

Efficient LLMs

Yatin Nandwani

Research Scientist, IBM Research



Semester 1,
2025-2026

Large Language Models: Introduction and Recent Advances

ELL881 · AIL821

IBM Research, India

Conversational-AI



**Yatin
Nandwani**



**Sonam
Mishra**



**Dinesh
Khandelwal**



**Gaurav
Pandey**



**Vineet
Kumar**



**Meghanadh
Pulivarthi**



**Kushagra
Bhushan**

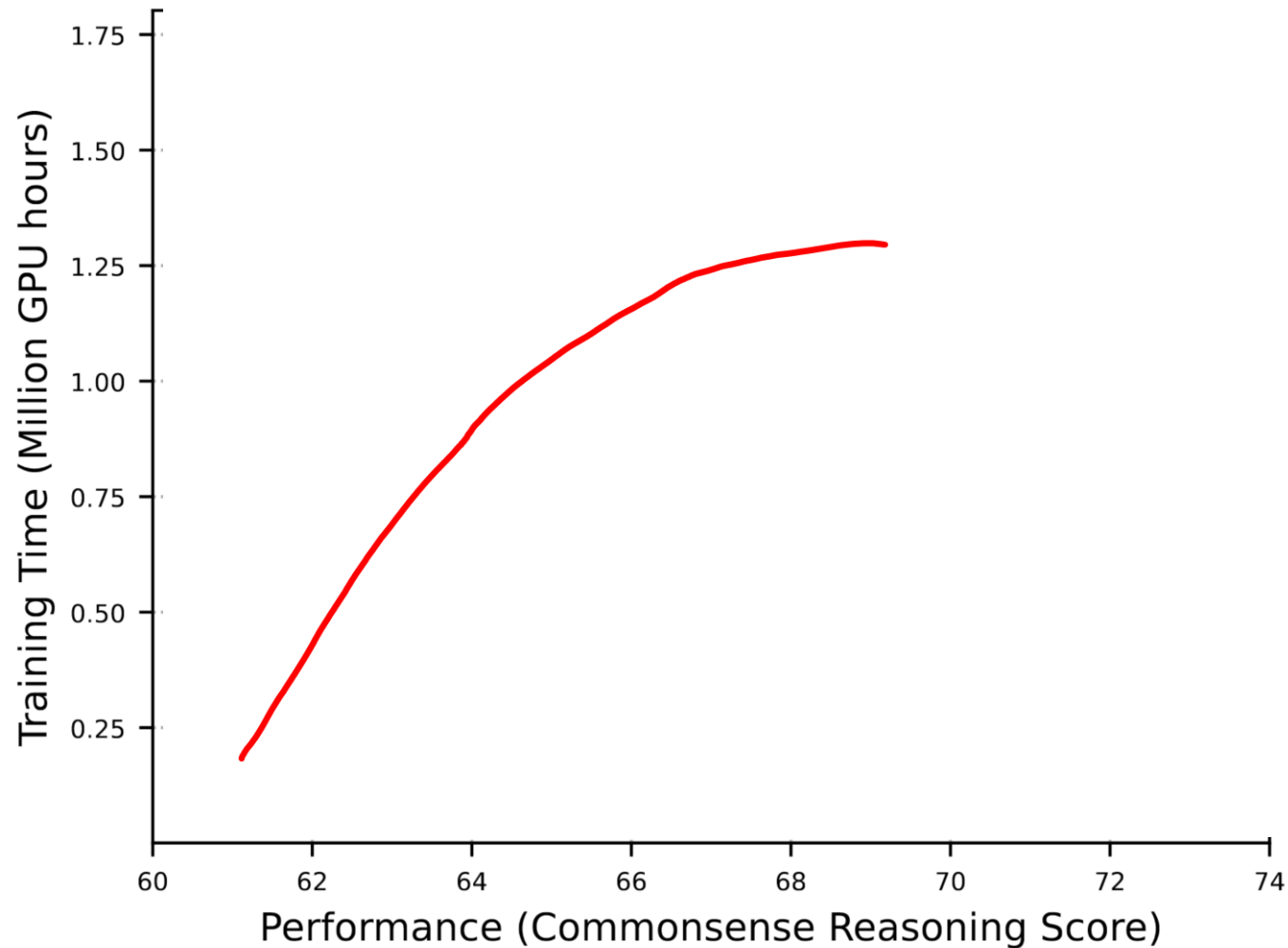


**Dinesh
Raghu**

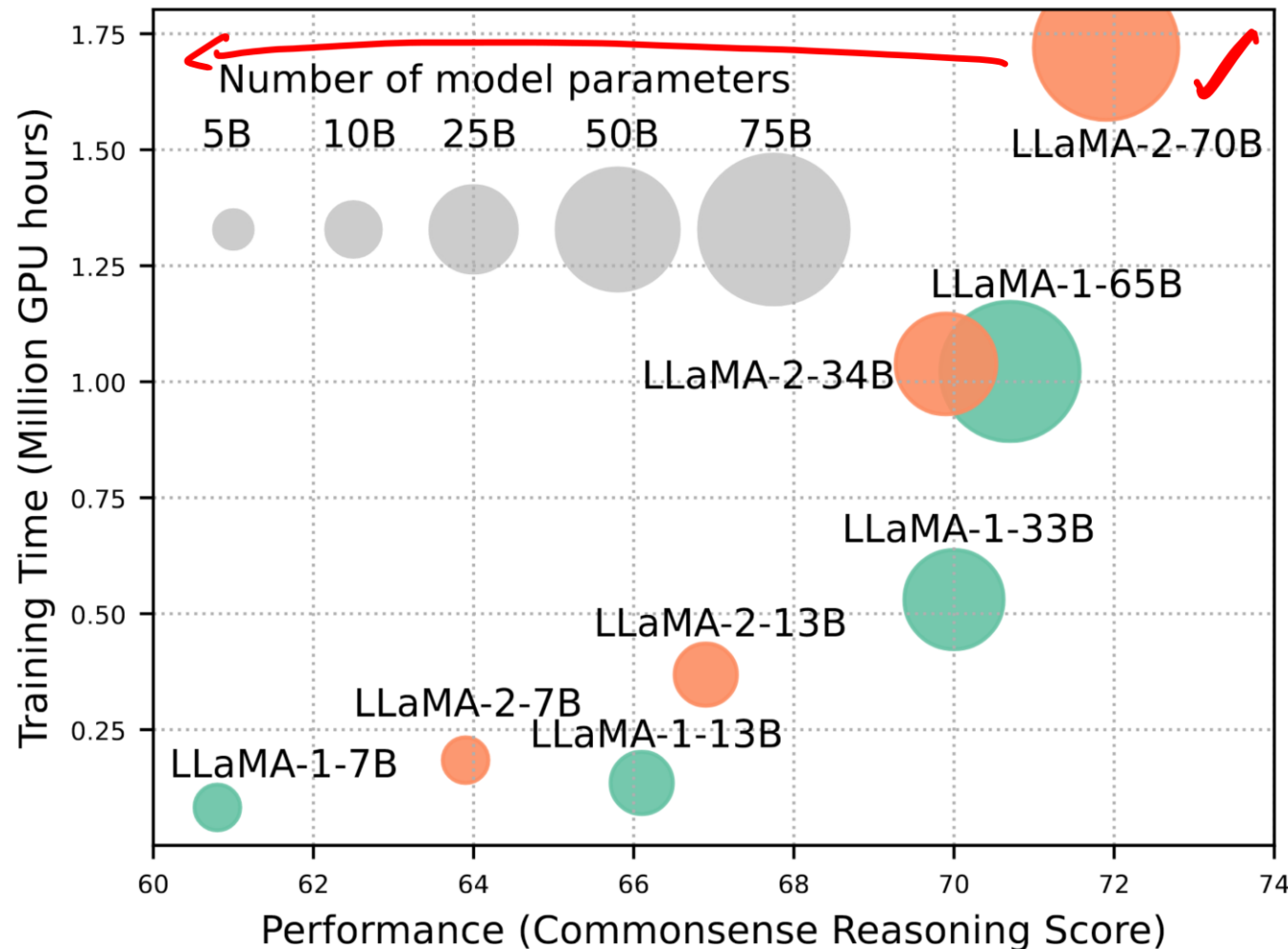


**Sachindra
Joshi**

Training Resources vs Performance



Training Resources vs Performance



- Based on Nvidia A100 80GB GPU

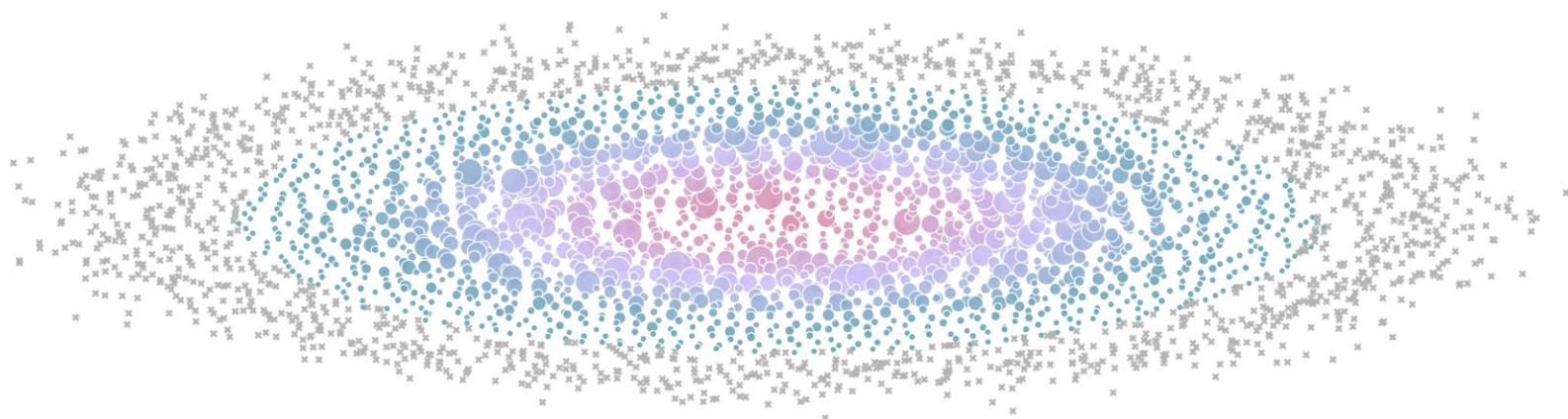
<https://huggingface.co/spaces/optimum/llm-perf-leaderboard>



Efficient LLMs

- How to scale training?
 - *Data Parallelism*
 - *Tensor Parallelism*
 - *Context Parallelism*
 - *Pipeline Parallelism*

The Ultra-Scale Playbook: Training LLMs on GPU Clusters

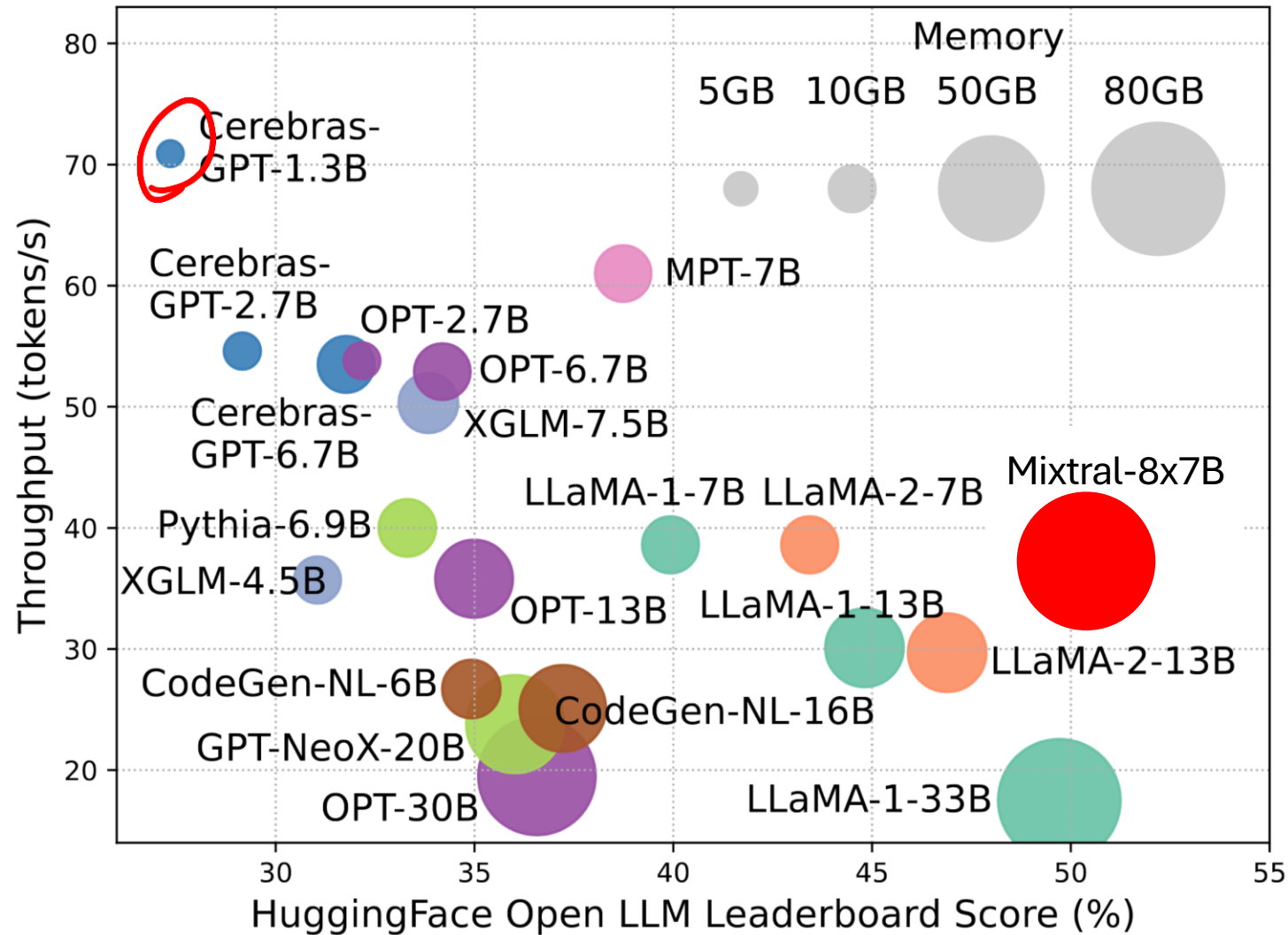


We ran over 4,000 scaling experiments on up to 512 GPUs and measured throughput (size of markers) and GPU utilization (color of markers). Note that both are normalized per model size in this visualization.

