

IN-MEMORY COMPUTING FOR LARGE LANGUAGE MODELS AT THE EDGE

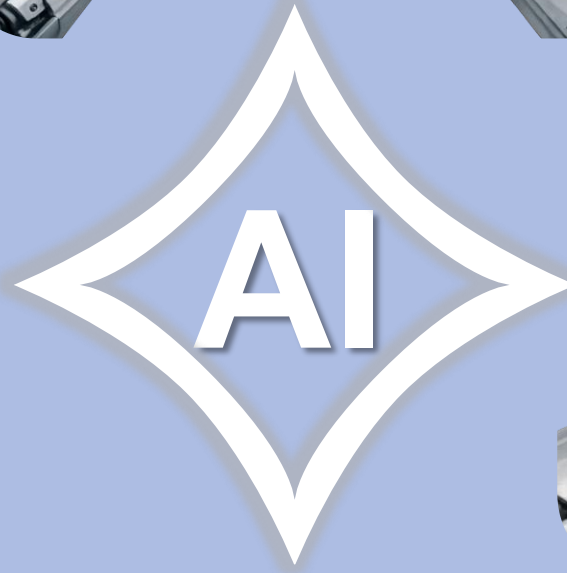
NEUROMORPHIC HARDWARE AND ALGORITHMS WORKSHOP AT THE UNIVERSITY OF SUSSEX

11TH OF JUNE, 2026 – DR.-ING. SEBASTIAN SIEGEL
PETER-GRÜNBERG-INSTITUTE (PGI-14)
FORSCHUNGSZENTRUM JÜLICH GMBH

Software development



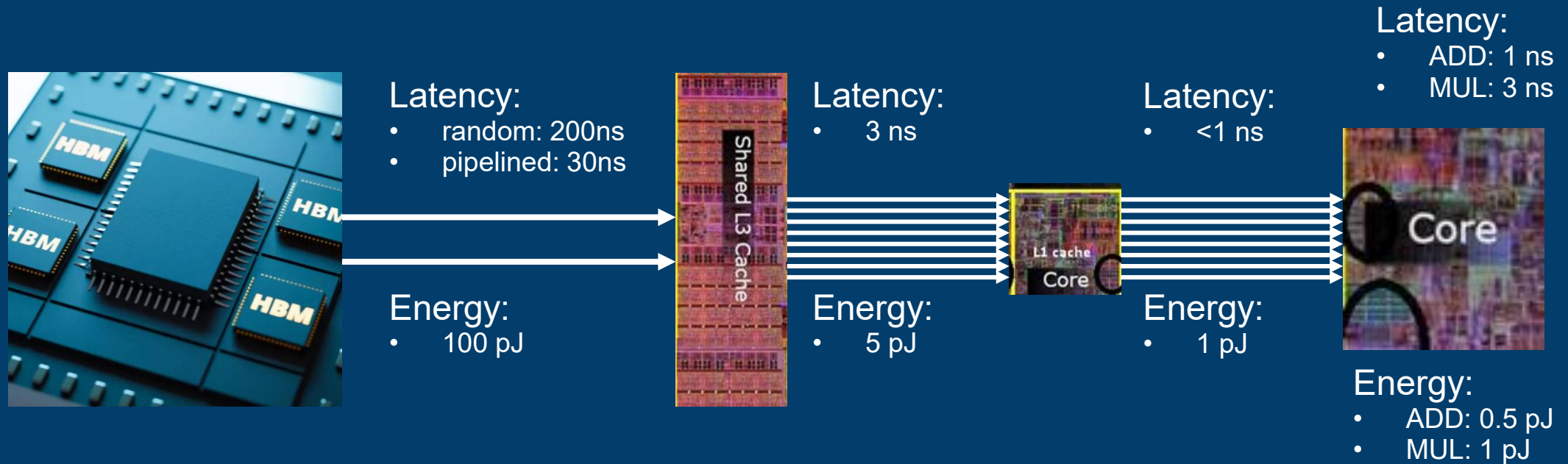
Knowledge retrieval



Artistic creation

NEURAL NETWORK INFERENCE

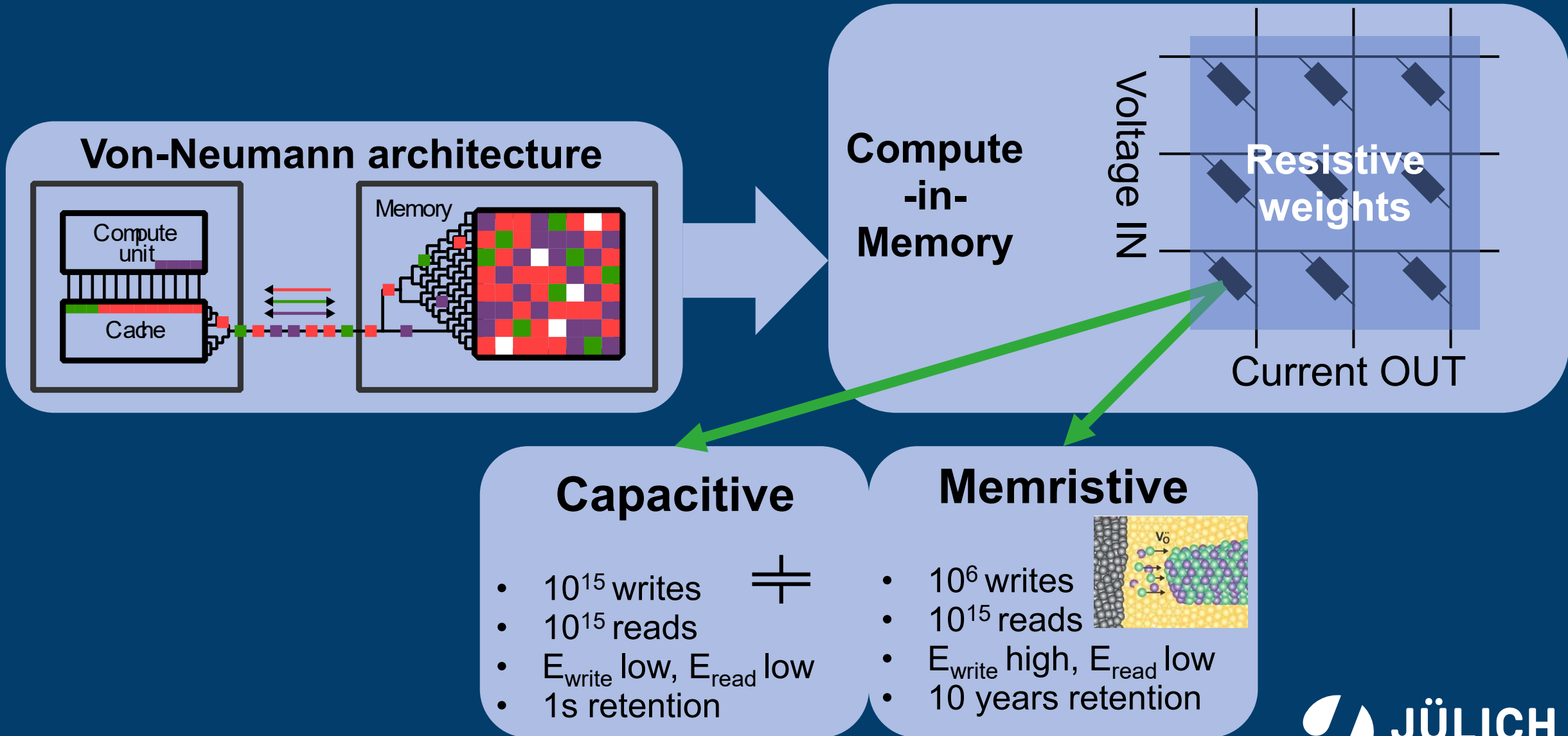
Recap of memory hierarchy in conventional “von-Neumann” computers



