



LoQT: Low-Rank Adapters for Quantized Pre-Training

Sebastian Loeschcke*, Mads Tofttrup*, Michael J. Kastoryano,
Serge Belongie, and Vésteinn Snæbjarnarson

NeurIPS 2024

UNIVERSITY OF
COPENHAGEN



AARHUS UNIVERSITY



PIONEER CENTRE FOR
ARTIFICIAL INTELLIGENCE

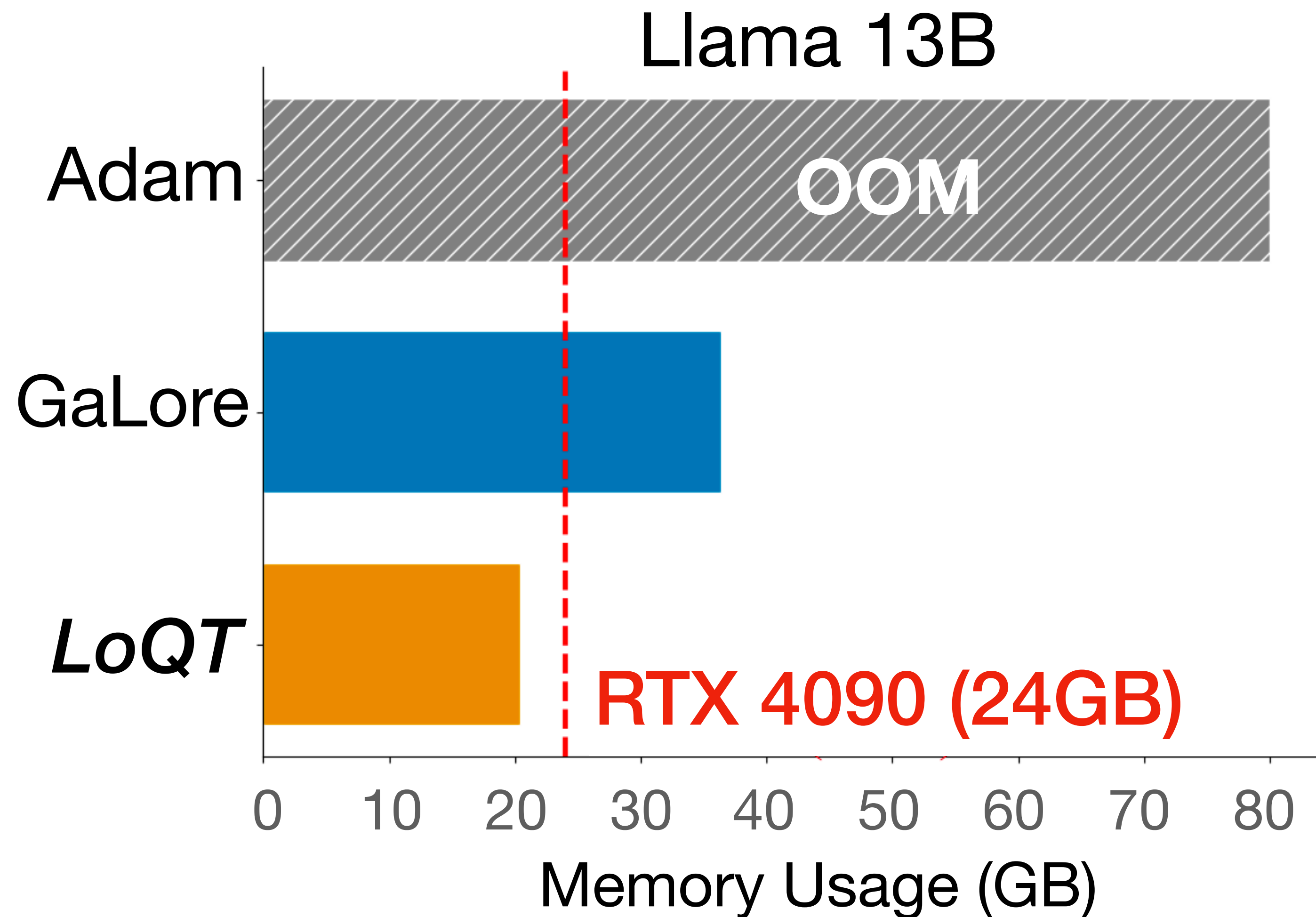
**Equal Contribution*

Pre-training LLMs require big GPU clusters