Inference Throughput vs Performance

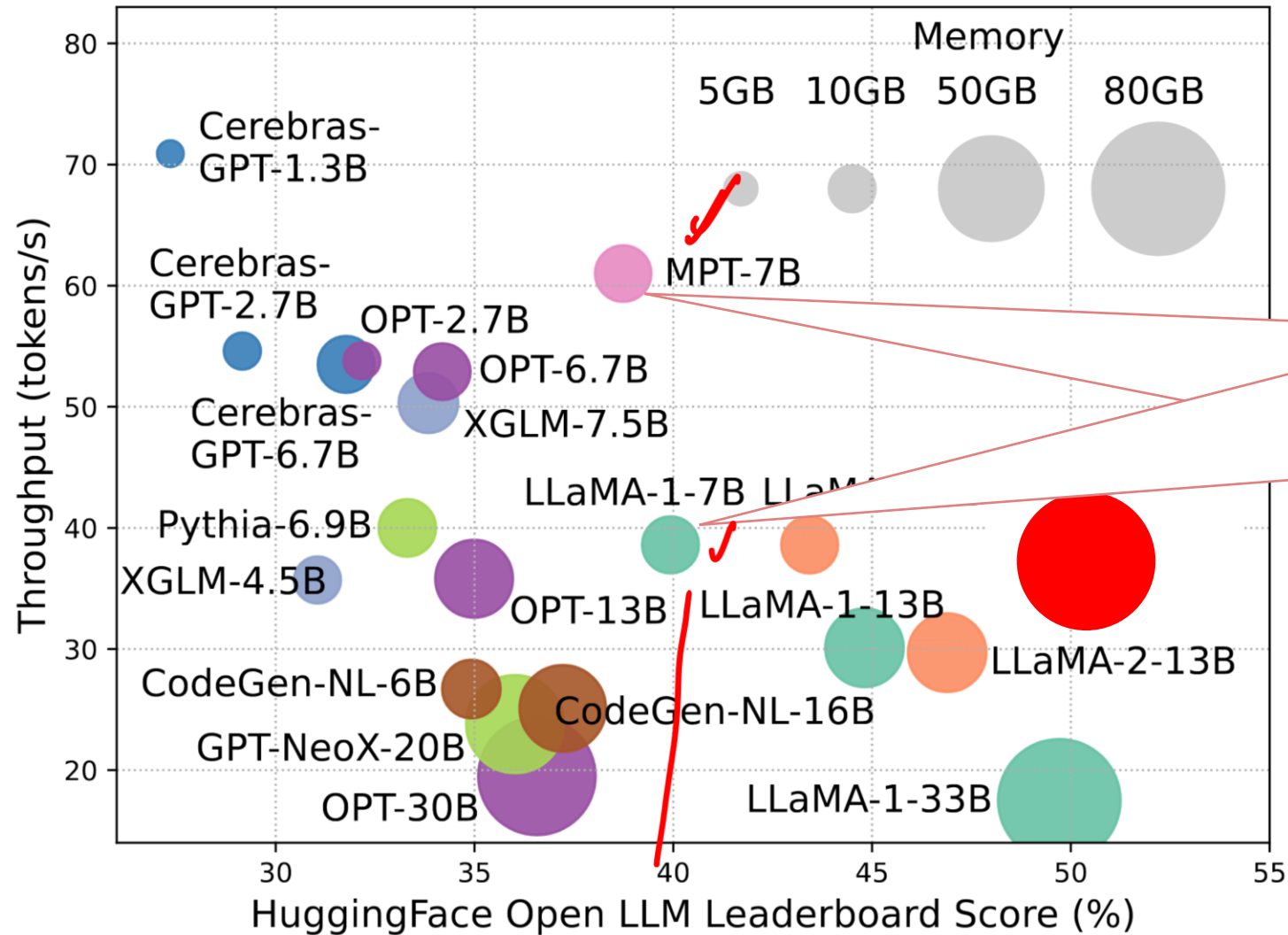


Inference Throughput vs Performance



- On Nvidia A100 80GB GPU;
- 16-bit quantized
- Batch Size - 1
- Prompt size of 256
- Generating 1000 tokens

Inference Throughput vs Performance



- Similar performance, different throughput! How?
- Efficient implementation –
 - Fused kernel for attention

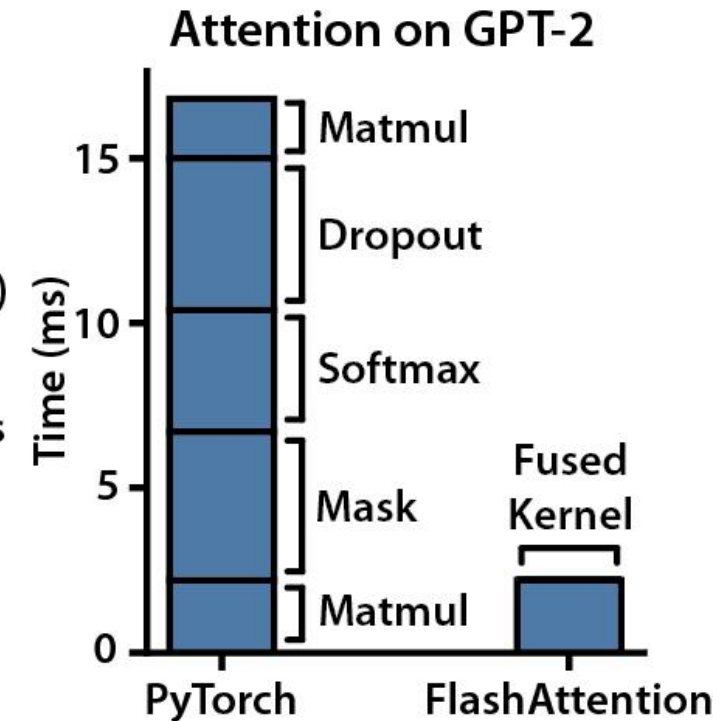
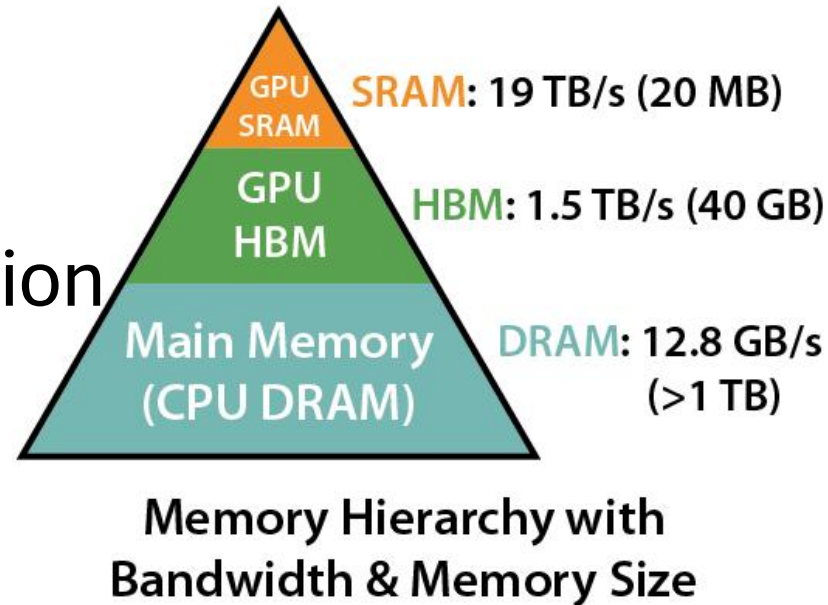
Efficient LLMs

- How to scale training?

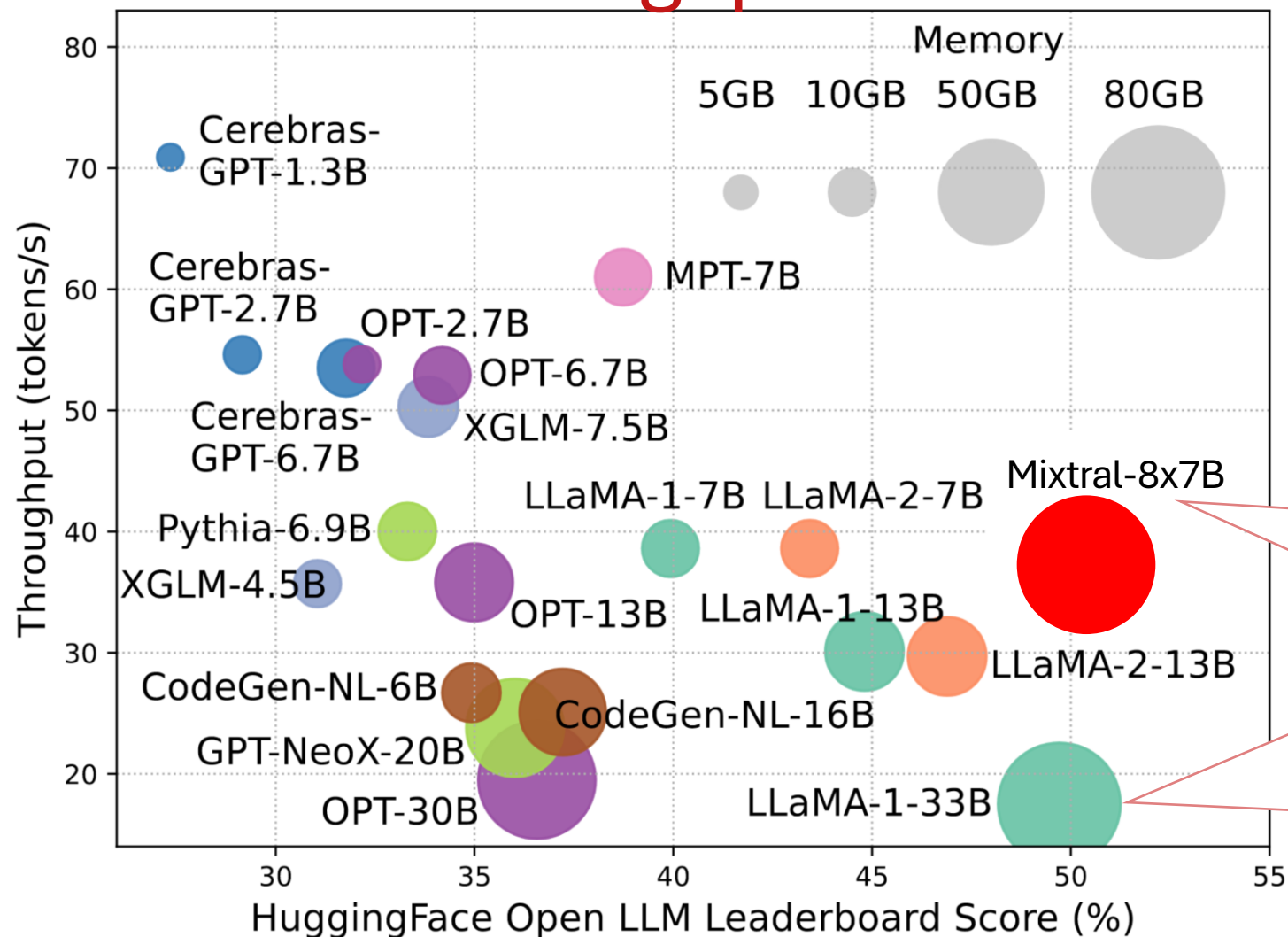
Parallelism ...

- Efficient implementation

- Flash Attention
- Paged Attention



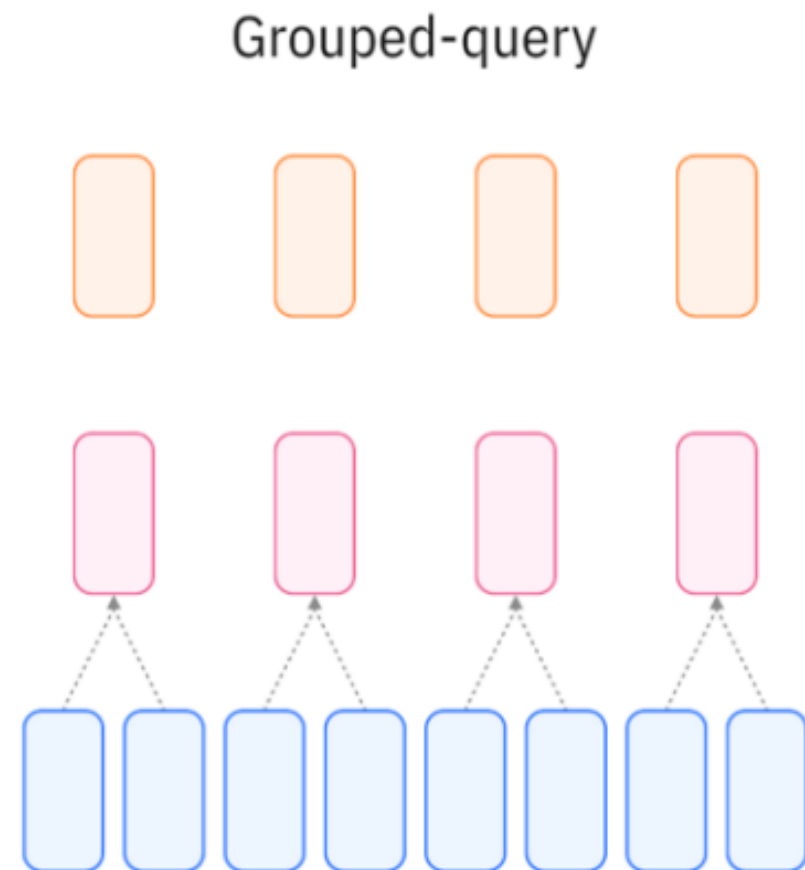
Inference Throughput vs Performance



- Similar performance, different throughput! How?
- Efficient design –
 - Grouped Query Attention
 - MoE

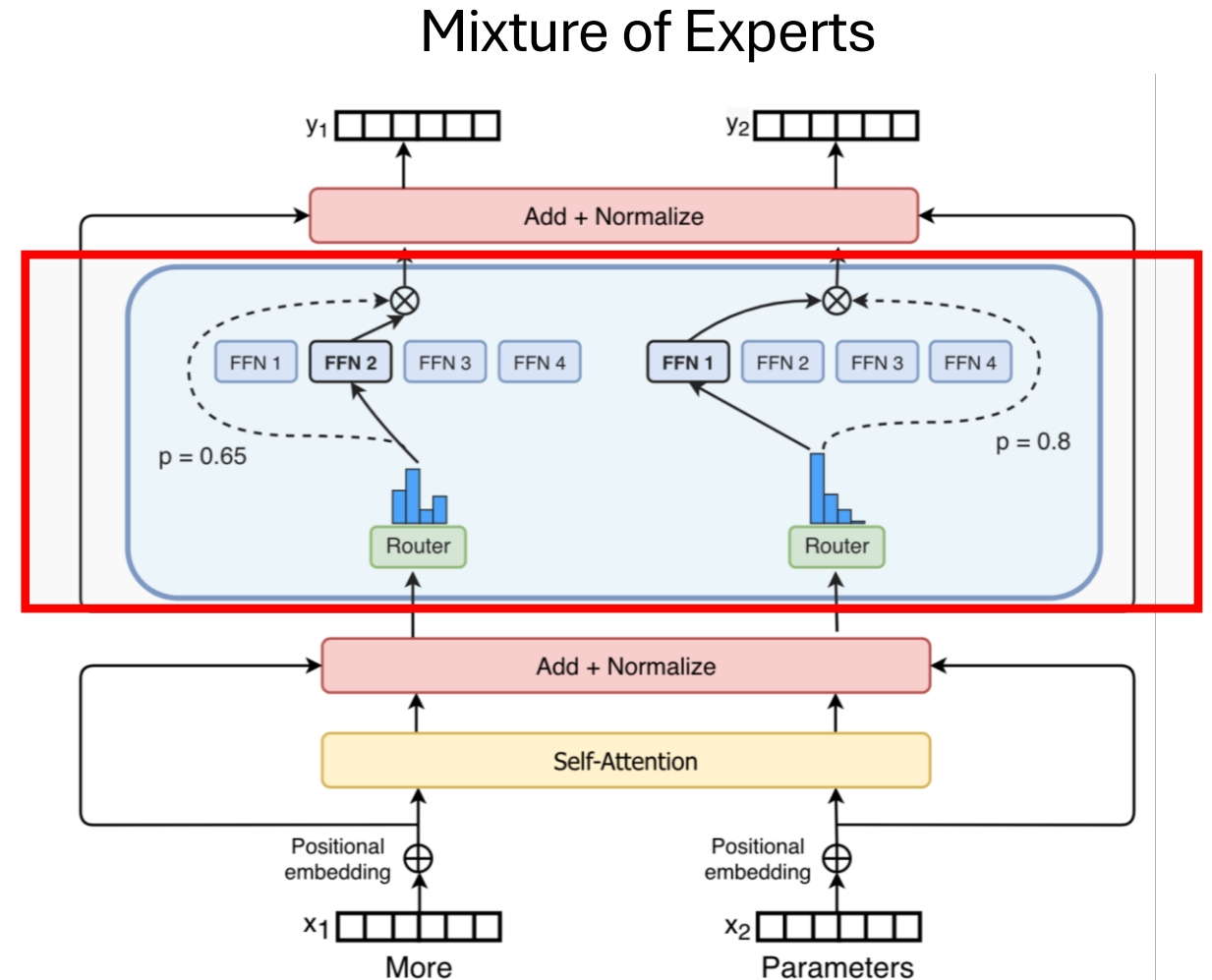
Efficient LLMs

- How to scale training?
Parallelism ...
- Efficient implementation
 - *Fused kernels ...*
- Efficient design
 - Grouped Query Attention
 - Mixture of Experts



Efficient LLMs

- How to scale training?
Parallelism ...
- Efficient implementation
 - *Fused kernels ...*
- Efficient design
 - Grouped Query Attention
 - Mixture of Experts



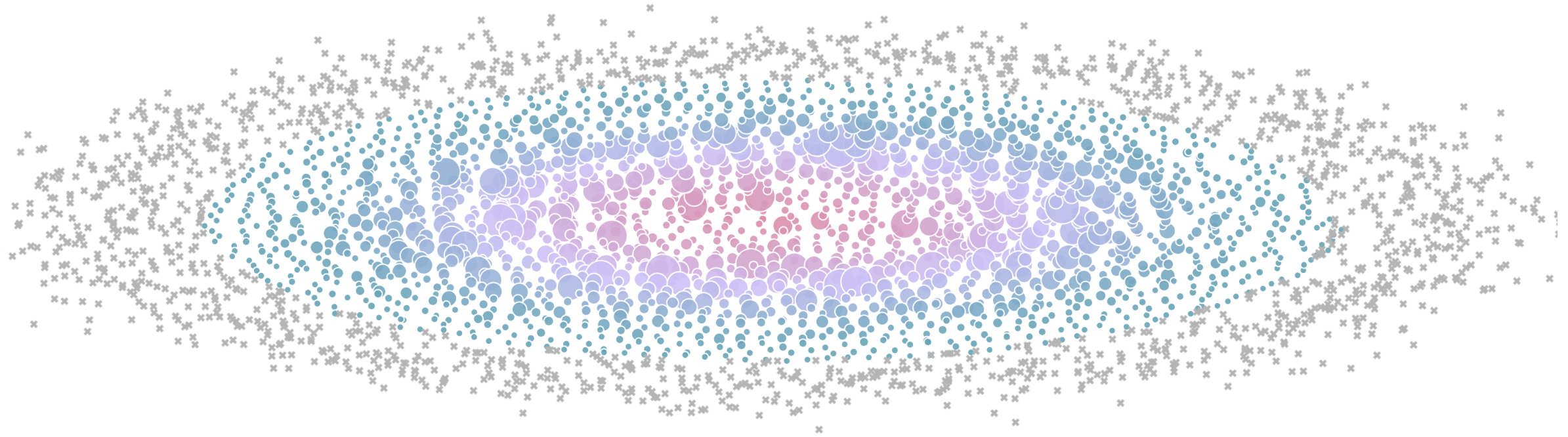
Efficient LLMs

- How to scale training?
Parallelism ...
- Efficient implementation
 - *Fused kernels ...*
- Efficient design
 - *Grouped Query Attention*
 - *Mixture of Experts*



