

Morphological variation and priming in Hungarian spontaneous dialogue

Péter Rácz, Budapest University of Technology and Economics, racz.peter.marton@ttk.bme.hu

I examine morphological priming in spontaneous dialogue by tracking the alternation between the long (-jál/-jél) and the short (-j) variant of the 2sg indefinite subjunctive in Hungarian task-oriented conversations, to see whether members of dyads persist and converge in their choice of suffixed variants. I combine conversational data from the Akaka Maptask and Budapest Games corpora with distributional estimates from the Hungarian Webcorpus, resulting in a dataset of 334 subjunctive tokens. Hierarchical Bayesian generalised linear models show robust within-verb priming: following a long prime, a long target is about 3x more likely, with no comparable effect across different verbs. Speaker identity and prime–target distance did not improve model fit. Independent of priming, conversational choices track corpus baselines, indicating an additive contribution of lexical distributions. The results provide naturalistic evidence for the priming of inflected morphological variants in Hungarian, and clarify how lexical distributions scaffold variant choice in conversation.



1. Introduction

1.1 Morphological priming

When language users encounter linguistic material, this facilitates the subsequent production of identical or related material. In experimental psycholinguistics, this facilitation effect is called *priming*: exposure to a prime stimulus reduces processing costs for a related target. In linguistic interaction, the same process maps onto how Speaker 1 and Speaker 2 use a given construction.

In conversation, multiple mechanisms affect the facilitation effect. *Self-priming* is the facilitation of a speaker's own subsequent production by their own prior output. *Reciprocal priming* occurs when one speaker's output facilitates the other speaker's production of the same or a related form. *Alignment* and *convergence* refer to the broader process by which interlocutors come to use increasingly similar forms over the course of an interaction, which may arise from reciprocal priming, shared situational pressures, or both (Garrod & Pickering, 2009; Pickering & Garrod, 2004). The absence of a speaker-identity effect in a conversational dataset is consistent with reciprocal priming, but does not rule out parallel self-priming under shared lexical biases.

A similar set of distinctions applies within morphology. The *lexical boost* is the facilitation of production by recent activation of the same lemma or word form (Pickering & Branigan, 1998). This is distinct from *structural priming*, which operates on abstract syntactic configurations independently of lexical content (Bock, 1986; Gries, 2005). Within morphological priming specifically, one can distinguish *stem-level* priming (facilitation of a lemma), *affix-level* priming (facilitation of a suffix or inflectional pattern across different stems), and *whole-inflected-form* priming (facilitation of a specific inflected token). Evidence for affix-level priming comes primarily from experimental work on derivational and inflectional decomposition (Amenta & Crepaldi, 2012; Marslen-Wilson, 2007). In naturally occurring speech, it is difficult to distinguish affix-level from whole-form priming, because the same speaker typically re-uses the same verb in the same inflected form. Whole-form priming is often referred to with a theoretically more neutral label, *persistence*, in conversational data – see Szmrecsanyi (2006) below.

In psycholinguistics, lexical boost is well established, and we have clear evidence for the priming effect of specific suffixes as well (see e.g. Amenta & Crepaldi, 2012). People access and produce stems and affixes more readily after recent exposure.

Rácz and Lukács (2024) and Rácz et al. (2020) demonstrated alignment between language users not only in the use of stems or affixes, but in entire distributions of variation in morphological inflection. Through a series of online word-picking games, they showed that players could converge to the choices of their co-player and that the effects of this interaction persisted after the game in post-testing. The effects themselves were not limited to across-word or across-suffix priming, but rather entailed a shift in lexical distributions at the subword level: the co-player's word choice influenced the behaviour of similar words and the strength of this effect was a function of their similarity to it.

In models of dialogue as joint action, convergence provides a core coordination mechanism and is a lynchpin of language variation and change (Beckner et al., 2009; Garrod & Pickering, 2004, 2009; Giles & Coupland, 1991; Pickering & Garrod, 2004, 2006). Convergence has been attested in naming preferences (Roberts, 2010), syntax (Gries, 2005), and phonetics (Babel, 2012). In comparison, morphological convergence in natural language has received limited attention. Weiner and Labov (1983) found that a strong predictor of passive use in English-language interviews was the presence of another passive in the previous five utterances. Szmrecsanyi (2006) used a corpus-based approach to look at the English comparative, where a word-formation pattern (*friendlier*) alternates with an analytical construction (*more friendly*). He found that the choice of either variant persisted in discourse, meaning that if speakers started using one, they would be more likely to keep using it.

The dearth of similar research work in natural language reflects the challenges of detecting morphological convergence in interactional data. Many analytic languages, like English, provide a shallow pool of morphological variation, which is, in turn, very hard to dredge in corpora. Lab studies, like Rácz et al. (2020), used an experimental battery that exposed participants to a rapid fire of nonwords and a framework that actively encouraged predicting the co-player's behaviour. This allowed these studies to explore more nuanced convergence effects across lexical distributions. On the flip side, the design was divorced from natural language use, in which convergence usually operates on familiar words and its incentives and mechanisms are more implicit. The natural-language studies, like Weiner and Labov (1983), relied on natural language and identified patterns that were frequent enough in English to constitute a robust, compact sample.

This study uses a set of dyadic interactions recorded in Hungarian. Interactive tasks entail the frequent use of instructions, requests, and commands. Hungarian uses the subjunctive to express these imperative functions (Tóth, 2007). The Hungarian 2SG subjunctive suffix varies between a long and a short form. This structured variation can be tracked across interactions and offers a unique window into how lexical distributions and convergence effects shape word use together in a language with rich inflection and complex word-formation processes.

