

L²M: Mutual Information Scaling Law for Long-Context Language Modeling

Zhuo Chen^{12*}, Oriol Mayné i Comas²³, Zhuotao Jin²⁴, Di Luo¹²⁴⁵, Marin Soljačić¹²

¹The NSF AI Institute for Artificial Intelligence and Fundamental Interactions

²Massachusetts Institute of Technology

³Polytechnic University of Catalonia

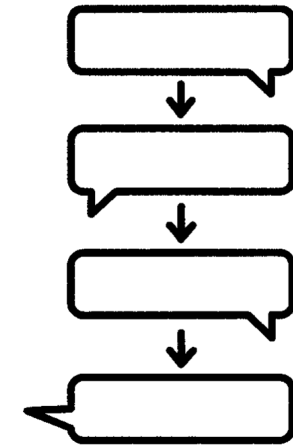
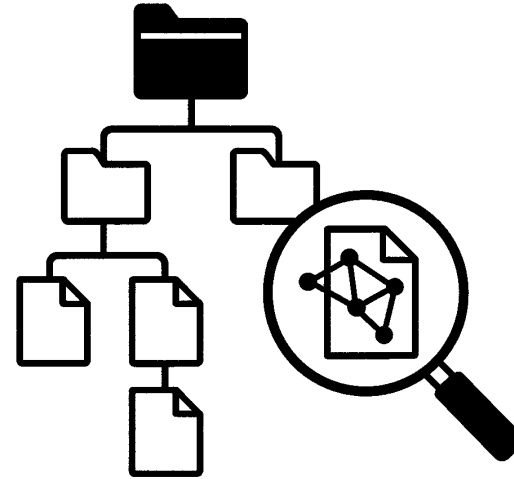
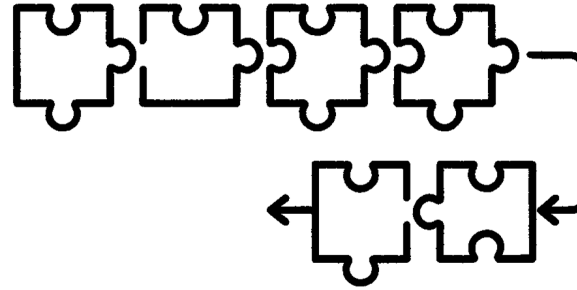
⁴Harvard University

⁵University of California, Los Angeles



Long-context Language Modeling is Crucial

- Reasoning through long chains of logic
- Maintaining consistency in long writing
- Understanding large codebases
- Keeping track of an extended conversation

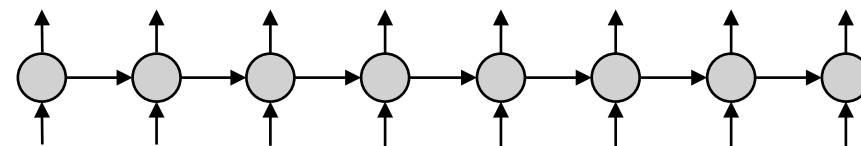
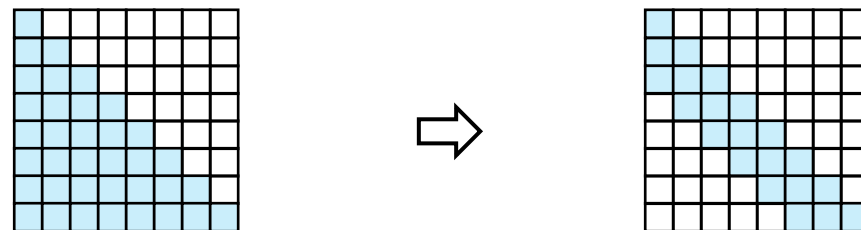
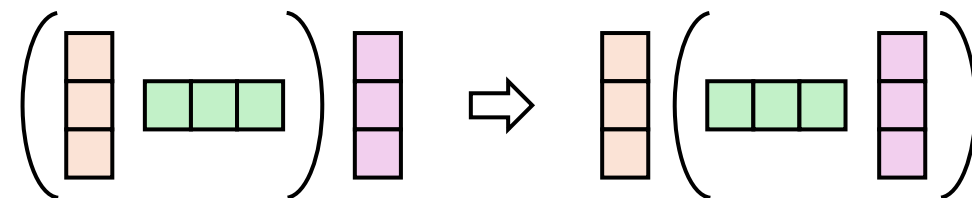


