

ON LINEAR ATTENTION

Theoretical Perspectives on Architecture Design

Philipp Nazari

June 25, 2026

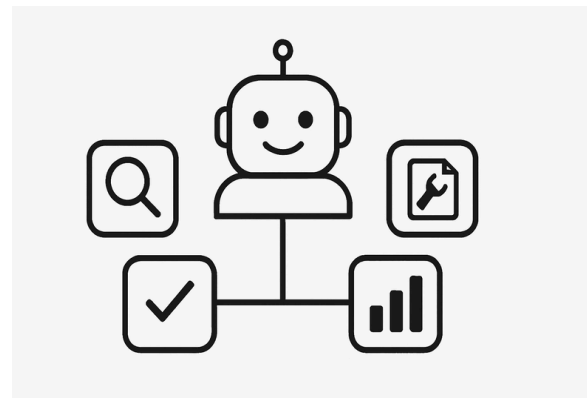
Deployment in resource-constrained environments



Liquid AI and Mercedes-Benz partner to scale embedded in-car intelligence



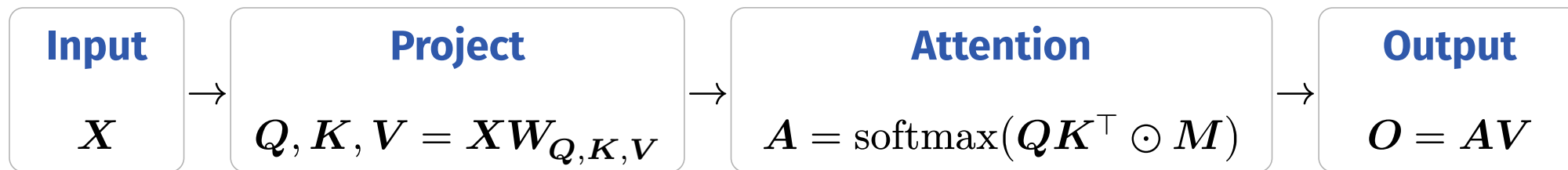
Long context scenarios



Receive a sequence $x_1, x_2, \dots, x_L$ of length L and map it to another sequence:

$$(x_1, \dots, x_L) \rightarrow (o_1, \dots, o_L).$$

A particular instance is **softmax attention**.



During training, use the parallel realization of the softmax operator:

$$\mathbf{O} = \text{softmax}(\mathbf{Q}\mathbf{K}^\top \odot \mathbf{M})\mathbf{V}$$

During autoregressive decode, need to use the recurrent formulation:

$$\mathbf{o}_t = \sum_{i=1}^{t-1} \left(\frac{\exp(\mathbf{q}_t^\top \mathbf{k}_i)}{\sum_{j=1}^t \exp(\mathbf{q}_t^\top \mathbf{k}_j)} \right) \mathbf{v}_i + \left(\frac{\exp(\mathbf{q}_t^\top \mathbf{k}_t)}{\sum_{j=1}^t \exp(\mathbf{q}_t^\top \mathbf{k}_j)} \right) \mathbf{v}_t.$$

During training, use the parallel realization of the softmax operator:

$$\mathbf{O} = \text{softmax}(\mathbf{Q}\mathbf{K}^\top)\mathbf{V}$$

During autoregressive decode, need to use the recurrent formulation:

$$\mathbf{o}_t = \sum_{i=1}^{t-1} \left(\frac{\exp(\mathbf{q}_t^\top \mathbf{k}_i)}{\sum_{j=1}^{t-1} \exp(\mathbf{q}_t^\top \mathbf{k}_j) + \exp(\mathbf{q}_t^\top \mathbf{k}_t)} \right) \mathbf{v}_i + \left(\frac{\exp(\mathbf{q}_t^\top \mathbf{k}_t)}{\sum_{j=1}^{t-1} \exp(\mathbf{q}_t^\top \mathbf{k}_j) + \exp(\mathbf{q}_t^\top \mathbf{k}_t)} \right) \mathbf{v}_t.$$

In practice, one keeps a **KV-cache** to avoid re-computation: $\mathcal{O}(L)$ memory

Can we do better? Yes, with **linear Attention** : $\mathcal{O}(1)$ memory

Remove the softmax operator:¹

$$O = (QK^{\top})V.$$

Softmax form.

$$\text{softmax}(QK^{\top})V$$

The nonlinear row-wise normalization blocks reassociation.

Linear form.

$$Q(K^{\top}V)$$

The history can be summarized before seeing the query.

¹Katharopoulos et al., *Transformers are RNNs: Fast Autoregressive Transformers with Linear Attention*, ICML 2020.