The Ultra-Scale Playbook: Training LLMs on GPU Clusters

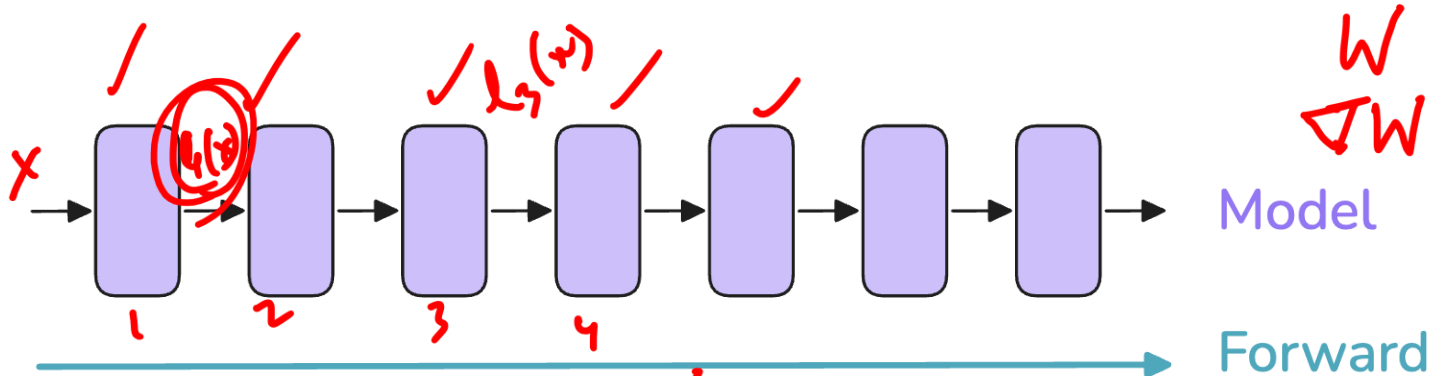
<https://huggingface.co/spaces/nanotron/ultrascale-playbook>



Over 4,000 scaling experiments on up to 512 GPUs and measured throughput (size of markers) and GPU utilization (color of markers) - Normalized per model size in this visualization.



First Steps: Training on One GPU



- ✓ (1) Model params. ✓
- ✓ (1) batch X
- ✓ (2) activations ✓
- ✓ (3) Gradients
- ✓ (4) optim state (2)

(1) MLP (nn. modules)

$f_1 = \text{nn. Linear}$

$f_2 = \text{LU}$

$f_3 = \text{nn. Linear}$

1. Forward Pass: pass inputs through the model to yield its outputs and compute Loss

forward(self, x)

for batch in dataloader:

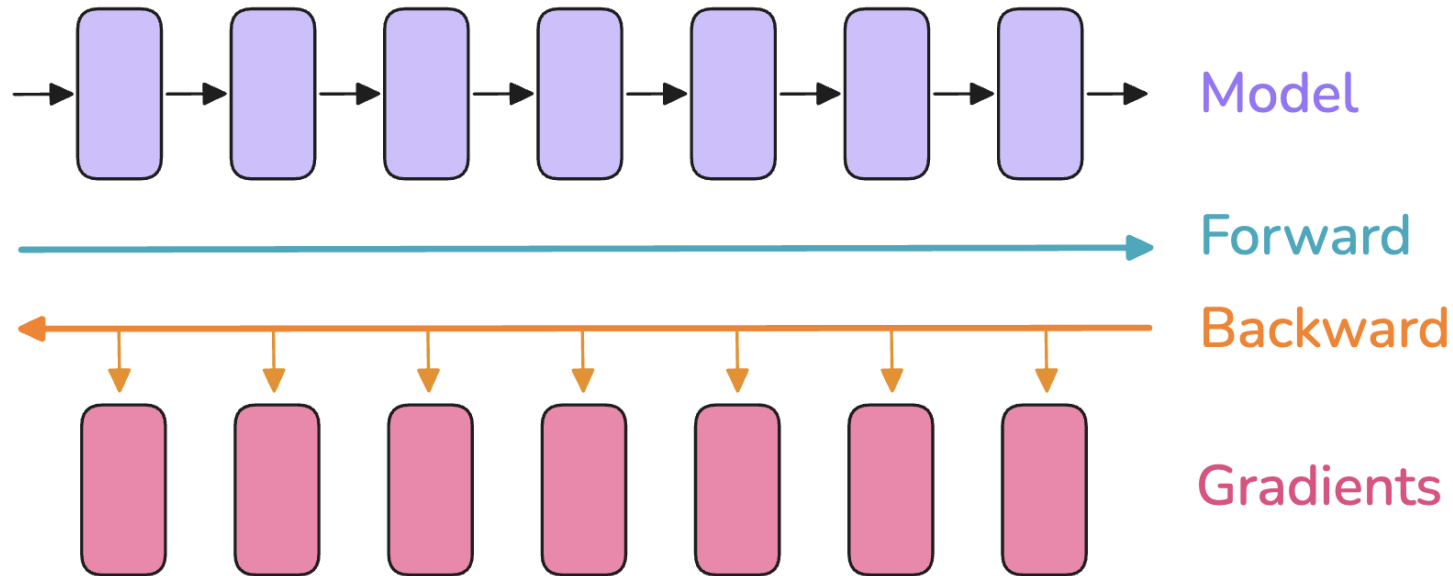
batch = batch.cuda()

→ loss = model(batch)

→ loss.backward()

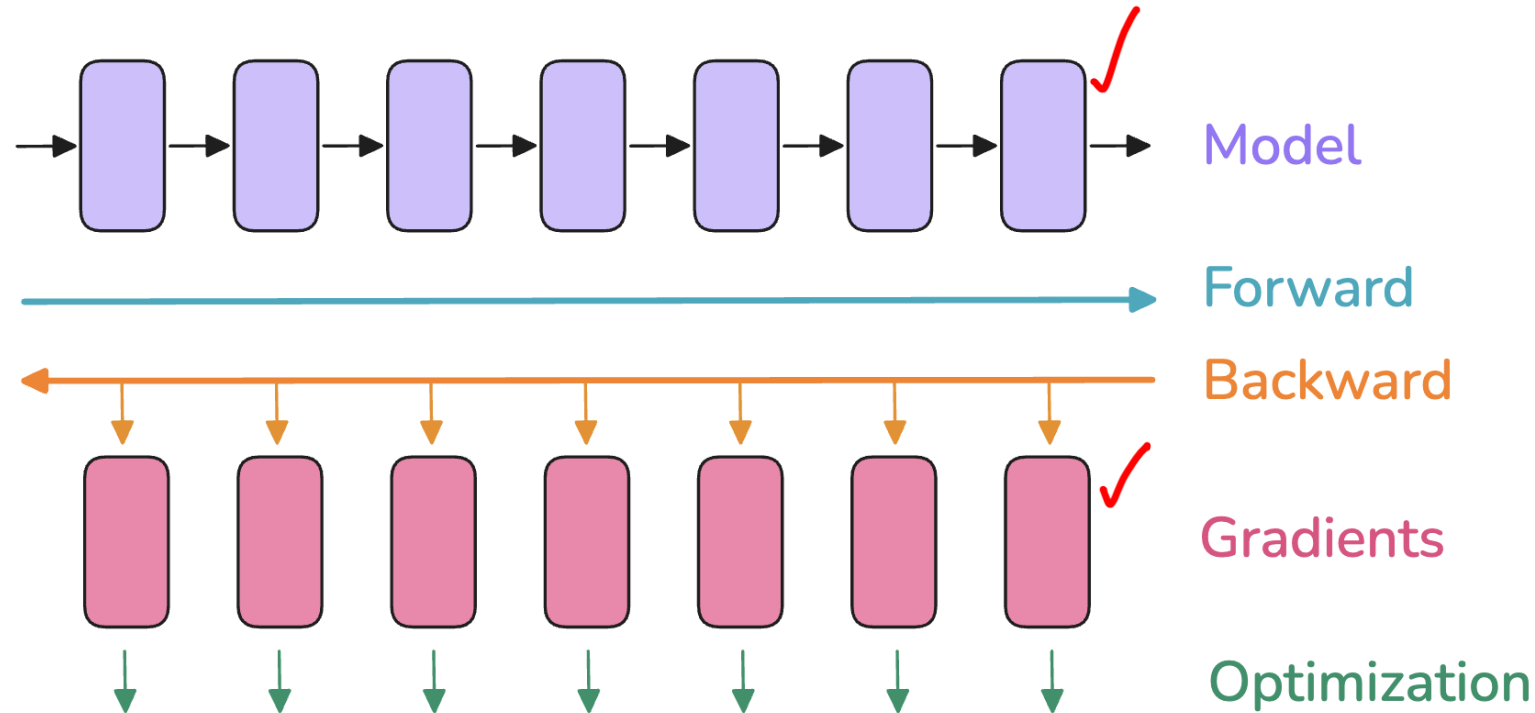
optim.step()

First Steps: Training on One GPU



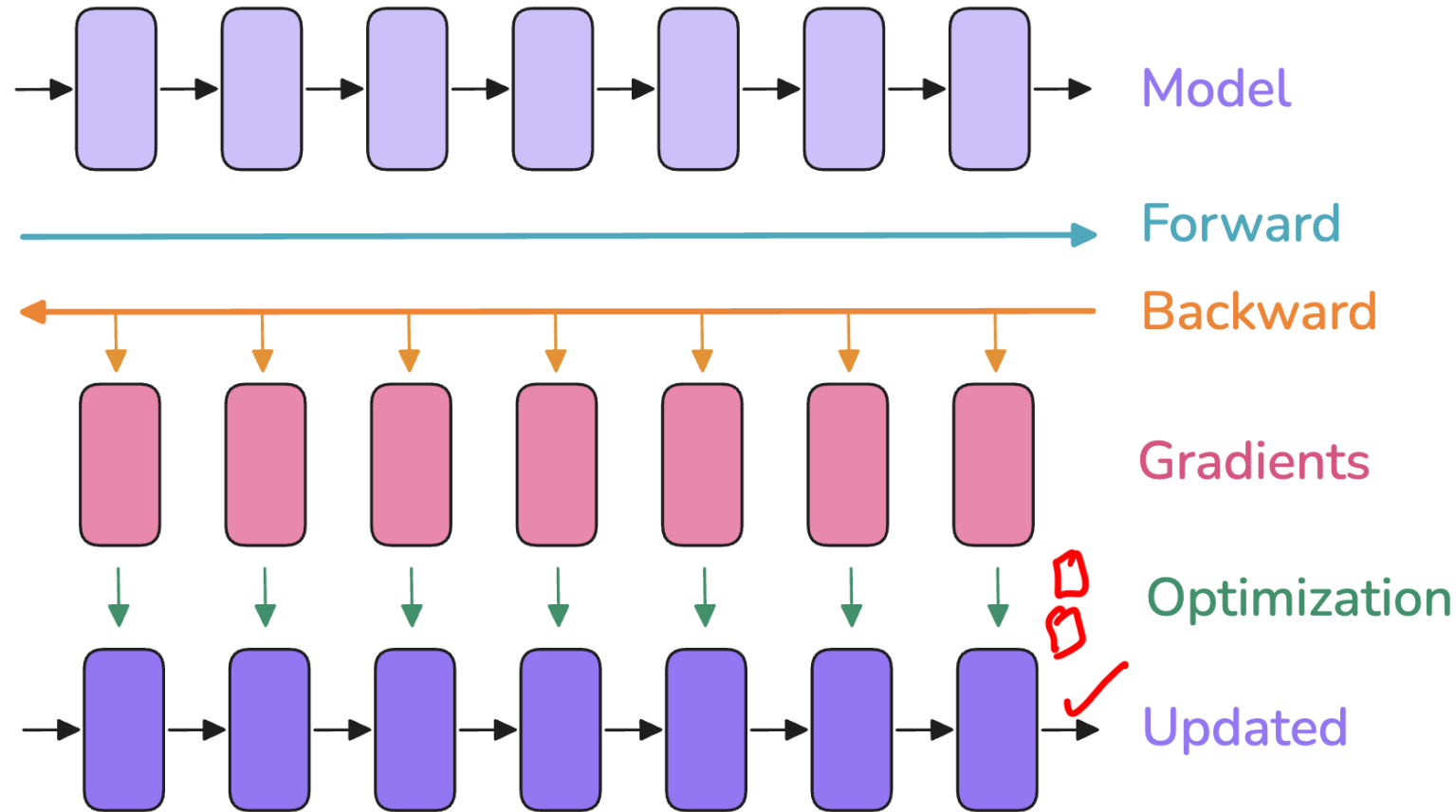
- 1. Forward Pass:** pass inputs through the model to yield its outputs and compute Loss
- 2. Backward Pass:** Compute the gradients

First Steps: Training on One GPU



- 1. Forward Pass:** pass inputs through the model to yield its outputs and compute Loss
- 2. Backward Pass:** Compute the gradients
- 3. Optimization Step:** Update the parameters using the gradients

First Steps: Training on One GPU

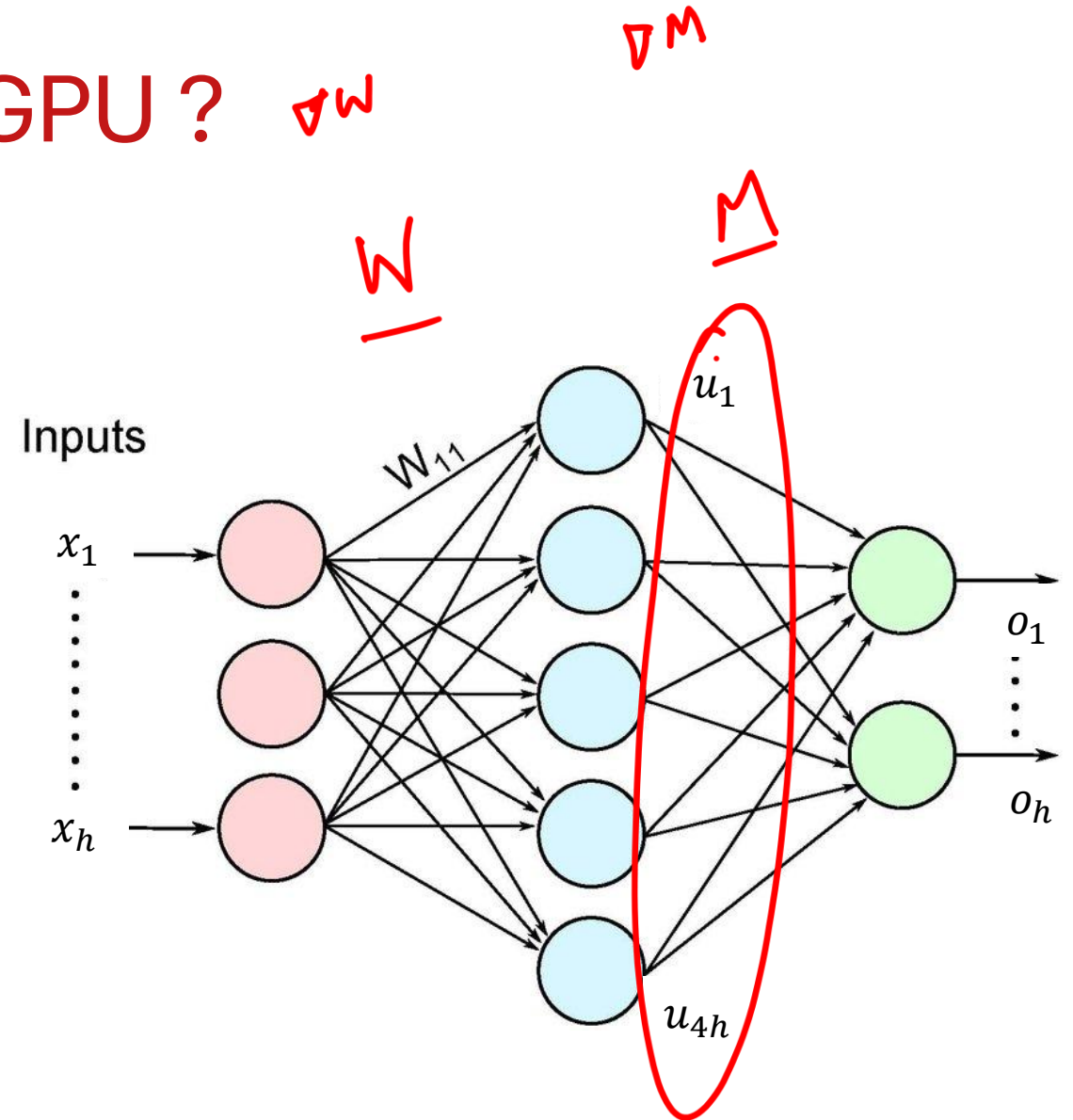


- 1. Forward Pass:** pass inputs through the model to yield its outputs and compute Loss
- 2. Backward Pass:** Compute the gradients
- 3. Optimization Step:** Update the parameters using the gradients

