

If Attention Serves as a Cognitive Model of Human Memory Retrieval, What is the Plausible Memory Representation?

Anonymous ACL submission

Abstract

Recent work in computational psycholinguistics has revealed intriguing parallels between attention mechanisms and human memory retrieval, focusing primarily on Transformer architectures that operate on token-level representations. However, computational psycholinguistic research has also established that syntactic structures provide compelling explanations for human sentence processing that word-level factors alone cannot fully account for. In this study, we investigate whether the attention mechanism of Transformer Grammar (TG), which uniquely operates on syntactic structures as representational units, can serve as a cognitive model of human memory retrieval, using Normalized Attention Entropy (NAE) as a linking hypothesis between model behavior and human processing difficulty. Our experiments demonstrate that TG's attention achieves superior predictive power for self-paced reading times compared to vanilla Transformer's, with further analyses revealing independent contributions from both models. These findings suggest that human sentence processing involves dual memory representations—one based on syntactic structures and another on token sequences—with attention serving as the general retrieval algorithm, while highlighting the importance of incorporating syntactic structures as representational units.

1 Introduction

Whether language models (LMs) developed in natural language processing (NLP) are plausible as cognitive models of human sentence processing is a central question in computational psycholinguistics. Over the past two decades, this question has been primarily addressed from the perspective of *expectation-based theories*—one of the two major classes of human sentence processing theory—examining whether LMs' next-word prediction can serve as a model of human predictive

processing (Hale, 2001; Levy, 2008; Wilcox et al., 2020; Merx and Frank, 2021; *inter alia*).

The recent success of Transformer architectures (Vaswani et al., 2017) in NLP has unexpectedly opened a new avenue of investigation from the perspective of *memory-based theories*, the other major class of sentence processing theory. Researchers have proposed that the attention mechanism, despite its engineering origins, can implement a human memory retrieval theory known as cue-based retrieval (Van Dyke and Lewis, 2003). Recent studies have revealed intriguing parallels between the weighted reference patterns exhibited by the attention mechanism and the elements that humans may retrieve during online sentence comprehension (Ryu and Lewis, 2021; Oh and Schuler, 2022; Timkey and Linzen, 2023).

Computational psycholinguistics has also established that human sentence processing cannot be fully explained by word-level factors alone; rather, *syntactic structures* have provided compelling explanations for it. For instance, next-word prediction from LMs that explicitly incorporate syntactic structure building demonstrates superior performance in accounting for human brain activity compared to vanilla RNNs and Transformers (Hale et al., 2018; Wolfman et al., 2024); the number of syntactic nodes hypothesized to be constructed per word correlates significantly with both reading times (Kajikawa et al., 2024) and neural activity patterns (Brennan et al., 2012).

Given these findings, if attention can serve as a general algorithm for memory retrieval in human sentence processing, human memory retrieval should be captured by the attention mechanism operating on syntactic structures as well as that operating on token sequences. In this study, we investigate whether the attention mechanism of Transformer Grammar (TG; Sartran et al., 2022), which uniquely operates on syntactic structures as representational units, can serve as a cognitive model of

human memory retrieval, using Normalized Attention Entropy (NAE; Oh and Schuler, 2022) as the linking hypothesis between model behavior and human processing difficulty. Our experiments demonstrate that TG’s attention achieves superior predictive power for self-paced reading times compared to vanilla Transformer’s, with further analyses revealing independent contributions from both models. These findings suggest that human sentence processing involves dual memory representations—one based on syntactic structures and another on token sequences—with attention serving as the general retrieval algorithm, while highlighting the importance of incorporating syntactic structures as representational units.

2 Background

2.1 Normalized Attention Entropy (NAE)

Many psycholinguistic studies assume that human sentence processing involves memory retrieval, where based on the various cues provided by the current input word (e.g., verbs), elements (e.g., their arguments) are retrieved from working memory. For example, in (1), which is taken from Van Dyke (2002), when the verb *was complaining* is input, its subject *the resident* must be retrieved from working memory.

- (1) a. The worker was surprised that the **resident**_[subj,anim] [who was living near the dangerous warehouse] *was complaining* about the investigation.
- b. The worker was surprised that the **resident**_[subj,anim] [who said that the warehouse_[subj] was dangerous] *was complaining* about the investigation.

According to the cue-based retrieval theory (Van Dyke and Lewis, 2003), such retrieval becomes more difficult when similar elements exist in the sentence because the cues are overloaded; for example, only in (1b), *warehouse* may interfere with *resident* since they both have the feature [subj] as a retrieval cue. Van Dyke (2002) showed that humans read *was complaining* more slowly in (1b) than in (1a), providing empirical support for the cue-based retrieval theory.

In recent computational psycholinguistics, attempts have been made to interpret the attention mechanism—a weighted reference of preceding tokens based on Query and Key vectors—as a computational implementation of cue-based retrieval.

Notably, Ryu and Lewis (2021) proposed Attention Entropy (AE) as a linking hypothesis, where the diffuseness of attention weights is assumed to quantify the degree of retrieval interference. While AE was initially proposed for modeling interference effects in specific constructions, Oh and Schuler (2022) extended it to naturally occurring text by introducing two normalizations: (i) division by the maximum entropy achievable given the number of preceding tokens, and (ii) sum-to-1 renormalization of attention weights over $1-(i-1)$ -th tokens (Normalized AE, NAE).¹

$$\text{NAE}_{l,h,i} = \frac{\mathbf{a}_{l,h,i[1:i-1]}^\top}{\log_2(i-1) \mathbf{1}^\top \mathbf{a}_{l,h,i[1:i-1]}} \left(\log_2 \frac{\mathbf{a}_{l,h,i[1:i-1]}}{\mathbf{1}^\top \mathbf{a}_{l,h,i[1:i-1]}} \right), \quad (1)$$

where $\mathbf{a}_{l,h,i}$ represents the attention weight vector when using the i -th token as the query in the h -th head of layer l .² In this paper, we employ this NAE as a linking hypothesis between attention mechanisms and human memory retrieval.³

2.2 Transformer Grammar (TG)

Transformer Grammar (TG; Sartran et al., 2022) is a type of syntactic LM, a generative model that jointly generates token sequences $\mathbf{x}$ and their corresponding syntactic structures $\mathbf{y}$. TG formulates the generation of $(\mathbf{x}, \mathbf{y})$ as modeling a sequence of actions, $\mathbf{a}$ (e.g., (S (NP The blue bird NP) (VP sings VP) S)), constructing both token sequences and their syntactic structures in a top-down, left-to-right manner. The action sequence $\mathbf{a}$ comprises three types of operations:

- (X: Generate a non-terminal symbol (X, where X represents a phrasal tag such as NP;
- w: Generate a terminal symbol w, where w represents a token such as bird;

¹Oh and Schuler (2022) showed that regression models for predicting reading times fail to converge with vanilla AE.

²Oh and Schuler (2022) explored NAE calculation using various attention weight formulations, but in this study, we adopt the norm-based attention weight formulation (Kobayashi et al., 2020), which achieved the highest predictive power on the self-paced reading time corpus.

³While Oh and Schuler (2022) also proposed other metrics based on distances between attention weights at consecutive time steps, we exclusively adopt NAE because (i) in TG, the number of preceding elements varies with time, making distance definition non-trivial, and (ii) Oh and Schuler (2022) demonstrated that NAE’s predictive power subsumes that of distance-based metrics in the self-paced reading time corpus.

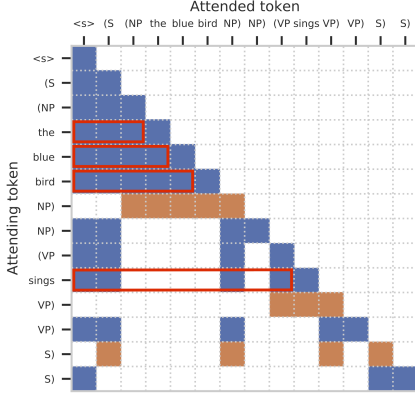


Figure 1: TG’s attention mask with **COMPOSE/STACK** attention mechanisms, adapted from Sartran et al. (2022). **COMPOSE** generates a vector representation of the closed phrase, while subsequent **STACK** operations reference this single vector as the representation of the closed phrase. Red boxes indicate the attention weights used to calculate NAE for each word.

