

(Manifest AI)

Scaling Context Requires Rethinking Attention

[ICLR 2025 preprint]

Hankyo Jeong

December 18, 2025

Seoul National University

- Autoregressive inference at step t uses past $\{k_1, \dots, k_t\}$ and $\{v_1, \dots, v_t\}$.
- Recomputing past K, V each step increases cost.
- Store past K, V as KV cache.

KV-cache state for length T : state $\sim O(cT)$, for some constant c

Causal Self-Attention

- $q_i, k_i \in \mathbb{R}^d, v_i \in \mathbb{R}^v$.

$$Q = \begin{bmatrix} q_1^\top \\ \vdots \\ q_t^\top \end{bmatrix} \in \mathbb{R}^{t \times d}, \quad K = \begin{bmatrix} k_1^\top \\ \vdots \\ k_t^\top \end{bmatrix} \in \mathbb{R}^{t \times d}, \quad V = \begin{bmatrix} v_1^\top \\ \vdots \\ v_t^\top \end{bmatrix} \in \mathbb{R}^{t \times v}.$$

$$M \in \mathbb{R}^{t \times t}, \quad M_{ij} = \mathbf{1}_{i \geq j}.$$

$$\text{attn}_{\text{exp}}(Q, K, V) = (\exp(QK^\top) \odot M) V.$$

- Compute $t \times t$ interactions $\Rightarrow$ cost scales with t^2 .
- KV cache grows with t .

Two Limits of Long Context

- Exponential attention:

compute $\sim O(t^2 \cdot d)$, state $\sim O(t)$.

- Linear attention:

compute $\sim O(t)$, state $\sim O(1)$ (fixed feature size).

- Fixed-size state can limit performance at long context.

Goal: larger state capacity + GPU-parallel execution form.

Linear Attention

Linear Attention

- Feature map: $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$.
- Row-wise application:

$$\phi(Q) = \begin{bmatrix} \phi(q_1)^\top \\ \vdots \\ \phi(q_t)^\top \end{bmatrix} \in \mathbb{R}^{t \times D} \quad \text{and} \quad \phi(K) \in \mathbb{R}^{t \times D}.$$

$$\text{attn}_{\text{lin}}^\phi(Q, K, V) = (\phi(Q)\phi(K)^\top \odot M) V.$$

Linear Attention: Recurrent Form and Linear Cost

- Index i denotes the i -th row (time step).

$$S_i := \sum_{j=1}^i v_j \phi(k_j)^\top \in \mathbb{R}^{v \times D}.$$

$$\text{attn}_{\text{lin}}^\phi(Q, K, V)_i = \sum_{j=1}^i \phi(q_i)^\top \phi(k_j) v_j = \left(\sum_{j=1}^i v_j \phi(k_j)^\top \right) \phi(q_i) = S_i \phi(q_i).$$

- $v_i \phi(k_i)^\top : O(vD), \quad S_i \phi(q_i) : O(vD).$
- Total: $O(tDv).$

State Size Limit in Linear Attention

- Linear attention replaces KV cache with $S_i \in \mathbb{R}^{v \times D}$.
- State capacity depends on D .
- If D is fixed, state capacity stays fixed as t increases.

Problem: fixed $D \rightarrow$ fixed state capacity at long context.

Power Attention

Power Attention: from $\exp(QK^\top)$ to $(QK^\top)^{\odot p}$

$$\text{attn}_{\text{exp}}(Q, K, V) = (\exp(S) \odot M) V = (\exp(QK^\top) \odot M) V,$$

$$\text{attn}_{\text{pow}}^p(Q, K, V) = (S^{\odot p} \odot M) V = ((QK^\top)^{\odot p} \odot M) V.$$

$A^{\odot p}$: elementwise p -th power.

Choose ϕ by Tensor Power (TPOW)

Lemma 1

Let $p \in \mathbb{N}$. define $\text{TPOW}_p : \mathbb{R}^d \rightarrow \mathbb{R}^{d^p}$ by

$$\phi = \text{TPOW}_p(x) = \begin{bmatrix} x_1 \cdots x_1 \\ x_1 \cdots x_2 \\ \vdots \\ x_d \cdots x_d \end{bmatrix} = \left[\prod_k x_{i_k} \right]_{(i_1, \dots, i_p) \in \mathbb{N}_d^{\times p}}.$$

Then for any $q, k \in \mathbb{R}^d$, $\text{TPOW}_p(q)^\top \text{TPOW}_p(k) = (q^\top k)^p$.

