minimal L1 effects with minimal exposure and feedback. Accuracy improved across blocks for both error types for both groups ($\beta = 0.431$, $p < .001$), reflecting increasing knowledge of Katopu. Notably, the learning trajectory for syntax was steeper than for semantics ($\beta = 0.376$, $p = .030$), suggesting that syntactic and semantic processing may rely on distinct cognitive mechanisms or develop at different rates.

The LMM showed similar RTs for the two groups ($\beta = -0.212$, $p = .135$). Mean RT declined across blocks ($\beta = -0.067$, $p = .011$), with a steeper decline for syntactic errors than

semantic errors ($\beta = -0.073$, $p = .012$), potentially indicating faster acquisition of syntax (Figure 4). Across all participants, semantic error detection was slower ($M = 3471$ ms, $SD = 728$) than syntactic error detection ($M = 2966$ ms, $SD = 788$), a difference confirmed by a post hoc test ($t = 4.917$, $df = 121$, $p < .001$), further supporting the involvement of distinct mechanisms in semantic versus syntactic processing.

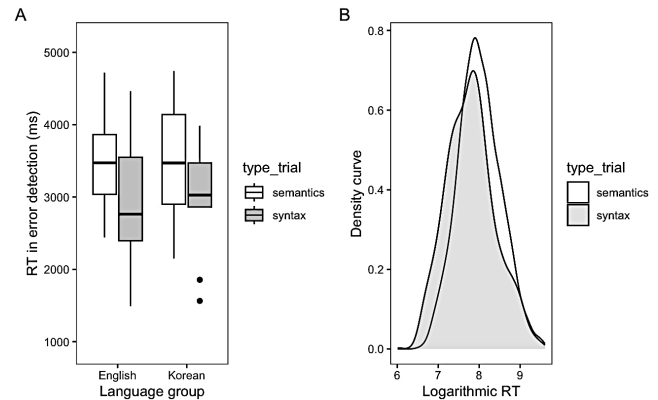


Figure 4: Box chart and density plot of RT in detecting semantic and syntactic errors in all participants.

Specific Error Types during Initial Learning

The previous analyses showed that substantial learning occurred in blocks 1–2, with the clearest increases in accuracy and decreases in RT, while performance in blocks 3–5 was more stable. Accuracy and RT in the first block likely reflect cognitive load and difficulty in detecting specific error types. To examine this, we analysed error detection for subtypes of semantic and syntactic anomalies during blocks 1–2, where response errors were more frequent. Additional GLMMs were fitted using the formula: *Status ~ Language * Type + (1 + Type | Participant) + (1 | Sentence Item)*. Here, we focused on mean performance rather than learning curves, with significance assessed via Type II Wald χ^2 tests for accuracy and F -tests for RT (Figure 5).

For semantic errors, English speakers detected object errors slightly better than subject errors ($M = 79.9\%$ vs. 72.2%), with a similar trend for Korean speakers ($M = 76.3\%$ vs. 74.6%). However, the main effect of error type was not significant ($\chi^2 = 1.551$, $p = 0.213$), nor were there effects of language group or interactions ($p > .05$). RTs showed no significant differences across error types, language groups, or interactions ($F = 0.416$, $p = 0.525$). These results indicate that word order did not affect semantic error detection, and both groups performed similarly during initial learning.

Syntactic error types showed a different pattern. Both English and Korean speakers detected double-marking errors more accurately than flipped-marking errors (English: $M = 81.2\%$ vs. 67.7% ; Korean: $M = 79.7\%$ vs. 67.4%). The effect of error type was significant ($\chi^2 = 13.164$, $p < 0.001$), while language group and interaction were not ($p > .05$). RTs mirrored this pattern: faster responses were observed for double-marking than flipped-marking errors (English: $M =$

3189 vs. 3645 ms; Korean: $M = 3350$ vs. 3920 ms), with a significant effect of error type ($F = 5.656$, $p = 0.027$) but no effect of language group or interaction.

In sum, unlike semantic errors, syntactic error types produced significantly different response patterns during initial learning, consistent with the DSMH, as participants engaged different levels of grammatical knowledge.