- X): Generate X) to close the most recent opened non-terminal symbol, where X matches the phrasal tag of the targeted non-terminal symbol.

The probability of action sequence $\mathbf{a} = (a_1, a_2, \dots, a_n)$ is decomposed using the chain rule. Formally, TG is defined as:

$$p(\mathbf{x}, \mathbf{y}) = p(\mathbf{a}) = \sum_{t=1}^n p(a_t | \mathbf{a}_{<t}). \quad (2)$$

TG’s key innovation lies in its handling of closed phrases: immediately after generating X), it computes a vector representation of the closed phrase, which subsequent next-action predictions use as the representation for that phrase. Technically, this operation is realized via two components: X) action duplication and a specialized attention mask. The duplication process transforms $\mathbf{a}$ into $\mathbf{a}'$ by duplicating all X) actions (e.g., (S (NP The blue bird NP) NP) (VP sings VP) VP) S) S)), while preserving the modeling space $p(\mathbf{a})$ by preventing predictions for duplicated positions. The attention mask implements two distinct attention mechanisms: **COMPOSE** and **STACK** (Figure 1). **COMPOSE** operates exclusively at the first occurrence of each X) to generate the phrasal representation by attending only to vectors between the corresponding (X and X) (without making predictions). **STACK** operates at all other positions to compute representations for next-action prediction, with attention

restricted to positions on the *stack* (comprising unclosed non-terminals, not-composed terminals, and closed phrases).

Previous research has demonstrated that TG’s probability estimates align more closely with human offline grammaticality judgments (Sartran et al., 2022) and online brain activity than vanilla Transformers (Wolfman et al., 2024). This study investigates whether the attention mechanism of TG, which uniquely operates on syntactic structures as representational units, can serve as a cognitive model of human memory retrieval.

3 Methods

3.1 NAE calculation with TG

The calculation of NAE with TG requires assumptions regarding two key perspectives:

1. What syntactic structures should be assumed for a given token sequence?
2. How should the cognitive load of attention from non-lexical tokens (i.e., (X and X)) be attributed to words?

In response to these considerations, we make the following assumptions:

- 1-A. We assume only the globally correct syntactic structure (i.e., “perfect oracle”; Brennan, 2016).
- 2-A. We consider only attention from words, excluding attention from non-lexical tokens.

The adoption of 1-A. is motivated by two factors. First, the self-paced reading time corpus we utilized here provides gold-standard syntactic structures for each sentence, and previous studies have developed predictors based on these annotations (Shain et al., 2020; Isono, 2024). Using the same structural assumptions enables fair comparison with these established predictors, considering the possibility of parsing errors. Second, TG’s current implementation lacks beam search procedure (Stern et al., 2017; Crabbé et al., 2019), an inference technique commonly used in cognitive modeling to handle local ambiguities through parallel parsing (Hale et al., 2018; Sugimoto et al., 2024).⁴

⁴As a proof of concept, we also conducted experiments using multiple syntactic structures generated by word-synchronous beam search with Recurrent Neural Network Grammar (Dyer et al., 2016; Kuncoro et al., 2017; Noji and Oseki, 2021), obtaining similar results (Appendix D).

Regarding 2-A., given the multiple possible approaches to attributing processing load from non-lexical tokens to words, we adopt the most straightforward and theoretically neutral approach. Figure 1 denotes the attention weights used to calculate NAE for each word, with red boxes.

3.2 Settings

Language models We used 16-layer, 8-head TG and Transformer (252M parameters).⁵ All hyperparameters followed the default settings described in Sartran et al. (2022) (see Appendix A). Following Oh and Schuler (2022), we computed NAE separately for each attention head at the topmost layers and then summed the values across heads.

Training data We used BLLIP-LG, a dataset containing 42M tokens (1.8M sentences) from the Brown Laboratory for Linguistic Information Processing (BLLIP) 1987–89 WSJ Corpus Release 1 (Charniak et al., 2000).⁶ The corpus was re-parsed using a state-of-the-art constituency parser (Kitaev and Klein, 2018) and split into train-val-test sets by Hu et al. (2020). BLLIP-LG has been widely used for training syntactic LMs, including TG. Following Sartran et al. (2022), we trained a 32K SentencePiece tokenizer (Kudo and Richardson, 2018) on the training set and segmented each sentence into subword units.

Both TG and Transformer were trained at the sentence level: TG maximized the joint probability $p(\mathbf{x}, \mathbf{y})$ on action sequences, while Transformer maximized the probability $p(\mathbf{x})$ on terminal subword sequences. For training hyperparameters, we largely followed the default settings in Sartran et al. (2022) but adjusted the batch size to fit within the memory constraints of our hardware (NVIDIA A100, 40GB). Accordingly, we tuned other hyperparameters (e.g., learning rate) to maintain training stability. We trained three models with different random seeds and selected the checkpoint with the lowest validation loss for each run.

Reading time data We used the Natural Stories corpus (Futrell et al., 2018)⁷ consisting of 10 stories (485 sentences, 10,245 words) with self-paced reading times collected from 181 anonymized native English speakers. Following Futrell et al.’s preprocessing, data points were removed if (i) a

participant scored less than 5/6 on comprehension questions for a story or (ii) individual reading times were less than 100 ms or greater than 3,000 ms. Following Oh and Schuler (2022), we also excluded sentence-initial and sentence-final data points. We further removed sentence-second data points, as they lack a log trigram frequency of the previous token required for our baseline regression model. After preprocessing, 724,883 data points from 180 participants remained for statistical analysis, out of the original 848,747 data points.

Statistical analysis We evaluated each LM’s NAE contribution to reading time prediction by measuring improvements in regression model fit when adding NAE as predictors. For each model (TG/Transformer), we included both the current word’s NAE (tg_nae/tf_nae) and the previous word’s NAE (tg_nae_so/tf_nae_so) to account for spillover effects.⁸ Model improvement was quantified as the increase in log-likelihood (ΔLogLik). This evaluation was conducted for each random seed, and we report the mean ΔLogLik with standard deviation.

The baseline regression model included standard predictors from the Natural Stories corpus:

- zone and position: word position in the story and sentence;
- wordlen: number of characters in the word;
- unigram, bigram, and trigram: log-transformed n-gram frequencies.

We additionally included the following predictors:

- tg_surp and tf_surp: surprisal from TG and Transformer;
- stack_count: number of elements in the *stack* (comprising unclosed non-terminals, not-composed terminals, and closed phrases).

Following Oh and Schuler (2022), we included surprisal to test NAE’s significance in the presence of surprisal predictors from the same LMs.¹⁰ Stack count was included to isolate the cost of holding elements (Joshi, 1990; Abney and Johnson, 1991; Resnik, 1992) from their interference effects, which TG’s NAE was designed to capture.

⁸_so indicates spillover.

⁹Following Oh and Schuler (2022), we summed the subword NAE values for each word.

¹⁰For an experiment on the predictive power of surprisal itself, see Appendix E.

⁵https://github.com/google-deepmind/transformer_grammars

⁶<https://catalog.ldc.upenn.edu/LDC2000T43>

⁷<https://github.com/languageMIT/naturalstories>

Model	ΔLogLik ($\uparrow$)	Predictor	Effect size [ms]	p -value range	Significant seeds
TG	96.3 (± 3.0)	tg_nae	2.91 (± 0.1)	<0.001	3/3
		tg_nae_so	0.655 (± 0.1)	<0.05	3/3
Transformer	35.2 (± 7.0)	tf_nae	1.77 (± 0.2)	<0.001	3/3
		tf_nae_so	0.293 (± 0.2)	0.04–0.38	1/3

Table 1: TG’s and Transformer’s NAE contribution to reading time prediction (ΔLogLik). The effect size per standard deviation is shown for each model-derived predictor, along with the p -value range across random seeds and the number of seeds showing significant contributions. Standard deviations across seeds for ΔLogLik and effect sizes are shown in parentheses. The mean reading time in the analysis is 335 ms.