1.2 The Hungarian subjunctive

The Hungarian verb is marked for person, number, tense, mood, and definiteness. The subjunctive, used both in the imperative and in certain subordinate constructions (Tóth, 2007), is marked in the present tense and has a different paradigm for definite and indefinite verbs. The 2SG form varies in both paradigms, with a long and a short form. In this article, I focus on the more frequent indefinite. Example (1) shows the long and the short form.

- (1) kelt-sél/kelt-s fel, ha szeptember véget ér
 wake-SBJV-2SG-INDEF PART if September end.ACC reach-3SG-INDEF
 ‘Wake me up when September ends.’

The short form is an underlying *-j* which can assimilate to the preceding stem, while the long form adds *-jál/-jél*, depending on the backness of the final stem vowel, following the general patterns of Hungarian vowel harmony (Siptár & Törkenczy, 2000) (hence, *keltsél*/**keltjél* wake-SBJV-2SG-INDEF). The long/short variation is attested, in some form, with all types of verbs, including irregular forms (see e.g. *légy/legyél* be-SBJV-2SG-INDEF, *gyere/jöjjél* come-SBJV-2SG-INDEF). This variation has been present in Hungarian since at least the 15th century (László, 2001), but no systematic, large-scale, quantitative study exists of the stylistic and social factors that condition it.

2. Current study

I used a Hungarian Webcorpus and a dataset of dyadic interactions between Hungarian speakers to see whether speakers show priming in the use of the long/short form of the indefinite subjunctive and whether their choices reflect variation in the ambient language. For each target variant of the subjunctive, I identified the last preceding subjunctive variant as its prime.

I tested the following research questions using a series of hierarchical Bayesian models (see the Appendix for the full model comparison). Each question maps onto a specific model comparison.

- RQ1.** Does the use of a subjunctive variant (the *target*) depend on the last used variant (the *prime*)? I compared a model with the prime effect to a baseline model with only lexical frequency as a predictor.
- RQ2.** Is this priming effect restricted to repetitions of the same verb, or does it generalise across verbs? I compared a model with a prime $\times$ same-verb interaction to a model with only a main effect of prime.
- RQ3.** Does priming differ depending on whether the same or a different speaker produced the prime? I compared a model with a prime $\times$ speaker interaction to the best model without it.
- RQ4.** Is a preference for the long/short form shaped by the verb’s lexical distribution in the webcorpus? The corpus frequency predictor is present in all models; I assess its contribution to the best model.

These questions are descriptive. The pattern of results can be interpreted in different theoretical frameworks. If priming is restricted to the same verb and the same inflected form, this is consistent with whole-word facilitation (lexical boost). If priming generalises across verbs sharing the same suffix, this would point to affix-level or paradigm-level convergence. I return to this in Section 3.

2.1 Sources

I drew corpus data from a frequency list based on the second Hungarian Webcorpus (Nemeskey, 2020; Rácz, 2025), which was built from Common Crawl and has a size of 9 billion words. I used transcribed conversation data from two sources. The Akaka Maptask Corpus consists of five hours of recordings of 24 task-oriented dialogues recorded using head-mounted microphones, with pairs of participants completing a map task (Molnár et al., 2023). The dataset has recordings from 46 participants (24 women, 22 men, mean age 20). The Budapest Games Corpus consists of 9 hours of 36 dialogues of pairs of participants. The dataset has recordings from 12 participants (five women, seven men), working together in an object identification task. Participants were recruited using convenience sampling, were matched in age group between the ages of 20 and 60 and knew each other (Mády et al., 2023). For more details, see Mihajlik et al. (2024).

The distributions of the raw data can be seen in **Figure 1**. We see the number of long/short forms used by speakers across conversations on the left and the number of long/short forms per verb lemma in total on the right. Both show a Pareto-distribution. While many conversations see only sporadic use of the subjunctive, in some, participants use the subjunctive extensively. Some verbs, like ‘wait’, ‘go’, or ‘watch’, are overrepresented in subjunctive use.

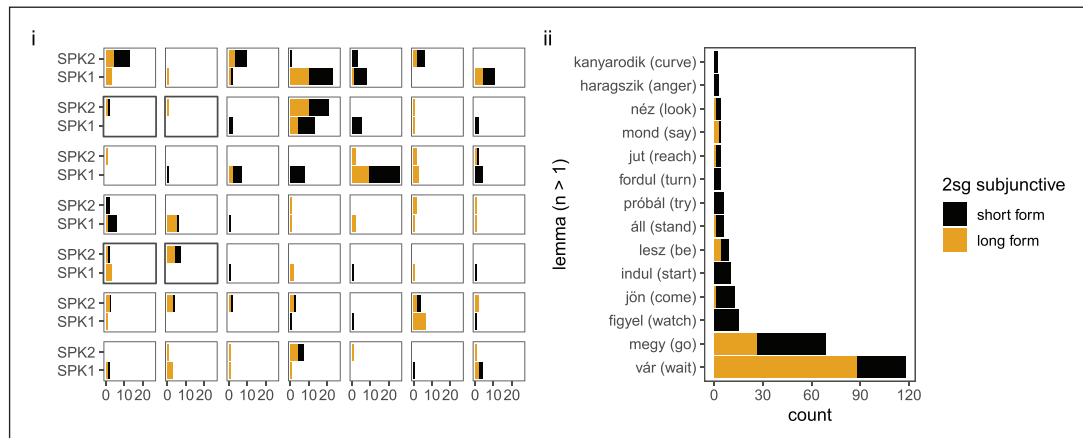


Figure 1: Counts of long/short indefinite subjunctive forms in the raw data. (i) Number of long/short forms used by Speaker 1/2 in each conversation, (ii) number of long/short forms per verb lemma across all conversations (excluding hapaxes).