Limited Understanding of Trade-off Between Cost and Performance

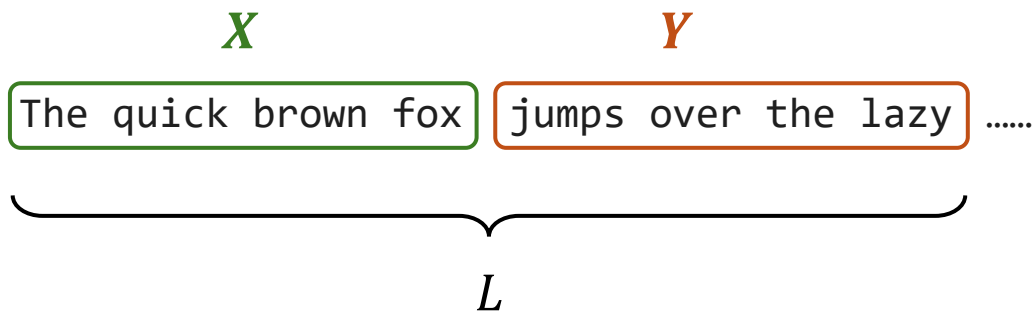
Transformers have quadratic computational cost

- Linear attention
- Sparse attention
- State space model
- Compression
- OCR
-

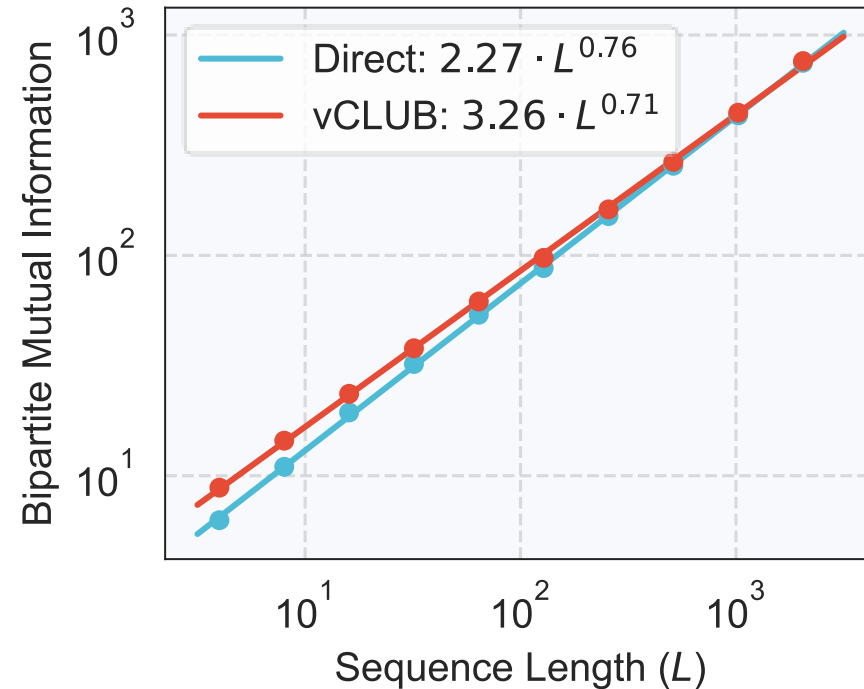
Lower cost without hurting performance?



