

CSCI 5541: Natural Language Processing

Lecture 15: LLM Compute efficiency and engineering

James Mooney

With slides borrowed from Song Han (MIT)

What Is Efficiency and Why Does It Matter?

- ❑ Efficiency for NLP is concerned with delivering faster, cheaper, smaller, less energy intensive solutions to problems involving natural language
- ❑ Faster models means LLM model services can meet the demands of many clients more quickly
- ❑ Cheaper models reduce costs for LLM model service providers
- ❑ Smaller model sizes allow for service providers to use fewer resources and can allow for individuals to deploy LLMs to their own (smaller) devices
- ❑ Less energy intensive means lower cost and easier to deploy at the edge, where energy is harder to come by

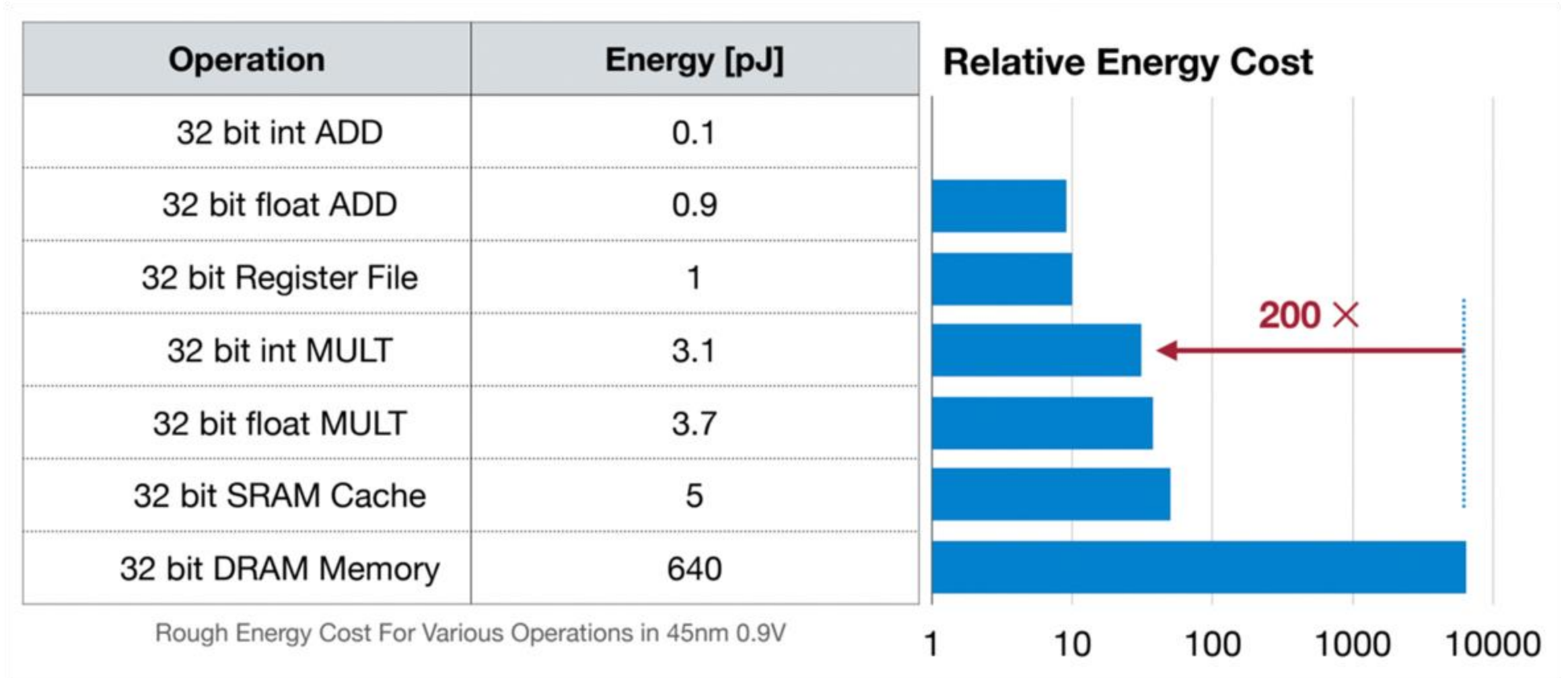


What Is Efficiency and Why Does It Matter?

- ❑ Efficiency for NLP is concerned with delivering **faster, cheaper, smaller, less energy** intensive solutions to problems involving natural language
- ❑ **Faster** models means LLM model services can meet the demands of many clients more quickly
- ❑ **Cheaper** models reduce costs for LLM model service providers
- ❑ **Smaller** model sizes allow for service providers to use fewer resources and can allow for individuals to deploy LLMs to their own (smaller) devices
- ❑ **Less energy** intensive means lower cost and easier to deploy at the edge, where energy is harder to come by



Model Energy Use

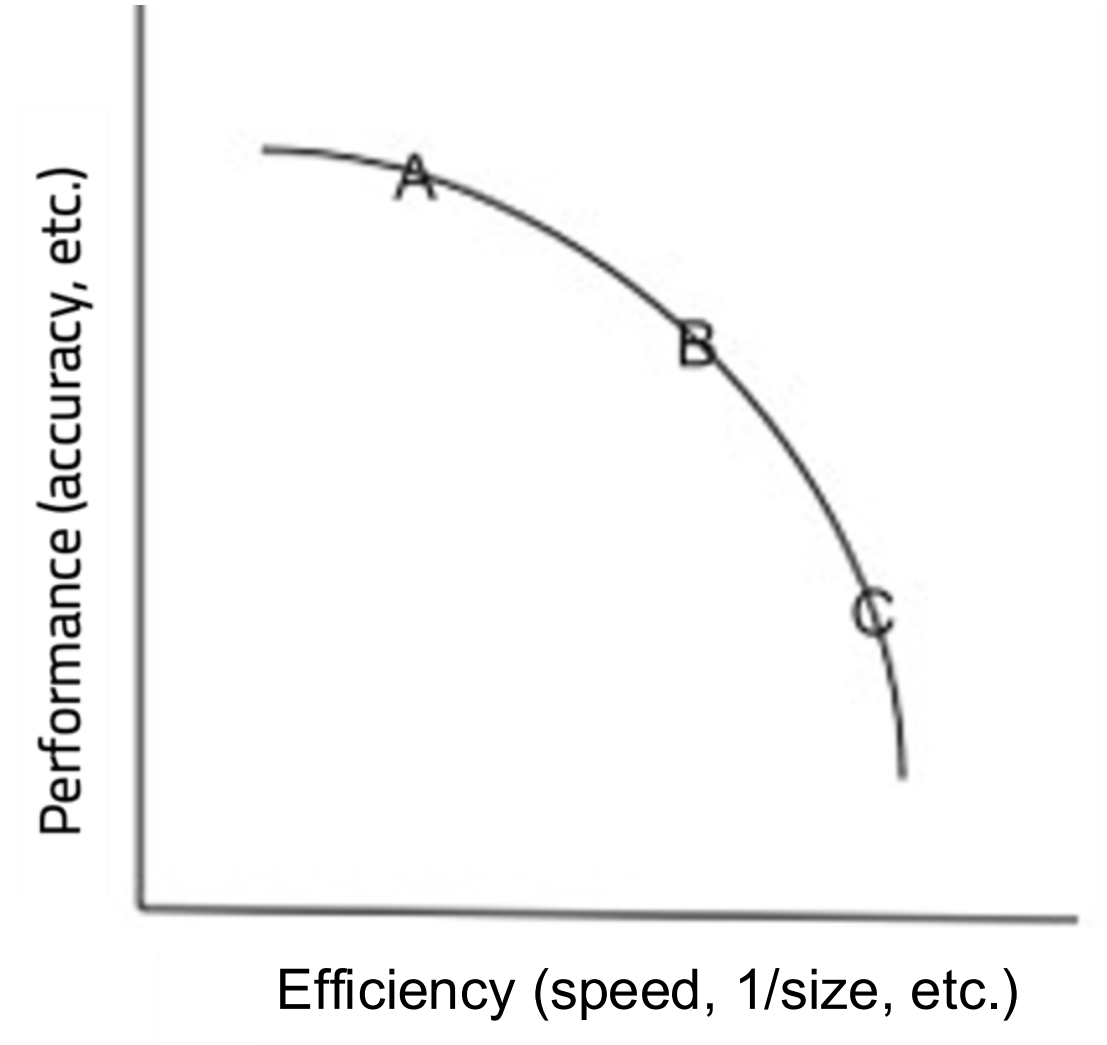


Computing's Energy Problem (and What We Can Do About it) [Horowitz, M., IEEE ISSCC 2014]



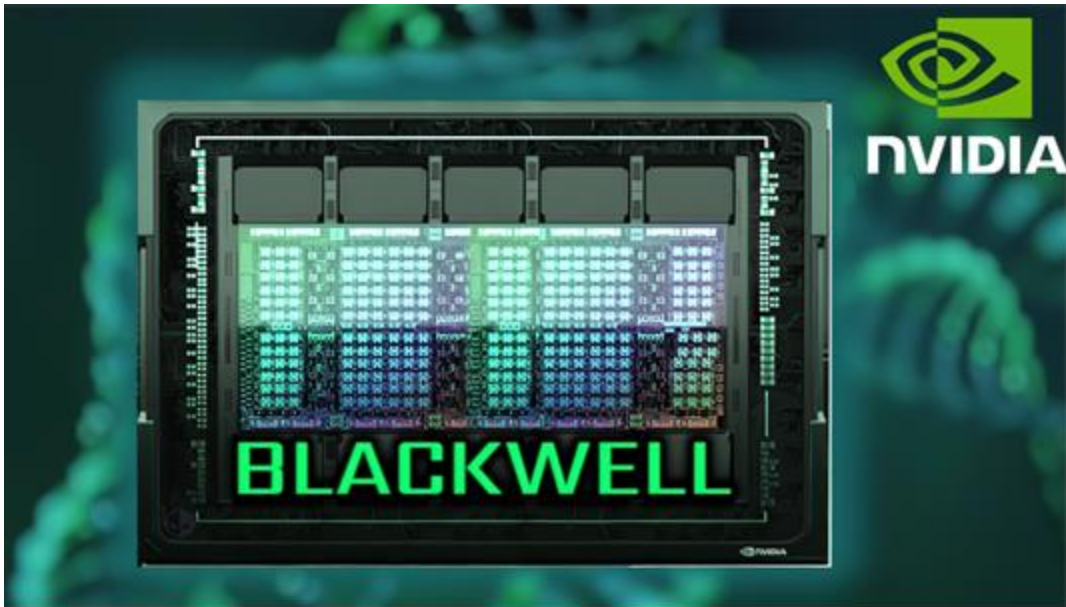
Efficiency Tradeoff

- ❑ More efficient models (smaller, faster) typically come at a cost of some performance of the model itself
- ❑ In the other direction, getting more performance from a model architecture likely means it will be larger, and require more computation

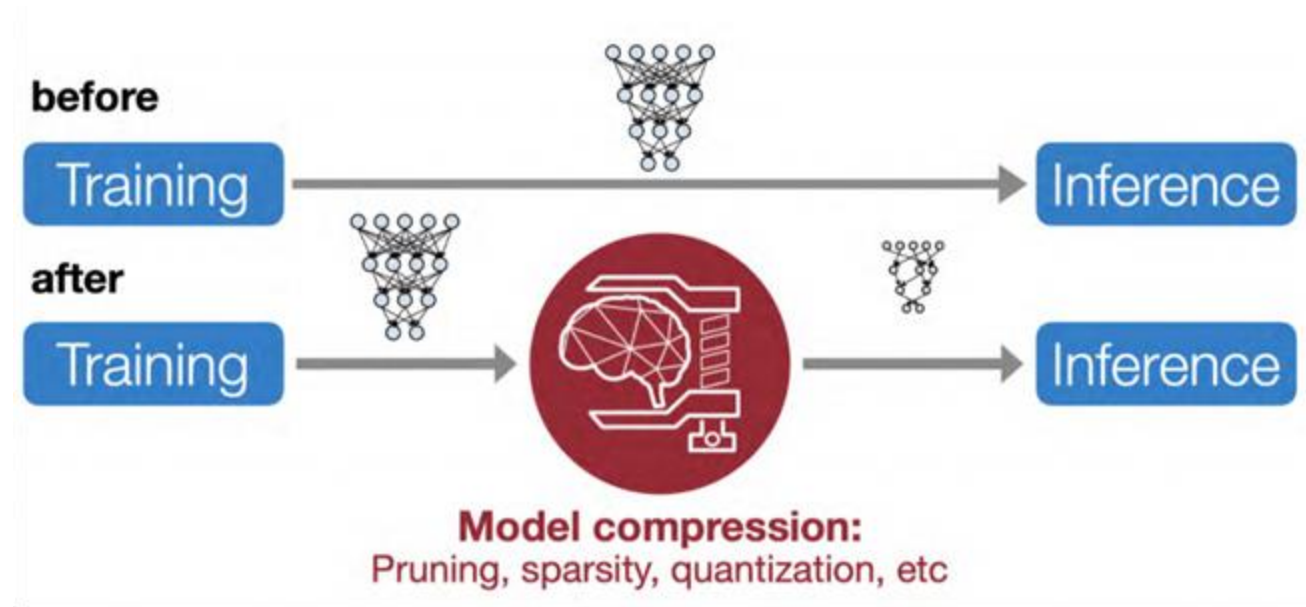


How to Improve Model Efficiency?

Hardware



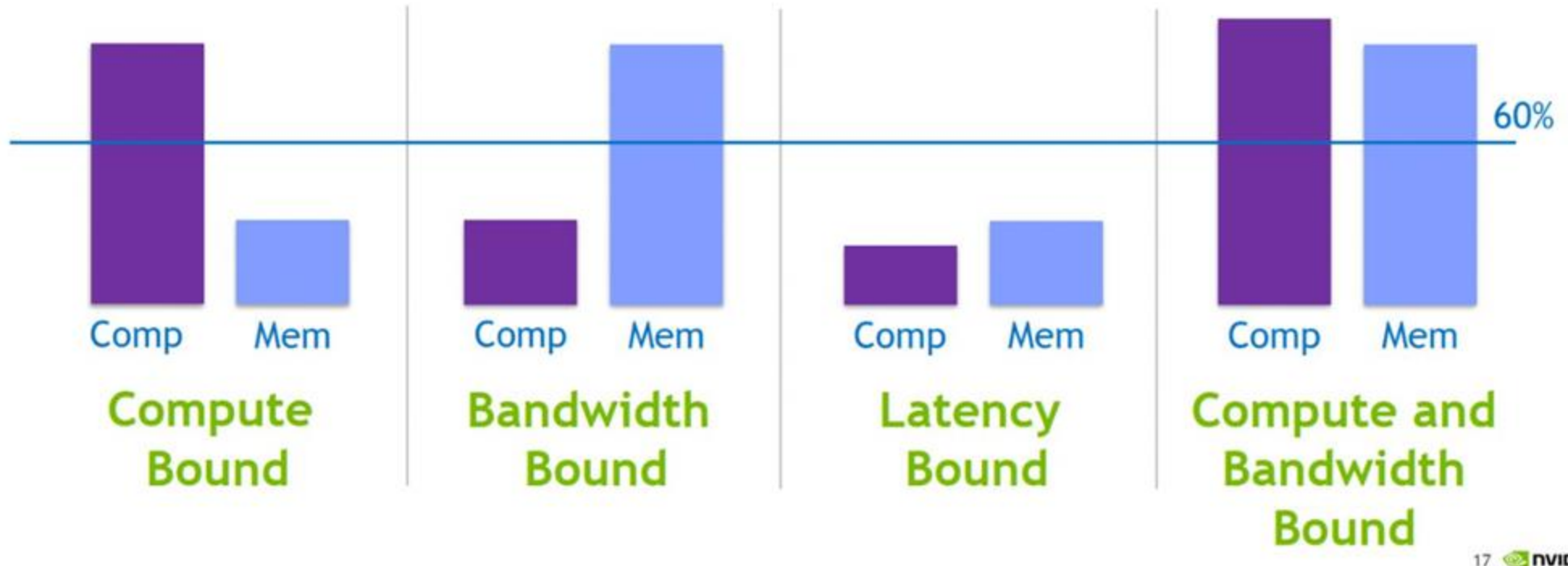
Software



What Makes a Language Model Slow

Memory Utilization vs Compute Utilization

Four possible combinations:



17 NVIDIA



Efficient LLMs

- ❑ Quantization
- ❑ Sparsity
- ❑ Long Context
- ❑ Serving & Systems
- ❑ Distillation
- ❑ Parameter Efficient Fine-Tuning
- ❑ Alternative Paradigms



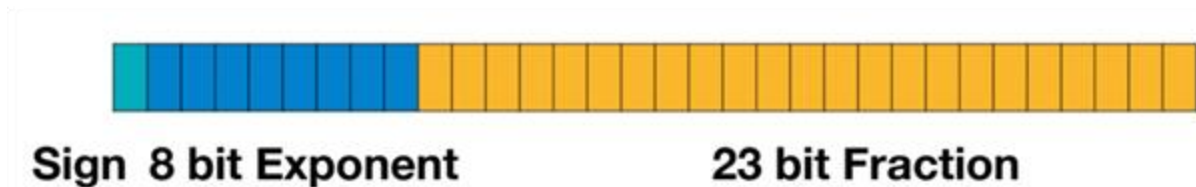
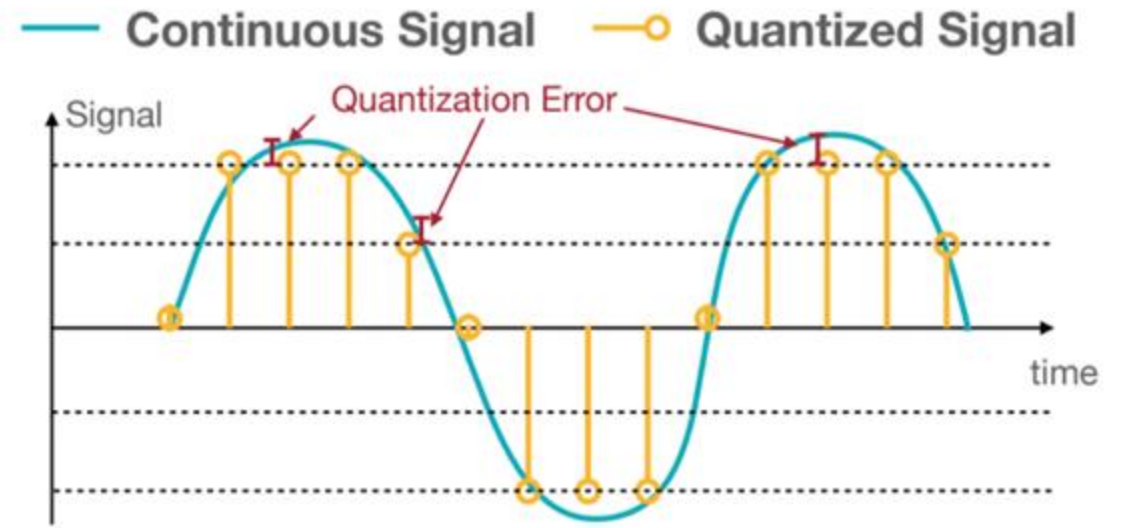
Efficient LLMs

- ❑ **Quantization**
- ❑ Sparsity
- ❑ Long Context
- ❑ Serving & Systems
- ❑ Distillation
- ❑ Parameter Efficient Fine-Tuning
- ❑ Alternative Paradigms

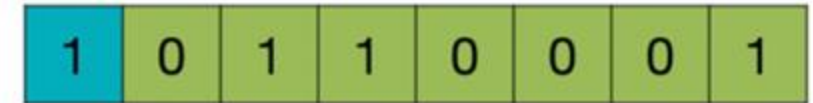


Quantization

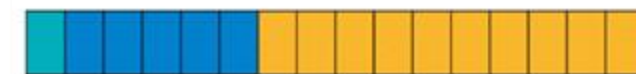
Reduce model size by
replacing high bit-
width
representations with
low bit-width
representations



Sign Bit



[IEEE 754](#) Half Precision 16-bit Float (IEEE FP16)



[Google](#) Brain Float (BF16)

How do we go from a high-bit width data type to a low-bit width data type?



K-Means Quantization vs Linear Quantization

2.09	-0.98	1.48	0.09
0.05	-0.14	-1.08	2.12
-0.91	1.92	0	-1.03
1.87	0	1.53	1.49

3	0	2	1	3:	2.00
1	1	0	3	2:	1.50
0	3	1	0	1:	0.00
3	1	2	2	0:	-1.00

1	-2	0	-1
-1	-1	-2	1
-2	1	-1	-2
1	-1	0	0

$(-1) \times 1.07$

**K-Means-based
Quantization**

**Linear
Quantization**

Storage	Floating-Point Weights	Integer Weights; Floating-Point Codebook	Integer Weights
Computation	Floating-Point Arithmetic	Floating-Point Arithmetic	Integer Arithmetic



K-Means Quantization vs Linear Quantization

2.09	-0.98	1.48	0.09
0.05	-0.14	-1.08	2.12
-0.91	1.92	0	-1.03
1.87	0	1.53	1.49

3	0	2	1	3:	2.00
1	1	0	3	2:	1.50
0	3	1	0	1:	0.00
3	1	2	2	0:	-1.00

1	-2	0	-1
-1	-1	-2	1
-2	1	-1	-2
1	-1	0	0

$(-1) \times 1.07$

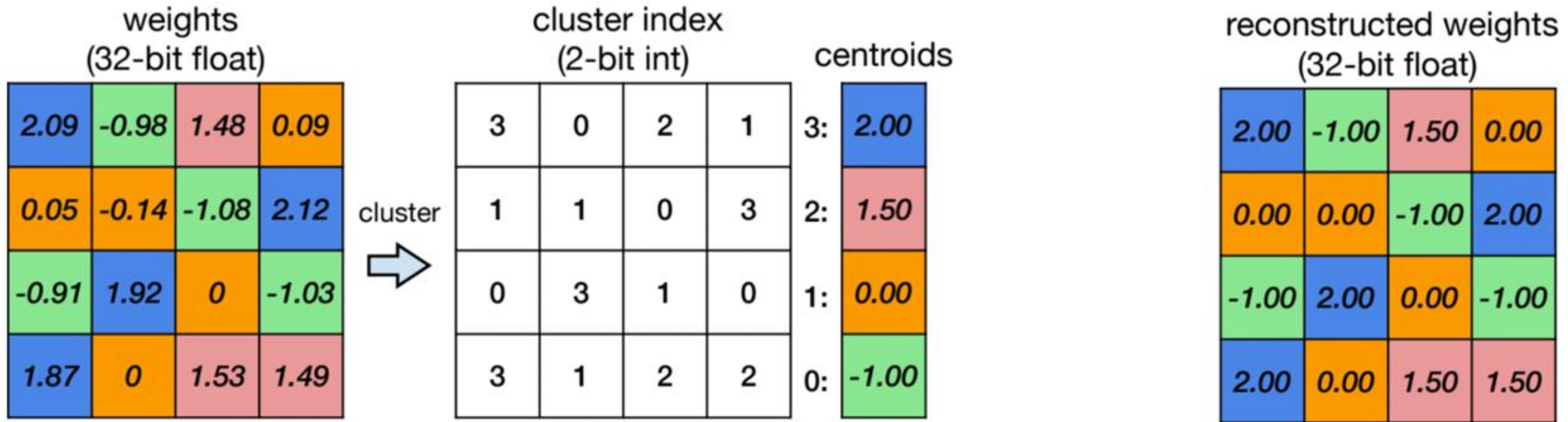
**K-Means-based
Quantization**

**Linear
Quantization**

Storage	Floating-Point Weights	Integer Weights; Floating-Point Codebook	Integer Weights
Computation	Floating-Point Arithmetic	Floating-Point Arithmetic	Integer Arithmetic



K-Means Quantization



Deep Compression [Han et al., ICLR 2016]



K-Means Quantization

Original weights

weights (32-bit float)			
2.09	-0.98	1.48	0.09
0.05	-0.14	-1.08	2.12
-0.91	1.92	0	-1.03
1.87	0	1.53	1.49

cluster



cluster index (2-bit int)				centroids	
3	0	2	1	3:	2.00
1	1	0	3	2:	1.50
0	3	1	0	1:	0.00
3	1	2	2	0:	-1.00

reconstructed weights
(32-bit float)

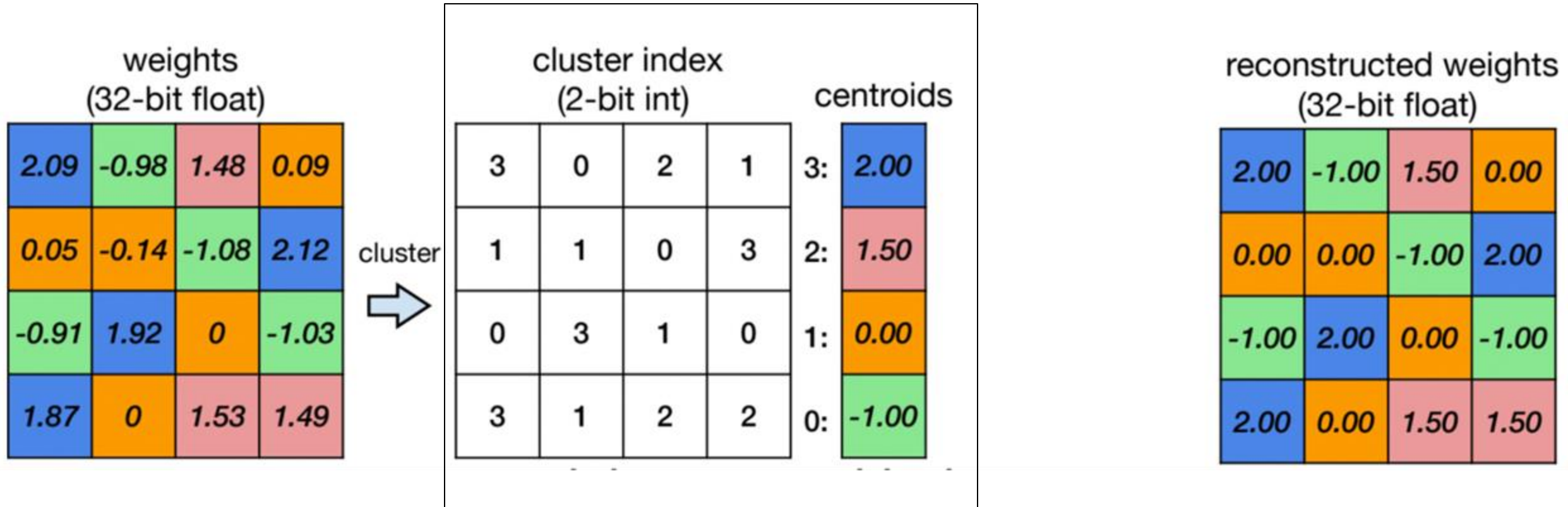
2.00	-1.00	1.50	0.00
0.00	0.00	-1.00	2.00
-1.00	2.00	0.00	-1.00
2.00	0.00	1.50	1.50

Deep Compression [Han et al., ICLR 2016]



K-Means Quantization

Stored weights after clustering

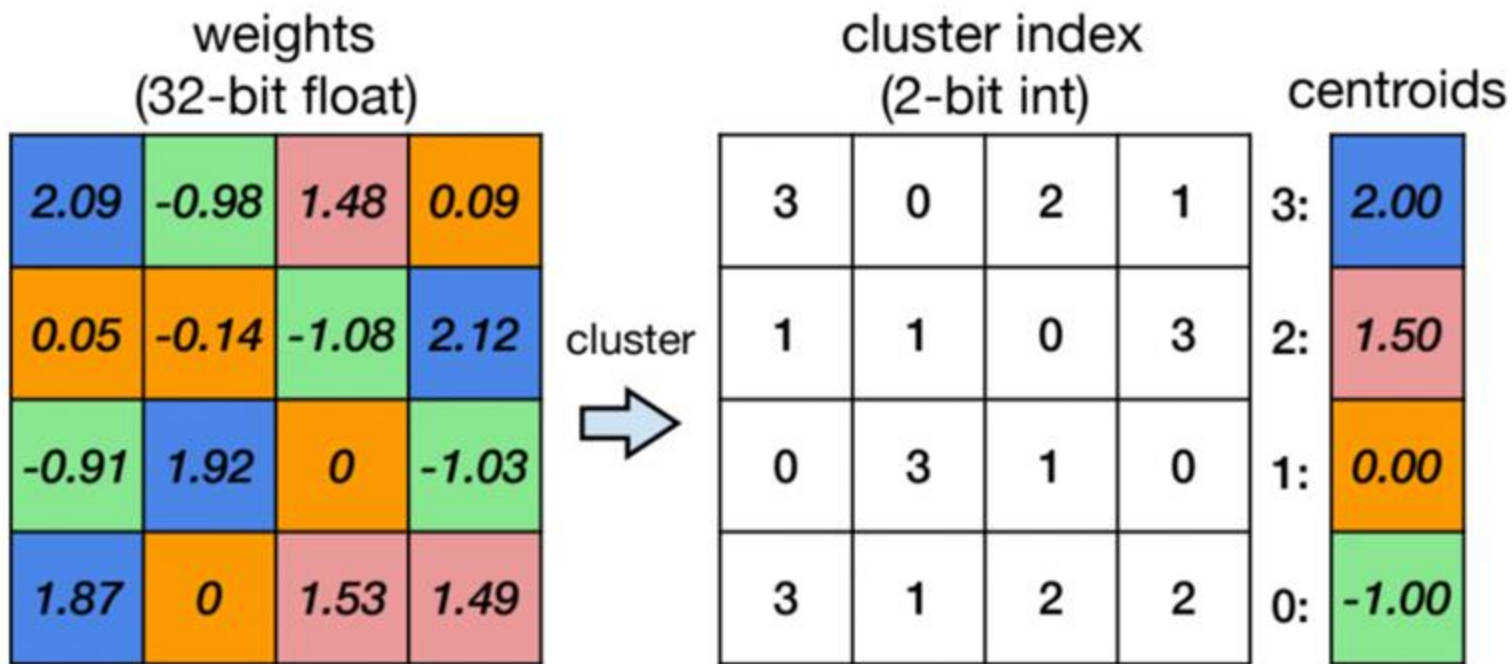


Deep Compression [Han et al., ICLR 2016]



K-Means Quantization

Retrieved weights to
be used at inference
time



reconstructed weights
(32-bit float)

2.00	-1.00	1.50	0.00
0.00	0.00	-1.00	2.00
-1.00	2.00	0.00	-1.00
2.00	0.00	1.50	1.50



K-Means Quantization vs Linear Quantization

2.09	-0.98	1.48	0.09
0.05	-0.14	-1.08	2.12
-0.91	1.92	0	-1.03
1.87	0	1.53	1.49

3	0	2	1	3:	2.00
1	1	0	3	2:	1.50
0	3	1	0	1:	0.00
3	1	2	2	0:	-1.00

1	-2	0	-1
-1	-1	-2	1
-2	1	-1	-2
1	-1	0	0

$(-1) \times 1.07$

**K-Means-based
Quantization**

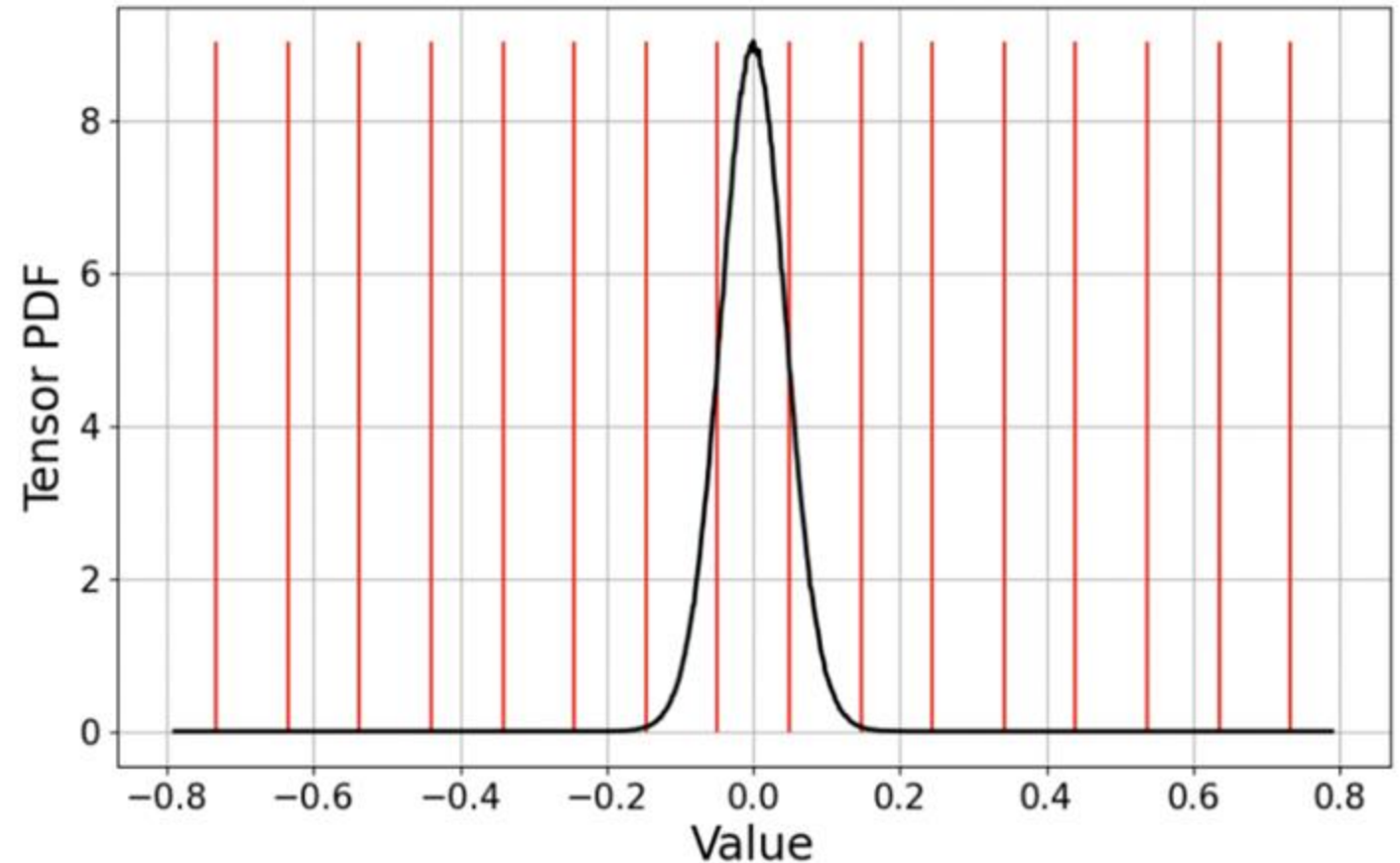
**Linear
Quantization**

Storage	Floating-Point Weights	Integer Weights; Floating-Point Codebook	Integer Weights
Computation	Floating-Point Arithmetic	Floating-Point Arithmetic	Integer Arithmetic



