

Efficient LLM Pretraining & Inference with Unlimited Context Length

Xuezhe Ma (Max)

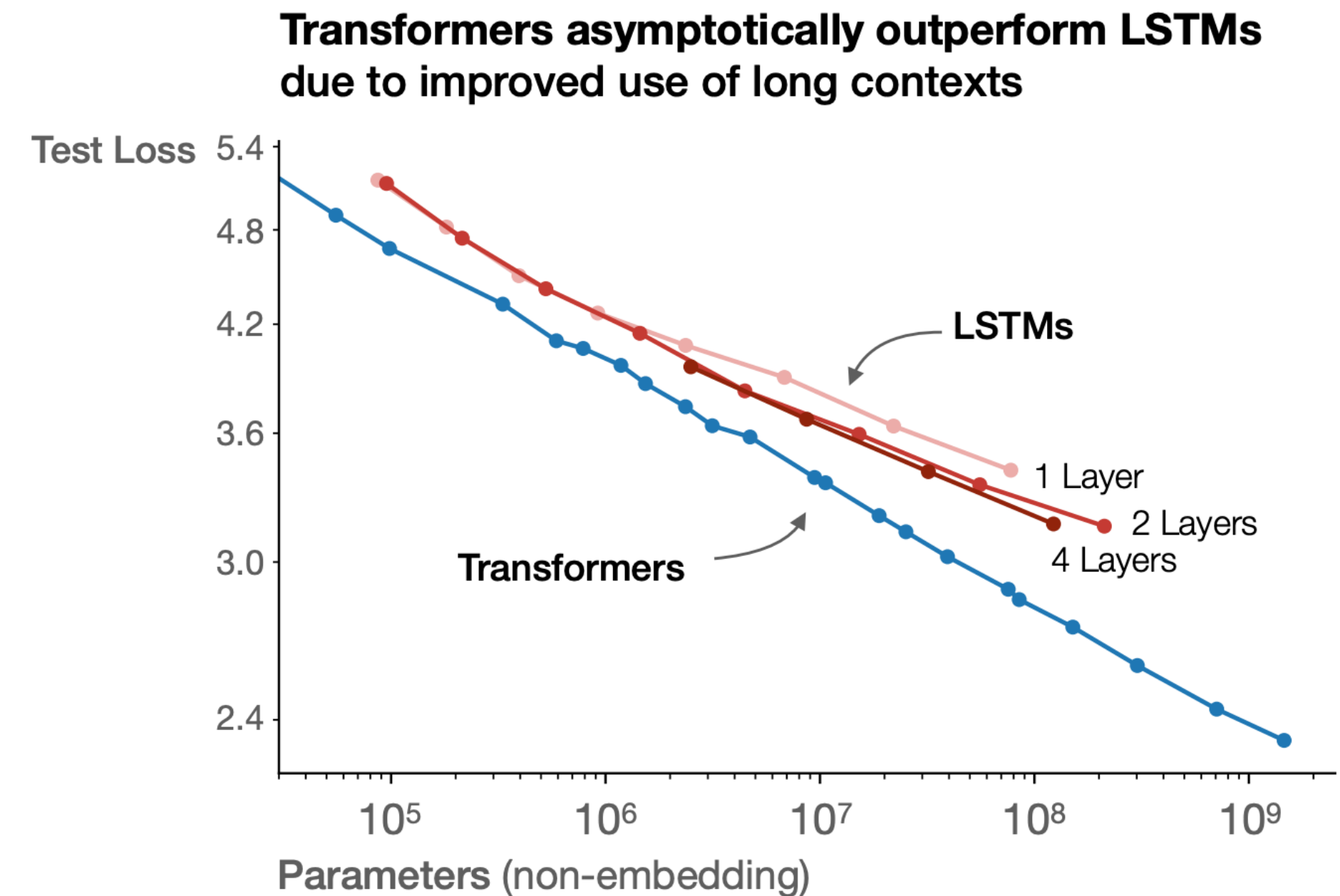
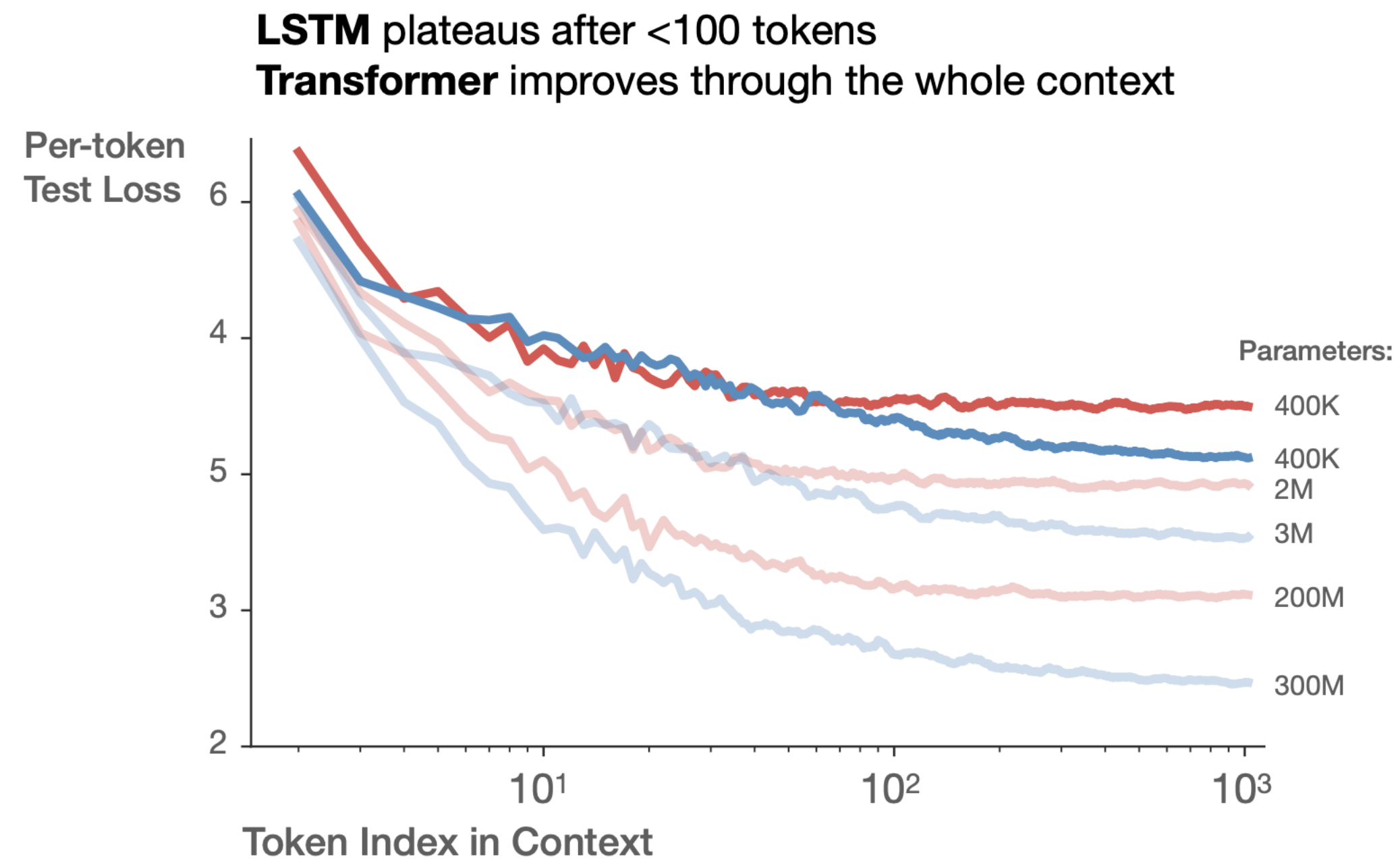
collaborators at

1. USC-ISI 2. Meta AI 3. CMU 4. UCSD 5. MIT 6. SJTU

Why Long Context?

- **Longer Context = Higher Intelligence + Better Performance**

Transformers Scale Better than LSTMs due to Improved Use of Context



“Better use of context”



“Better model performance”