2.2 Methods

For details on the transcription and annotation of the original data, see Mády et al. (2023) and Molnár et al. (2023). I used force-aligned Praat textgrids and extracted the word tier and speaker identifier from each textgrid file. I filtered the resulting dataset to only include indefinite subjunctives and checked the data to avoid overlaps between subjunctive forms.

I drew a dataset of subjunctive forms from the webcorpus and calculated $\log(\text{freq}(\text{long variant})/\text{freq}(\text{short variant}))$ for each lemma.

I joined these data with the subjunctive dataset from the conversations to add word-level information. Three observations, for three different low-frequency verbs (*tenyerel* ‘palm’ *kalauzol* ‘guide’, *tologat* ‘push habitually’) in the conversation data had no long form in the corpus. I excluded these from the final analysis.

The final dataset consisted of 334 subjunctives from 22 recordings in the Akaka Maptask Corpus and 27 recordings in the Budapest Games Corpus. This is 0.04% of all transcribed words in the recordings. The sample was hand-checked for accuracy.

2.3 Data analysis

I analysed the data in R (R Core Team, 2025), fit models in brms (Bürkner, 2021), and used the ggplot2 and sjPlot packages for visualisations (Lüdtke, 2025; Wickham, 2016).

The analysis compares each subjunctive verb form (the target) with the last subjunctive verb form (the prime) in the conversation. The distance between the prime and the target can vary. Out of the 334 observations, 3 are infrequent verbs that have no long form in the webcorpus and 52 are single mentions or first mentions in a conversation, meaning that there is no preceding variant to compare them with. This leaves 279 observations which belong to 26 verbs. The top five verbs (*megy* ‘go’, *vár* ‘wait’, *figyel* ‘watch’, *jön* ‘come’, and *indul* ‘start’) are responsible for 80% of all observations, with a Gini coefficient of 0.76 (see **Figure 1**). The Gini coefficient quantifies distributional inequality: here, it indicates that most tokens belong to few types.

I fit a series of eight hierarchical generalised linear models with a binomial error distribution and a logit link function, predicting $p(\text{target is long})$. All models included conversation, speaker, and verb lemma as grouping factors (random intercepts). Fixed effects varied across models: the baseline model included only the verb’s corpus log-odds of the long form; subsequent models added the prime (long/short), pairwise interactions of the prime with same/different verb, same/different speaker, distance, and corpus frequency, as well as a three-way interaction and a model with all two-way interactions. The full set of models and their formulae are reported in the Appendix (Table A1).

I placed Normal(0, 3) priors on intercepts, Normal(0, 2) priors on fixed-effect coefficients, and Exponential(1) priors on random-effect standard deviations. These are weakly informative priors that allow large effects on the log-odds scale, while regularising against implausible extremes, following standard recommendations for logistic regression (Gelman et al., 2008). I compared models using approximate leave-one-out cross-validation (LOO-CV; Vehtari et al., 2017) and Bayes Factors computed via bridge sampling (Gronau et al., 2020). Residual autocorrelation was inspected manually and was not detected in any model.

I report the best-fitting model below. Where relevant, I convert log-odds estimates to probabilities using the inverse logit function ($p = \exp(x)/(1 + \exp(x))$, implemented as `plogis()` in R) to aid interpretation. I use three significant figures throughout (unless further precision is informative), matching **Table 1**.

Table 1: Estimates of the best model, with 95% credible intervals.

term	estimate	95% CI	
intercept	−4.85	−7.12	−2.67
prime is long	1.68	0.70	2.66
prime is different verb	1.65	0.76	2.58
target log(long/short) in webcorpus	4.15	1.31	6.87
prime long form $\times$ different verb	−2.00	−3.29	−0.71

2.4 Results

The estimates of the best model can be seen in **Table 1**. Both priming effects and lexical effects are visible.

The best-fitting model (fit3 in the Appendix) includes the prime, whether the prime is the same or a different verb, their interaction, and the verb’s corpus log-odds. Model comparison supports the following answers to the research questions.

RQ1: Priming. The model with within-verb priming is strongly favoured over the baseline (BF = 62.26; Table A1). A long prime increases the probability of a long target.

RQ2: Same verb vs. across verbs. The interaction model is strongly favoured over a model with only a main effect of prime (BF = 106.72). Priming is restricted to the same verb lemma. When the previous verb was different, the priming effect was not credibly different from zero (95% CI: [−1.32; 0.63], calculated from the joint posterior).

RQ3: Speaker identity. Adding a prime $\times$ speaker interaction did not improve model fit (BF = 0.008 in favour of the more complex model). The data provide no evidence that priming differs depending on who produced the prime. I return to the interpretation of this null result in Section 3.

RQ4: Lexical distribution. The verb’s corpus log-odds of the long form is a strong predictor in all models (estimate: 4.15, 95% CI: [1.31; 6.87]). Additionally, prime $\times$ distance (BF = 0.002) and prime $\times$ corpus frequency (BF = 0.006) did not improve fit over the best model. The priming and lexical effects are additive. A three-way interaction of prime, same verb, and corpus frequency yielded indeterminate evidence (BF = 0.82). Full model comparisons are reported in the Appendix.

Going back to the best model, the credible intervals capture how the groups are related to the intercept (short primes). How they are related to each other is best understood looking at the raw data in **Figure 2** and the model predictions in **Figure 3**.