All predictors were z -transformed, and we also included the previous word’s values as predictors to model spillover, except for the positional information. The baseline regression model was a linear mixed-effects model (Baayen et al., 2008) with these fixed effects and by-subjects and by-story random intercepts:

$$\begin{aligned} \log(\text{RT}) \sim & \text{zone} + \text{position} + \text{wordlen} + \\ & \text{unigram} + \text{bigram} + \text{trigram} + \\ & \text{tf_surp} + \text{tg_surp} + \\ & \text{stack_count} + \text{wordlen_so} + \\ & \text{unigram_so} + \text{bigram_so} + \\ & \text{trigram_so} + \text{tf_surp_so} + \\ & \text{tg_surp_so} + \text{stack_count_so} + \\ & (1 | \text{participant}) + (1 | \text{story}) \quad (3) \end{aligned}$$

To assess each LM’s independent contribution to reading time prediction, we also conducted likelihood ratio tests by extending Equation 3 in two ways: adding both LMs’ NAE versus adding only one LM’s NAE. Note that a larger ΔLogLik from one LM does not necessarily indicate that it contributes above and beyond the other LM, nor does a smaller ΔLogLik indicate no unique contribution. Following Aurnhammer and Frank (2019), we used NAE and surprisal values averaged across random seeds for these nested model comparisons.

4 Results

4.1 Does TG’s NAE have predictive power for reading times?

Table 1 presents the contributions of TG’s and Transformer’s NAE to reading time prediction. First, Transformer’s NAE exhibited significant predictive power for reading times, independent of baseline predictors such as surprisal. The effect size was in the expected positive direction (higher

NAE values corresponding to longer reading times), primarily showing the immediate effect. This corroborates the arguments of Ryu and Lewis and Oh and Schuler that the attention mechanism—the weighted reference of preceding tokens—functions as a cognitive model of human memory retrieval, despite its engineering-oriented origins.

Second, TG’s NAE exhibited robust predictive power, demonstrating significant positive effects in both immediate and spillover contexts. This finding not only provides additional evidence for incremental construction of syntactic structures in human sentence processing (e.g., Fossum and Levy, 2012), but also suggests that TG’s attention mechanism effectively models memory retrieval from these constructed syntactic representations.

Finally, TG’s NAE made a substantially stronger contribution to reading time prediction ($\Delta\text{LogLik}=96.3$) compared to Transformer’s NAE ($\Delta\text{LogLik}=35.2$). This finding suggests that retrieval from syntactic memory representations plays a more dominant role in human sentence processing than retrieval from lexical memory representations. This underscores the importance of incorporating syntactic structures as a unit of memory representation, which we implemented through the integration of TG and NAE here.

4.2 Do TG’s and Transformer’s NAE have independent contributions?

Figure 2 presents the results of likelihood ratio tests examining the independence of TG’s and Transformer’s NAE contributions. The regression model incorporating NAE from both LMs (‘TG & Transformer’) demonstrated significantly higher predictive power than the models containing NAE from either LM alone (‘TG’ or ‘Transformer’). This reveals that TG’s NAE certainly captures variance in reading times that Transformer’s NAE cannot explain, while Transformer’s NAE, despite its lower

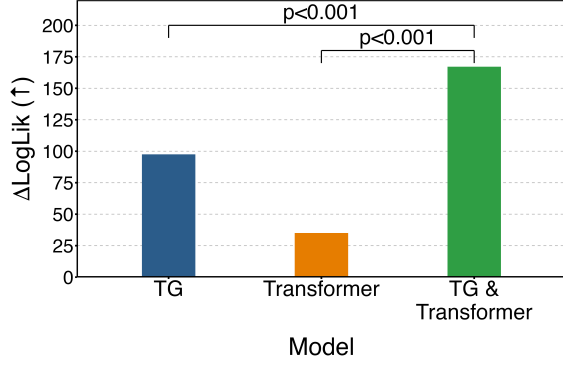


Figure 2: Likelihood ratio test results examining the independence of NAE’s predictive power

overall predictive power, accounts for unique variance not captured by TG’s NAE. This finding aligns with psycholinguistic literature, where cognitive models of memory retrieval encompass both syntax-based approaches (e.g., verb-argument relationships; [Lewis and Vasishth, 2005](#)) and semantic-based approaches (e.g., bag-of-words-like similarity; [Brouwer et al., 2012](#)), suggesting that the attention mechanisms of TG and Transformer serve as complementary cognitive models, each capturing distinct aspects of human memory retrieval.

4.3 What aspect of memory retrieval do TG’s and Transformer’s NAE capture?

To investigate the aspects of human memory retrieval captured by TG’s and Transformer’s NAE, we analyzed differences in prediction improvement across part-of-speech (POS) tags annotated in the Natural Stories corpus. Our analysis followed three steps: (i) selecting POS tags with more than 1,000 occurrences, (ii) for each POS tag, testing the significance of improvement from the baseline regression model (measured in Δ Root Mean Squared Error, Δ RMSE) when adding NAE of the current and previous word as fixed effects,¹¹ and (iii) examining the significance of *differences* in Δ RMSE between TG and Transformer for POS tags where either model showed significant improvement. We assessed significance using Wilcoxon signed-rank tests with Bonferroni correction ($p < 0.05$).

Figure 3 presents the differences in prediction improvement across POS tags.¹² Consistent with the larger Δ LogLik value, TG’s NAE

¹¹We used the same regression models as in Section 4.2, where surprisal and NAE values were averaged across seeds.

¹²For a complete list of POS tags in the Natural Stories corpus, see Appendix C.

Model	Δ LogLik	Predictor	p -value
TG	96.1 (± 15.9)	*_nae	*** (3/3)
		*_nae_so	** (3/3)
TG _{comp}	86.3 (± 31.8)	*_nae	*** (3/3)
		*_nae_so	<i>n.s.</i> (0/3)

Table 2: TG’s and TG_{comp}’s contribution to reading time prediction. The rightmost column shows the p -value range across random seeds (*** $p < 0.001$, ** $p < 0.01$, and *n.s.* not significant), along with the number of seeds showing significant contributions. Due to the potential multicollinearity between the Transformer’s NAE and TG/TG_{comp}’s NAE, the column of the effect size is omitted.

demonstrated advantages over Transformer’s NAE across a broader range of POS tags. Notably, TG’s NAE exhibited superior improvement across verbs (VB, VBD, VBG, VBP), while Transformer’s NAE excelled only for possessive pronouns (\$PRP). These findings indicate that different types of retrieval operations—verb-triggered retrieval (e.g., subject retrieval) and possessive pronoun-triggered retrieval (e.g., antecedent retrieval)—are better modeled by distinct cognitive mechanisms: attention with syntactic and token memory representations, respectively. This pattern supports our earlier argument regarding the complementary nature of these models (Section 4.2).