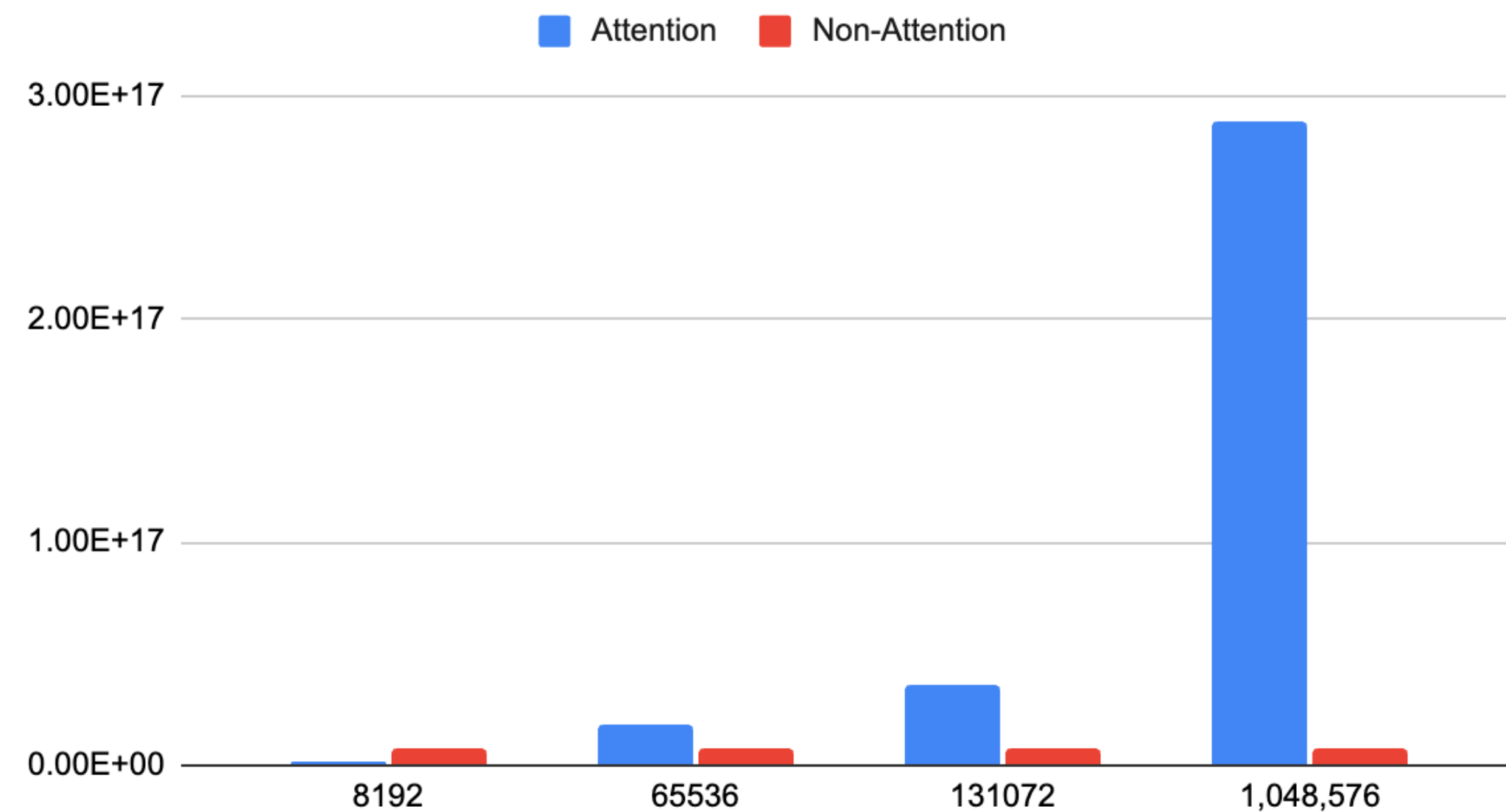
Long Context is hard for Transformer

- Complexity

- Expensive: quadratic complexity
 - Both time and space
 - $O(hn^2)$: h heads and sequence length of n

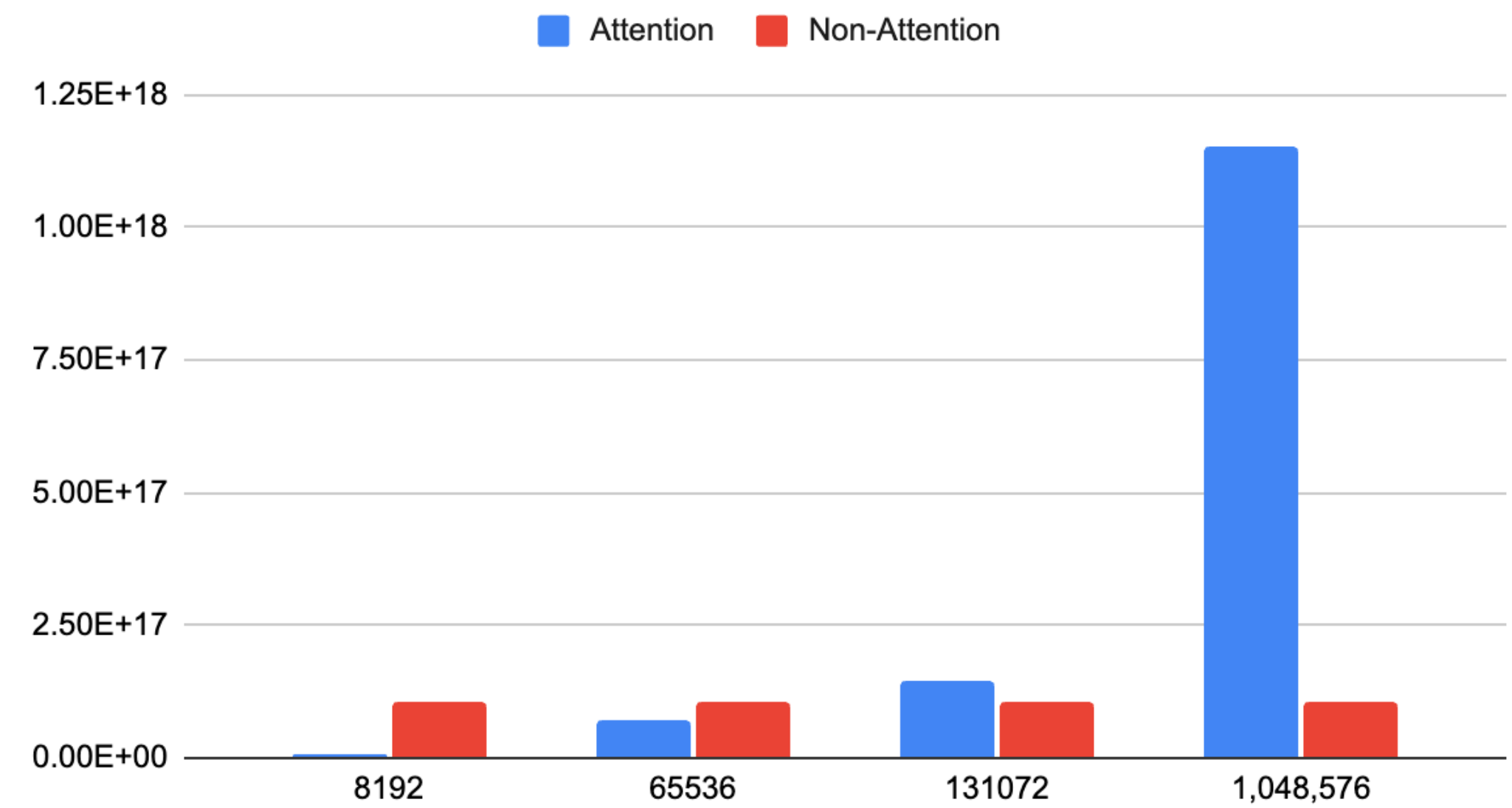
FLOPs of Attention and Non-attention Operations

FLOPs of Attention and Non-Attention (7B)



sequence length

FLOPs of Attention and Non-Attention (400B)

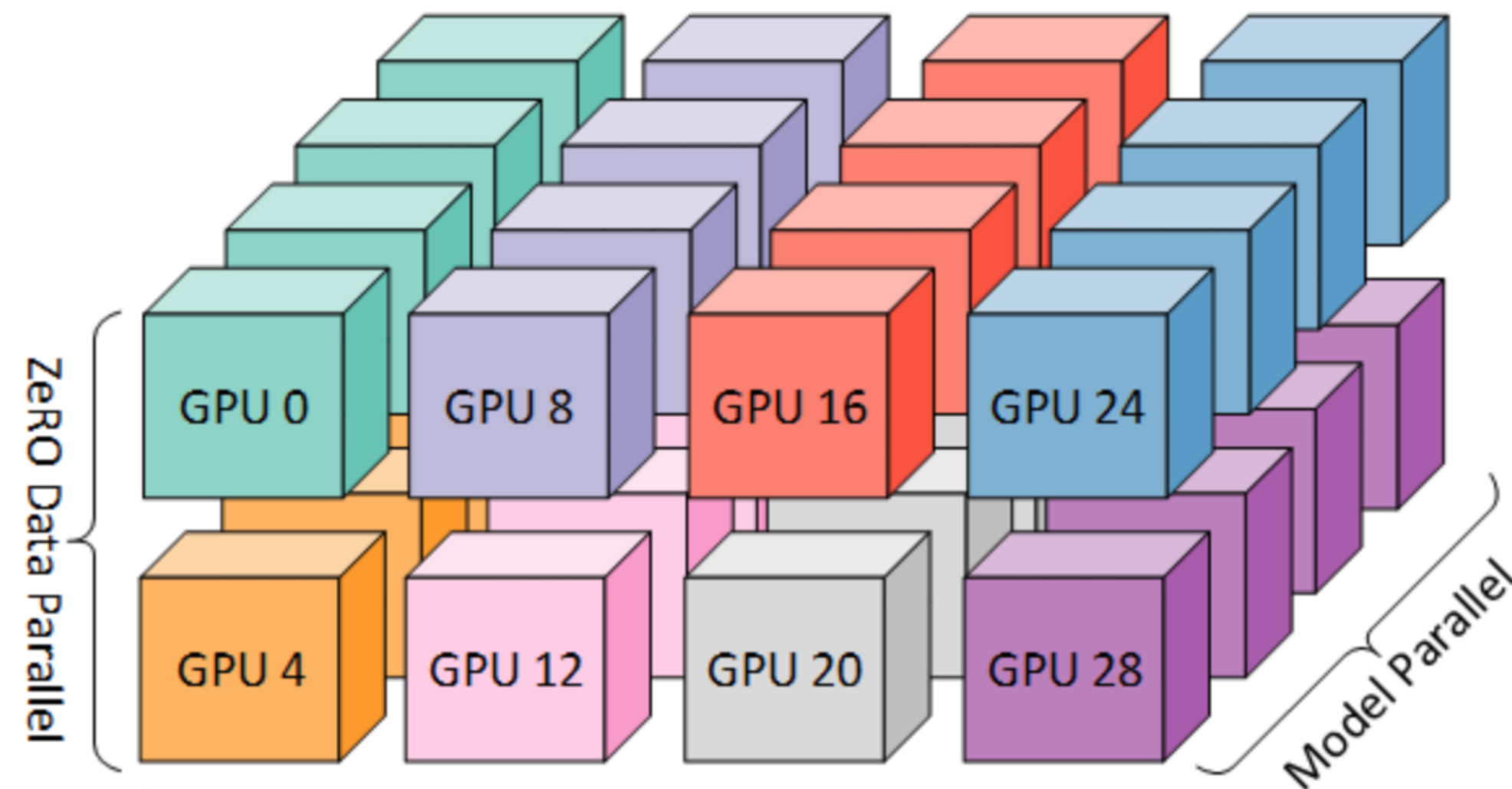


sequence length

FLOPs of attention is not negligible even for very large models!

