

# ON THE STATE OF STATE SPACE MODELS

---

By Philipp Nazari

# **1. WHAT ARE STATE-SPACE MODELS?**

---

- Continuous dynamical system:

$$\dot{\mathbf{h}}(t) = \mathbf{A}\mathbf{h}(t) + \mathbf{B}\mathbf{x}(t)$$

$$\mathbf{y}(t) = \mathbf{C}\mathbf{h}(t)$$

- Discretization yields a **Linear RNN**:

$$\mathbf{h}_t = \tilde{\mathbf{A}}\mathbf{h}_{t-1} + \tilde{\mathbf{B}}\mathbf{x}_t$$

- where  $\tilde{\mathbf{A}} = \exp(\Delta\mathbf{A})$  and  $\tilde{\mathbf{B}} = \mathbf{A}^{-1}(\tilde{\mathbf{A}} - \mathbf{I})\mathbf{B}$ .

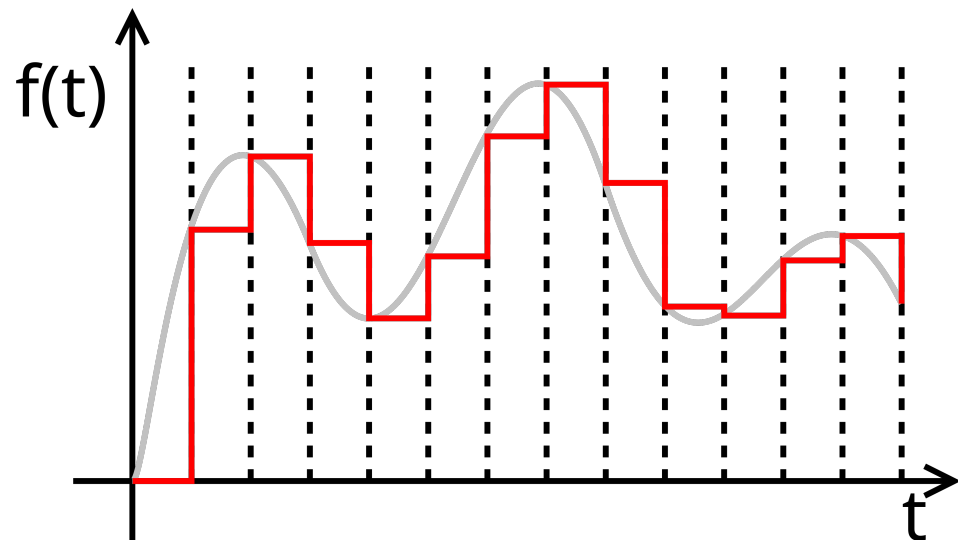


Figure 1: Zero-Order Hold discretization

LRU<sup>1</sup>:

$$\mathbf{h}_t = \text{diag}(e^{i\theta_k + \lambda_k}) \mathbf{h}_{t-1} + \gamma \mathbf{B} \mathbf{x}_t \in \mathbb{C}^d$$

$\mathbf{h}_t$  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.

---

<sup>1</sup>A. Orvieto et al. Resurrecting Recurrent Neural Networks for Long Sequences.

### **3. COMPRESSM: IN-TRAINING COMPRESSION**

---

- **identify** and **truncate** unused state dimensions during training<sup>2</sup>.
- Solve the discrete Lyapunov equations for the **Gramians**:

$$\tilde{A}P\tilde{A}^\top + \tilde{B}\tilde{B}^\top = P, \quad \tilde{A}^\top Q\tilde{A} + C^\top C = Q$$

- $P = \mathbf{Controllability}$  (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.

---

<sup>2</sup>Makram 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:

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

- **Balanced realization:** transform coordinates. States with small  $\sigma_i$  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.