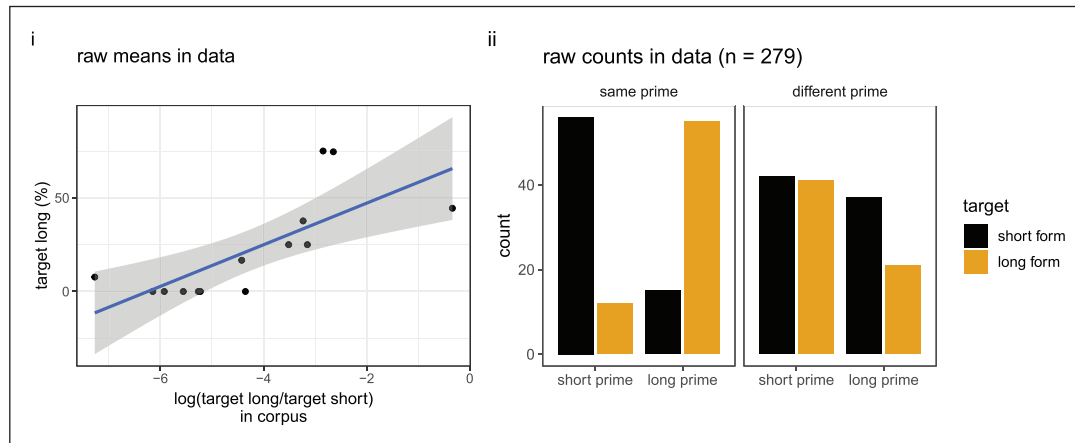


Figure 2: Verb long/short preference in the raw data. (i) Targets are more likely to be long in conversations (y axis) if the verb is more likely to be long in the webcorpus (x axis – the widening of the credible interval reflects gaps in the distribution of verb types along the corpus frequency axis). (ii) Long primes increase the probability of long targets; in the absence of a long prime, the short form predominates. This effect is absent if it was a different lemma.

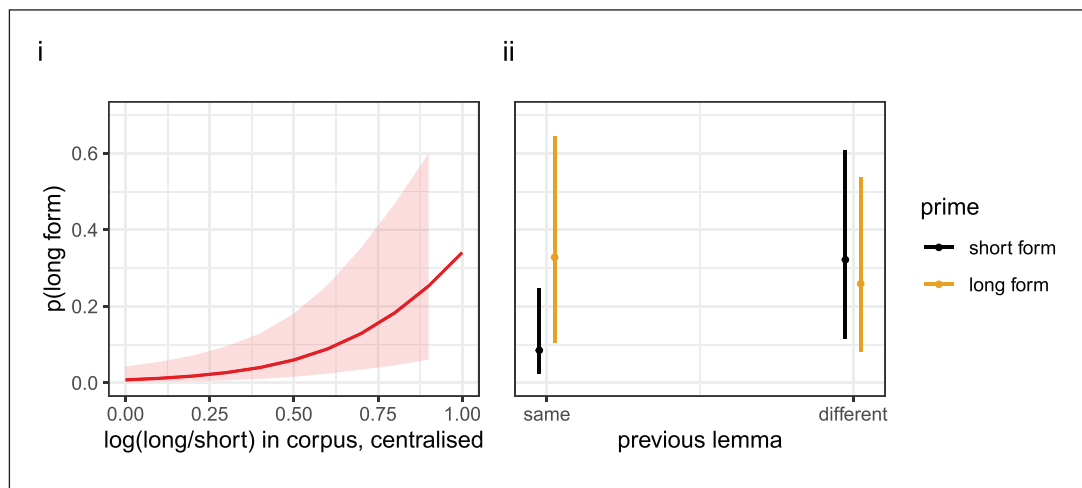


Figure 3: Verb long/short preference based on best model. (i) Targets are more likely to be long in conversations (y axis) if the verb is more likely to be long in the webcorpus (x axis), (ii) Long primes increase the probability of long targets if the prime was the same verb lemma. In the absence of a long prime, short forms predominate. This effect is absent if it was a different lemma.

If a verb prefers the long form in the webcorpus (x axis), it will do so in the conversations (y axis). We see a strong priming effect within the same verb: a long prime is followed by a long target. In the absence of a long prime, the short form predominates. There is no clear priming pattern across verbs.

Both the lexical effect and the priming effect persist in the best model. If we transform log odds to probabilities using the inverse logit function, we can say that verbs at the highest end of the (scaled) corpus distribution of long forms are about 30x more likely to pick a long form over verbs at the lowest end ($plogis(-4.85 + 4.15) = .332$ over $plogis(-4.85) = .008$). Within verbs, a long prime variant means that a long target variant is about 3x more likely ($plogis(-4.85 + 1.68) = .040$ over $plogis(-4.85) = .008$). The two effects are additive, not interactive, in the model. Model predictions can be seen in **Figure 3**.

Based on model selection, the best-fitting model has two plausible alternatives. First, target verb length may reflect lexical preferences alone, with no priming effect present. There is robust evidence against this model ($BF = 62.61$).

A second alternative model includes an interaction between prime and lexical preference. This is, again, best understood by looking at the predictions of this alternative model (**Figure 4**). In this model, it remains true that the prime only has a clear effect on the target if it is the same verb. However, this effect interacts with the verb's lexical preference: if a verb is unlikely to be long based on the webcorpus, it will remain short in the conversation, irrespective of the prime. If the verb is more likely to be long in the corpus, this is more prominent in the conversation. Evidence for this interaction is indeterminate ($BF = 0.82$ relative to the simpler model), meaning the data neither support nor rule out a role for lexical expectations in modulating priming. If this interaction proved robust in larger samples, it would be consistent with expectation-based models of production. I return to this in Section 3.