5 Follow-up analysis

5.1 Do TG’s advantages stem from the COMPOSE attention?

As described in Section 2.2, TG’s key feature is the COMPOSE attention, which explicitly generates single vector representations for closed phrases. Here, we investigate whether TG’s predictive power derives from merely considering syntactic structures or from explicitly treating closed phrases as single representations (see [Hale et al., 2018](#); [Brennan et al., 2020](#)). To address this question, we developed TG_{comp}, a TG variant that processes each action in the action sequence a as an individual token (i.e., [Choe and Charniak’s](#) approach), without the COMPOSE attention. We trained TG_{comp} with identical hyperparameters as TG. The baseline regression model (Equation 3) was augmented with (i) TG_{comp}’s surprisal and (ii) Transformer’s NAE to (i) ensure a fair comparison between TG and TG_{comp} and (ii) distinguish

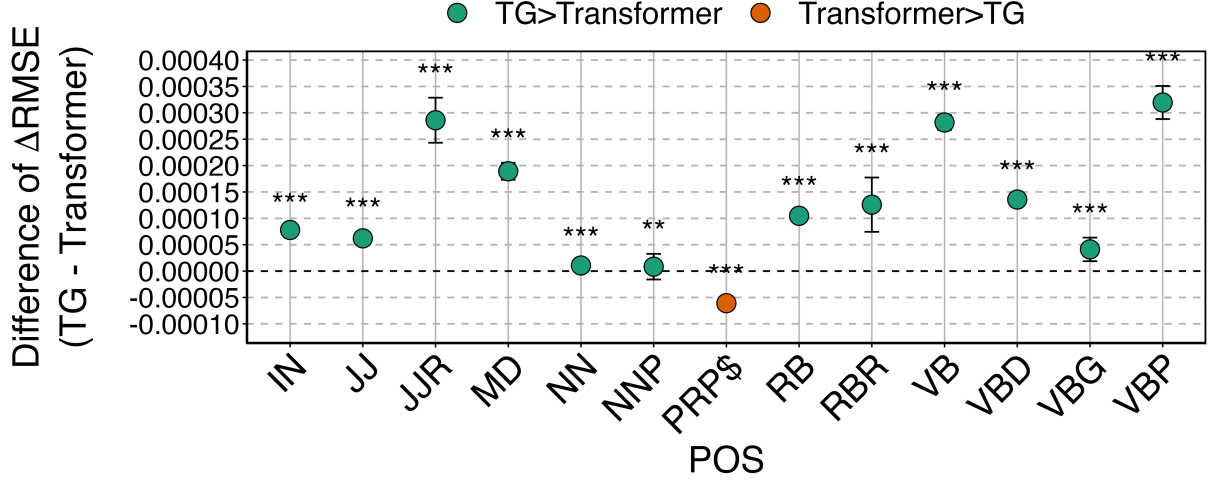


Figure 3: Differences in reading time prediction improvement (ΔRMSE) between TG and Transformer across POS tags (TG - Transformer). The y-axis shows the mean differences per word, with the error bars representing standard errors. Only POS tags showing significant improvement in either model and significant differences between models are displayed. Statistical significance after Bonferroni correction: ** $p < 0.01$, *** $p < 0.001$.

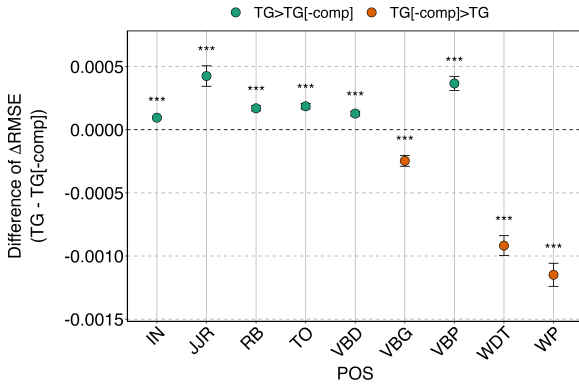


Figure 4: Differences in ΔRMSE between TG and TG_{comp} across POS tags (TG - TG_{comp})

between the effects of direct terminal token access and syntactic structure consideration in TG_{comp} .

Table 2 presents the ΔLogLik values obtained when incorporating either TG’s or TG_{comp} ’s NAE as fixed effects into the baseline regression model. Note that due to the potential multicollinearity between Transformer’s NAE and TG/ TG_{comp} ’s NAE, we focus on the ΔLogLik values and significance of the contribution rather than individual effect sizes. Our analysis reveals two key findings. First, TG_{comp} ’s NAE demonstrates significant predictive power for reading times, even in the presence of Transformer’s NAE, implying that consideration of syntactic structures alone captures certain memory retrievals based on syntactic information. Second, although we should note that the standard deviation across seeds is relatively large, TG’s

NAE outperforms TG_{comp} ’s, suggesting that the attention mechanism operating on syntactic memory representations more effectively captures variance in syntax-based memory retrieval. Likelihood ratio tests further revealed that TG’s NAE captured reading time patterns unexplainable by TG_{comp} ($\text{‘TG \& TG}_{\text{comp}}’ > \text{‘TG}_{\text{comp}}’$, $p < 0.001$). However, we also found that TG_{comp} , despite its lower overall predictive power, also explained unique variance ($\text{‘TG \& TG}_{\text{comp}}’ > \text{‘TG’}$, $p < 0.001$).

To investigate the underlying mechanisms of this complementary relationship, we analyzed the ΔRMSE differences across POS tags (Figure 4). The results showed that TG’s NAE demonstrated advantages over TG_{comp} ’s NAE across most POS tags, with notable exceptions for VBG, VBP, and WP—POS tags that typically involve immediate modification of the previously composed phrase (e.g., (NP (NP ... NP) (SBAR (WHNP (WP who). This pattern leads us to hypothesize that memory retrieval based on syntactic cues typically operates on composed representations for efficiency, but in contexts where immediate modification of a composed phrase occurs (e.g., relative clauses), humans may strategically retrieve individual lexical elements within the composed phrases.

5.2 Does TG’s NAE capture interference effects?

Psycholinguistic research has identified two primary types of memory retrieval costs: *interference* effects, which NAE aims to capture, and *decay*

effects—the cognitive load associated with accessing elements at greater linear distances (e.g., [Gibson, 1998, 2000](#)). Here, we examine whether TG’s NAE genuinely captures interference effects by testing its independence from variables that model memory decay effects. For modeling decay effects, we employed Category Locality Theory (CLT; [Isono, 2024](#)),¹³ which treats phrases in syntactic structure¹⁴ as representational units of memory and quantifies decay effects using the distance (measured in content words) between an input and the phrases to be composed with it.

To assess independence, we tested whether TG’s NAE and CLT maintain their contributions when simultaneously included in the baseline regression model (Equation 3), and examined their independence through likelihood ratio tests.¹⁵ The results (Table 3) show that TG’s NAE exhibited significant effects in both immediate and spillover conditions, and CLT demonstrated a significant immediate effect (with a marginally significant spillover effect). A nested model comparison confirmed that these effects were independent (‘TG & CLT’ > ‘CLT’, $p < 0.001$; ‘TG & CLT’ > ‘TG’, $p < 0.05$).

These results provide empirical evidence that NAE quantifies interference rather than decay in memory retrieval—extending beyond previous studies on NAE ([Ryu and Lewis, 2021](#); [Oh and Schuler, 2022](#)). This finding is significant because, as far as we are aware, while psycholinguistics has developed various implementations of memory decay effects, it has lacked broad-coverage implementations of interference effects applicable to naturally occurring texts. Our results suggest that NAE represents a promising approach for quantifying interference effects in a broad-coverage manner.

6 Level of description

In cognitive modeling studies based on surprisal theory, explanations typically follow the form “if these LMs were models of human prediction, the difficulty of next-word disambiguation that humans solve would be approximated as follows.” Such ex-

¹³Although Dependency Locality Theory (DLT; [Gibson, 1998, 2000](#)) is widely recognized as one of the most prominent models for capturing decay effects, we opted for CLT in this study, following [Isono](#)’s finding that DLT-based predictors fail to achieve statistical significance in explaining reading times in the Natural Stories corpus.

¹⁴CLT assumes syntactic structure based on Combinatory Categorical Grammar ([Steedman, 2000](#)).

¹⁵As in other likelihood ratio tests, we used surprisal and NAE values averaged across random seeds.

Model	Predictor	Effect size [ms]
TG & CLT	tg_nae	2.93***
	tg_nae_so	0.69**
	clt	0.32*
	clt_so	0.25·

Table 3: Effect sizes per standard deviation are shown for TG’s NAE and CLT predictors. Significance levels: *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$.

planations typically operate at the most abstract of [Marr](#)’s three levels of description—the computational level. Recently, [Futrell et al. \(2020\)](#) proposed lossy-context surprisal to integrate memory representation perspectives into surprisal theory. However, as they explicitly stated, this theory remains at the computational level, relaxing assumptions about memory representations in human predictive processing. In contrast, cognitive models of human memory, such as cue-based retrieval, generally provide explanations about *mechanisms* that deal with specific mental representations. These explanations typically move down one level to the algorithmic level of description. While not explicitly stated by the authors, we argue that the work of [Ryu and Lewis](#) and [Oh and Schuler](#)—conceptualizing attention mechanisms as implementations of cue-based retrieval—also operates at the algorithmic level, similar to cue-based retrieval itself. Our research can be characterized as an investigation into the nature of memory representations, addressing fundamental questions at this level (see [Hale, 2014](#)).