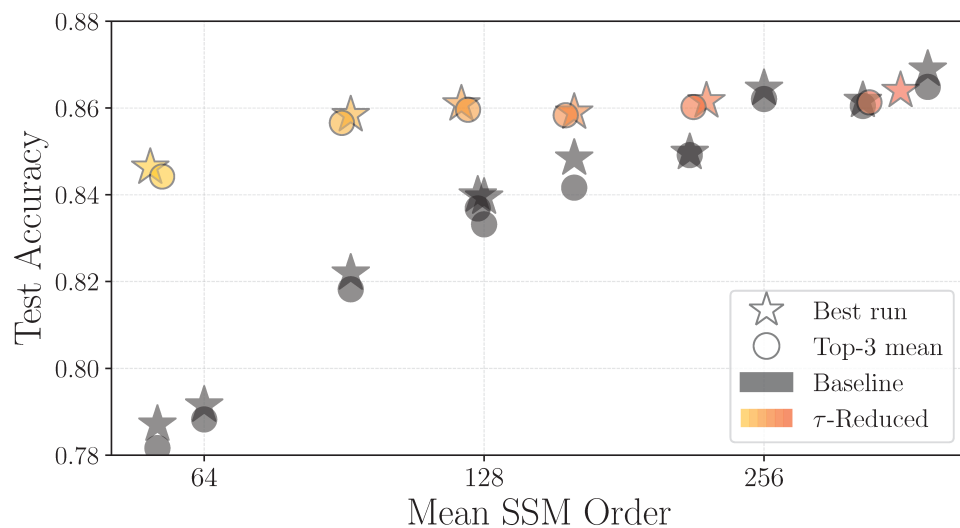


Figure 2: Test Acc vs. Final Dim

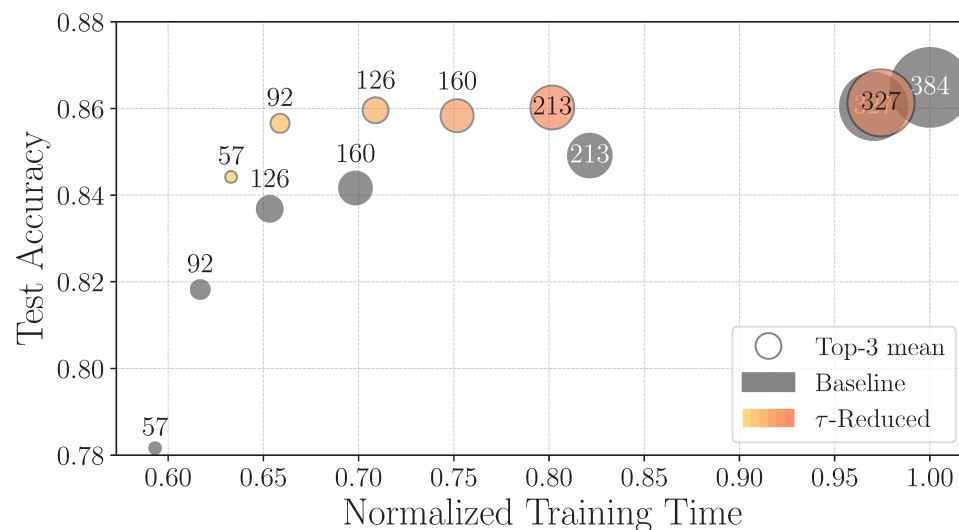


Figure 3: Test Acc vs. Train Time

# **5. THE KEY TO STATE REDUCTION IN LINEAR ATTENTION**

---

- **Drop the softmax.** The output at position  $t$  becomes:

$$o_t = \sum_{s=1}^t (q_t^\top k_s) v_s = q_t^\top \underbrace{\sum_{s=1}^t k_s v_s^\top}_{S_t \in \mathbb{R}^{d_k \times d_v}}$$

- The past is compressed into a **fixed-size state**  $S_t$ :

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

- **Observation:**  $S_t \in \mathbb{R}^{d_k \times d_v}$  exhibit a distinct **low-rank structure**<sup>3</sup>.
- **Post-training reduction:** faster and more memory-efficient models.
- **Method:** structured pruning method that **jointly selects query- and key-channels**.