First Steps: Training on One GPU ? ∇W ∇M

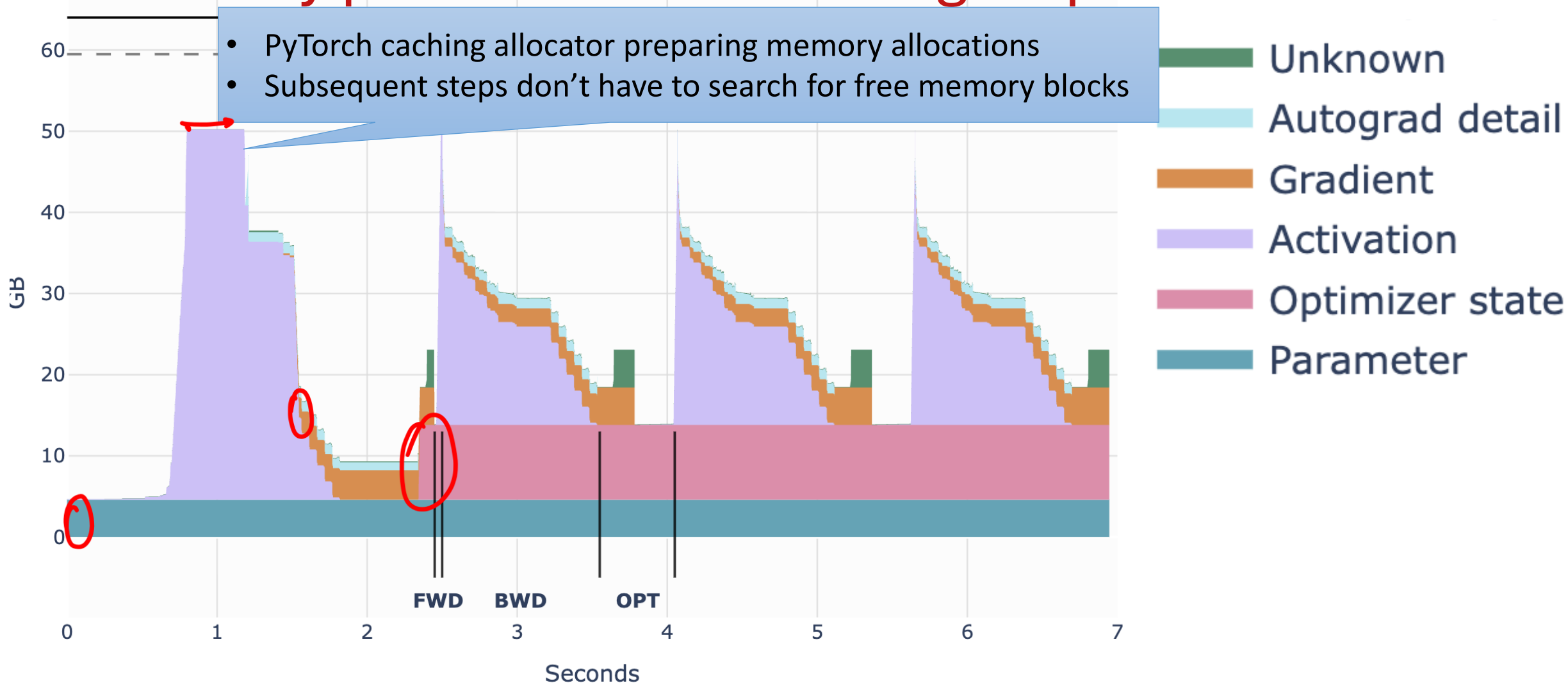
- **Memory usage in transformers**

- Model weights ✓
- Model gradients ✓
- Optimizer states ✓
- Activations needed to compute the gradients ✓



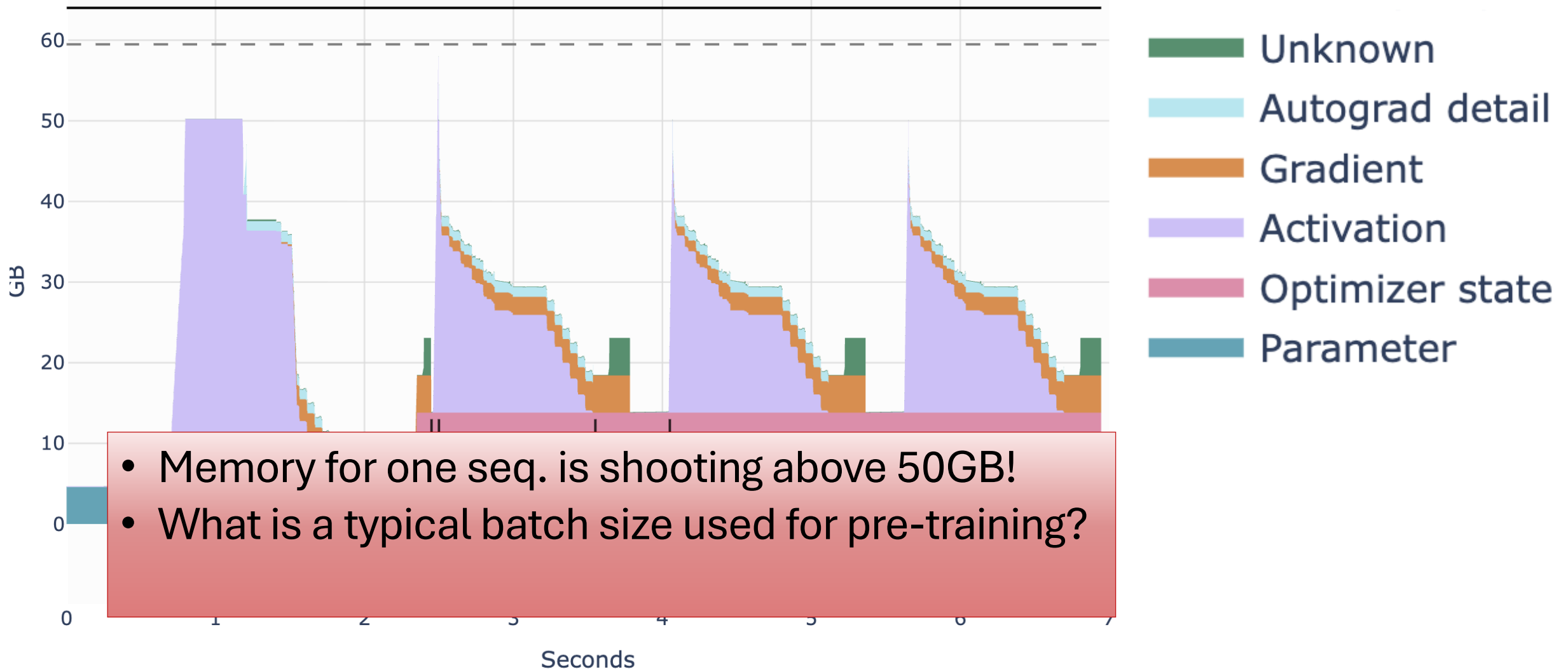
Memory profile of first 4 training steps of Llama 1B

- PyTorch caching allocator preparing memory allocations
- Subsequent steps don't have to search for free memory blocks



this = seq. x mbs
1024 x 400 4 1024 x 400

Memory profile of first 4 training steps of Llama 1B



- Memory for one seq. is shooting above 50GB!
- What is a typical batch size used for pre-training?

Batch Size

- Reported in **#Tokens** = Batch Size Tokens = $bst = bs * seq$
Llama1 – batch size ~ 4M for 1.4T tokens; DeepSeek - batch size ~ 60M for 14T tokens
- Affects both model convergence and throughput

Small batch size:

-  Noisy gradient estimates
-  Quick to compute => each step is fast
-  More optimizer steps for fixed data
-  May evade optima due to noisy gradients

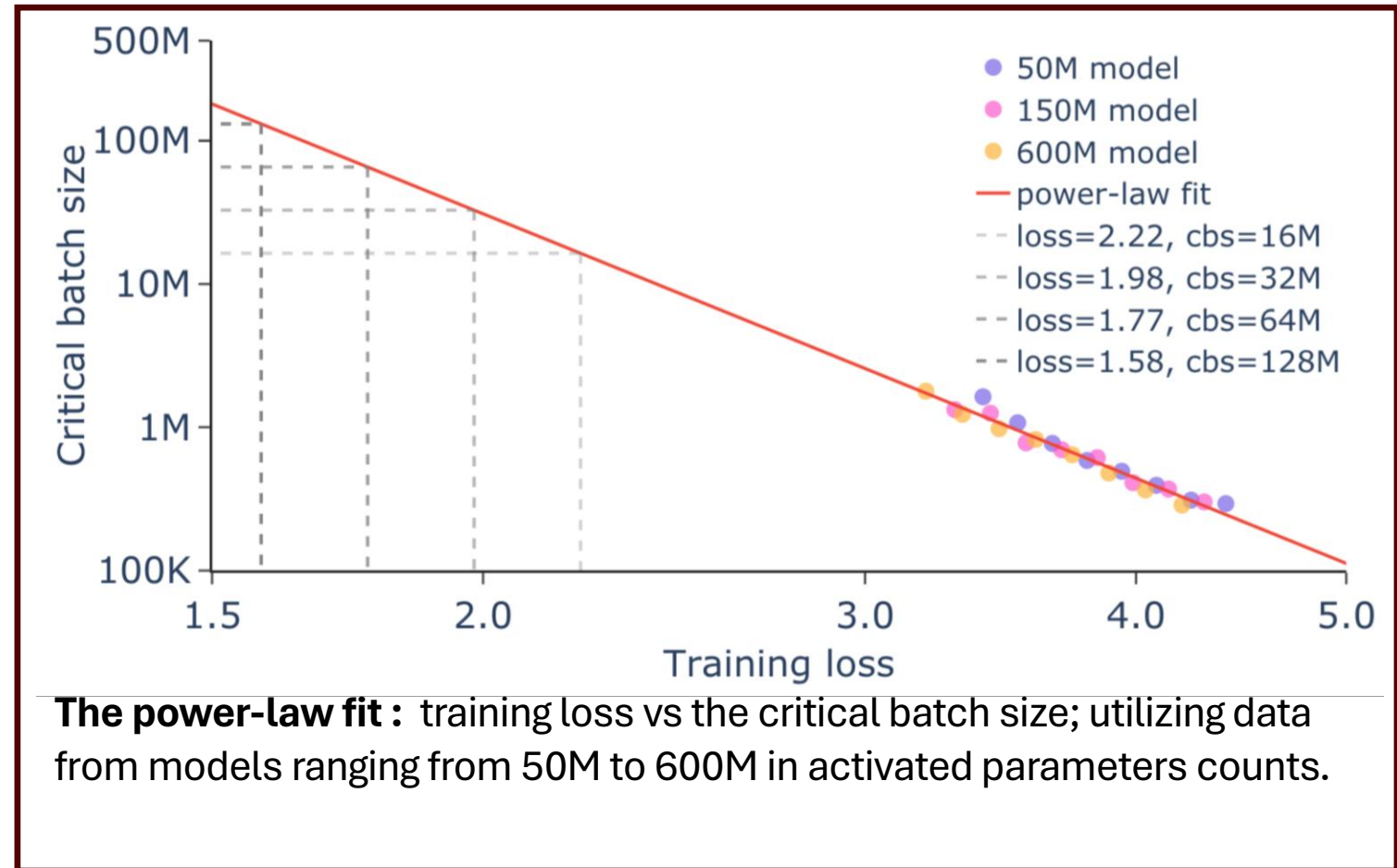
Large batch size:

-  Stable gradient estimates
-  More time to compute => each step is slow
-  Less optimizer steps for fixed data
-  Arrive at better optima due to stable gradients

Batch Size

- Training at the **critical batch size** yields a near-optimal balance between training time and data efficiency

[McCandlish et al. 2018, An Empirical Model of Large-Batch Training](#)



Batch Size

Towards the beginning...



Small Batch Size

- Quickly move through the loss landscape (multiple noisy gradient update steps)



Large Batch Size

- Accurate gradient, but slower convergence, potentially wasting compute

Towards the end...



Small Batch Size

Model may not converge to the most optimal performance (due to noisy gradients)



Large Batch Size

Accurate gradient, helps in reaching optimal performance

Batch Size: 16M

32M

64M

128M

Tokens Seen: 0

69B

790B

4700B

10,400B

Gradient update steps →

https://filecdn.minimax.chat/_Arxiv_MiniMax_01_Report.pdf

Kaplan et al. 2020 Scaling laws for neural language models.



Memory profile of first 4 training steps of Llama 1B



- Memory for one seq. is shooting above 50GB!
- What is a typical batch size used for pre-training?
- How much memory a transformer layer uses?

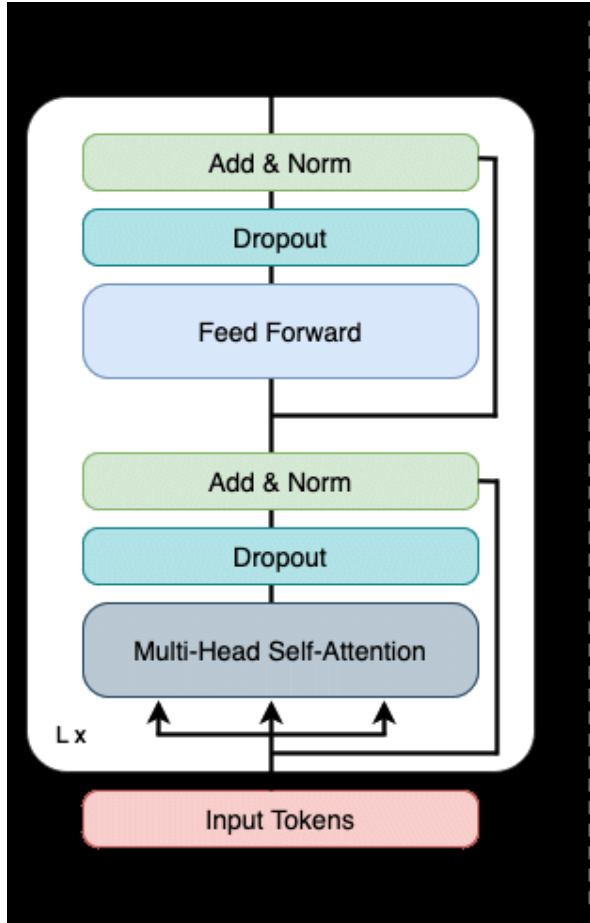
$$\mu + \sigma \left(\frac{x - \mu}{\sigma} \right)$$

$$h < \frac{1}{h} > h$$

Memory for Weights, Grads, and Optim States

- ✓ h : hidden dimension,
- ✓ v : vocabulary size,
- L : number of layers.

Total	
Layer Norm	$2h$
MLP	$(h \times 4h + 4h) + 4h^2 + h$
Layer Norm	$2h$
Projection	$h^2 + h$
MHA	$3(h^2 + h)$
Token Emb.	$h \times v$



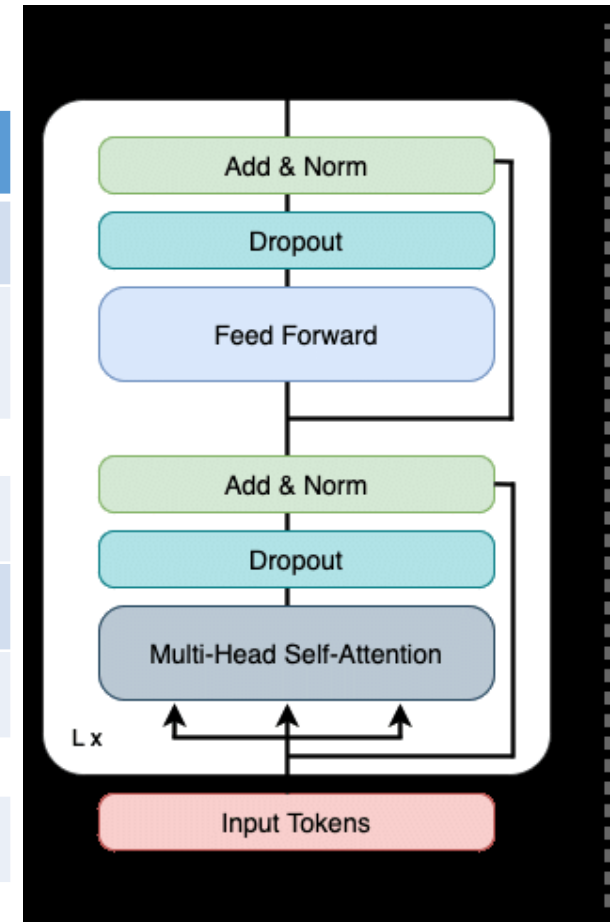
<https://michaelwornow.net/2024/01/18/counting-params-in-transformer>

Memory for Weights, Grads, and Optim States

- $N = \underline{h * v} + \underline{L * (12 * h^2 + 13 * h)} + \underline{2 * h}$

h : hidden dimension,
 v : vocabulary size,
 L : number of layers.