Universal Framework Based on Bipartite Mutual Information

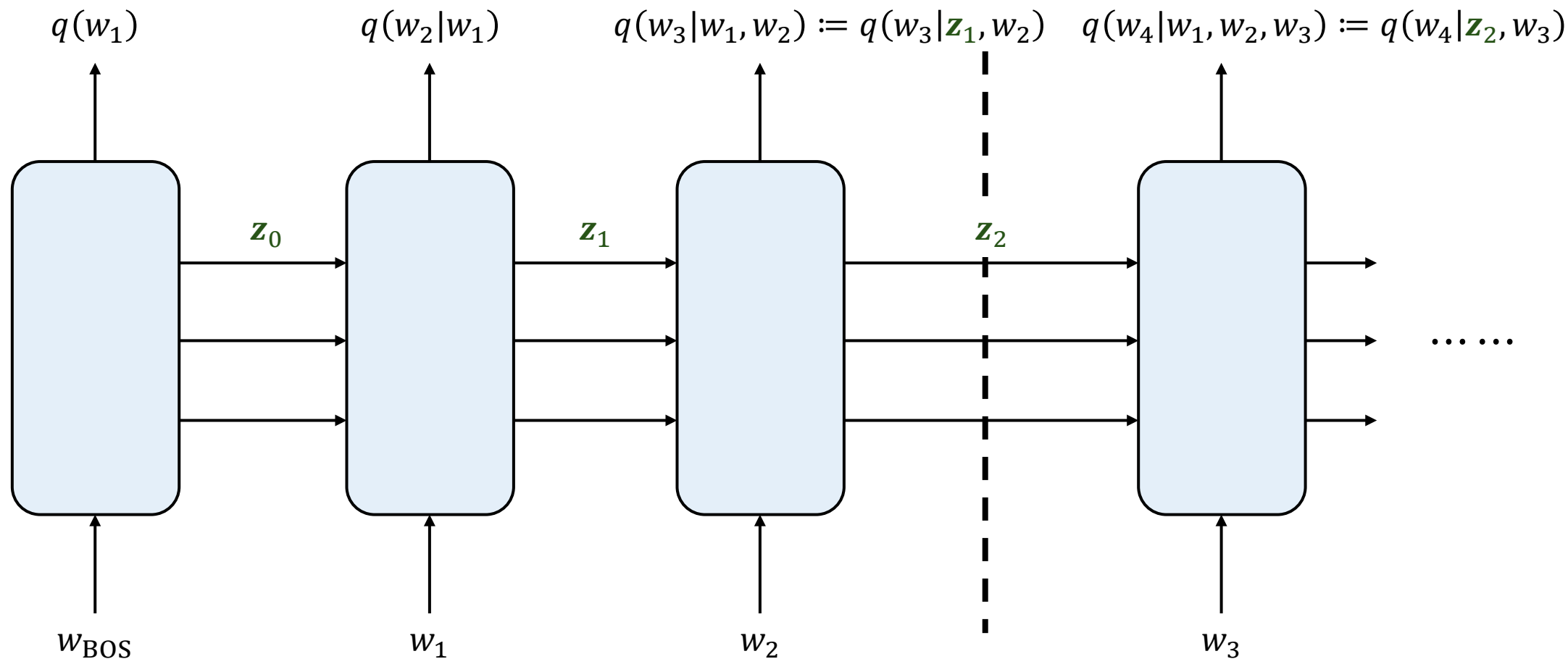


Bipartite mutual information $I(X; Y) \sim L^\beta$



