

Word Surprisal Correlates with Sentential Contradiction in LLMs

Ning Shi, Bradley Hauer, David Basil, John Zhang and Grzegorz Kondrak

{ning.shi,gkondrak}@ualberta.ca

Introduction

Contradiction: A premise P contradicts a hypothesis H when P entails the negation of H . (MacCartney and Manning, 2009)

$$\text{CON}(P, H) \Leftrightarrow P \models \neg H$$

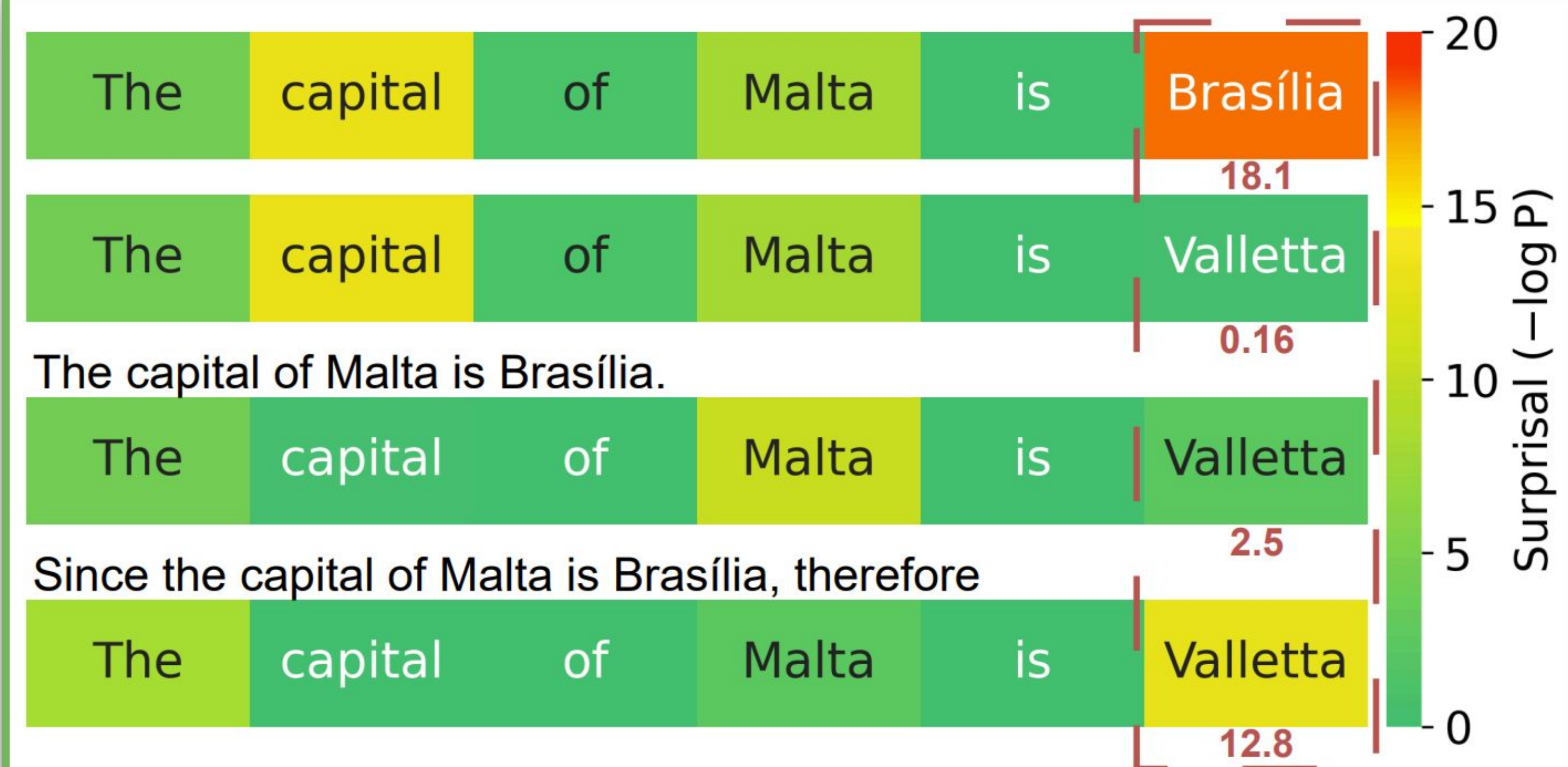
Surprisal measures the unpredictability of a word given a preceding context. (Hale, 2001)

$$\text{Surprisal}(w \mid C) = -\log P(w \mid C)$$

Open question: How does a language model respond to semantic inconsistency between two given sentences?