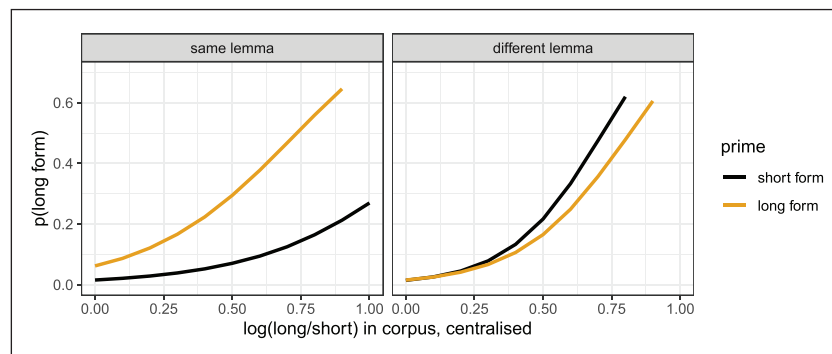


Figure 4: Predictions of alternate model: interaction between lexical preference, prime long/short, and verb lemma. If the lemma is the same (left), we see a stronger lexical effect for long primes versus short primes. If the lemma is different, this effect is absent (right).

Leave-one-out cross-validation fails to discriminate between models (identical LOO scores across all three candidates), likely reflecting data scarcity: with relatively few observations and few verb types, out-of-sample prediction is uninformative. As noted above, Bayes Factor comparison favours the priming model over the null, but remains indeterminate between the priming model and the more complex alternative which assumes a lexical preference : priming interaction. This indeterminacy may also reflect insufficient data to justify additional parameters and could be tested on a larger interactional dataset.

Since it was flagged by cross-validation, I used a permutation test to further confirm the priming effect (the reported model). I subsetting the data to observations where the target was the same verb as the prime (138 observations). I randomly shuffled the prime (long/short) 10,000 times and calculated a simulated ratio for the number of long targets over short targets. The mean of these ratios was 0.49. The true mean of long targets following long primes was much higher: 0.78. No simulated mean was higher than this. This means that the priming effect across identical primes and targets is unlikely to be due to chance alone ($p < 0.001$). This simple test does not replicate the entire interaction in the best model, but it does lend further support to the robustness of the priming effect.

3. Conclusion

I analysed spoken interactions in Hungarian-speaking dyads working together in simple game-like tasks, looking at variation in the 2SG indefinite subjunctive, which is largely used to express requests and commands in Hungarian (Tóth, 2007). I wanted to see whether the choice of the short or long variant of subjunctive forms persisted across spoken dialogue.

I found that a long prime variant made the use of a long target variant about 3x more likely for the same verb. In the absence of a long prime, the short form predominated. The effect was restricted to repetitions of the same verb lemma and did not generalise across verbs. In addition, the choice of the short or long subjunctive for a given verb strongly correlated with its general preference for short or long subjunctives in a large webcorpus of Hungarian. The priming and lexical effects were additive.

The speaker distinction did not improve model fit ($BF = 0.008$). This means that the data do not distinguish between reciprocal priming, in which one speaker's output facilitates the other's production, and parallel self-priming, in which both speakers are independently influenced by the same lexical distributions. Under the strong frequency effects observed here, both speakers are likely to gravitate toward the same form, regardless of interpersonal coordination. Distinguishing these accounts would require designs that manipulate speaker exposure independently, such as confederate paradigms or asymmetric information tasks.

The finding that priming operates exclusively within lemmata is consistent with a lexical boost account, in which facilitation is driven by recent activation of the same word form (Pickering &

Branigan, 1998). The priming observed here operates on specific inflected tokens: the same verb in the same subjunctive variant. There is no evidence of generalisation across verbs, which would be expected under affix-level or paradigm-level convergence. We therefore cannot conclude that the effect operates at the level of subword representations. It may equally reflect whole-word storage and retrieval.

At the same time, the alternation under study is morphologically conditioned: the long and short forms are distinct inflectional variants of the same exponent. The strong corpus-frequency effect shows that distributional properties of the inflectional paradigm shape conversational choices independently of local priming. This is a necessary, though not sufficient, condition for morphological-level processes. A stronger test would require evidence that priming of the long form of one verb increases the probability of the long form for a different verb with a similar distributional profile. The data do not show this, but the sample is small and heavily concentrated in a few high-frequency verbs.

We discussed an alternate, more complex model of the results, in which lexical preference and prime jointly shape target choice in conversation. Evidence for this interaction is indeterminate in the data ($BF = 0.82$). If the interaction between lexical preference and prime length proves robust in larger samples, it would match the broader consensus regarding the relationship between priming effects and lexical structure in morphological processes (Hay & Baayen, 2005) and support expectation-based models of production in which speakers adjust form selection to adapt to the statistics of the current environment (Jaeger & Snider, 2013). Such models track expectation violations and converge towards higher-probability forms. However, evidence for these mechanisms is contested (Fazekas et al., 2024), and the conclusions of the data are indeterminate on this point, potentially due to sample size limitations.

A caveat on the interpretation of additive effects in logistic regression: because the logit link function is nonlinear, effects that are additive on the log-odds scale are not additive on the probability scale (Ai & Norton, 2003). In the model, the prime effect and the corpus-frequency effect do not interact on the log-odds scale. On the probability scale, however, the same prime coefficient produces a larger shift in probability for verbs with a higher baseline probability of the long form. This means that the prime's influence on variant choice is already modulated by lexical frequency, even without an explicit interaction term. This does not amount to a surprisal-based mechanism, but it does mean that the “additive” result is less theoretically flat than it may appear: the model already captures a form of frequency-dependent priming through the geometry of the link function. This may also explain why the explicit interaction model (fit7) is indeterminate, rather than clearly disfavoured: it is trying to capture variance that the logit link already partially accounts for.

Large-scale online morphological convergence tasks reported in Rácz and Lukács (2024) and Rácz et al. (2020) find this interaction between lexical distribution and linguistic convergence.