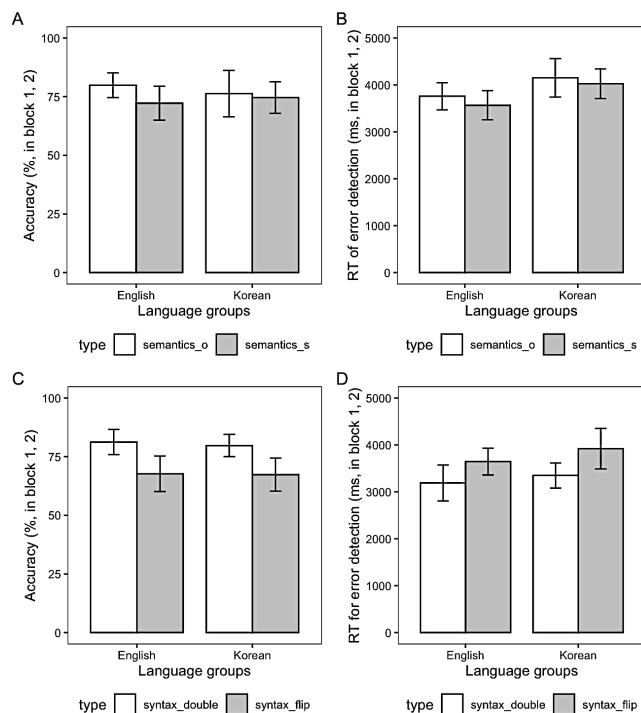


Figure 5: Accuracy and RT data in detecting sub-types of semantic and syntactic errors during the training blocks 1–2.

Conclusion

The way we process and acquire new linguistic knowledge is shaped by our prior knowledge. For example, the well-documented difficulties that L2 learners face in acquiring grammatical features are often attributed to the learner's L1, particularly, the degree of L1–L2 structural or conceptual (dis)similarity. Motivated by this link, the present study examined whether L1 case-marking strategy (Korean vs. English) affects the acquisition of a [+case] artificial language, *Katopu*. Overall, both groups acquired *Katopu* vocabulary and the case system through input and feedback, showed comparable learning trajectories (accuracy; RTs), and processed syntactic errors faster than semantic errors.

We first tested two contrasting hypotheses concerning L1–L2 structural overlap: FFFH (Hawkins & Chan, 1997) attributes non-native-like representations to the absence of a structural equivalent in the learner's L1 corresponding to the target feature, whereas the FT/FA Model (Schwartz & Sprouse, 1996) posits that, given sufficient input, learners can ultimately attain native-like representations regardless of structural disparity. Our results indicate that English speakers, despite the absence of overt case morphology in the

L1, and Korean speakers, despite differences in the configuration of case system, can map a novel case system to their meanings and grammatical functions when provided with sufficient exposure and feedback (see also Frinsel et al., 2024), consistent with Monaghan et al. (2021). Crucially, these facilitative conditions appear to mitigate L1 transfer effects, thereby supporting the FT/FA Model.

Furthermore, based on DT (Sabourin et al., 2006), which posits that not only structural, but also conceptual overlap of features could be facilitative, we predicted that Korean speakers would exhibit steeper and greater extent of learning. However, no such advantage was observed. Although this null transfer effect could suggest limited robustness of DT, it may also be attributable to *Katopu*'s fixed word order. Research has shown that speakers of [-case] languages rely on word order to assign thematic roles, which may impede the learning of case markers (Frenck-Mestre et al., 2019; Hopp, 2015; Mitsugi & Macwhinney, 2016). In this context, English speakers may have relied on word order exclusively to infer the thematic roles and the function of case markers, rendering case markers a redundant cue. Additionally, unlike prior studies using a similar paradigm which found that English speakers failed to acquire case markers (e.g., Rebuschat et al., 2021; Walker et al., 2020), our participants had access to both spoken and written forms of *Katopu*. This multimodal input may have reduced processing demands and promoted greater attentional allocation to case markers.

Finally, we explored mechanisms underlying language processing by measuring RTs in semantic and syntactic error detection. Both groups showed slower RTs in detecting semantic than syntactic violations, and for sentences with reversed subject-object marking than for sentences with a single case marker marking both subject and object roles. These patterns suggest that processing semantic information is more costly than computation of syntactic regularity and that, consistent with DSMH (Ullman, 2001, 2020; Ullman et al., 1997), learners rely on distinct systems for lexical-semantic and morphosyntactic processing, irrespective of L1.