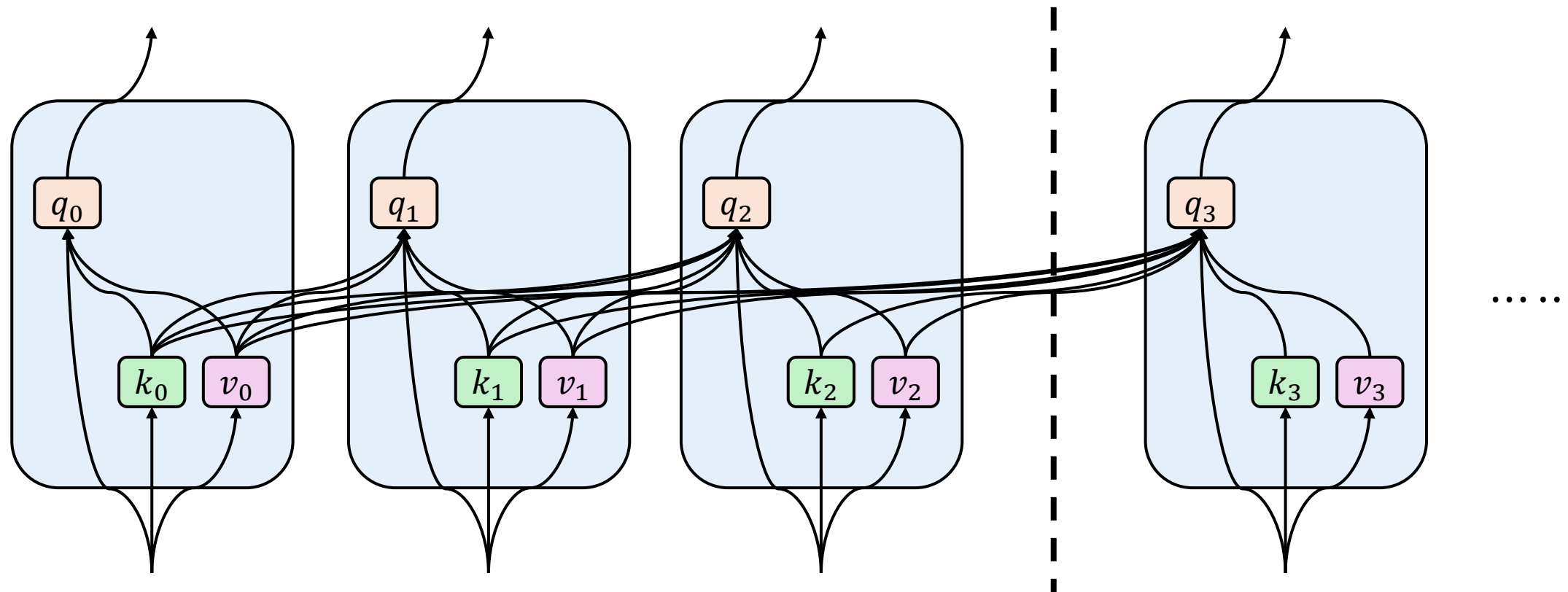
Measures the amount of information that must be retained from X to correctly predict $p(Y | X)$