Full Attention: quadratic Complexity

- It is hard to do distributed training across the sequence dimension on multiple devices with Transformers



Even with efficient attention or advanced sequence parallel implementations:

- Memory cost ($2bnh$) scales linearly with sequence length n (e.g. flash attention)
- Huge communication overhead and FLOPs (e.g. ring attention, context parallelism, etc)
- KV cache takes huge memory at inference time when the prompt is very long

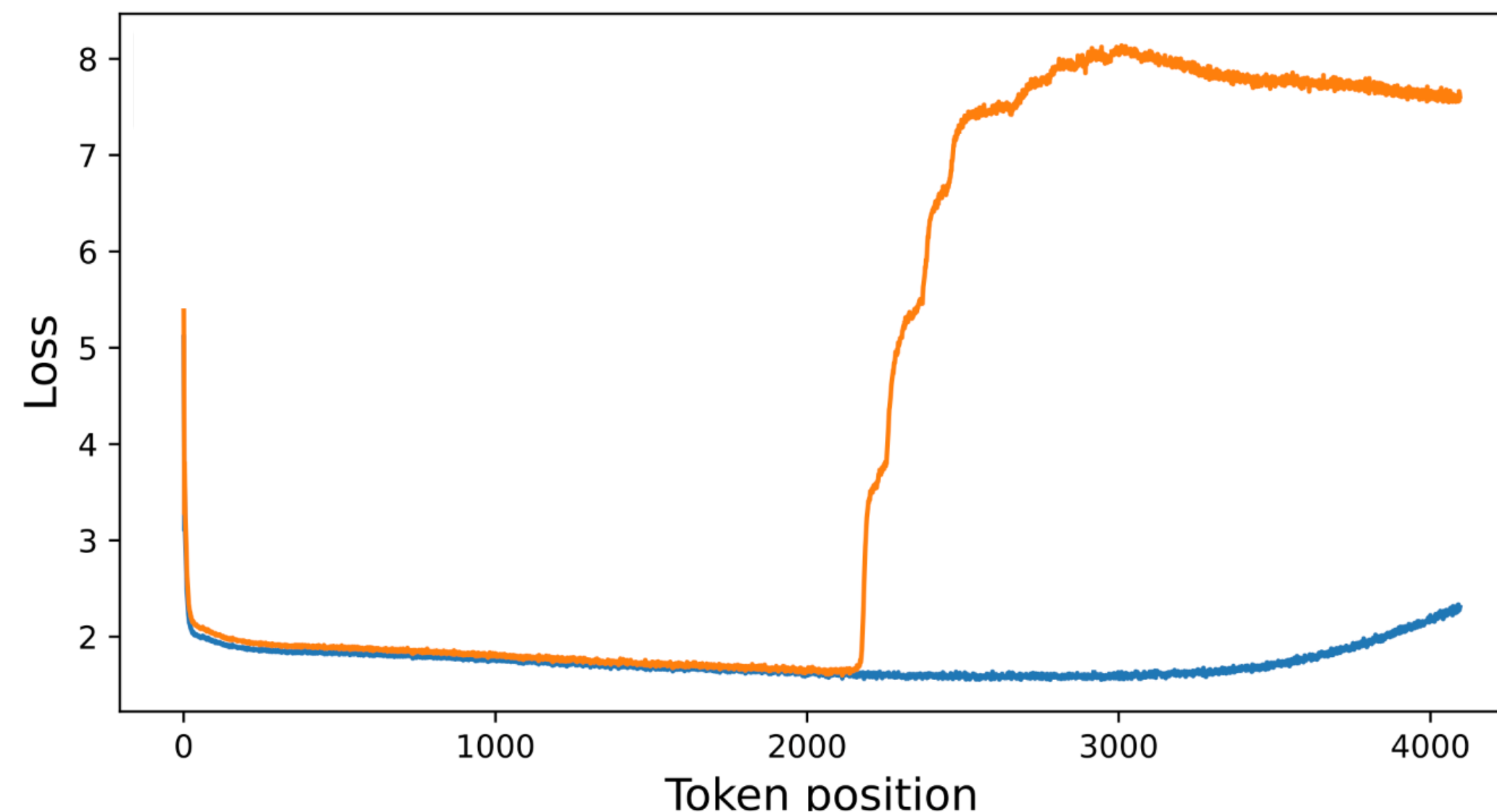
Long Context is hard for Transformer

- Complexity

- Expensive: quadratic complexity
 - Both time and space
 - $O(hn^2)$: h heads and sequence length of n

- Inductive Bias on Length Generalization

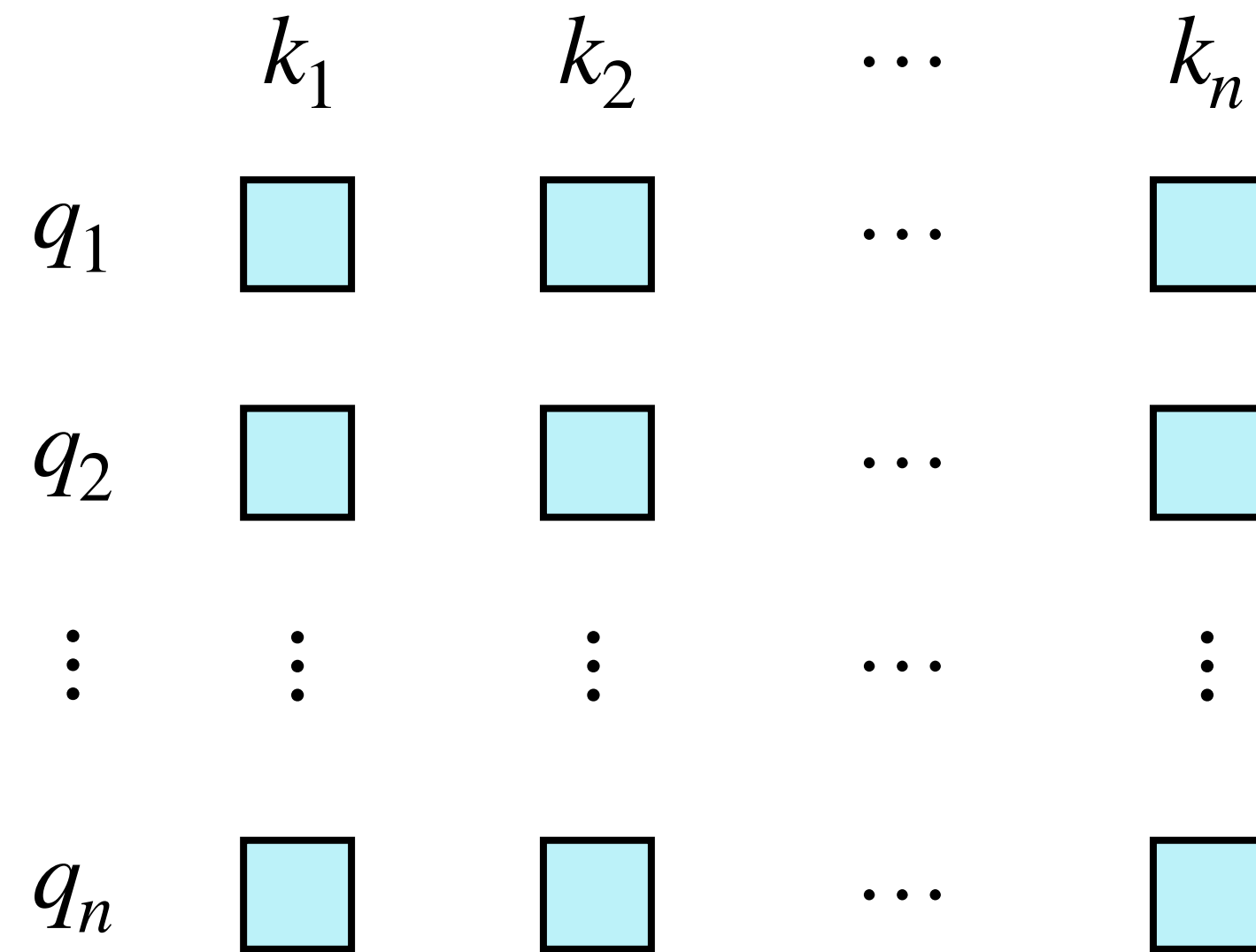
- Hard to be extended to longer contexts than pre-training



Weak Inductive Bias on Length Generalization

- Transformer Hard to be extended to longer contexts than pre-training

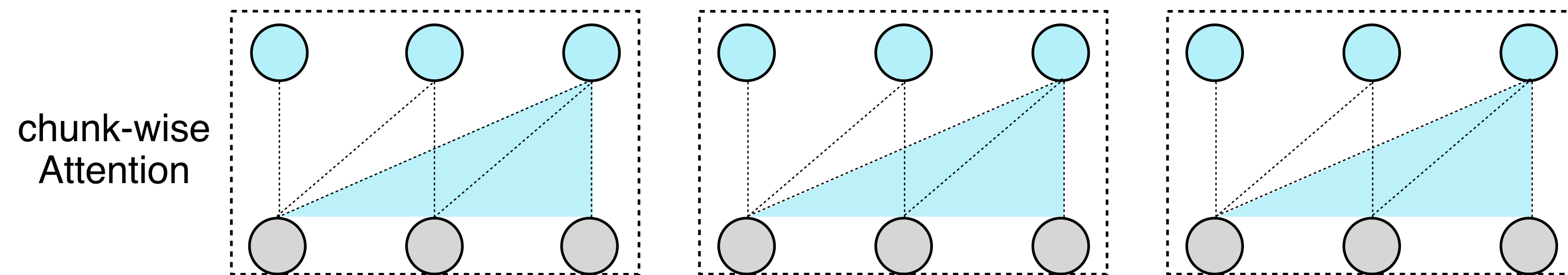
Attention score $A \in \mathbb{R}^{n \times n}$



- Pair-wise interaction
- Order-invariant if no positional embeddings

Why not Chunk-wise Attention?