7 Conclusion

In this paper, we have demonstrated that attention can serve as the general algorithm for modeling human memory retrieval from two representational systems. Furthermore, we have shown that among the LMs examined in this study (TG, TG_{comp}, and Transformer), TG—whose attention mechanism uniquely operates on syntactic structures as representational units—best captures dominant factors in human sentence processing. Our results suggest that the integration of attention mechanisms (developed in NLP) with syntactic structures (theorized in linguistics) constitutes a broad-coverage candidate implementation for human memory retrieval. We hope these findings will foster greater collaboration between these two fields.

Limitations

Our NAE calculation comprised three steps: (i) computing NAE for each attention head in the top-most layers, (ii) adding the values across heads, and (iii) summing subword-level values into word level. While this procedure strictly adhered to Oh and Schuler (2022), alternative approaches to handling layers, attention heads (Ryu and Lewis, 2021), and subword tokens (Oh and Schuler, 2024; Giulianelli et al., 2024) warrant investigation.

While our study provides an in-depth investigation using the Natural Stories corpus—an English self-paced reading time dataset—the breadth of our analysis has certain limitations. The generalizability of our findings to different languages (e.g., Japanese self-paced reading time corpus from Asahara, 2022) and other cognitive load (e.g., gaze duration from Kennedy et al., 2003 or EEG and fMRI from Bhattasali et al., 2020) remains to be investigated.

As discussed in Section 3.1, we employed “perfect oracles” as syntactic structures behind token sequences. This idealization leaves the resolution of local ambiguities, which humans encounter during actual sentence processing, outside the scope of our study (for a conceptual case study, see Appendix D). By incorporating these kinds of entirely new factors, more detailed models could emerge.

Following prior work (Wolfman et al., 2024), we adopted the default TG implementation of a top-down parsing strategy. However, psycholinguistic literature has suggested that a left-corner parsing strategy might be more plausible for human sentence processing (Abney and Johnson, 1991; Resnik, 1992). While previous studies have primarily evaluated the plausibility of parsing strategies from a memory capacity perspective (cf. `stack_count`), TG’s NAE might offer a new opportunity to revisit the question of cognitively plausible parsing strategies from the perspective of memory interference.

Finally, while this paper focused on investigating the attention mechanism of TG through the lens of memory-based theory, exploring TG as an integrated implementation for expectation-based theory (via surprisal) and memory-based theory (via NAE) represents a promising future direction (Michaelov et al., 2021; Ryu and Lewis, 2022). Specifically, future work could investigate the attention mechanism of TG as the underlying driver of surprisal’s predictive power (Appendix E), ana-

lyzing the relationship between surprisal and NAE.

Ethical considerations

We employed AI-based tools (Claude, ChatGPT, GitHub Copilot, and Grammarly) for writing and coding assistance. These tools were used in compliance with the ACL Policy on the Use of AI Writing Assistance.

References

- Steven P. Abney and Mark Johnson. 1991. *Memory requirements and local ambiguities of parsing strategies*. *Journal of Psycholinguistic Research*, 20(3):233–250.
- Masayuki Asahara. 2022. *Reading Time and Vocabulary Rating in the Japanese Language: Large-Scale Japanese Reading Time Data Collection Using Crowdsourcing*. In *Proceedings of the Thirteenth Language Resources and Evaluation Conference*, pages 5178–5187, Marseille, France. European Language Resources Association.
- Christoph Aurnhammer and Stefan L. Frank. 2019. *Comparing Gated and Simple Recurrent Neural Network Architectures as Models of Human Sentence Processing*. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 41(0).
- R. H. Baayen, D. J. Davidson, and D. M. Bates. 2008. *Mixed-effects modeling with crossed random effects for subjects and items*. *Journal of Memory and Language*, 59(4):390–412.
- Douglas Bates, Martin Mächler, Ben Bolker, and Steve Walker. 2015. *Fitting linear mixed-effects models using lme4*. *Journal of Statistical Software*, 67(1):1–48.
- Shohini Bhattasali, Jonathan Brennan, Wen-Ming Luh, Berta Franzluebbers, and John Hale. 2020. *The Alice Datasets: fMRI & EEG Observations of Natural Language Comprehension*. In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 120–125, Marseille, France. European Language Resources Association.
- Jonathan Brennan. 2016. *Naturalistic Sentence Comprehension in the Brain*. *Language and Linguistics Compass*, 10(7):299–313.
- Jonathan Brennan, Yuval Nir, Uri Hasson, Rafael Malach, David J. Heeger, and Liina Pykkänen. 2012. *Syntactic structure building in the anterior temporal lobe during natural story listening*. *Brain and Language*, 120(2):163–173.
- Jonathan R. Brennan, Chris Dyer, Adhiguna Kuncoro, and John T. Hale. 2020. *Localizing syntactic predictions using recurrent neural network grammars*. *Neuropsychologia*, 146:107479.