At timestep t ,

$$o_t = \sum_{i \leq t} v_i (k_i^\top q_t) = \left(\sum_{i \leq t} v_i k_i^\top \right) q_t.$$

Define the hidden state

$$S_t := \sum_{i \leq t} v_i k_i^\top \in \mathbb{R}^{d_v \times d_k}.$$

which yields the recursion

$$S_t = S_{t-1} + v_t k_t^\top, \quad o_t = S_t q_t.$$

Linear Attention is a **linear RNN!**

Linear attention writes

$$S_t = S_{t-1} + v_t k_t^\top.$$

What can go wrong?

Memory conflict! Two timesteps writing into the same memory

If a later token has a similar key, the model writes another value into nearly the same memory direction.

Read the current memory at key k_t :

$$\hat{v}_t = S_{t-1} k_t.$$

The prediction error is

$$e_t = v_t - S_{t-1} k_t.$$

DeltaNet writes this residual:

$$S_t = S_{t-1} + e_t k_t^\top = S_{t-1} (I - k_t k_t^\top) + v_t k_t^\top.$$

Plain linear attention

$$S_t = S_{t-1} + v_t k_t^\top$$

Model	Memory update	Adds
Linear Attention ¹	$S_t = S_{t-1} + v_t k_t^\top$	write
Mamba-2 ³	$S_t = \alpha_t S_{t-1} + v_t k_t^\top$	forget
DeltaNet ²	$S_t = S_{t-1} (I - k_t k_t^\top) + v_t k_t^\top$	erase
Gated DeltaNet ⁴	$S_t = \alpha_t S_{t-1} (I - k_t k_t^\top) + v_t k_t^\top$	forget + erase

¹Katharopoulos et al., ICML 2020.

³Dao & Gu, *Transformers are SSMS*, 2024.

²Schlag et al., ICML 2021.

⁴Yang, Kautz, Hatamizadeh, *Gated Delta Networks*, 2024/2025.

OUR RESEARCH: EFFICIENT EFFICIENT ARCHITECTURES

Linear Recurrent Unit⁵:

$$h_t = Ah_{t-1} + Bx_t \in \mathbb{R}^d, \quad y_t = Ch_t$$

h_t acts as a fixed-size **summary of the past**

- **Questions:**
 1. How well does model **utilize** this fixed-size hidden state?
 2. Can we make it more **efficient**?
- We want **large states** for better memory, but **small states** for efficiency.

⁵A. Orvieto et al. Resurrecting Recurrent Neural Networks for Long Sequences.

- **identify** and **truncate** unused state dimensions during training

6

- Solve the discrete Lyapunov equations for the **Gramians**:

$$APA^\top + BB^\top = P, \quad A^\top QA + C^\top C = Q$$

- $P = \mathbf{Controllability}$ (input - hidden state)
- $Q = \mathbf{Observability}$ (hidden state - output)

A state dimension that is hard to reach **and** hard to observe is essentially dead weight.

⁶Makram Chahine, **Philipp Nazari**, Daniela Rus, T. Konstantin Rusch. *CompreSSM: The Curious Case of In-Training Compression of State Space Models*. arXiv preprint arXiv:2510.02823 (2025).

- **Hankel Singular Values (HSVs)** combine both perspectives into a single importance ranking:

$$\sigma_i = \sqrt{\lambda_i(PQ)}, \quad \sigma_1 \geq \sigma_2 \geq \cdots \geq \sigma_n$$

- **Balanced realization:** transform coordinates. States with small σ_i are both hard to control and hard to observe.
- **Truncation:** discard the tail of the balanced system.

- **The Pipeline.** During the **first 10%** of training:
 1. Extract matrices and compute HSVs periodically.
 2. Sort and truncate the state dimension r based on an energy threshold.
 3. Replace weights and continue training.

- **performance close to a full-size model** at the **speed and memory footprint** of a much smaller one.