History State in Large Language Models



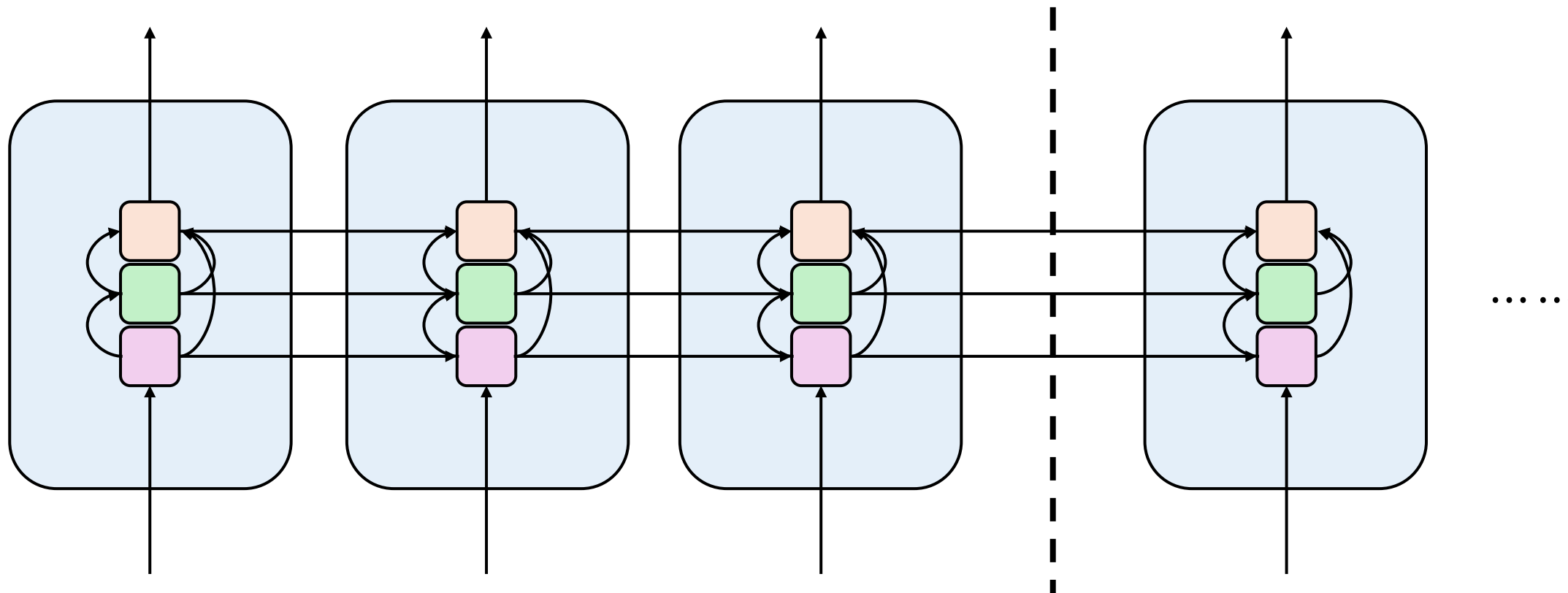
History state = the latent variables responsible for storing past information