Power attention is linear attention (with $\phi = \text{TPOW}_p$)

- Linear attention with feature map $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$:

$$\text{attn}_{\text{lin}}^{\phi}(Q, K, V) := (\phi(Q)\phi(K)^{\top} \odot M) V, \quad \phi(Q) = \begin{bmatrix} \phi(q_1)^{\top} \\ \vdots \\ \phi(q_t)^{\top} \end{bmatrix}.$$

- With $\phi = \text{TPOW}_p$, Lemma1 gives, entrywise for all i, j :

$$(\phi(Q)\phi(K)^{\top})_{ij} = \phi(q_i)^{\top} \phi(k_j) = (q_i^{\top} k_j)^p = ((QK^{\top})^{\odot p})_{ij}.$$

$$\begin{aligned} \text{attn}_{\text{pow}}^p(Q, K, V) &= ((QK^{\top})^{\odot p} \odot M) V \\ &= (\phi(Q)\phi(K)^{\top} \odot M) V \\ &= \text{attn}_{\text{lin}}^{\phi}(Q, K, V). \end{aligned}$$

From TPOW to SPOW: permutation duplicates

- **Example** ($p = 2$, $x = [a, b, c]^T$).

TPOW produces a symmetric tensor

$$\text{TPOW}_2(x) = \begin{bmatrix} aa & ab & ac \\ ab & bb & bc \\ ac & bc & cc \end{bmatrix}$$

SPOW produces unique elements of that tensor

$$\text{SPOW}_2(x) = [aa, ab, ac, bb, bc, cc]^T.$$

- scaled by a coefficient that depends on the permutation count.

$$\text{SPOW}_2\left(\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}\right) = \begin{bmatrix} x_1 x_1 \\ \sqrt{2} x_1 x_2 \\ x_2 x_2 \end{bmatrix}, \quad \text{SPOW}_3\left(\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}\right) = \begin{bmatrix} x_1 x_1 x_1 \\ \sqrt{3} x_1 x_1 x_2 \\ \sqrt{3} x_1 x_2 x_2 \\ x_2 x_2 x_2 \end{bmatrix}.$$

Which gives:

1. For $x, y \in \mathbb{R}^d$:

$$\langle \text{SPOW}_p(x), \text{SPOW}_p(y) \rangle = (x^\top y)^p = \langle \text{TPOW}_p(x), \text{TPOW}_p(y) \rangle.$$

2. Dimension:

$$D = \binom{d+p-1}{p}.$$

Recap: Power attention via ϕ

- Power attention:

$$\text{attn}_{\text{pow}}^p(Q, K, V) = ((QK^\top)^{\odot p} \odot M) V.$$

- If ϕ satisfies $\phi(q)^\top \phi(k) = (q^\top k)^p$, then

$$\text{attn}_{\text{pow}}^p(Q, K, V) = \text{attn}_{\text{lin}}^\phi(Q, K, V).$$

$$\phi = \text{TPOW}_p \Rightarrow D = d^p, \quad \phi = \text{SPOW}_p \Rightarrow D = \binom{d+p-1}{p}.$$

GPU Parallelism via Chunking

Chunking for GPU Execution

- Recurrent update: $S_i \leftarrow S_{i-1}$ creates a dependency chain.
- GPU throughput is tied to batched matmul kernels.
- Chunking rewrites computation into chunk-level operations.
- Let chunk size be c and chunk index be n .
- Use only boundary states: $S_0, S_c, S_{2c}, \dots$

Chunked Form

- Treat Q, K as feature-space inputs (after ϕ) to match the chunk derivation.

$$Y_{nc+m} = S_{nc} Q_{nc+m} + \sum_{j=nc+1}^{nc+m} (Q_{nc+m}^{\top} K_j) V_j, \quad m \in \{1, \dots, c\}.$$

$$S_{c(n+1)} = S_{cn} + V_{(n)c} K_{(n)c}^{\top}.$$

- Cost: $O(tDv + tcd)$.

Conclusion

- Exponential attention: compute scales with t^2 , KV-cache state scales with t .
- Linear attention: recurrent state $S_i \in \mathbb{R}^{v \times D}$, compute $O(tDv)$.
- Chunking: rewrite recurrent computation into GPU matmul form, cost $O(tDv + tcd)$.
- Power attention: score $(q^\top k)^p$ with a feature map satisfying $\phi(q)^\top \phi(k) = (q^\top k)^p$.
- TPOW gives $D = d^p$; SPOW gives $D = \binom{d+p-1}{p}$.

Thank you!

Appendix

Appendix A: Chunked Form Setup

- Chunk size c , chunk index i .
- Chunk matrices: $Q_{(i)c}, K_{(i)c}, V_{(i)c}$.
- Boundary state: S_{ci} .