Linear Quantization

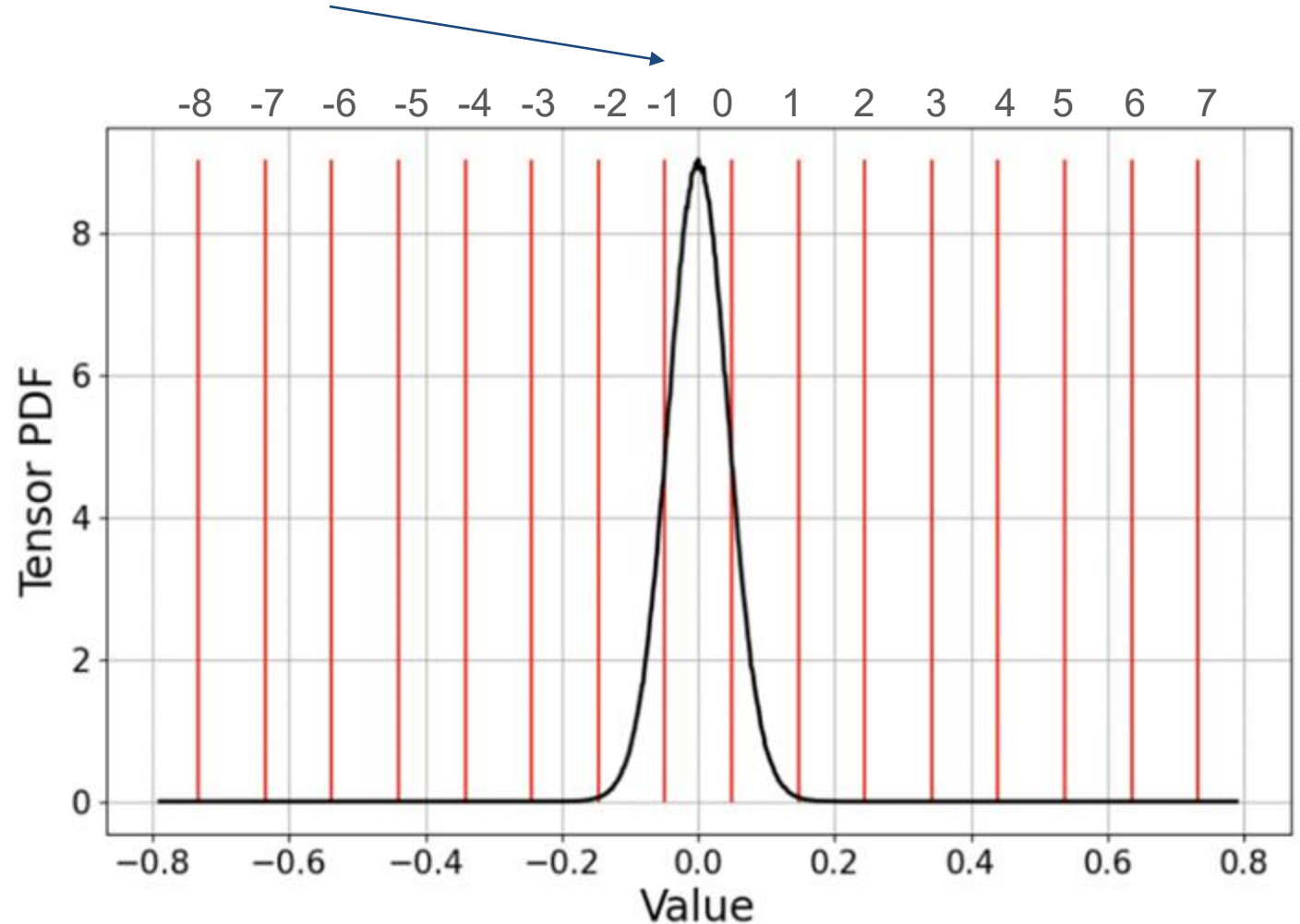
- ❑ Apply linear function on weights and hidden state activations from floating point values (r) to integer values (q)
- ❑ Original weights (black), Quantized bins (red)
- ❑ Black weights are mapped to one of the vertical red lines



Linear Quantization

- ❑ Apply linear function on weights and hidden state activations from floating point values (r) to integer values (q)
- ❑ Original weights (black), Quantized bins (red)
- ❑ Black weights are mapped to one of the vertical red lines

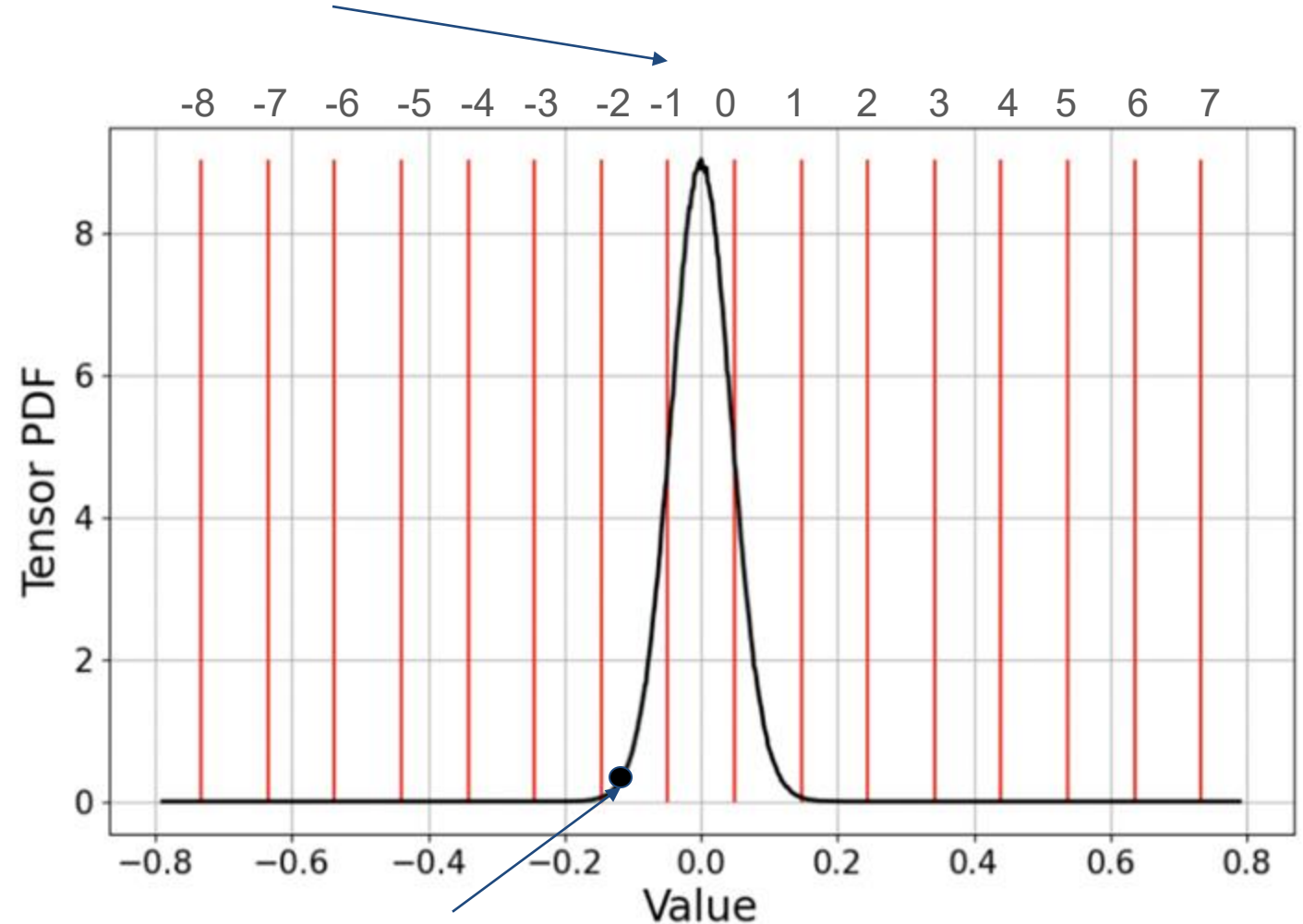
32-bit float to 4-bit int



Linear Quantization

- ❑ Apply linear function on weights and hidden state activations from floating point values (r) to integer values (q)
- ❑ Original weights (black), Quantized bins (red)
- ❑ Black weights are mapped to one of the vertical red lines

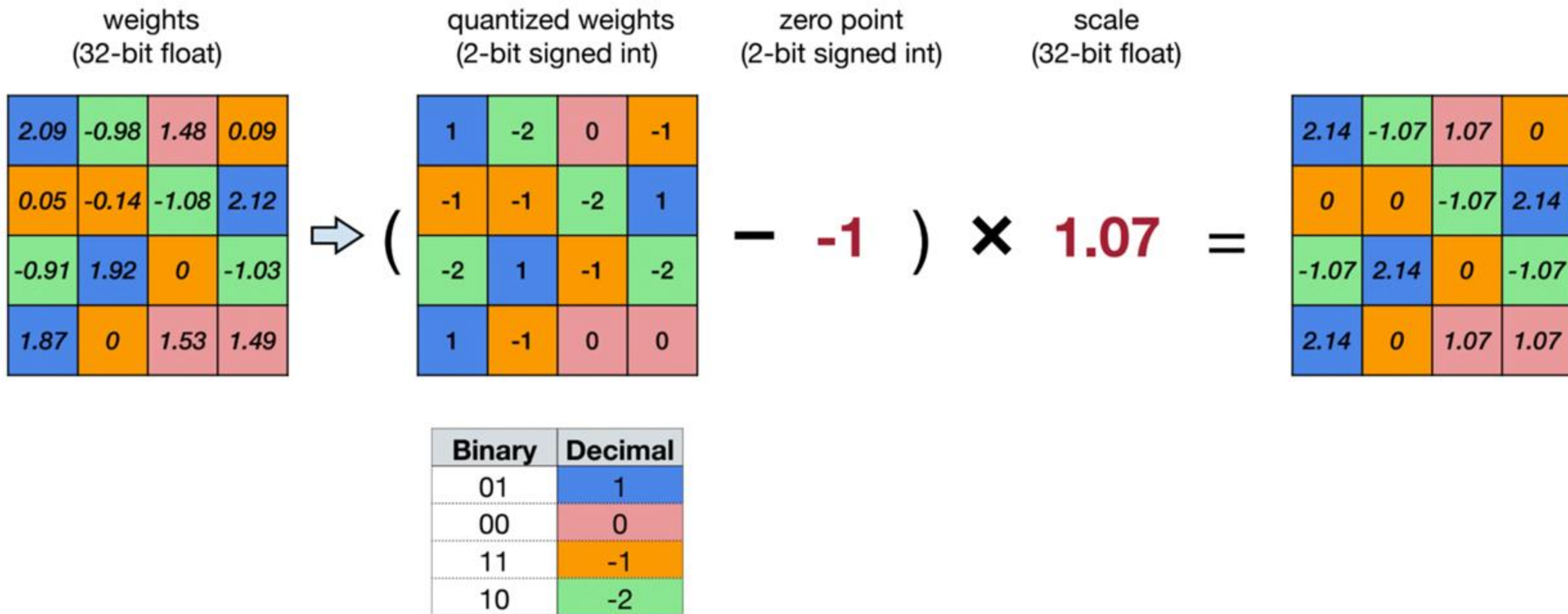
32-bit float to 4-bit int



Before Quantization: -0.14

After Quantization: -2

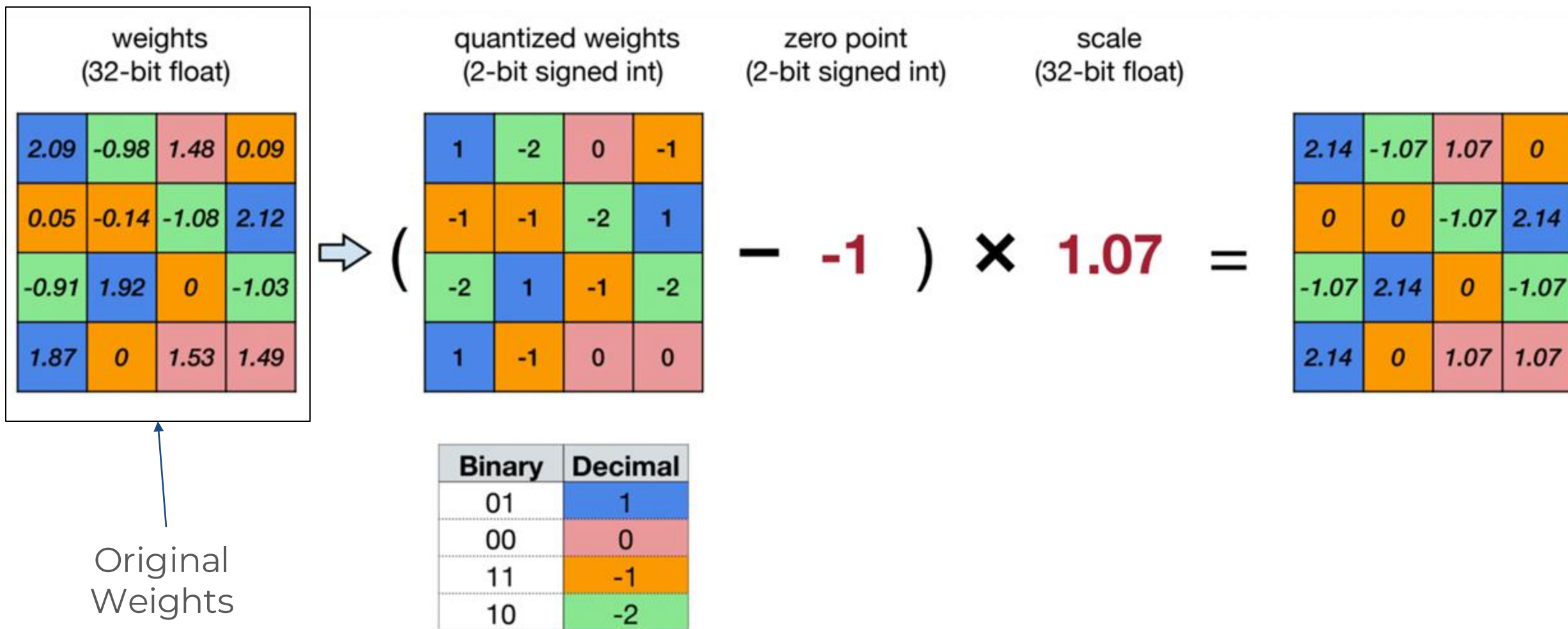
Linear Quantization



Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference [Jacob et al., CVPR 2018]

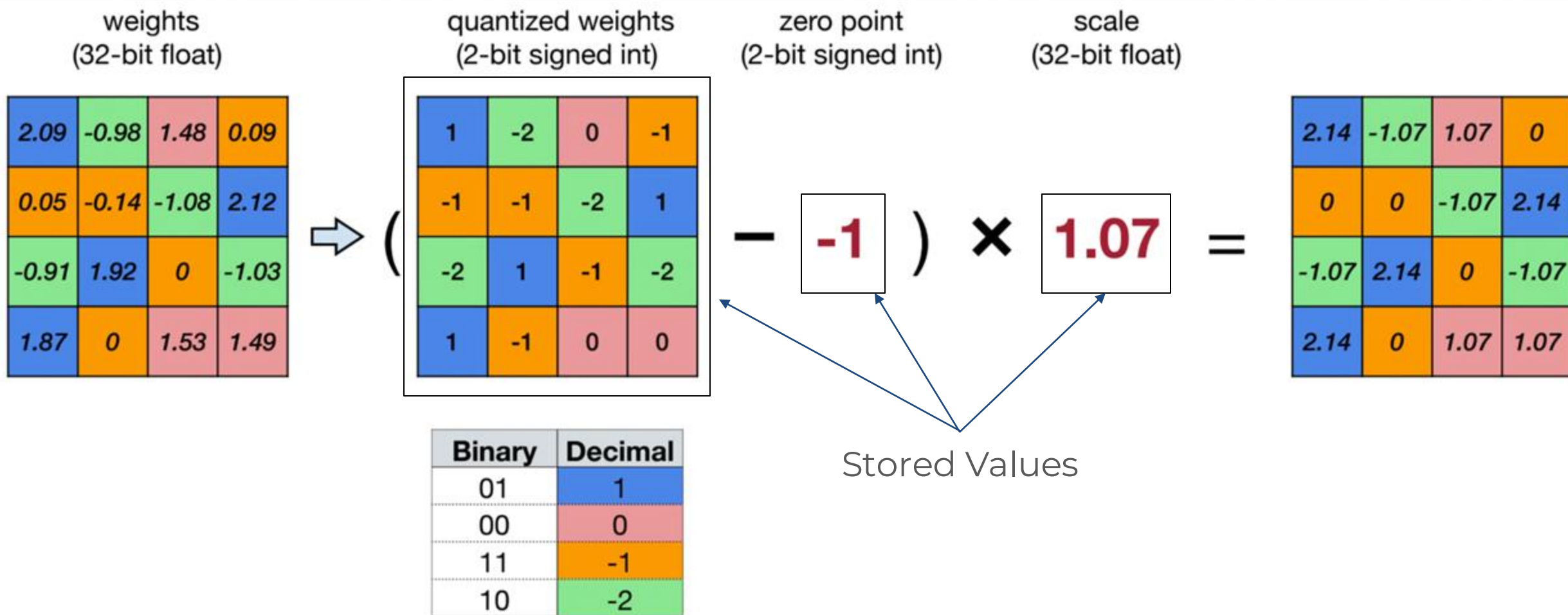


Linear Quantization



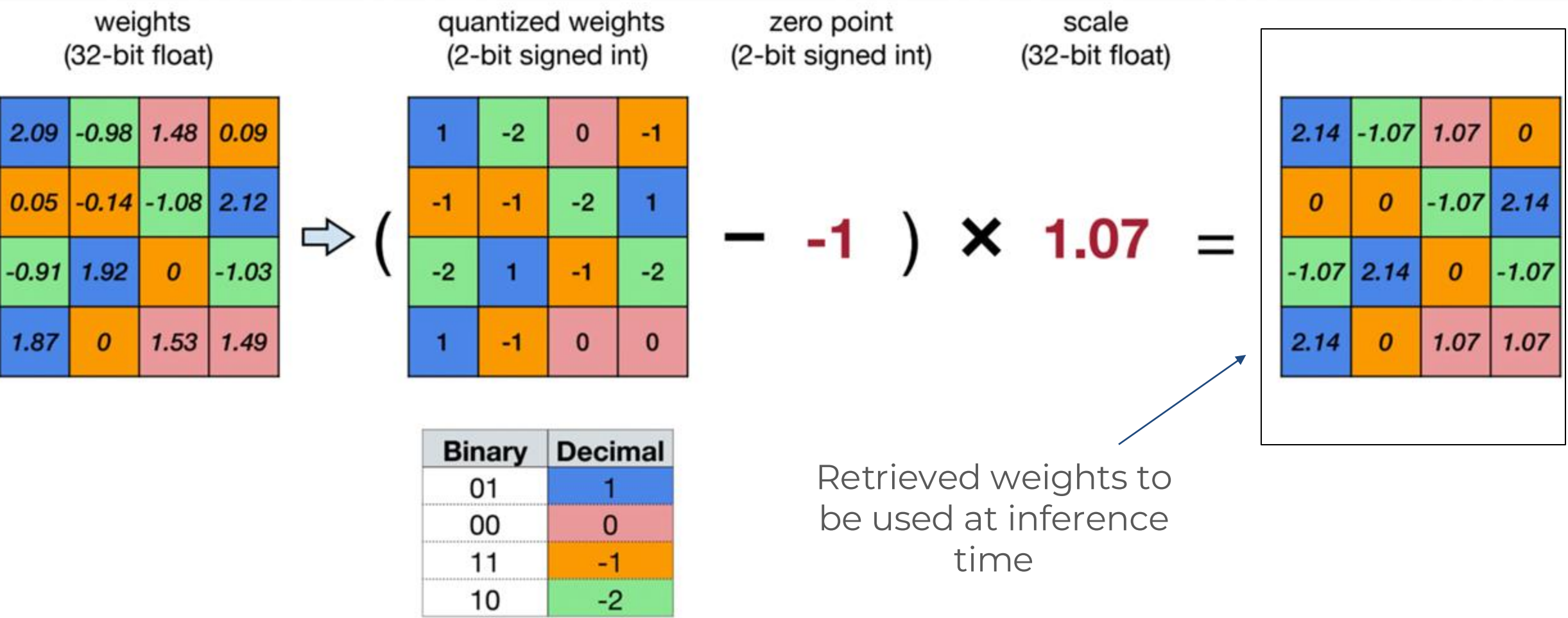
Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference [Jacob et al., CVPR 2018]

Linear Quantization



Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference [Jacob et al., CVPR 2018]

Linear Quantization

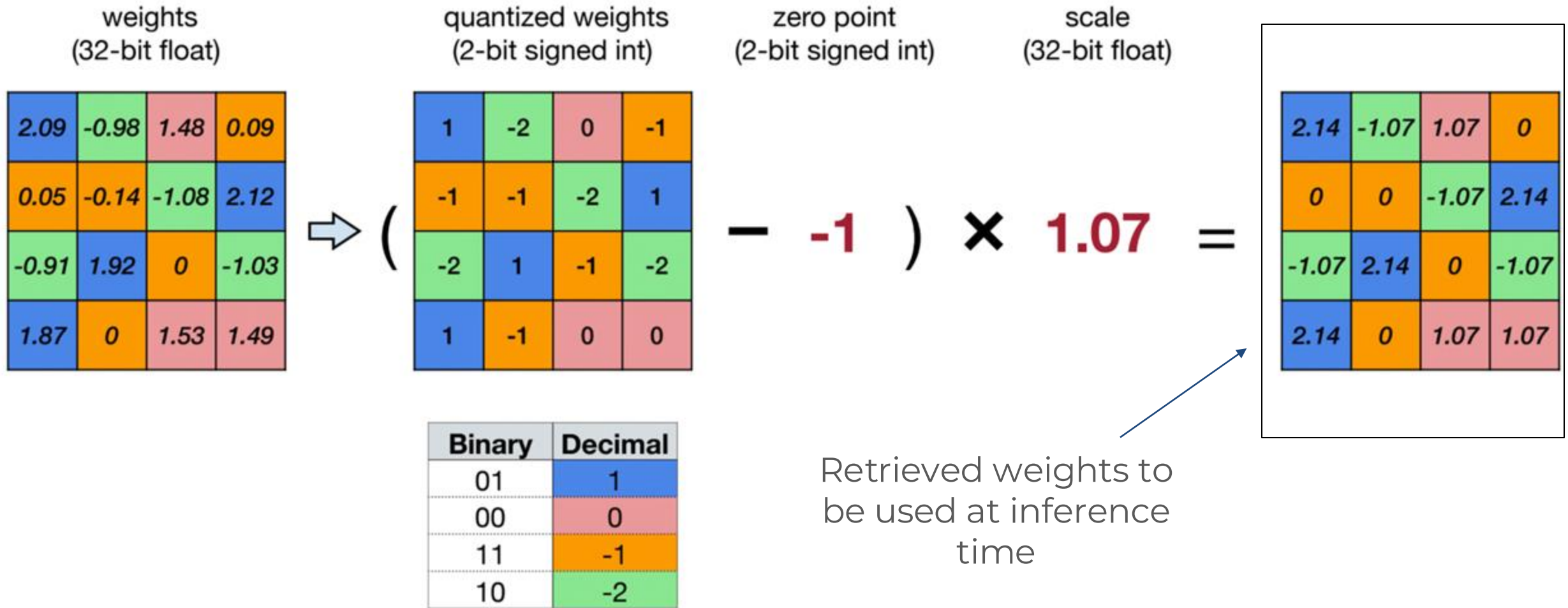


Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference [Jacob et al., CVPR 2018]



Linear Quantization

Virtually all modern methods use this
(as opposed to KMeans Quantization)



Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference [Jacob et al., CVPR 2018]



Should we use KMeans/Linear Quantization on all weights at the same time?



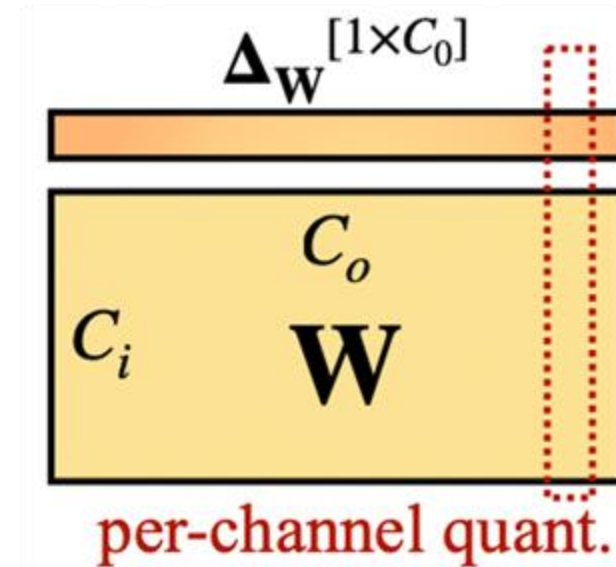
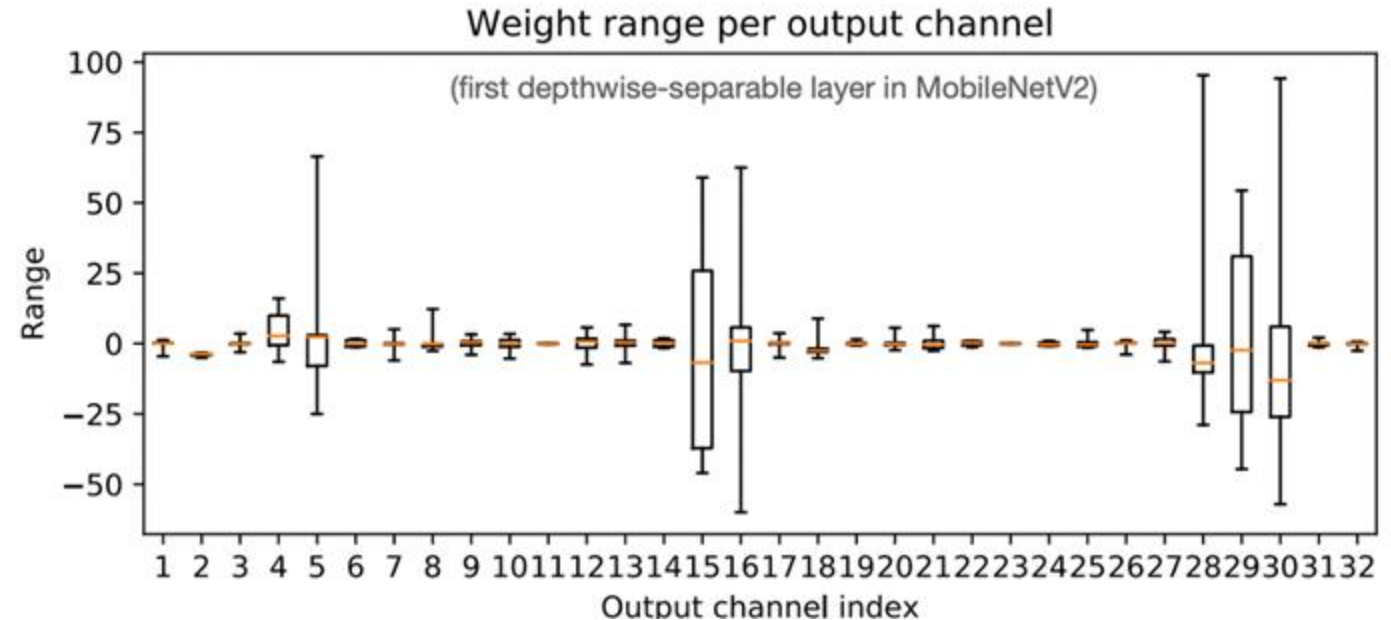
Should we use KMeans/Linear Quantization on all weights at the same time?

No. We should group them.
The size of the grouping we choose is denoted as the granularity.



Weight Granularity

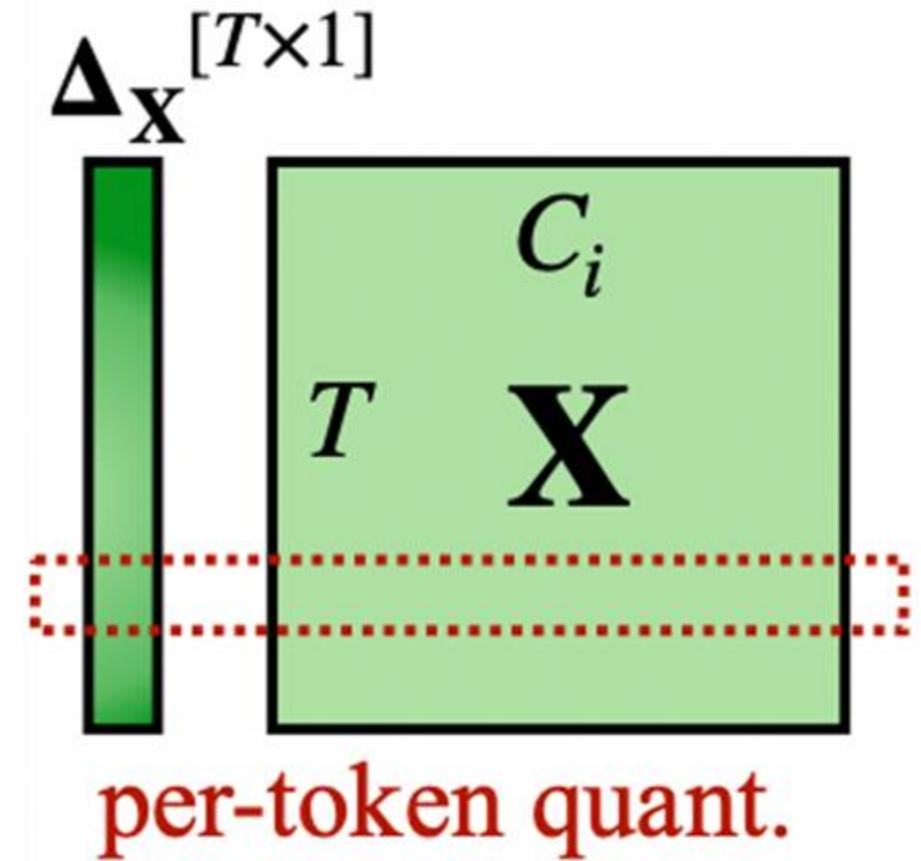
- ❑ Weight matrices will often have different variances along each output channel
- ❑ High variance in weights means that applying linear quantization will result in large performance degradation
- ❑ To fix this, we can perform linear quantization along each channel of the weight tensor separately



SmoothQuant: Accurate and Efficient Post-Training Quantization for Large Language Models[Xiao et. al., ICML 2023]

Activation Granularity

- ❑ Activations can have a similar problem whereby the variance by channel can be quite different
- ❑ The variance by token can also differ dramatically
- ❑ When applying quantization, we should split up channels, tokens to take this into account



SmoothQuant: Accurate and Efficient Post-Training Quantization for Large Language Models[Xiao et. al., ICML 2023]

When do we perform this quantization?



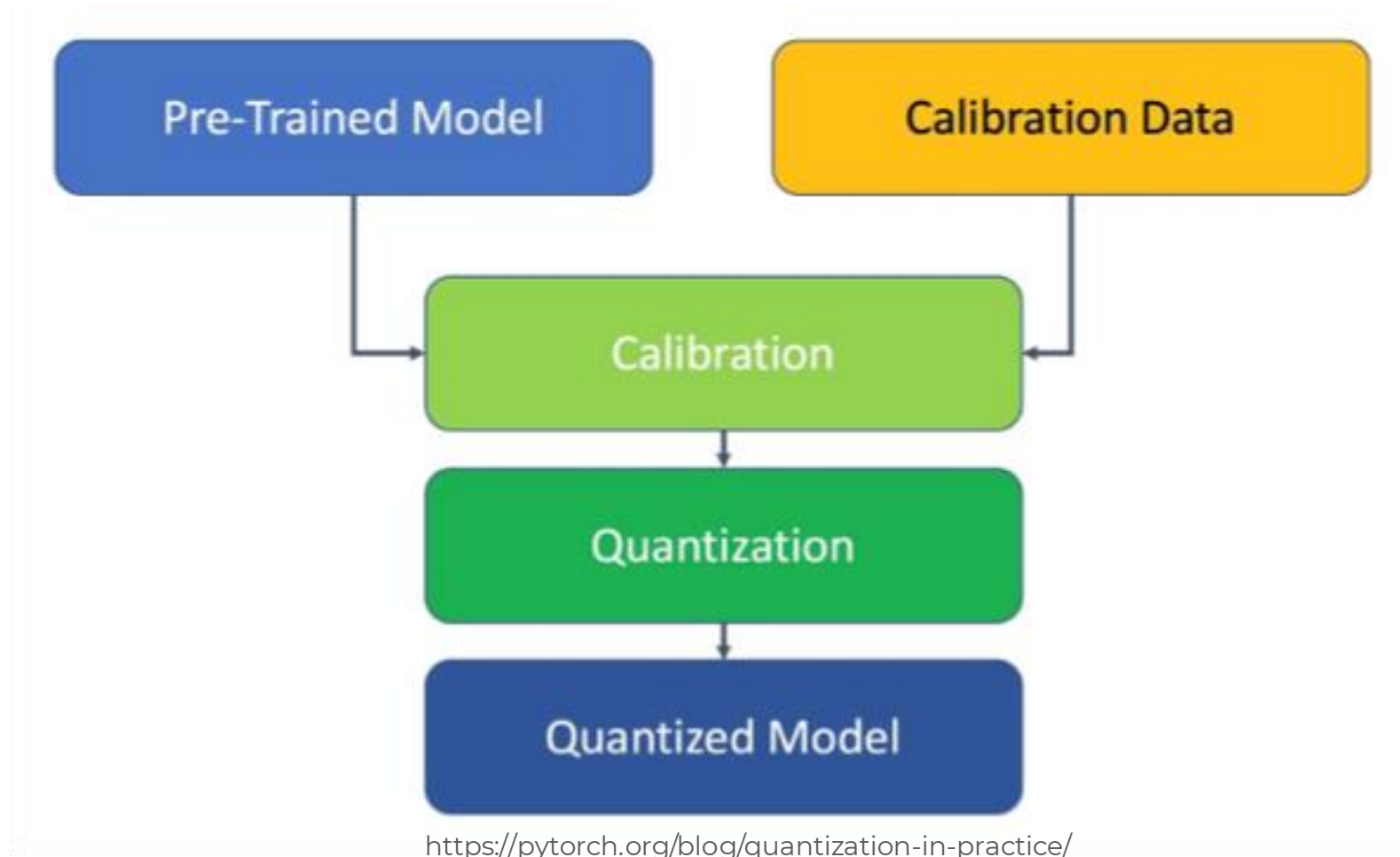
When do we perform this quantization?