All assumptions for single FP16 value

Member of the Helmholtz Association

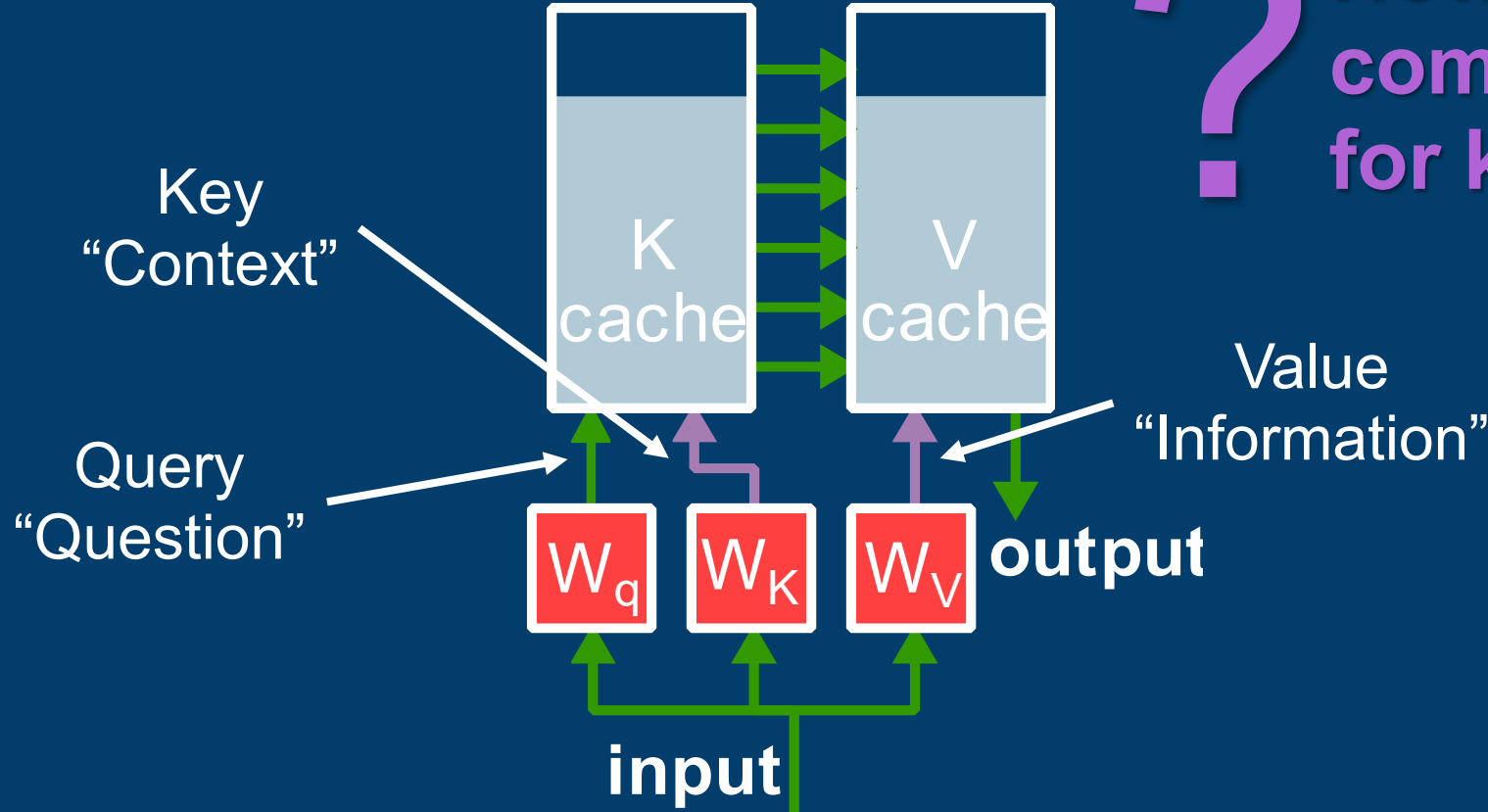
COMPUTE-IN-MEMORY IS ALL YOU NEED?



THE TRANSFORMER MODEL

The backbone of modern AI algorithms

How can we realize
compute-in-memory
for key / value caches?



COMPUTE-IN-MEMORY IS ALL YOU NEED?

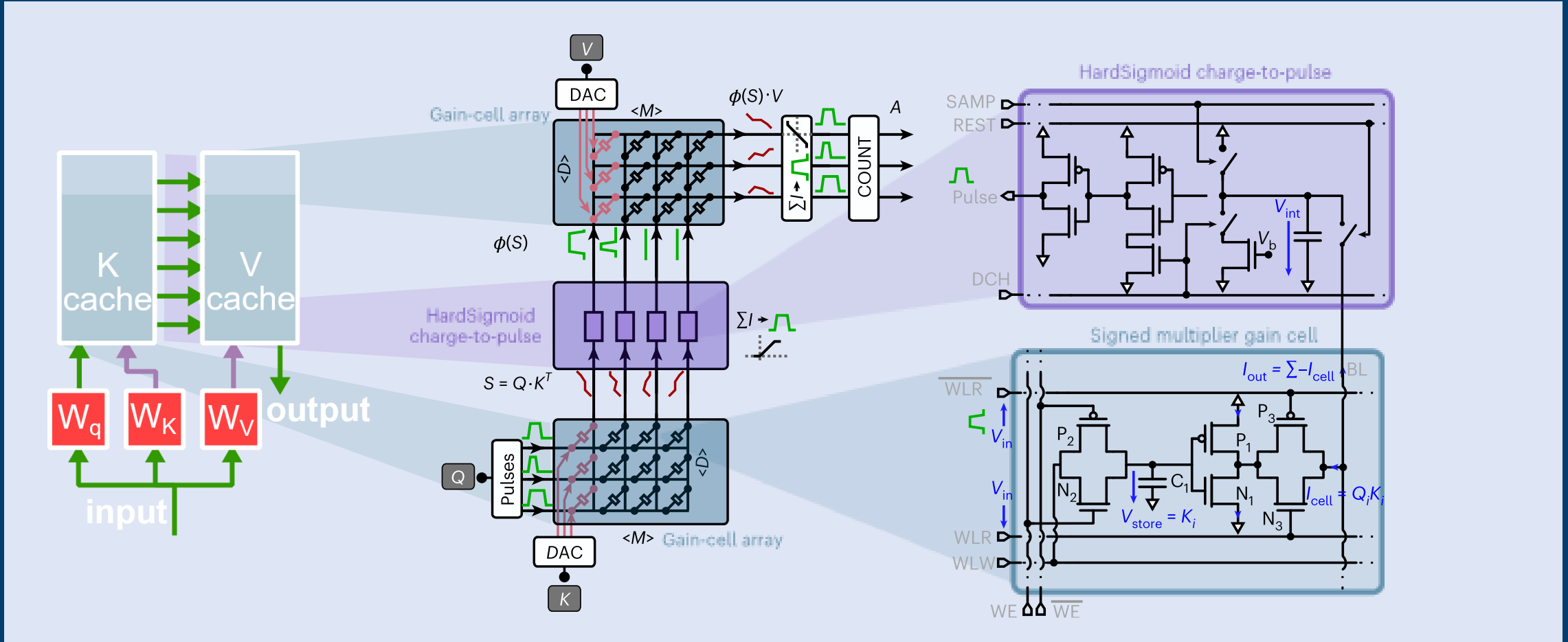
MLP / CNN			Self-Attention		RNN / SSM	
Memory write: Once			Once per sample		Once per token	
	capacitive	memristive	capacitive	memristive	capacitive	memristive
Write lifetime	190B years	10 years	190B years	1.9 years	19M years	17 hours
Read energy	Low	Low	Low	Low	Low	Low
Write energy	High	Very Low	Low to moderate	Moderate	Moderate	Extremely High
Limiting factor	Retention	None! Match!	None! Match!	Possible, but write limited	Potential match!	Write limited

