

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Markedness as Surprisal: Information Theory and Language Inefficiency

Permalink

<https://escholarship.org/uc/item/8sq4699j>

Journal

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

Authors

Kilpatrick, Alexander

Bundgaard-Nielsen, Rikke L.

Iida, Hinano

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

Markedness as Surprisal: Information Theory and Language Inefficiency

Alexander Kilpatrick (alex@u-aizu.ac.jp)¹, Rikke Bundgaard-Nielsen² & Hinano Iida³

¹University of Aizu

²Melbourne University

³Nagoya University

Abstract

Information-theoretic approaches to language predict pressure toward efficient, low-surprisal encoding, yet languages systematically retain forms that violate these pressures. Such departures are traditionally described as marked. In this study, we model markedness as localized phonological inefficiency, quantified using phonemic bigram surprisal. Focusing on iconic words in English (e.g., *vroom*) and ideophones in Japanese (e.g., *fuyafuya*), we examine how surprisal is distributed within words relative to language-specific phonotactic baselines. Across both languages, iconic forms exhibit reliably higher surprisal than non-iconic controls, indicating increased phonological unpredictability. This inefficiency is not uniform: surprisal is concentrated at specific word-internal positions, particularly at word onsets. In Japanese, baseline asymmetries in CV structure strongly shape surprisal, but ideophones retain distinct information-structural profiles once these constraints are controlled. The results presented here are limited in that they are correlational and do not by themselves establish that phonological inefficiency is functionally deployed, but nonetheless demonstrate a systematic association between iconicity and elevated surprisal.

Keywords: Language Inefficiency; Information Theory; Phonological Surprisal; Markedness

Introduction

A central assumption in cognitive science is that cognition is shaped by pressures toward efficiency: mental systems are expected to minimize processing cost, reduce uncertainty, and make optimal use of the available and finite computational resources (Lieder & Griffiths, 2020; Lieder et al., 2018; Musslick & Cohen, 2020; Giudice & Crespi, 2018). In language and language processing, this assumption is sometimes formalized using information-theoretic constructs such as Surprisal (Shannon, 1948), leading to the observation that linguistic systems are characterized by a prevalence of predictable, low-cost forms while maintaining sufficient variability to express a full range of meanings (Gibson et al., 2019; Hahn et al., 2020; Mollica et al., 2021). Within this framework, deviations from efficiency are typically treated as novelty, noise, constraints, or historical residue.

However, and despite these assumptions, linguists have long engaged with systematic and recurrent departures from measures of efficiency, conceptualizing these departures as *markedness*. Marked forms are often less frequent, less predictable, and more complex than unmarked forms, yet persist across languages (Mollica et al., 2021; Haspelmath, 2021). Despite its descriptive usefulness, markedness has resisted formalization: it lacks a clear computational

definition, cannot be straightforwardly quantified, and generalizes poorly across linguistic systems (see discussions in Haspelmath [2006] and Berent [2017]).

The Attentional Optimization Hypothesis (AOH; Kilpatrick & Bundgaard-Nielsen, 2026) provides a potential framework for understanding such processing inefficiencies—including those arising from the processing of *marked* forms—as functional rather than problematic. The AOH proposes that linguistic systems are optimized not only for efficient information transmission, but also for the selective allocation of limited attentional resources. From this perspective, processing *inefficiency* can be a design feature: forms that are costly to process may attract attention to communicatively salient meanings.

Here, we reconceptualize linguistic markedness in explicitly information-theoretic terms by examining how information is distributed within English iconic words and Japanese ideophones—word types that are described as marked across languages (Thompson et al., 2021). We show that the internal information density of ideophones and iconic words systematically departs from phonological and distributional expectations in both languages. This is important because studies have shown increases processing difficulty locally while enhancing attention, and memorability (Kilpatrick & Bundgaard-Nielsen, 2025). By taking this approach, markedness emerges as a graded, model-relative property of intra-word surprisal distributions that reflects a trade-off between processing efficiency and attentional gain, consistent with the AOH.

The present study

The present study explores the hypothesis that linguistic markedness is a graded, model-relative property of information density rather than a categorical feature of a given lexical item. The relevant unit of markedness in this account is the word (or morphophonological string), and the baseline is given by language-specific expectations (Japanese and English) about how information is typically distributed across positions. We do so by employing surprisal as a measure of predictability in iconic language—where sound and meaning are linked—in two typologically distant and unrelated languages, English and Japanese. This approach allows us to assess whether markedness-as-information-structure generalizes across fundamentally different linguistic systems.

Background

Surprisal is an information-theoretic measure that quantifies how unexpected a given (here, linguistic) unit is in its context, defined as the negative logarithm of the conditional probability of that unit given the preceding material (Boeve & Bogaerts, 2025; Hale, 2001; Levy, 2008; Shannon, 1948; Venhuizen et al., 2019). Within surprisal theory, processing difficulty is assumed to increase with surprisal: the less predictable a word, structure, or sound is, the more cognitive effort is required to integrate it, a relationship robustly observed in behavioral measures such as reading times and in electrophysiological responses including the N400 and P600 across multiple languages (Aurnhammer et al., 2021; Michaelov et al., 2023; Staub, 2024; Wilcox et al., 2023).