Typically, this is done after first
completely training a model
(pretraining $\rightarrow$ SFT $\rightarrow$ post-
training)



Post Training Quantization (PTQ)

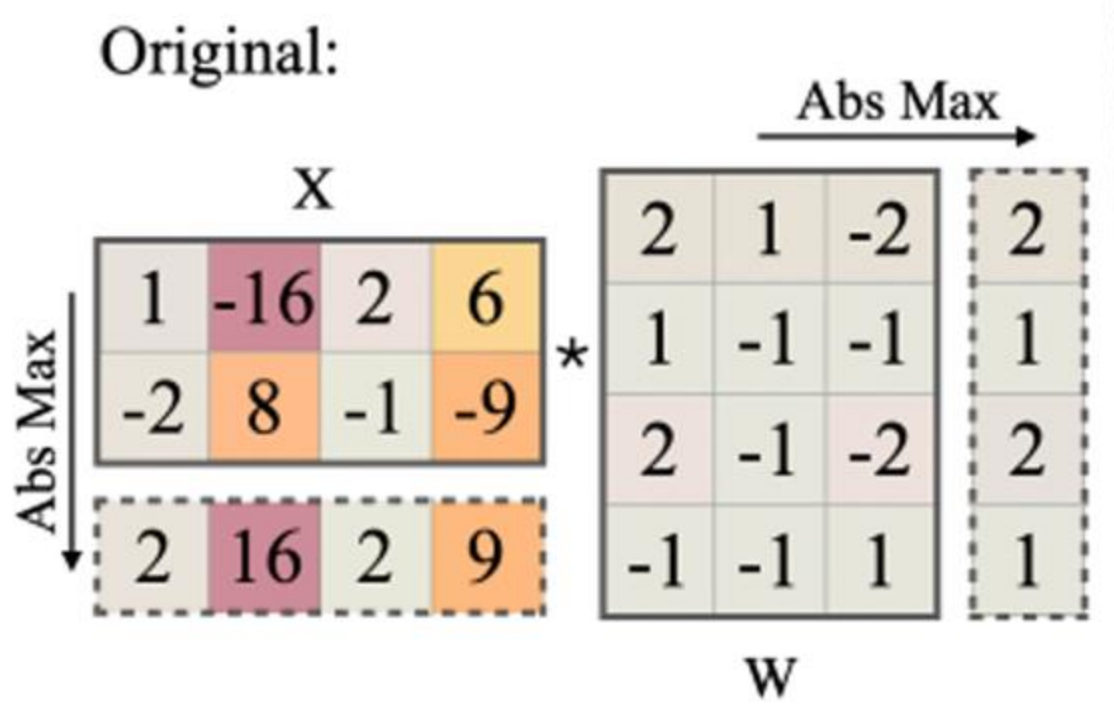
Quantize after training



What are some of the modern methods for performing quantization?



AWQ (W4A16)

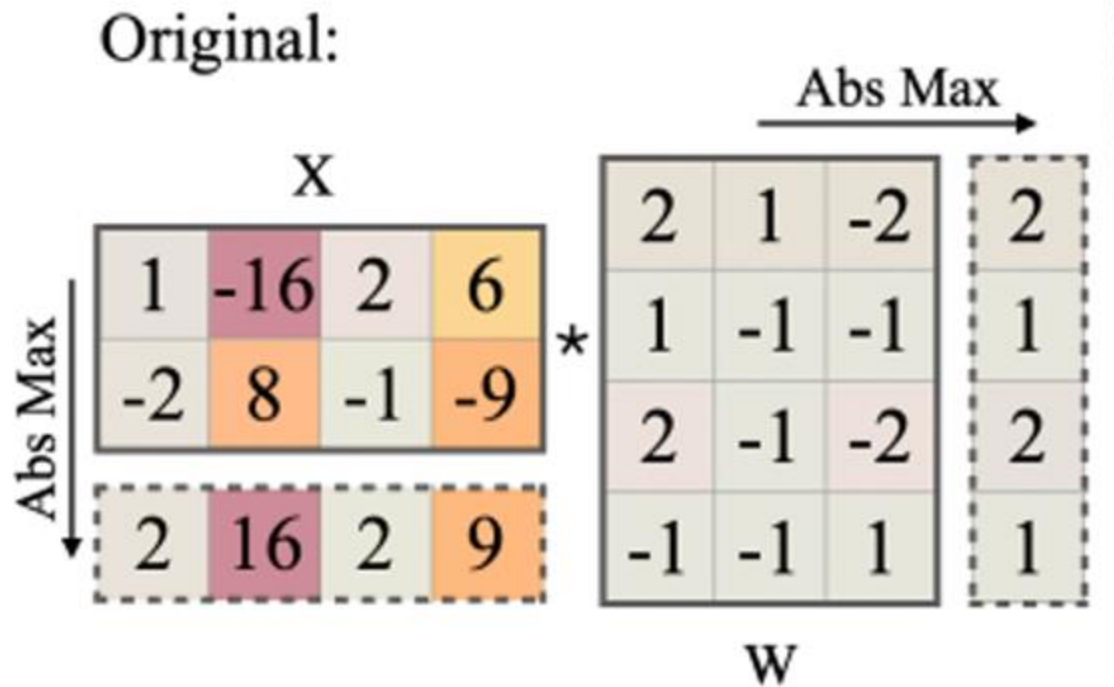


Observation: ?

SmoothQuant: Accurate and Efficient Post-Training Quantization for Large Language Models[Xiao et. al., ICML 2023]



AWQ (W4A16)

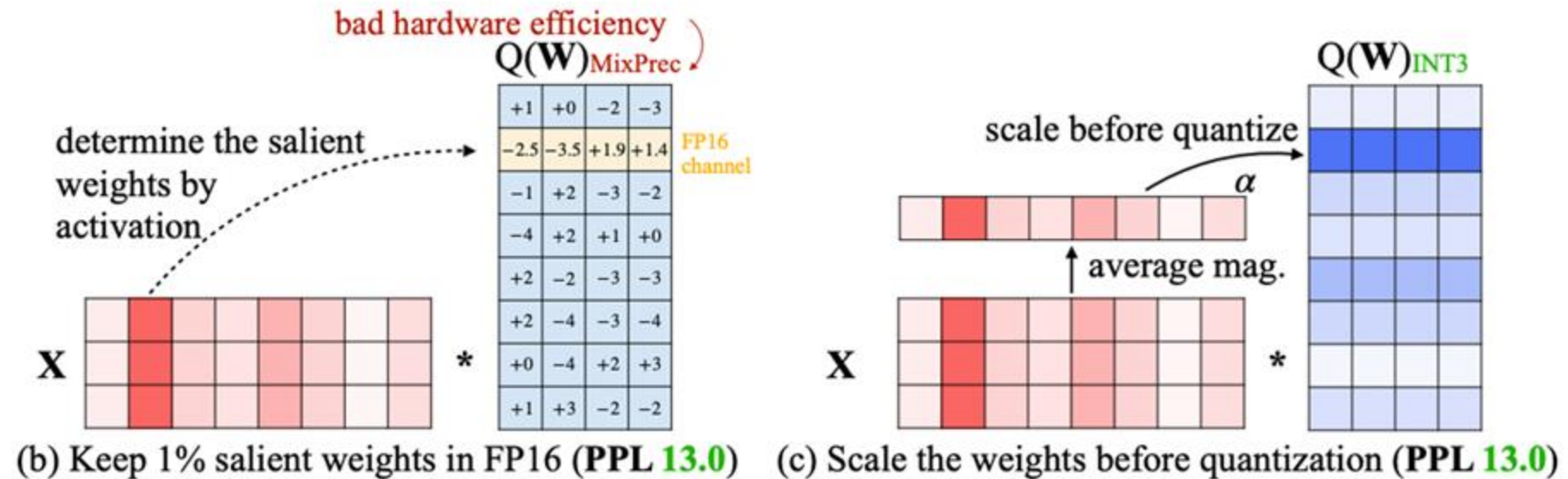
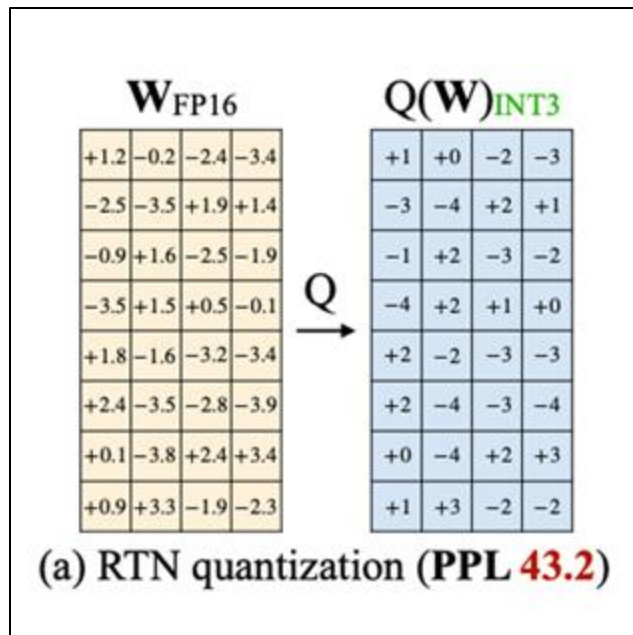


Observation: High variance channels are fixed in activations in LLM FFN layers- weights have relatively little difference in variance

SmoothQuant: Accurate and Efficient Post-Training Quantization for Large Language Models[Xiao et. al., ICML 2023]

AWQ (W4A16)

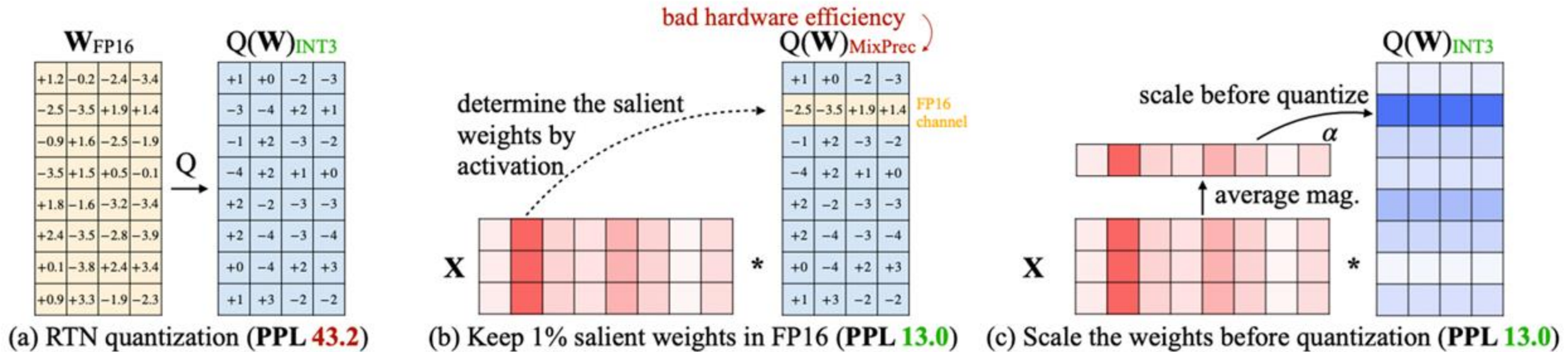
Normal quantization on LLMs performs poorly due to outliers in the model's hidden state



AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration [Lin et. al., arxiv 2023]

AWQ (W4A16)

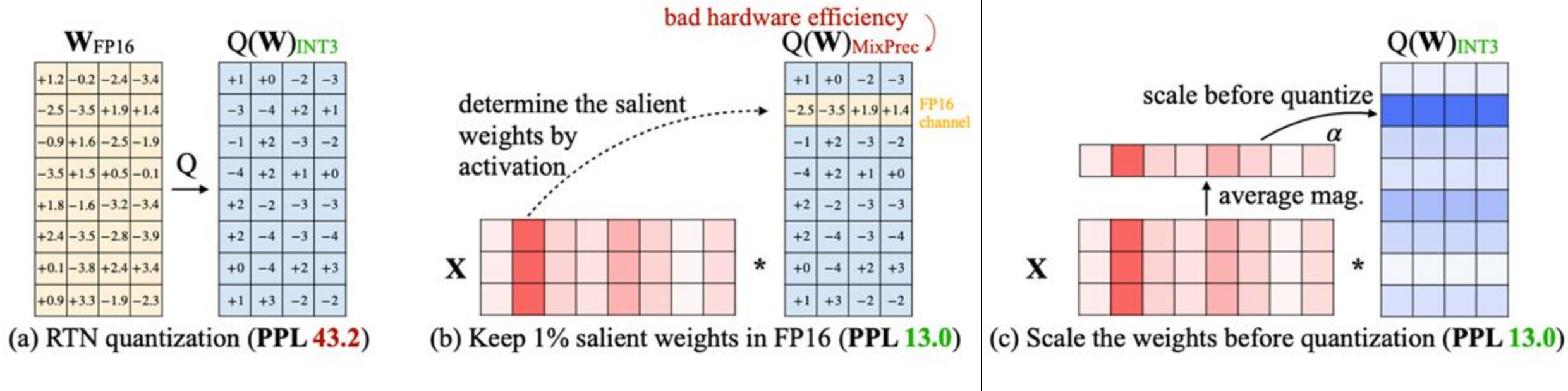
AWQ (c below) scales weights to handle the channels separately



AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration [Lin et. al., arxiv 2023]

AWQ (W4A16)

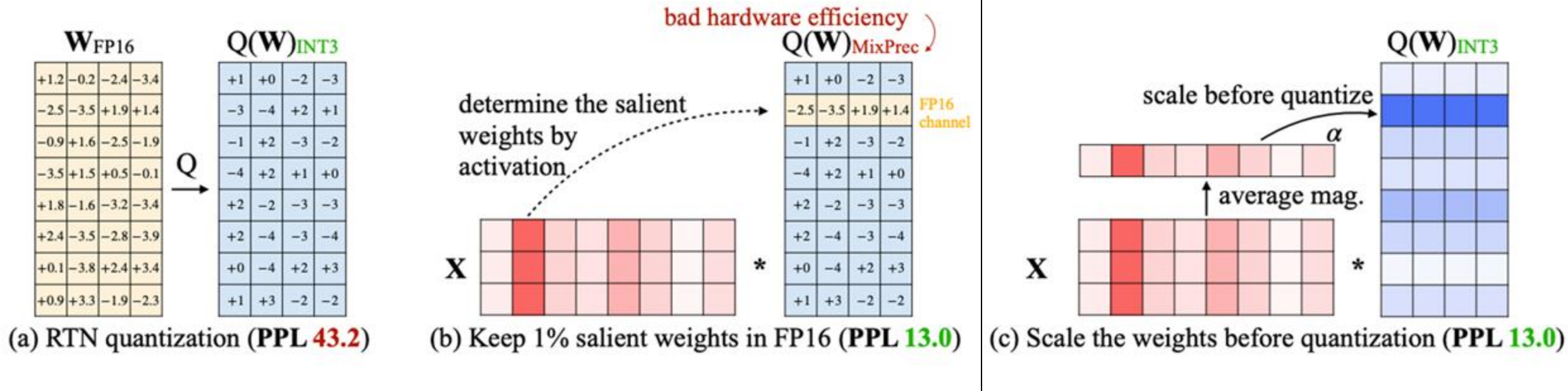
Uses a calibration batch of data to determine what the outlier channels are along with a per-channel scaling factor for these channels



AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration [Lin et. al., arxiv 2023]

AWQ (W4A16)

AWQ is weight-only →
multiplication is still
performed in BF16/FP16.

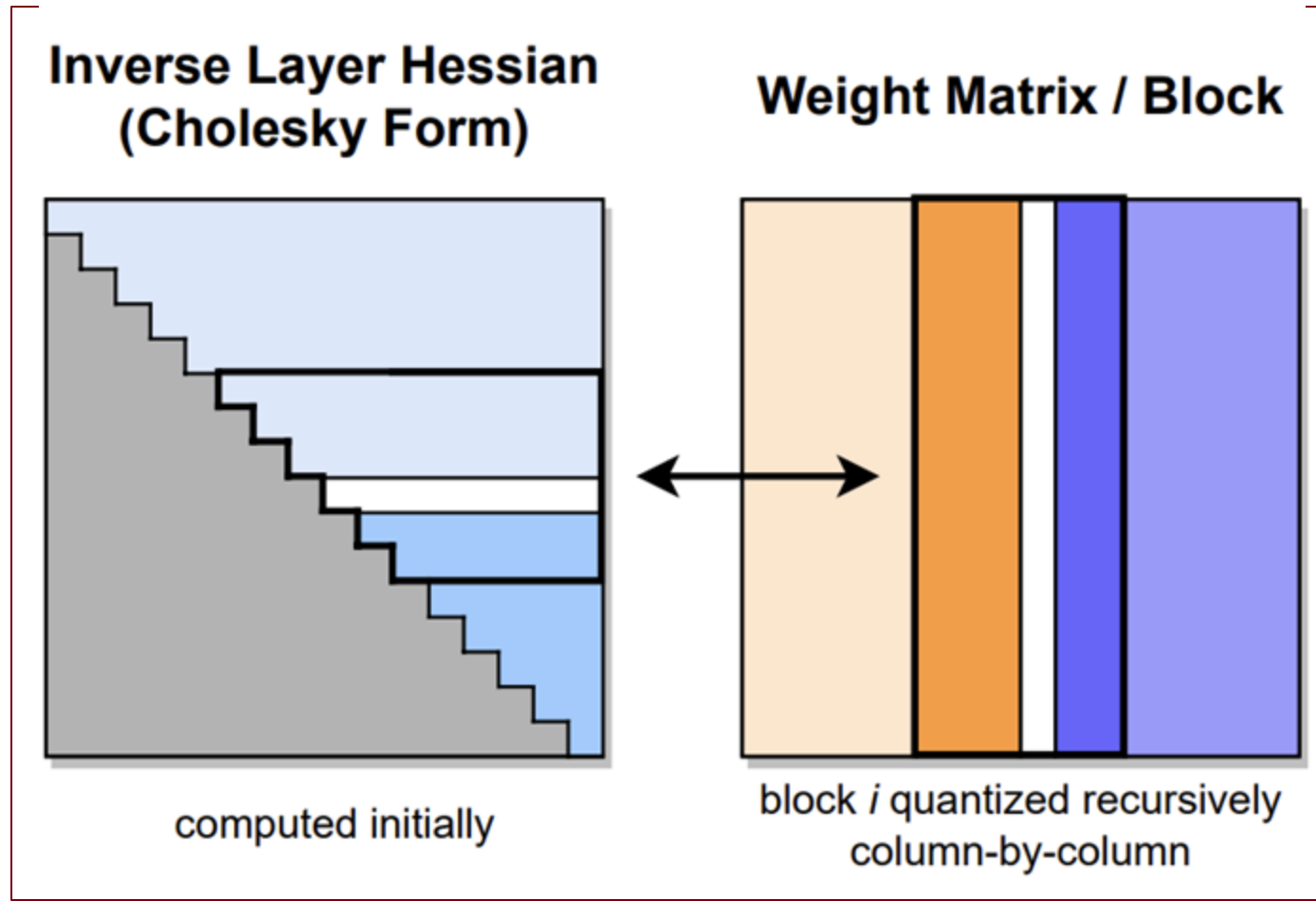


AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration [Lin et. al., arxiv 2023]

GPTQ

- Quantizes a block of each weight matrix at a time.
- Uses a loss term + the inverse hessian for the layer to update the non-quantized weights
- Takes a *long* time

$$\operatorname{argmin}_{\widehat{\mathbf{W}}} ||\mathbf{W}\mathbf{X} - \widehat{\mathbf{W}}\mathbf{X}||_2^2.$$



GPTQ: Accurate Post-Training Quantization for Generative Pre-Trained Transformers [Frantar et al., arXiv 2022]

GGUF Quantization

- No model retraining as is done with GPTQ
- A model is quantized by taking blocks of the input matrix, then performing linear quantization on each block separately
- Uses several different methods based on level of quantization
- Will come with names such as IQ2_XS–IQ4_XS, Q2_K_S–Q5_K_S
- You can use [unsloth](#) to create models with this type



GPTQ: Accurate Post-Training Quantization for Generative Pre-Trained Transformers [Frantar et al., arXiv 2022]



FP8 Quantization (Nvidia)

- Offered by the [NVIDIA TensorRT Model Optimizer](#) library
- Performs a separate per-channel post-training quantization
- Supported by hardware
- This library also performs a number of other functions for improving model efficiency



Modern Quantization (Open Source Models)

Oftentimes, open source LLM releases from model providers will release versions of models that use several of the above quantization techniques

 Qwen/Qwen3-235B-A22B-Thinking-2507-FP8

 Text Generation •  235B • Updated Jul 30 •  17.9k •  64

 Qwen/Qwen3-235B-A22B-Thinking-2507

 Text Generation •  235B • Updated Aug 17 •  47.3k •  •  377

 Qwen/Qwen3-235B-A22B-Instruct-2507-FP8

 Text Generation •  235B • Updated Sep 17 •  97.2k •  134

 Qwen/Qwen3-235B-A22B-Instruct-2507

 Text Generation •  235B • Updated Sep 17 •  108k •  •  714

 Qwen/Qwen3-30B-A3B-Thinking-2507-FP8

 Text Generation •  31B • Updated Jul 30 •  20.2k •  45

 Qwen/Qwen3-30B-A3B-Thinking-2507

 Text Generation •  31B • Updated Aug 17 •  267k •  •  315

 Qwen/Qwen3-30B-A3B-Instruct-2507-FP8

 Text Generation •  31B • Updated Sep 17 •  486k •  92



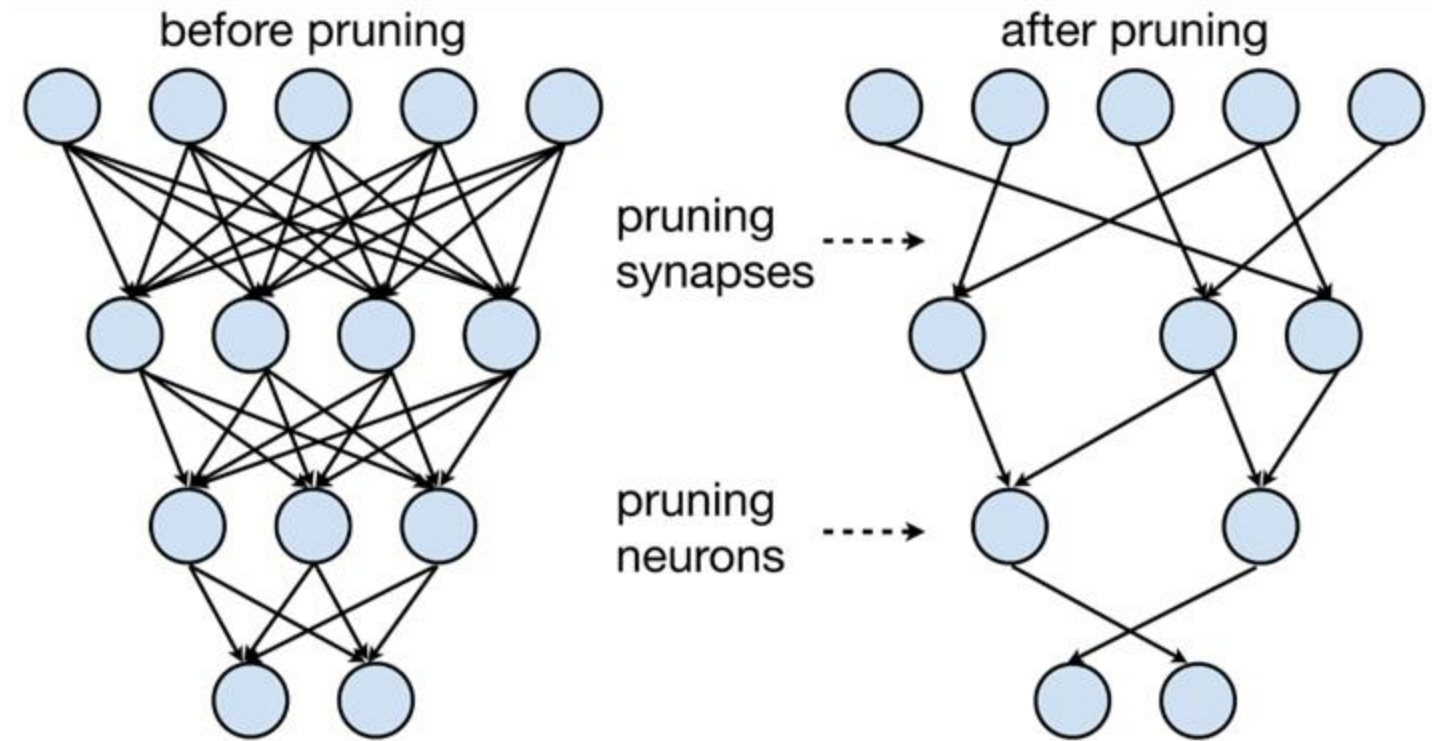
Efficient LLMs

- ❑ Quantization
- ❑ **Sparsity**
- ❑ Long Context
- ❑ Serving & Systems
- ❑ Distillation
- ❑ Parameter Efficient Fine-Tuning
- ❑ Alternative Paradigms



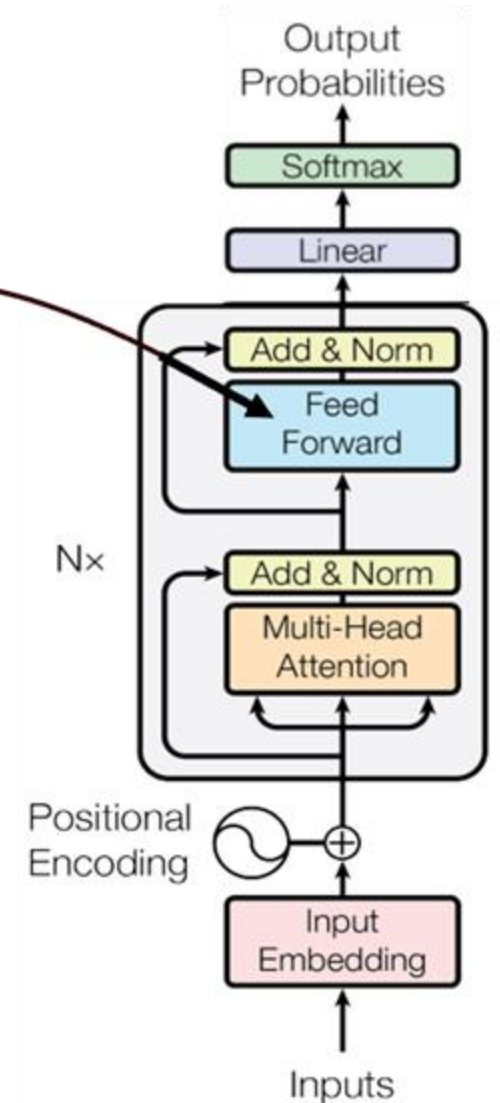
Sparsity

Even though our model may have many parameters, we can get speedups by only using a much smaller number of those parameters for a given instance

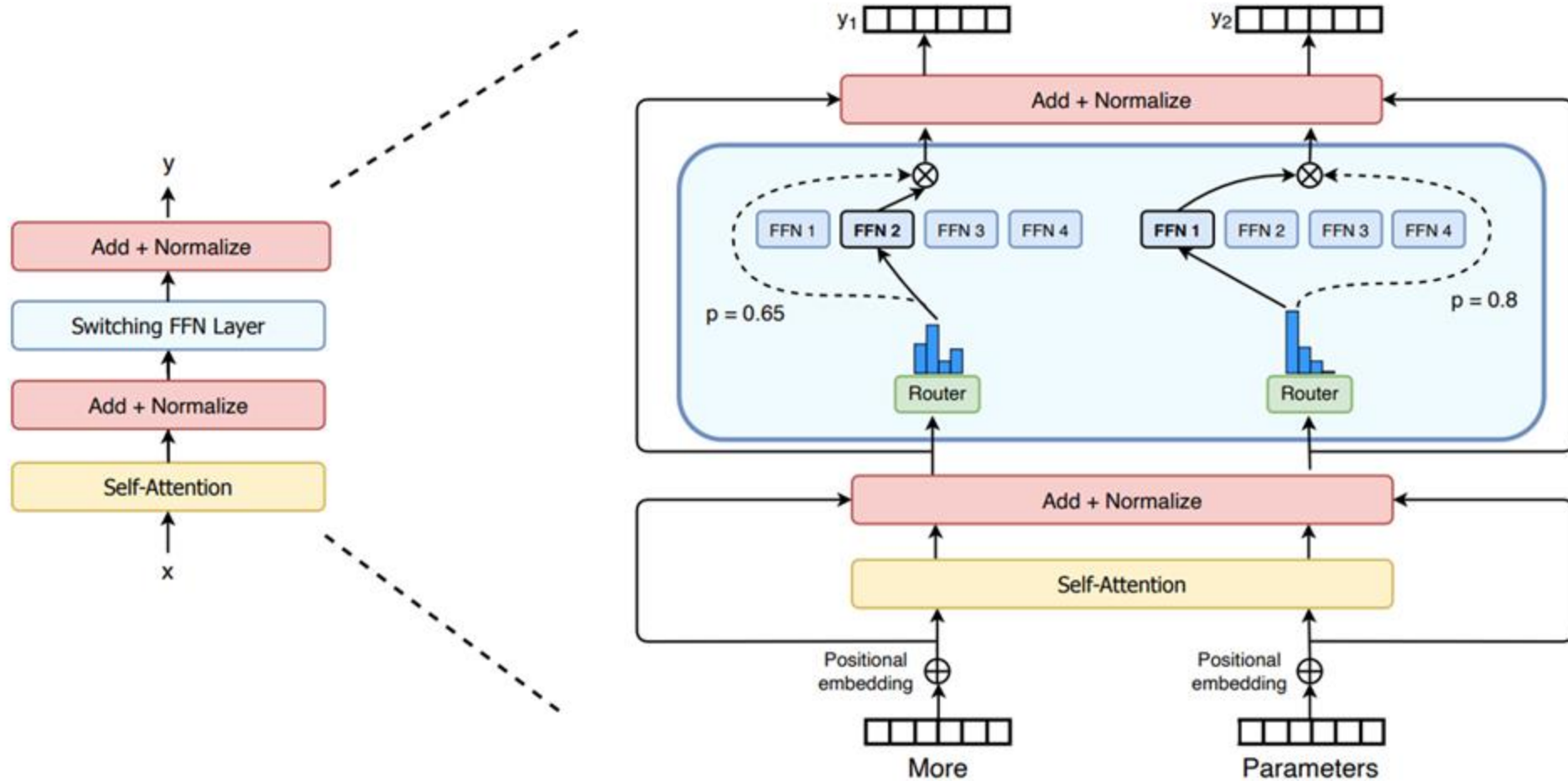


Mixture of Experts (MoE)

Replace FFN layers in traditional transformers with a switching FFN layer (more generally called an MoE layer)



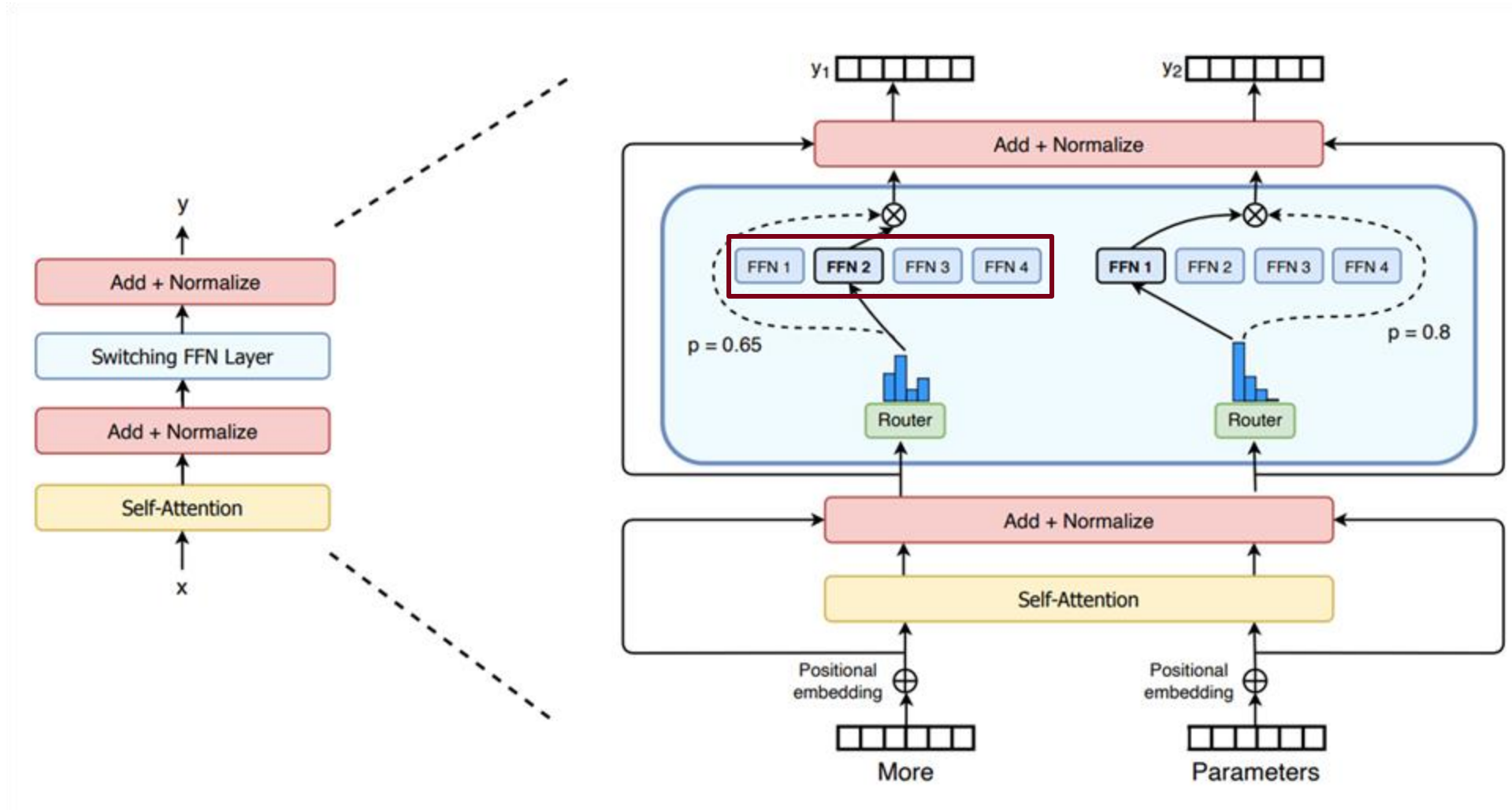
Mixture of Experts (MoE)



Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity [Fedus et al., CoRR 2021]

Mixture of Experts (MoE)

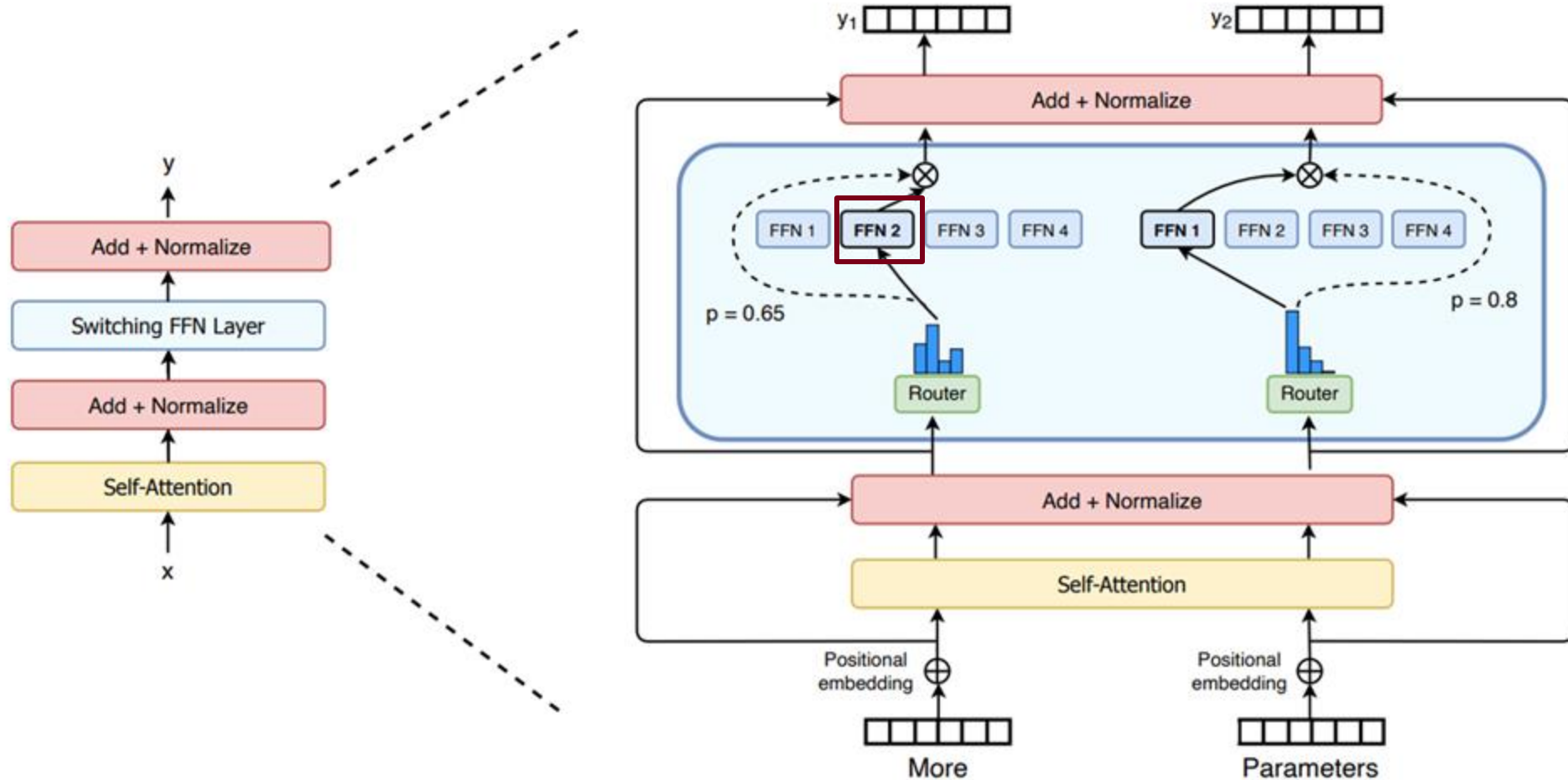
Four FFN layers



Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity [Fedus et al., CoRR 2021]

Mixture of Experts (MoE)

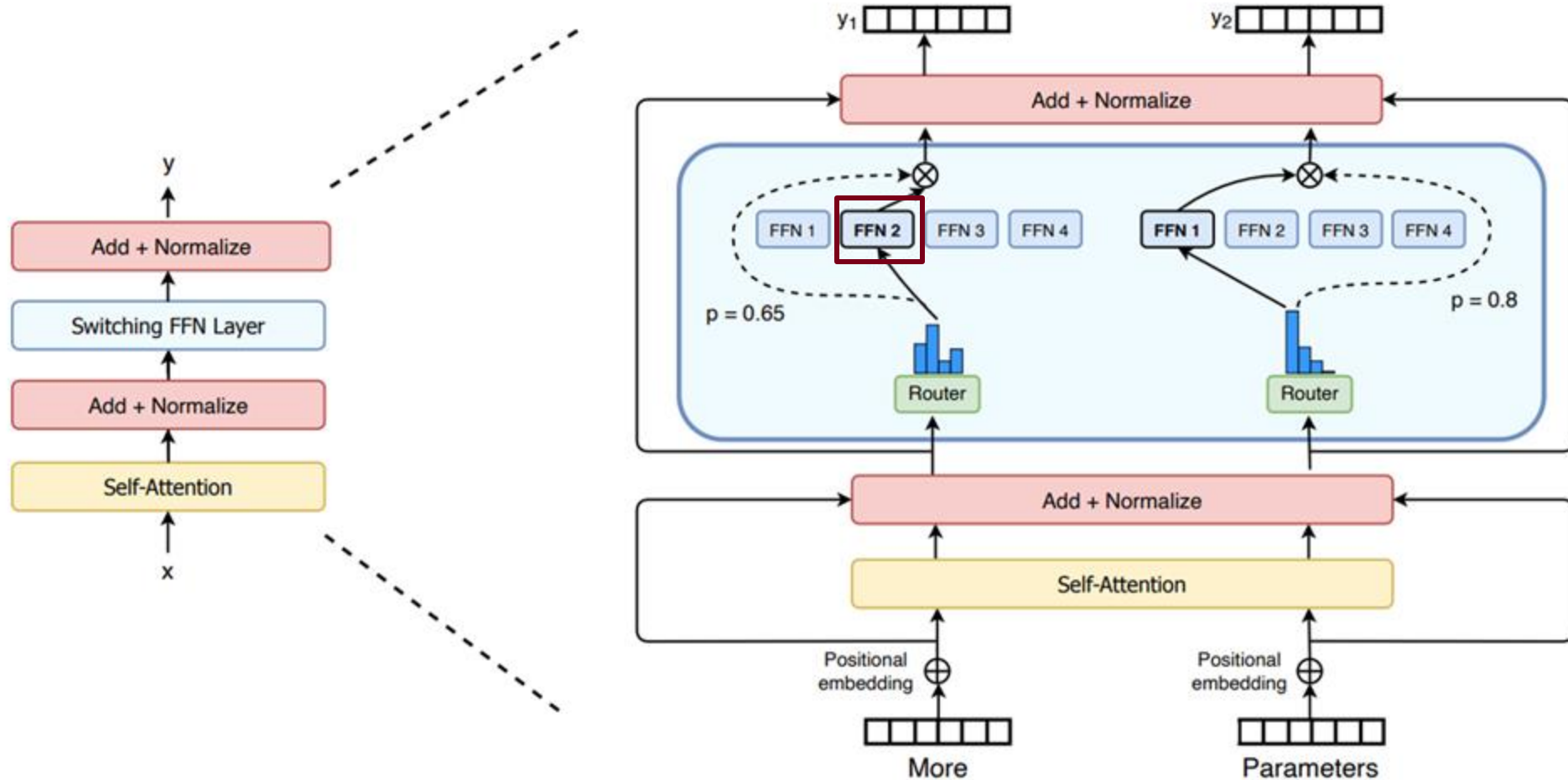
Only one is used per token



Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity [Fedus et al., CoRR 2021]

Mixture of Experts (MoE)

Only 25% of the FFN parameters are used for a single token



Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity [Fedus et al., CoRR 2021]

Mixture of Experts (MoE)

Most of the best-performing open source models produced today are mixture of experts models

 Qwen/Qwen3-235B-A22B-Thinking-2507-FP8

 Text Generation •  235B • Updated Jul 30 •  17.9k •  64

 Qwen/Qwen3-235B-A22B-Thinking-2507

 Text Generation •  235B • Updated Aug 17 •  47.3k •  •  377

 Qwen/Qwen3-235B-A22B-Instruct-2507-FP8

 Text Generation •  235B • Updated Sep 17 •  97.2k •  134

 Qwen/Qwen3-235B-A22B-Instruct-2507

 Text Generation •  235B • Updated Sep 17 •  108k •  •  714

 Qwen/Qwen3-30B-A3B-Thinking-2507-FP8

 Text Generation •  31B • Updated Jul 30 •  20.2k •  45

 Qwen/Qwen3-30B-A3B-Thinking-2507

 Text Generation •  31B • Updated Aug 17 •  267k •  •  315

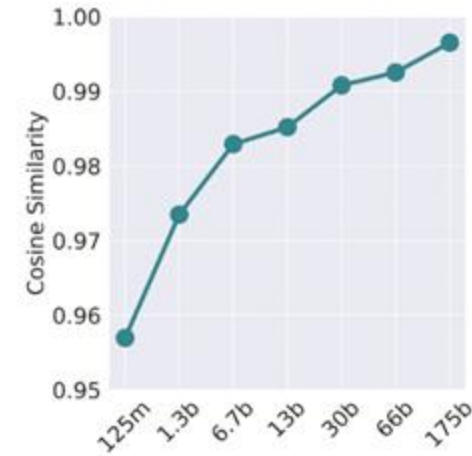
 Qwen/Qwen3-30B-A3B-Instruct-2507-FP8

 Text Generation •  31B • Updated Sep 17 •  486k •  92

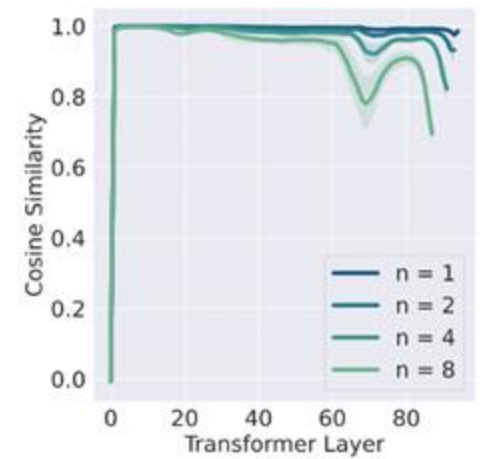


Deja Vu: Contextual Sparsity

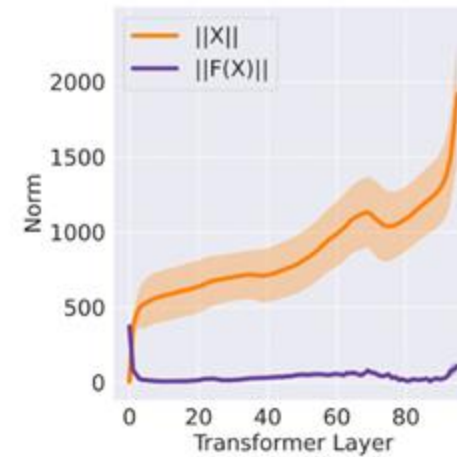
Observation 1: Model activations change very little between consecutive layers of a network



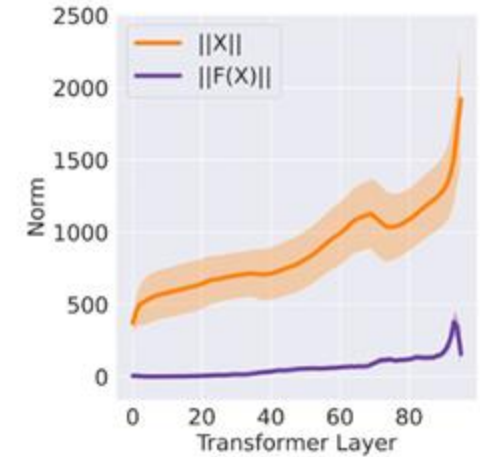
(a) Model Comparison



(b) Across Layer



(c) Residual Around Attention



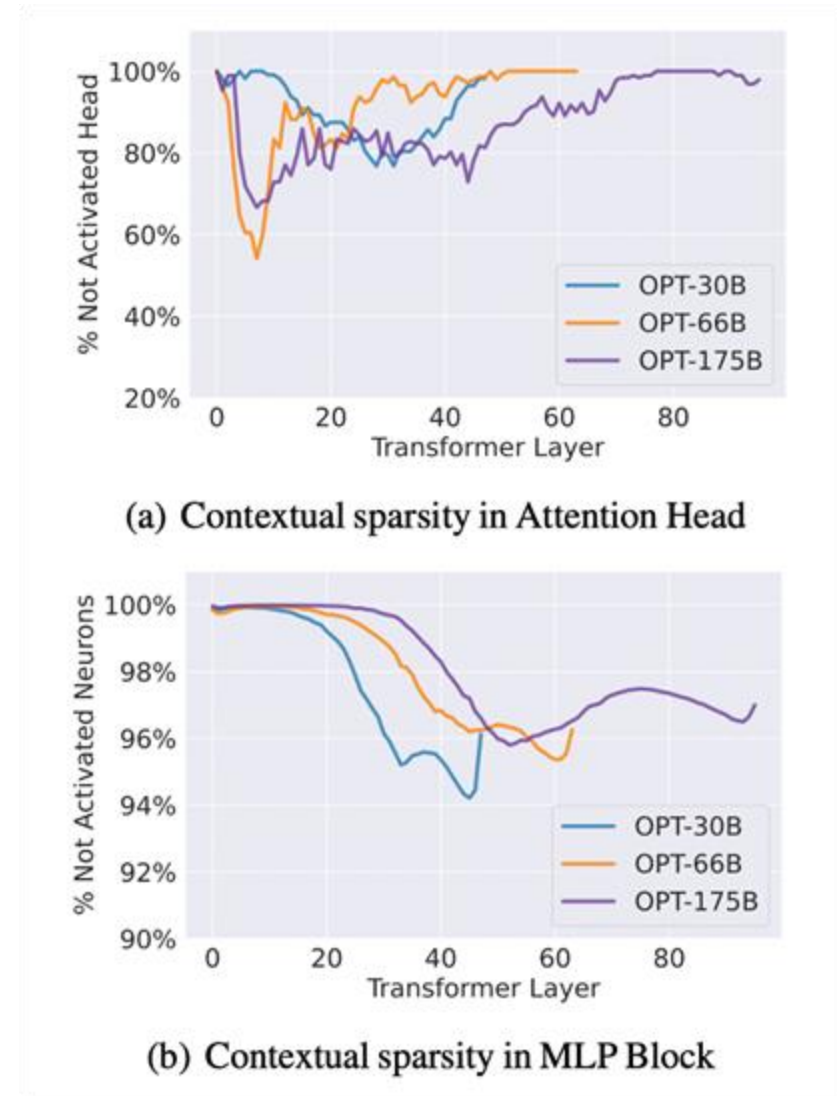
(d) Residual Around MLP

Deja Vu: Contextual Sparsity for Efficient LLMs at Inference Time [Liu et al., 2023]



Deja Vu: Contextual Sparsity

Observation 2: Most attention heads and most neurons are not used

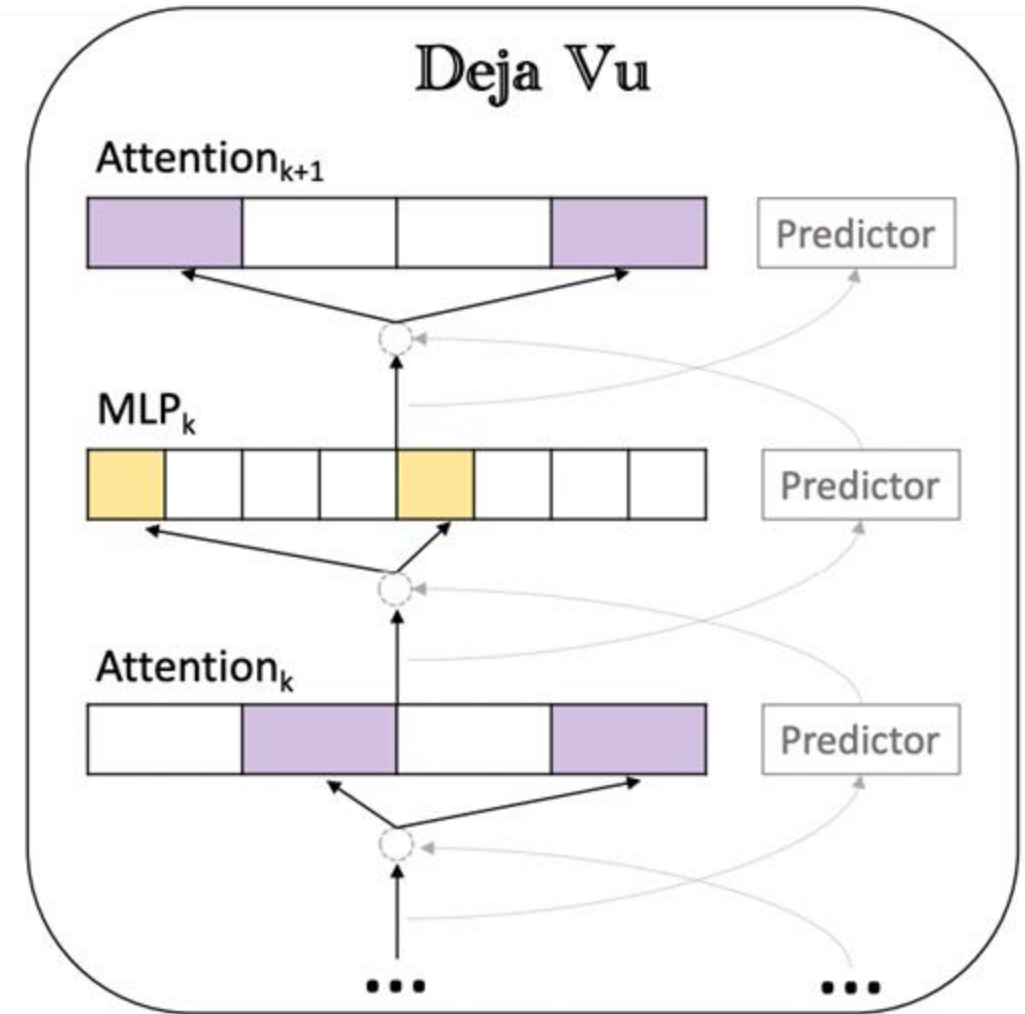


Deja Vu: Contextual Sparsity for Efficient LLMs at Inference Time [Liu et al., 2023]



Deja Vu: Contextual Sparsity

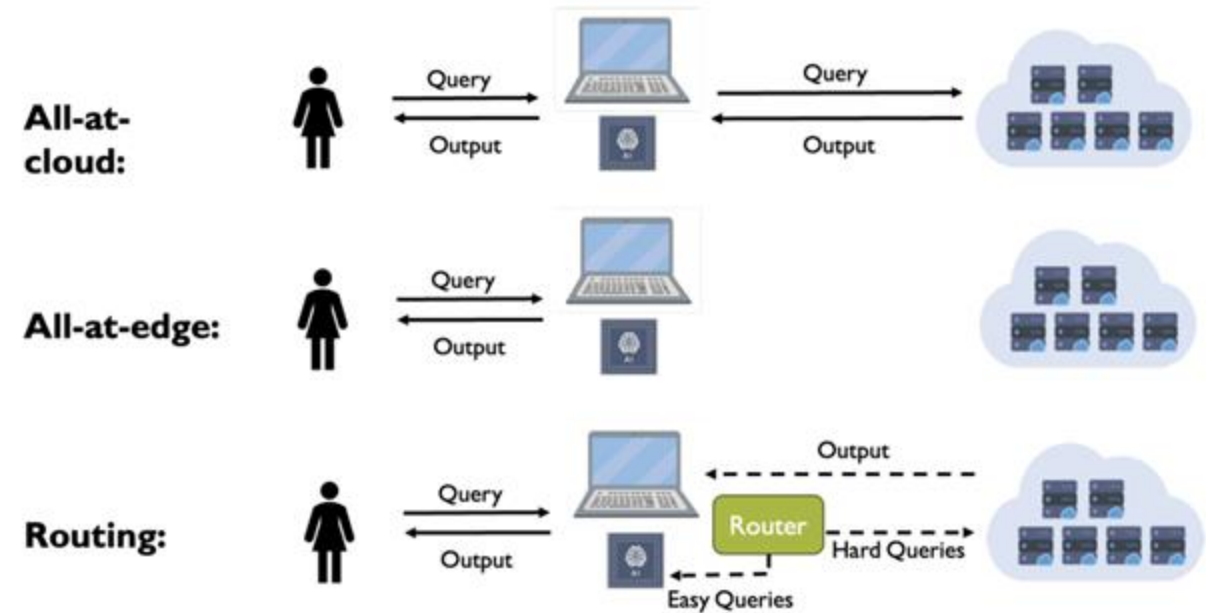
Sparsification: Use predictors in each layer to determine which neurons to activate and which attention heads to use – ignore all unpredicted heads/neurons



Deja Vu: Contextual Sparsity for Efficient LLMs at Inference Time [Liu et al., 2023]

Model Routing

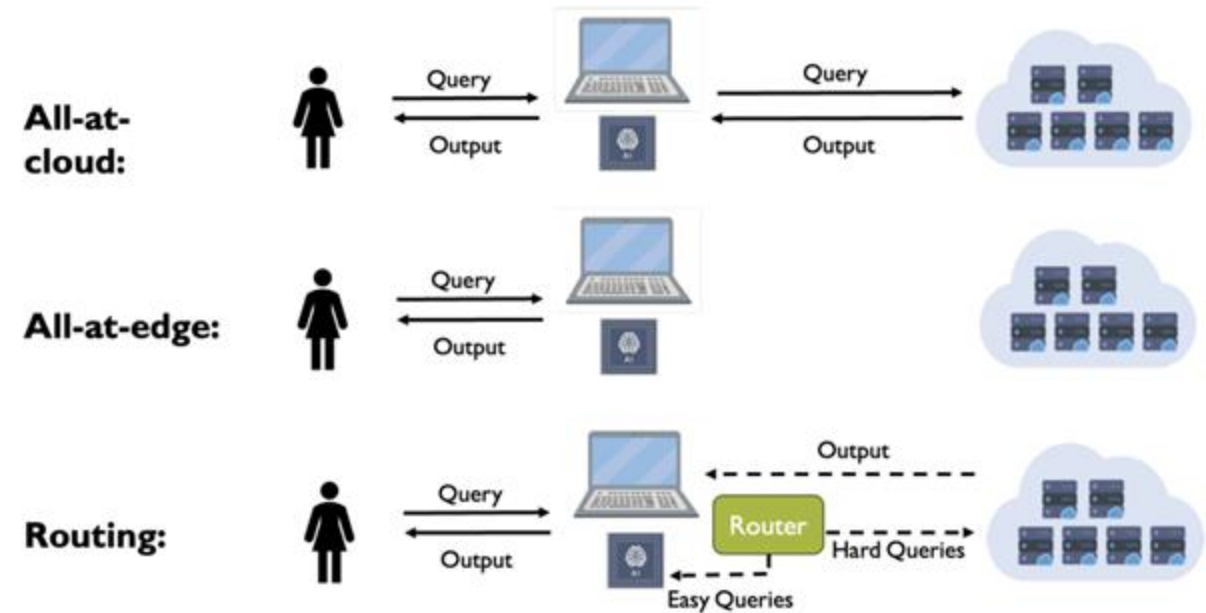
Route easier queries to smaller models, harder queries to larger models



Model Routing

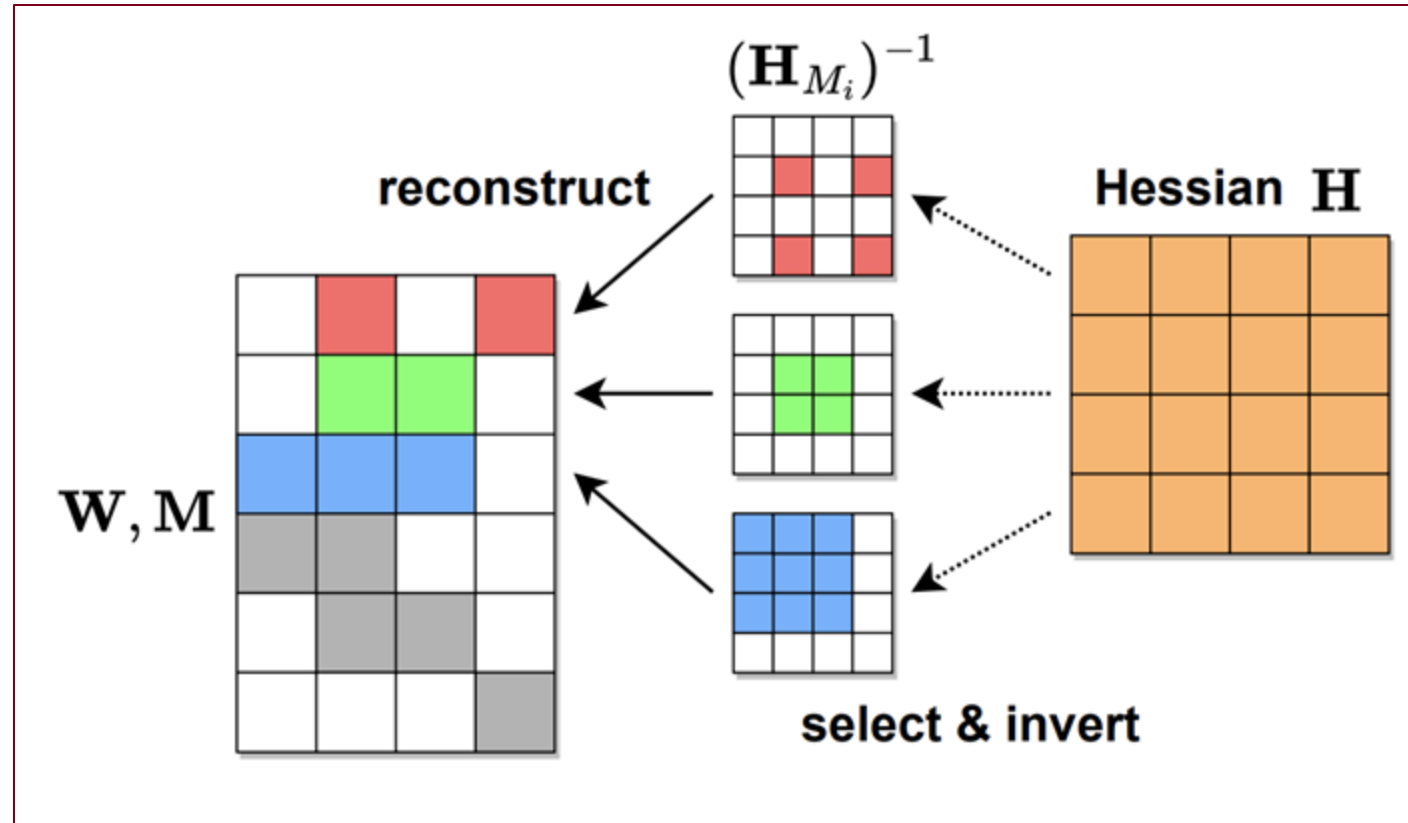
Valuable for companies like OpenAI/Anthropic when they serve multiple models of different capabilities & sizes

Route easier queries to smaller models, harder queries to larger models



SparseGPT

- Similar to GPTQ, uses a reconstruction loss term to prune entries from each weight matrix
- This can be taken advantage of with either 2:4, 4:8 sparsity on Nvidia GPUs



Efficient LLMs

- ☐ Quantization
- ☐ Sparsity
- ☐ **Long Context**
- ☐ Serving & Systems
- ☐ Distillation
- ☐ Parameter Efficient Fine-Tuning
- ☐ Alternative Paradigms



What is a central issue with having longer context models?



$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d}}\right) V$$

What is a central issue with having longer context models?

The storage costs are linear in generation length and quadratic in computation



$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d}}\right) V$$

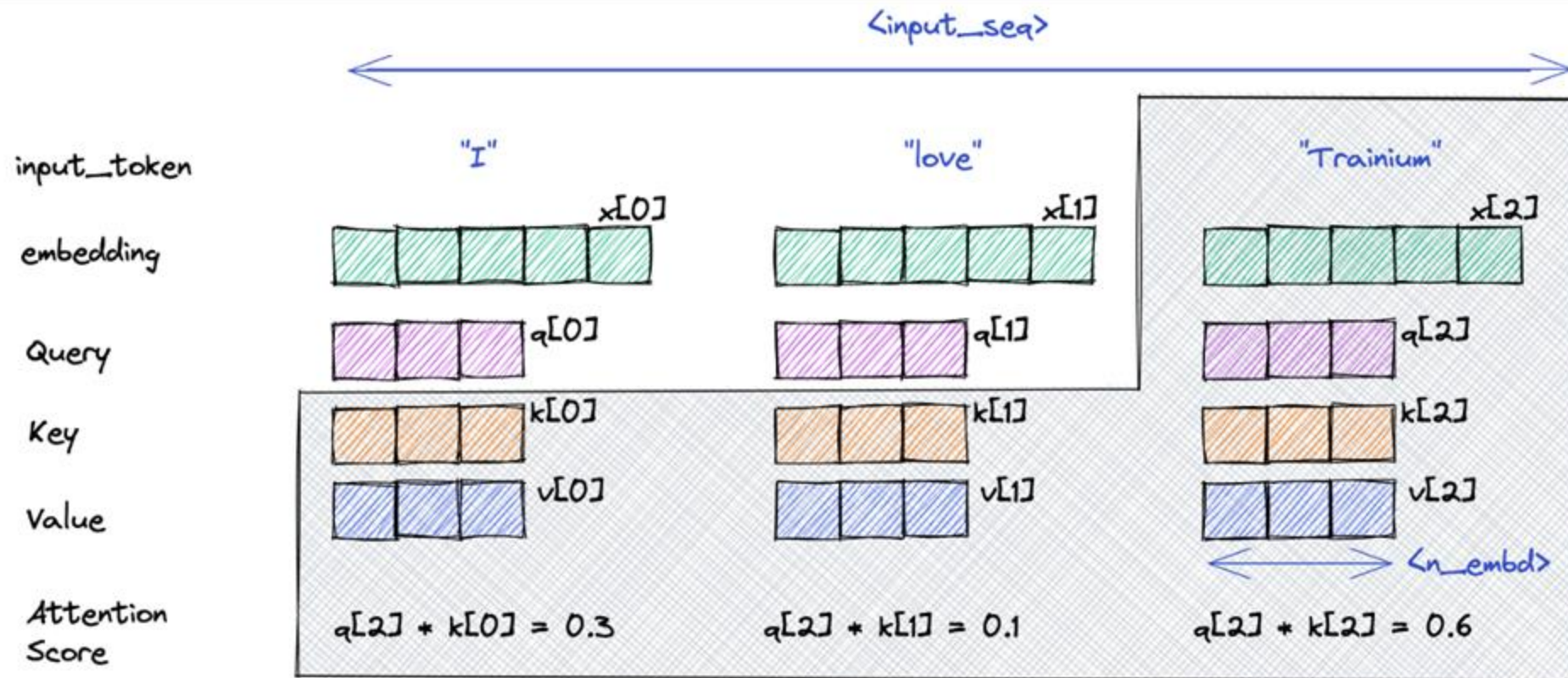
What is a central issue with having longer context models?