- Applying attention individually to each chunk with fixed length



Pros

- **Efficiency**: Linear complexity
- **Inductive bias**: No need for length extension

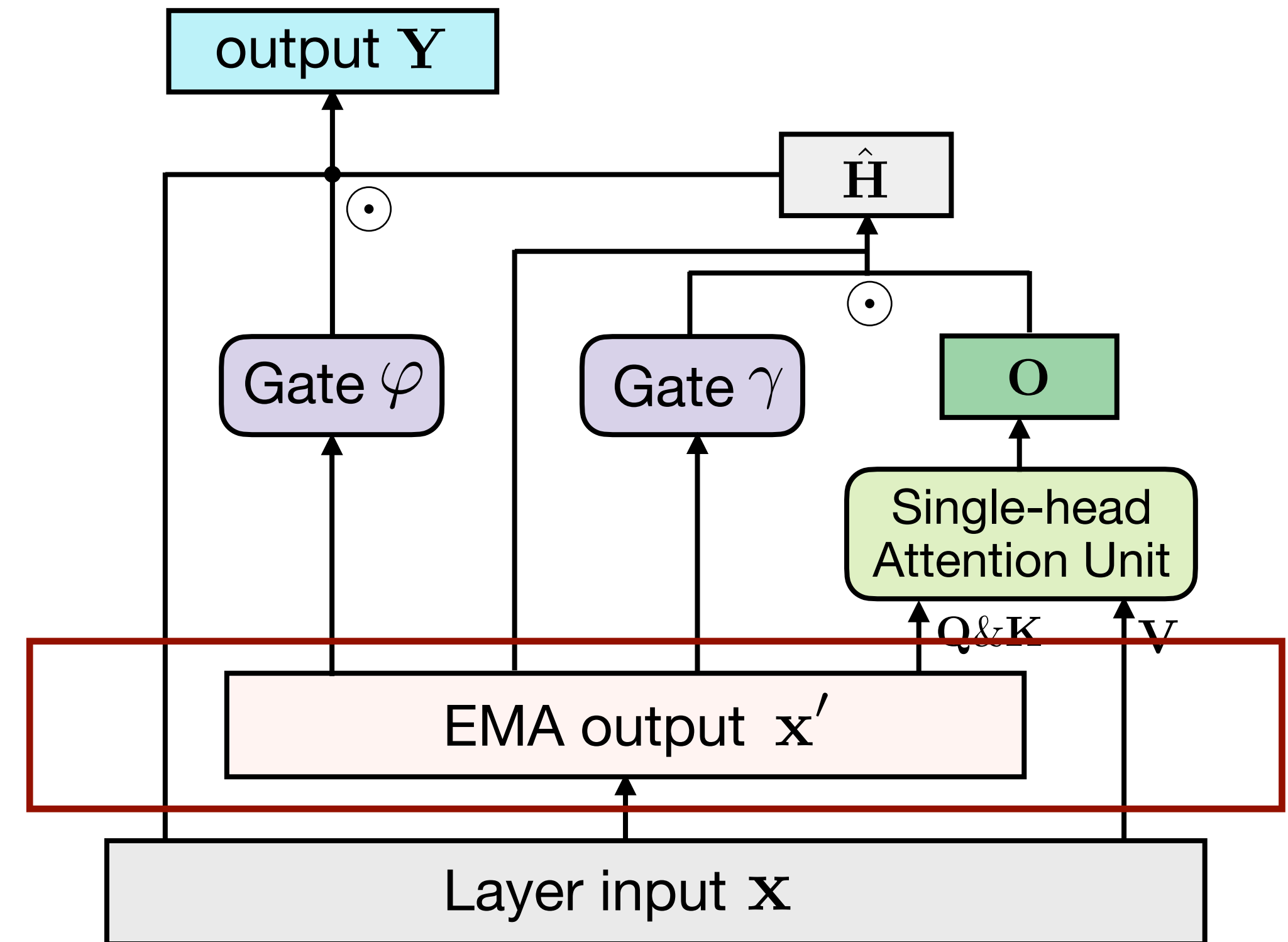
Cons

- **Capacity**: Limited contexts

Mega: Executive Summary

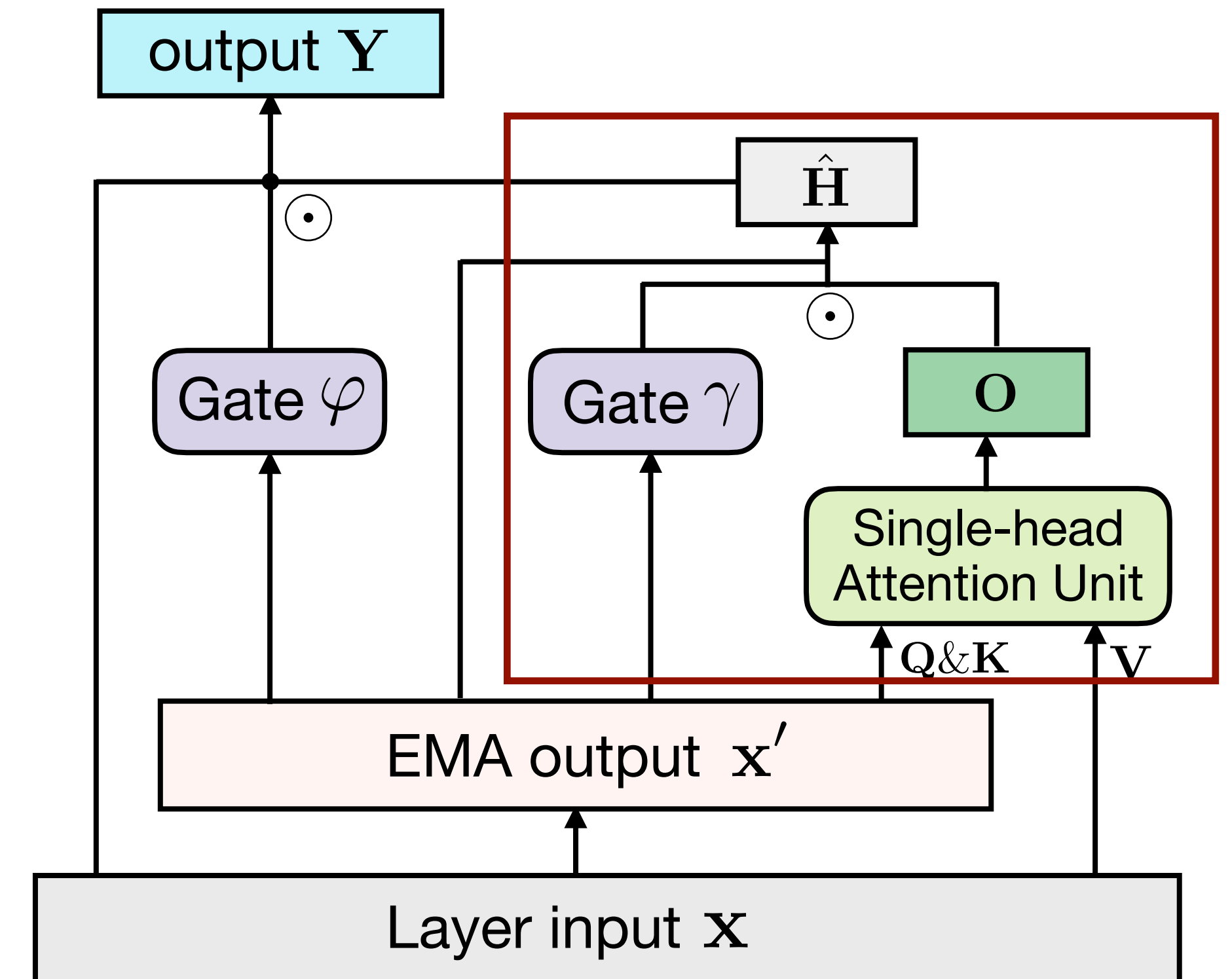
Mega Architecture: Outline

- **Exponential Moving Average (EMA)**
 - Local dependencies that decaying exponentially over time



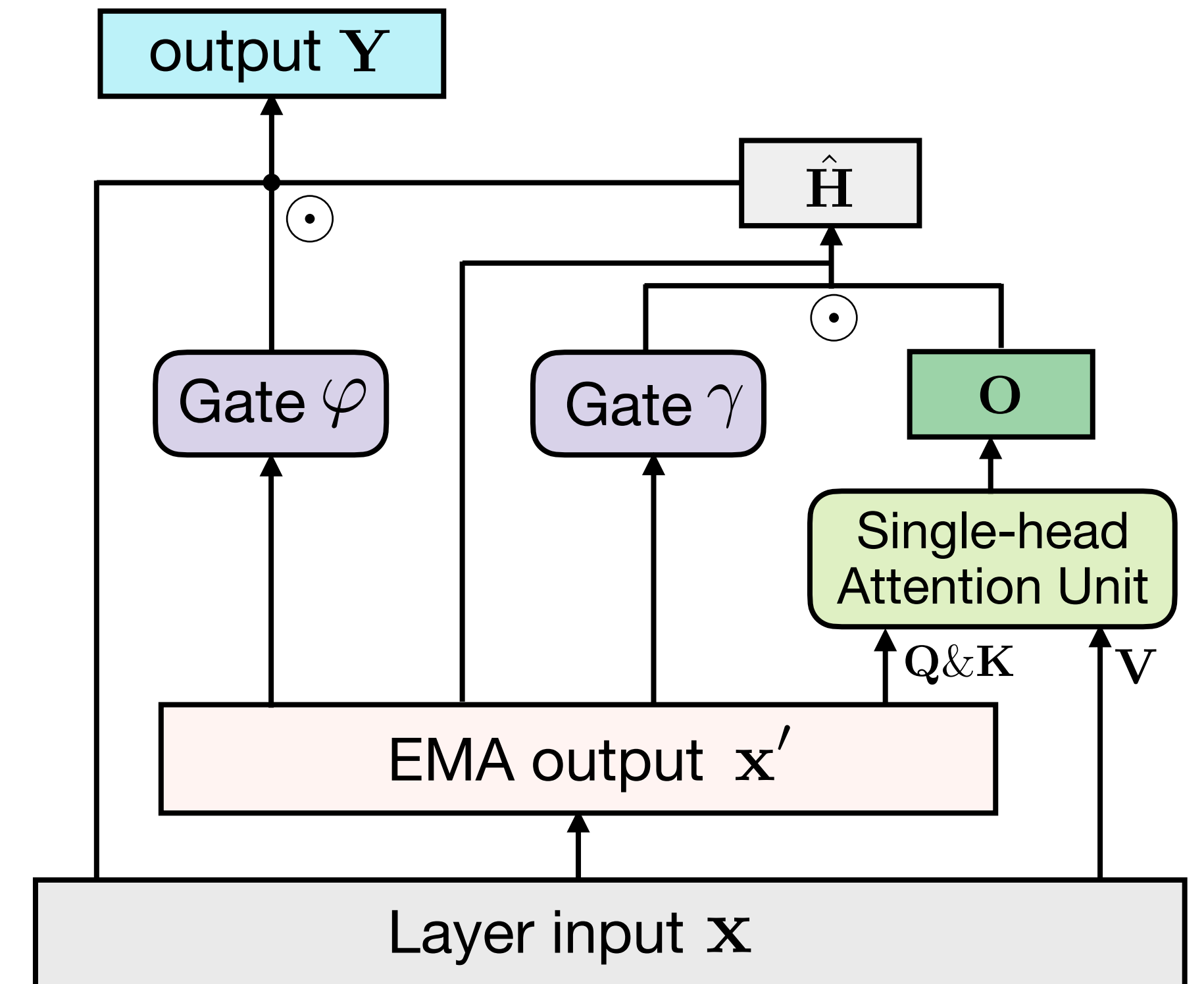
Mega Architecture: Outline

- **Exponential Moving Average (EMA)**
 - Local dependencies that decaying exponentially over time
- **Single-head Gated Attention**
 - Adding a reset gate to the attention output
 - Theoretically proving that **single-head** gated attention is as expressive as **multi-head** one