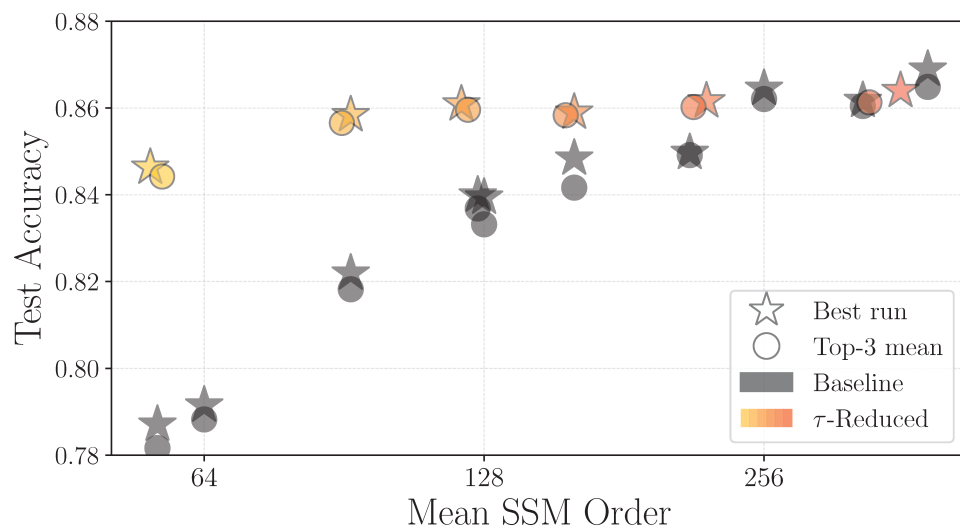


Figure 1: Test Acc vs. Final Dim

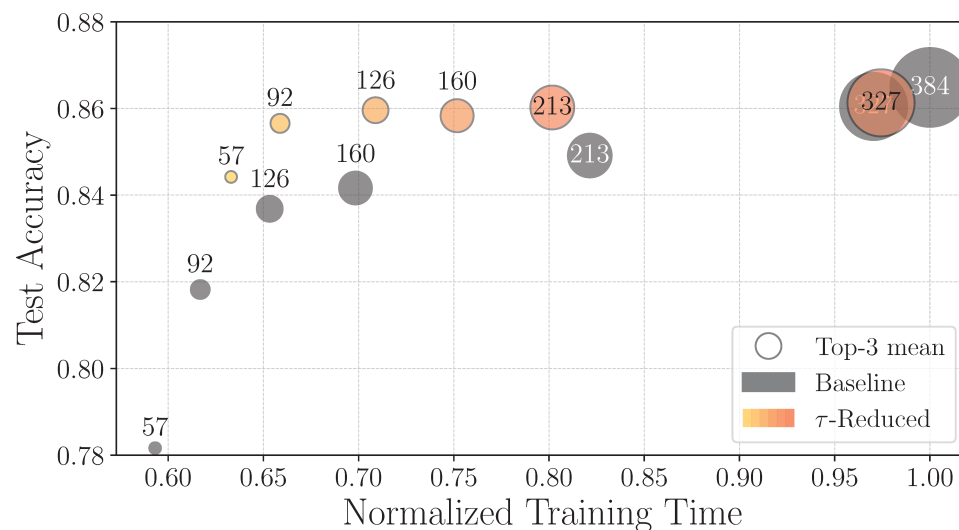


Figure 2: Test Acc vs. Train Time

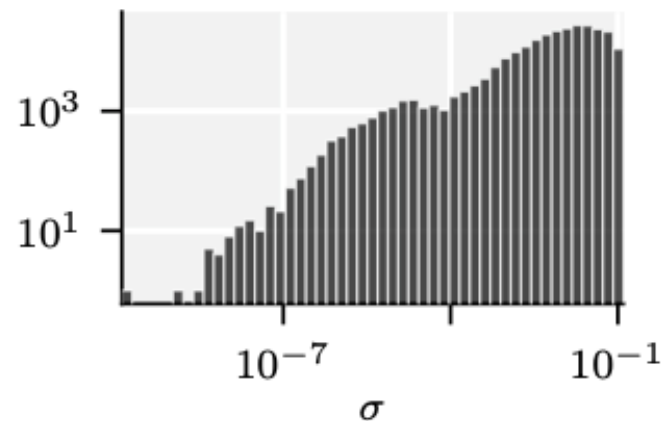
In trained models, spectrum of S_t is often heavy-tailed.⁷