The **storage** costs are linear in generation length and quadratic in **computation**



The KV-Cache

The transformer needs to have access to the keys and values for all previous tokens in all layers for all heads when



<https://awsdocs-neuron.readthedocs-hosted.com/en/latest/general/appnotes/transformers-neuronx/generative-llm-inference-with-neuron.html>

The KV-Cache

In total, we must store

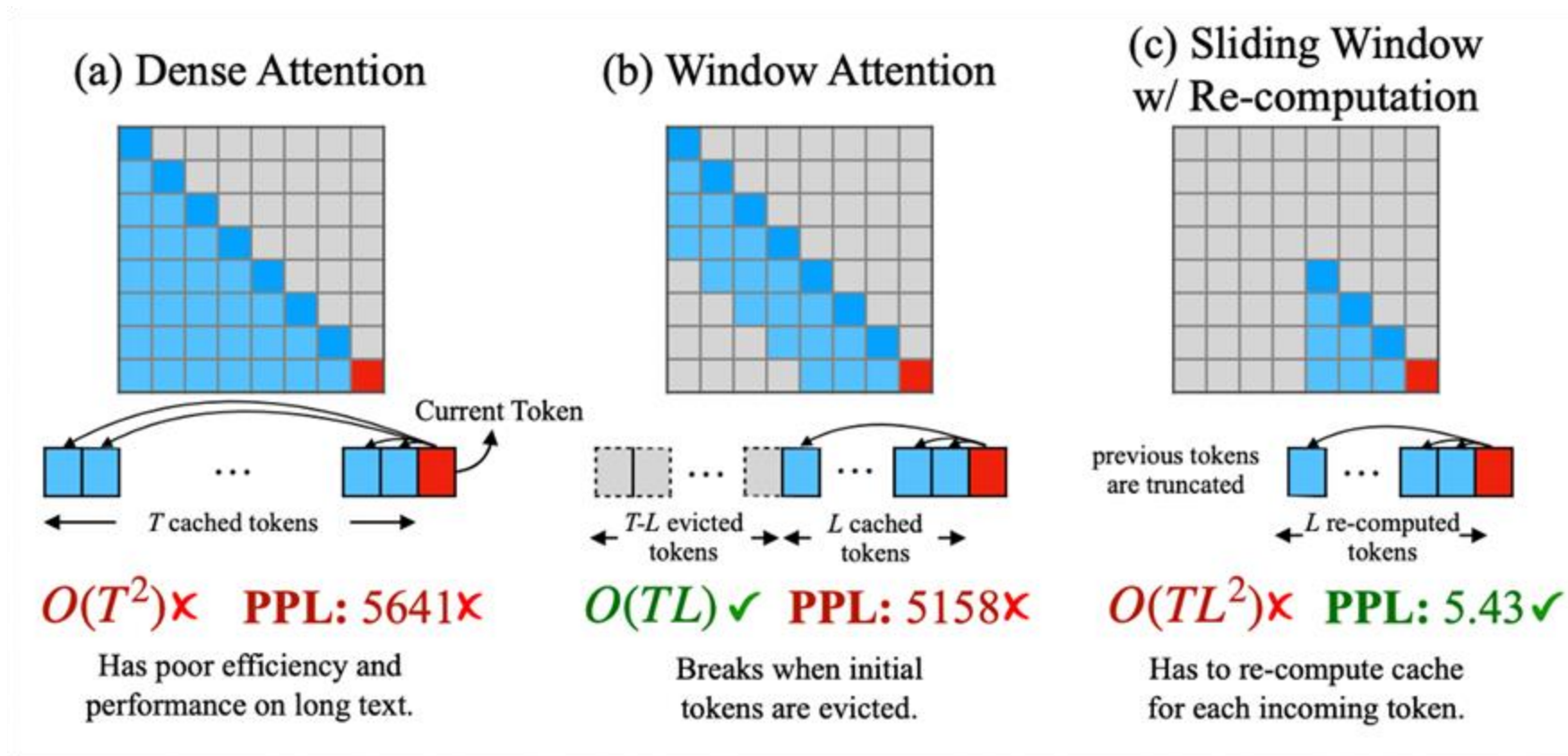
$$\text{Batch_size} * \text{seq_len} * \text{num_heads} * \text{num_layers} * \text{emb_dim} * 2$$

separate values in the kv cache



StreamingLLM

How can we extend models to have much longer context length at minimal cost?

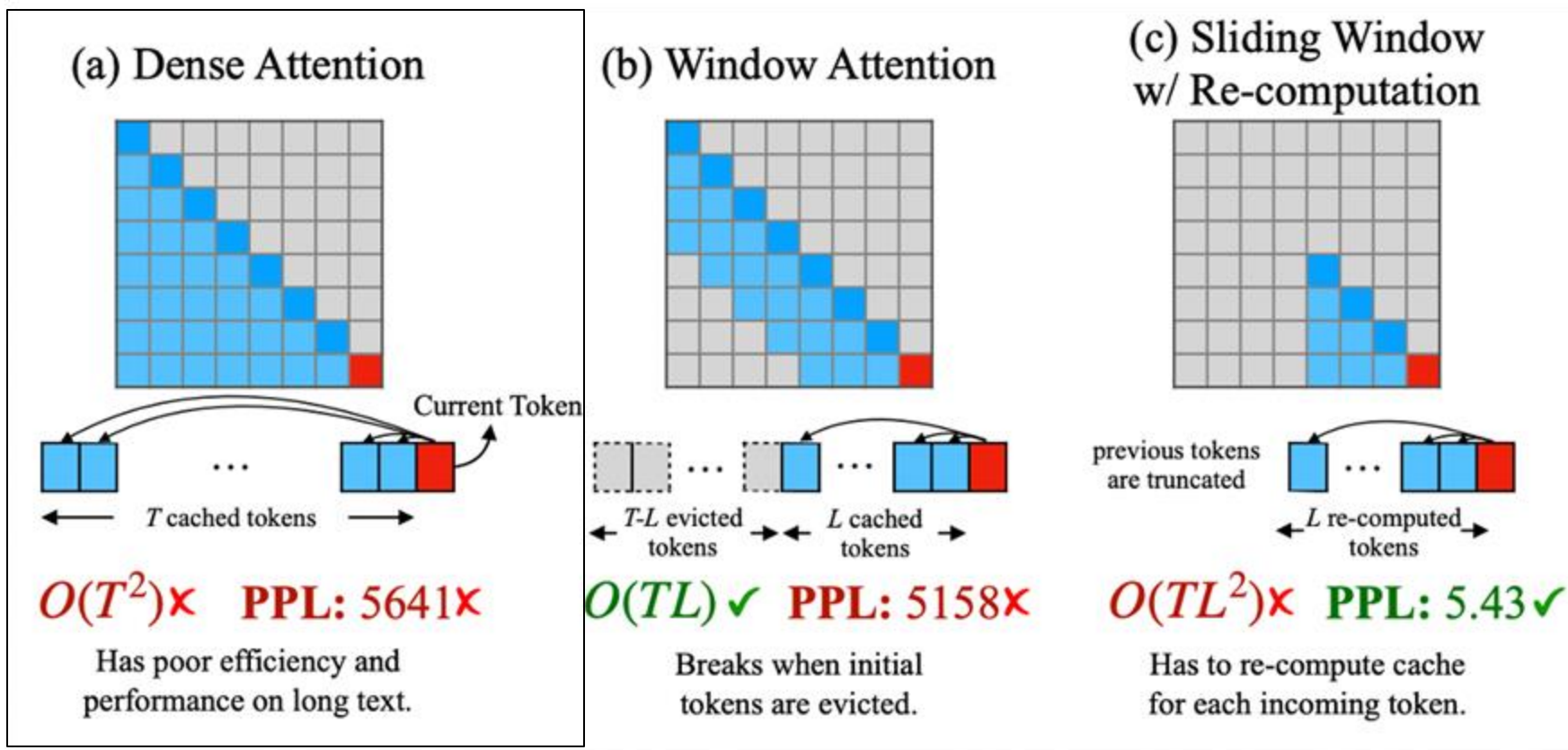


Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

StreamingLLM

How can we extend models to have much longer context length at minimal cost?

Too
much
storage

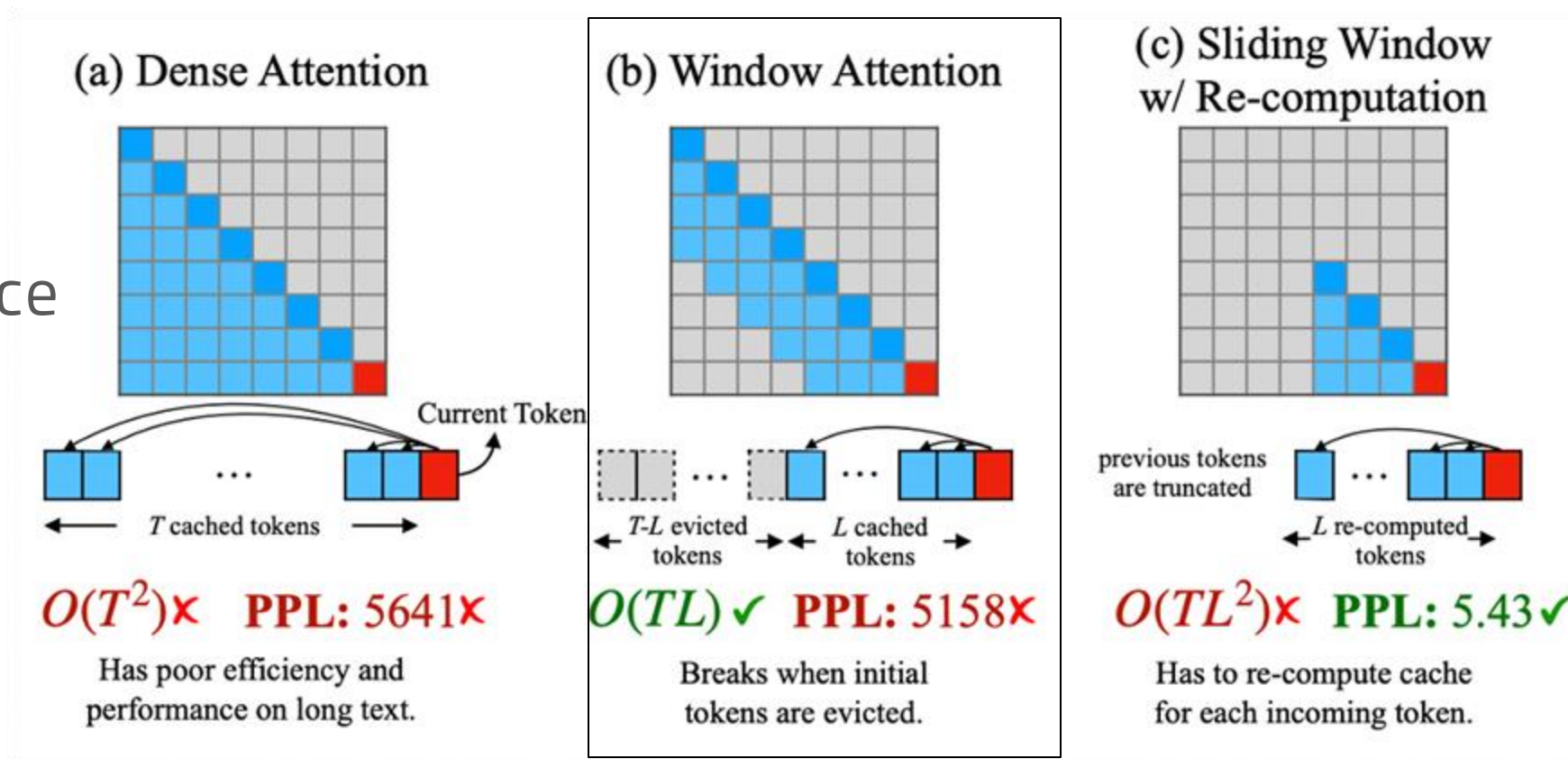


Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

StreamingLLM

How can we extend models to have much longer context length at minimal cost?

Bad
performance

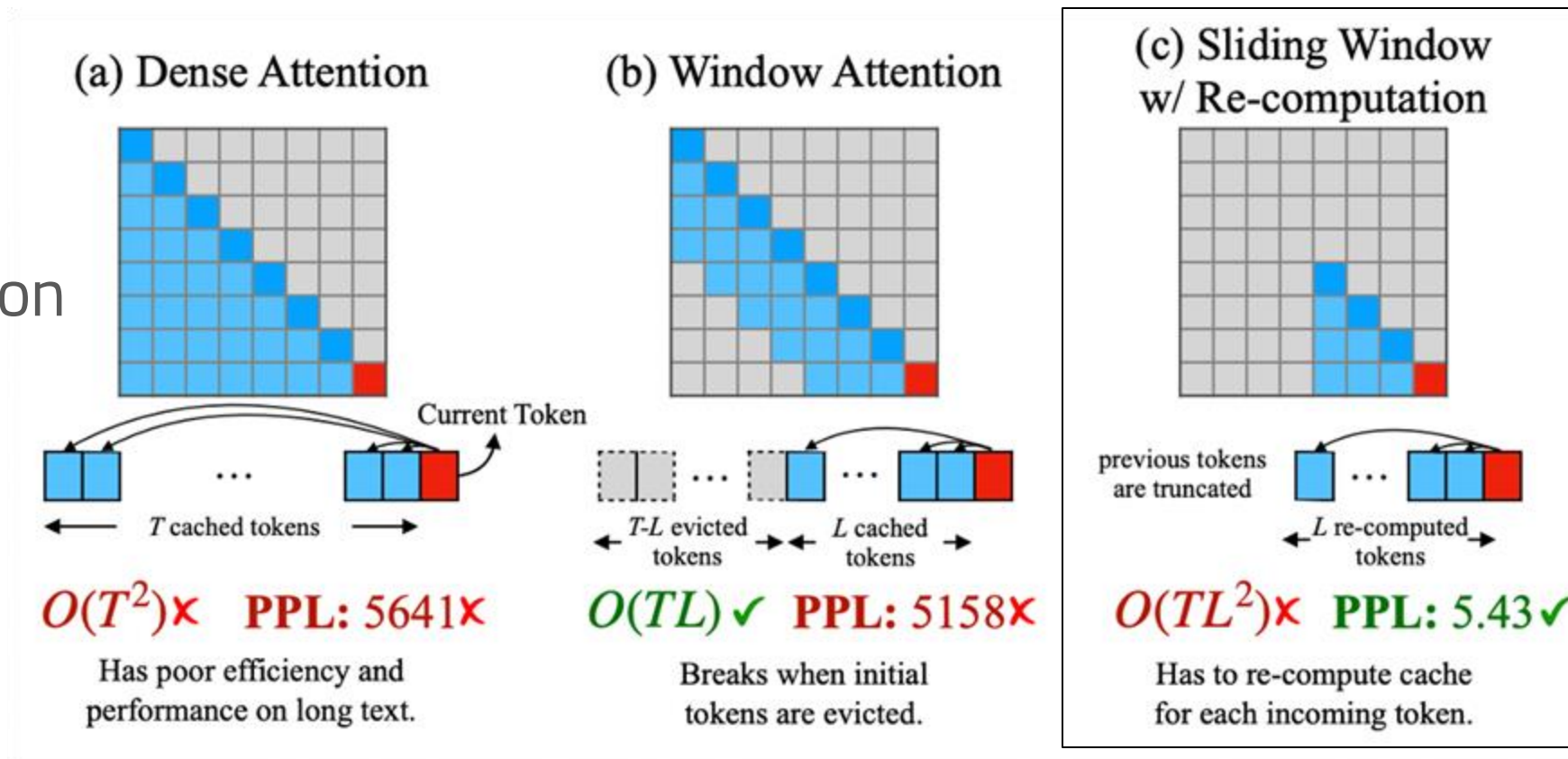


Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

StreamingLLM

How can we extend models to have much longer context length at minimal cost?

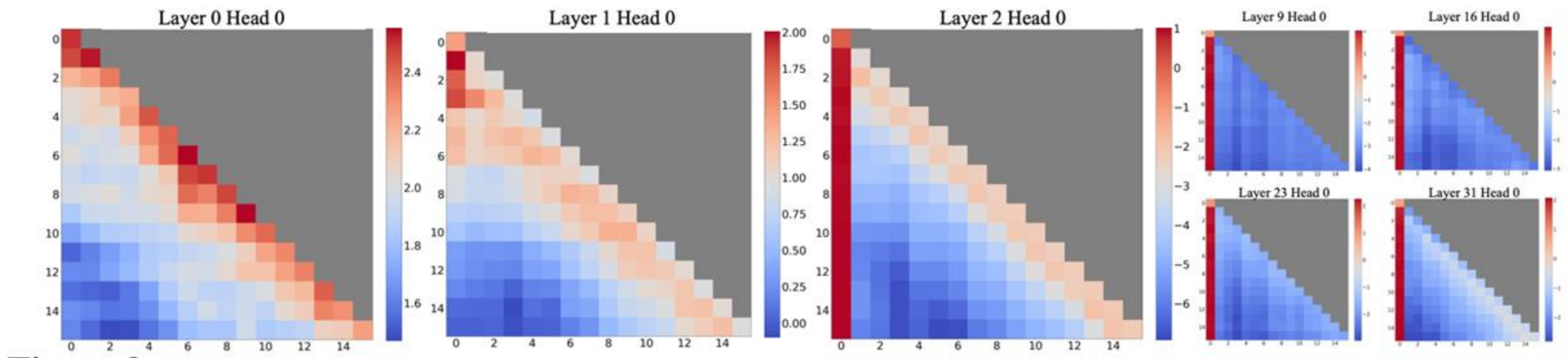
Too much
recomputation



Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

StreamingLLM

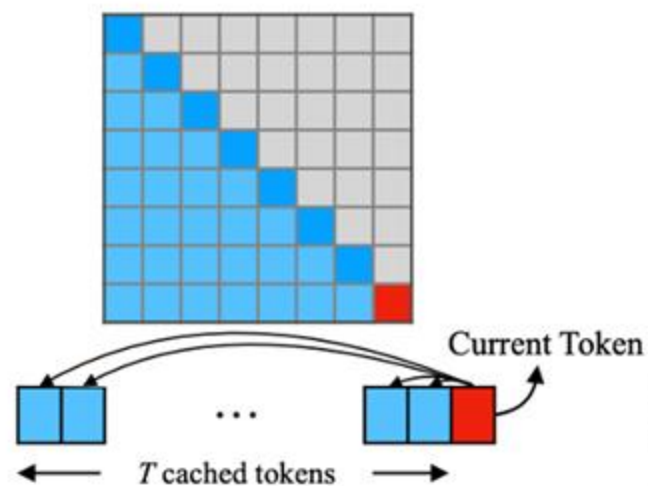
Observation: Most attention is either placed on the first token or to tokens that the model has recently seen.



Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

StreamingLLM

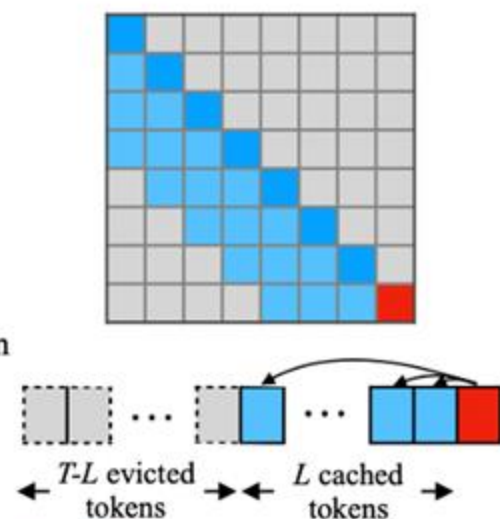
(a) Dense Attention



$O(T^2)$ ✗ PPL: 5641✗

Has poor efficiency and performance on long text.

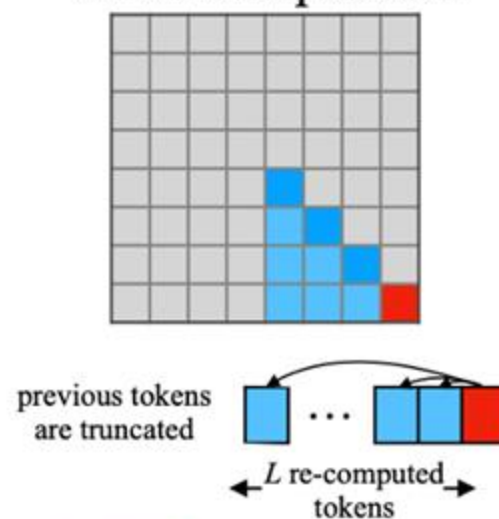
(b) Window Attention



$O(TL)$ ✓ PPL: 5158✗

Breaks when initial tokens are evicted.

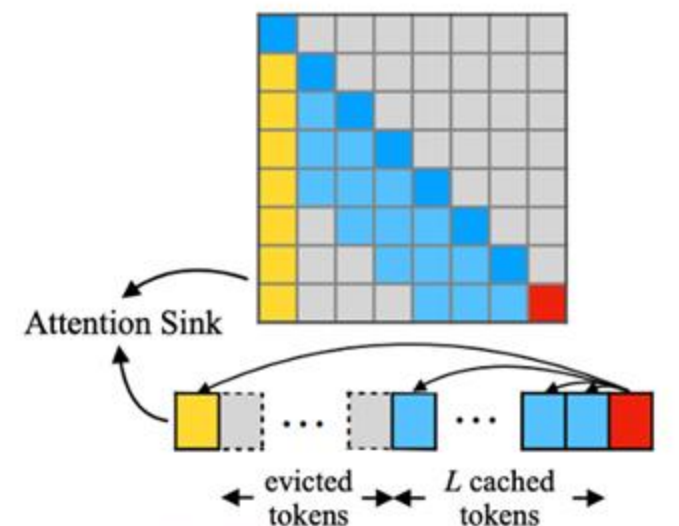
(c) Sliding Window w/ Re-computation



$O(TL^2)$ ✗ PPL: 5.43✓

Has to re-compute cache for each incoming token.

(d) StreamingLLM (ours)

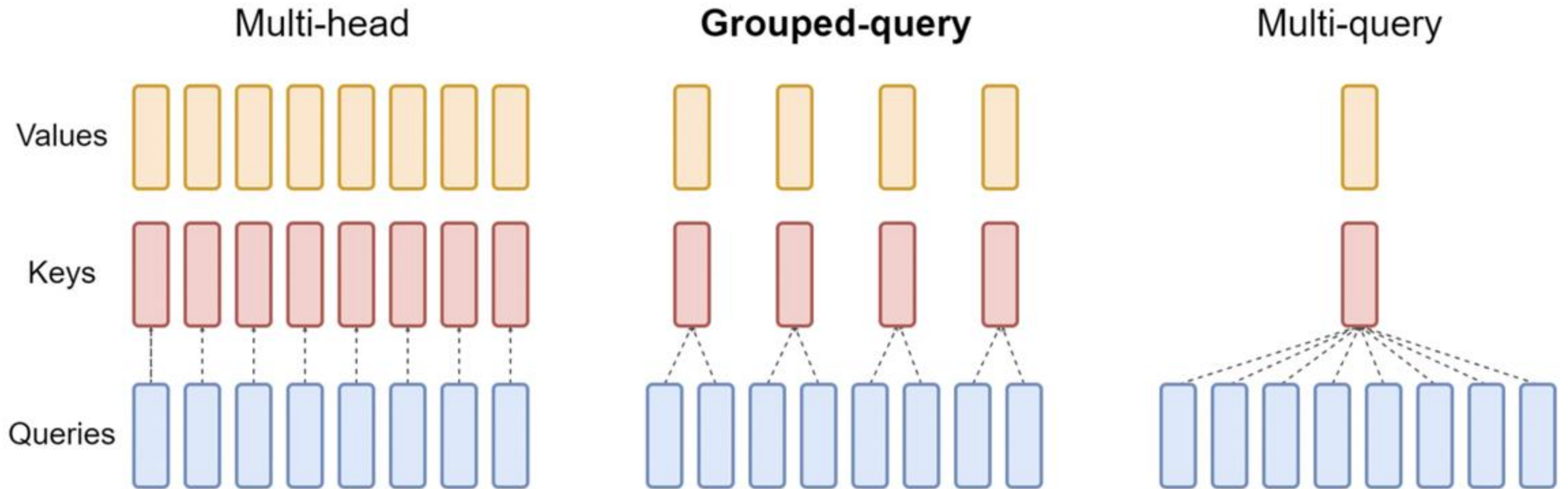


$O(TL)$ ✓ PPL: 5.40✓

Can perform efficient and stable language modeling on long texts.

Efficient Streaming Language Models with Attention Sinks [Xiao et al., 2023]

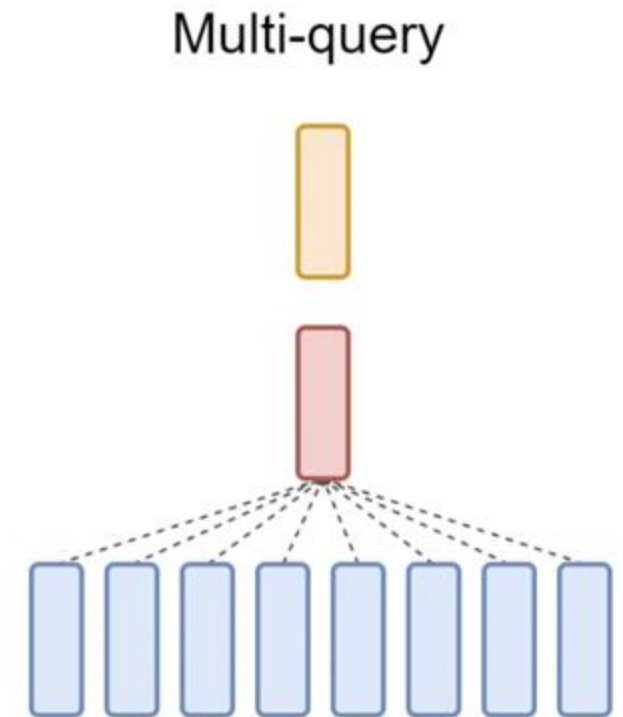
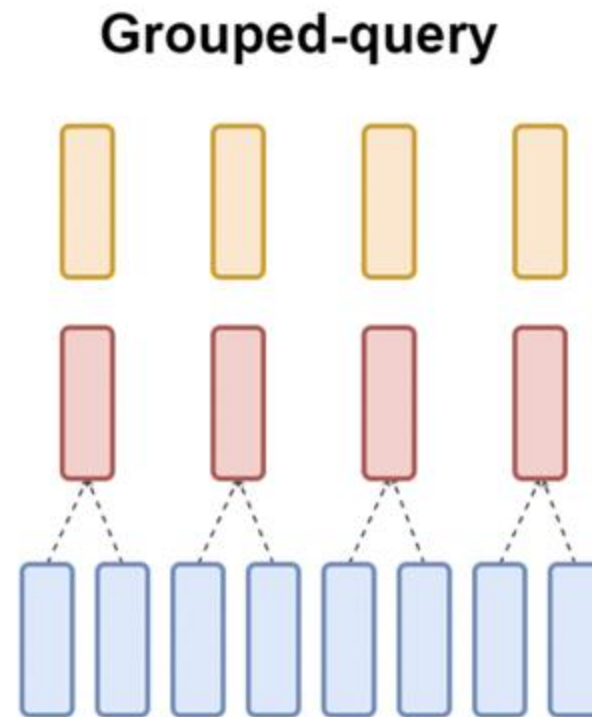
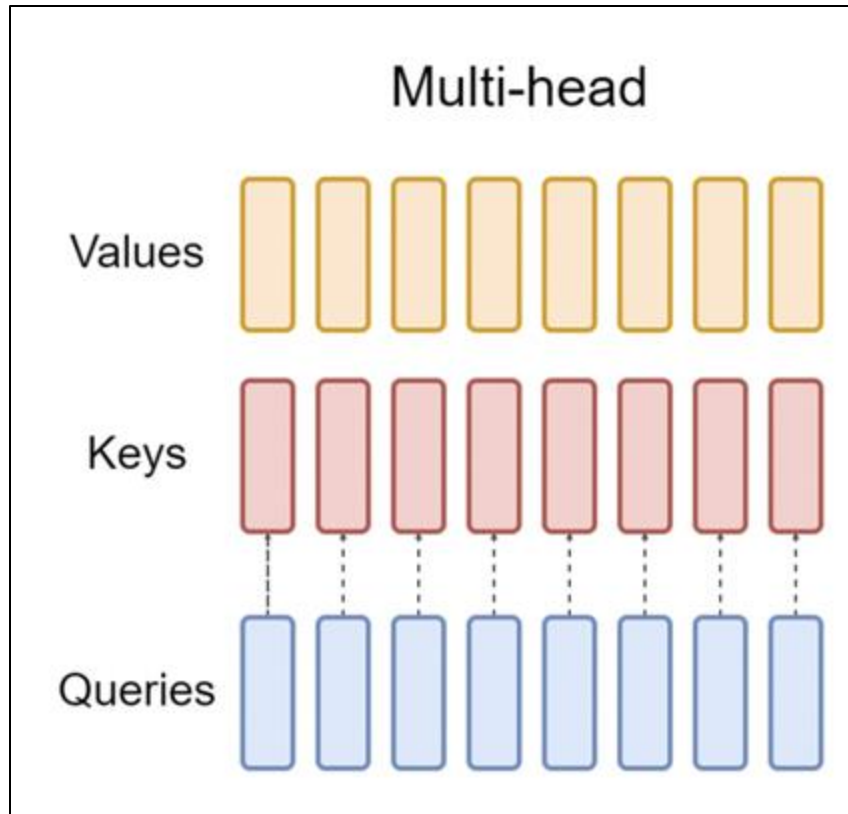
MHA/GQA/MQA



GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints

MHA/GQA/MQA

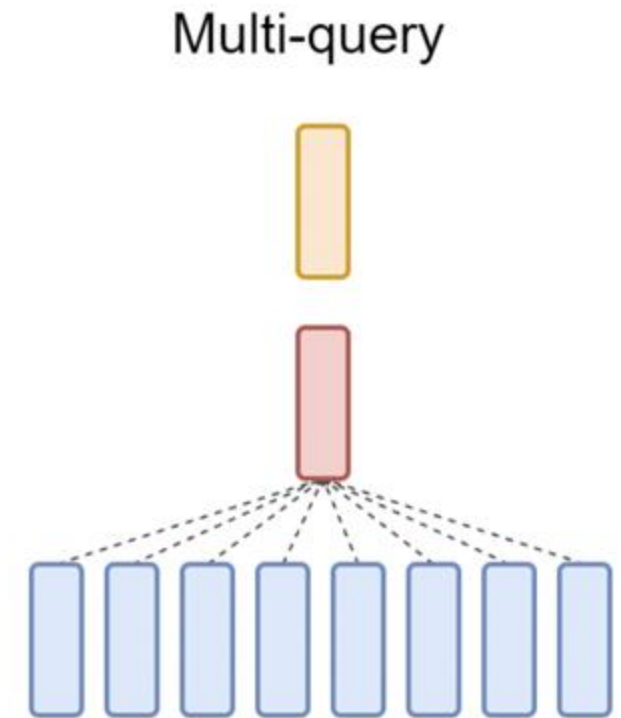
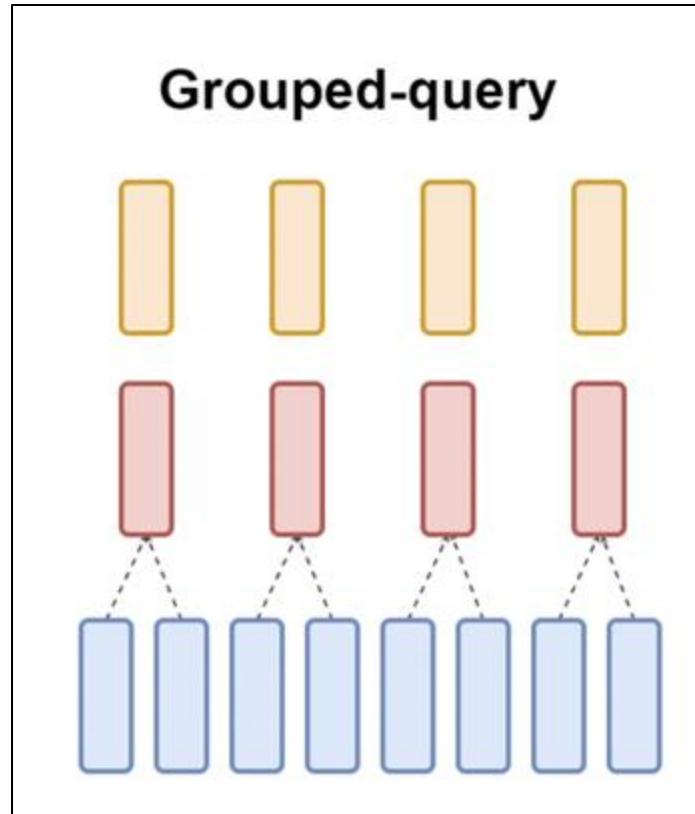
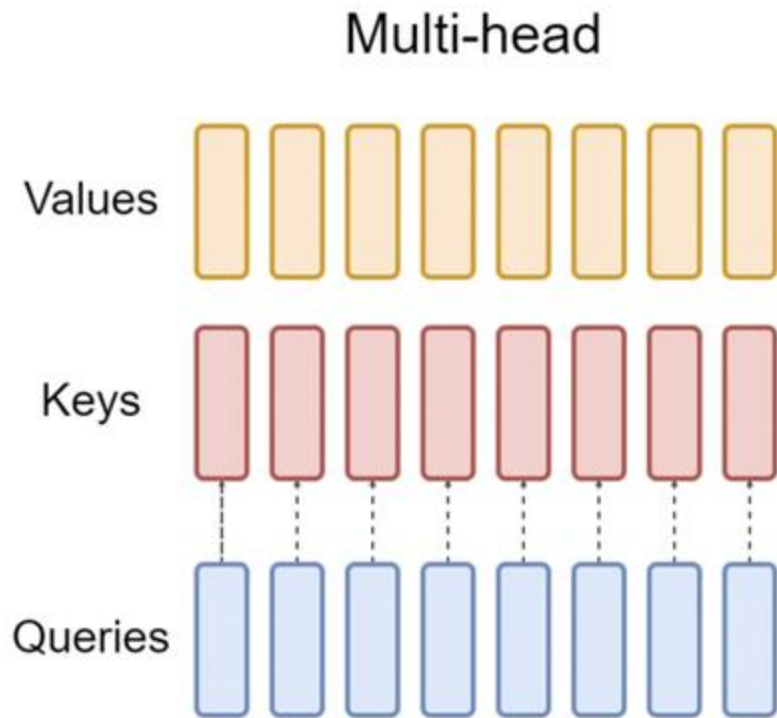
Each attention head calculates separate keys and values for each token



GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints

MHA/GQA/MQA

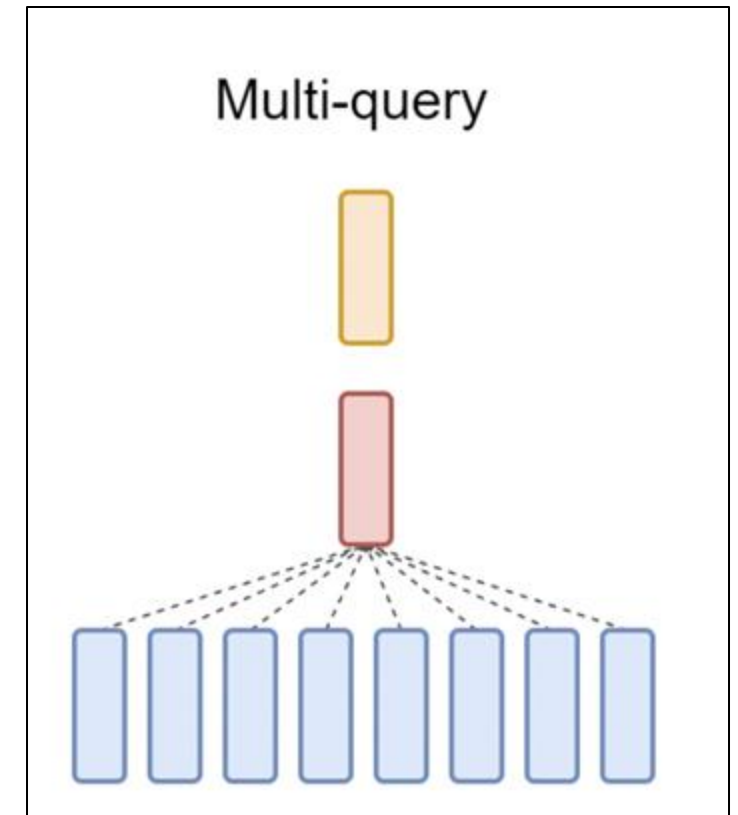
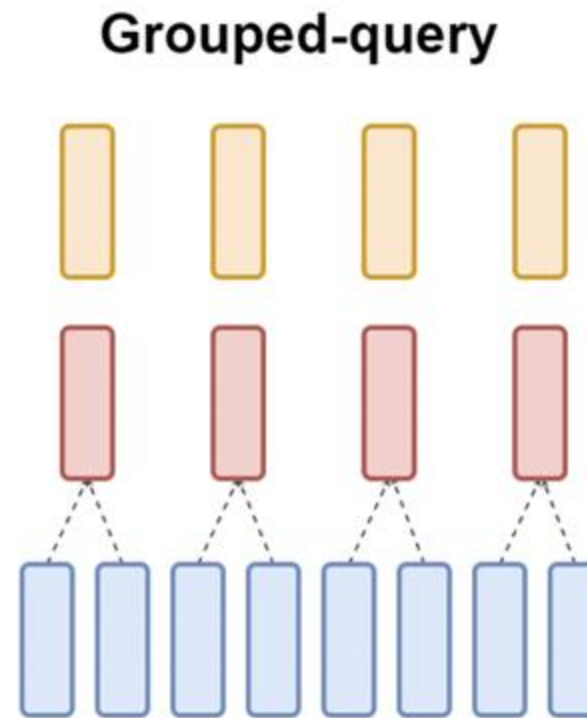
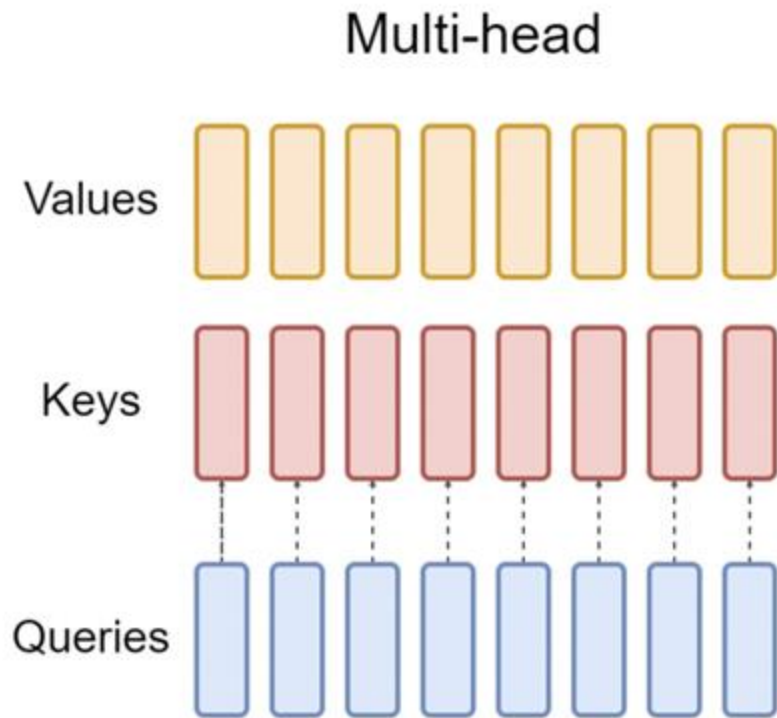
Attention heads are split into groups.
Each group has one key/value per token.



GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints

MHA/GQA/MQA

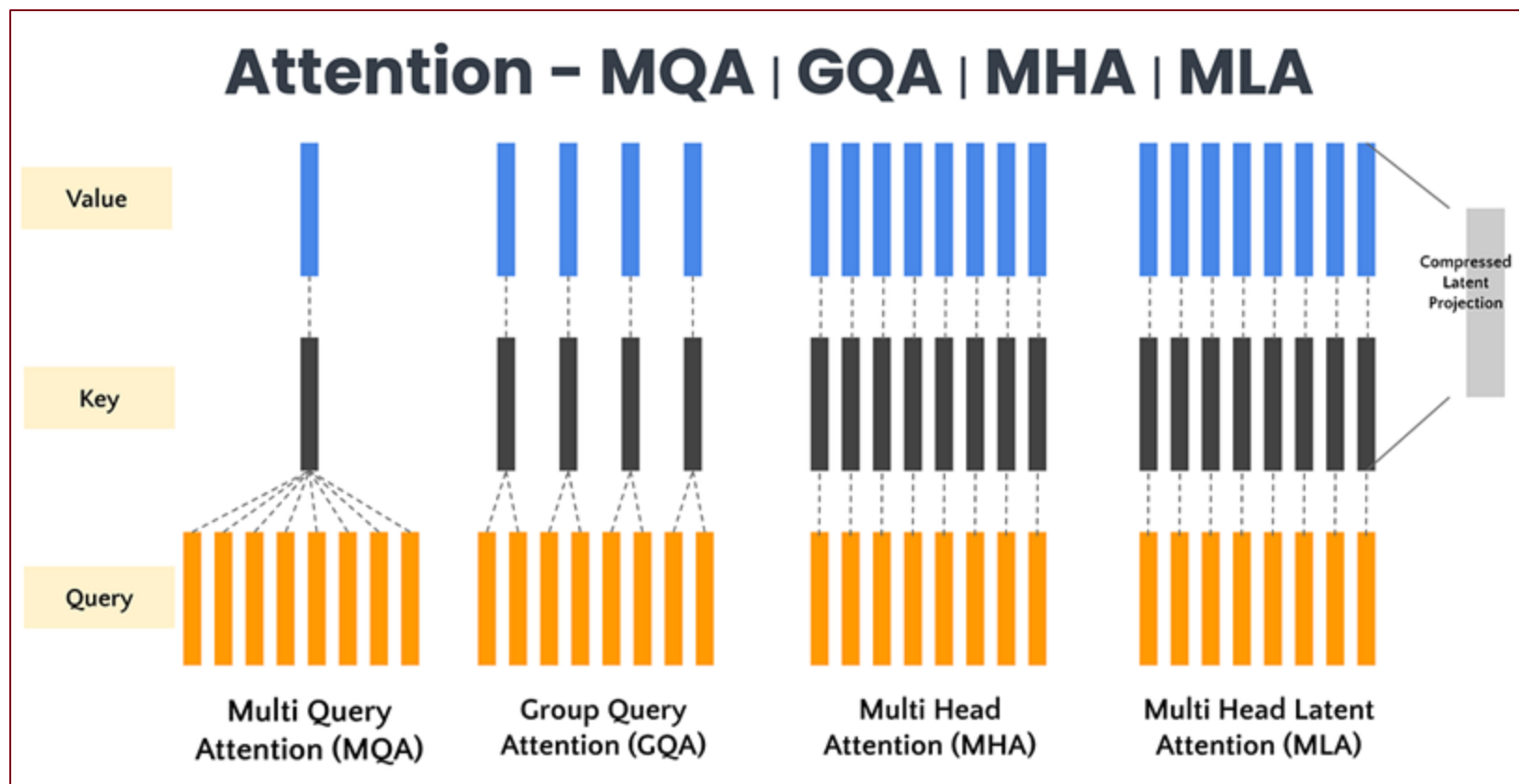
Attention heads share the same keys and values for each token



GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints

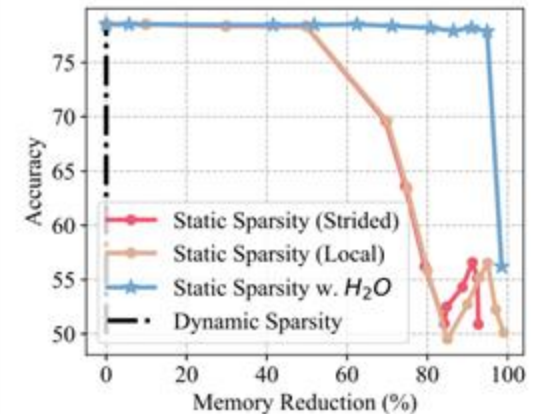
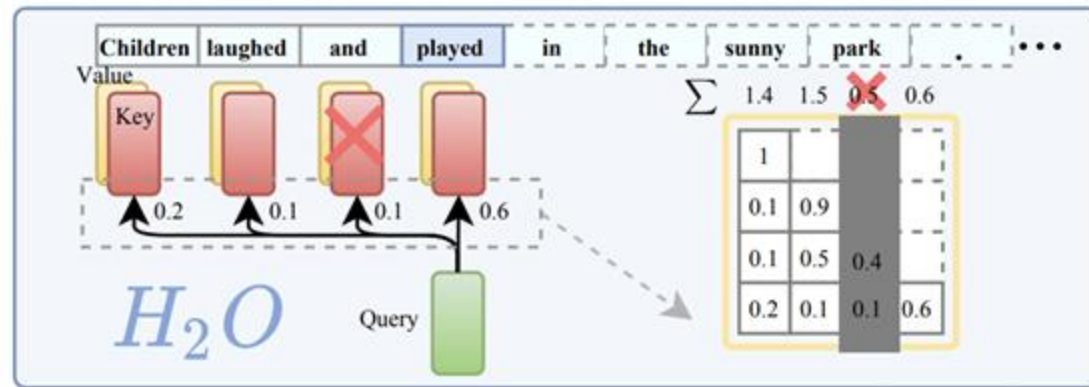
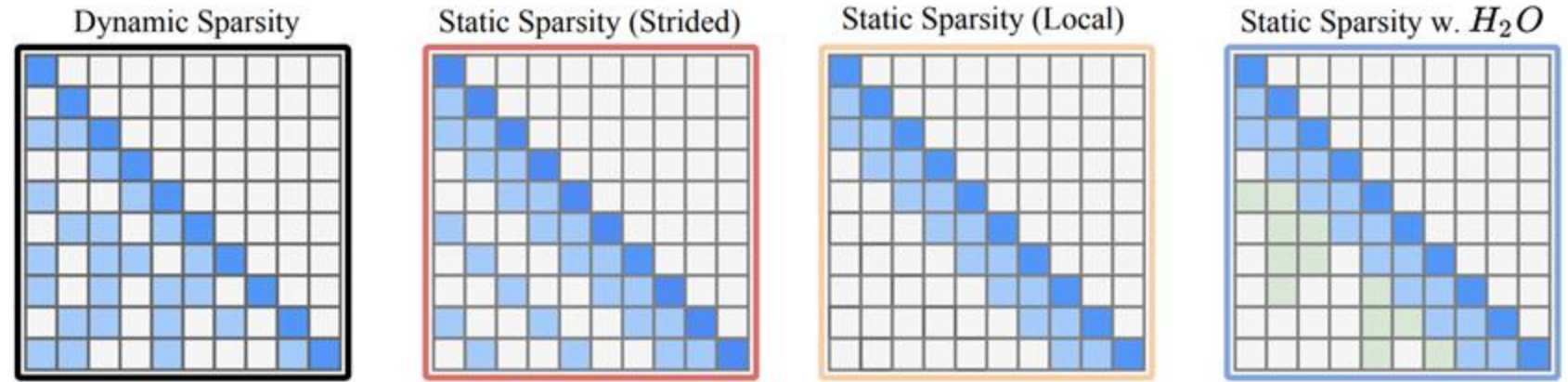
MHA Latent Query Attention

Some modern models (such as DeepSeek) train with what is called latent query attention → Keys + Values are compressed to a lower dimension to reduce storage costs



H₂O: Heavy Hitter Oracle

Evict all but k-highest cumulative attention tokens from cache

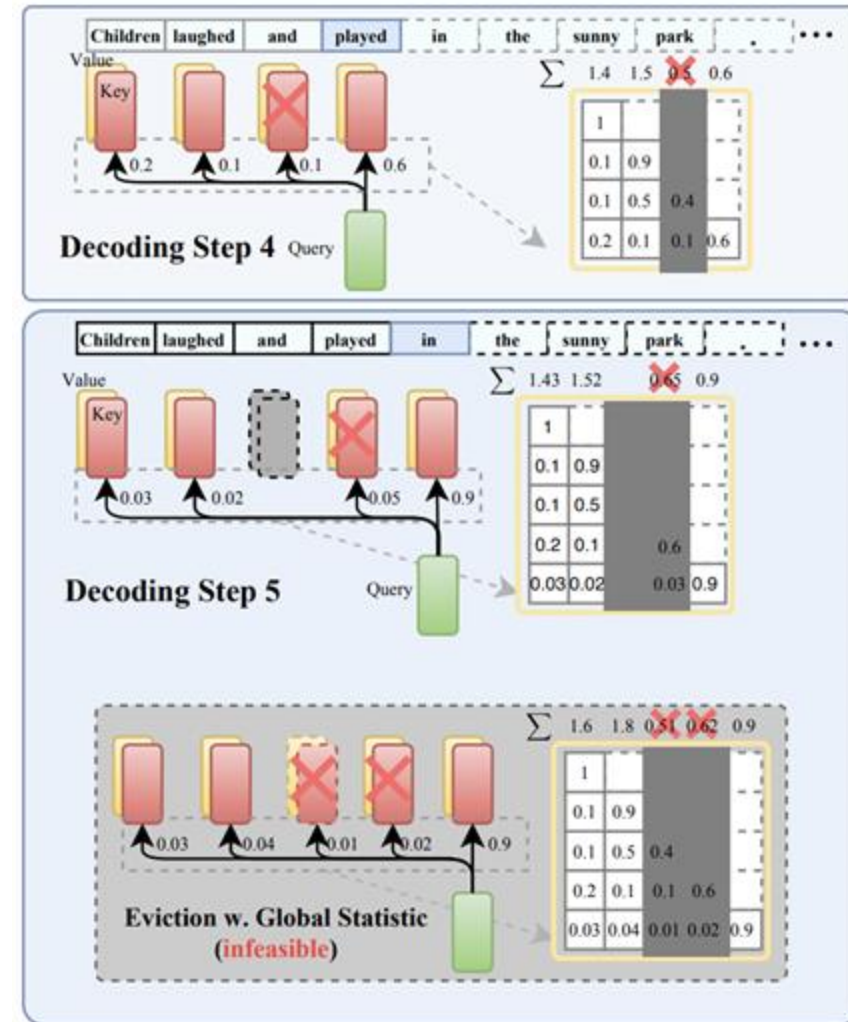


H2O: Heavy-Hitter Oracle for Efficient Generative Inference of Large Language Models [Zhang et. al, 2023]



H₂O: Heavy Hitter Oracle

Evict all but k-
highest
cumulative
attention tokens
from cache



H2O: Heavy-Hitter Oracle for Efficient Generative Inference of Large Language Models [Zhang et. al, 2023]

Efficient LLMs

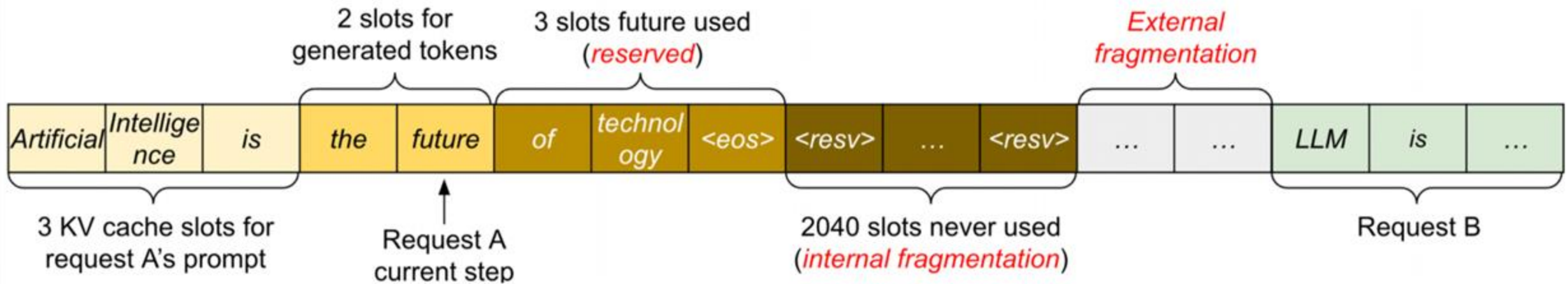
- ❑ Quantization
- ❑ Sparsity
- ❑ Long Context
- ❑ **Serving & Systems**
- ❑ Distillation
- ❑ Parameter Efficient Fine-Tuning
- ❑ Alternative Paradigms



PagedAttention

How does a large LLM service (large ChatGPT) handle multiple incoming requests with respect to the KV-cache?

-Originally, most systems just assign fixed sized blocks of memory to each incoming request. How to improve?

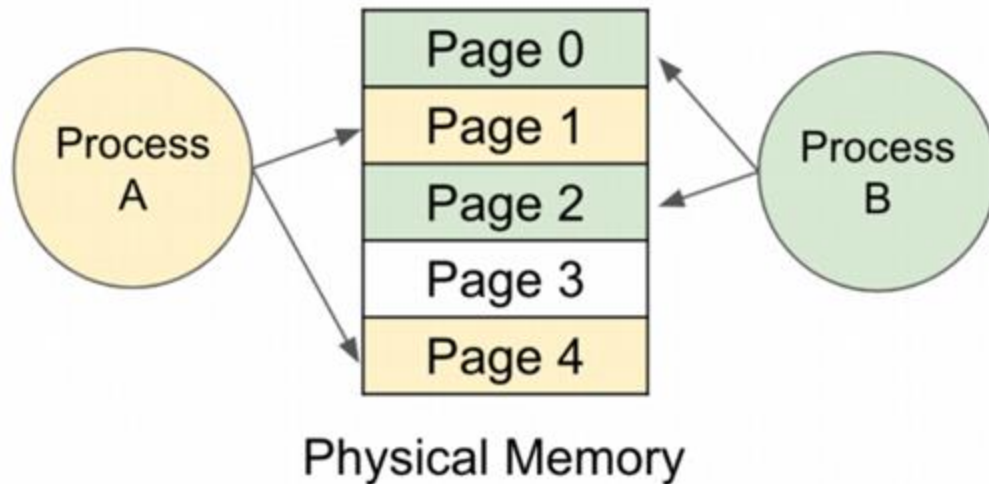


Efficient Memory Management for Large Language Model Serving with PagedAttention (Kwon et al., 2023)

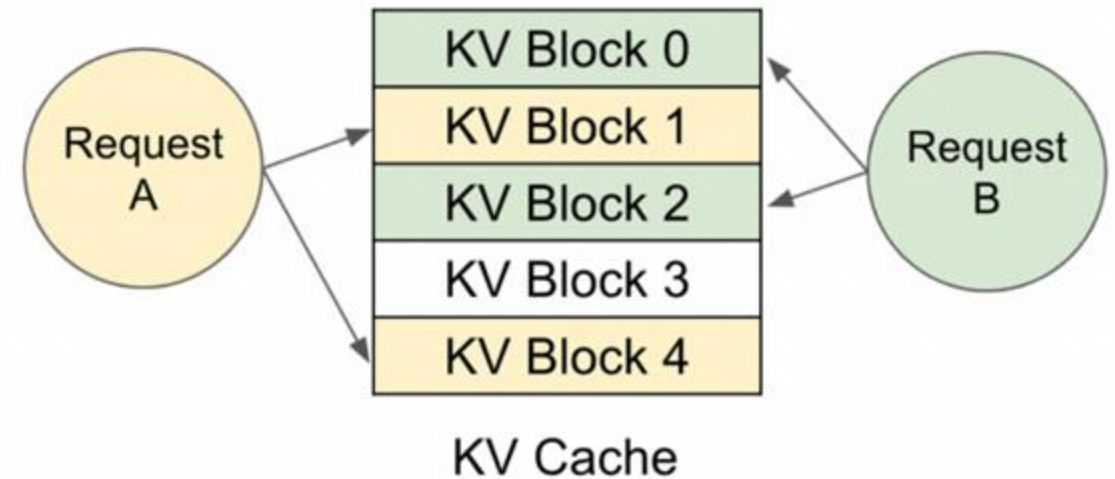
PagedAttention

Let's adopt a similar approach to that found in virtual memory!

Memory management in OS

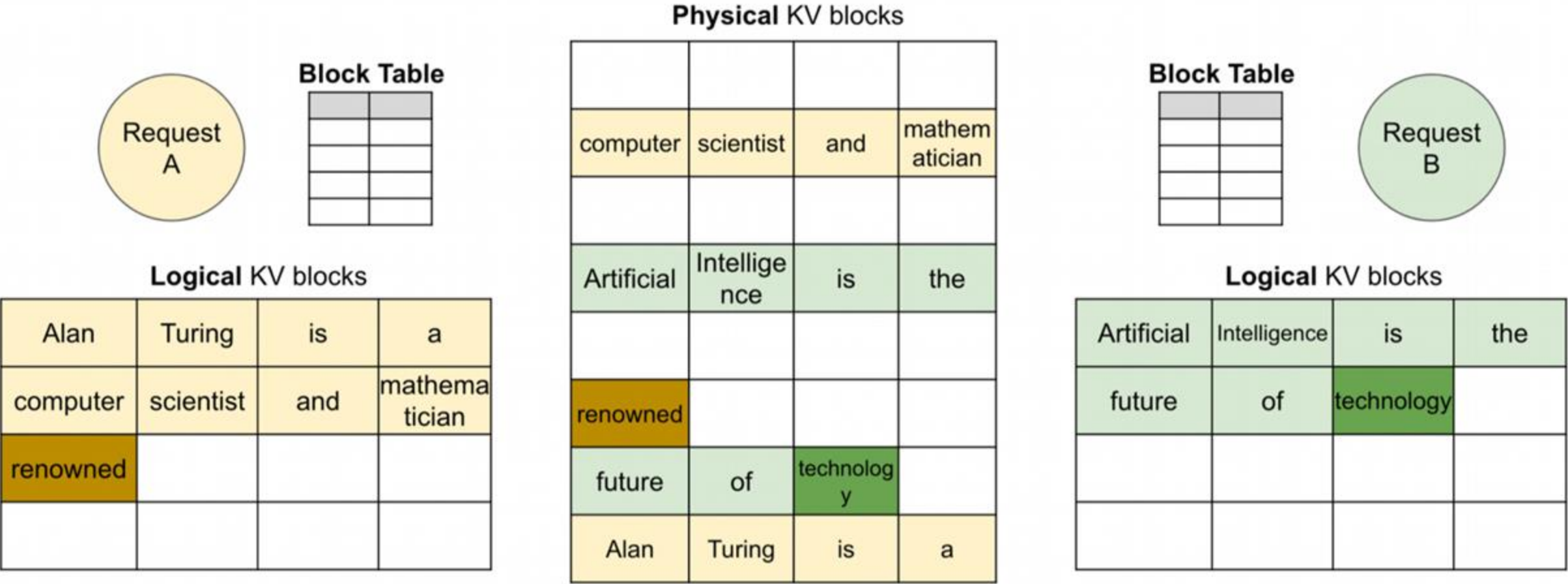


Memory management in vLLM



Efficient Memory Management for Large Language Model Serving with PagedAttention (Kwon et al., 2023)

PagedAttention



Efficient Memory Management for Large Language Model Serving with PagedAttention (Kwon et al., 2023)

Examples of Modern Inference Systems

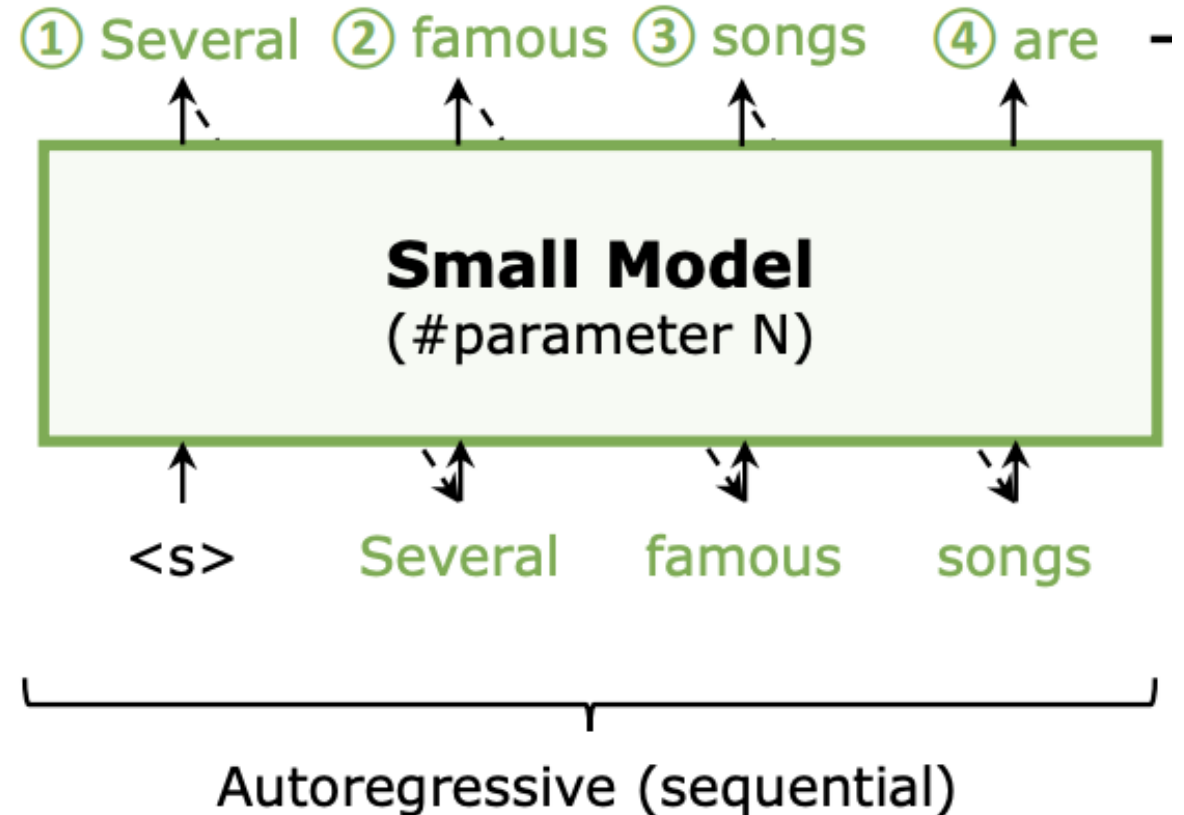
[vLLM](#), [sglang](#), and [TensorRT-LLM](#) are all examples of inference engines which take advantage of this approach.

If you are serving models in any setting, you should use one of the above (rather than directly instantiating some huggingface or pytorch model)



Speculative Decoding

- In standard autoregressive decoding, we are only using each parameter **one time** when the batch size is 1
- This means standard decoding has a **low** arithmetic intensity and is memory bound
- We have a bunch more compute we could be getting for free given how massively parallel GPUs are

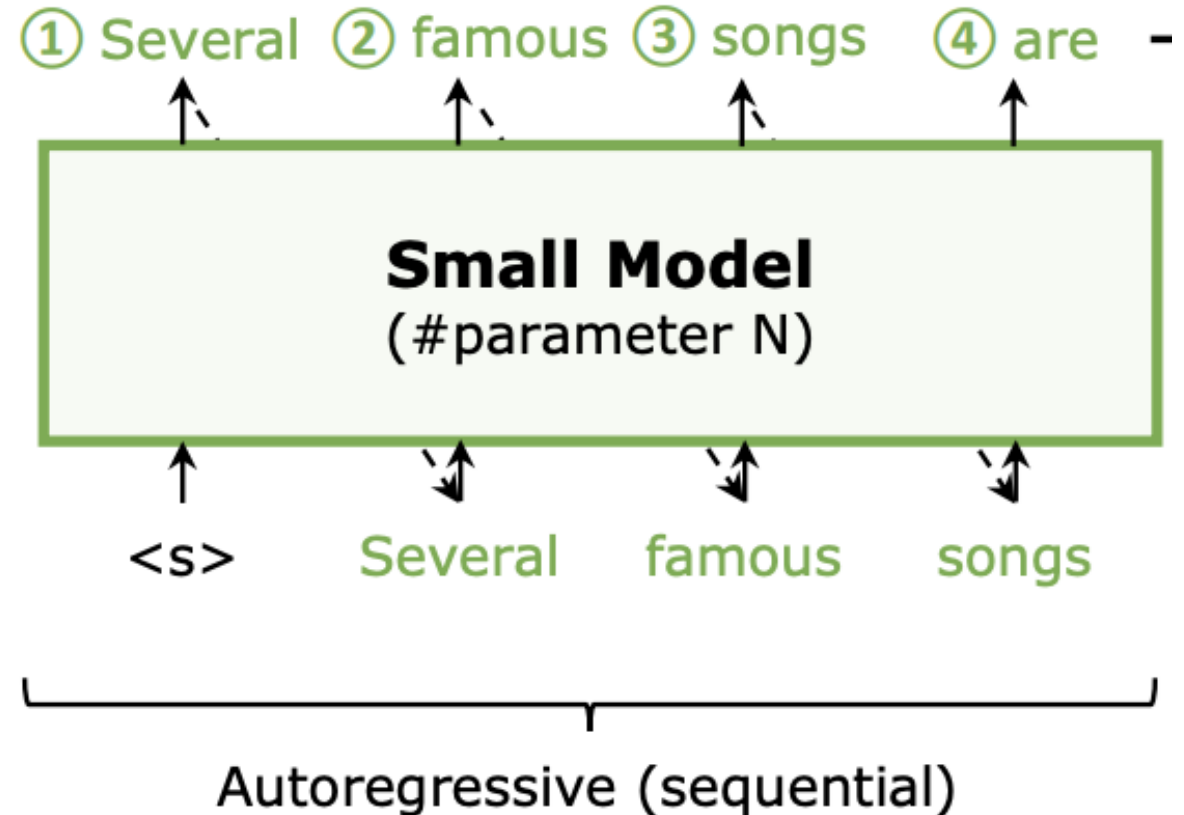


Big Little Decoder [Kim, et. al]



Speculative Decoding

- Speculative decoding resolves this by 'speculating' multiple tokens into the future with a smaller, cheaper model

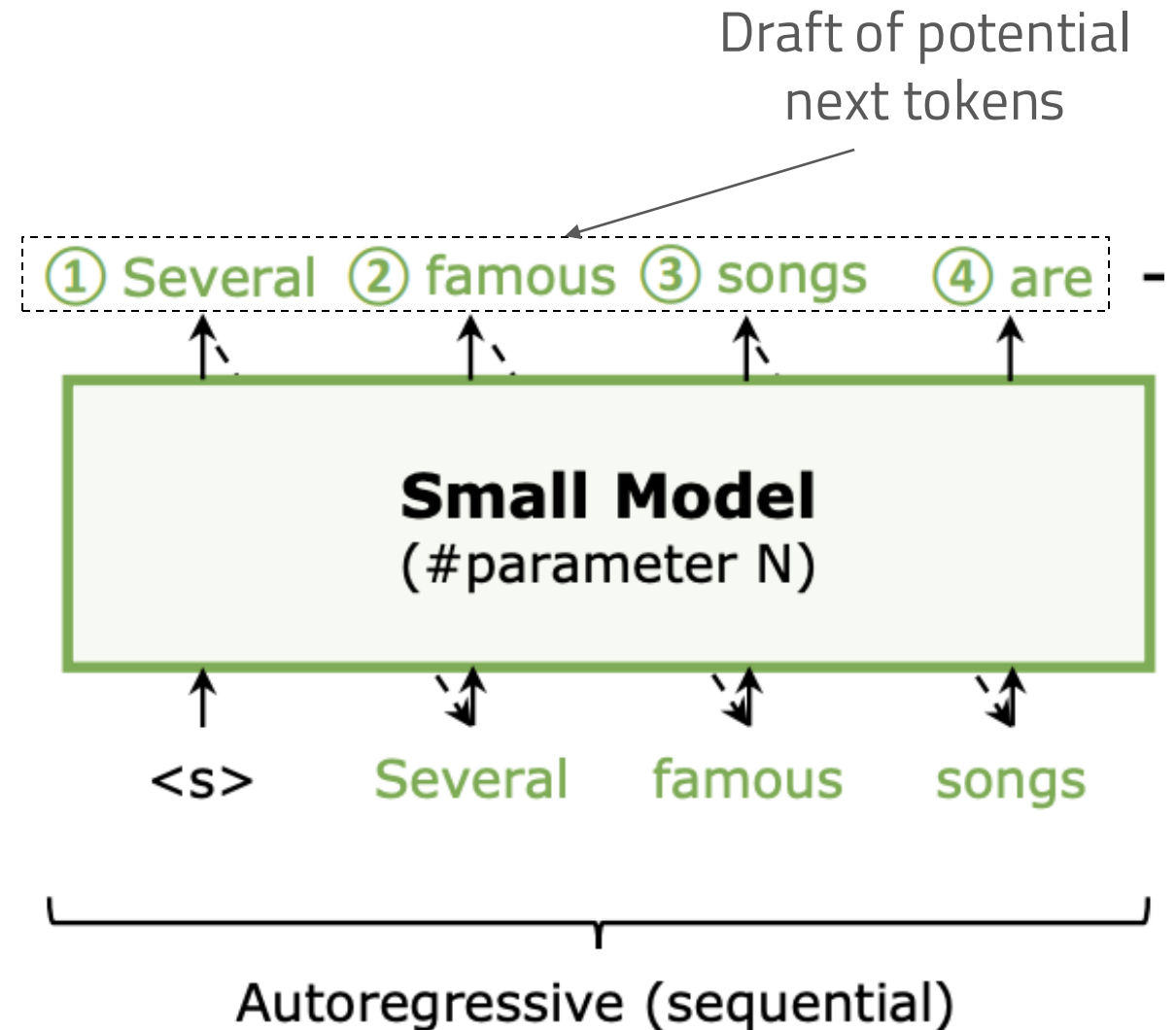


Big Little Decoder [Kim, et. al]