The present study demonstrates that Korean and English speakers alike can learn a novel [+case] language through input and feedback, despite the lack of conceptual or structural overlap, and offers evidence for the differences in processing operations involved in semantics and syntax. It contributes to a small body of work examining early stages of morphosyntactic acquisition using artificial language paradigms (Andringa, 2020; Finley, 2023; Zhu et al., 2025) and extends this literature by incorporating behavioural measures, such as RTs, in the learning task.

As a single-experiment study, several limitations should be noted, including the imbalance in participants' language backgrounds—Korean speakers were bilinguals and English speakers were monolinguals—the use of fixed word order in stimuli, and finally the inclusion of orthographic input. Future studies should address these limitations to further clarify the roles of L1 background, input and feedback, and the effects of auditory versus orthographic input in language processing and acquisition.

References

- Ahn, H. (2015). Second language acquisition of Korean case by learners with different first languages [Doctoral dissertation, University of Washington]. <http://hdl.handle.net/1773/34000>
- Andringa, S. (2020). The emergence of awareness in uninstructed L2 learning: A visual world eye tracking study. *Second Language Research*, 36(3), 335–357. <https://doi.org/10.1177/0267658320915502>
- Cutler, A. (1996). Prosody and the Word Boundary Problem. In *Signal to Syntax*. Psychology Press.
- DeKeyser, R. M. (2005). What Makes Learning Second-Language Grammar Difficult? A Review of Issues. *Language Learning*, 55(S1), 1–25. <https://doi.org/10.1111/j.0023-8333.2005.00294.x>
- Ellis, N. C. (2022). Second language learning of morphology. *Journal of the European Second Language Association*, 6(1), 34–59. <https://doi.org/10.22599/jesla.85>
- Finley, S. (2023). Morphological cues as an aid to word learning: A cross-situational word learning study. *Journal of Cognitive Psychology*, 35(1), 1–21. <https://doi.org/10.1080/20445911.2022.2113087>
- Frenck-Mestre, C., Choo, H., Zappa, A., Herschensohn, J., Kim, S.-K., Ghio, A., & Koh, S. (2022). The Online Processing of Korean Case by Native Korean Speakers and Second Language Learners as Revealed by Eye Movements. *Brain Sciences*, 12(9), 1230. <https://doi.org/10.3390/brainsci12091230>
- Frenck-Mestre, C., Kim, S. K., Choo, H., Ghio, A., Herschensohn, J., & Koh, S. (2019). Look and listen! The online processing of Korean case by native and non-native speakers. *Language, Cognition and Neuroscience*, 34(3), 385–404. <https://doi.org/10.1080/23273798.2018.1549332>
- Frinsel, F. F., Trecca, F., & Christiansen, M. H. (2024). The Role of Feedback in the Statistical Learning of Language-Like Regularities. *Cognitive Science*, 48(3), e13419. <https://doi.org/10.1111/cogs.13419>
- Hawkins, R., & Chan, C. Y. (1997). The partial availability of Universal Grammar in second language acquisition: The ‘failed functional features hypothesis.’ *Second Language Research*, 13(3), 187–226. <https://doi.org/10.1191/026765897671476153>
- Hawkins, R., & Franceschina, F. (2004). Explaining the acquisition and non-acquisition of determiner-noun gender concord in French and Spanish. In P. Prévost & J. Paradis (Eds.), *Language Acquisition and Language Disorders* (Vol. 32, pp. 175–205). John Benjamins Publishing Company. <https://doi.org/10.1075/lald.32.10haw>
- He, Y., Baek, S., & Legge, G. E. (2018). Korean reading speed: Effects of print size and retinal eccentricity. *Vision Research*, 150, 8–14. <https://doi.org/10.1016/j.visres.2018.06.013>
- Henry, N., Hopp, H., & Jackson, C. N. (2017). Cue additivity and adaptivity in predictive processing. *Language, Cognition and Neuroscience*, 32(10), 1229–1249. <https://doi.org/10.1080/23273798.2017.1327080>
- Hopp, H. (2015). Semantics and morphosyntax in predictive L2 sentence processing. *International Review of Applied Linguistics in Language Teaching*, 53(3). <https://doi.org/10.1515/iral-2015-0014>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82(13). <https://doi.org/10.18637/jss.v082.i13>
- Lee, H. (2007). Case ellipsis at the grammar/pragmatics interface: A formal analysis from a typological perspective. *Journal of Pragmatics*, 39(9), 1465–1481. <https://doi.org/10.1016/j.pragma.2007.04.012>
- Mitsugi, S., & Macwhinney, B. (2016). The use of case marking for predictive processing in second language Japanese. *Bilingualism: Language and Cognition*, 19(1), 19–35. <https://doi.org/10.1017/S1366728914000881>
- Monaghan, P., Ruiz, S., & Rebuschat, P. (2021). The role of feedback and instruction on the cross-situational learning of vocabulary and morphosyntax: Mixed effects models reveal local and global effects on acquisition. *Second Language Research*, 37(2), 261–289. <https://doi.org/10.1177/0267658320927741>
- Rebuschat, P., Monaghan, P., & Schoetensack, C. (2021). Learning vocabulary and grammar from cross-situational statistics. *Cognition*, 206, 104475. <https://doi.org/10.1016/j.cognition.2020.104475>
- Sabourin, L., & Haverkort, M. (2003). 8. Neural substrates of representation and processing of a second language. In R. Van Hout, A. Hulk, F. Kuiken, & R. J. Towell (Eds.), *Language Acquisition and Language Disorders* (Vol. 30, pp. 175–195). John Benjamins Publishing Company. <https://doi.org/10.1075/lald.30.09sab>
- Sabourin, L., & Stowe, L. A. (2008). Second language processing: When are first and second languages processed similarly? *Second Language Research*, 24(3), 397–430. <https://doi.org/10.1177/0267658308090186>
- Sabourin, L., Stowe, L. A., & De Haan, G. J. (2006). Transfer effects in learning a second language grammatical gender system. *Second Language Research*, 22(1), 1–29. <https://doi.org/10.1191/0267658306sr259oa>
- Schwartz, B. D., & Sprouse, R. A. (1996). L2 cognitive states and the Full Transfer/Full Access model. *Second Language Research*, 12(1), 40–72. <https://doi.org/10.1177/026765839601200103>
- Stoet, G. (2010). PsyToolkit: A software package for programming psychological experiments using Linux. *Behavior Research Methods*, 42(4), 1096–1104. <https://doi.org/10.3758/BRM.42.4.1096>
- Stoet, G. (2017). PsyToolkit: A Novel Web-Based Method for Running Online Questionnaires and Reaction-Time Experiments. *Teaching of Psychology*, 44(1), 24–31. <https://doi.org/10.1177/0098628316677643>
- Tyler, M. D., & Cutler, A. (2009). Cross-language differences in cue use for speech segmentation. *The Journal of the Acoustical Society of America*, 126(1), 367–376. <https://doi.org/10.1121/1.3129127>

