

The Role of Context Gating in Predictive Sentence Processing

Yasemin Gokcen

How do we process sentences?

During language comprehension, humans routinely use information from preceding context to implicitly predict upcoming input.

This is thought to be a core computation in language comprehension.

Surprisal Theory

I'm going to return my book to
the...

Library → expected, low
surprisal, high probability
with context



Dog → not expected, high
surprisal, lower probability
with context



Surprisal

- Measures of how unexpected a word is, given context
- $\text{Surprisal} = -\log P(\text{word} \mid \text{context})$

Behavioral Measure of Surprisal

- Surprisal estimates from language models associated with reading times
 - Smith & Levy (2008), Frank (2009), Shain et al. (2024)
 - Higher surprisal positively correlated with higher reading times/eye fixations

We have established the
importance of prediction in
sentence processing...

What mechanisms *enable*
us to predict?

Context Memory

- Need some memory for previous words to predict the next word
- Memory representation of the context is lossy
- Not every word contributes to next-word prediction equally

I'm going to return my book to the...



I'm going to **return** my **book** to **the**...



?? ?? ?? return ?? book ?? the...



Memory for Context

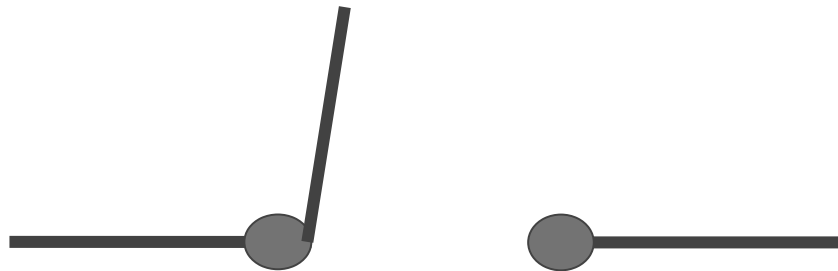
- Working memory
 - Adaptive
 - Selecting information to track
 - Based on current situation
 - Ex: driving at a busy intersection
 - Help us keep track of *context*
 - From “I’m going to return my book to the...”
 - Hold onto **return** and **book**



Working Memory and Computational Models

- Computational models of working memory (WM) mechanisms associated with prefrontal cortex (PFC) use gating
 - Capture ability to actively maintain and rapidly update contextual/task-relevant information in memory.
- Anatomy and physiology of PFC is thought to support gating

Similar gating
mechanisms may be
involved in context
memory for
sentence processing



Return

Book To The...

Research Question

To what degree are the gating mechanisms of prefrontal cortex used to maintain context for the language network?

Methodology Outline

Comparing gating vs non-gating models

- Ability to predict the next word
- Architecture difference
- Compare to WM demand measures

Language Learning Models

- Compare performance across language models
 - Trained to predict next word
 - Different architectures
 - Recurrent Neural Networks (RNN) and Long Short-Term Memory (LSTM)

LSTM has PFC-like gating mechanisms, RNN does not

- Trained on the Wikitext 2 dataset and tested on Natural Stories (Futrell et al., 2018)
- Natural Stories corpus has sentences of various linguistic properties; used in neuroimaging studies (Shain et al., 2022; Whebe et al., 2021)

Model Architectures

If you were to journey to the North of England, you would come to a valley that is surrounded by moors as high as mountains.