We posit that:

When H is conditioned on P , the surprisal magnitude of content words in H is positively correlated with the presence of a contradiction, $\text{CON}(P, H)$.



Token-to-Word Decoding

Token $\neq$ Word: “My name is ____”

$P(\text{“John”})$ includes probability mass for “Johnathan”.

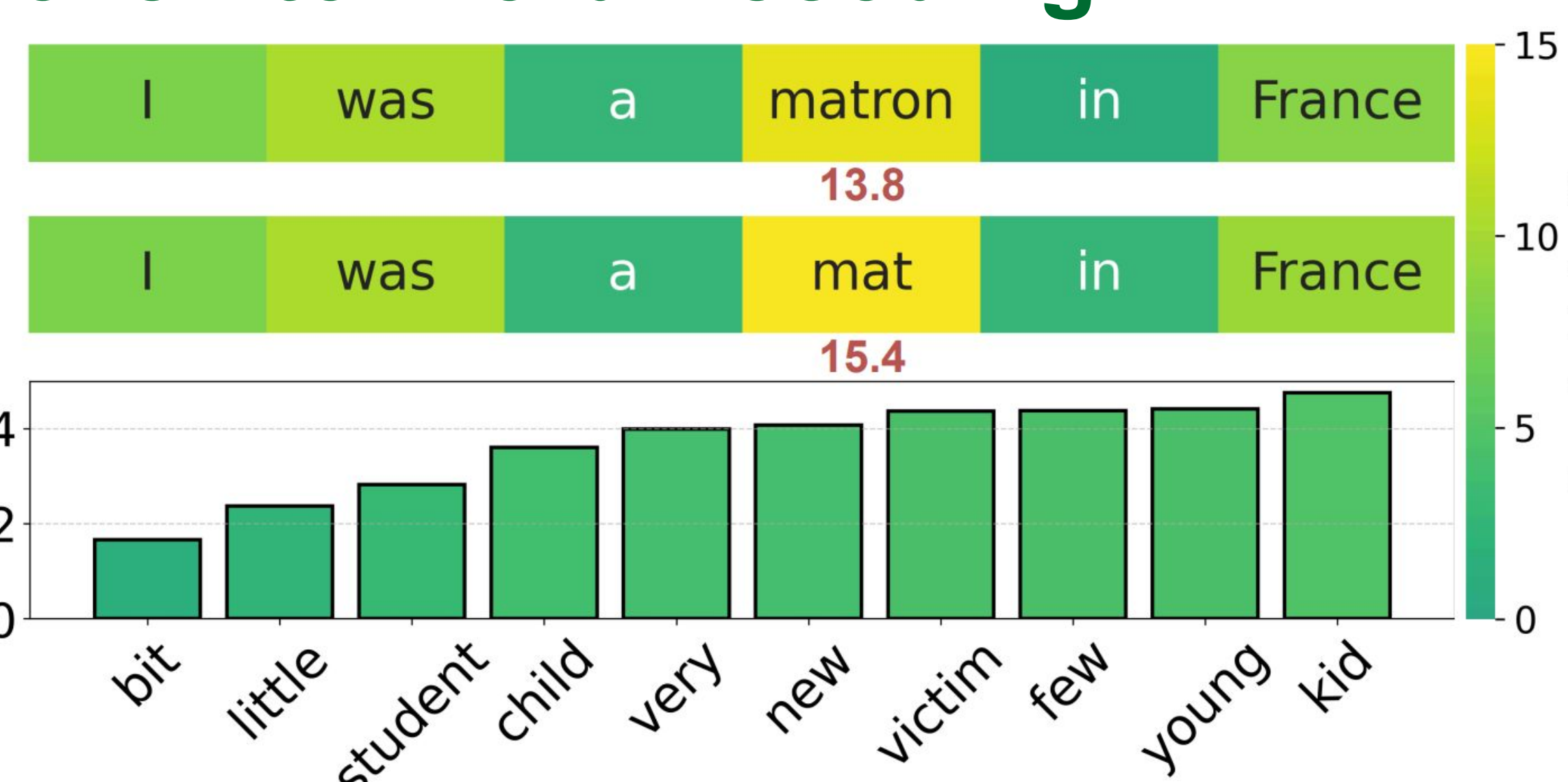
Constrained Expansion:

Only expand token sequences that form valid word prefixes.

Dynamic Normalization: At each decoding step, renormalize probability mass over the valid token subset.