History State in Transformers



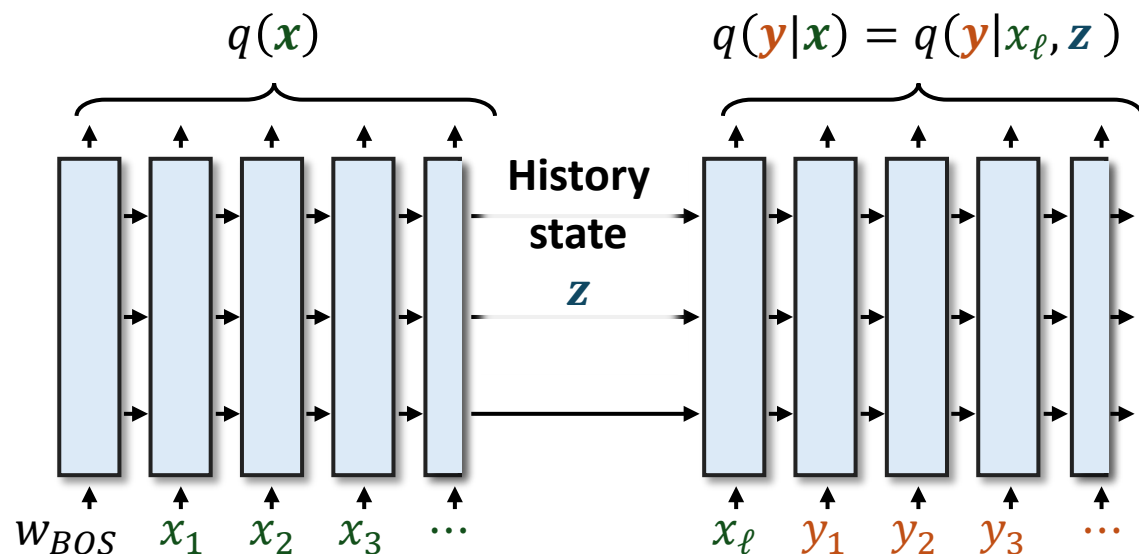
History state = all key-value pairs

History State in State Space Models

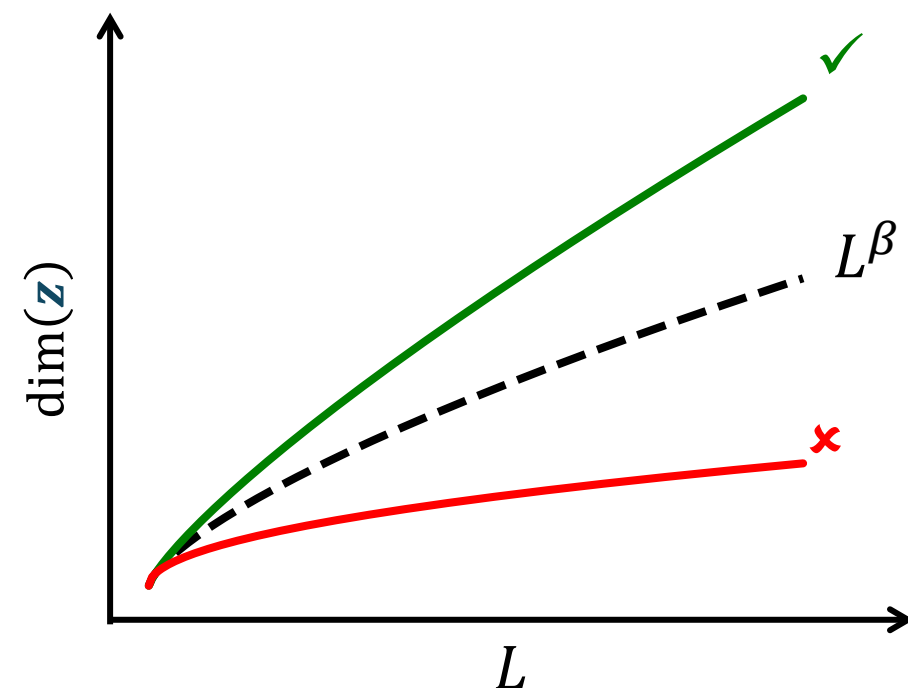


History state = recurrent state

Necessary Scaling of History State

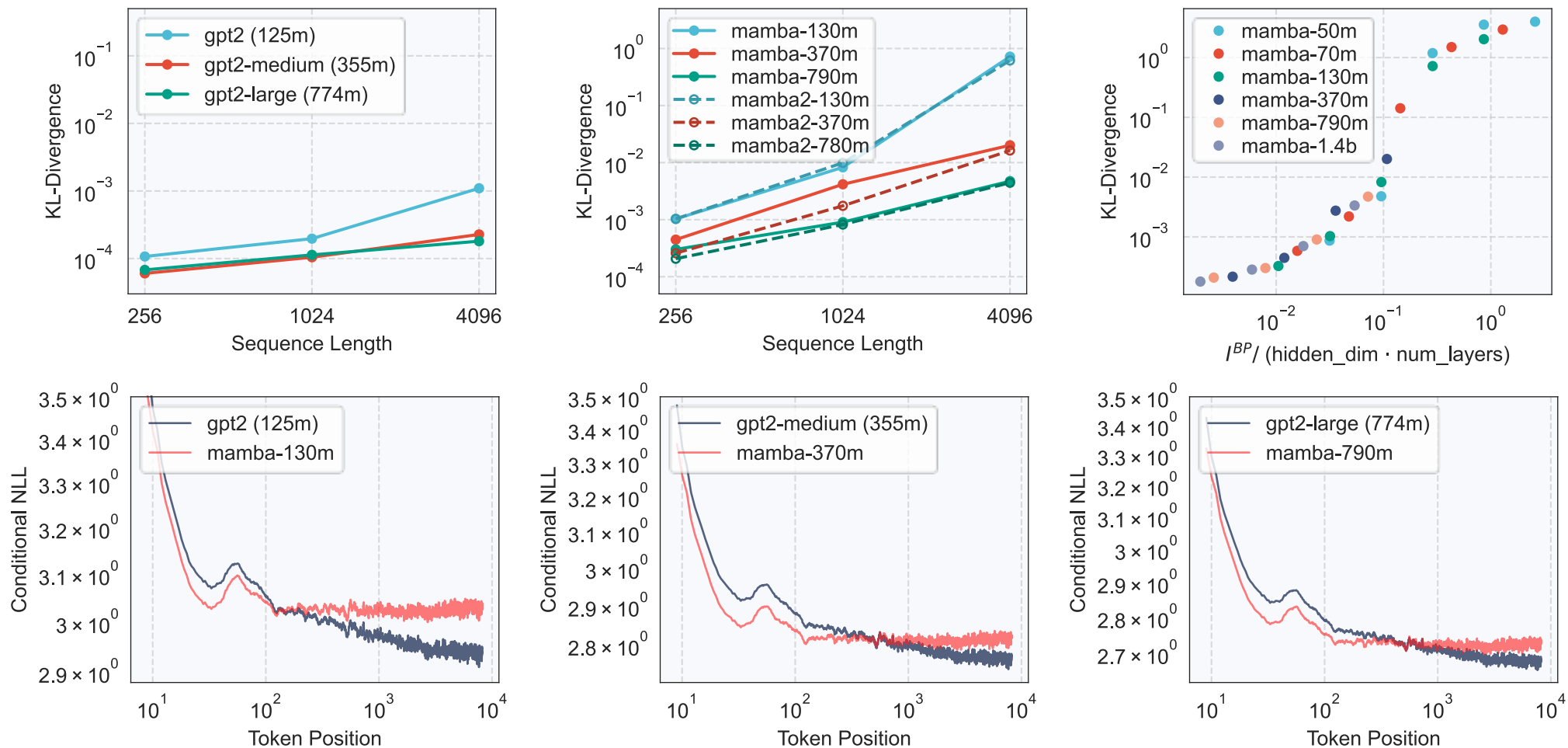