695	Harm Brouwer, Hartmut Fitz, and John Hoeks. 2012.	<i>Image, Language, Brain: Papers from the First Mind</i>	751
696	Getting real about Semantic Illusions: Rethinking the	<i>Articulation Project Symposium</i> , pages 94–126. The	752
697	functional role of the P600 in language comprehen-	MIT Press, Cambridge, MA, US.	753
698	sion. <i>Brain Research</i> , 1446:127–143.		
699	Eugene Charniak, Don Blaheta, Niyu Ge, Keith Hall,	Mario Giulianelli, Luca Malagutti, Juan Luis Gastaldi,	754
700	John Hale, and Mark Johnson. 2000. BLLIP 1987-89	Brian DuSell, Tim Vieira, and Ryan Cotterell. 2024.	755
701	WSJ Corpus Release 1 .	On the Proper Treatment of Tokenization in Psy-	756
702		cholingistics. In <i>Proceedings of the 2024 Confer-</i>	757
703	Do Kook Choe and Eugene Charniak. 2016. Parsing	<i>ence on Empirical Methods in Natural Language Pro-</i>	758
704	as Language Modeling . In <i>Proceedings of the 2016</i>	<i>cessing</i> , pages 18556–18572, Miami, Florida, USA.	759
705	<i>Conference on Empirical Methods in Natural Lan-</i>	Association for Computational Linguistics.	760
706	<i>guage Processing</i> , pages 2331–2336, Austin, Texas.		
707	Association for Computational Linguistics.	John Hale. 2001. A Probabilistic Earley Parser as a	761
708		Psycholinguistic Model . In <i>Second Meeting of the</i>	762
709	Benoit Crabbé, Murielle Fabre, and Christophe Pallier.	<i>North American Chapter of the Association for Com-</i>	763
710	2019. Variable beam search for generative neural	<i>putational Linguistics</i> .	764
711	parsing and its relevance for the analysis of neuro-		
712	imaging signal . In <i>Proceedings of the 2019 Confer-</i>	John Hale, Chris Dyer, Adhiguna Kuncoro, and	765
713	<i>ence on Empirical Methods in Natural Language Pro-</i>	Jonathan Brennan. 2018. Finding syntax in human	766
714	<i>cessing and the 9th International Joint Conference</i>	encephalography with beam search . In <i>Proceedings</i>	767
715	<i>on Natural Language Processing (EMNLP-IJCNLP)</i> ,	<i>of the 56th Annual Meeting of the Association for</i>	768
716	pages 1150–1160, Hong Kong, China. Association	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	769
717	for Computational Linguistics.	pages 2727–2736, Melbourne, Australia. Association	770
718		for Computational Linguistics.	771
719			
720	Zihang Dai, Zhilin Yang, Yiming Yang, Jaime Car-	John T. Hale. 2014. <i>Automaton Theories of Human</i>	772
721	bonell, Quoc Le, and Ruslan Salakhutdinov. 2019.	<i>Sentence Comprehension</i> . CSLI Studies in Computa-	773
722	Transformer-XL: Attentive Language Models beyond	tional Linguistics. CSLI Publications, Stanford, CA.	774
723	a Fixed-Length Context . In <i>Proceedings of the 57th</i>		
724	<i>Annual Meeting of the Association for Computational</i>	Jennifer Hu, Jon Gauthier, Peng Qian, Ethan Wilcox,	775
725	<i>Linguistics</i> , pages 2978–2988, Florence, Italy. Asso-	and Roger Levy. 2020. A Systematic Assessment of	776
726	ciation for Computational Linguistics.	Syntactic Generalization in Neural Language Mod-	777
727		els . In <i>Proceedings of the 58th Annual Meeting of</i>	778
728	Chris Dyer, Adhiguna Kuncoro, Miguel Ballesteros, and	<i>the Association for Computational Linguistics</i> , pages	779
729	Noah A. Smith. 2016. Recurrent Neural Network	1725–1744, Online. Association for Computational	780
730	Grammars . In <i>Proceedings of the 2016 Conference</i>	Linguistics.	781
731	<i>of the North American Chapter of the Association</i>		
732	<i>for Computational Linguistics: Human Language</i>	Shinnosuke Isono. 2024. Category Locality Theory: A	782
733	<i>Technologies</i> , pages 199–209, San Diego, California.	unified account of locality effects in sentence com-	783
734	Association for Computational Linguistics.	prehension . <i>Cognition</i> , 247:105766.	784
735			
736	Victoria Fossum and Roger Levy. 2012. Sequential vs.	Aravind K. Joshi. 1990. Processing crossed and nested	785
737	Hierarchical Syntactic Models of Human Incremental	dependencies: An automation perspective on the psy-	786
738	Sentence Processing . In <i>Proceedings of the 3rd</i>	cholingistic results . <i>Language and Cognitive Pro-</i>	787
739	<i>Workshop on Cognitive Modeling and Computational</i>	<i>cesses</i> , 5(1):1–27.	788
740	<i>Linguistics (CMCL 2012)</i> , pages 61–69, Montréal,		
741	Canada. Association for Computational Linguistics.	Kohei Kajikawa, Ryo Yoshida, and Yohei Oseki. 2024.	789
742		Dissociating Syntactic Operations via Composition	790
743	Richard Futrell, Edward Gibson, and Roger P. Levy.	Count . <i>Proceedings of the Annual Meeting of the</i>	791
744	2020. Lossy-Context Surprisal: An Information-	<i>Cognitive Science Society</i> , 46(0).	792
745	Theoretic Model of Memory Effects in Sentence Pro-		
746	cessing . <i>Cognitive Science</i> , 44(3):e12814.	Alan Kennedy, Robin Hill, and Joël Pynte. 2003. The	793
747		dundee corpus. In <i>Proceedings of the 12th European</i>	794
748	Richard Futrell, Edward Gibson, Harry J. Tily, Idan	<i>Conference on Eye Movement</i> .	795
749	Blank, Anastasia Vishnevetsky, Steven Piantadosi,		
750	and Evelina Fedorenko. 2018. The Natural Stories	Diederik P. Kingma and Jimmy Ba. 2015. Adam: A	796
	Corpus . In <i>Proceedings of the Eleventh International</i>	method for stochastic optimization . In <i>3rd Inter-</i>	797
	<i>Conference on Language Resources and Evaluation</i>	<i>national Conference on Learning Representations,</i>	798
	<i>(LREC 2018)</i> , Miyazaki, Japan. European Language	<i>ICLR 2015, San Diego, CA, USA, May 7-9, 2015,</i>	799
	Resources Association (ELRA).	<i>Conference Track Proceedings</i> .	800
	Edward Gibson. 1998. Linguistic complexity: Locality	Nikita Kitaev and Dan Klein. 2018. Constituency Pars-	801
	of syntactic dependencies . <i>Cognition</i> , 68(1):1–76.	ing with a Self-Attentive Encoder . In <i>Proceedings</i>	802
		<i>of the 56th Annual Meeting of the Association for</i>	803
	Edward Gibson. 2000. The dependency locality theory:	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	804
	A distance-based theory of linguistic complexity. In	pages 2676–2686, Melbourne, Australia. Association	805
		for Computational Linguistics.	806

807	Goro Kobayashi, Tatsuki Kuribayashi, Sho Yokoi, and	GPT-2 Surprisal. In <i>Proceedings of the 2022 Con-</i>	862
808	Kentaro Inui. 2020. Attention is Not Only a Weight:	<i>ference on Empirical Methods in Natural Language</i>	863
809	Analyzing Transformers with Vector Norms . In	<i>Processing</i> , pages 9324–9334, Abu Dhabi, United	864
810	<i>Proceedings of the 2020 Conference on Empirical</i>	Arab Emirates. Association for Computational Lin-	865
811	<i>Methods in Natural Language Processing (EMNLP)</i> ,	guistics.	866
812	pages 7057–7075, Online. Association for Computa-		
813	tional Linguistics.		
814	Taku Kudo and John Richardson. 2018. SentencePiece:	Byung-Doh Oh and William Schuler. 2024. Leading	867
815	A simple and language independent subword token-	Whitespaces of Language Models’ Subword Vocab-	868
816	izer and detokenizer for Neural Text Processing .	ulary Pose a Confound for Calculating Word Prob-	869
817	In <i>Proceedings of the 2018 Conference on Empirical</i>	abilities . In <i>Proceedings of the 2024 Conference on</i>	870
818	<i>Methods in Natural Language Processing: System</i>	<i>Empirical Methods in Natural Language Processing</i> ,	871
819	<i>Demonstrations</i> , pages 66–71, Brussels, Belgium.	pages 3464–3472, Miami, Florida, USA. Association	872
820	Association for Computational Linguistics.	for Computational Linguistics.	873
821	Adhiguna Kuncoro, Miguel Ballesteros, Lingpeng	R Core Team. 2024. <i>R: A Language and Environment</i>	874
822	Kong, Chris Dyer, Graham Neubig, and Noah A.	<i>for Statistical Computing</i> . Vienna, Austria.	875
823	Smith. 2017. What Do Recurrent Neural Network	Philip Resnik. 1992. Left-Corner Parsing and Psycho-	876
824	Grammars Learn About Syntax? In <i>Proceedings of</i>	logical Plausibility . In <i>COLING 1992 Volume 1: The</i>	877
825	<i>the 15th Conference of the European Chapter of the</i>	<i>14th International Conference on Computational Lin-</i>	878
826	<i>Association for Computational Linguistics: Volume</i>	<i>guistics</i> .	879
827	<i>1, Long Papers</i> , pages 1249–1258, Valencia, Spain.		
828	Association for Computational Linguistics.	Soo Hyun Ryu and Richard Lewis. 2021. Accounting	880
829	Alexandra Kuznetsova, Per B. Brockhoff, and Rune	for Agreement Phenomena in Sentence Comprehen-	881
830	H. B. Christensen. 2017. lmerTest package: Tests in	sion with Transformer Language Models: Effects of	882
831	linear mixed effects models . <i>Journal of Statistical</i>	Similarity-based Interference on Surprisal and Atten-	883
832	<i>Software</i> , 82(13):1–26.	tion . In <i>Proceedings of the Workshop on Cognitive</i>	884
833	Roger Levy. 2008. Expectation-based syntactic compre-	<i>Modeling and Computational Linguistics</i> , pages 61–	885
834	hension . <i>Cognition</i> , 106(3):1126–1177.	71, Online. Association for Computational Linguis-	886
835	Richard L. Lewis and Shravan Vasishth. 2005. An	tics.	887
836	activation-based model of sentence processing as	Soo Hyun Ryu and Richard L. Lewis. 2022. Using	888
837	skilled memory retrieval . <i>Cognitive Science</i> ,	Transformer language model to integrate surprisal,	889
838	29(3):375–419.	entropy, and working memory retrieval accounts of	890
839	David Marr. 1982. <i>Vision: A Computational Investiga-</i>	sentence processing . In <i>Proceedings of the 35th</i>	891
840	<i>tion into the Human Representation and Processing</i>	<i>Annual Conference on Human Sentence Processing</i> ,	892
841	<i>of Visual Information</i> . W. H. Freeman and Company,	Santa Cruz, CA, USA.	893
842	San Francisco.	Laurent Sartran, Samuel Barrett, Adhiguna Kuncoro,	894
843	Danny Merx and Stefan L. Frank. 2021. Human Sen-	Miloš Stanojević, Phil Blunsom, and Chris Dyer.	895
844	tence Processing: Recurrence or Attention? In <i>Pro-</i>	2022. Transformer Grammars: Augmenting Trans-	896
845	<i>ceedings of the Workshop on Cognitive Modeling</i>	former Language Models with Syntactic Inductive	897
846	<i>and Computational Linguistics</i> , pages 12–22, Online.	Biases at Scale . <i>Transactions of the Association for</i>	898
847	Association for Computational Linguistics.	<i>Computational Linguistics</i> , 10:1423–1439.	899
848	James A. Michaelov, Megan D. Bardolph, Seana Coul-	Cory Shain, Idan Asher Blank, Marten van Schijn-	900
849	son, and Benjamin Bergen. 2021. Different kinds of	del, William Schuler, and Evelina Fedorenko. 2020.	901
850	cognitive plausibility: Why are transformers better	fMRI reveals language-specific predictive coding dur-	902
851	than RNNs at predicting N400 amplitude? <i>Proceed-</i>	ing naturalistic sentence comprehension . <i>Neuropsy-</i>	903
852	<i>ings of the Annual Meeting of the Cognitive Science</i>	<i>chologia</i> , 138:107307.	904
853	<i>Society</i> , 43(43).	Mark Steedman. 2000. <i>The Syntactic Process</i> . MIT	905
854	Hiroshi Noji and Yohei Oseki. 2021. Effective Batching	Press, Cambridge, MA, USA.	906
855	for Recurrent Neural Network Grammars . In <i>Find-</i>	Mitchell Stern, Daniel Fried, and Dan Klein. 2017.	907
856	<i>ings of the Association for Computational Linguis-</i>	Effective Inference for Generative Neural Parsing .	908
857	<i>tics: ACL-IJCNLP 2021</i> , pages 4340–4352, Online.	In <i>Proceedings of the 2017 Conference on Empiri-</i>	909
858	Association for Computational Linguistics.	<i>cal Methods in Natural Language Processing</i> , pages	910
859	Byung-Doh Oh and William Schuler. 2022. Entropy-	1695–1700, Copenhagen, Denmark. Association for	911
860	and Distance-Based Predictors From GPT-2 Atten-	Computational Linguistics.	912
861	tion Patterns Predict Reading Times Over and Above	Yushi Sugimoto, Ryo Yoshida, Hyeonjeong Jeong,	913
		Masatoshi Koizumi, Jonathan R. Brennan, and Yohei	914
		Oseki. 2024. Localizing Syntactic Composition with	915
		Left-Corner Recurrent Neural Network Grammars .	916
		<i>Neurobiology of Language</i> , 5(1):201–224.	917