Because of these clear predictions for processing, measuring surprisal has become a central tool in psycholinguistics and neurolinguistics for linking probabilistic language models to observable human behavior (Boeve & Bogaerts, 2025; Venhuizen et al., 2019; Wilcox et al., 2023). Here, surprisal-based accounts of processing are typically situated within a broader efficiency perspective, according to which languages must encode a rich space of meanings while respecting constraints on memory, attention, and processing time.

Efficiency-based theories propose that grammars and usage patterns evolve toward trade-offs that balance processing cost against expressivity and informativeness, as illustrated by lossy-context surprisal models showing how memory limitations and predictability jointly shape incremental processing (Futrell et al., 2020; Hahn et al., 2020). For example, cross-linguistic corpus work suggests that attested word orders are optimized relative to counterfactual alternatives, favoring structures that minimize processing difficulty under bounded memory while preserving necessary distinctions (Hahn & Futrell, 2019).

Information-theoretic approaches to lexical organization further suggest that word length and information content are systematically related. Several studies show that the average information a word conveys in context (its contextual informativity) is a better predictor of length than raw frequency, such that less predictable, more informative words tend to be longer (Piantadosi et al., 2011). In large multilingual corpora, word length in characters, phonemes, or syllables correlates positively with average surprisal, indicating that shorter words are typically used in more predictive contexts and thus carry less information, whereas longer words occur in less predictive contexts and compensate by encoding more information overall (Levshina, 2021). At the same time, work at the sublexical level demonstrates that, as words grow longer, the information density per segment tends to decrease: surprisal is often highest near the beginnings of words and lower in later positions, implying that longer words distribute their information more thinly across more segments (King & Wedel, 2020; Kilpatrick & Bundgaard-Nielsen, 2025).

Taken together, this body of work portrays language as shaped by information-theoretic principles that balance the

communicative value of high surprisal against its processing costs, positioning surprisal as a central tool for understanding both efficient encoding and systematic departures from efficiency (Futrell et al., 2020; Hahn et al., 2020; Pimentel et al., 2021; Wilcox et al., 2023).

Although initially applied at the lexical and syntactic levels, information theoretic analyses have been extended to the phonological domain, where phonological or phonemic surprisal is computed from transitional probabilities between sounds, capturing how unexpected a given segment or sound sequence is within a language's phonotactic system (Brandt et al., 2021; Gaston et al., 2021; Malisz et al., 2018). At this sublexical level, higher-surprisal segments tend to be produced with more expanded, hyper-articulated phonetic realizations. They also elicit stronger neural responses, indicating increased processing demands (Brandt et al., 2021; Gaston et al., 2021; Malisz et al., 2018; Puffay et al., 2023).

At the phonological level, efficiency pressures, like those described for sentence processing above, are also reflected in a surprisal–duration trade-off. Across languages, more surprising phones are produced with longer durations (e.g., Shaw & Kawahara, 2019), while highly predictable phones are compressed, consistent with a tendency to maintain a relatively stable information rate (Pimentel et al., 2021). Cultural-evolution simulations further show that core properties of language, including combinatorial sound systems and compositional structure, can emerge from tensions between compressibility and expressivity, mirroring the same efficiency trade-offs assumed in surprisal-based accounts (Kirby et al., 2015).

Crucially, surprisal-based inefficiency on the phonological level need not be treated as a failure of optimization. The AOH (Kilpatrick & Bundgaard-Nielsen, 2026) proposes that linguistic systems are shaped by trade-offs between, on one hand, maximizing understanding and minimizing processing cost, and on the other hand, by modulating attentional engagement, such that departures from efficiency can be functional rather than incidental, inhibitive or disruptive.

Indeed, under the AOH, while some linguistic forms are shaped for predictability or information compression (a pressure to reduce surprisal), others may be shaped or preserved in a form that requires the recruitment of additional attentional resources, particularly at moments of communicative importance (a pressure to introduce or preserve surprisal). Such a pattern of elevated levels of surprisal for *some* wordforms has been established for words which score highly on measures of vividness in English (i.e., *concreteness*, *imaginability* and *specificity*: Kilpatrick & Bundgaard-Nielsen, 2026). The AOH thus reframes efficiency not as a *global* objective to be maximized uniformly across a linguistic system, in a manner that aligns naturally with surprisal-based models of language processing but extends them by explaining why languages systematically tolerate—and even exploit—*localized* inefficiencies rather than eliminating them. One such potential domain is ideophonic words.

Ideophones are commonly defined as marked words that depict sensory imagery (Dingemanse, 2012). Their markedness is consistently observed cross-linguistically in phonological and morphological structure: ideophones tend to exhibit skewed phonotactic distributions, expressive morphology, and non-canonical configurations relative to ordinary lexical items (Akita & Dingemanse, 2019). Ideophones are systematically associated with iconicity, exhibiting phonological structures that depart from baseline expectations through non-canonical and expressive patterning. Consequently, iconic words also tend to deviate from language-wide phonotactic distributions. These deviations predict elevated surprisal for iconic forms across languages.