$$\max I^q(\mathbf{X}; \mathbf{Y}) \sim \dim(\mathbf{z})$$



An LLM architecture is effective for long context modeling only if $\dim(\mathbf{z}) \gtrsim L^\beta$

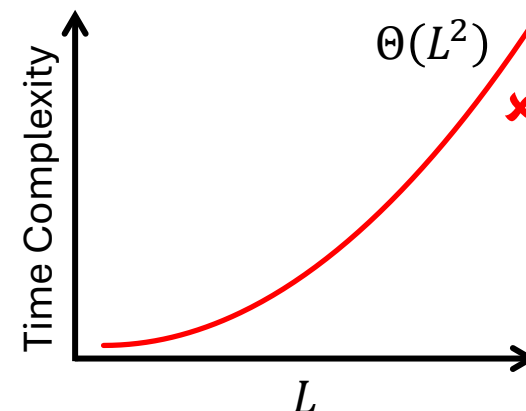
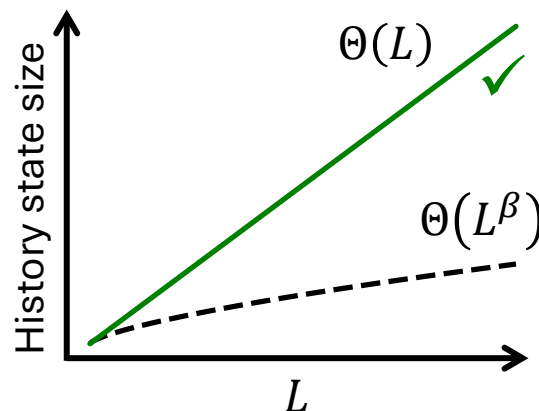
Experimental Confirmation of Our Theory