In those tasks, a nonword's behaviour is shaped by the interaction of the nonword's baseline distribution and the behaviour of the co-player. At the same time, these online tasks expose a larger number of participants to a much wider, well-balanced array of nonwords, making it much easier to identify relevant complex patterns in the data. It is entirely possible that a larger conversational dataset would also support the otherwise intuitive result that priming effects and lexical effects interact.

The main strength of the analysis is that it demonstrates priming of morphological variants for existing words within a realistic language interaction. Work on morphological variation and priming in naturalistic settings is scarce in the literature. This study provides important confirmation for existing corpus-based (Szmrecsanyi, 2006; Weiner & Labov, 1983) and artificial language studies (Rácz & Lukács, 2024; Rácz et al., 2020) on morphological convergence, and extends these findings to a morphologically rich language.

The study contends with a number of limitations. The interactions play out between familiar participants in an informal setting. Interactive tasks also carry a lot of emphasis, with participants repeatedly telling one another to look at something or go somewhere, and emphasis might explain some repetition in the long/short form of the subjunctive in these conversations. The stylistic and social factors driving variation are not known. The relatively small set of verbs and the structure of the conversations likely translates to more limited insights into the relationship between priming effects and lexical effects on morphological convergence. This is not unusual, as state-of-the-art theories assume robust but hard-to-detect pressures acting on language variation and change (Hay et al., 2015). The results remain important as a demonstration of inflected-form priming in a naturalistic interaction, with lexical distributions independently scaffolding variant choice.

Any quantitative analysis of convergence effects based on audio transcripts carries more general limitations, and these reference two important questions in current theory.

First, speaker convergence across variable linguistic output is related to more general work on variation and systematising in language and cognition. Schumacher and Pierrehumbert (2021) make the point that linguistic behaviour that looks like probability matching in the aggregate might stem from pooling heterogeneous speakers, who modulate their choices based on socially learned baselines and associations. The aggregate analysis of the linguistic data, despite hierarchical modelling, might pave over a similar set of individual strategies that determine morphological convergence above and beyond repetition and lexical pressure. Schumacher and Pierrehumbert also make the observation that formulating a rule might be a separate, subsequent, step after learning from example, and this fits in with the result that priming and lexical effects are additive, rather than interactive, in the dataset. It is possible that participants do use rules in selecting the subjunctive for various verbs, but these rules are not directly built on corpus frequencies. This would imply that there might be an interaction between lexical preference and

priming in the dataset, but this is poorly represented by aggregate frequency counts drawn from a webcorpus.

Second, there is an emerging consensus in linguistics that the efficient communication of information moulds language variation down to the micro-level (Piantadosi et al., 2011). Analysing dyadic interactions should be a stark reminder that much of this information is non-linguistic and hard to operationalise or quantify. Joint gaze, attention, non-verbal turn-taking, and the larger social context of any interaction could have as much influence on morphological convergence as lexical distributions or word-to-word priming. This highlights the lessons that quantitative psycholinguistics can learn from linguistic anthropology and ethnographic fieldwork, which can capture the nonverbal and social dynamics that shape human interactions (Goodwin, 2018; Stivers et al., 2009).

Appendix

Table A1: Models fit on the data.

Description	Formula	Fit	elpd_diff	se_diff	bfs	bf_labels
Baseline (lexical)	~ corpus log(long/short)	fit1	-2.35	4.00	62.26	fit3/fit1
Prime main effect	~ long prime + corpus log(long/short)	fit2	-4.62	3.57	106.72	fit3/fit2
Prime × same verb	~ long prime * prime diff. lemma + corpus log(long/short)	fit3	0.00	0.00	–	–
Prime × speaker	~ long prime * prime diff. speaker + corpus log(long/short) + prime diff. lemma	fit4	-3.02	2.73	0.008	fit4/fit3
Prime × distance	~ long prime * distance to prime + corpus log(long/short) + prime diff. lemma + prime diff. speaker	fit5	-5.13	3.15	0.002	fit5/fit3
Prime × frequency	~ long prime * corpus log(long/short) + prime diff. lemma + prime diff. speaker	fit6	-4.83	3.08	0.006	fit6/fit3
Three-way interaction	~ long prime * prime diff. lemma * corpus log(long/short) + prime diff. speaker	fit7	-0.15	1.20	0.82	fit7/fit3
All two-way int.	~ long prime * prime diff. lemma + long prime * corpus log(long/short) + long prime * prime diff. speaker + long prime * distance to prime	fit8	-2.39	1.15	0.05	fit8/fit3

Table A1 summarises the eight hierarchical generalised linear models fit to the data. All models predict $p(\text{target is long})$ with a binomial error distribution and a logit link. All models include conversation, speaker, and verb lemma as grouping factors (random intercepts). Priors are weakly informative: $\text{Normal}(0, 3)$ on the intercept, $\text{Normal}(0, 2)$ on fixed-effect coefficients, and $\text{Exponential}(1)$ on random-effect standard deviations. These follow standard recommendations for logistic regression (Gelman et al., 2008) and are intended to regularise estimates without imposing strong directional constraints.

Models were compared using approximate leave-one-out cross-validation (LOO-CV; Vehtari et al., 2017) and Bayes Factors (BF), computed via bridge sampling as implemented in the `bayestestR` package (Makowski et al., 2019). The `elpd_diff` column reports the difference in expected log pointwise predictive density, relative to the best-fitting model (fit3); negative values indicate worse out-of-sample prediction. The `se_diff` column gives the standard error of this difference. LOO-CV differences are small, relative to their standard errors across all models, which reflects limited discriminability, given sample size ($n = 279$).