Total	
Layer Norm	$2 * h$
MLP	$h * 4h + 4h + 4h * h + h$
Layer Norm	$2 * h$
Projection	$h * h + h$
MHA	$3 * (h * h + h)$
Token Emb.	$h * v$



<https://michaelwornow.net/2024/01/18/counting-params-in-transformer>

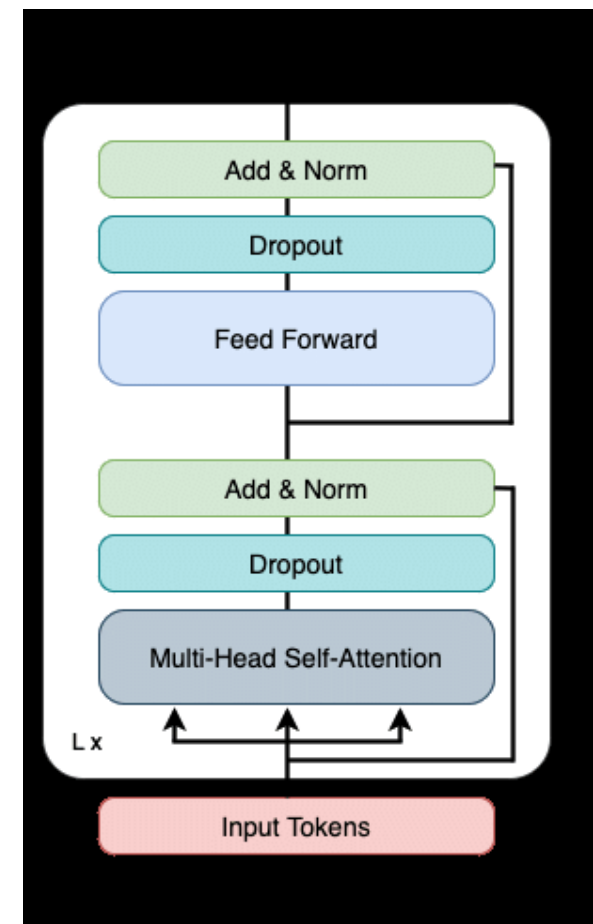
Memory for Weights, Grads, and Optim States

- $N = h * v + L * (12 * h^2 + 13 * h) + 2 * h$
- FP32 representation (4 bytes per param)
 - $m_{params} = 4 * N$ ✓
 - $m_{grad} = 4 * N$ ✓
 - $m_{opt} = (4 + 4) * N$ ✓

Total: $16 * N$ bytes
- Mixed precision: computation in bf16 (2 bytes); storage in FP32
 - $m_{params} = 2 * N$ }
 - $m_{grad} = 2 * N$ }
 - $m_{opt} = (4 + 4) * N$ 8N
 - $m_{params_fp32} = 4 * N$ }

Total: $16 * N$ bytes

Why use mixed precision if total memory is same?



Memory for Weights, Grads, and Optim States

- $N = h * v + L * (12 * h^2 + 13 * h) + 2 * h$

- FP32 representation (4 bytes per param)

- $m_{params} = 4 * N$
- $m_{grad} = 4 * N$
- $m_{opt} = (4 + 4) * N$

Total: $16 * N$ bytes

- Mixed precision: computation in bf16 (2 bytes); storage in FP32

- $m_{params} = 2 * N$
- $m_{grad} = 2 * N$
- $m_{opt} = (4 + 4) * N$
- $m_{params_fp32} = 4 * N$
(master weights)

Total: $16 * N$ bytes

Why use mixed precision if total memory is same?

1. Allows us to use optimized lower precision operations on the GPU, which are faster
2. Reduces the activation memory requirements during the forward pass



Memory for Weights, Grads, and Optim States

- $N = h * v + L * (12 * h^2 + 13 * h) + 2 * h$

- FP32 representation (4 bytes per param)

- $m_{params} = 4 * N$
- $m_{grad} = 4 * N$
- $m_{opt} = (4 + 4) * N$

Total: $16 * N$ bytes

- Mixed precision: computation in bf16 (2 bytes); storage in FP32

- $m_{params} = 2 * N$
- $m_{grad} = 2 * N$
- $m_{opt} = (4 + 4) * N$
- $m_{params_fp32} = 4 * N$
(master weights)

Total: $16 * N$ bytes

Why use mixed precision if total memory is same?

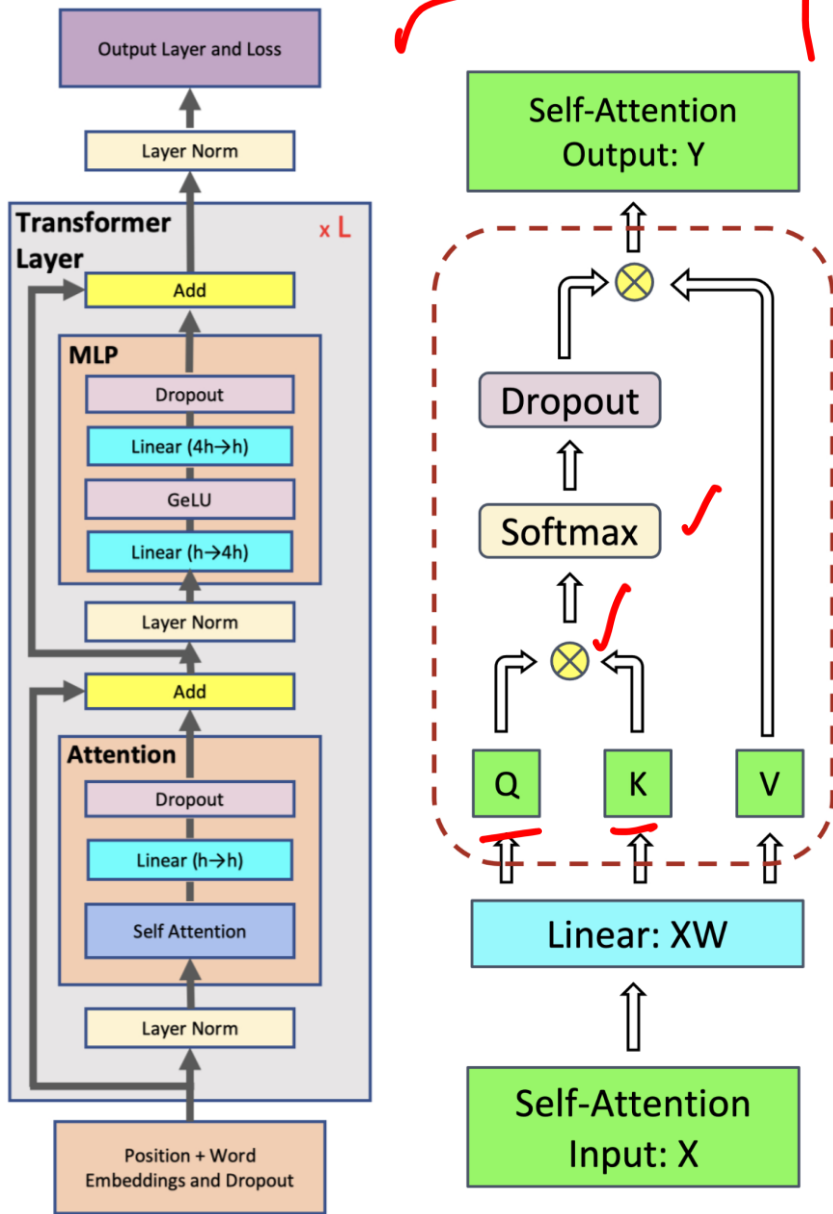
W
1st moment
2nd m

$(\nabla W)_t \rightarrow m_1$
 $(\nabla W)_t \rightarrow m_2$

Model parameters	Total Mem.
1B	16 GB
7B	112 GB
70B	1120 GB
405B	6480 GB



Memory for Activations

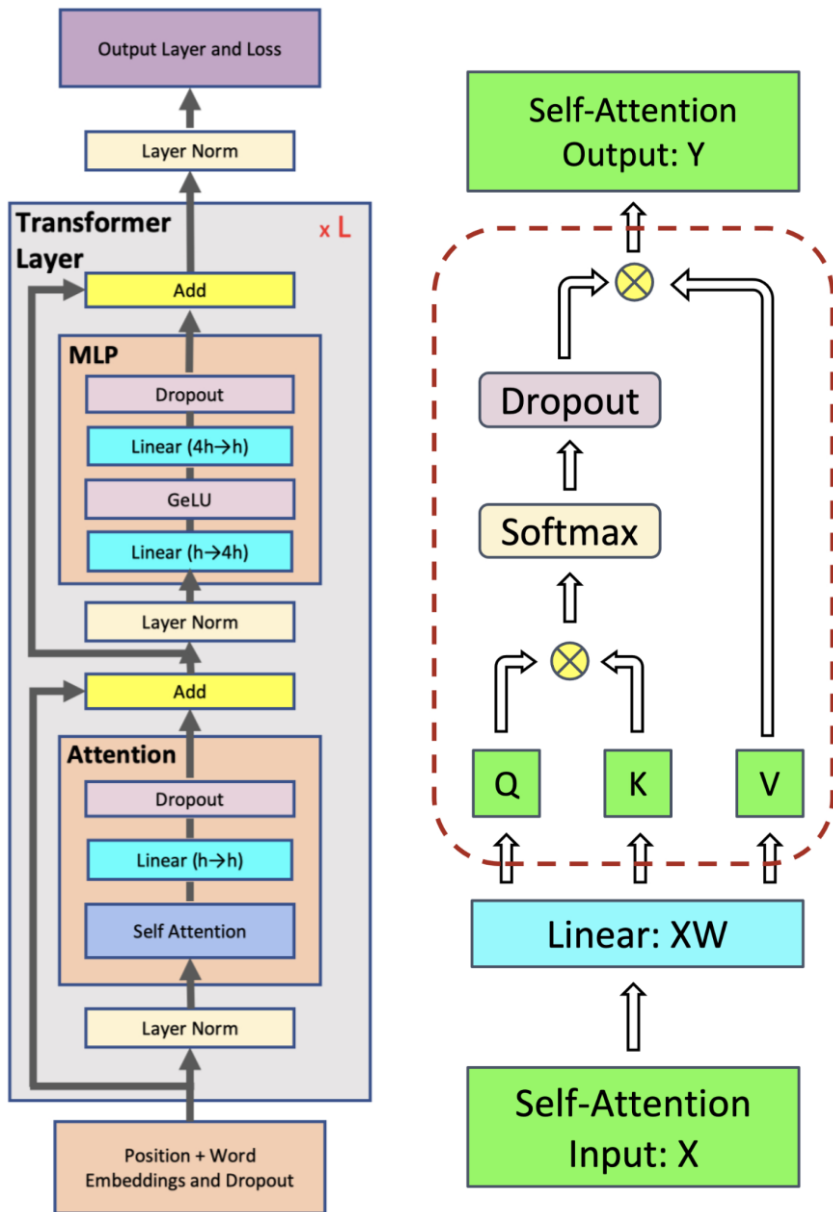


MHA	
Q, K, V input	$bs \times seq \times h$
QK^T	$o(b \times h)$
Softmax	$\alpha(n_{heads} \times seq^2 \times bs)$
Softmax d/o mask	
Prob*Values	

Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models



Memory for Activations

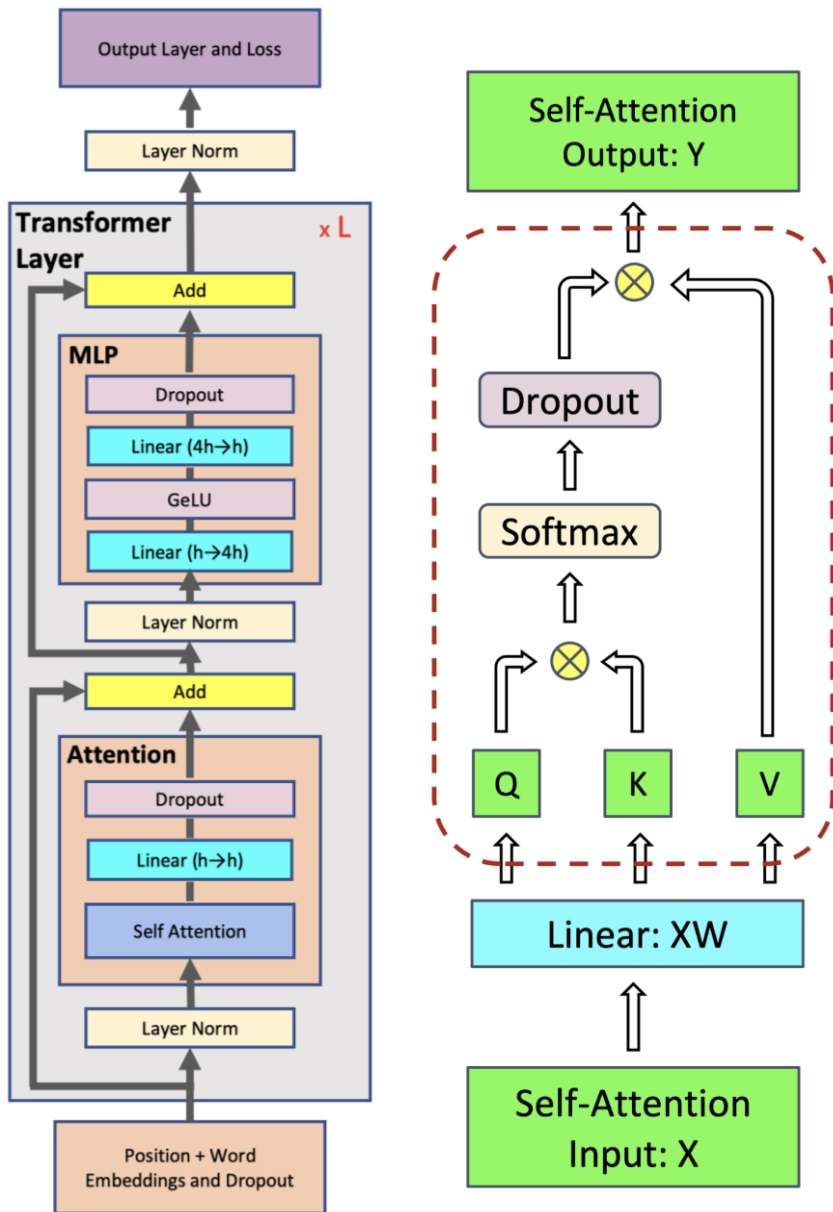


Attention Block	
Dropout mask	
Linear Proj:	
MHA	$8 * seq * bs * h + 5 * n_{heads} * seq^2 * bs$
Q, K, V input	$2 * seq * bs * h$
QK^T	$4 * seq * bs * h$
Softmax	$2 * n_{heads} * bs * seq * seq$
Softmax d/o mask	$n_{heads} * bs * seq * seq$
Prob*Values	$2 * n_{heads} * bs * seq * seq + 2 * seq * bs * h$

Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models



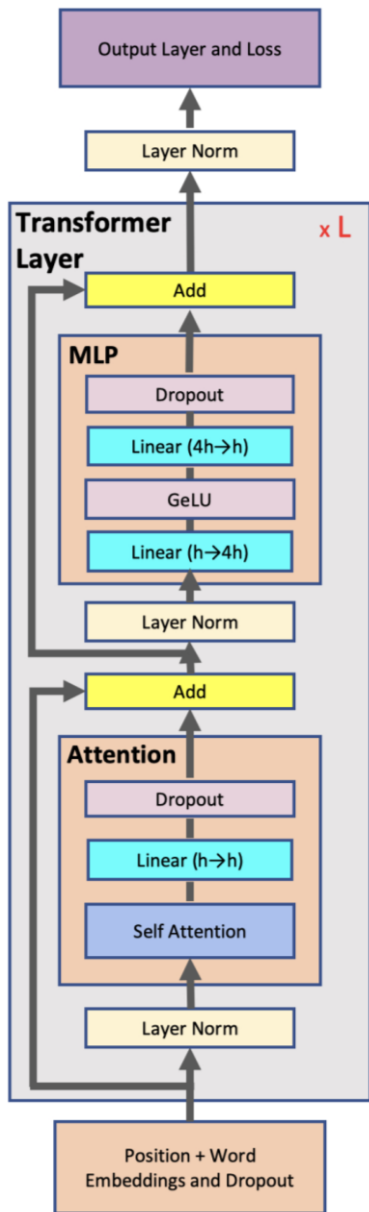
Memory for Activations



Attention Block	$11 * seq * bs * h + 5 * n_{heads} * seq^2 * bs$
Dropout mask	$seq * bs * h$
Linear Proj:	$2 * seq * bs * h$
MHA	$8 * seq * bs * h + 5 * n_{heads} * seq^2 * bs$
Q, K, V input	$2 * seq * bs * h$
QK^T	$4 * seq * bs * h$
Softmax	$2 * n_{heads} * bs * seq * seq$
Softmax d/o mask	$n_{heads} * bs * seq * seq$
Prob*Values	$2 * n_{heads} * bs * seq * seq + 2 * seq * bs * h$

Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models



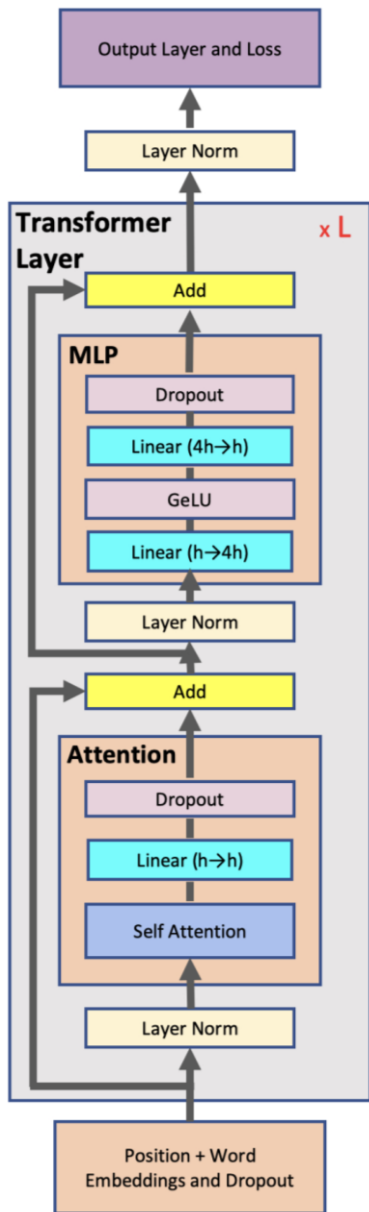


Memory for Activations

MLP Block	
D/o mask	
Linear ($4h \rightarrow h$)	
GELU	
Linear ($h \rightarrow 4h$)	
Layer Norm	
Attention Block	$11 * seq * bs * h + 5 * n_{heads} * \underline{seq^2} * bs$
Layer Norm	

[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)





Total	$34 * seq * bs * h + 5 * n_{heads} * seq^2 * bs$
MLP Block	$19 * seq * bs * h$
D/o mask	$1 * seq * bs * h$
Linear ($4h \rightarrow h$)	$8 * seq * bs * h$
GELU	$8 * seq * bs * h$
Linear ($h \rightarrow 4h$)	$2 * seq * bs * h$
Layer Norm	$2 * seq * bs * h$
Attention Block	$11 * seq * bs * h + 5 * n_{heads} * seq^2 * bs$
Layer Norm	$2 * seq * bs * h$

Memory for Activations

m_{act}

$$= L * \left(34 * seq * bs * h + 5 * n_{heads} * seq^2 * bs \right)$$

- Scales linearly with batch size
- Quadratically with the sequence length

[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)



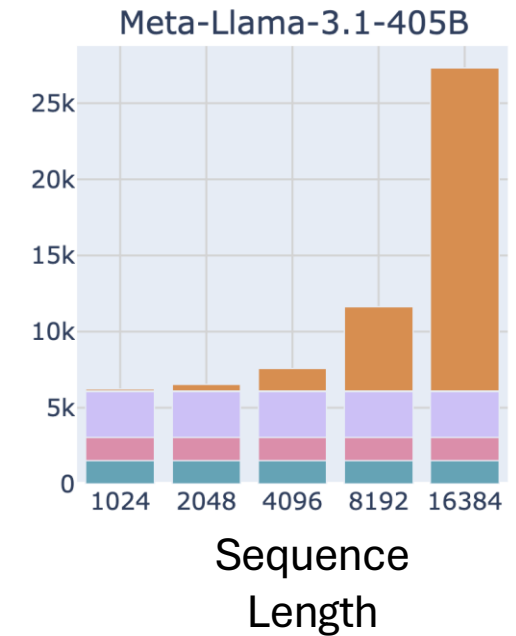
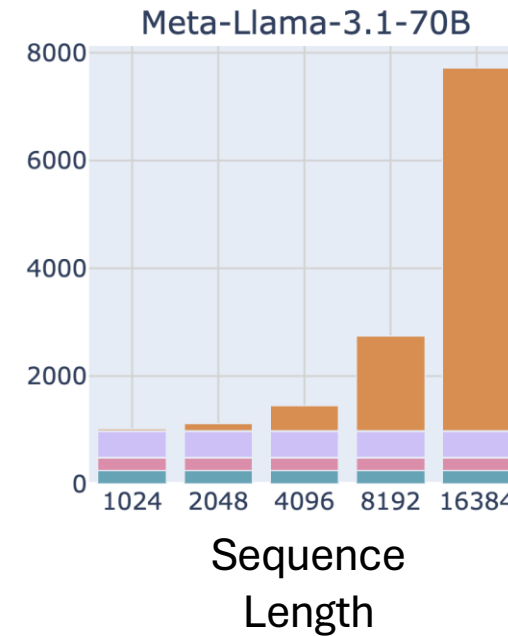
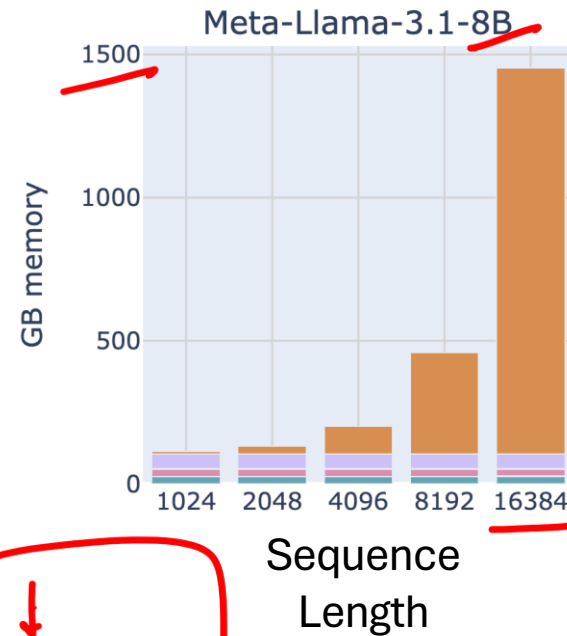
$$y = Wx$$



Memory usage in Transformers

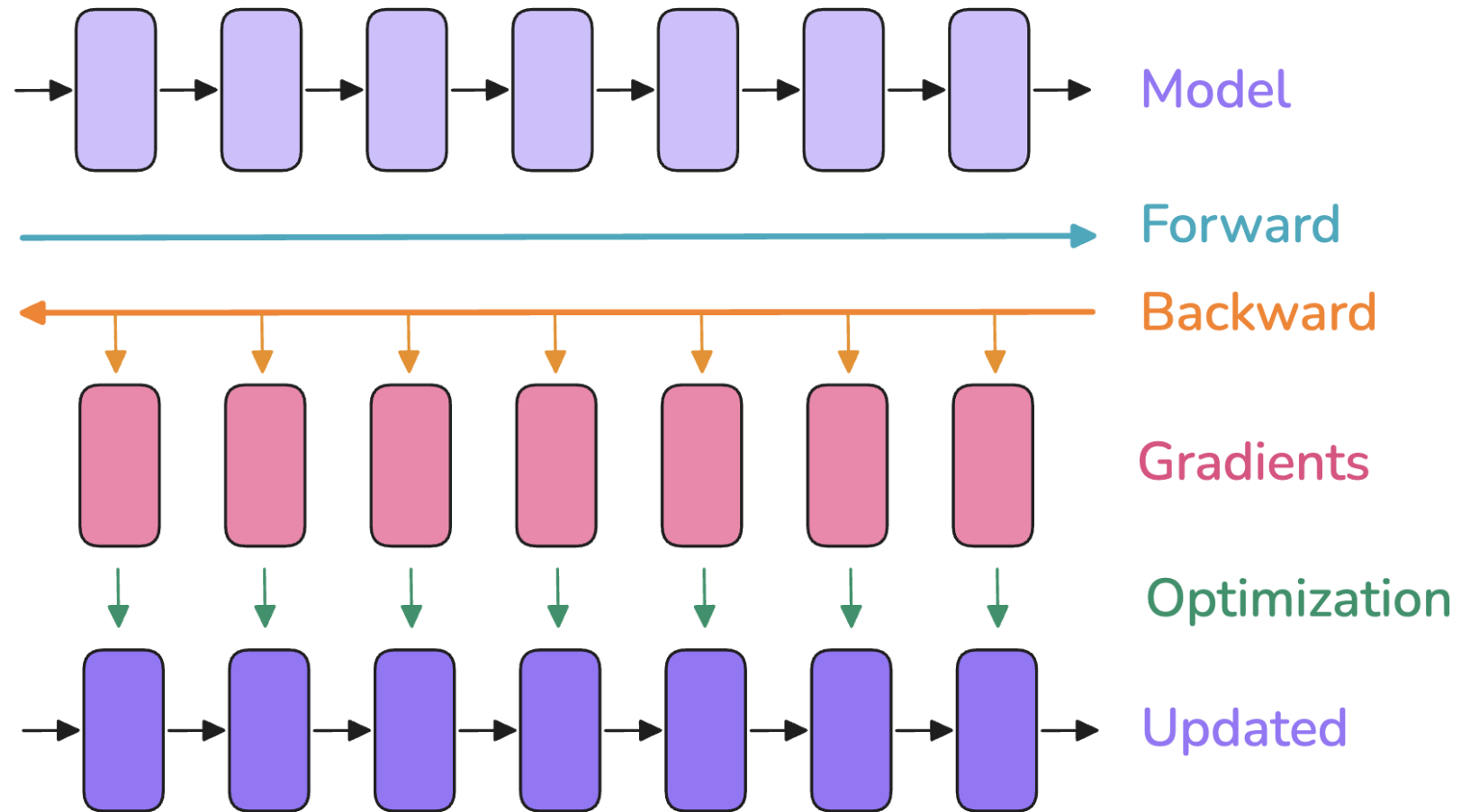
Batch Size: 1

Parameters 
 Gradients 
 Optimizer states 
 Activations 



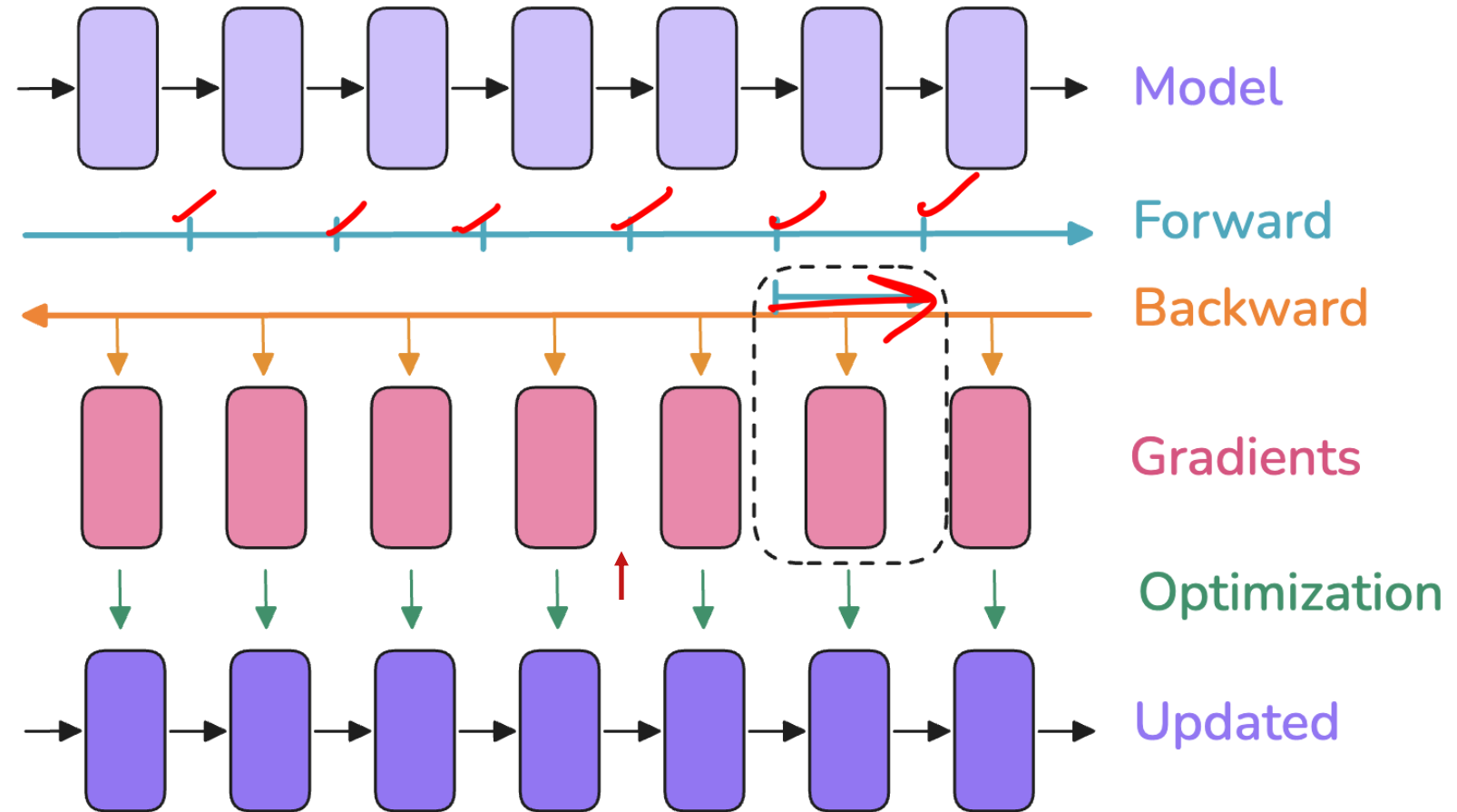
For large numbers of input tokens (i.e., large batch sizes / sequences), activations become by far the largest memory burden.

Activation re-computation *aka* Gradient Checkpointing



Activation re-computation *aka* Gradient Checkpointing

- Discard some activations during the forward pass to save memory
- Spend some extra compute to recompute these on the fly during the backward pass.



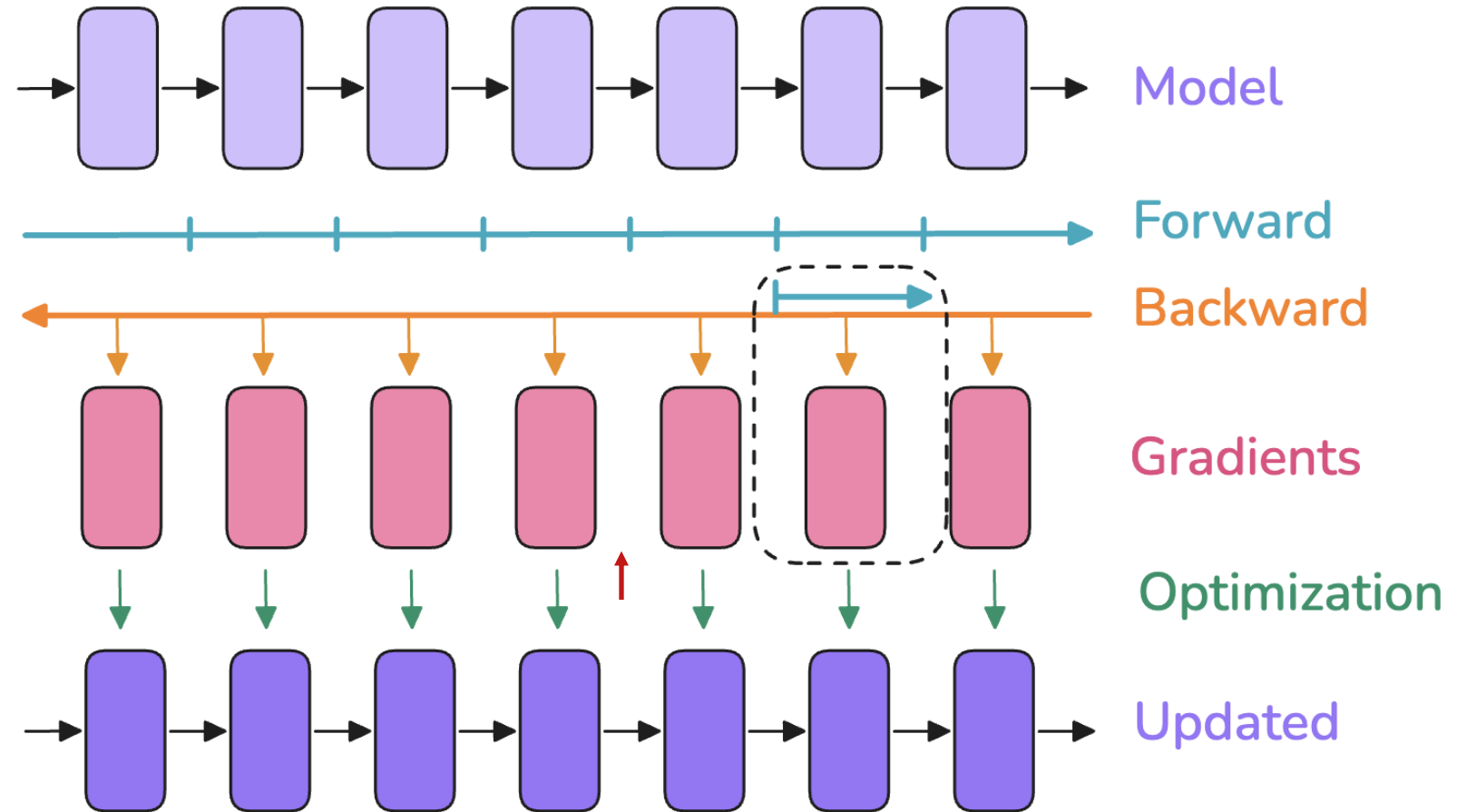
[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)



Activation re-computation *aka* Gradient Checkpointing

- Discard some activations during the forward pass to save memory
- Spend some extra compute to recompute these on the fly during the backward pass.