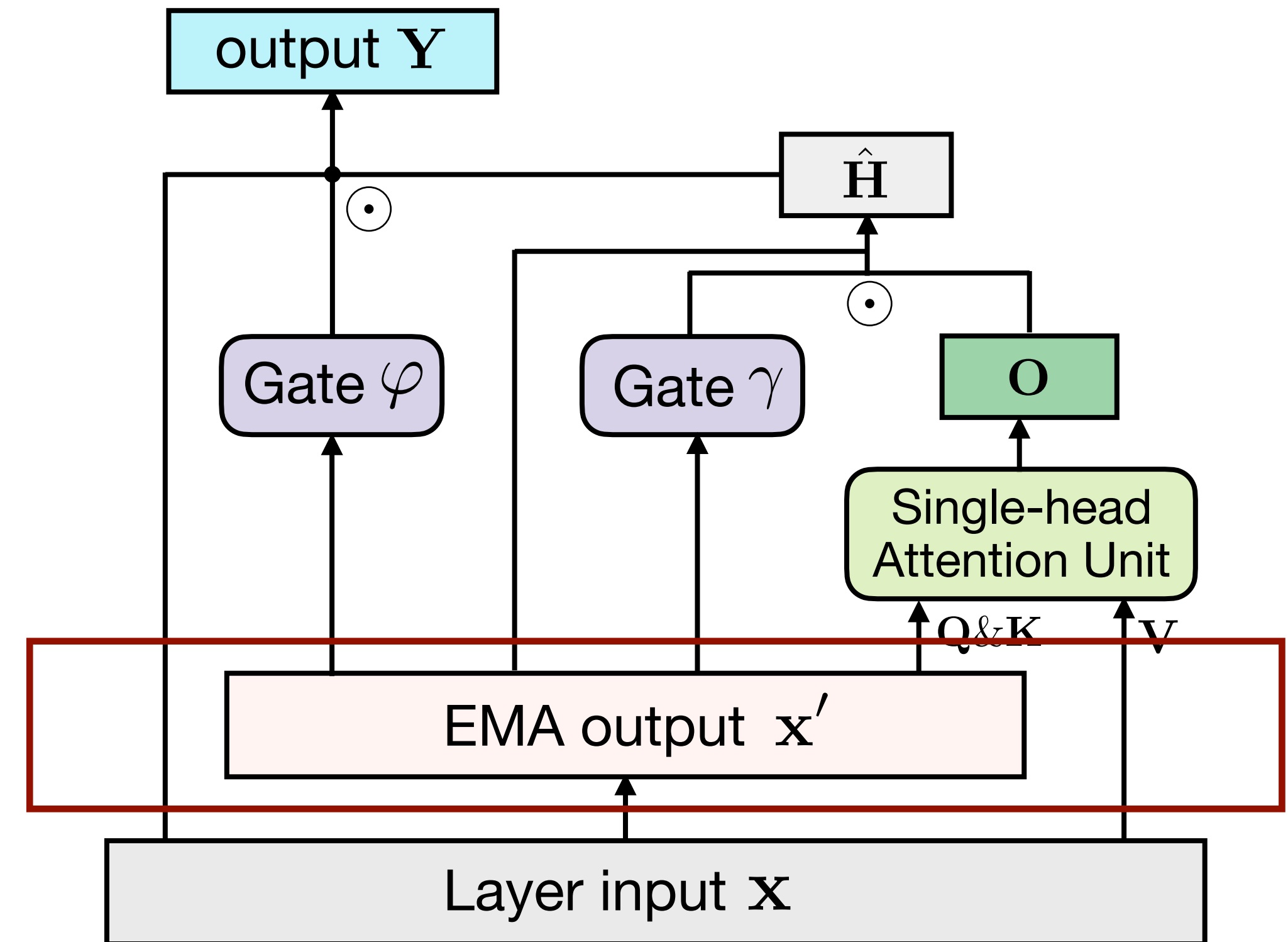
Mega Architecture: Outline

- **Exponential Moving Average (EMA)**
 - Local dependencies that decaying exponentially over time
- **Single-head Gated Attention**
 - Adding a reset gate to the attention output
 - Theoretically proving that **single-head** gated attention is as expressive as **multi-head** one
- **Mega-Chunk**
 - Applying attention to local chunks of fixed length
 - Reducing quadratic complexity to linear



Mega Architecture: Outline

- **Exponential Moving Average (EMA)**
 - Local dependencies that decaying exponentially over time
- **Single-head Gated Attention**
 - Adding a reset gate to the attention output
 - Theoretically proving that single-head gated attention is as expressive as multi-head one
- **Mega-Chunk**
 - Applying attention to local chunks of fixed length
 - Reducing quadratic complexity to linear



Exponential Moving Average (EMA)

Notations: Assuming 1-dim input sequence $\mathbf{x} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$, $\mathbf{x}_t \in \mathbb{R}$

EMA

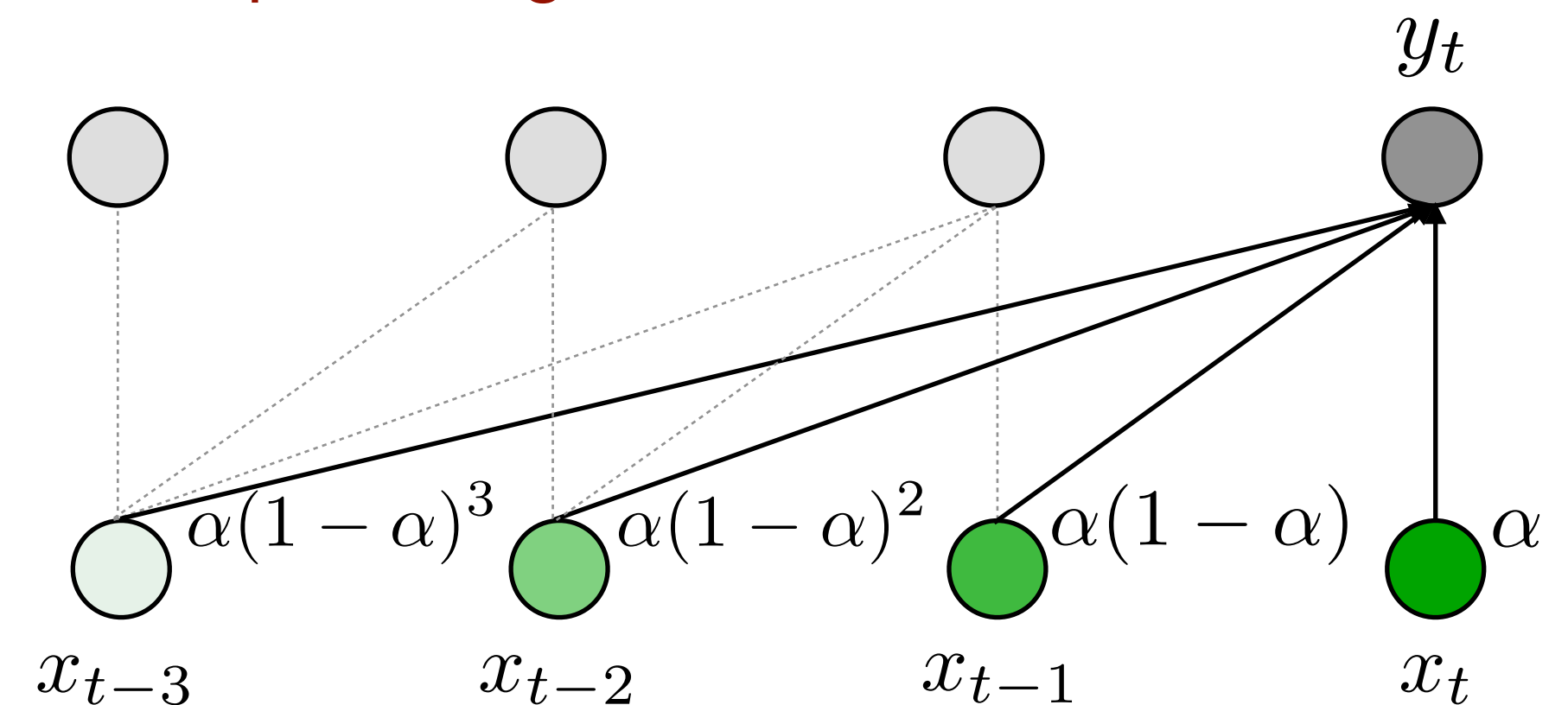
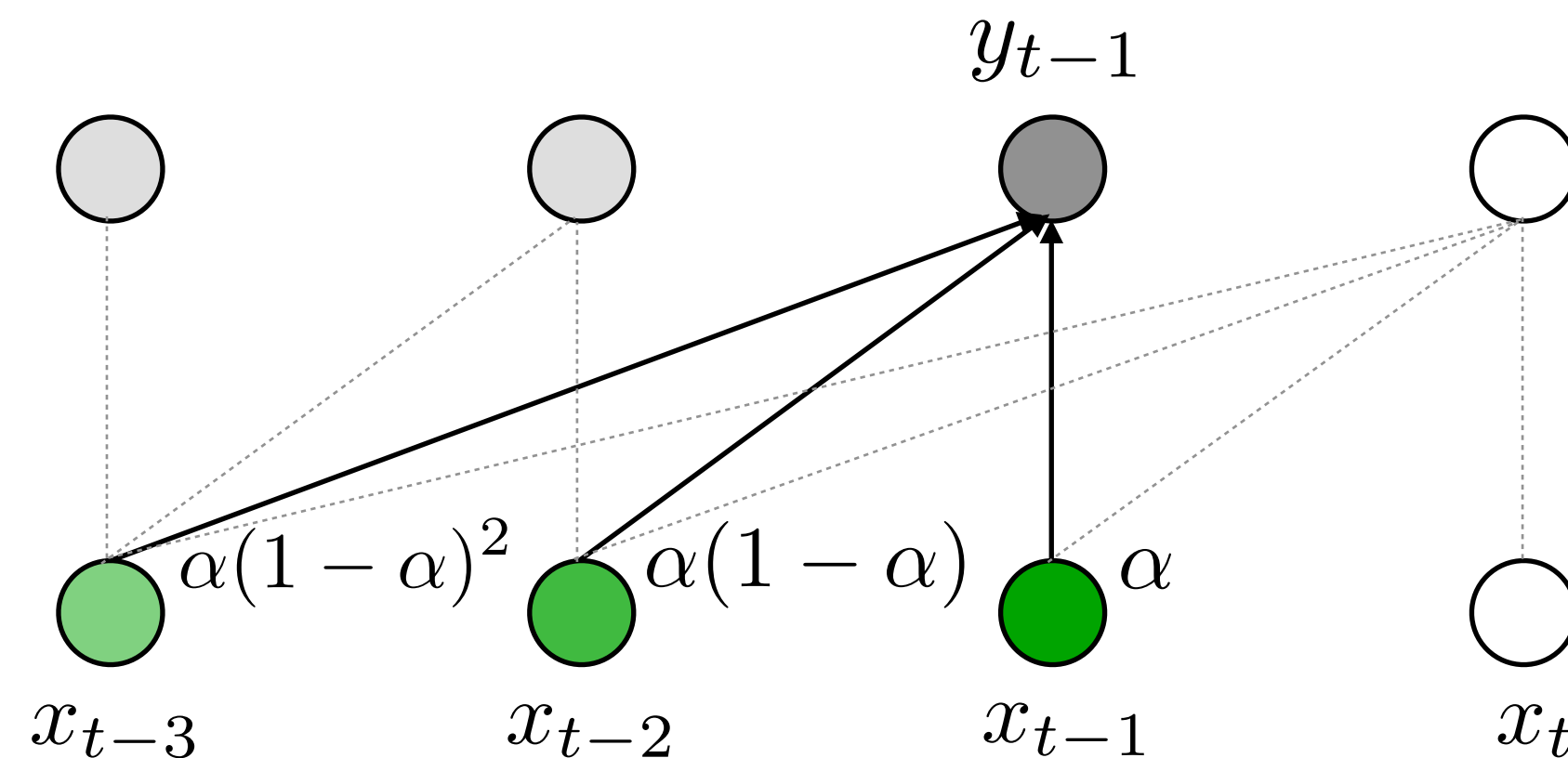
$$\mathbf{y}_t = \boxed{\alpha} \odot \mathbf{x}_t + \boxed{1 - \alpha} \odot \mathbf{y}_{t-1}, \quad \alpha \in (0, 1)$$