Neither Transformer nor State Space Model is Optimal

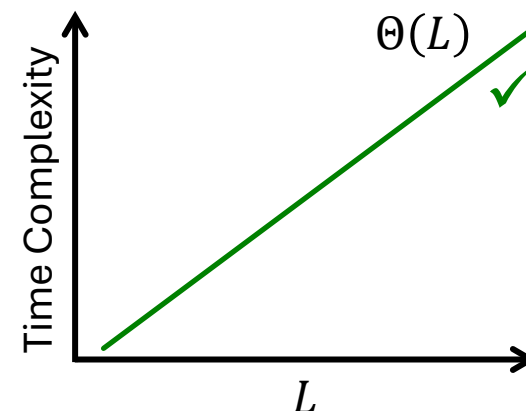
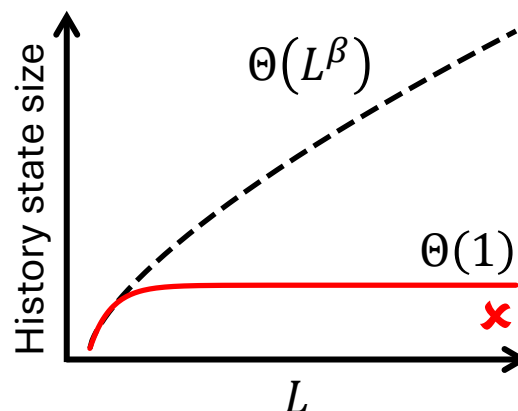
Transformers:

- ✓ History state size $\sim \Theta(L)$
- ✗ Time complexity $\sim \Theta(L^2)$



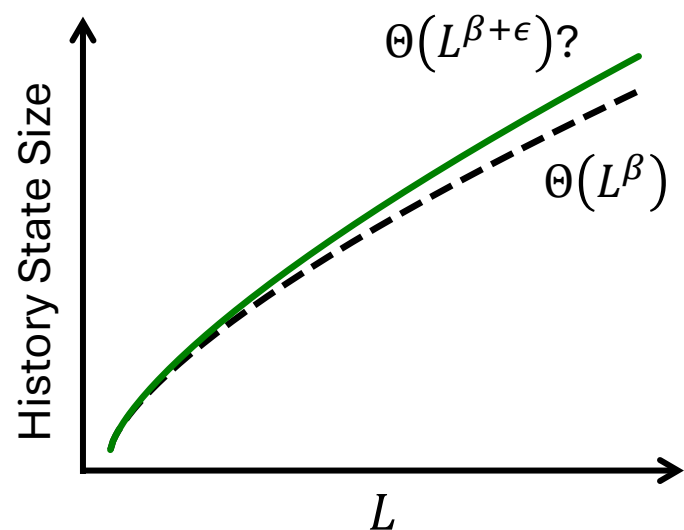
State space models:

- ✗ History state size $\sim \Theta(1)$
- ✓ Time complexity $\sim \Theta(L)$

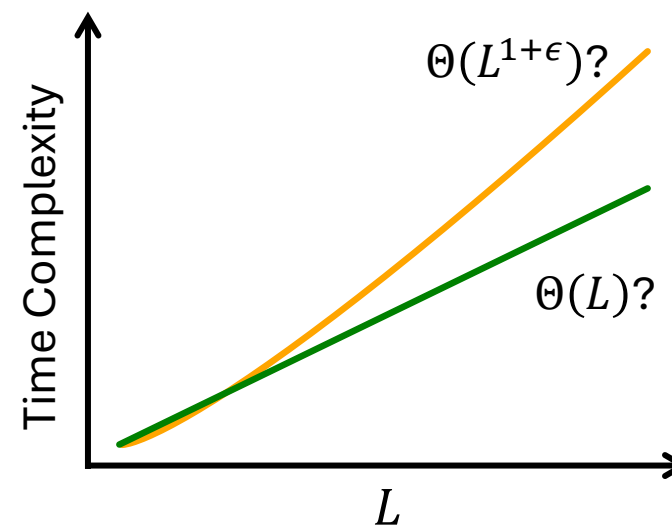


Optimal Level of Sparsity?

We don't need linearly scaling history size



and we want low time complexity



Future direction: to engineer the optimal MI-capable sparsity