Method

Operationlisation

In this study, we operationalize markedness using the measure of phonemic bigram surprisal, and we examine its distribution in Japanese ideophonic word forms, and English iconic word forms, respectively. *Japanese* has a rich and well-described inventory of mimetic (i.e., iconic) words. These forms exhibit several marked phonological properties that align with cross-linguistic tendencies observed for ideophones (e.g., Hamano, 1998).

Although *English* lacks a dedicated ideophone class comparable to that of Japanese, English *iconic* words also display structural properties that set them apart from ordinary lexical items (Dingemanse, 2019; Winter et al., 2023). Iconic markedness strategies in English include violations of canonical phonotactic constraints (e.g., *vroom* /vɹu:m/), the use of marginal or non-native speech sounds (e.g., *ugh* /ʌχ/), consonant gemination or expressive prolongation (e.g., *grrrr* /gɹ:/), and vowel lengthening (e.g., *whaaat* /wa:t/) (Voeltz & Kilian-Hatz, 2001). While English and Japanese differ in how they mark iconic forms, such specific feature-based descriptions do not address how markedness is distributed across word-internal structure. In what follows, we approach this issue from an information-theoretic perspective in an analysis of marked words in two typologically different languages: English and Japanese.

Japanese ideophones are disproportionately associated with phonological features that are otherwise rare or restricted in the language, such as word-initial /p/, voiced obstruents, medial /ɾ/, and vowel harmony (Akita, 2021). For instance, while word-initial /p/ is generally prohibited in native Japanese vocabulary, approximately 15.25% of Japanese ideophones begin with /p/, as in /patʃin/ ‘*slap*’ and /pikapika/ ‘*glittering*’ (Akita, 2021). These properties distinguish mimetics from ordinary lexical items and have long been taken as evidence of their structural markedness within the Japanese lexicon.

Calculating Surprisal

We calculated bigram surprisal on diphone transitional probabilities for both English and Japanese. For *English*,

lexical frequency data were drawn from SUBTLEX-US, a corpus of approximately 50 million tokens covering roughly 50,000 unique word types derived from subtitle data of spoken American English (Brysbaert & New, 2009). These forms were cross-referenced with the Carnegie Mellon University Pronouncing Dictionary to obtain phonemic transcriptions (Weide, 1998); items without a match were excluded. Bigram surprisal was computed over the resulting phonemic dataset. Morpheme counts were obtained by cross-referencing an additional morphological database (Sánchez-Gutiérrez et al., 2018). Iconicity was operationalized using mean iconicity judgments, indicating how much each word sounds like its meaning (Winter et al., 2023). Words in the top 10% of the iconicity distribution were measured against the remaining 90%. This percentile-based split isolates strongly iconic forms while retaining sufficient statistical power in the non-iconic reference set.

The *Japanese* dataset was drawn from the Corpus of Spontaneous Japanese, which provides phonemic transcriptions of spoken language (Maekawa, 2003). Japanese lexical items were classified using existing annotations (Iida & Akita, 2023). Words labeled as ideophones were analyzed as a distinct category and compared against non-ideophonic lexical items. Japanese has a five-vowel system with a length distinction and a highly constrained phonotactic system, favouring /CV/ syllables. As a result, the transitional probability of consonant–vowel (/CV/) bigrams is systematically higher than that of vowel–consonant (/VC/) bigrams, independent of any language-specific markedness effects. Failing to control for this baseline asymmetry would inflate surprisal estimates for /VC/ transitions and obscure whether observed differences reflect genuine phonotactic markedness or merely inventory-driven distributional constraints. We address this issue by including even/odd parity in the Japanese models.

Hypotheses

If language *efficiency* is the only goal structuring linguistic systems (here wordforms), surprisal would be expected to be distributed smoothly across all wordforms in accordance with probabilistic regularities. If, however, linguistic systems are shaped by pressures beyond efficiency alone, potentially including attentional factors, to focus attentional resources—our information theoretic operationalization of markedness would reveal atypical form-internal information distributions. Specifically, marked forms (e.g., Japanese ideophones and English iconic words) would present localized spikes or elevated variance in informativity. The present data allows us to formulate the following hypotheses:

H1: Iconic words in English and ideophones in Japanese will exhibit higher average phonemic surprisal than non-iconic and non-ideophonic lexical items from the same language.

H2: Markedness is expressed as a systematic redistribution of information within ideophonic/iconic words that mirrors established length–surprisal effects.

Results

English Iconicity

To test **H1**, whether iconicity predicts phonemic surprisal independent of length, a linear regression model was fitted with surprisal as the dependent variable and iconicity group (top 10% vs. remaining 90%) and bigram position as predictors. The model revealed significant effects of both predictors ($F(2, 68149) = 641.3, p < .001$), though overall explained variance was modest ($R^2 = 0.018$). There was a significant main effect of iconicity group: highly iconic words exhibited higher surprisal than non-iconic words ($\beta = 0.192, SE = 0.026, t = 7.37, p < .001$). This indicates that segments in highly iconic words are less predictable relative to the language-wide phonotactic model. Position also exhibited a strong negative effect on surprisal ($\beta = -0.111, SE = 0.003, t = -34.07, p < .001$), reflecting a general tendency for surprisal to decrease as a word unfolds. This reflects the incremental nature of phonological processing, whereby each successive segment constrains the probability space for subsequent ones, rendering later positions more predictable on average.