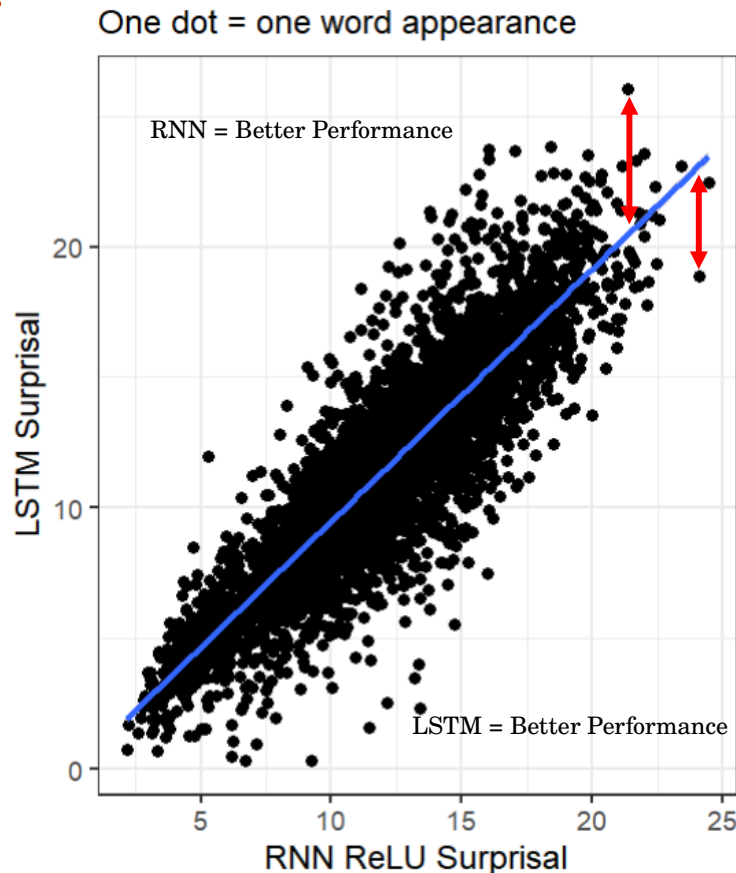


Word	LSTM	RNN(Tanh)	RNN(ReLU)
If	10.3523	8.3332	9.0970
You	4.5668	6.1161	6.7887
Were	5.1052	5.3894	4.8747
To	5.5413	5.2790	6.4274
Journey	17.3029	15.7554	14.2847

- We trained and tested 3 different model types
 - LSTM
 - RNN
 - 2 different activation functions: Tanh and ReLU
- Each model produces surprisal values for each word

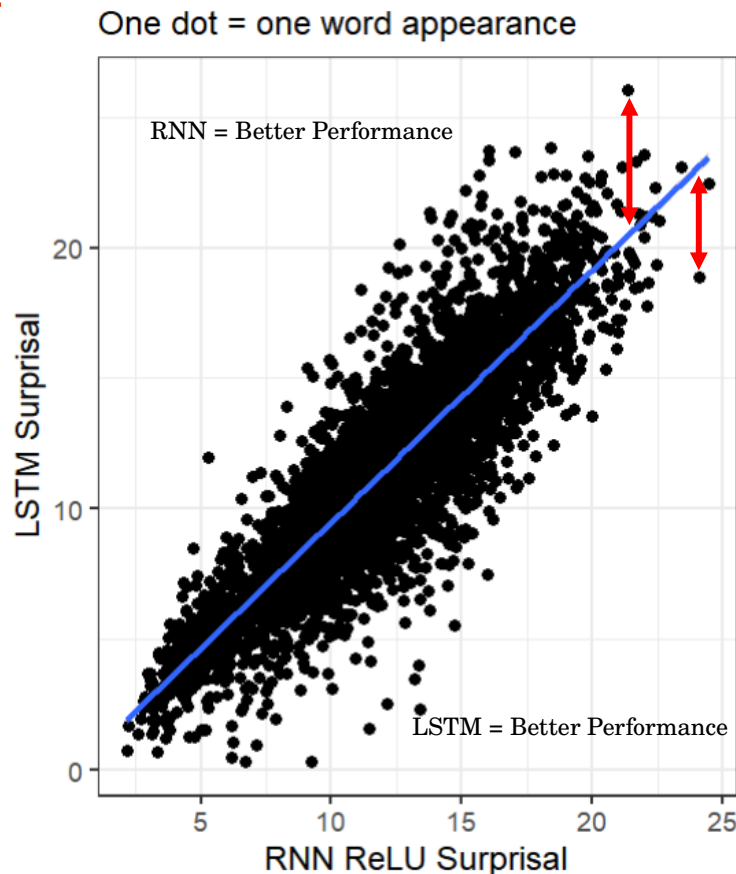
How do we compare model performance?

- Compare all of the surprisal values for each word
- Compute the *residual*
 - How well does the RNN predict the LSTM value?



How do we compare model performance?

- For each model and word we have:
 - Residual values that represent performance difference
- Negative $\rightarrow$ LSTM is performing better
 - Gating may be playing an important role



Working Memory Demand Measures

- 3 theoretical frameworks to access WM demand (Shain et al., 2022)
 - Dependency Locality Theory
 - How far words are away from each other
 - ACT-R model
 - Temporal costs
 - Memory decay
 - Left-Corner Parsing
 - e.g., How many phrases are between words?