Use case:

- 1 sample per minute, 1000 token per sample

IN-MEMORY COMPUTING FOR TRANSFORMERS

Gain-cell memory for KV-caches

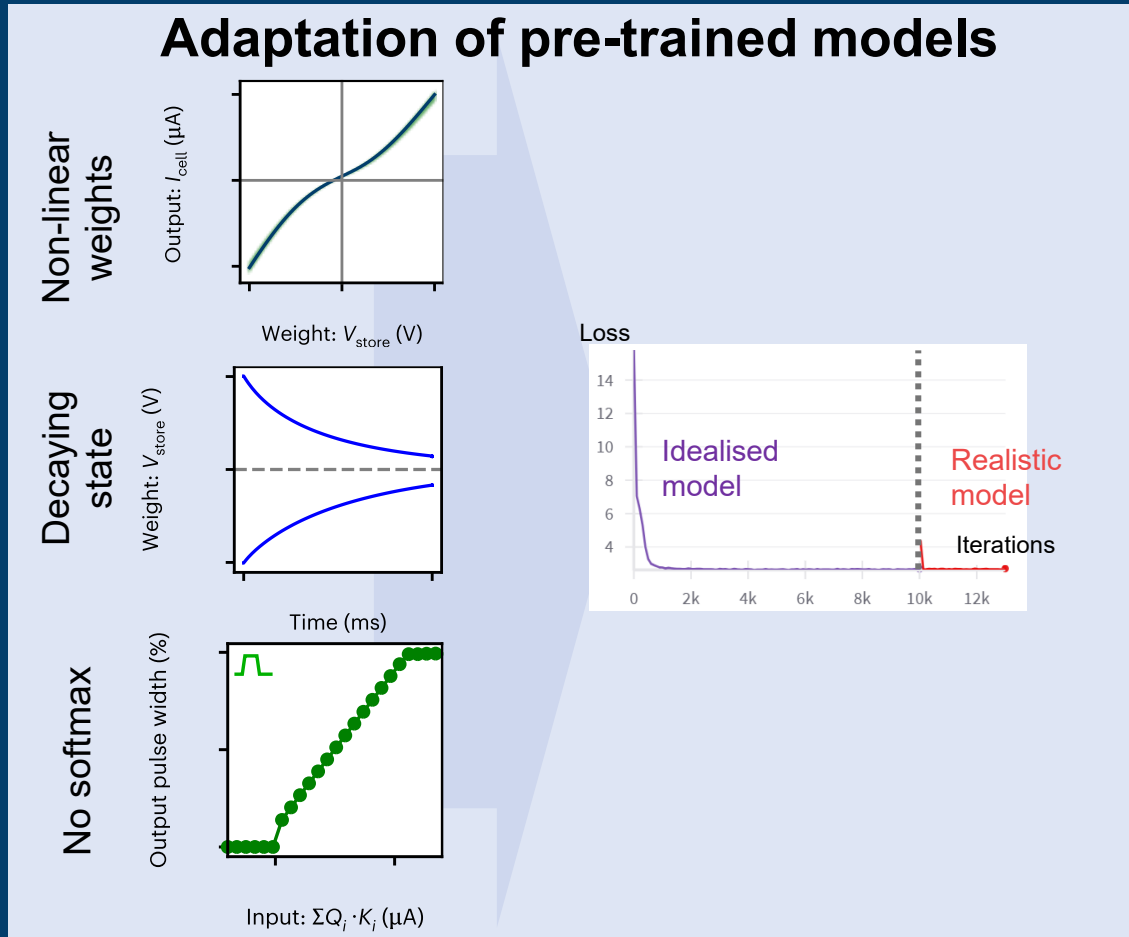


Leroux, N., Manea, P. P., Sudarshan, C., Finkbeiner, J., Siegel, S., Strachan, J. P., & Neftci, E. (2025). Analog in-memory computing attention mechanism for fast and energy-efficient large language models. *Nature Computational Science*, 5(9), 813-824.