$$Y_{(i)c} = S_{ci} Q_{(i)c} + V_{(i)c} (Q_{(i)c} K_{(i)c}^\top \odot M), \quad (1)$$

$$S_{c(i+1)} = S_{ci} + V_{(i)c} K_{(i)c}^\top. \quad (2)$$

Appendix A: Position-wise Expansion

$$Y_i = S_{c\lfloor i/c \rfloor} Q_i + \sum_{j=\lfloor i/c \rfloor c+1}^i (Q_i K_j^\top) V_j \quad (3)$$

$$= (S_{c(\lfloor i/c \rfloor - 1)} + V_{(i)c} K_{(i)c}^\top) Q_i + \sum_{j=\lfloor i/c \rfloor c+1}^i (Q_i K_j^\top) V_j \quad (4)$$

$$= S_{c(\lfloor i/c \rfloor - 1)} Q_i + \sum_{j=\lfloor (i-1)/c \rfloor c+1}^{\lfloor i/c \rfloor c} V_j K_j^\top Q_i + \sum_{j=\lfloor i/c \rfloor c+1}^i (Q_i K_j^\top) V_j \quad (5)$$

$$= S_{c(\lfloor i/c \rfloor - 1)} Q_i + \sum_{j=\lfloor (i-1)/c \rfloor c+1}^{\lfloor i/c \rfloor c} (K_j^\top Q_i) V_j + \sum_{j=\lfloor i/c \rfloor c+1}^i (Q_i K_j^\top) V_j. \quad (6)$$

Appendix A: Telescoping to Attention Form

$$Y_i = S_{c(\lfloor i/c \rfloor - 1)} Q_i + \sum_{j=\lfloor (i-1)/c \rfloor c + 1}^{\lfloor i/c \rfloor c} (Q_i K_j^\top) V_j + \sum_{j=\lfloor i/c \rfloor c + 1}^i (Q_i K_j^\top) V_j \quad (7)$$

$$= S_{c(\lfloor i/c \rfloor - 1)} Q_i + \sum_{j=\lfloor (i-1)/c \rfloor c + 1}^i (Q_i K_j^\top) V_j \quad (8)$$

$$= \dots = S_0 Q_i + \sum_{j=1}^i (Q_i K_j^\top) V_j \quad (9)$$

$$= \sum_{j=1}^i (Q_i K_j^\top) V_j \quad (\text{if } S_0 = 0). \quad (10)$$

Appendix B: Tensor Power Identity

$$\text{TPOW}(x, p) = \left[\prod_{k=1}^p x_{i_k} \right]_{(i_1, \dots, i_p) \in [d]^p} \quad (11)$$

$$\text{TPOW}(x, p)^\top \text{TPOW}(y, p) = \sum_{(i_1, \dots, i_p) \in [d]^p} x_{i_1} \cdots x_{i_p} y_{i_1} \cdots y_{i_p} \quad (12)$$

$$= \left(\sum_{i_1 \in [d]} x_{i_1} y_{i_1} \right) \left(\sum_{i_2 \in [d]} x_{i_2} y_{i_2} \right) \cdots \left(\sum_{i_p \in [d]} x_{i_p} y_{i_p} \right) \quad (13)$$

$$= (x^\top y)^p. \quad (14)$$

- With $\phi = \text{TPOW}(\cdot, p)$, $(q^\top k)^p = \phi(q)^\top \phi(k)$.
- Power attention fits linear attention with state expansion $D = d^p$.

Appendix B.2: Symmetric power (SPOW) definition

Index SPOW by a multi-index $\alpha = (\alpha_1, \dots, \alpha_d) \in \mathbb{N}_0^d$ with $|\alpha| := \sum_{i=1}^d \alpha_i = p$.

$$\text{SPOW}_p(x)_\alpha := \sqrt{\frac{p!}{\alpha_1! \cdots \alpha_d!}} \prod_{i=1}^d x_i^{\alpha_i}, \quad \alpha \in \mathbb{N}_0^d, |\alpha| = p. \quad (27)$$

- The number of such α is $D = \binom{d+p-1}{p}$.
- This removes permutation duplicates of ordered indices.

Appendix B.2: Inner product identity (multinomial theorem)

Let $x, y \in \mathbb{R}^d$. Then

$$\langle \text{SPOW}_p(x), \text{SPOW}_p(y) \rangle = \sum_{|\alpha|=p} \left(\sqrt{\frac{p!}{\alpha!}} \prod_{i=1}^d x_i^{\alpha_i} \right) \left(\sqrt{\frac{p!}{\alpha!}} \prod_{i=1}^d y_i^{\alpha_i} \right) \quad (28)$$

$$= \sum_{|\alpha|=p} \frac{p!}{\alpha!} \prod_{i=1}^d (x_i y_i)^{\alpha_i} \quad (29)$$

$$= \left(\sum_{i=1}^d x_i y_i \right)^p \quad (30)$$

$$= (x^\top y)^p. \quad (31)$$

Here $\alpha! := \alpha_1! \cdots \alpha_d!$.

Appendix B.2: Conclusion

$$\langle \text{SPOW}_p(x), \text{SPOW}_p(y) \rangle = (x^\top y)^p = \langle \text{TPOW}_p(x), \text{TPOW}_p(y) \rangle. \quad (32)$$

- SPOW_p can replace TPOW_p in power attention.
- SPOW_p reduces the feature dimension from d^p to $\binom{d+p-1}{p}$.