Bayes Factors are reported as the ratio of marginal likelihoods for the row model over its comparison model (`bf_labels` column). Fit3, the $\text{prime} \times \text{same-verb}$ model, serves as the reference. $\text{BF} > 10$ indicates strong evidence in favour of the numerator model; $\text{BF} < 0.1$ indicates strong evidence against it; values near 1 are indeterminate. Fit3 is strongly favoured over the baseline ($\text{BF} = 62.26$) and over the prime-only model ($\text{BF} = 106.72$), confirming that within-verb priming improves on both lexical effects alone and an undifferentiated prime effect. Adding speaker identity (fit4, $\text{BF} = 0.008$), prime-target distance (fit5, $\text{BF} = 0.002$), $\text{prime} \times \text{corpus frequency}$ (fit6, $\text{BF} = 0.006$), or all two-way interactions (fit8, $\text{BF} = 0.05$) does not improve on fit3 and is actively disfavoured. The three-way interaction model (fit7, $\text{BF} = 0.82$) is indeterminate: the data neither support nor rule out an additional interaction between priming, lemma identity, and corpus frequency.

Residual autocorrelation was checked for all models and was not detected.

Abbreviations

2SG	second person singular
3SG	third person singular
ACC	accusative
INDEF	indefinite
PART	particle
SBJV	subjunctive

Data accessibility statement

Data and code are available at <https://doi.org/10.5281/zenodo.20204995>.

Ethics and consent

This analysis uses secondary data sources. Data made available for the replication of this analysis is fully anonymised.

Funding

Bolyai János Research Scholarship of the Hungarian Academy of Sciences.

Acknowledgements

I would like to thank my Editor, my Reviewers, Tekla Gráci, and Márton Sóskuthy.

Competing interests

The author has no competing interests to declare.

ORCID IDs

Péter Rácz: <https://orcid.org/0000-0001-7896-4801>

References

- Ai, C., & Norton, E. C. (2003). Interaction terms in logit and probit models. *Economics Letters*, 80(1), 123–129. [https://doi.org/10.1016/S0165-1765\(03\)00032-6](https://doi.org/10.1016/S0165-1765(03)00032-6)
- Amenta, S., & Crepaldi, D. (2012). Morphological processing as we know it: An analytical review of morphological effects in visual word identification. *Frontiers in Psychology*, 3, 232. <https://doi.org/10.3389/fpsyg.2012.00232>
- Babel, M. (2012). Evidence for phonetic and social selectivity in spontaneous phonetic imitation. *Journal of Phonetics*, 40(1), 177–189. <https://doi.org/10.1016/j.wocn.2011.09.001>
- Beckner, C., Blythe, R., Bybee, J., Christiansen, M. H., Croft, W., Ellis, N. C., Holland, J., Ke, J., Larsen-Freeman, D., et al. (2009). Language is a complex adaptive system: Position paper. *Language Learning*, 59, 1–26. <https://doi.org/10.1111/j.1467-9922.2009.00533.x>
- Bock, J. K. (1986). Syntactic persistence in language production. *Cognitive Psychology*, 18(3), 355–387. [https://doi.org/10.1016/0010-0285\(86\)90004-6](https://doi.org/10.1016/0010-0285(86)90004-6)
- Bürkner, P.-C. (2021). Bayesian item response modeling in R with brms and Stan. *Journal of Statistical Software*, 100(5), 1–54. <https://doi.org/10.18637/jss.v100.i05>
- Fazekas, J., Sala, G., & Pine, J. (2024). Prime surprisal as a tool for assessing error-based learning theories: A systematic review. *Languages*, 9(4), 147. <https://doi.org/10.3390/languages9040147>
- Garrod, S., & Pickering, M. J. (2004). Why is conversation so easy? *Trends in Cognitive Sciences*, 8(1), 8–11. <https://doi.org/10.1016/j.tics.2003.10.016>
- Garrod, S., & Pickering, M. J. (2009). Joint action, interactive alignment, and dialog. *Topics in Cognitive Science*, 1(2), 292–304. <https://doi.org/10.1111/j.1756-8765.2009.01020.x>
- Gelman, A., Jakulin, A., Pittau, M. G., & Su, Y.-S. (2008). A weakly informative default prior distribution for logistic and other regression models. *The Annals of Applied Statistics*, 2(4), 1360–1383. <https://doi.org/10.1214/08-AOAS191>