William Timkey and Tal Linzen. 2023. *A Language Model with Limited Memory Capacity Captures Interference in Human Sentence Processing*. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 8705–8720, Singapore. Association for Computational Linguistics.

Julie A Van Dyke and Richard L Lewis. 2003. *Distinguishing effects of structure and decay on attachment and repair: A cue-based parsing account of recovery from misanalyzed ambiguities*. *Journal of Memory and Language*, 49(3):285–316.

Julie Ann Van Dyke. 2002. *Retrieval Effects in Sentence Parsing and Interpretation*. University of Pittsburgh ETD, University of Pittsburgh.

Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. *Attention is All you Need*. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc.

Ethan G. Wilcox, Jon Gauthier, Jennifer Hu, Peng Qian, and Roger P. Levy. 2020. *On the Predictive Power of Neural Language Models for Human Real-Time Comprehension Behavior*. In *Proceedings of the Annual Meeting of the Cognitive Science Society*, volume 42.

Michael Wolfman, Donald Dunagan, Jonathan Brennan, and John Hale. 2024. *Hierarchical syntactic structure in human-like language models*. In *Proceedings of the Workshop on Cognitive Modeling and Computational Linguistics*, pages 72–80, Bangkok, Thailand. Association for Computational Linguistics.

A Hyperparameters

The hyperparameters are shown in Table 4. All model hyperparameters follow Sartran et al. (2022); Wolfman et al. (2024), while training hyperparameters were adjusted to accommodate the batch size suitable for our computational resources (NVIDIA A100, 40GB). The total computational cost required for all experiments was approximately 225 GPU hours.

B Correlations between predictors

The correlations between predictors in our statistical analysis are shown in Table 5. While the NAE from different LMs shows a very high correlation with each other, their predictive power for the self-paced reading times remains independent (see Section 4.2 and 5.1).¹⁶

¹⁶_mcomp indicates –comp.

C Part-of-speech tags

Table 5 presents the complete list of part-of-speech (POS) and symbol tags in the Natural Stories corpus. As reading times are annotated for each whitespace-delimited region, for data points containing symbol tags (e.g., NNP.), we used the stripped version (e.g., NNP) in our analysis. Additionally, we excluded from our analysis any data points containing multiple POS tags (e.g., NNP POS).

D Parallel parsing experiment

As a conceptual case study for the local ambiguity resolution in syntactic structures behind token sequences, we implemented TG’s NAE calculation using 10 syntactic structures obtained through word-synchronous beam search (Stern et al., 2017) with Recurrent Neural Network Grammar (RNNG; Dyer et al., 2016; Kuncoro et al., 2017; Noji and Oseki, 2021).^{17,18} NAE was computed individually for each syntactic structure and then aggregated as a weighted average:

$$\text{NAE_TG}_{l,h,i} := \frac{\sum_{t \in \text{Beam}_i} p(t) \cdot \text{NAE}_{l,h,i}^t}{\sum_{t \in \text{Beam}_i} p(t)}, \quad (4)$$

where Beam represents the set of syntactic structures synchronized at the i -th word ($|\text{Beam}| = 10$).¹⁹

The analysis revealed patterns consistent with those observed when considering only the globally correct syntactic structure: both LMs’ NAE demonstrated significant predictive power for reading times, with TG’s NAE showing stronger contributions compared to Transformer’s (Table 6).²⁰ The likelihood ratio test further confirmed independent contributions from both LMs ($p < 0.001$ for both comparisons: ‘TG & Transformer’ > ‘Transformer’ and ‘TG & Transformer’ > ‘TG’).

E Surprisal experiment

We analyzed each LM’s surprisal contribution to reading time prediction using a baseline regression model that excluded both LMs’ surprisal from Equation 3 but included their NAE (Table 7). While

¹⁷<https://github.com/aistairc/rnng-pytorch>

¹⁸RNNG was trained on BLLIP-LG using default hyperparameters. For inference, action beam size and fast track were set to 100 and 1, respectively.

¹⁹stack_count was similarly calculated as the weighted average across syntactic structures in Beam.

²⁰_bs indicates beam search.