InjectEOW: Insert an explicit End-of-Word (EOW)

option so completed words receive probability mass independently of longer continuations.



Algorithm 1 Token-to-Word Decoding

Input: tokenized context S^C , beam width W , beam depth D , language model $F(\cdot)$

```
1:  $B_0 \leftarrow \{\langle 0, S^C \rangle\}$ 
2: for  $d \in \{1, \dots, D\}$  :
3:    $B \leftarrow \emptyset$ 
4:   for  $(p, S) \in B_{d-1}$  :
5:     if  $S.\text{last}() = \text{EOW}$  :
6:        $B.\text{add}(\langle p, S \rangle)$ ; continue
7:    $\mathcal{P} \leftarrow F(S)$ 
8:   if  $d = 1$  :
9:      $\mathcal{P} \leftarrow \text{Normalize}(\mathcal{P}, S_{\text{bow}})$ 
10:  else:
11:     $\mathcal{P} \leftarrow \text{InjectEOW}(\mathcal{P}, S_{\text{bow}})$ 
12:     $\mathcal{P} \leftarrow \text{Normalize}(\mathcal{P}, S_{\text{mid}})$ 
13:  for  $(p', s) \in \text{Top}(\mathcal{P}, W)$  :
14:     $B.\text{add}(\langle p \cdot p', S \circ s \rangle)$ 
15:   $B_d \leftarrow B.\text{top}(W)$ 
16: return  $B_D.\text{sort}()$ 
```

Experiments

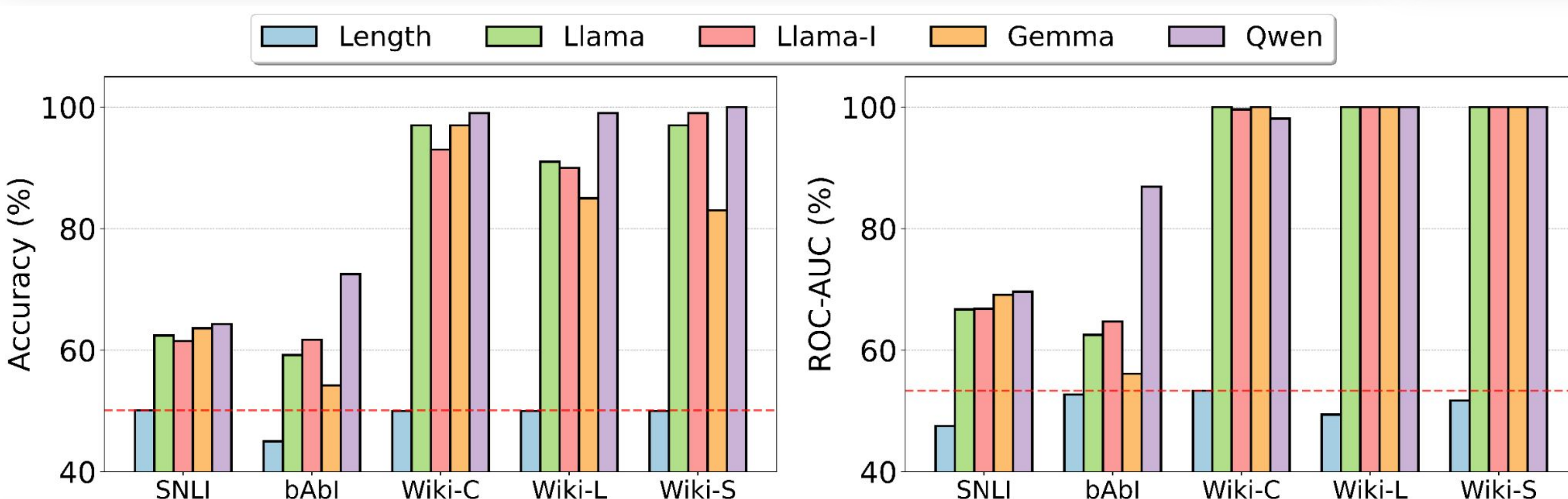
TEMP Context: “Since $\{P\}$, therefore $\{H\}$ ”

Direct Surprisal: $-\log P(w \mid C)$

Mean Aggregation: Average surprisal across hypothesis content words

Classifier: Tune a threshold on the validation set and predict contradiction when the mean surprisal exceeds.

Word-level surprisal achieves statistically significant performance in both Accuracy and ROC-AUC across datasets (e.g., SNLI) and model families (e.g., Qwen).