Damped EMA

$$\mathbf{y}_t = \alpha \odot \mathbf{x}_t + (1 - \alpha \odot \delta) \odot \mathbf{y}_{t-1}, \quad \delta \in (0, 1)$$

Learnable
Parameters

Relaxing the coupled weights

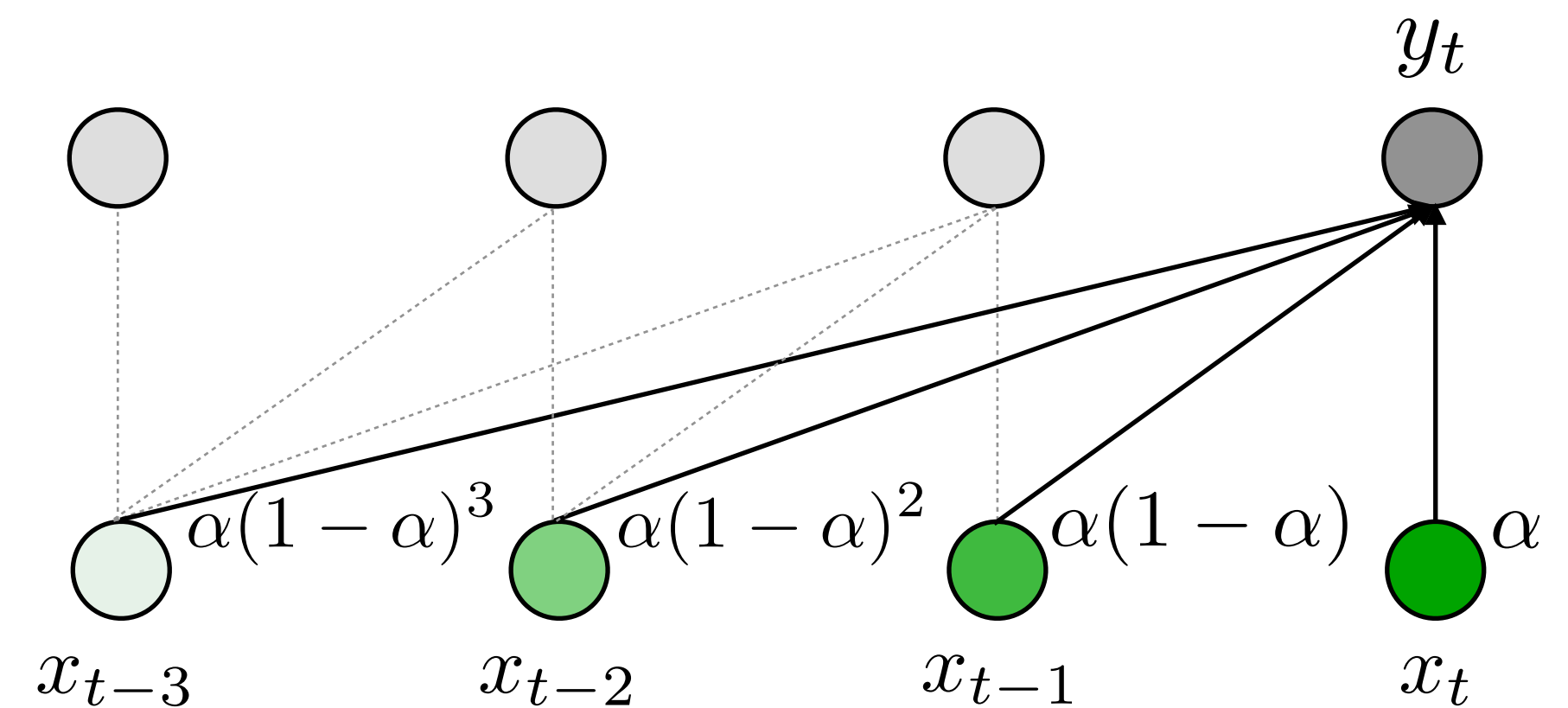
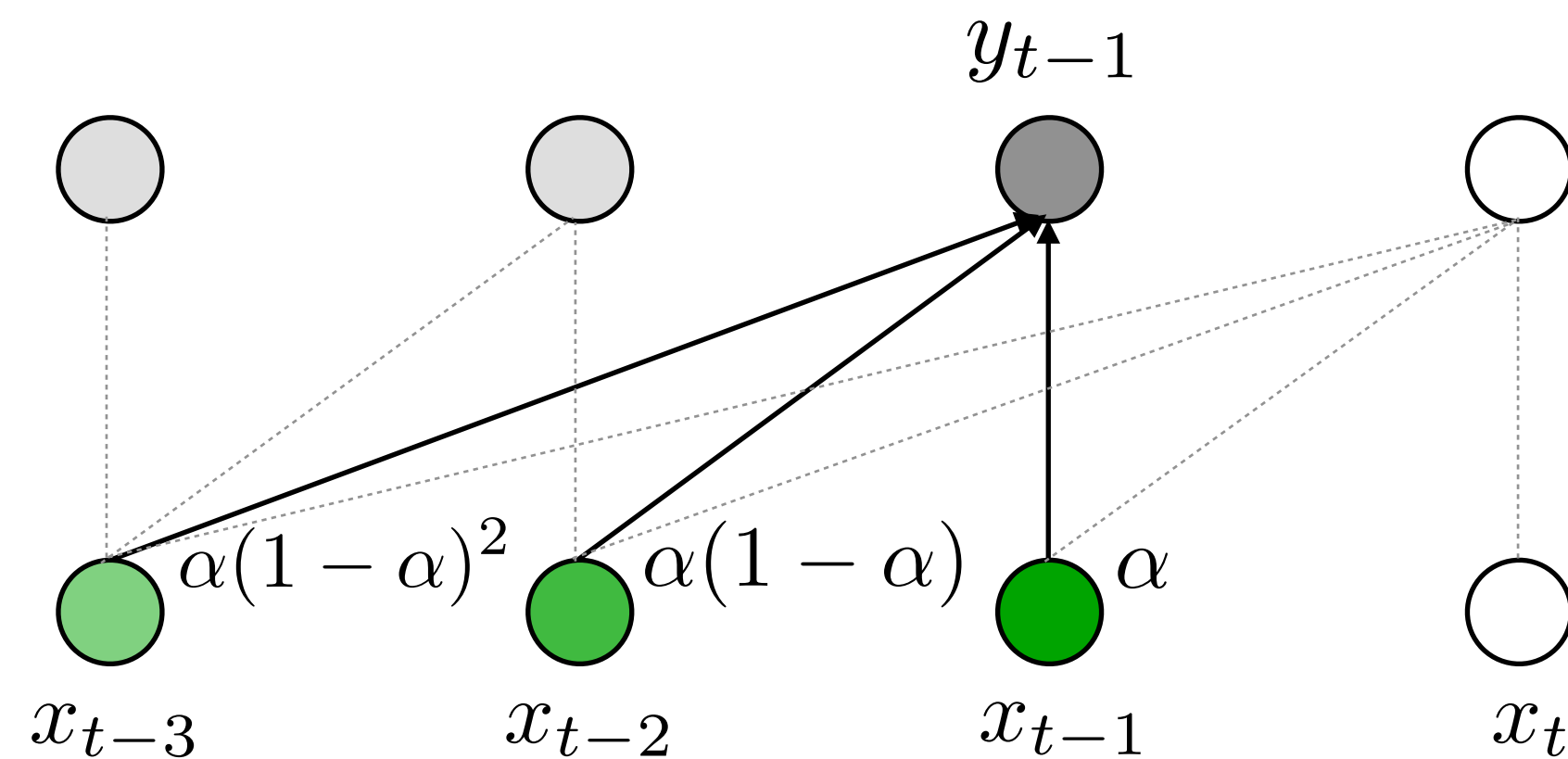


Efficient Algorithm for EMA

- Efficiently compute EMA outputs of all tokens in parallel

$$\mathbf{y}_t = \boxed{\alpha} \odot \mathbf{x}_t + \boxed{(1 - \alpha \odot \delta)} \odot \mathbf{y}_{t-1}$$