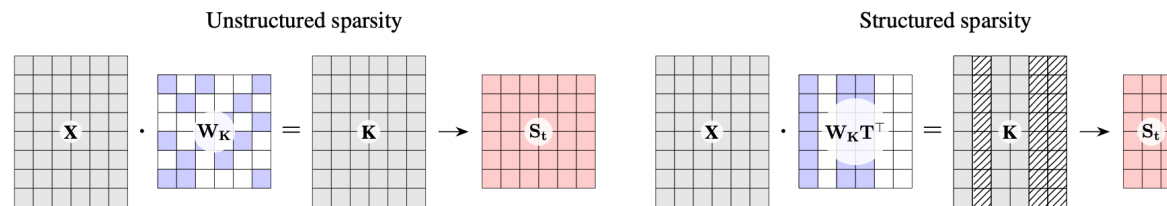


Figure 1. **Left:** Unstructured pruning yields a sparse weight matrix  $\mathbf{W}_K$  yet preserves the column dimension of  $\mathbf{K}$ , leaving the domain of the state matrix  $\mathbf{S}_t \in \mathbb{R}^{d_v, d_k}$  invariant. **Right:** Structured pruning eliminates basis vectors, mapping keys to a lower-dimensional space  $\mathbb{R}^{d'_k}$  where  $d'_k < d_k$ . This results in a compressed state  $\mathbf{S}_t \in \mathbb{R}^{d_v, d'_k}$ . This reduction strictly decreases the FLOP count required to compute the recurrence. This figure is inspired by Ashkboos et al. (2024a, Figure 1).

---

<sup>3</sup>**Philipp Nazari**, T. Konstantin Rusch. *The Key to State Reduction in Linear Attention: A Rank-based Perspective*. arXiv preprint arXiv:2602.04852 (2026).

## Can remove $\approx 50\%$ of the query and key channels

Table 1. Comprehensive post-RFT evaluation of Gated DeltaNet 1.3B. We report Wikitext and Lambada perplexities, the average zero-shot generation accuracy (**ZS**, Gao et al. (2024)) as well as average retrieval accuracy (**Ret**, Arora et al. (2024)) across varying compression ratios (best results in **bold**, second best underlined).

| METHOD  | 75% COMPRESSION |             |             |             | 50% COMPRESSION |             |             |             | 40% COMPRESSION |             |             |             | 30% COMPRESSION |             |             |             |
|---------|-----------------|-------------|-------------|-------------|-----------------|-------------|-------------|-------------|-----------------|-------------|-------------|-------------|-----------------|-------------|-------------|-------------|
|         | WIKI<br>↓       | LMB<br>↓    | ZS<br>↑     | RET<br>↑    | WIKI<br>↓       | LMB<br>↓    | ZS<br>↑     | RET<br>↑    | WIKI<br>↓       | LMB<br>↓    | ZS<br>↑     | RET<br>↑    | WIKI<br>↓       | LMB<br>↓    | ZS<br>↑     | RET<br>↑    |
| RAND    | 19.3            | 18.1        | 54.6        | 23.9        | 16.9            | <u>11.5</u> | 57.8        | 32.1        | 16.5            | <u>10.7</u> | 58.3        | 34.6        | <u>16.1</u>     | <u>10.2</u> | 58.7        | <u>37.3</u> |
| L1      | 18.6            | 16.8        | 55.3        | 25.7        | 16.8            | 13.0        | 57.2        | 32.9        | 16.5            | 11.9        | 57.8        | 34.6        | 16.3            | 11.3        | 58.3        | 35.4        |
| S-WANDA | 18.2            | 15.1        | 55.4        | 26.2        | 16.6            | 11.7        | 57.8        | <u>33.2</u> | 16.4            | 11.5        | 58.1        | 34.4        | 16.1            | 11.0        | 58.6        | 37.2        |
| GRAD    | <b>17.8</b>     | <b>12.5</b> | <b>56.8</b> | <b>26.9</b> | <u>16.4</u>     | <b>10.5</b> | <b>58.3</b> | <b>34.3</b> | <b>16.1</b>     | <b>10.1</b> | <u>58.6</u> | <b>35.8</b> | <b>15.9</b>     | <b>9.9</b>  | <b>59.0</b> | <b>38.0</b> |
| DRRQR   | <u>17.9</u>     | <u>14.7</u> | <u>56.2</u> | <u>26.1</u> | <b>16.3</b>     | 12.0        | <u>58.0</u> | 33.0        | <b>16.1</b>     | 11.0        | <b>58.7</b> | <u>35.0</u> | <b>15.9</b>     | 10.7        | <u>58.8</u> | 36.5        |
| BASLINE | 16.8            | 9.7         | 59.4        | 40.3        | —————           |             |             |             | —————           |             |             |             | —————           |             |             |             |

Figure 4: Perplexity and zero-shot performance

Thank You! Questions?