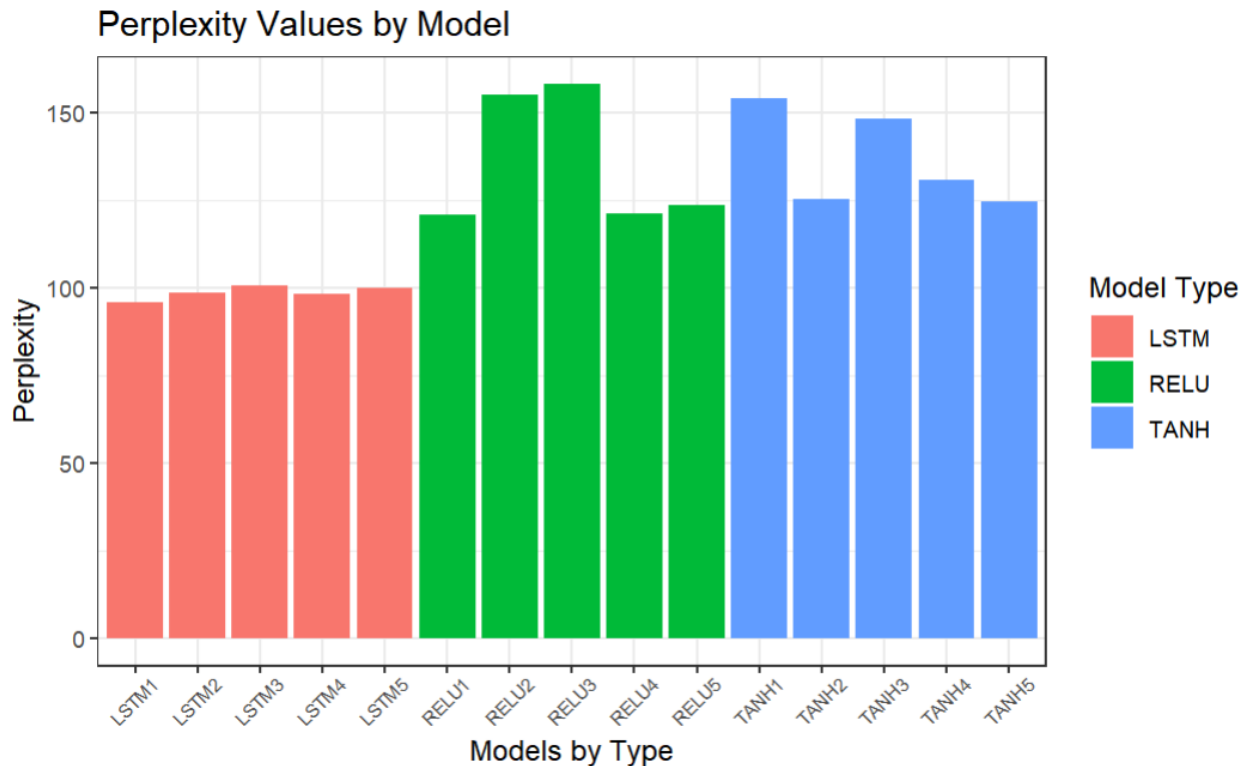
**Calculated and implemented for each word
independent of the LSTM/RNN models**