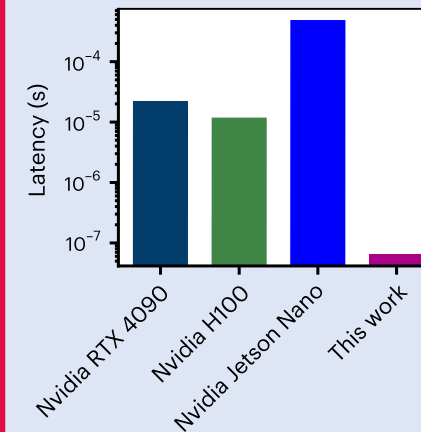
Member of the Helmholtz Association

IN-MEMORY COMPUTING FOR TRANSFORMERS

Gain-cell memory for KV-caches

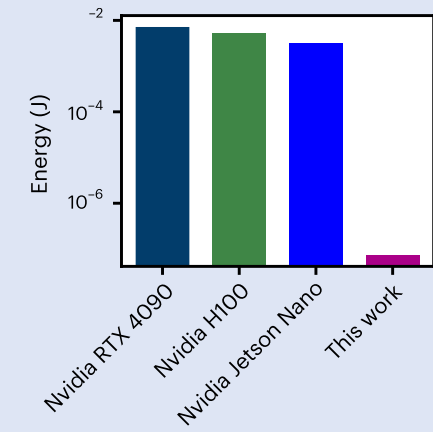


Latency gains



1000x faster

Energy gains



40.000x less energy consumption

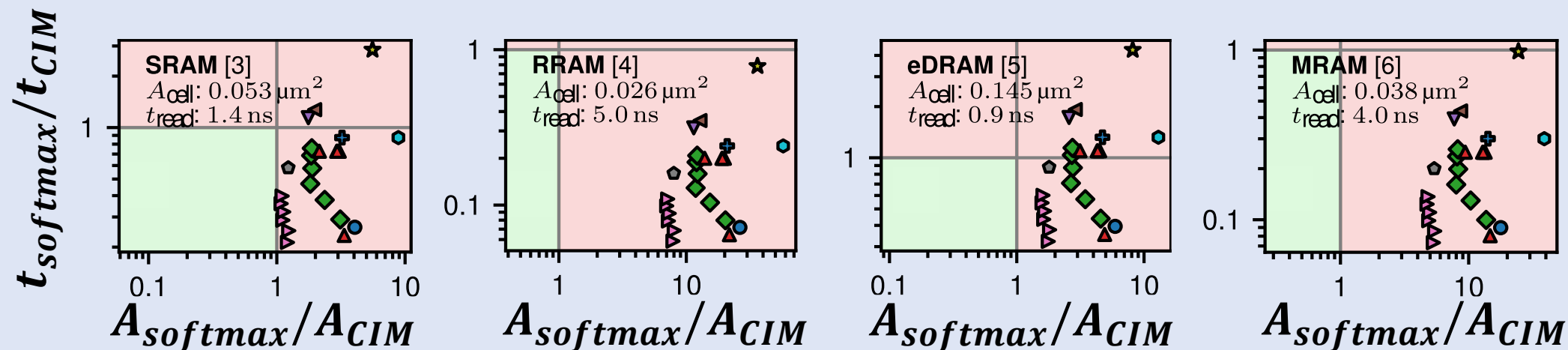
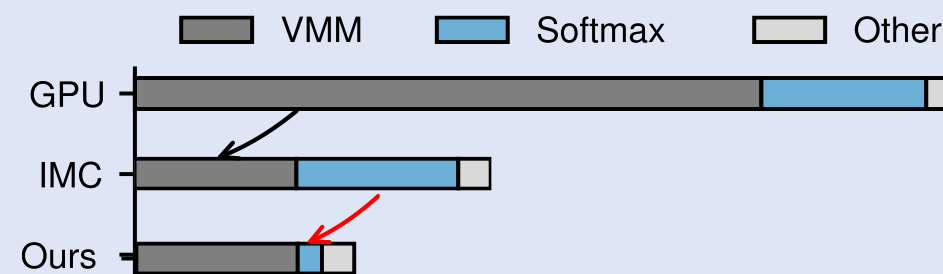
Leroux, N., Manea, P. P., Sudarshan, C., Finkbeiner, J., Siegel, S., Strachan, J. P., & Neftci, E. (2025). Analog in-memory computing attention mechanism for fast and energy-efficient large language models. *Nature Computational Science*, 5(9), 813-824.

Member of the Helmholtz Association

IN-MEMORY COMPUTING FOR TRANSFORMERS

Exact softmax attention for direct Transformer model deployment

$$attention(\vec{q}, K, V) = \frac{\exp(\vec{q}K)}{\sum_{m=0}^N \exp(\vec{q}K)_m} V$$



Finkbeiner, J., Siegel, S., Sudarshan, C., Zhao, Y., Strachan, J. P., & Neftci, E. (2026). Algorithm-Hardware Implications of Softmax Approximations for In-Memory Computing based LLM Accelerators. *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*.