To examine how iconicity-related differences in surprisal vary across word-internal position, a series of simple linear regression models were fitted separately for each bigram position, with surprisal as the dependent variable and iconicity group as the sole predictor (Figure 1). Table 1 reports the estimated effect of iconicity for positions 1–7. A strong effect was observed at the initial position: highly iconic words showed substantially higher surprisal than non-iconic words at position 1 ($\beta = 0.506, SE = 0.043, t = 11.8, p < .001$). A smaller but still significant effect was also present at position 2 ($\beta = 0.123, SE = 0.057, t = 2.17, p = .03$). From position 3 onward, no reliable differences between iconic and non-iconic words were observed.

Table 1: Front-to-Back English Position models.

Position	Estimate	Std. Error	t-value	p-value
1	0.506	0.043	11.8	< 0.001
2	0.123	0.057	2.17	0.03
3	0.041	0.065	0.62	0.533
4	0.116	0.074	1.57	0.117
5	-0.063	0.091	-0.69	0.489
6	-0.044	0.123	-0.36	0.721
7	0.015	0.172	0.09	0.929

To assess whether iconicity-related differences in surprisal are also present toward the ends of words, a series of linear regression models were fitted for each position counted from the word-final edge (Figure 2). In these *Back-to-Front* models, surprisal was predicted from iconicity group (top 10% vs. remaining 90%), with word length included as a covariate to control for baseline length-related differences in information distribution. Table 2 reports the estimated effect of iconicity for each word-final position. Significant effects of iconicity were observed at the word-final position ($\beta = 0.235, SE = 0.046, t = 5.12, p < .001$), the second-to-last position ($\beta = 0.127, SE = 0.049, t = 2.58, p < .01$), and the

third-to-last position ($\beta = 0.226, SE = 0.066, t = 3.41, p < .001$). In each case, highly iconic words exhibited higher surprisal than non-iconic words, even after controlling for word length. No reliable differences were observed at more distant positions from the word-final edge.

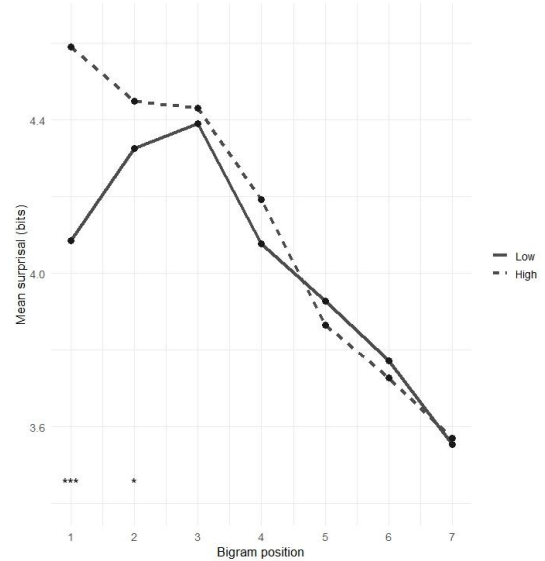


Figure 1: Front-to-Back Surprisal in English words. Top 10% iconicity scores against all other words. Asterisks represent significant differences in the position models.

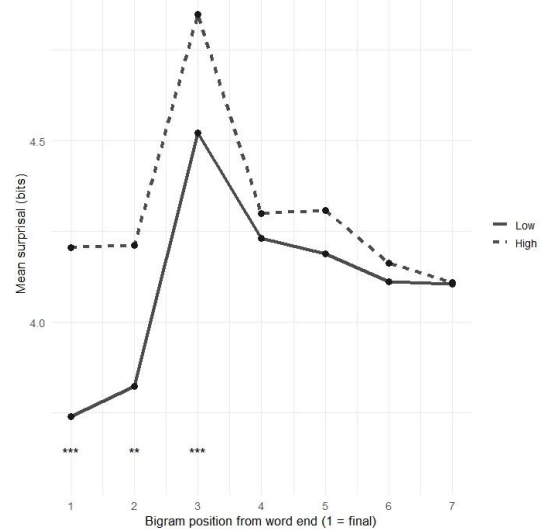


Figure 2: Back-to-Front Surprisal in English words. Top 10% iconicity scores against all other words. Asterisks represent significant differences in the position models.

Table 2: Back-to-Front English Position models.

Position	Estimate	Std. Error	t-value	p-value
Final	0.235	0.046	5.12	< 0.001
2 nd last	0.127	0.049	2.58	< 0.01
3 rd last	0.226	0.066	3.41	< 0.001
4 th last	0.063	0.075	0.84	0.4
5 th last	0.142	0.094	1.51	0.131
6 th last	0.048	0.134	0.36	0.722
7 th last	0.056	0.191	0.29	0.769

Japanese Ideophones

As discussed in the Method section, the analysis of Japanese ideophones is complicated by a systematic even–odd parity effect arising from Japanese phonotactics. Because Japanese strongly favours consonant–vowel structure, bigram position is tightly coupled with segment type: odd positions tend to correspond to consonant–vowel transitions, while even positions tend to correspond to vowel–consonant transitions. As a result, surprisal differs sharply between even and odd positions independent of lexical class. This structural asymmetry is visible in Figure 3, where alternating peaks and troughs in surprisal reflect parity rather than word-internal information density per se.