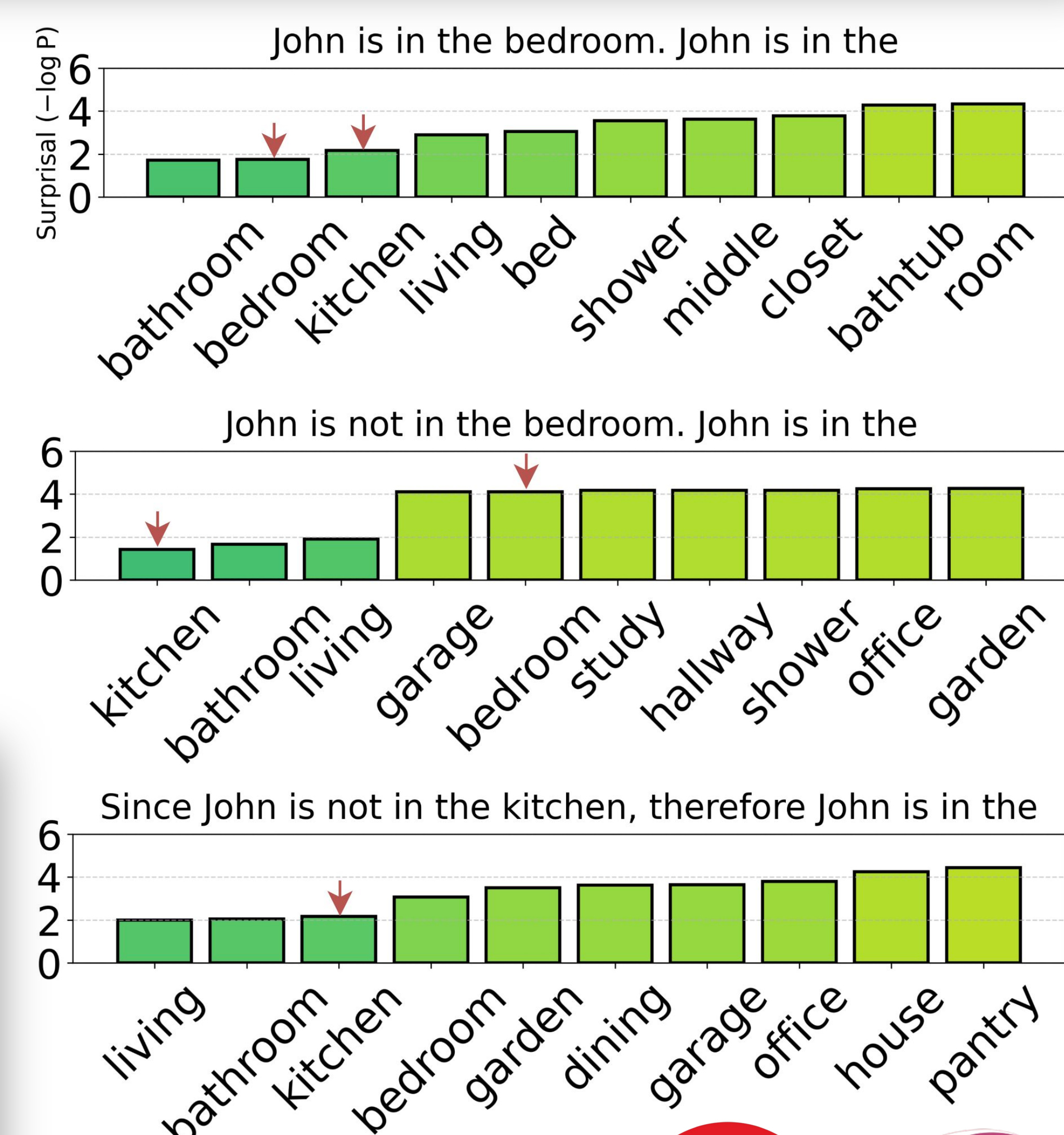
Analysis

Correlation over Causation:

The model (e.g., Llama) favors distributional correlation over causal reasoning (“kitchen” over “bedroom”).

Negation Failure:

The model also struggles with negation. It fails to raise surprise for words that should be logically inconsistent (“not in the”).



Conclusion

- Introduced **token-to-word decoding** for reliable word surprisal.
- Found a clear **positive correlation** between surprisal and contradiction.
- Demonstrated **robustness** across methods, models, and datasets.
- Revealed key model **limitations**:
 - dependency on word correlation
 - weak negation understanding



<https://github.com/ShiningLab/CON2LM>