IN-MEMORY COMPUTING FOR TRANSFORMERS

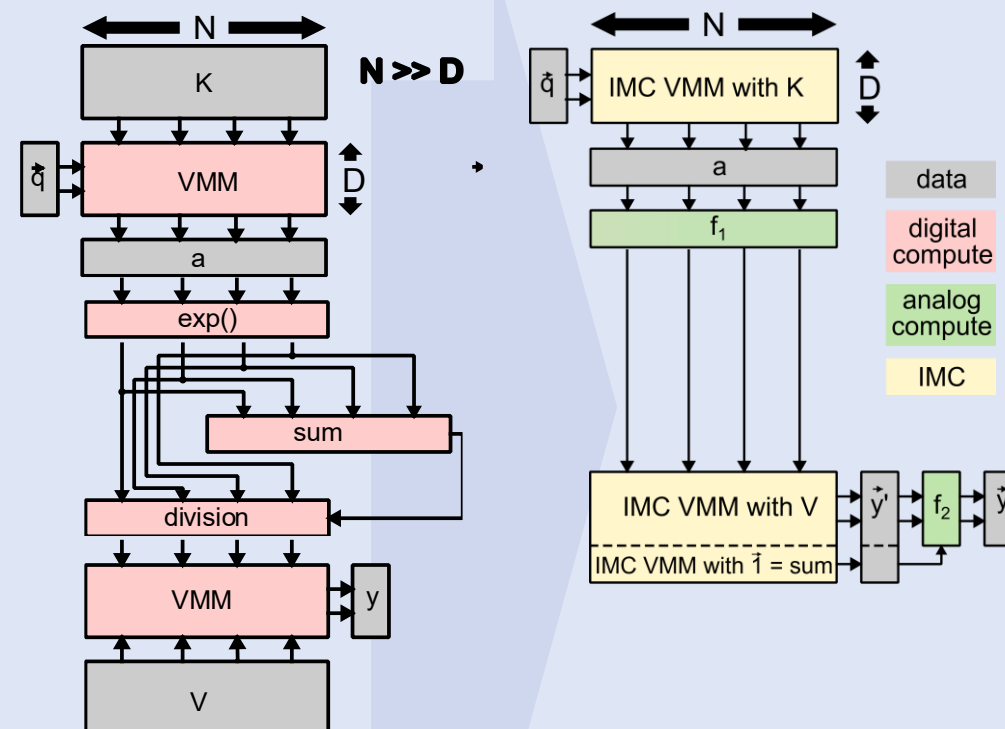
Exact softmax attention for direct Transformer model deployment

Non-linearity

$$attention(\vec{q}, K, V) = \frac{\exp(\vec{q}K)}{\sum_{m=0}^N \exp(\vec{q}K)_m} V$$

Normalization over sequence length

$$attention(\vec{q}, K, V) = f_2(f_1(\vec{q}K) V)$$



IN-MEMORY COMPUTING FOR TRANSFORMERS

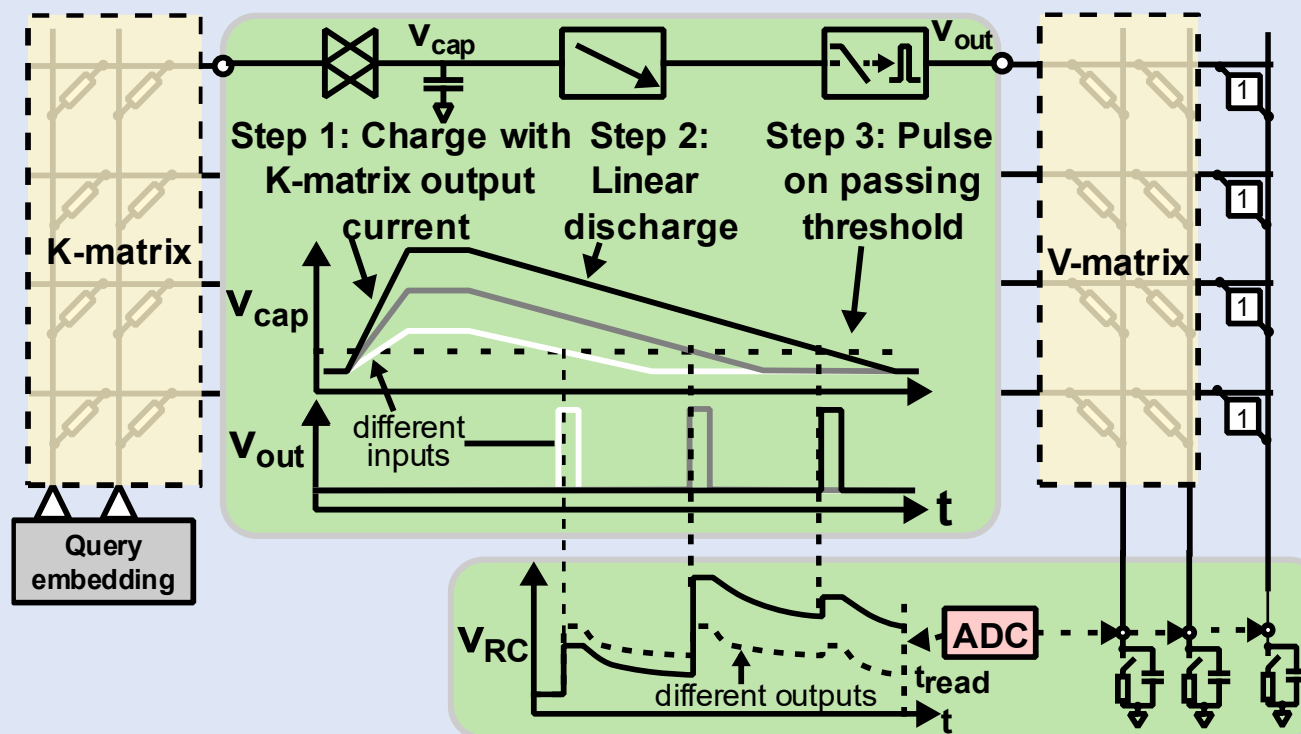
Exact softmax attention for direct Transformer model deployment