Which activations to store?



[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)

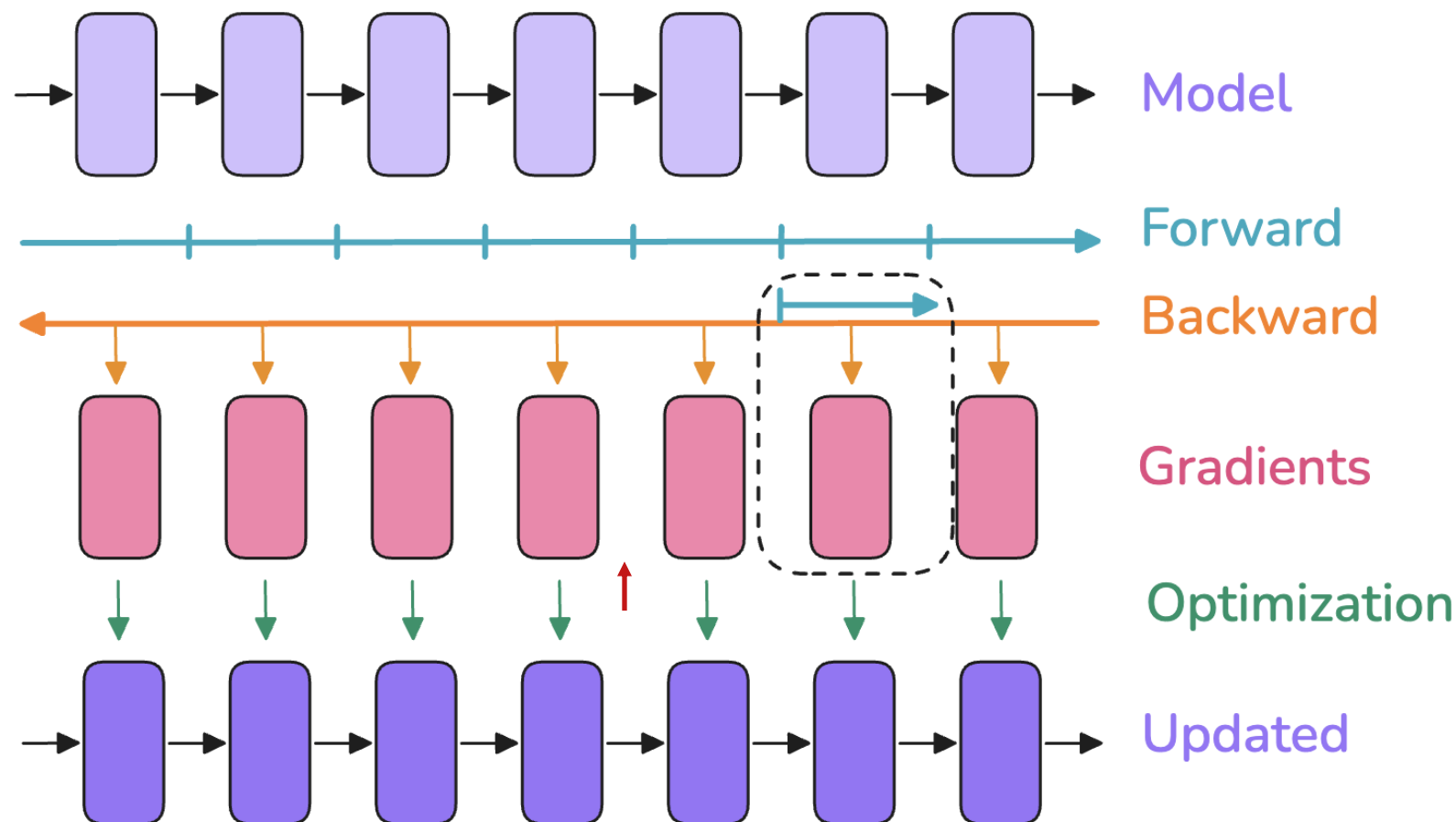


Activation re-computation *aka* Gradient Checkpointing

- Discard some activations during the forward pass to save memory
- Spend some extra compute to recompute these on the fly during the backward pass.

Which activations to store?

- **Full:** Save at the transition of layers
 - 👍 Saves most memory
 - 👎 Most expensive – 30-40%



[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)

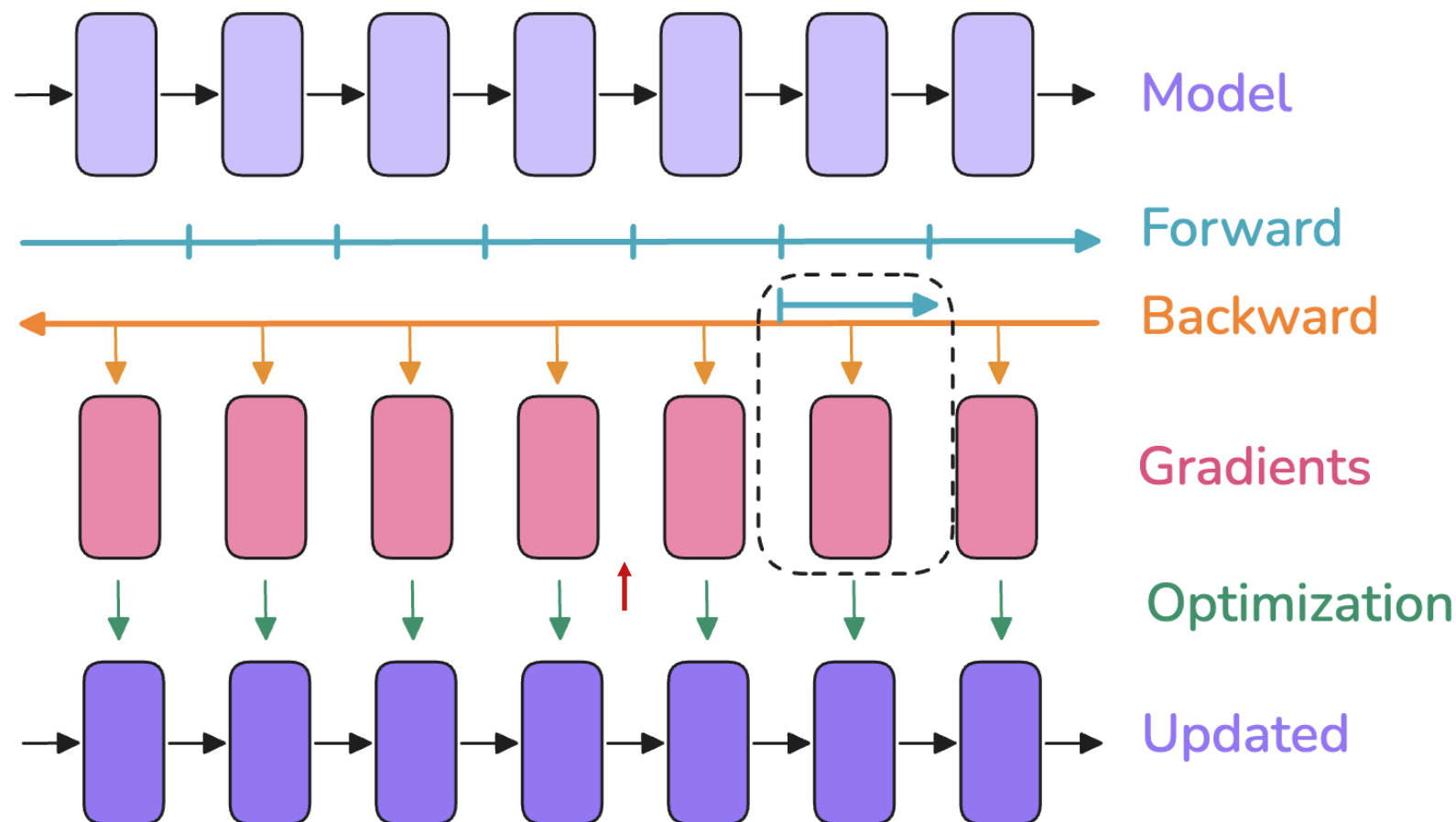


Activation re-computation *aka* Gradient Checkpointing

- Discard some activations during the forward pass to save memory
- Spend some extra compute to recompute these on the fly during the backward pass.

Which activations to store?

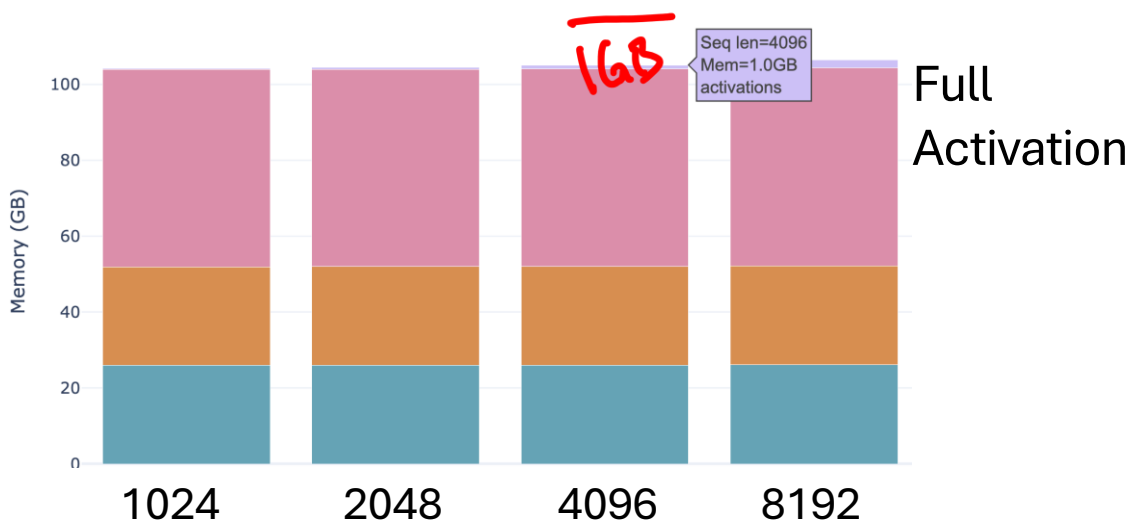
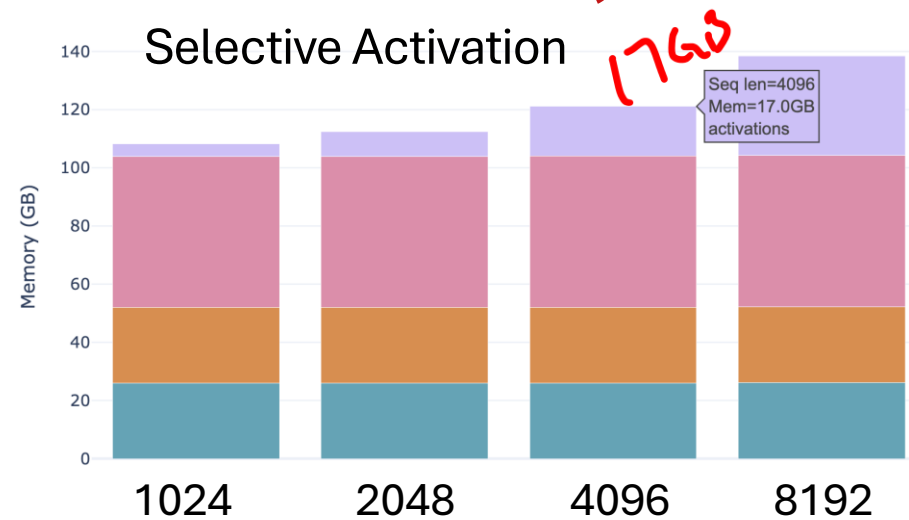
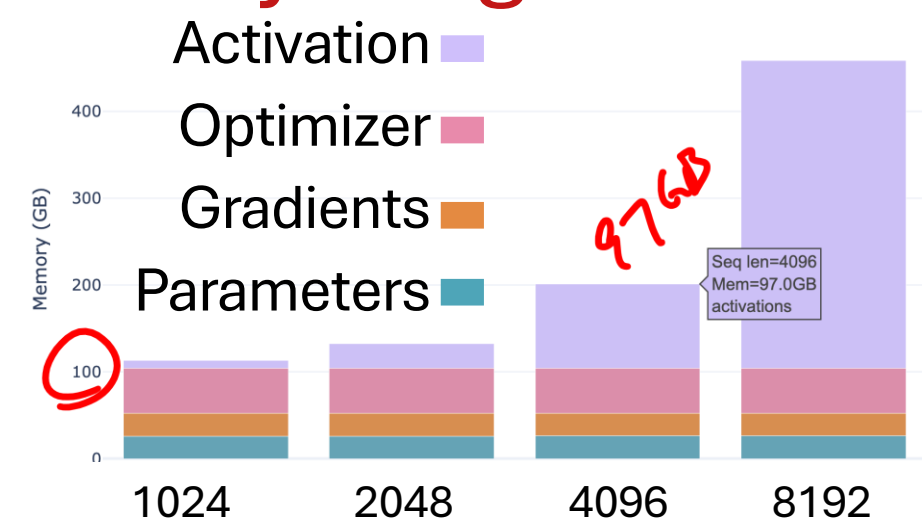
- **Full:** Save at the transition of layers
 - 👍 Saves most memory
 - 👎 Most expensive – 30-40%
- **Selective:** Save MLP; discard attn
 - 70% memory reduction ✓
 - 2.7% compute cost ✓



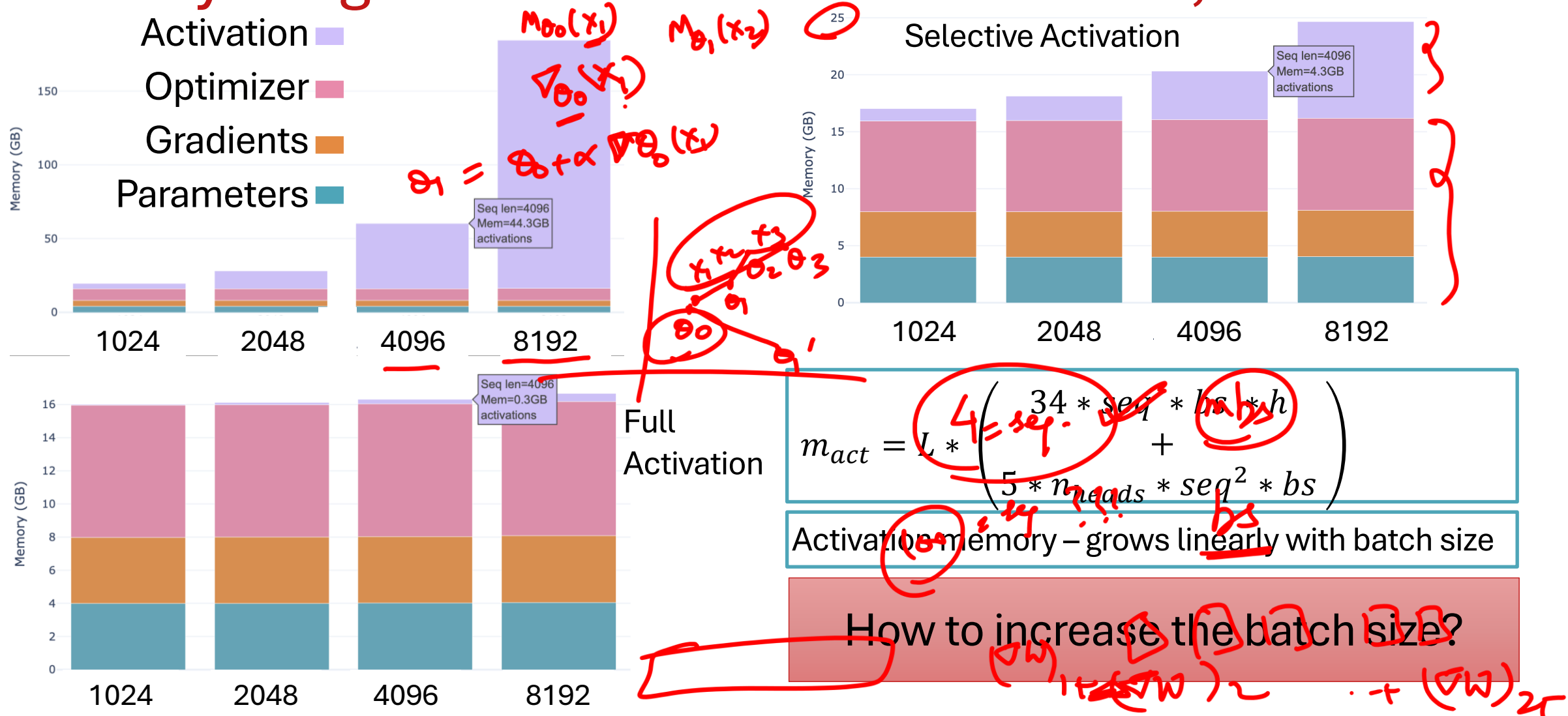
[Korthikanti et al. 2022, Reducing Activation Recomputation in Large Transformer Models](#)