Speculative Decoding

- Speculative decoding resolves this by 'speculating' multiple tokens into the future with a smaller, cheaper model

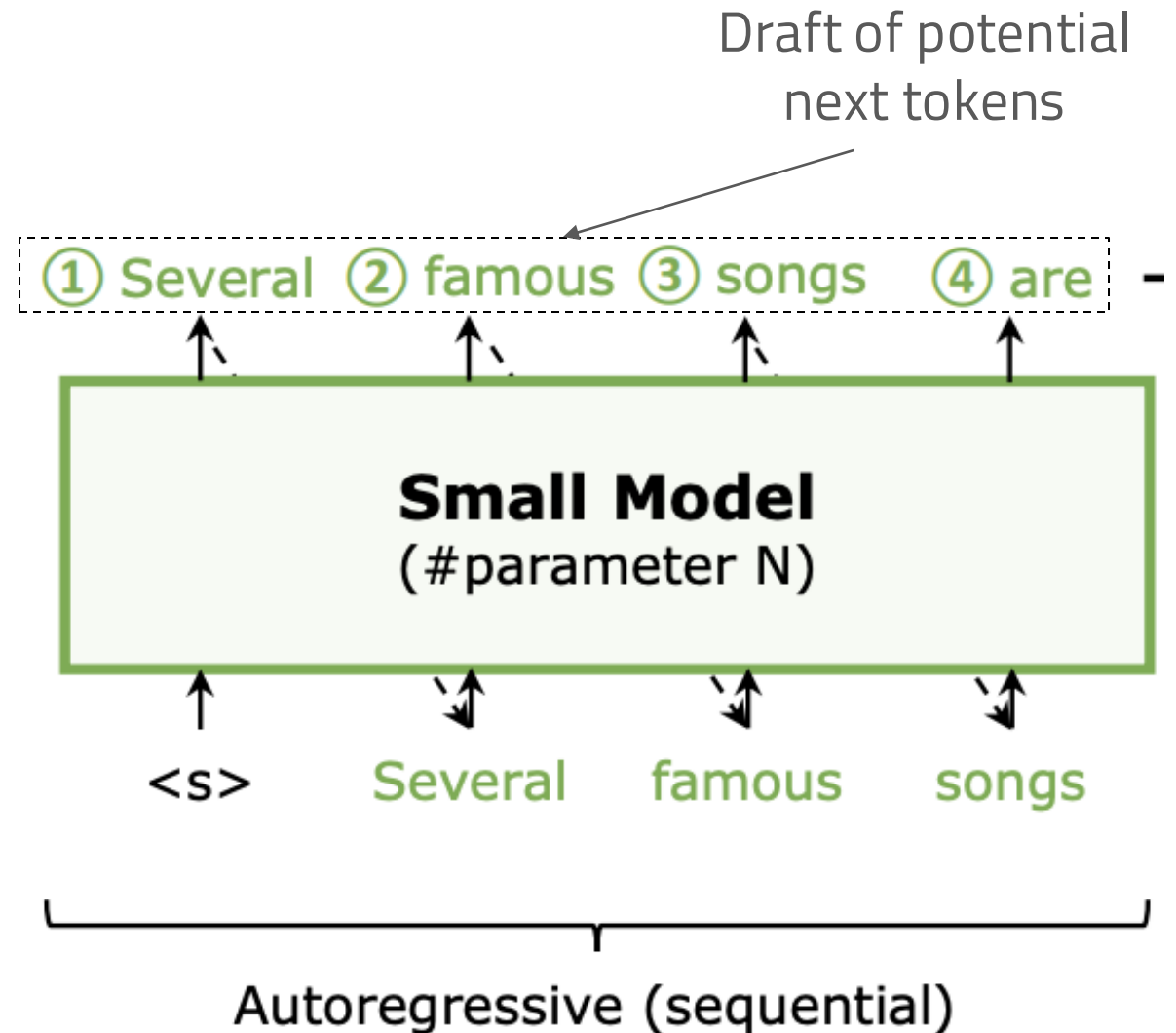


Big Little Decoder [Kim, et. al]



Speculative Decoding

- Speculative decoding resolves this by 'speculating' multiple tokens into the future with a smaller, cheaper model
- We can now send this set of tokens on to a much larger model to verify the sequence

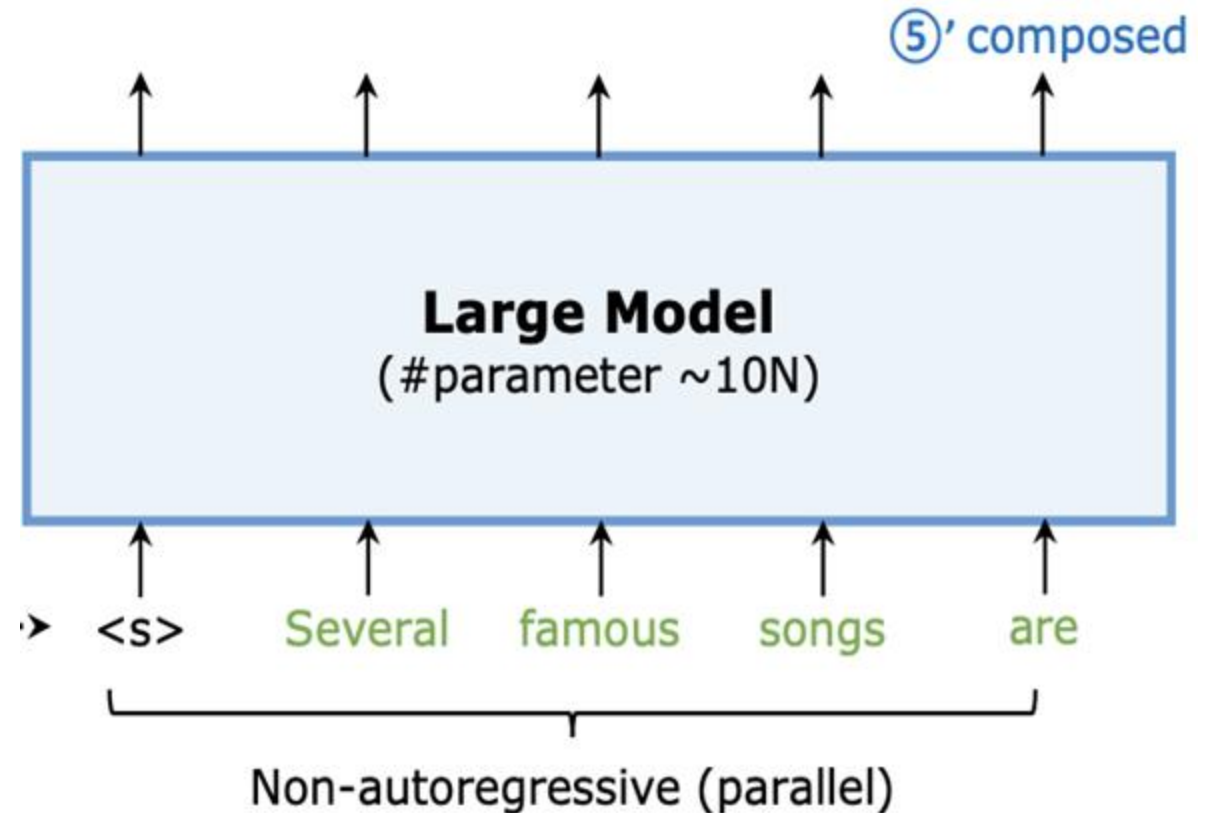


Big Little Decoder [Kim, et. al]



Speculative Decoding

- Because the sequence will be run in parallel the arithmetic intensity will be proportional to the number of draft tokens
- We run each token through and see if the output of the large model matches that of the smaller, draft model

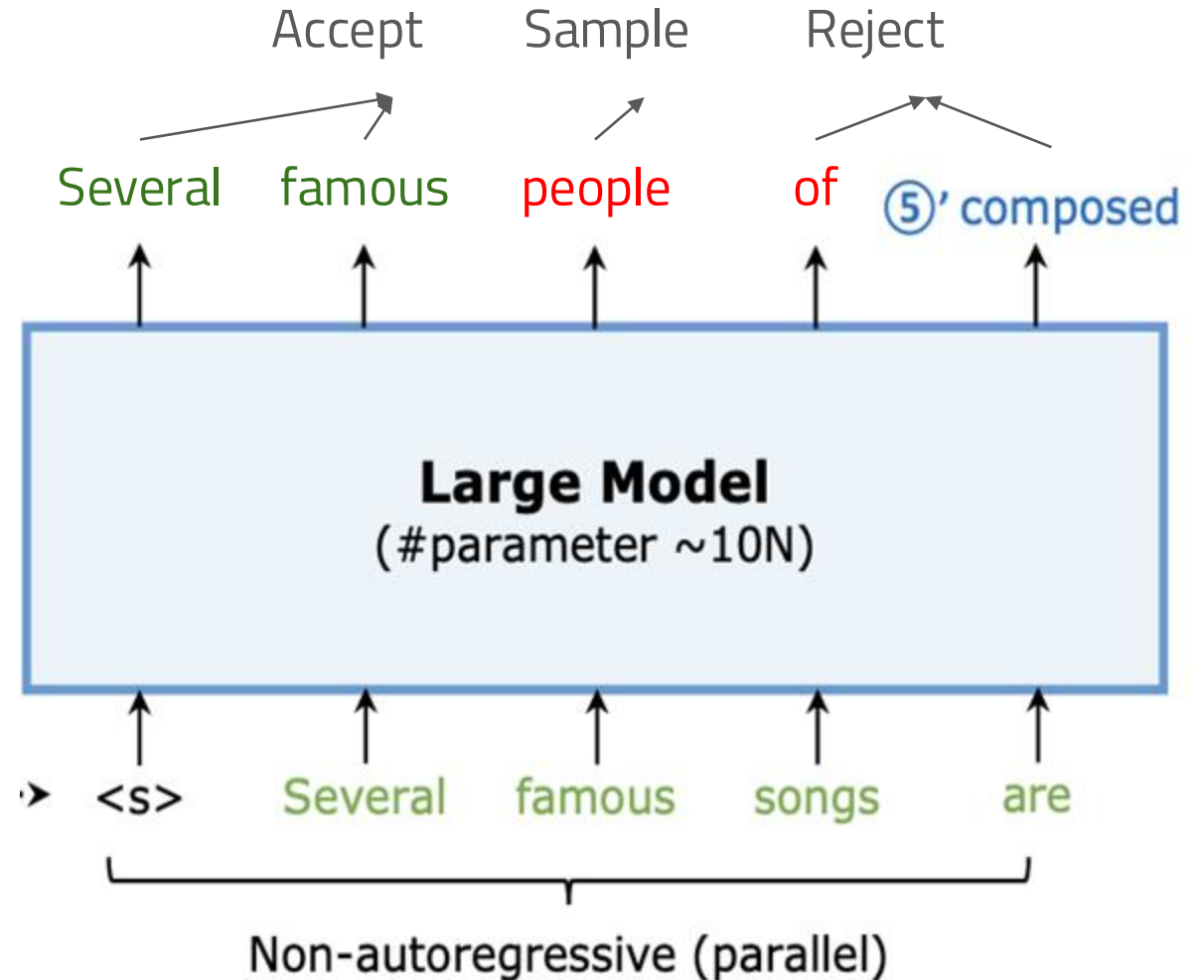


Big Little Decoder [Kim, et. al]



Speculative Decoding

- Because the sequence will be run in parallel the arithmetic intensity will be proportional to the number of draft tokens
- We run each token through and see if the output of the large model matches that of the smaller, draft model
- We accept the matching tokens

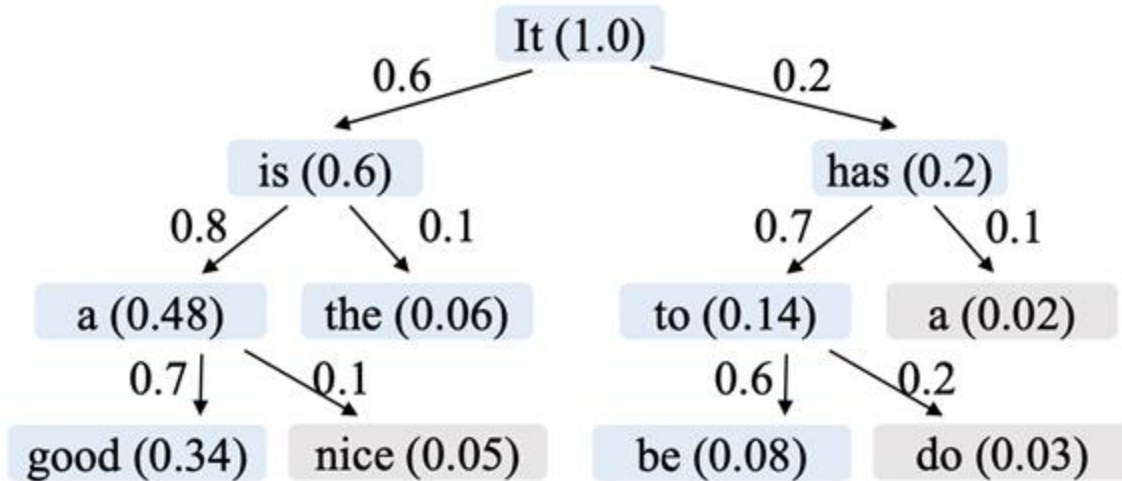


Big Little Decoder [Kim, et. al]



Advanced Approaches

More advanced approaches will use draft trees, rather than draft sequences



		Attention mask							
		It	is	has	a	the	to	good	be
It	✓								
is	✓	✓							
has	✓			✓					
a	✓	✓			✓				
the	✓	✓				✓			
to	✓			✓			✓		
good	✓	✓			✓			✓	
be	✓			✓			✓		✓

Big Little Decoder [Kim, et. al]



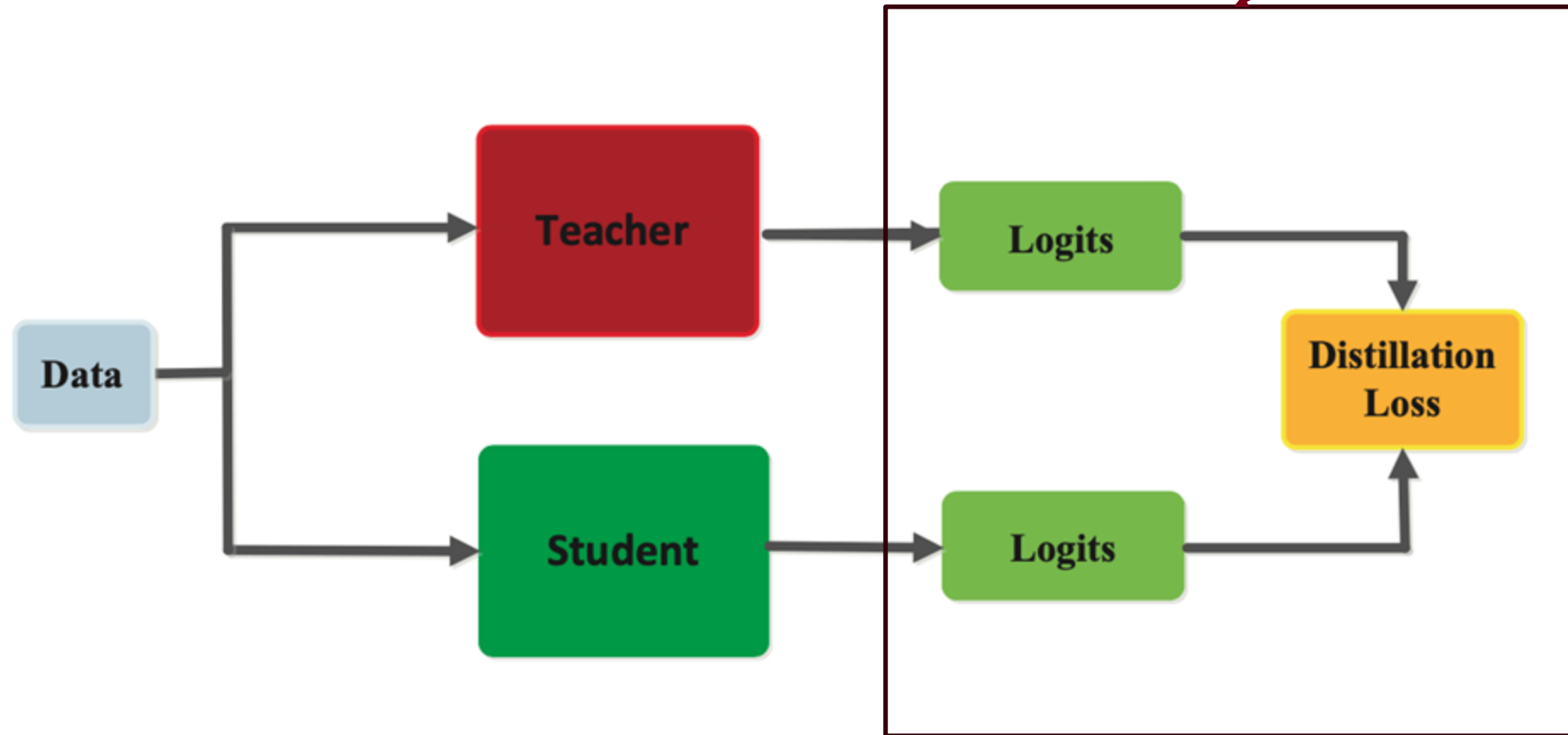
Efficient LLMs

- ❑ Quantization
- ❑ Sparsity
- ❑ Long Context
- ❑ Serving & Systems
- ❑ **Distillation**
- ❑ Parameter Efficient Fine-Tuning
- ❑ Alternative Paradigms



Logit Distillation

Don't use labels, use the modeling distribution of a better 'Teacher' model to train a smaller, 'Student' model



Logit Distillation



Gemma 3



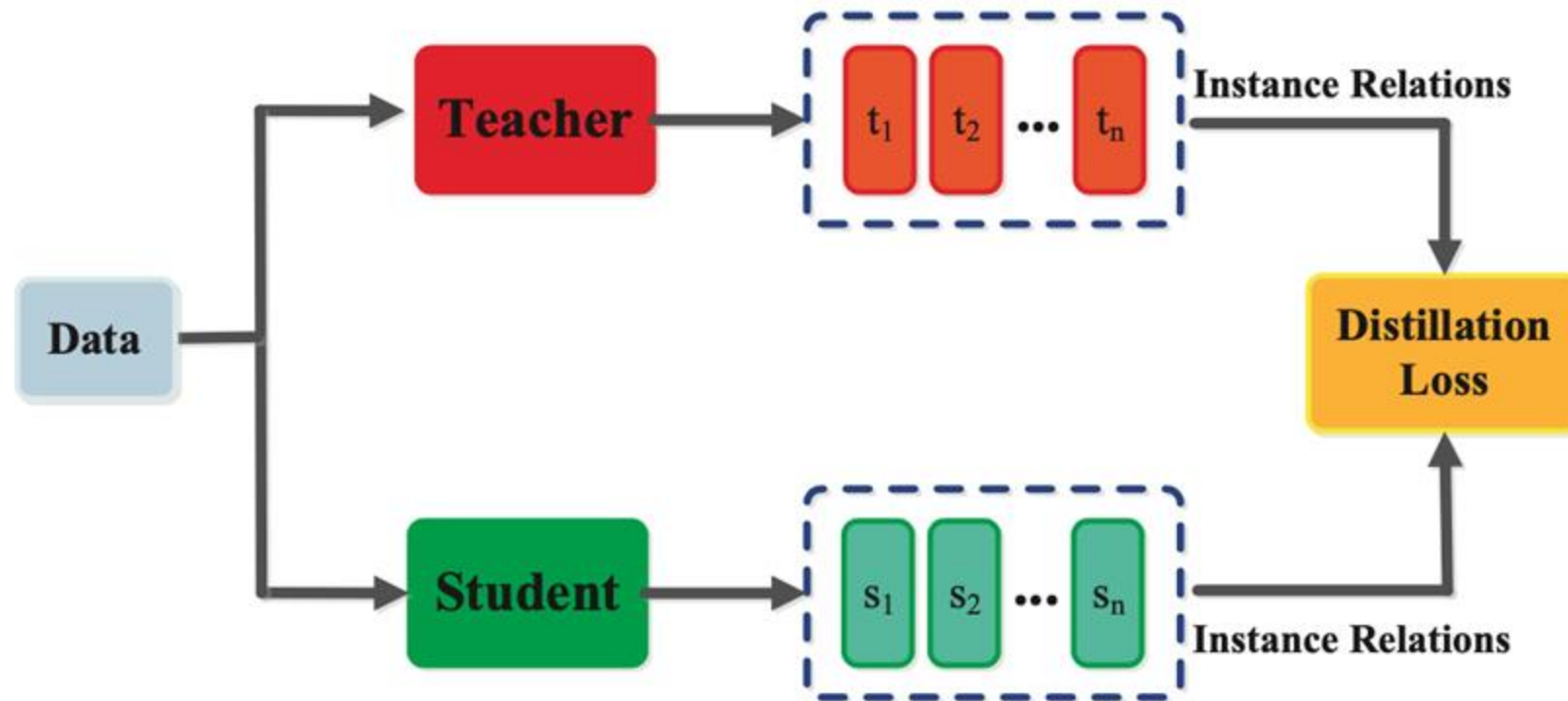
LLAMA 4

Both Gemma3 &
Llama4 use logit
distillation during
pretraining

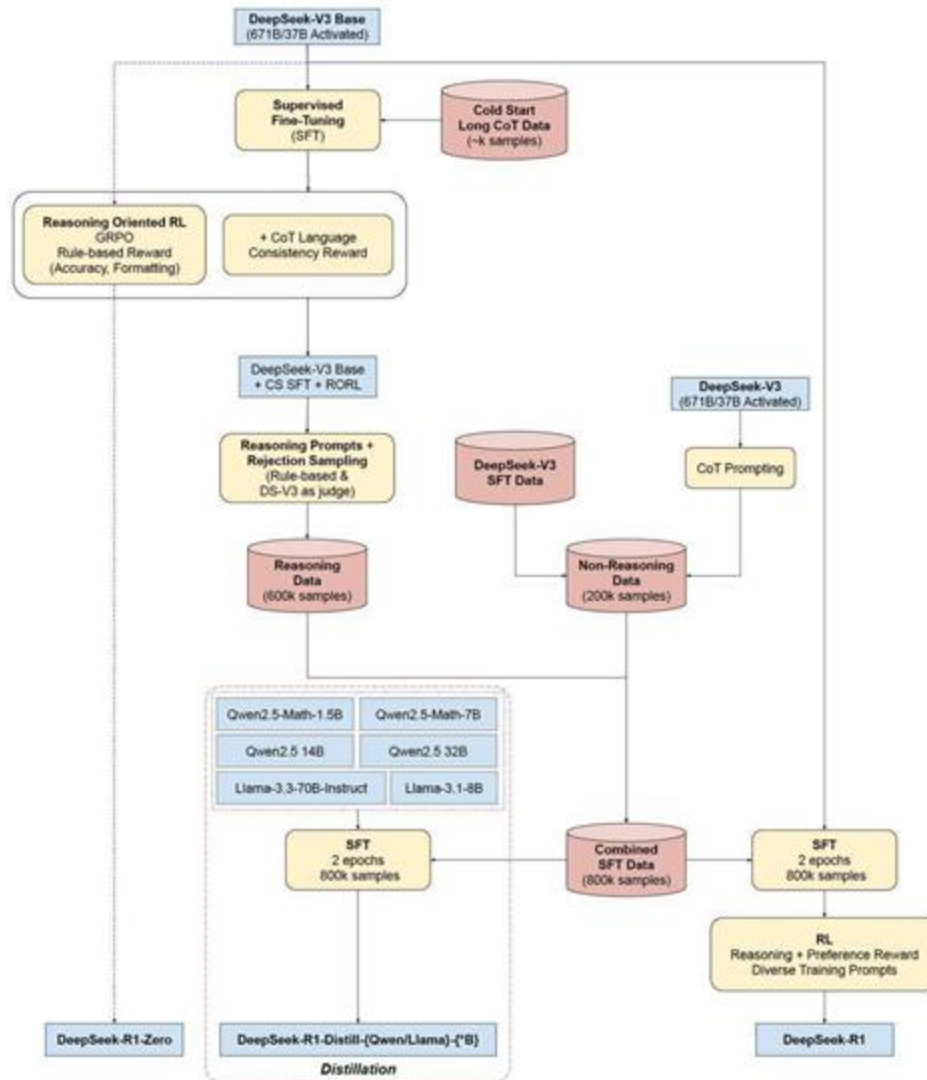


Layer Wise Distillation (LWD)

This can be expanded to trying to match the hidden states of the student and teacher model as well

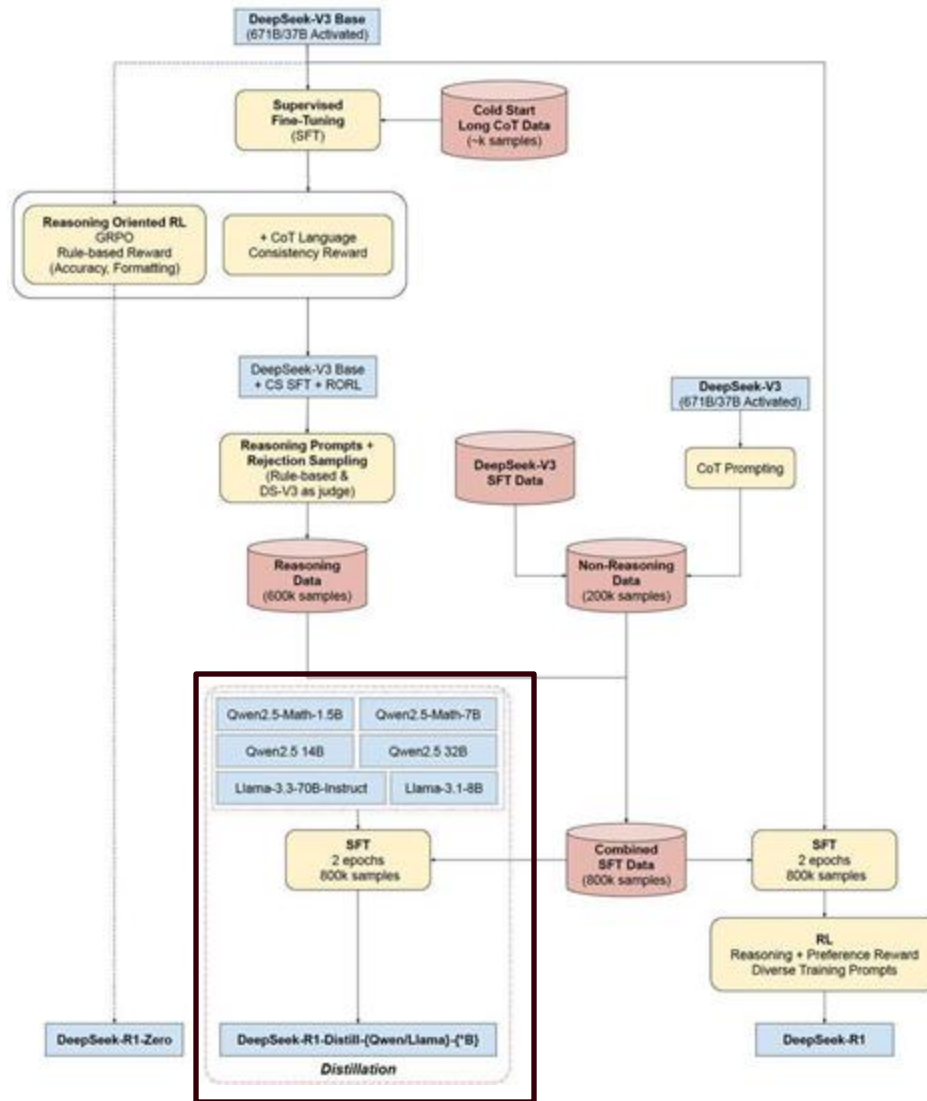


Output Tokens as Distillation (Synthetic Data)



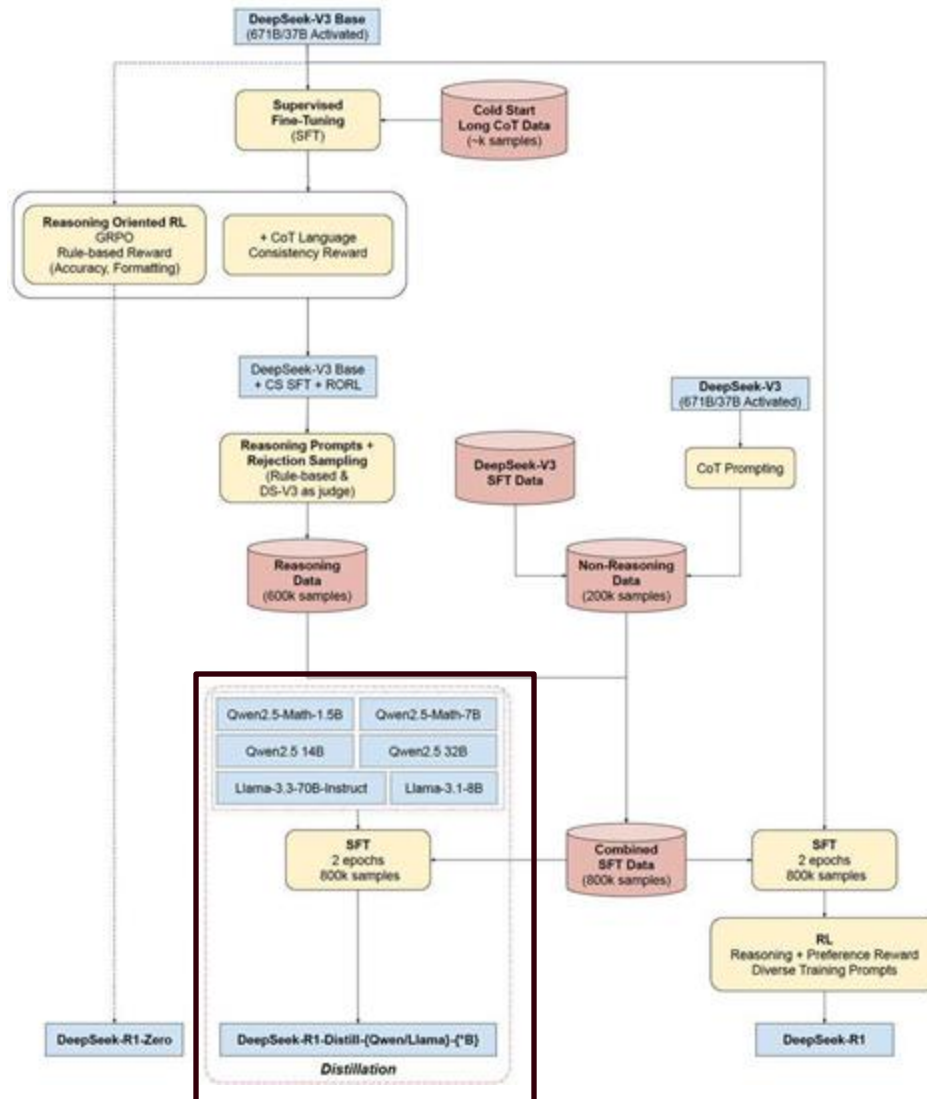
DeepSeekR1 outputs tokens for a given set of problems. In this case 'distillation', just means performing Supervised fine tuning on the tokens create by the larger V3 Models

Output Tokens as Distillation (Synthetic Data)



DeepSeekR1 outputs tokens for a given set of problems. In this case 'distillation', just means performing Supervised fine tuning on the tokens create by the larger V3 Models

Output Tokens as Distillation (Synthetic Data)



DeepSeekR1 outputs tokens for a given set of problems. In this case 'distillation', just means performing Supervised fine tuning on the tokens create by the larger V3 Models

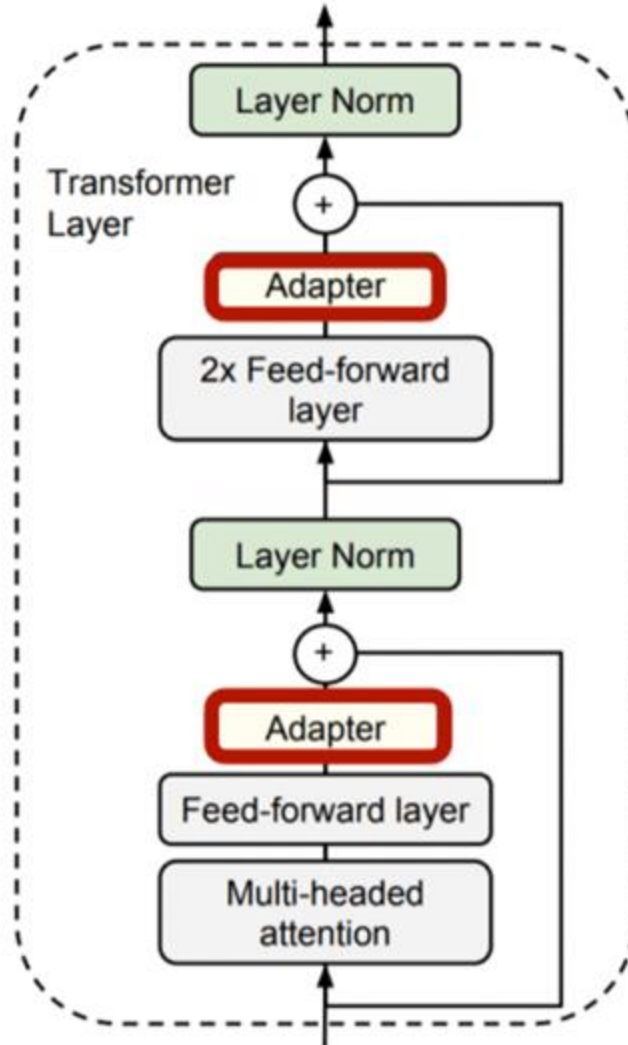
Most small-model training today employs similar methods as these