To address this issue, surprisal was modelled as a function of lexical group (ideophone vs. non-ideophone), parity (even vs. odd position), and position, including an interaction between group and parity (Figure 4). This model allows the strong parity-driven alternation in surprisal to be explicitly controlled, while testing whether ideophones differ from other lexical items beyond this baseline structural effect.

In line with **H1**, the model revealed a robust main effect of lexical group: non-ideophonic words exhibited significantly lower surprisal than ideophones ($\beta = -0.529$, $SE = 0.102$, $t = -5.17$, $p < .001$), indicating that ideophones are characterized by higher overall phonemic unpredictability. There was also a very strong main effect of parity, with odd positions showing substantially lower surprisal than even positions ($\beta = -1.867$, $SE = 0.117$, $t = -15.94$, $p < .001$), confirming the magnitude of the /CV/–/VC/ asymmetry in Japanese phonotactics. The effect of linear position was small and marginal ($\beta = -0.031$, $SE = 0.016$, $t = -1.93$, $p = .053$), suggesting that once parity is accounted for, front-to-back position contributes relatively little additional variance. Critically, the interaction between lexical group and parity was significant ($\beta = 0.290$, $SE = 0.135$, $t = 2.15$, $p = .031$), indicating that the parity effect differs between ideophones and non-ideophonic words. This interaction shows that ideophones do not simply inherit the baseline /CV/–/VC/ alternation of Japanese phonotactics, but modulate it, exhibiting elevated surprisal at specific parity-defined positions: after controlling for parity, ideophones retain distinct internal information density.

To further examine how ideophonic status affects surprisal across word-internal position in Japanese, a series of front-to-back simple linear regression models were fitted separately for each position, with surprisal as the dependent variable and lexical group (ideophone vs. non-ideophone) as the sole predictor. The front-to-back analyses revealed significant group differences at multiple early and mid-positions. Ideophones exhibited higher surprisal than non-ideophonic words at position 1 ($\beta = 0.414$, $SE = 0.132$, $t = 3.14$, $p = .002$), position 2 ($\beta = 0.749$, $SE = 0.191$, $t = 3.93$, $p < .001$), position 3 ($\beta = 0.366$, $SE = 0.121$, $t = 3.01$, $p = .003$), position 4 ($\beta = 0.362$, $SE = 0.19$, $t = 1.9$, $p = .058$), position 5 ($\beta = -0.251$, $SE = 0.17$, $t = -1.48$, $p = .141$), position 6 ($\beta = 0.434$, $SE = 0.297$, $t = 1.46$, $p = .145$), and position 7 ($\beta = 0.365$, $SE = 0.128$, $t = 2.86$, $p = .005$).

.005). Table 3 reports the output of these models. *Back-to-front* models were not conducted for Japanese because reversing position conflates word-final distance with even–odd parity, making the resulting estimates uninterpretable with respect to lexical class.

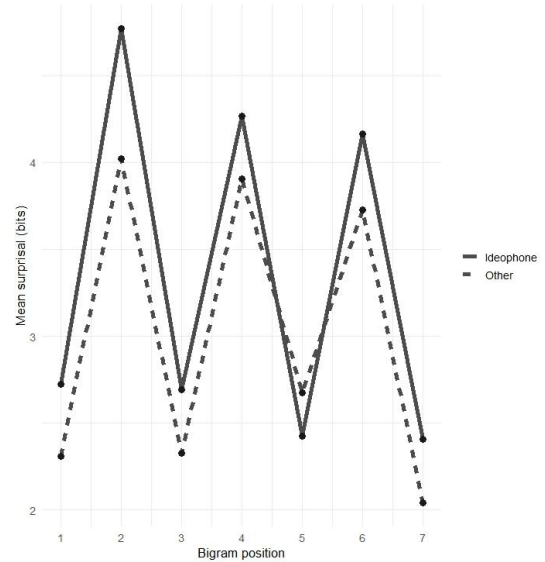


Figure 3: Unmodified Front-to-Back Surprisal in Japanese words.

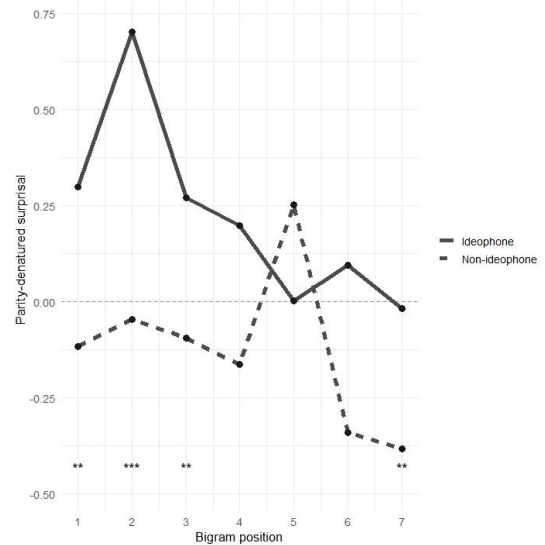


Figure 4: Parity-Denatured Front-to-Back Surprisal in Japanese words. Asterisks represent significant differences according to the position models.