For plain linear attention,

$$S_t = \sum_{i \leq t} v_i k_i^\top.$$

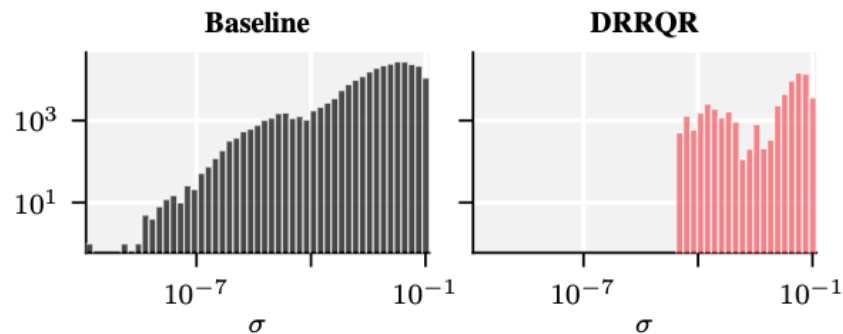
Therefore

$$\text{rank}(S_t) \leq \min(\text{rank}(V_t), \text{rank}(K_t)).$$

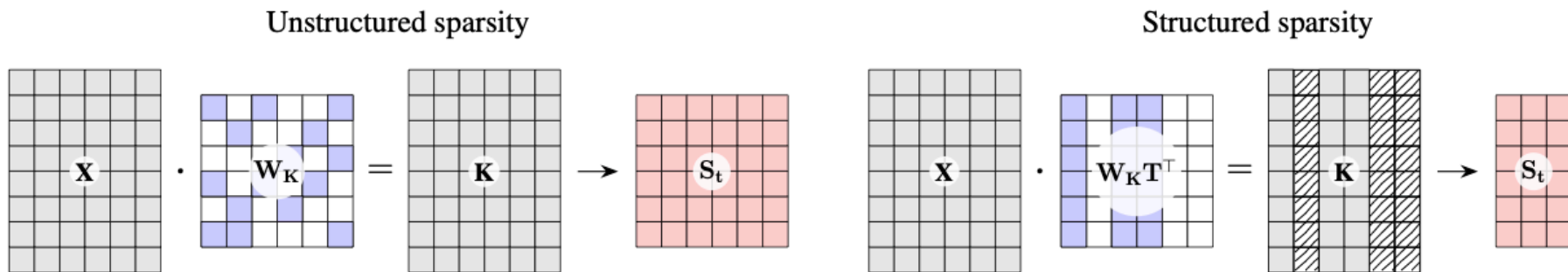


⁷**Nazari** & Rusch, *The Key to State Reduction in Linear Attention: A Rank-Based Perspective*, 2026.

Prune query/key dimension in a rank-aware manner: choose a subset of key/query channels shrinking the state to



$$\mathbf{S}_t \in \mathbb{R}^{d_v \times d'_k} \text{ with } d_{k'} < d_k$$



Result

Around **50% of key/query channels** can be removed with only marginal perplexity cost, while reducing state size and improving sequence-mixer throughput.

Table 1: Comprehensive post-RFT evaluation of Gated DeltaNet 1.3B. We report Wikitext and Lambada perplexities, the average zero-shot generation accuracy on language-model evaluations (**LM evals**, [16]) across varying compression ratios (best results in **bold**, second best underlined).

METHOD	75%			50%			40%			30%		
	WIKI ↓	LMB ↓	LM EVALS ↑	WIKI ↓	LMB ↓	LM EVALS ↑	WIKI ↓	LMB ↓	LM EVALS ↑	WIKI ↓	LMB ↓	LM EVALS ↑
RAND	19.3	18.1	54.6	16.9	<u>11.5</u>	57.8	16.5	<u>10.7</u>	58.3	<u>16.1</u>	<u>10.2</u>	58.7
L1	18.6	16.8	55.3	16.8	13.0	57.2	16.5	11.9	57.8	16.3	11.3	58.3
S-WANDA	18.2	15.1	55.4	16.6	11.7	57.8	16.4	11.5	58.1	16.1	11.0	58.6
GRAD	17.8	12.5	56.8	<u>16.4</u>	10.5	58.3	16.1	10.1	<u>58.6</u>	15.9	9.9	59.0
DRRQR	<u>17.9</u>	<u>14.7</u>	<u>56.2</u>	16.3	12.0	<u>58.0</u>	16.1	11.0	58.7	15.9	10.7	<u>58.8</u>
BASLINE	16.8	9.7	59.4	16.8	9.7	59.4	16.8	9.7	59.4	16.8	9.7	59.4

THANKS! QUESTIONS?

- Vaswani et al., *Attention Is All You Need*, NeurIPS 2017.
- Katharopoulos et al., *Transformers are RNNs: Fast Autoregressive Transformers with Linear Attention*, ICML 2020.
- Schlag, Irie, Schmidhuber, *Linear Transformers Are Secretly Fast Weight Programmers*, ICML 2021.
- Dao & Gu, *Transformers are SSMS: Generalized Models and Efficient Algorithms through Structured State Space Duality*, 2024.
- Yang, Kautz, Hatamizadeh, *Gated Delta Networks: Improving Mamba2 with Delta Rule*, 2024/2025.
- Chahine, Nazari, Rus, Rusch, *The Curious Case of In-Training Compression of State Space Models*, ICLR 2026.
- Nazari & Rusch, *The Key to State Reduction in Linear Attention: A Rank-Based Perspective*, 2026.