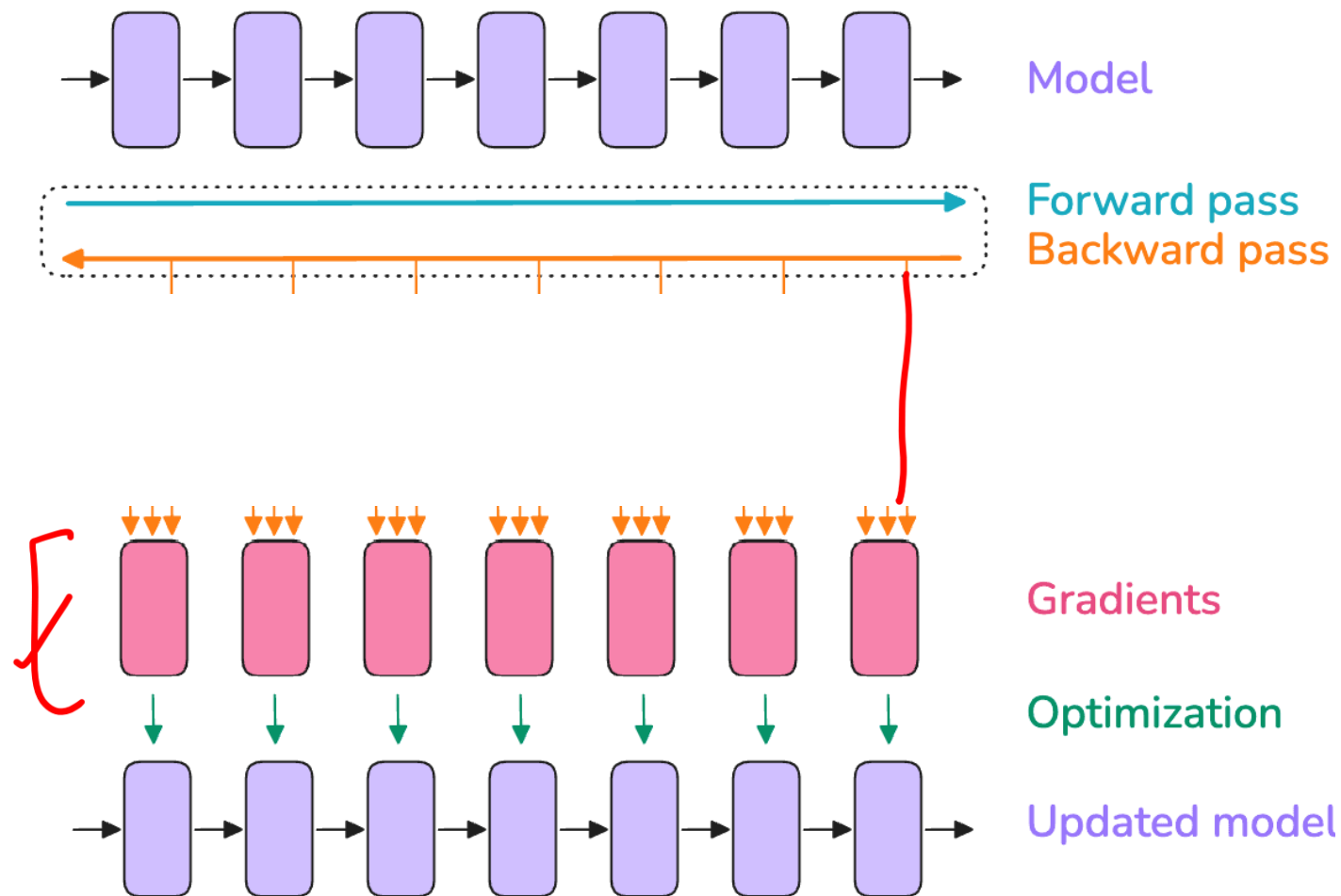
Memory usage in Transformers – 8B Model , $bs = 1$



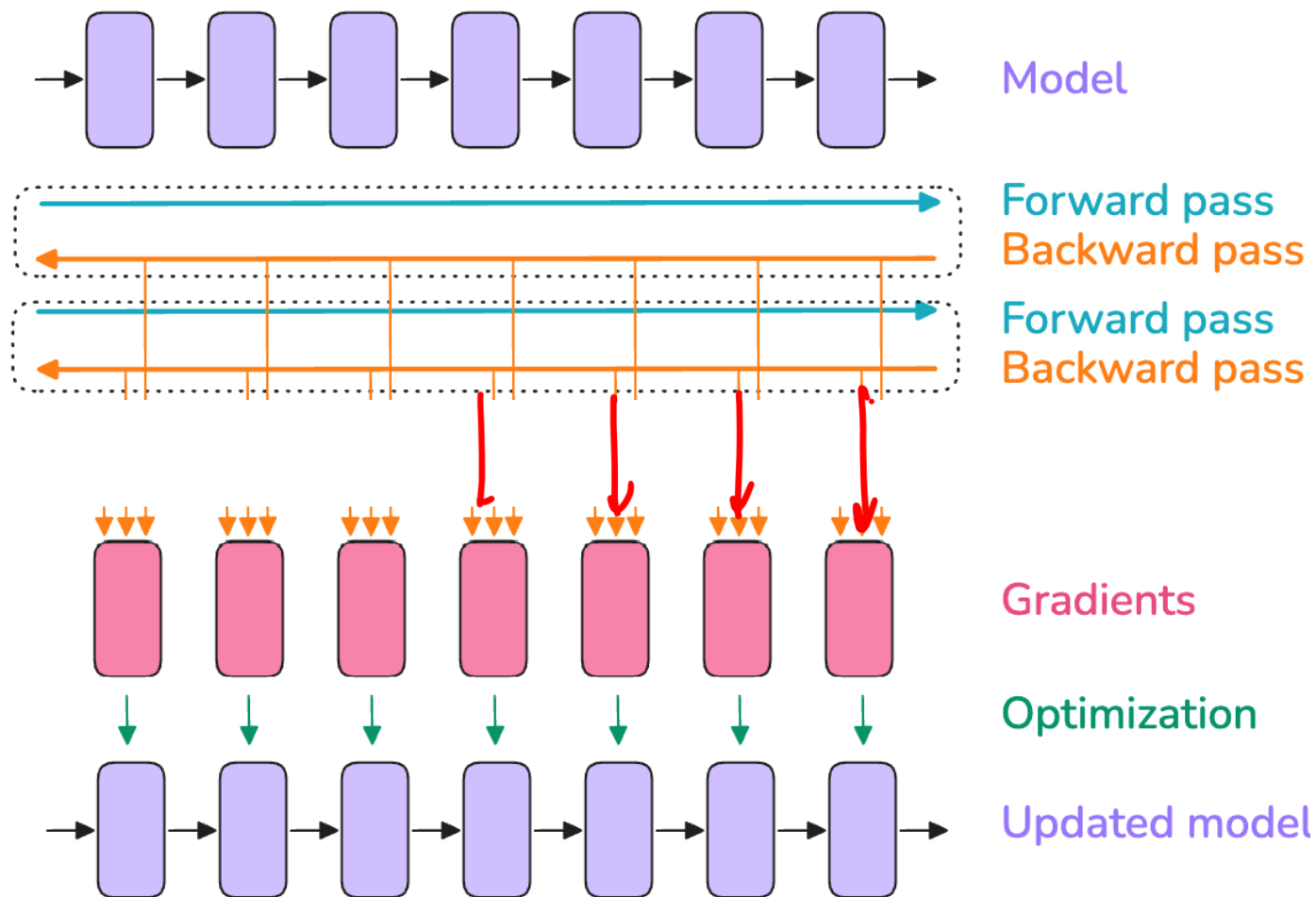
Memory usage in Transformers – 1B Model , $bs = 1$



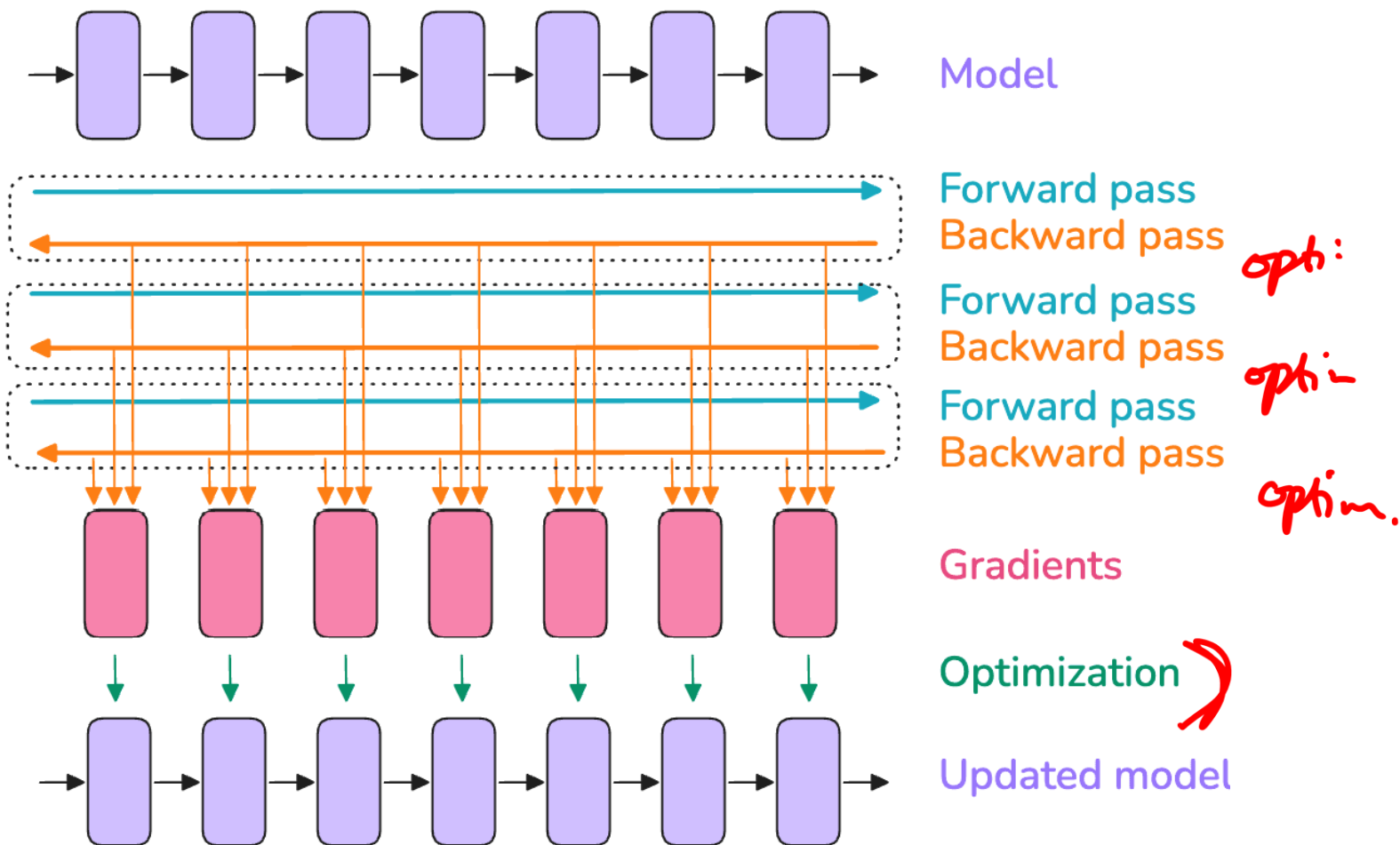
How to increase batch size? – *Gradient Accumulation*



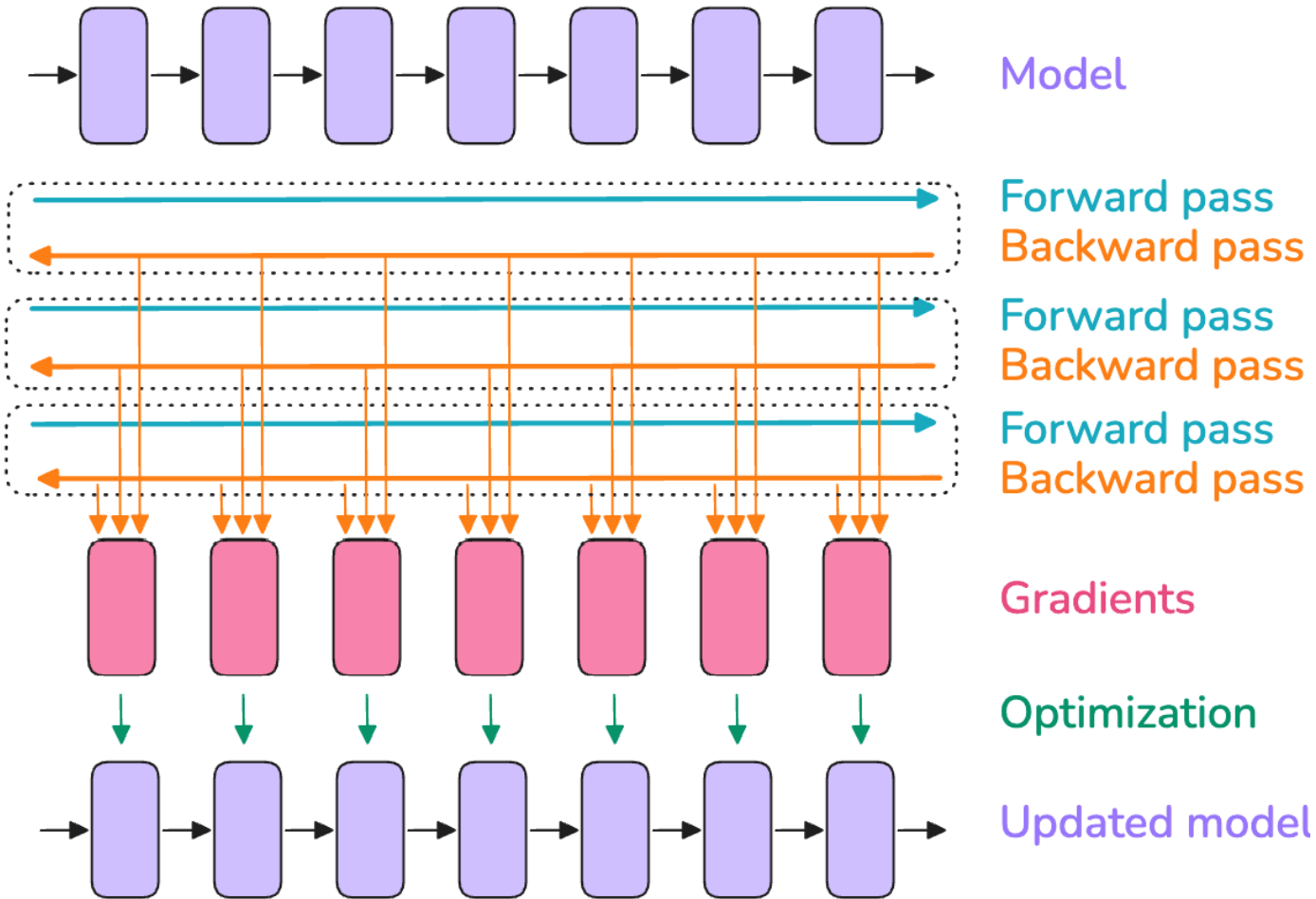
How to increase batch size? – *Gradient Accumulation*



How to increase batch size? – Gradient Accumulation



How to increase batch size? – *Gradient Accumulation*



- Process smaller micro-batches sequentially

$$bs = gbs = mbs * grad_acc$$

Only one micro-batch's worth of activations needs to be kept in memory at a time

Requires multiple consecutive forward/backward passes per optimization step

How to speedup?