Table 3: Front-to-Back Japanese Position models.

Position	Estimate	Std. Error	t-value	p-value
1	0.414	0.132	3.14	0.002
2	0.749	0.191	3.93	< 0.001
3	0.366	0.121	3.01	0.003
4	0.362	0.19	1.9	0.058
5	-0.251	0.17	-1.48	0.141
6	0.434	0.297	1.46	0.145
7	0.365	0.128	2.86	0.005

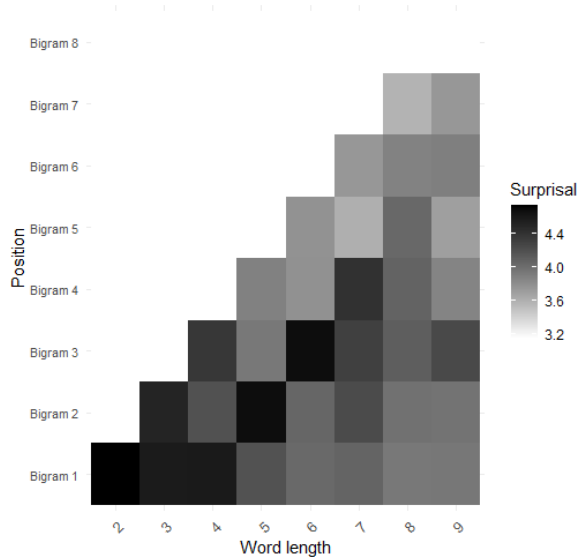


Figure 5: Heatmap of surprisal in monomorphemic English words according to word length and position.

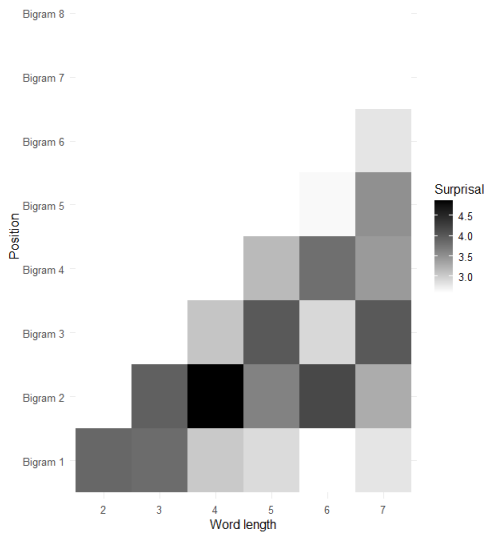


Figure 6: Heatmap of surprisal in Japanese morphemes according to word length and position.

Word Length and Surprisal

To explore **H2**, we constructed two simple linear regression models, one for English and one for Japanese. Both the English ($t(1,44558) = -42.45, p < .001, R^2 = 0.039$) and Japanese ($t(1,12775) = -15.62, p < .001, R^2 = 0.019$) models revealed significant negative correlations. To further explore **H2**, we examined whether the canonical length-surprisal relationship manifests differently across word-internal positions, testing the prediction that inefficiency is concentrated at perceptually salient boundaries rather than distributed uniformly across the word. Longer words and morphemes tended to have lower informativity at their onset ($\beta = -0.139, p < .001, R^2 = 0.033$), and similarly at their end ($\beta = -0.171, p < .001, R^2 = 0.043$), highlighting the significance of word boundaries in information expression.

This length-driven reduction in boundary surprisal reflects a dilution effect: as words grow longer, information is spread across more segments, lowering average density even at edges. This is compatible with the boundary-salience pattern observed for iconic words, which holds within length-matched comparisons; the two effects operate at different levels of analysis. The Word Middle model showed a smaller effect size ($\beta = -0.017, p = .027, R^2 < 0.001$), suggesting a weaker influence of word length on informativity in the middle of words. Figures 5 and 6 presents heatmaps of surprisal by word length and bigram position.

Discussion

This study reconceptualizes linguistic markedness as a property of phonemic surprisal, operationalized as deviations from baseline expectations about information distribution within words. English iconic words and Japanese ideophones exhibit reliably higher surprisal than comparison forms, indicating that markedness can be captured as increased unpredictability relative to a listener's probabilistic model rather than as a categorical lexical distinction. It is important to note, however, that the present results are correlational and do not uniquely support the AOH over alternatives. A natural alternative is that iconic forms are subject to additional phonological constraints arising from sound-meaning mapping. If iconic words must jointly satisfy iconicity-driven and efficiency-driven pressures while non-iconic words satisfy only the latter, elevated surprisal in iconic forms follows without appeal to functional inefficiency. Distinguishing between these accounts requires evidence that links phonological surprisal in iconic words directly to downstream consequences for attention and memory rather than form-meaning correspondence constraints alone.

Regardless, our approach provides a quantifiable alternative to traditional notions of markedness, which have lacked a formal computational definition. And critically, markedness effects are not distributed uniformly across words. Elevated surprisal associated with iconicity or ideophonic status is concentrated at word boundaries, particularly at word onset and, to a lesser extent, near word-final positions. Interior positions show weaker effects. These patterns closely parallel established findings on the relationship between word length and surprisal, where information density is highest at word edges. Markedness may thus exploit the same incremental processing, concentrating information where it is perceptually salient.