$$attention(\vec{q}, \mathbf{K}, \mathbf{V}) = f_2(f_1(\vec{q}\mathbf{K}) \mathbf{V}) = \frac{\exp(\vec{q}\mathbf{K})}{\sum_{m=0}^N \exp(\vec{q}\mathbf{K})_m} \mathbf{V}$$

$$t_{pulse} = t_{max} - d\vec{q}\mathbf{K}^T$$

$$\begin{aligned} \vec{v}_{RC} &= \exp(-(t_{read} + t_{spike})) \mathbf{V} \\ &= \exp(-t_{read} - t_{max} + d\vec{q}\mathbf{K}^T) \mathbf{V} \\ &= const. \exp(d\vec{q}\mathbf{K}^T) \mathbf{V} \\ &= f_1(\vec{q}\mathbf{K}^T) \mathbf{V} \end{aligned}$$

$$f_2 = \frac{f_1 \mathbf{V}}{f_1 * \vec{1}} = attention(\vec{q}, \mathbf{K}, \mathbf{V})$$



Finkbeiner, J., Siegel, S., Sudarshan, C., Zhao, Y., Strachan, J. P., & Neftci, E. (2026).

Algorithm-Hardware Implications of Softmax Approximations for In-Memory Computing based LLM Accelerators.

IEEE Journal on Emerging and Selected Topics in Circuits and Systems.