- Ullman, M. T. (2001). The neural basis of lexicon and grammar in first and second language: The declarative/procedural model. *Bilingualism: Language and Cognition*, 4(2), 105–122. <https://doi.org/10.1017/S1366728901000220>
- Ullman, M. T. (2020). The declarative/procedural model. In B. VanPatten, G. d. Keating, & S. Wulff (Eds.), *Theories in Second Language Acquisition: An Introduction* (3rd ed., pp. 159–190). Routledge.
- Ullman, M. T., Corkin, S., Coppola, M., Hickok, G., Growdon, J. H., Koroshetz, W. J., & Pinker, S. (1997). A Neural Dissociation within Language: Evidence that the Mental Dictionary Is Part of Declarative Memory, and that Grammatical Rules Are Processed by the Procedural System. *Journal of Cognitive Neuroscience*, 9(2), 266–276. <https://doi.org/10.1162/jocn.1997.9.2.266>
- Walker, N., Monaghan, P., Schoetensack, C., & Rebuschat, P. (2020). Distinctions in the Acquisition of Vocabulary and Grammar: An Individual Differences Approach. *Language Learning*, 70(S2), 221–254. <https://doi.org/10.1111/lang.12395>
- Zhu, L., Rebuschat, P., Nixon, J. S., & Monaghan, P. (2025). Learning morphology from cross-situational statistics. *Studies in Second Language Acquisition*, 1–27. <https://doi.org/10.1017/S027226312510106X>