Efficient LLMs

- ❑ Quantization
- ❑ Sparsity
- ❑ Long Context
- ❑ Serving & Systems
- ❑ Distillation
- ❑ **Parameter Efficient Fine-Tuning**
- ❑ Alternative Paradigms



Adapter



Add trainable
layers after each
feedforward layer

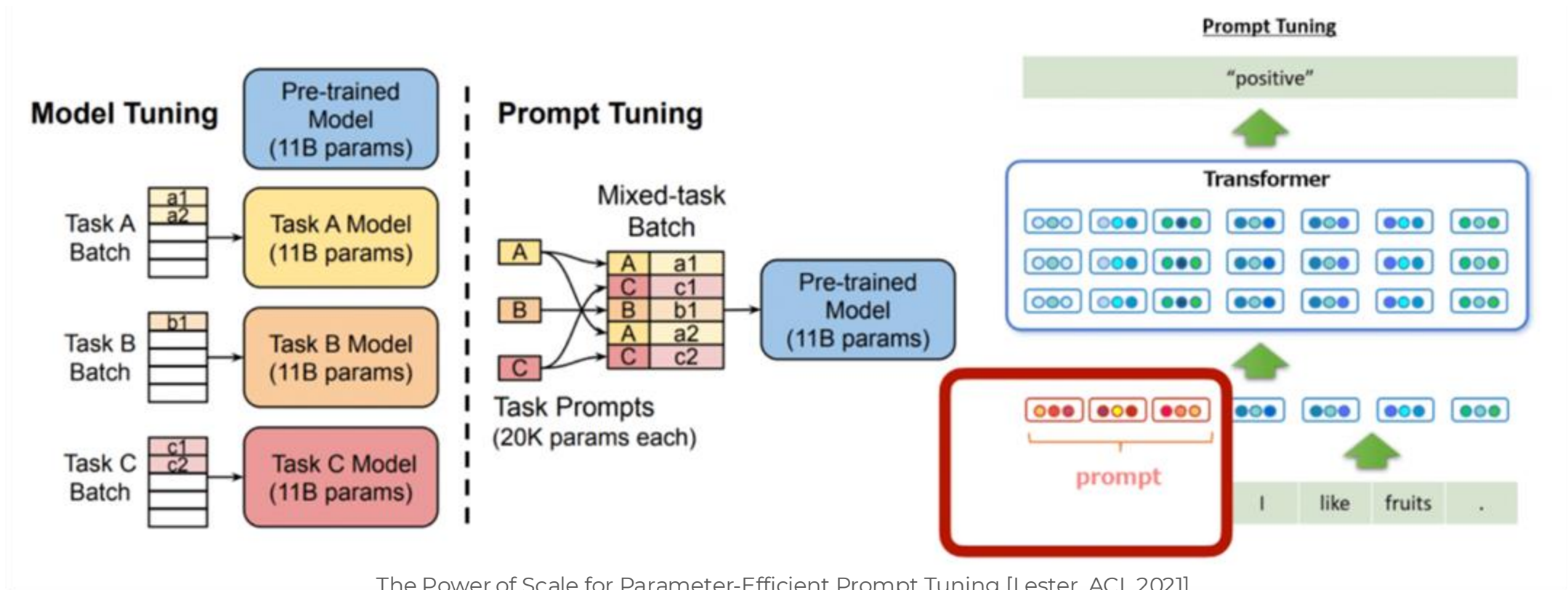
	Total num params	Trained params / task	CoLA	SST	MRPC	STS-B	QQP	MNLI _m	MNLI _{mm}	QNLI	RTE	Total
BERT _{LARGE}	9.0×	100%	60.5	94.9	89.3	87.6	72.1	86.7	85.9	91.1	70.1	80.4
Adapters (8-256)	1.3×	3.6%	59.5	94.0	89.5	86.9	71.8	84.9	85.1	90.7	71.5	80.0
Adapters (64)	1.2×	2.1%	56.9	94.2	89.6	87.3	71.8	85.3	84.6	91.4	68.8	79.6

Parameter-Efficient Transfer Learning for NLP [Houlsby et al, ICML 2019]



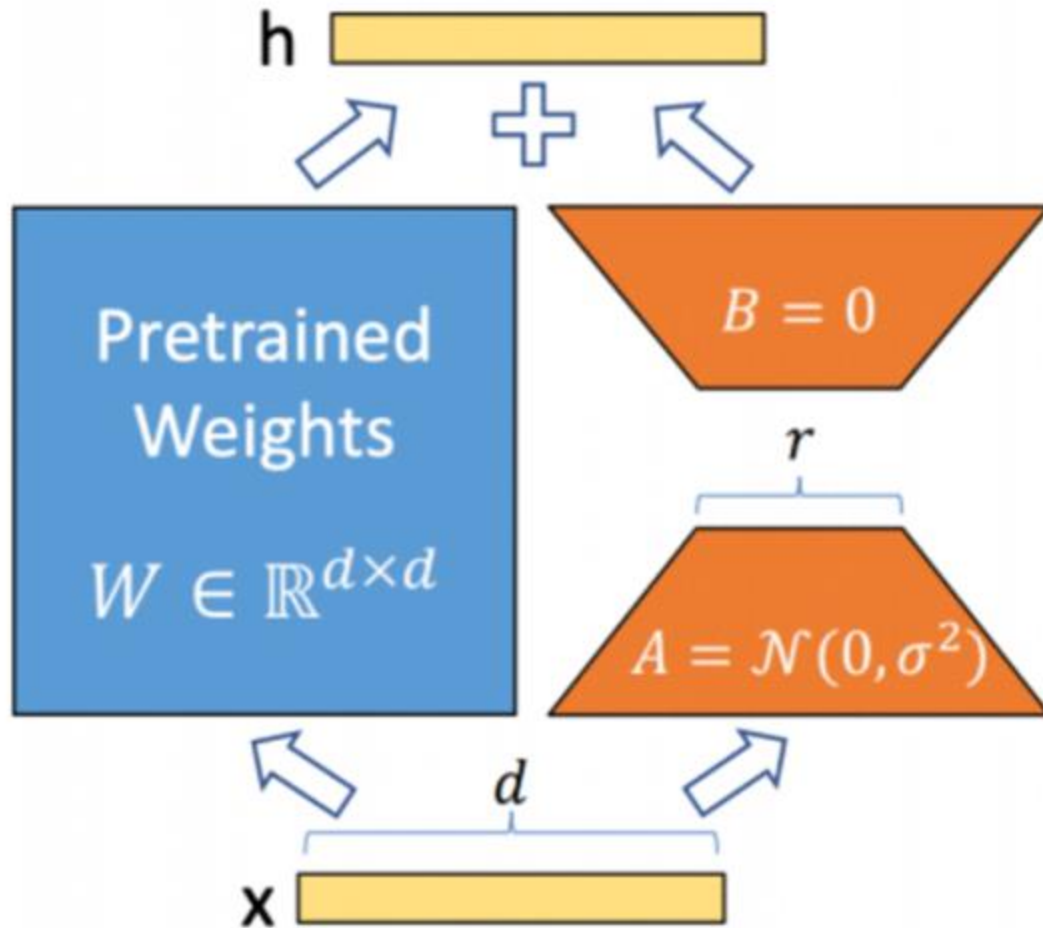
Prompt Tuning (Soft Prompting)

Train a continuous, learnable prompt in embedding space for each task we are training on



The Power of Scale for Parameter-Efficient Prompt Tuning [Lester, ACL 2021]

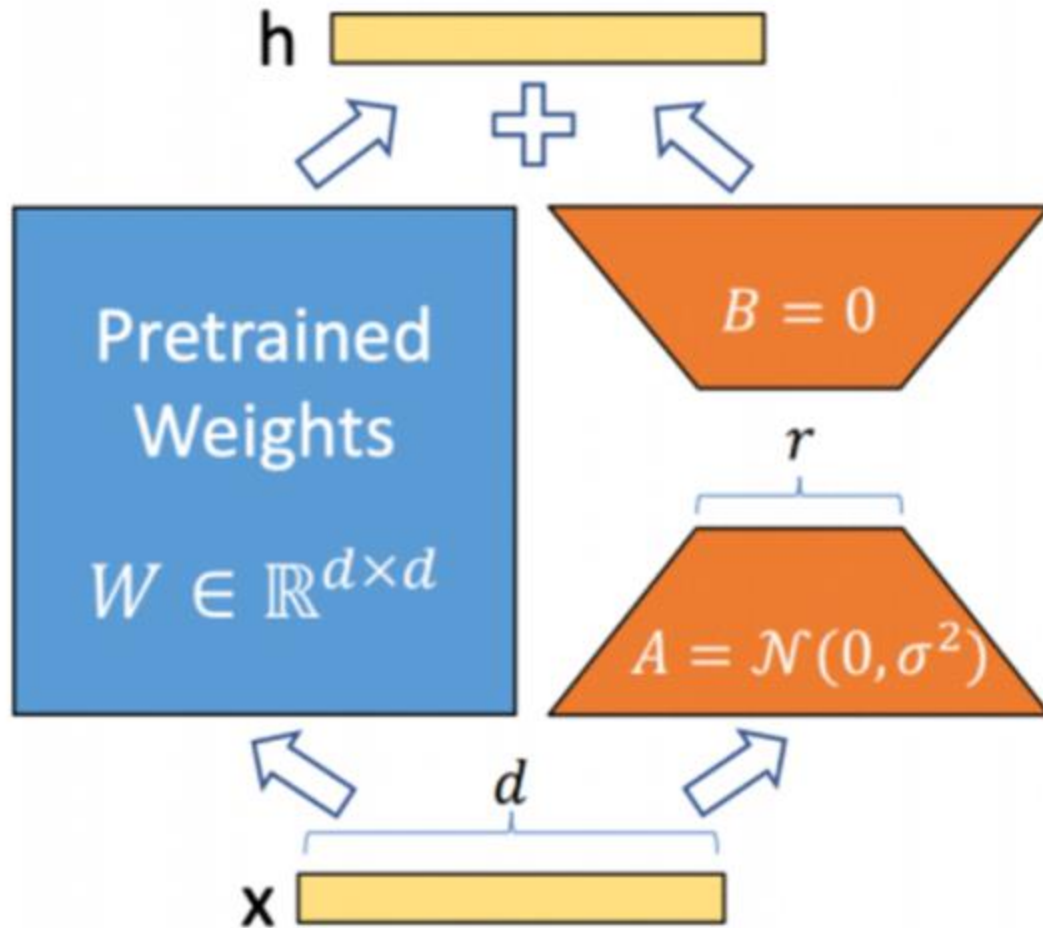
LoRA



The Power of Scale for Parameter-Efficient Prompt Tuning [Lester, ACL 2021]

- Hypothesizes that fine-tuning results in only low rank updates
- Thus, we may approximate the updates themselves as low-rank and train on this low-rank approximation directly

LoRA



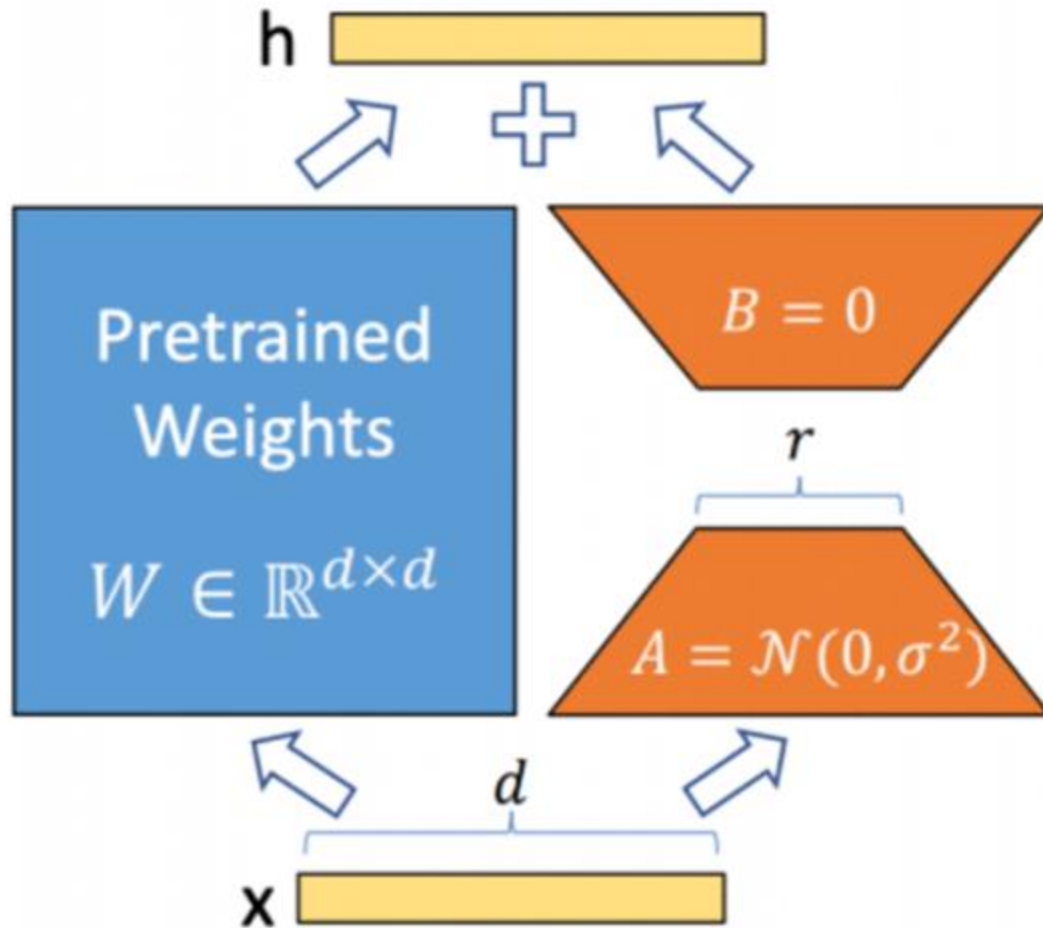
The Power of Scale for Parameter-Efficient Prompt Tuning [Lester, ACL 2021]

Normal Finetuning:

$$h = Wx$$

Update W

LoRA



The Power of Scale for Parameter-Efficient Prompt Tuning [Lester, ACL 2021]

Normal Finetuning:

$$h = Wx$$

Update W

LoRA Finetuning:

$$h = Wx + BAx$$

Update B, A
Leave W unchanged

LoRA

Model & Method	# Trainable Parameters	MNLI	SST-2	MRPC	CoLA	QNLI	QQP	RTE	STS-B	Avg.
RoB _{base} (FT)*	125.0M	87.6	94.8	90.2	63.6	92.8	91.9	78.7	91.2	86.4
RoB _{base} (BitFit)*	0.1M	84.7	93.7	92.7	62.0	91.8	84.0	81.5	90.8	85.2
RoB _{base} (Adpt ^D)*	0.3M	87.1 $\pm$.0	94.2 $\pm$.1	88.5 $\pm$ 1.1	60.8 $\pm$.4	93.1 $\pm$.1	90.2 $\pm$.0	71.5 $\pm$ 2.7	89.7 $\pm$.3	84.4
RoB _{base} (Adpt ^D)*	0.9M	87.3 $\pm$.1	94.7 $\pm$.3	88.4 $\pm$.1	62.6 $\pm$.9	93.0 $\pm$.2	90.6 $\pm$.0	75.9 $\pm$ 2.2	90.3 $\pm$.1	85.4
RoB _{base} (LoRA)	0.3M	87.5 $\pm$.3	95.1$\pm$.2	89.7 $\pm$.7	63.4 $\pm$ 1.2	93.3$\pm$.3	90.8 $\pm$.1	86.6$\pm$.7	91.5$\pm$.2	87.2
RoB _{large} (FT)*	355.0M	90.2	96.4	90.9	68.0	94.7	92.2	86.6	92.4	88.9
RoB _{large} (LoRA)	0.8M	90.6$\pm$.2	96.2 $\pm$.5	90.9$\pm$1.2	68.2$\pm$1.9	94.9$\pm$.3	91.6 $\pm$.1	87.4$\pm$2.5	92.6$\pm$.2	89.0
RoB _{large} (Adpt ^P)†	3.0M	90.2 $\pm$.3	96.1 $\pm$.3	90.2 $\pm$.7	68.3$\pm$1.0	94.8$\pm$.2	91.9$\pm$.1	83.8 $\pm$ 2.9	92.1 $\pm$.7	88.4
RoB _{large} (Adpt ^P)†	0.8M	90.5$\pm$.3	96.6$\pm$.2	89.7 $\pm$ 1.2	67.8 $\pm$ 2.5	94.8$\pm$.3	91.7 $\pm$.2	80.1 $\pm$ 2.9	91.9 $\pm$.4	87.9
RoB _{large} (Adpt ^H)†	6.0M	89.9 $\pm$.5	96.2 $\pm$.3	88.7 $\pm$ 2.9	66.5 $\pm$ 4.4	94.7 $\pm$.2	92.1 $\pm$.1	83.4 $\pm$ 1.1	91.0 $\pm$ 1.7	87.8
RoB _{large} (Adpt ^H)†	0.8M	90.3 $\pm$.3	96.3 $\pm$.5	87.7 $\pm$ 1.7	66.3 $\pm$ 2.0	94.7 $\pm$.2	91.5 $\pm$.1	72.9 $\pm$ 2.9	91.5 $\pm$.5	86.4
RoB _{large} (LoRA)†	0.8M	90.6$\pm$.2	96.2 $\pm$.5	90.2$\pm$1.0	68.2 $\pm$ 1.9	94.8$\pm$.3	91.6 $\pm$.2	85.2$\pm$1.1	92.3$\pm$.5	88.6
DeB _{XXL} (FT)*	1500.0M	91.8	97.2	92.0	72.0	96.0	92.7	93.9	92.9	91.1
DeB _{XXL} (LoRA)	4.7M	91.9$\pm$.2	96.9 $\pm$.2	92.6$\pm$.6	72.4$\pm$1.1	96.0$\pm$.1	92.9$\pm$.1	94.9$\pm$.4	93.0$\pm$.2	91.3

Table 2: RoBERTa_{base}, RoBERTa_{large}, and DeBERTa_{XXL} with different adaptation methods on the GLUE benchmark. We report the overall (matched and mismatched) accuracy for MNLI, Matthew’s correlation for CoLA, Pearson correlation for STS-B, and accuracy for other tasks. Higher is better for all metrics. * indicates numbers published in prior works. † indicates runs configured in a setup similar to Houlsby et al. (2019) for a fair comparison.

The Power of Scale for Parameter-Efficient Prompt Tuning [Lester, ACL 2021]



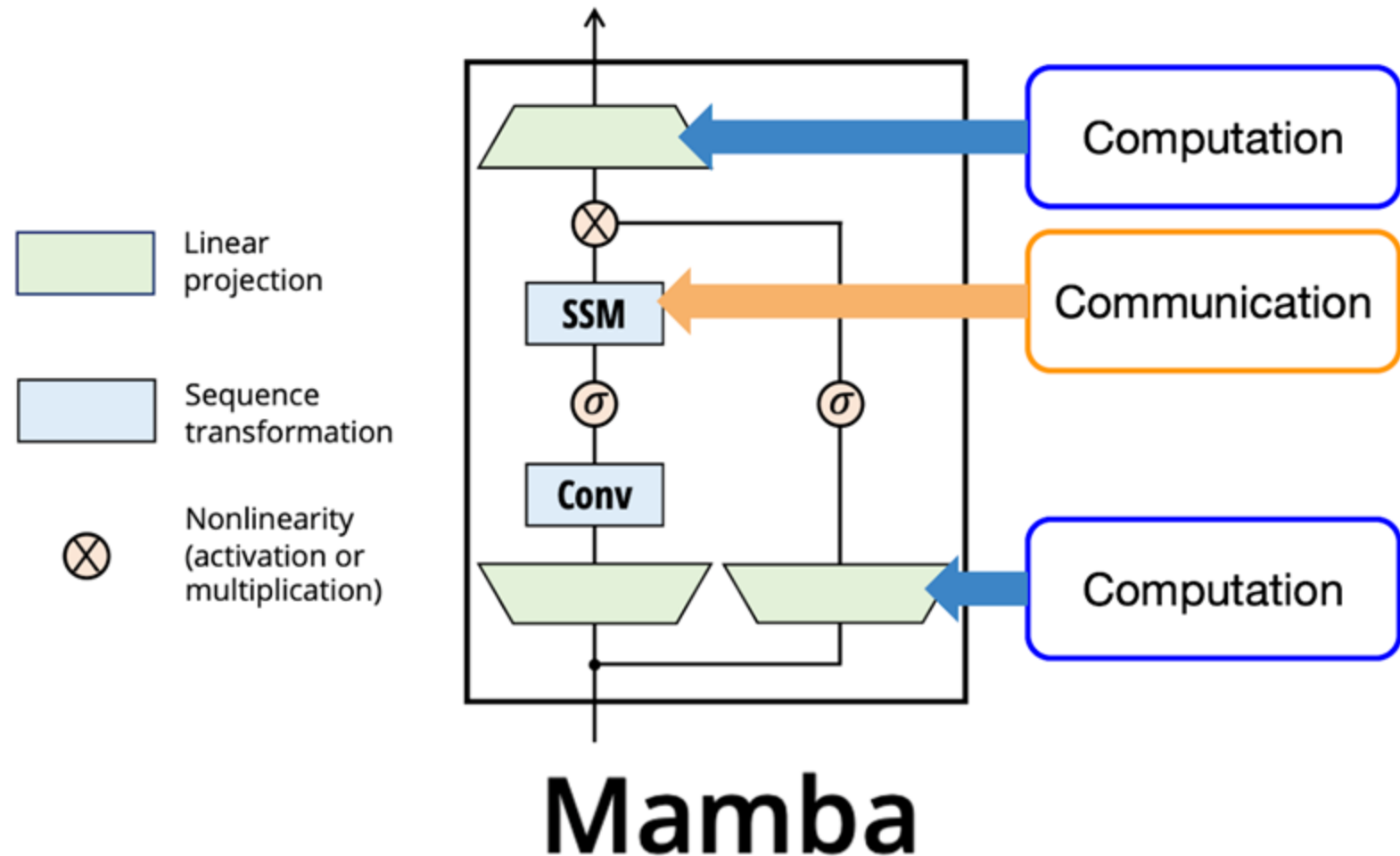
Efficient LLMs

- ❑ Quantization
- ❑ Sparsity
- ❑ Long Context
- ❑ Serving & Systems
- ❑ Distillation
- ❑ Parameter Efficient Fine-Tuning
- ❑ **Alternative Paradigms**



State Space Model (SSM)

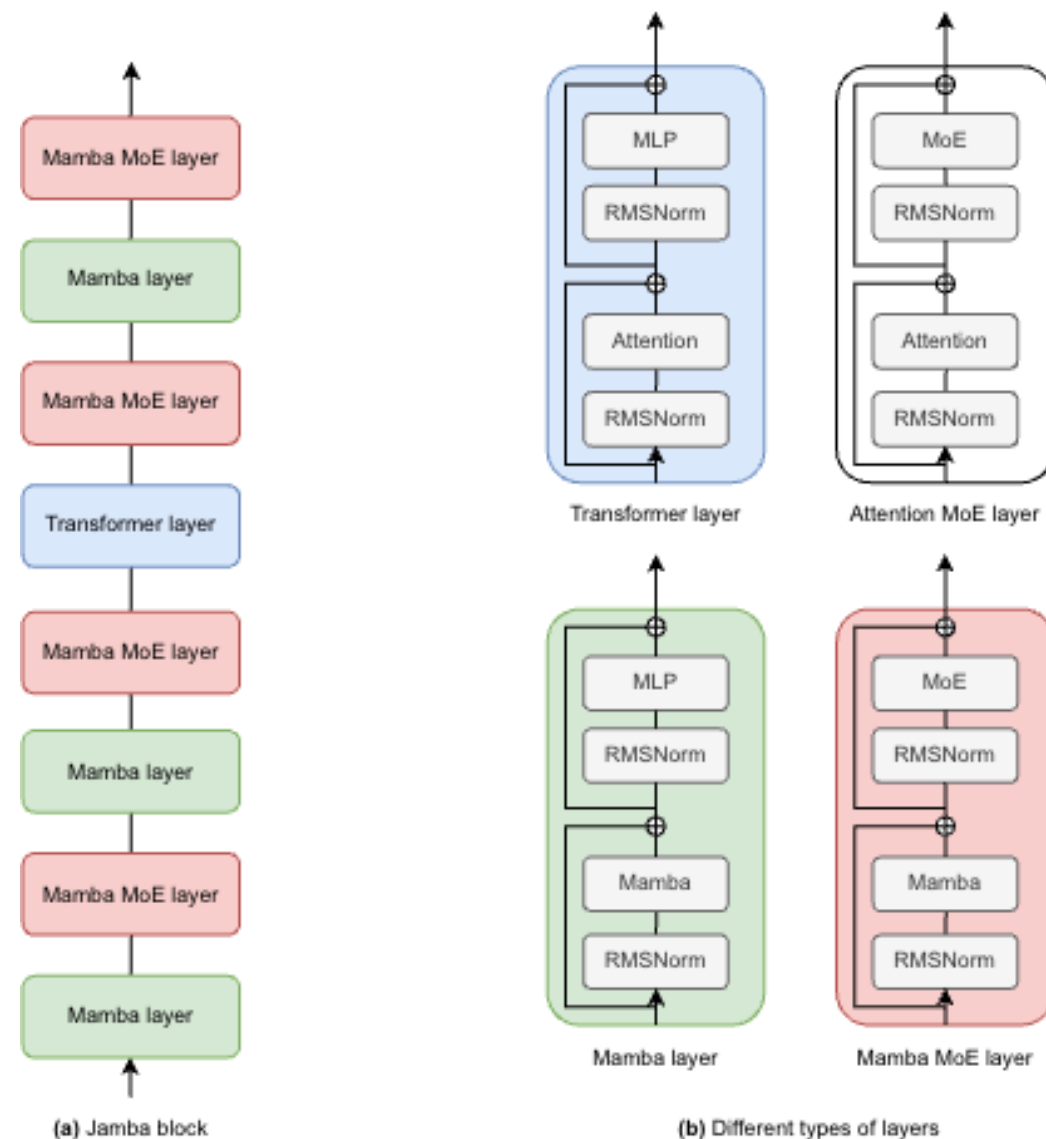
Linearity Strikes back



Mamba: Linear-Time Sequence Modeling with Selective State Spaces (Gu et. Al)

SSM

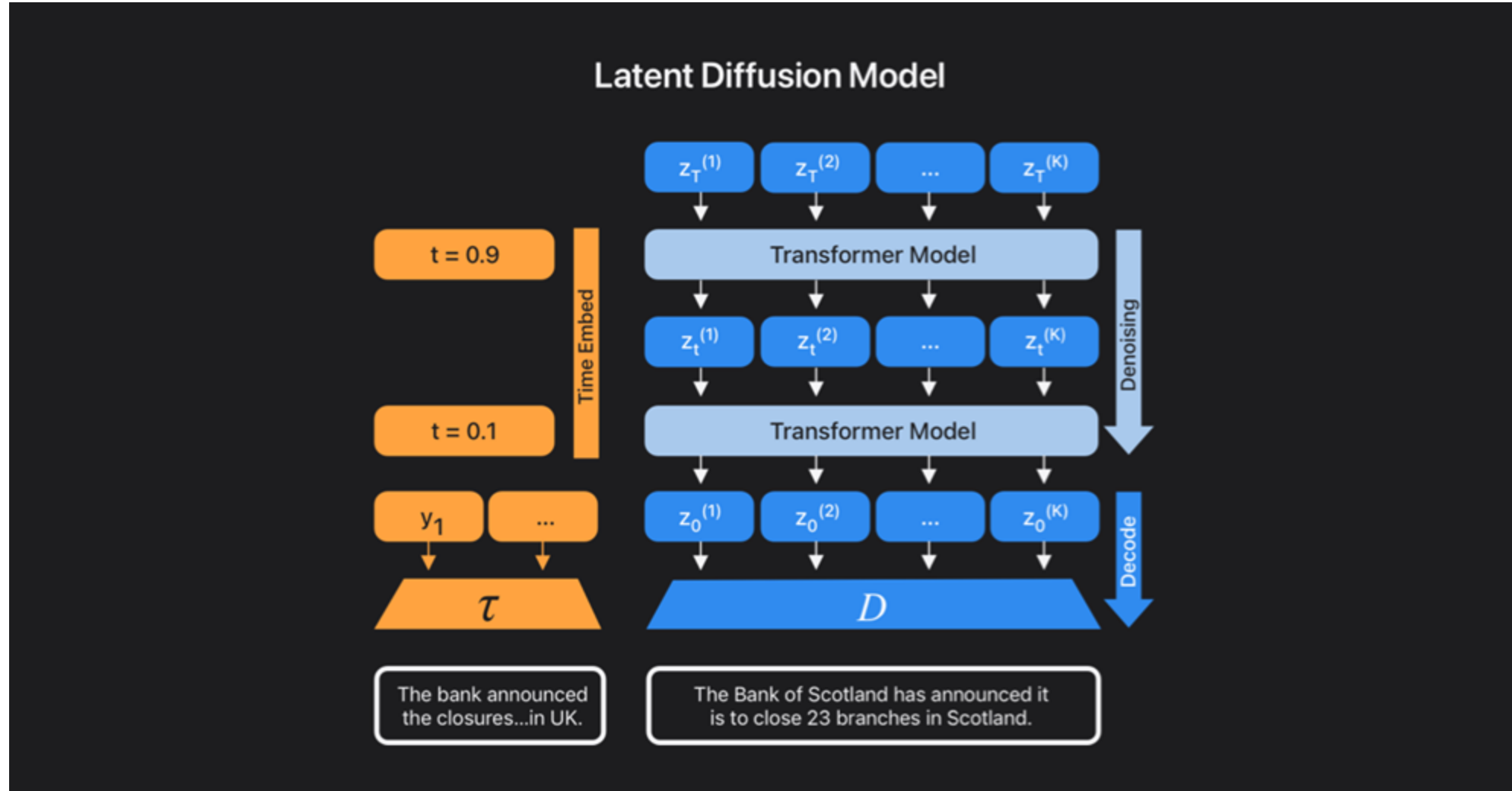
- SSMs are growing in popularity
- Linear complexity makes them excellent candidates for very long tasks
- If they can begin outperforming transformers in practical settings, these could displace the transformer architecture



Jamba: A Hybrid Transformer-Mamba Language Model

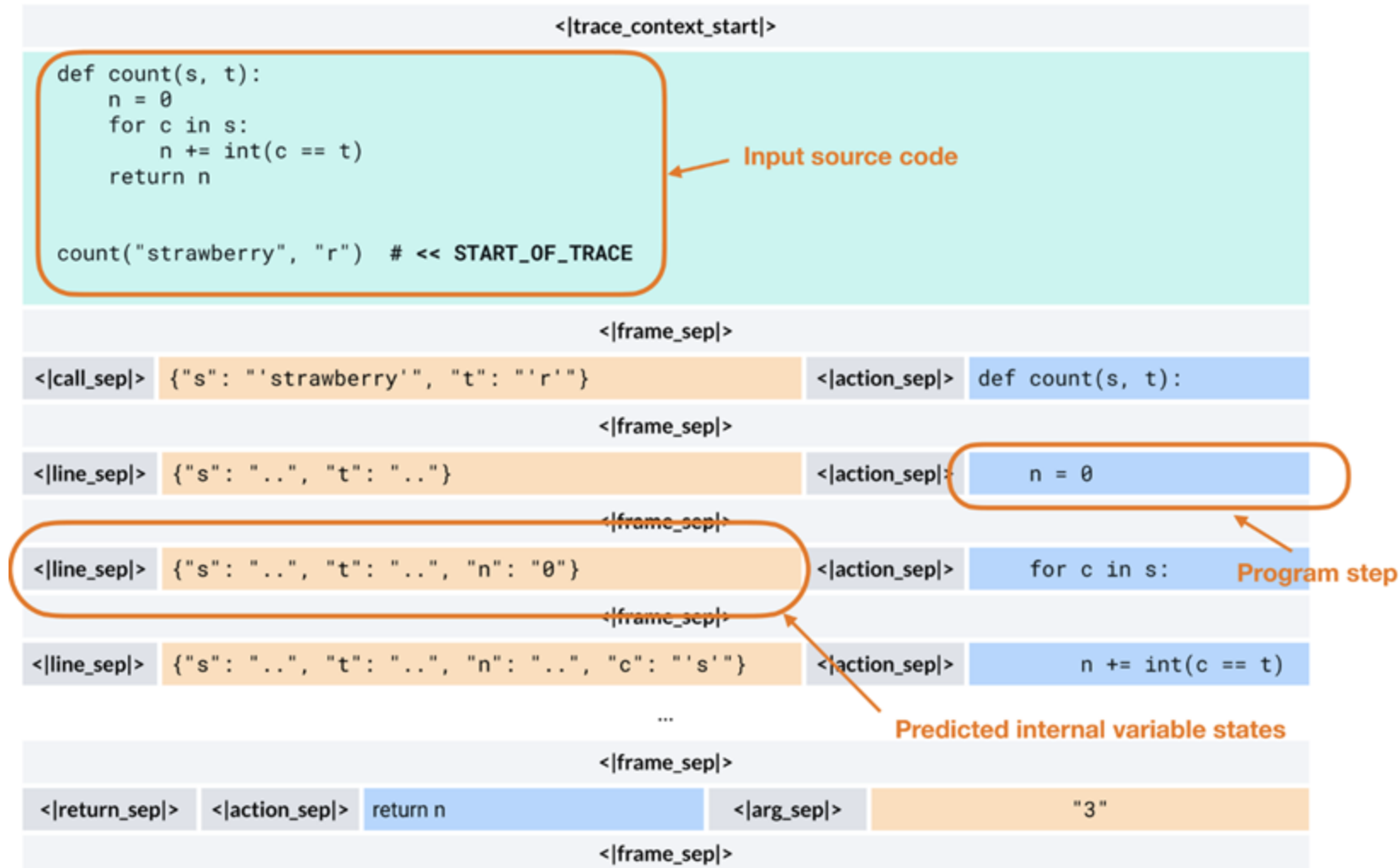
Text Diffusion

Predict multiple tokens at the same time using a diffusion model



World Models

Use some dense environment signals to predict what the consequences of each token prediction is



Recursive Transformers

- Use transformers recursively to iteratively refine model outputs
- This approach is still mostly useful in toy settings at the moment