The Japanese results further highlight the importance of language-specific baselines. Strong even-odd parity effects driven by CV phonotactics substantially shape surprisal, but once these are controlled, ideophones retain distinct surprisal profiles. This demonstrates that markedness emerges relative to a language's predictive structure rather than from raw phonotactic rarity. Overall, these findings support an information-theoretic view of markedness as *structured inefficiency*: marked forms introduce localized processing costs at word boundaries rather than globally violating efficiency.

References

- Akita, K. (2021). Onomatopoeic “husizen”-na onshoochoo (“Unnatural” sound symbolism in ideophones). In Waseda Bungaku Kai (Ed.), *Waseda Bungaku 2021-nen Harugoo* (Waseda literature, spring 2021 issue), 24–36. Tokyo: Chikuma Shobo.
- Akita, K., & Dingemanse, M. (2019). Ideophones (Mimetics, Expressives). *Oxford Research Encyclopedia of Linguistics*.
<https://doi.org/10.1093/acrefore/9780199384655.013.477>.
- Aurnhammer, C., Delogu, F., Schulz, M., Brouwer, H., & Crocker, M. (2021). Retrieval (N400) and integration (P600) in expectation-based comprehension. *PLOS ONE*, 16(9), e0257430.
- Berent, I. (2017). Is markedness a confused concept?. *Cognitive Neuropsychology*, 34, 493–499.
<https://doi.org/10.1080/02643294.2017.1422485>.
- Boeve, S., & Bogaerts, L. (2025). A systematic evaluation of Dutch large language models’ surprisal estimates in sentence, paragraph and book reading. *Behavior Research Methods*.
- Brandt, E., Möbius, B., & Andreeva, B. (2021). Dynamic formant trajectories in German read speech: Impact of predictability and prominence. *Speech Communication*, 133, 1–14.
- Brysbaert, M., & New, B. (2009). Moving beyond Kučera and Francis: A critical evaluation of current word frequency norms and the introduction of a new and improved word frequency measure for American English. *Behavior Research Methods*, 41(4), 977–990.
<https://doi.org/10.3758/BRM.41.4.977>
- Del Giudice, M., & Crespi, B. J. (2018). Basic functional trade-offs in cognition: An integrative framework. *Cognition*, 179, 56–70.
- Dingemanse, M. (2012). Advances in the cross-linguistic study of ideophones. *Language and Linguistics compass*, 6(10), 654–672.
- Dingemanse, M. (2019). Playful iconicity: structural markedness underlies the relation between funniness and iconicity. *Language and Cognition*, 12, 203–224.
<https://doi.org/10.1017/langcog.2019.49>.
- Flaksman, M., & Kilpatrick, A. (2025). Against the tide: How language-specificity of imitative words increases with time (as evidenced by Surprisal). *SKASE Journal of Theoretical Linguistics*, 22(2).
- Futrell, R., Gibson, E., & Levy, R. (2020). Lossy-context surprisal: An information-theoretic model of memory effects in sentence processing. *Cognitive Science*, 44(3), e12814.
- Futrell, R., & Hahn, M. (2022). Information theory as a bridge between language function and language form. *Frontiers in Communication*, 7, 657725.
- Gaston, P., Brodbeck, C., Phillips, C., & Lau, E. F. (2021). Auditory word comprehension is less incremental in isolated words. *Neurobiology of Language*, 2(4), 567–596.
- Hahn, M., & Futrell, R. (2019). Crosslinguistic word orders enable an efficient tradeoff of memory and surprisal. In *Proceedings of NAACL-HLT 2019* (pp. 381–390).
- Hahn, M., Degen, J., & Futrell, R. (2020). Modeling word and morpheme order in natural language as an efficient trade-off of memory and surprisal. *Psychological Review*, 127(4), 684–732.
- Hahn, M., Mathew, R., & Degen, J. (2021). Morpheme ordering across languages reflects optimization for processing efficiency. *Open Mind: Discoveries in Cognitive Science*, 5, 365–393.
- Hale, J. (2001). A probabilistic Earley parser as a psycholinguistic model. In *Proceedings of the 2nd Meeting of the North American Chapter of the Association for Computational Linguistics* (pp. 1–8).
- Hamano, S. (1998). *The sound-symbolic system of Japanese*. CSLI Publications.
- Haspelmath, M. (2006). Against markedness (and what to replace it with). *Journal of Linguistics*, 42(1), 25–70.
- Haspelmath, M. (2021). Explaining grammatical coding asymmetries: Form–frequency correspondences and predictability. *Journal of Linguistics*, 57, 605–633.
<https://doi.org/10.1017/s0022226720000535>.
- Iida, H., & Akita, K. (2023). Perceptual strength norms for 510 Japanese words, including ideophones: A comparative study with English. In *Proceedings of the Annual Meeting of the Cognitive Science Society (CogSci 2023)* (pp. 2201–2207).
- King, A., & Wedel, A. (2020). Greater Early Disambiguating Information for Less-Probable Words: The Lexicon Is Shaped by Incremental Processing. *Open Mind : Discoveries in Cognitive Science*, 4, 1–12.
https://doi.org/10.1162/opmi_a.00030.
- Kilpatrick, A., & Bundgaard-Nielsen, R. (2025). Exploring the dynamics of Shannon’s information and iconicity in language processing and lexeme evolution. *PLOS ONE*, 20(4).
- Kilpatrick, A., & Bundgaard-Nielsen, R. (2026). Say it like you mean it: Linguistic vividness and the attentional optimization hypothesis. *Cognition*, 269, 106406.
- Kirby, S., Tamariz, M., Cornish, H., & Smith, K. (2015). Compression and communication in the cultural evolution of linguistic structure. *Cognition*, 141, 87–102.
- Levshina, N. (2021). Frequency, informativity and word length: Insights from typologically diverse corpora. *Entropy*, 23(12).
- Levy, R. (2008). Expectation-based syntactic comprehension. *Cognition*, 106(3), 1126–1177.
- Lieder, F., & Griffiths, T. L. (2020). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43, e1.
- Lieder, F., Shenhav, A., Musslick, S., & Griffiths, T. L. (2018). Rational metareasoning and the plasticity of cognitive control. *PLoS Computational Biology*, 14(4), e1006043.