- Giles, H., & Coupland, N. (1991). *Language: Contexts and consequences*. Thomson Brooks/Cole Publishing Co.
- Goodwin, C. (2018). *Co-operative action*. Cambridge University Press. <https://doi.org/10.1017/9781139016735>
- Gries, S. T. (2005). Syntactic priming: A corpus-based approach. *Journal of Psycholinguistic Research*, 34(4), 365–399. <https://doi.org/10.1007/s10936-005-6139-3>
- Gronau, Q. F., Singmann, H., & Wagenmakers, E.-J. (2020). Bridgesampling: An R package for estimating normalizing constants. *Journal of Statistical Software*, 92(10), 1–29. <https://doi.org/10.18637/jss.v092.i10>
- Hay, J. B., & Baayen, R. H. (2005). Shifting paradigms: Gradient structure in morphology. *Trends in Cognitive Sciences*, 9(7), 342–348. <https://doi.org/10.1016/j.tics.2005.04.002>
- Hay, J. B., Pierrehumbert, J. B., Walker, A. J., & LaShell, P. (2015). Tracking word frequency effects through 130 years of sound change. *Cognition*, 139, 83–91. <https://doi.org/10.1016/j.cognition.2015.02.012>
- Jaeger, T. F., & Snider, N. E. (2013). Alignment as a consequence of expectation adaptation: Syntactic priming is affected by the prime's prediction error given both prior and recent experience. *Cognition*, 127(1), 57–83. <https://doi.org/10.1016/j.cognition.2012.10.013>
- László, H. (2001). A felszólító módú igealakok kettősségének történetéhez [Regarding the history of the dual pattern of imperative verb forms]. In L. Búky & T. Forgács (Eds.), *A nyelvtörténeti kutatások újabb eredményei ii. magyar és finnugor alaktan [the latest results in diachronic linguistics ii: Finno-ugric morphology]* (pp. 55–65). Szegedi Tudományegyetem Magyar Nyelvészeti Tanszék.
- Lüdecke, D. (2025). *Sjplot: Data visualization for statistics in social science* [R package version 2.9.0]. <https://CRAN.R-project.org/package=sjPlot>
- Mády, K., Kohári, A., Reichel, U. D., Szalontai, Á., & Mihajlik, P. (2023). The Budapest Games Corpus. *Speech Research Conference*, 75. <https://doi.org/10.18135/BeszKutKonf.2023>
- Makowski, D., Ben-Shachar, M. S., & Lüdecke, D. (2019). Bayestestr: Describing effects and their uncertainty, existence and significance within the Bayesian framework. *Journal of Open Source Software*, 4(40), 1541. <https://doi.org/10.21105/joss.01541>
- Marslen-Wilson, W. D. (2007). Morphological processes in language comprehension. In M. Gareth Gaskell (Ed.), *The oxford handbook of psycholinguistics* (pp. 175–193). Oxford University Press. <https://doi.org/10.1093/oxfordhb/9780198568971.013.0011>
- Mihajlik, P., Mády, K., Kohári, A., Fruzsina, F. S., Kiss, G., Grácsi, T. E., & Doğruöz, A. S. (2024). Is spoken Hungarian low-resource?: A quantitative survey of Hungarian speech data sets. *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, 9382–9388. <https://doi.org/10.63317/4qrohu7pa37j>
- Molnár, C. S., Mády, K., Mihajlik, P., & Gyuris, B. (2023). The Akaka Maptask Corpus. *Beszédkutatás-Speech Research Conference*, 81–83. <https://doi.org/10.18135/BeszKutKonf.2023>
- Nemeskey, D. M. (2020). *Natural language processing methods for language modeling* [Doctoral dissertation]. Eötvös Loránd University [ELTE Digital Institutional Repository]. <https://doi.org/10.15476/ELTE.2020.066>

- Piantadosi, S. T., Tily, H., & Gibson, E. (2011). Word lengths are optimized for efficient communication. *Proceedings of the National Academy of Sciences*, 108(9), 3526–3529. <https://doi.org/10.1073/pnas.1012551108>
- Pickering, M. J., & Branigan, H. P. (1998). The representation of verbs: Evidence from syntactic priming in language production. *Journal of Memory and Language*, 39(4), 633–651. <https://doi.org/10.1006/jmla.1998.2592>
- Pickering, M. J., & Garrod, S. (2004). Toward a mechanistic psychology of dialogue. *Behavioral and Brain Sciences*, 27(2), 169–190. <https://doi.org/10.1017/S0140525X04000056>
- Pickering, M. J., & Garrod, S. (2006). Alignment as the basis for successful communication. *Research on Language and Computation*, 4(2–3), 203–228. <https://doi.org/10.1007/s11168-006-9004-0>
- R Core Team. (2025). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rácz, P. (2025). *Word frequency list from the Hungarian Webcorpus* (Version 1.0). Zenodo. <https://doi.org/10.5281/zenodo.17508385>
- Rácz, P., Beckner, C., Hay, J. B., & Pierrehumbert, J. B. (2020). Morphological convergence as on-line lexical analogy. *Language*, 96(4), 735–770. <https://doi.org/10.1353/lan.2020.0061>
- Rácz, P., & Lukács, Á. (2024). Lexical and social effects on the learning and integration of inflectional morphology. *Cognitive Science*, 48(8), e13483. <https://doi.org/10.1111/cogs.13483>
- Roberts, G. (2010). An experimental study of social selection and frequency of interaction in linguistic diversity. *Interaction Studies*, 11(1), 138–159. <https://doi.org/10.1075/is.11.1.06rob>
- Schumacher, R. A., & Pierrehumbert, J. B. (2021). Familiarity, consistency, and systematizing in morphology. *Cognition*, 212, 104512. <https://doi.org/10.1016/j.cognition.2020.104512>
- Siptár, P., & Törkenczy, M. (2000). *The phonology of Hungarian*. Oxford University Press. <https://doi.org/10.1093/oso/9780198238416.001.0001>
- Stivers, T., Brown, P., Englert, C., Hayashi, M., Heinemann, T., Hoymann, G., Rossano, F., De Ruiter, J. P., Yoon, K.-E., et al. (2009). Universals and cultural variation in turn-taking in conversation. *Proceedings of the National Academy of Sciences*, 106(26), 10587–10592. <https://doi.org/10.1073/pnas.0903616106>
- Szmrecsanyi, B. (2006). *Morphosyntactic persistence in spoken English: A corpus study at the intersection of variationist sociolinguistics, psycholinguistics, and discourse analysis*. Walter de Gruyter. <https://doi.org/10.1515/9783110197808>
- Tóth, E. (2007). The imperative and the subjunctive proper in Hungarian. *Sprachtheorie und germanistische Linguistik*, 17(2), 125–145.
- Vehtari, A., Gelman, A., & Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. <https://doi.org/10.1007/s11222-016-9696-4>
- Weiner, E. J., & Labov, W. (1983). Constraints on the agentless passive. *Journal of Linguistics*, 19(1), 29–58. <https://doi.org/10.1017/S0022226700007441>
- Wickham, H. (2016). *ggplot2: Elegant graphics for data analysis*. Springer-Verlag. <https://doi.org/10.1007/978-3-319-24277-4>