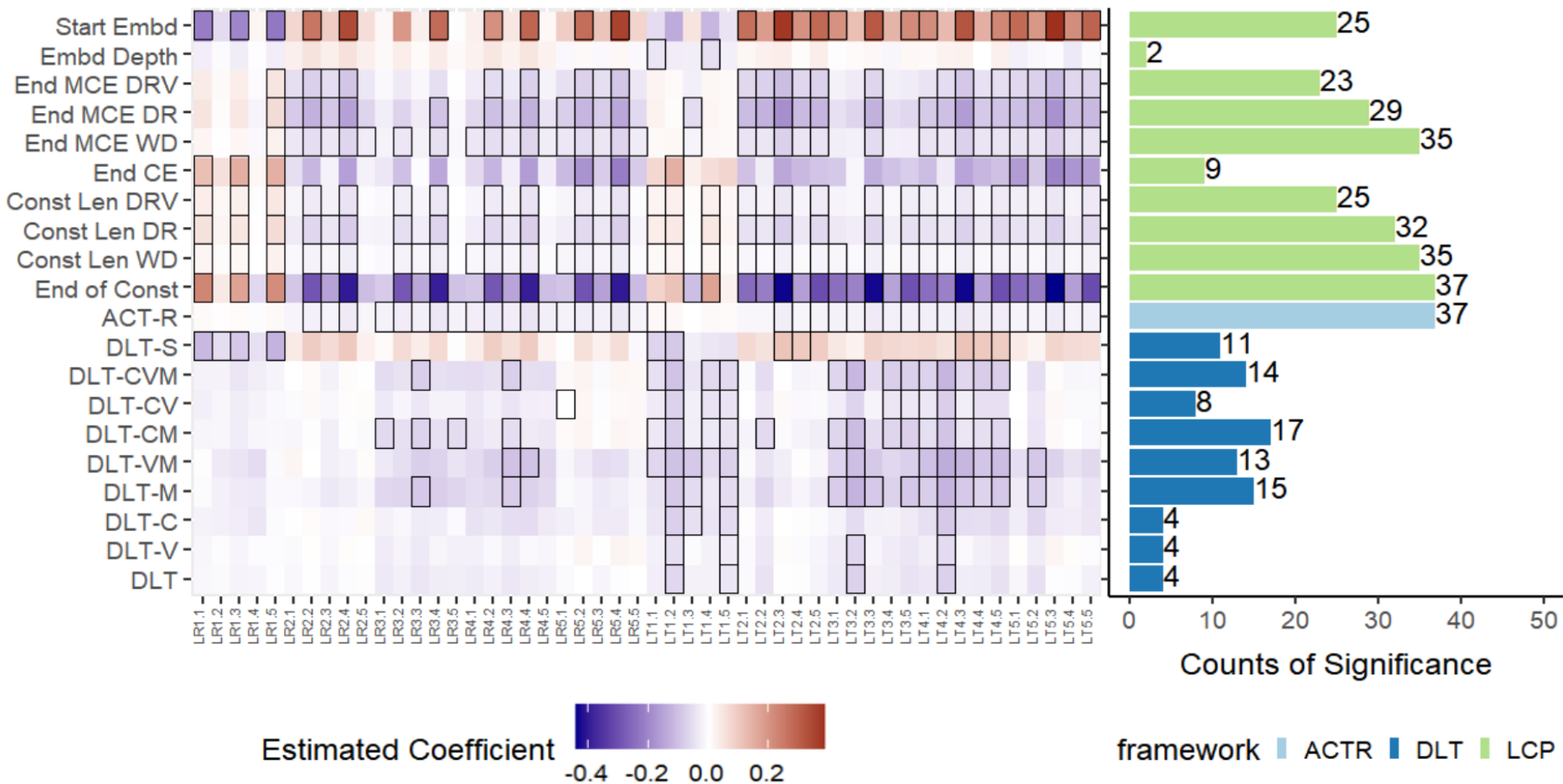
WM Measures & Residuals

- We computed 5 iterations of each model (all with different seeds)
 - 650 hidden units, 128 batch size
 - Total of 15 models
- Computed for each combination (LSTM ~ Tanh and LSTM ~ ReLU)
- LSTM₁ vs Tanh₁, LSTM₂ vs Tanh₁, LSTM₃ vs Tanh₁
 - 25 residuals for LSTM ~ Tanh + 25 residuals for LSTM ~ ReLU
- **Is gating related to WM demands?**
- Correlate Residuals with WM demand measures for each

Comparing Models by Surprisal

- All three models are performing similarly
- LSTMs generally performing better
- Slight differences across iterations





Summary

- We can use model surprisal values as estimates for human processing difficulty with processing words
- Gating appears to be important for contexts with higher WM demand
- DLT measures were rarely significantly correlated with residuals
- Constituent and many embedding measures more often significantly correlated with residuals
- LSTM does relatively better (compared to RNN) as these measures of WM demand increase
 - Regardless of activation function

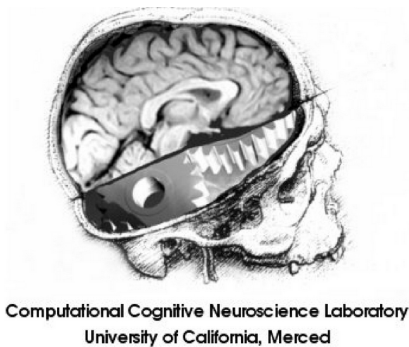
Back to our research question

- To what degree are the gating mechanisms of prefrontal cortex used to maintain context for the language network?
 - There are types of sentence structures in which gating is helpful in lowering surprisal.
 - The language network in the brain could make use of PFC-like gating when WM demand is high.

Future Directions

- Analyze changes in hidden layer activation from word-to-word
- Examine gating units for each cell in the LSTM
 - Can we identify where/when active maintenance and rapid updating occurs in a sentence?
- EEG data collection in progress! → correlate with model predictions

Thanks!



Special thanks to Dr. David Noelle and Dr. Rachel Ryskin
for mentorship and guidance throughout this project!

Further thanks to the CCN lab and LInC lab for extra
feedback and camaraderie 😊