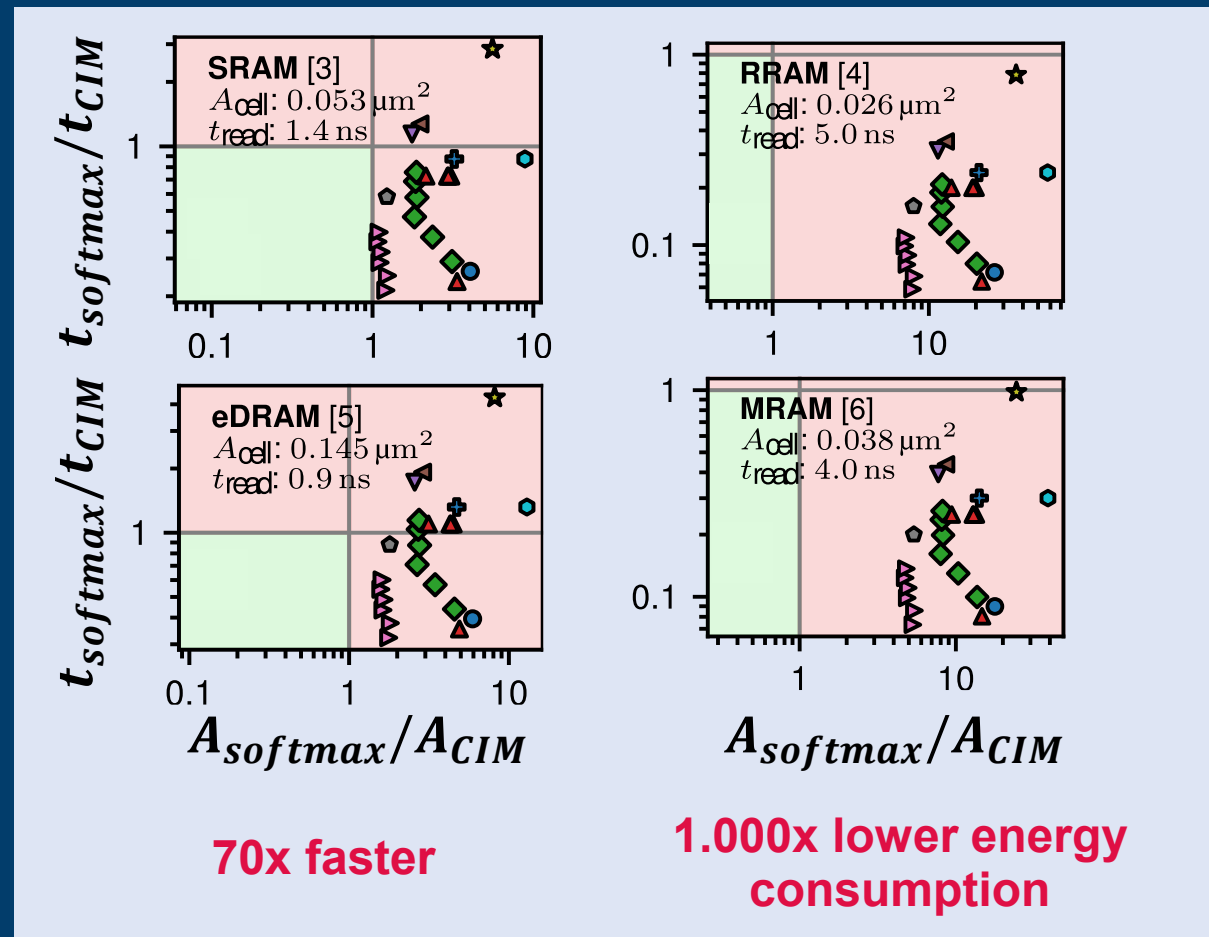
Member of the Helmholtz Association

Page 11

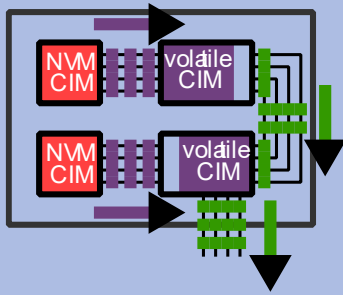
IN-MEMORY COMPUTING FOR TRANSFORMERS

Exact softmax attention for direct Transformer model deployment

Model	Softmax Implementation	WikiText2 ppl ↓	LAMBADA ppl ↓
OPT 125M	Softmax	33.837	27.696
	Smoothquant [38]	48.817	26.679
	Temporal (Analog)	34.354	27.145
	LUT + Shift	33.842	27.588
	1st Order + Shift (Digital)	33.842	27.769
OPT 30B	Softmax	10.757	3.616
	Smoothquant [38]	22.078	3.717
	Temporal (Analog)	10.766	3.640
	LUT + Shift	10.762	3.617
	1st Order + Shift (Digital)	10.758	3.615
Llama3 1B	Softmax	8.630	5.745
	Smoothquant [38]	11.593	5.744
	Temporal (Analog)	8.638	5.888
	LUT + Shift	8.628	5.752
	1st Order + Shift (Digital)	8.636	5.759
Llama3 405B	Softmax	1.257	2.379
	Smoothquant [38]	177.772	2.407
	Temporal (Analog)	1.258	2.376
	LUT + Shift	1.257	2.376
	1st Order + Shift (Digital)	1.256	2.377

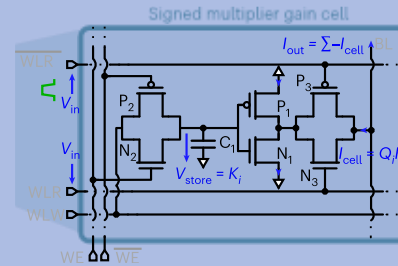


Finkbeiner, J., Siegel, S., Sudarshan, C., Zhao, Y., Strachan, J. P., & Neftci, E. (2026). Algorithm-Hardware Implications of Softmax Approximations for In-Memory Computing based LLM Accelerators. *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*.

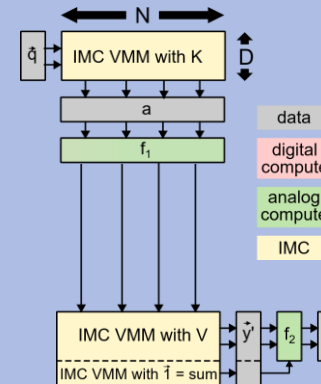


**Compute-in-Memory
can benefit Large
Language Model
deployment**

Leroux, N., Manea, P. P., Sudarshan, C., Finkbeiner, J., Siegel, S., Strachan, J. P., & Neftci, E. (2025). Analog in-memory computing attention mechanism for fast and energy-efficient large language models. *Nature Computational Science*, 5(9), 813-824.



**Capacitive gain-cell
memory
device can
effectively
implement
volatile weight**



**Peripheral
functions must be
merged with
compute-in-
memory to gain
maximum
efficiency**

Finkbeiner, J., Siegel, S., Sudarshan, C., Zhao, Y., Strachan, J. P., & Neftci, E. (2026). Algorithm-Hardware Implications of Softmax Approximations for In-Memory Computing based LLM Accelerators. *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*.

**Thank you for
your attention!**