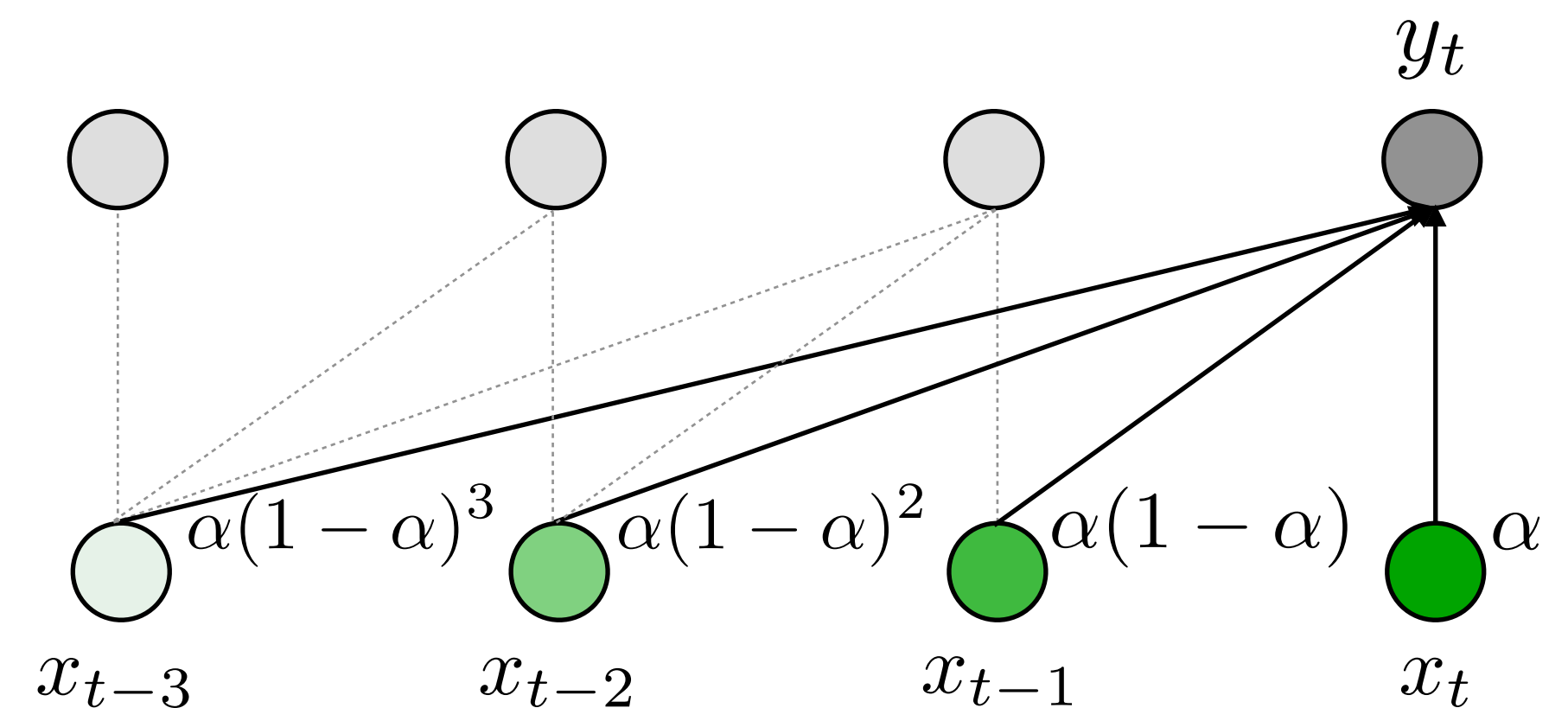
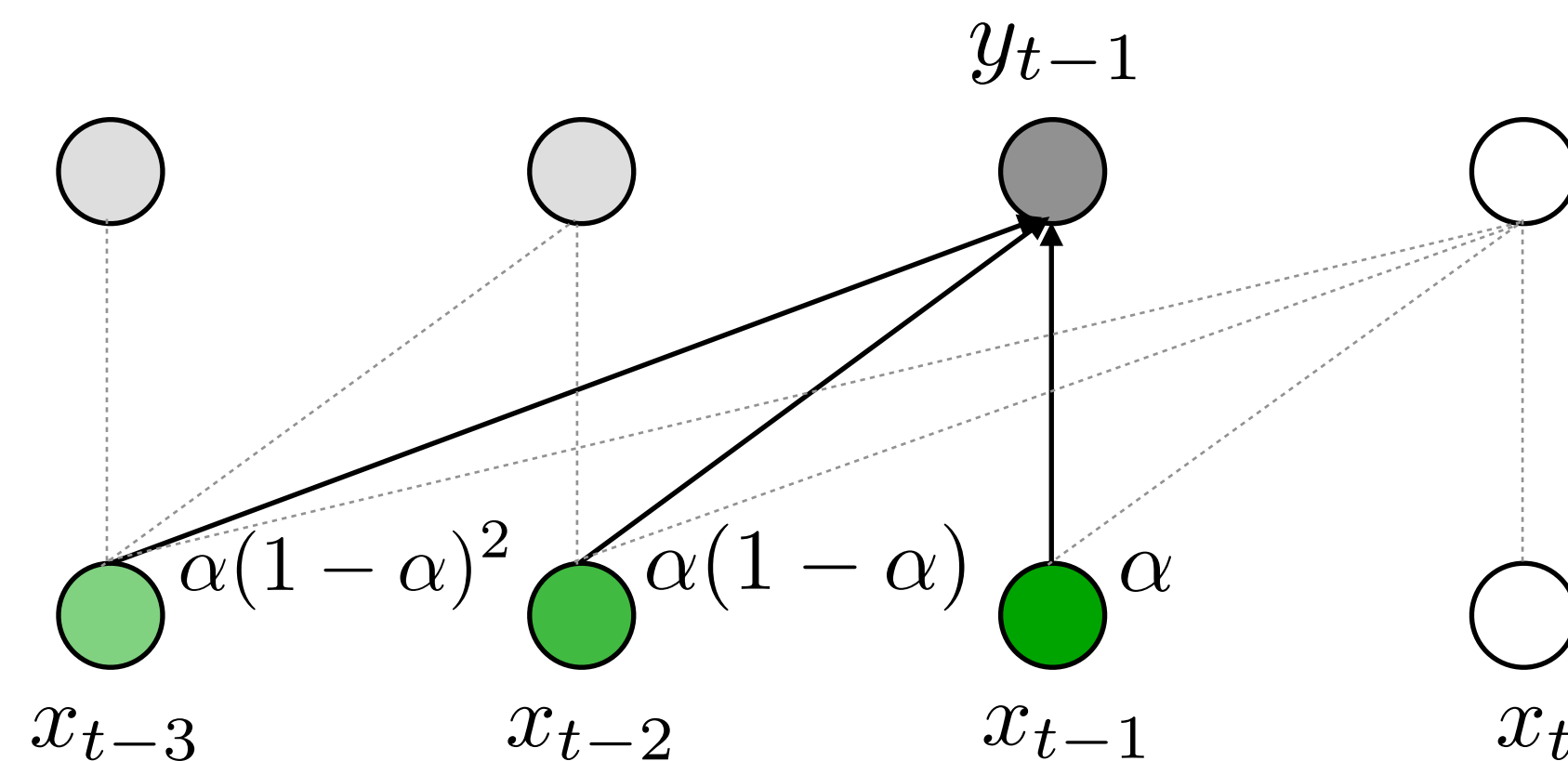
EMA weights are input independent



We can compute the weights for each input tokens in advance and compute EMA with FFTs.

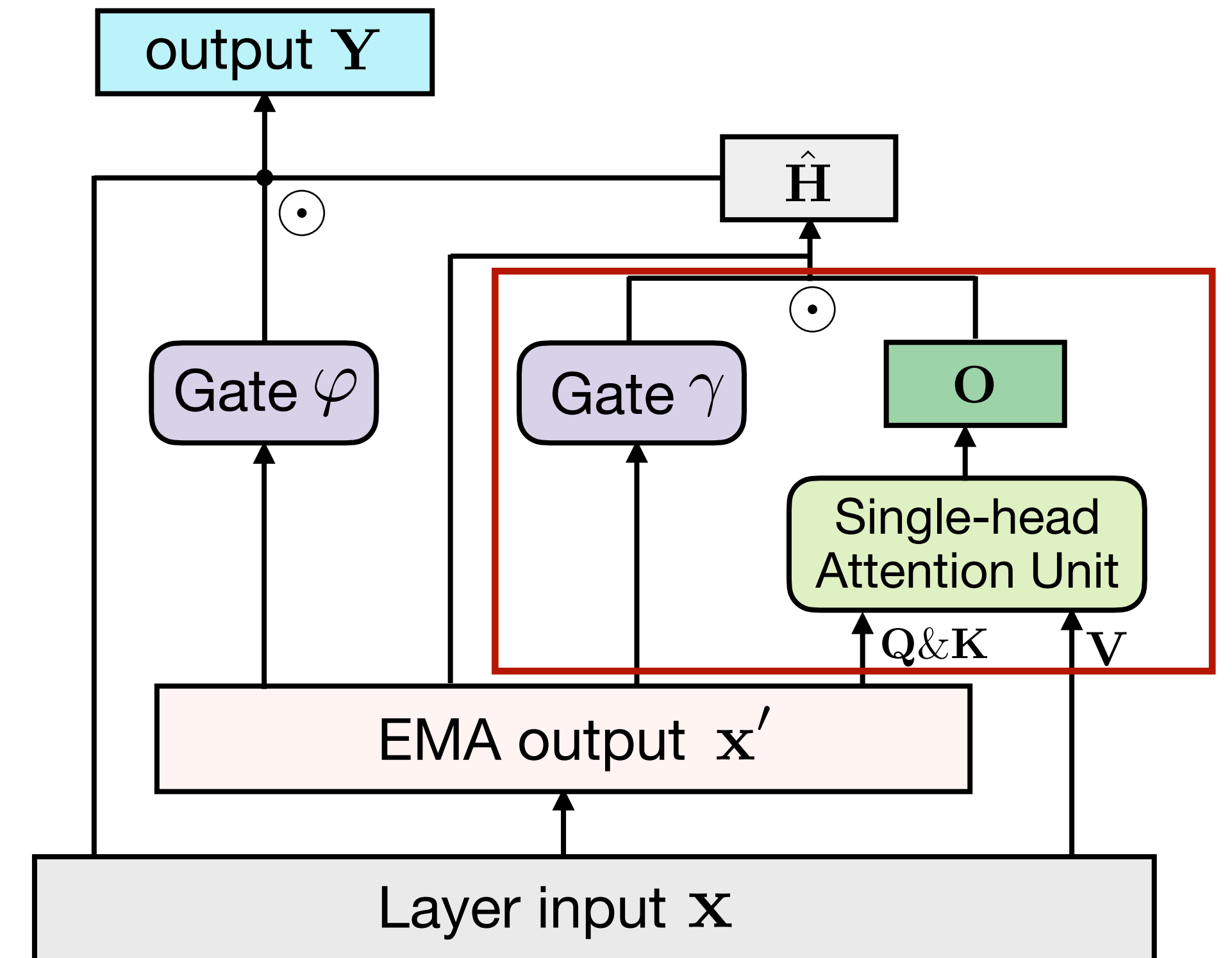
EMA: Summary

- **Strong Inductive Bias**
 - Favors recent, local contexts
 - In principle, can peek **unbounded** context history
- **Efficient to compute and easy to learn**
 - Parallel computation
 - Less parameters & capacity
 - Extremely simplified state space model