- Maekawa, K. (2003). Corpus of spontaneous Japanese: Its design and evaluation. In *Proceedings of the ISCA & IEEE Workshop on Spontaneous Speech Processing and Recognition* (pp. 7–12).
- Malisz, Z., Brandt, E., Möbius, B., Oh, Y., & Andreeva, B. (2018). Dimensions of segmental variability: Interaction of prosody and surprisal in six languages. *Frontiers in Communication*, 3, 25.
- Michaelov, J. A., Bardolph, M. D., Van Petten, C., Bergen, B., & Coulson, S. (2023). Strong prediction: Language model surprisal explains multiple N400 effects. *Neurobiology of Language*, 4(2), 199–232.
- Mollica, F., & Piantadosi, S. T. (2019). Humans store about 1.5 megabytes of information during language acquisition. *Royal Society Open Science*, 6(3), 181393.
- Mollica, F., Bacon, G., Zaslavsky, N., Xu, Y., Regier, T., & Kemp, C. (2021). The forms and meanings of grammatical markers support efficient communication. *Proceedings of the National Academy of Sciences*, 118. <https://doi.org/10.17605/osc.io/s5b7h>.
- 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.
- Pimentel, T., Meister, C., Salesky, E., Teufel, S., Blasi, D. E., & Cotterell, R. (2021). A surprisal–duration trade-off across and within the world’s languages. *arXiv preprint arXiv:2109.12052*.
- Puffay, C., Vanthornhout, J., Gillis, M., Accou, B., Van Hamme, H., & Francart, T. (2023). Robust neural tracking of linguistic speech representations using a convolutional neural network. *Journal of Neural Engineering*, 20(5), 056007.
- Rubio-Fernández, P., Mollica, F., & Jara-Ettinger, J. (2020). Speakers and listeners exploit word order for communicative efficiency: A cross-linguistic investigation. *Journal of Experimental Psychology: General*, 149(9), 1762–1783.
- Sánchez-Gutiérrez, C. H., Mailhot, H., Deacon, S. H., & Wilson, M. A. (2018). MorphoLex: A derivational morphological database for 70,000 English words. *Behavior Research Methods*, 50, 1568–1580. <https://doi.org/10.3758/s13428-017-0981-8>
- Shannon, C. E. (1948). A mathematical theory of communication. *Bell System Technical Journal*, 27, 379–423.
- Shaw, J. A., & Kawahara, S. (2019). Effects of surprisal and entropy on vowel duration in Japanese. *Language and speech*, 62(1), 80–114.
- Staub, A. (2024). Predictability in language comprehension: Prospects and problems for surprisal. *Annual Review of Linguistics*, 10, 1–21.
- Thompson, A. L., Van Hoey, T., & Do, Y. (2021). Articulatory features of phonemes pattern to iconic meanings: Evidence from cross-linguistic ideophones. *Cognitive Linguistics*, 32(4), 563–608.
- Venhuizen, N. J., Crocker, M. W., & Brouwer, H. (2019). Expectation-based comprehension: Modeling the interaction of world knowledge and linguistic experience. *Discourse Processes*, 56(3), 229–255.
- Voeltz, F. F., & Kilian-Hatz, M. (Eds.). (2001). *Ideophones*. John Benjamins Publishing Company. <https://doi.org/10.1075/tsl.45>
- Weide, R. (1998). *The Carnegie Mellon pronouncing dictionary* (Release 0.6). Carnegie Mellon University. <https://www.cs.cmu.edu/~cmusphinx/pronouncing-dictionary/>
- Wilcox, E. G., Pimentel, T., Meister, C., Cotterell, R., & Levy, R. (2023). Testing the predictions of surprisal theory in 11 languages. *Transactions of the Association for Computational Linguistics*, 11, 144–161.
- Winter, B., Dingemanse, M., Perlman, M., Lupyan, G., & Perry, L. K. (2023). Iconicity ratings for 14,000+ English words. *Behavior Research Methods*. Advance online publication. <https://doi.org/10.3758/s13428-023-02112-6>