Model architecture	Transformer-XL (Dai et al., 2019)
Vocabulary size	32,768
Model dimension	1,024
Feed-forward dimension	4,096
Number of layers	16
Number of heads	8
Segment length	256
Memory length	256
Optimizer	Adam ($\beta_1 = 0.9$, $\beta_2 = 0.999$) (Kingma and Ba, 2015)
Batch size	16
Number of training steps	400,000
Learning rate scheduler	Linear warm-up & cosine annealing
Number of warm-up steps	32,000
Initial learning rate	2.5×10^{-8}
Maximum learning rate	3.75×10^{-5}
Final learning rate	7.5×10^{-8}
Dropout rate	0.1

Table 4: Model and training hyperparameters

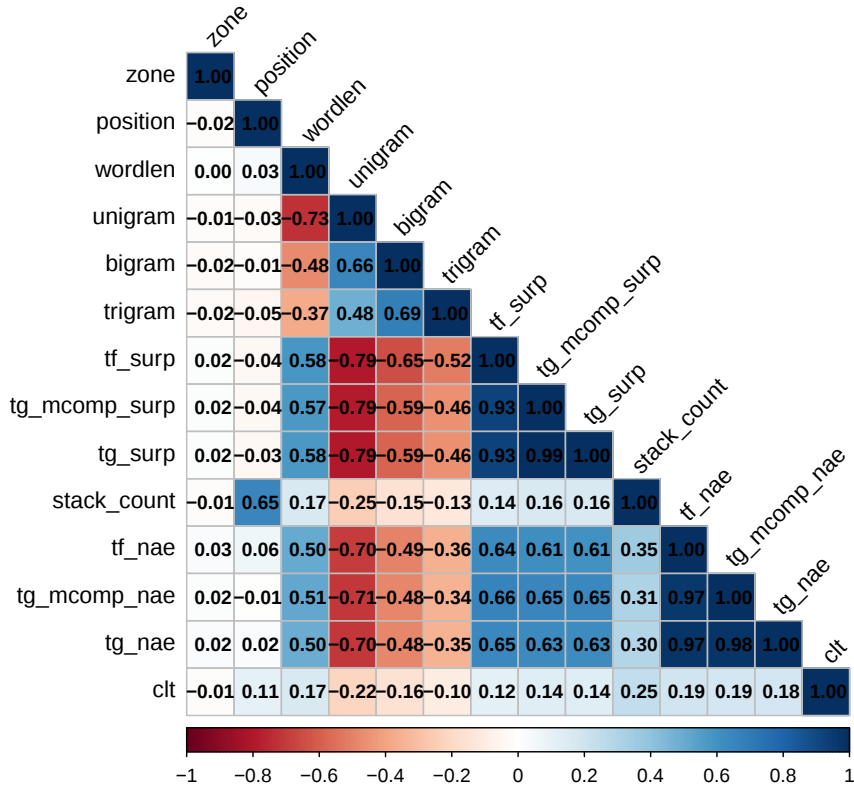


Figure 5: Correlations between predictors in our statistical analysis

CC	Coordinating conjunction	PRP\$	Possessive pronoun
CD	Cardinal number	RB	Adverb
DT	Determiner	RBR	Adverb, comparative
EX	Existential <i>there</i>	RBS	Adverb, superlative
FW	Foreign word	RP	Particle
IN	Preposition or subordinating conjunction	TO	<i>to</i>
JJ	Adjective	UH	Interjection
JJR	Adjective, comparative	VB	Verb, base form
JJS	Adjective, superlative	VBD	Verb, past tense
MD	Modal	VBG	Verb, gerund or present participle
NN	Noun, singular or mass	VBN	Verb, past participle
NNS	Noun, plural	VBP	Verb, non-3rd person singular present
NNP	Proper noun, singular	VBZ	Verb, 3rd person singular present
NNPS	Proper noun, plural	WDT	Wh-determiner
PDT	Predeterminer	WP	Wh-pronoun
POS	Possessive ending	WP\$	Possessive wh-pronoun
PRP	Personal pronoun	WRB	Wh-adverb
-LRB-	Left round bracket	,	Comma
-RRB-	Right round bracket	.	Period
“	Open double quotes	:	Colon
”	Closing double quotes		

Table 5: POS and symbol tags in the Natural Stories corpus

Model	ΔLogLik ($\uparrow$)	Predictor	Effect size [ms]	p -value range	Significant seeds
TG	76.6 (± 6.6)	tg_bs_nae	2.53 (± 0.2)	<0.001	3/3
		tg_bs_nae_so	0.849 (± 0.1)	<0.001	3/3
Transformer	28.0 (± 6.9)	tf_nae	1.58 (± 0.2)	<0.001	3/3
		tf_nae_so	0.457 (± 0.2)	0.006–0.13	1/3

Table 6: TG’s and Transformer’s NAE contribution to reading time prediction, where TG’s NAE was calculated with multiple syntactic structures generated by word-synchronous beam search with RNNG

Model	ΔLogLik ($\uparrow$)	Predictor	Effect size [ms]	p -value range	Significant seeds
TG	265 (± 11)	tg_surp	4.63 (± 0.1)	<0.001	3/3
		tg_surp_so	1.64 (± 0.1)	<0.001	3/3
Transformer	299 (± 30)	tf_surp	5.31 (± 0.2)	<0.001	3/3
		tf_surp_so	1.68 (± 0.2)	<0.001	3/3

Table 7: TG’s and Transformer’s surprisal contribution to reading time prediction

both LMs’ surprisal demonstrated significant predictive power for reading times, Transformer’s surprisal exhibited a stronger contribution compared to TG’s. Additionally, our likelihood ratio test using the averaged surprisal revealed that the regression model incorporating both LMs’ surprisal showed significantly higher predictive power compared to models with only one LM ($p < 0.001$ for both comparisons: ‘TG & Transformer’ > ‘Transformer’ and ‘TG & Transformer’ > ‘TG’). These findings suggest that (i) unlike attention mechanisms, next-word prediction based solely on token sequences more effectively captures dominant factors of human prediction processing, but (ii) similar to attention mechanisms, both types of next-word prediction—those based on token sequences alone and those leveraging both syntactic structures and token sequences—may coexist as models that capture distinct aspects of human predictive processing.

F Comparison between TG and TG_{comp} with a weak baseline regression model

In Section 5.1, we explored the advantage of treating closed phrases as single representations beyond explicit syntactic structure consideration. Our analysis incorporated Transformer’s NAE in the baseline regression model to distinguish between two effects in TG_{comp}: syntactic structure consideration and direct terminal token access.

To evaluate which model—TG or TG_{comp}—better captures more dominant factors in human sentence processing as a single model, we assessed their predictive power without Transformer’s NAE in the baseline regression model (Table 8). The analysis revealed TG’s superior predictive power ($\Delta\text{LogLik}=97.3$) compared to TG_{comp} ($\Delta\text{LogLik}=92.2$). Additionally, TG demonstrated both immediate and spillover effects, while TG_{comp} primarily showed an immediate effect. These results highlight that TG, which explicitly treats closed phrases as single representations, outperforms TG_{comp}, even when considering TG_{comp}’s advantage in direct terminal token access. Consistent with findings in Section 5.1, likelihood tests confirmed TG’s independent predictive power from TG_{comp} (‘TG & TG_{comp}’ > ‘TG_{comp}’, $p < 0.001$); TG_{comp}, despite its lower overall predictive power, accounted for unique variance (‘TG & TG_{comp}’ > ‘TG’, $p < 0.01$).

G License

Table 9 summarizes the licenses of the data and tools employed in this paper. All data and tools were used under their respective license terms.

Model	ΔLogLik ($\uparrow$)	Predictor	Effect size [ms]	p -value range	Significant seeds
TG	97.3 (± 4.0)	tg_nae	2.93 (± 0.1)	<0.001	3/3
		tg_nae_so	0.657 (± 0.1)	<0.01	3/3
TG _{-comp}	92.2 (± 10)	tg_mcomp_nae	2.88 (± 0.2)	<0.001	3/3
		tg_mcomp_nae_so	0.416 (± 0.1)	0.01–0.17	1/3

Table 8: TG’s and TG_{-comp}’s NAE contribution to reading time prediction with Transformer’s NAE excluded from the regression baseline model

Dataset/Tool	License
BLLIP (Charniak et al., 2000)	BLLIP 1987–89 WSJ Corpus Release 1
Natural Stories corpus (Futrell et al., 2018)	CC BY-NC-SA 4.0
transformer_grammar (Sartran et al., 2022)	Apache 2.0
rnng-pytorch (Noji and Oseki, 2021)	MIT License
SentencePiece (Kudo and Richardson, 2018)	Apache 2.0
R (version 4.4.2) (R Core Team, 2024)	GNU GPL ≥ 2
lme4 (version 1.1.34) (Bates et al., 2015)	GNU GPL ≥ 2
lmerTest (version 3.1.3) (Kuznetsova et al., 2017)	GNU GPL ≥ 2

Table 9: Licenses of datasets and tools