Mega Architecture: Outline

- **Exponential Moving Average (EMA)**
 - Local dependencies that decaying exponentially over time
 - Connection with State Space Models (e.g S4)
- **Single-head Gated Attention**
 - Adding a reset gate to the attention output
 - Theoretically proving that **single-head** gated attention is as expressive as multi-head one
- **Mega-Chunk**
 - Applying attention to local chunks of fixed length
 - Reducing quadratic complexity to linear

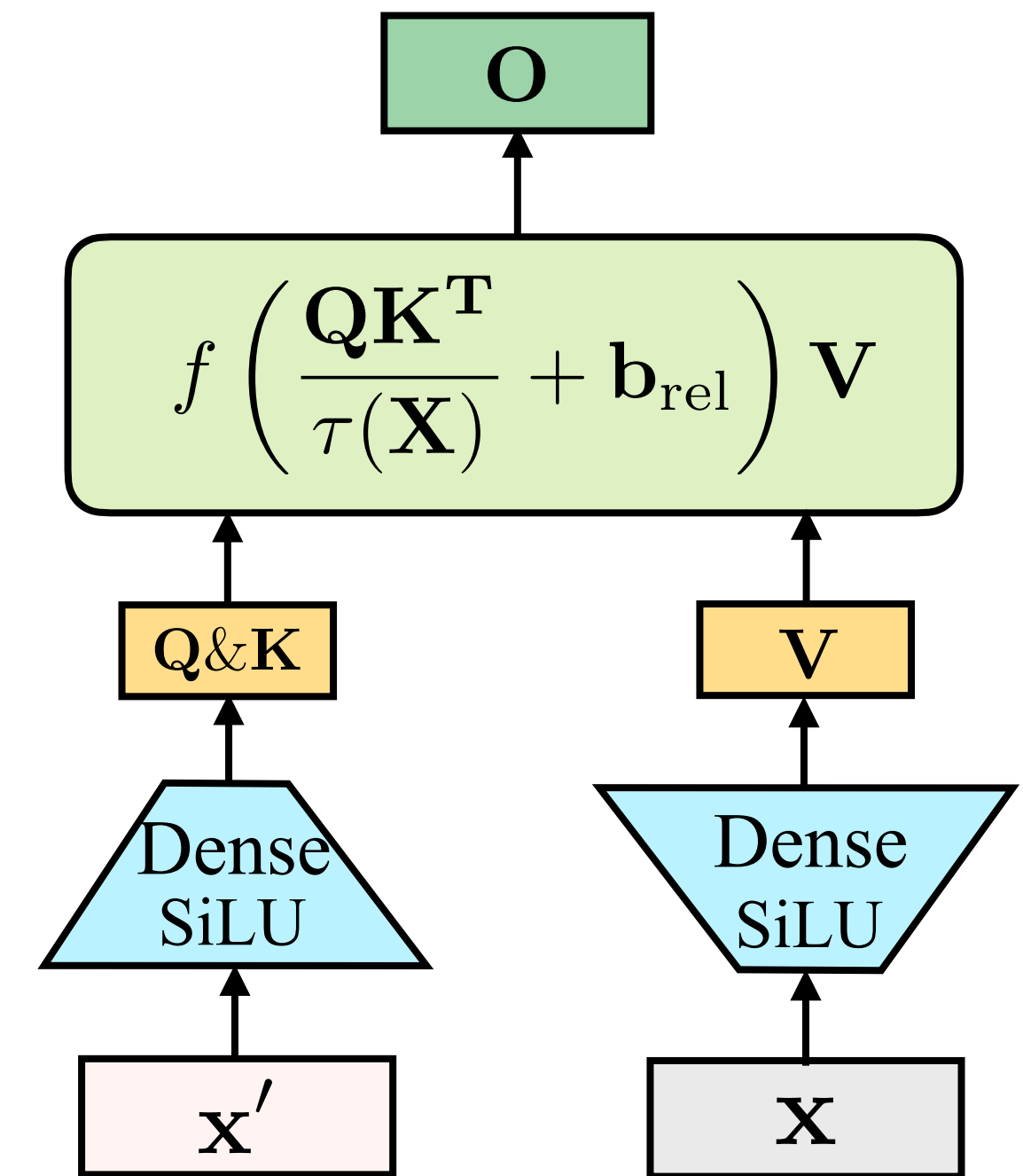


Single-head Gated Attention

- Adding a reset gate to the attention output

Attention:

$$O = \text{softmax} \left(\frac{QK^T}{\tau(X)} + b_{\text{rel}} \right) V$$



Single-head Gated Attention

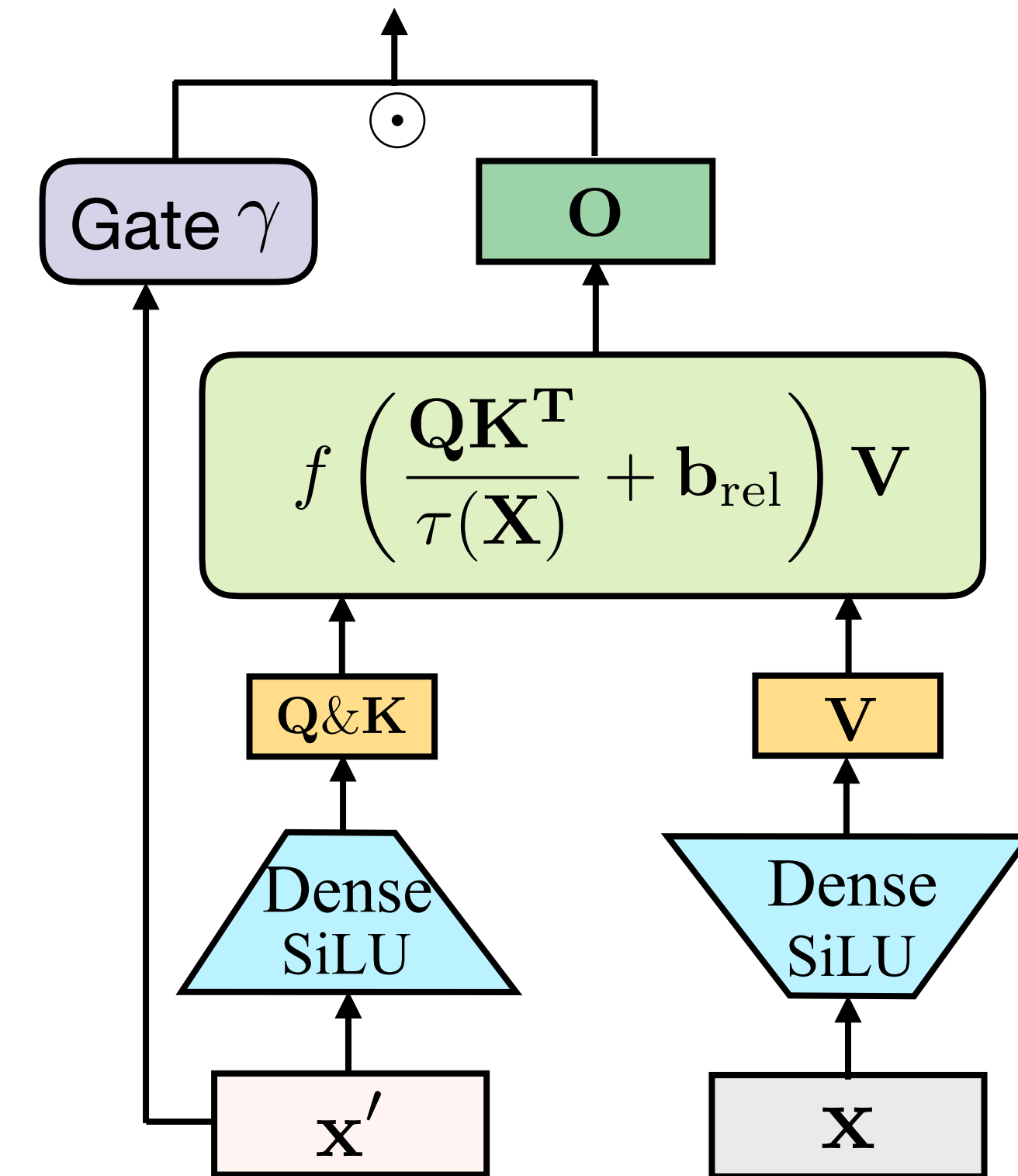
- Adding a reset gate to the attention output

Attention: $O = \text{softmax} \left(\frac{QK^T}{\tau(X)} + b_{\text{rel}} \right) V$

Gated Attention: $\gamma = \mathcal{G}(X)$

$$O_{\text{SHGA}} = O \odot \gamma$$

$\mathcal{G}()$ is a transform, e.g. an MLP



Experimental Results

	LRA	WMT'14	WikiText-103	ImageNet	Raw-SC
XFM	59.24	27.68	18.66	81.80	31.24
S4	85.86	—	20.95	—	97.50
Mega	88.21	29.01	18.07	82.31	97.30

Limitations of Mega

- **Mode Capacity vs. Full Attention**
 - Limited capacity of EMA sub-layer
 - Mega-chunk fails behind Mega w. full attention