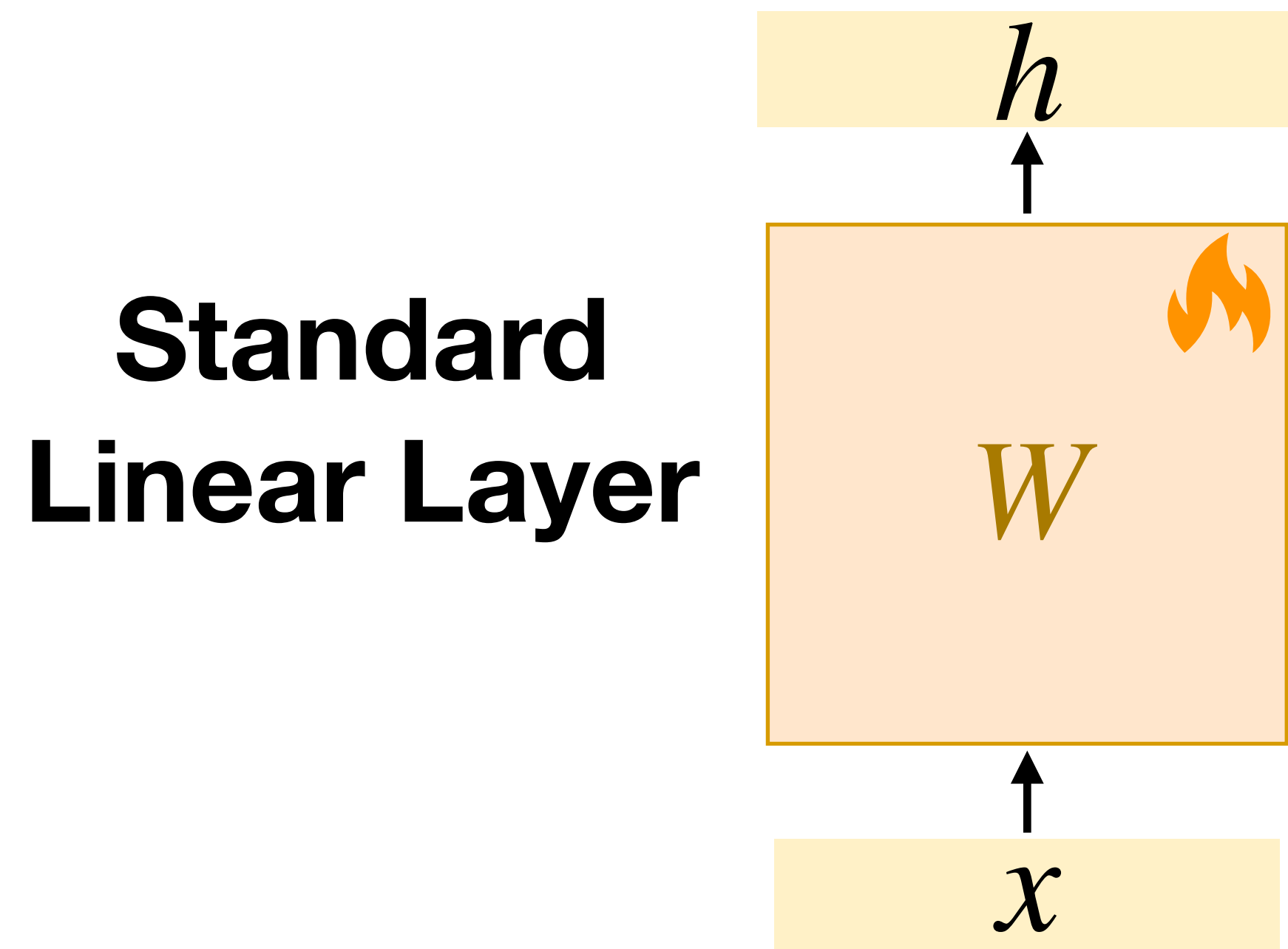
Pre-training a **LLaMA 13B** model with Adam and a *batch size of one* requires over **100 GB** of memory.

Pre-training LLMs require big GPU clusters

Pre-training a **LLaMA 13B** model with Adam and a *batch size of one* requires over **100 GB** of memory.

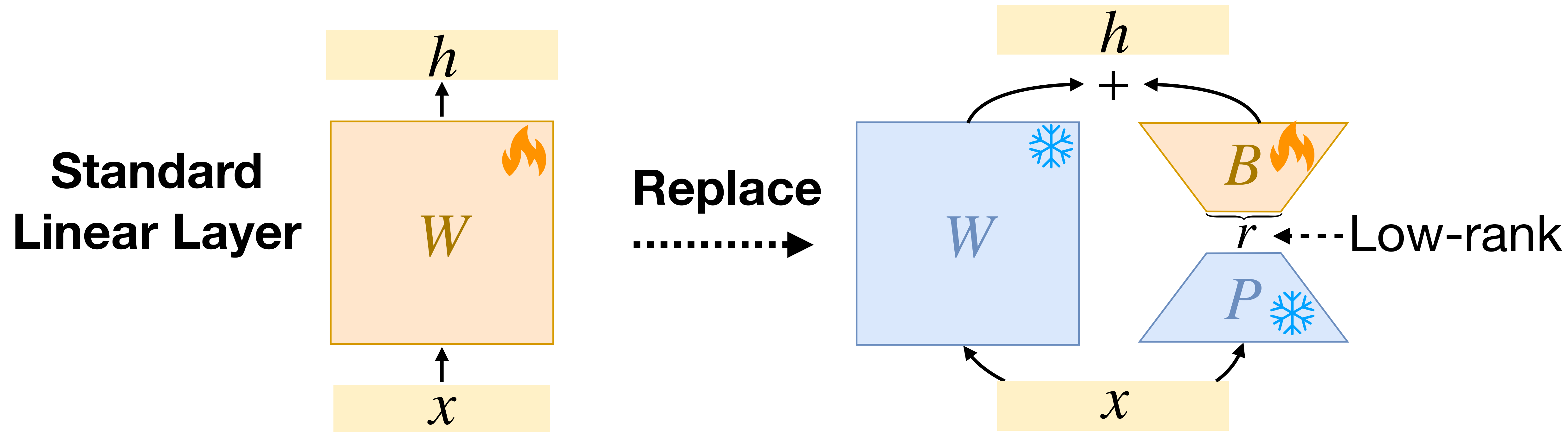


LoQT: Low-Rank Quantized Pre-Training



 Optimize  Freeze

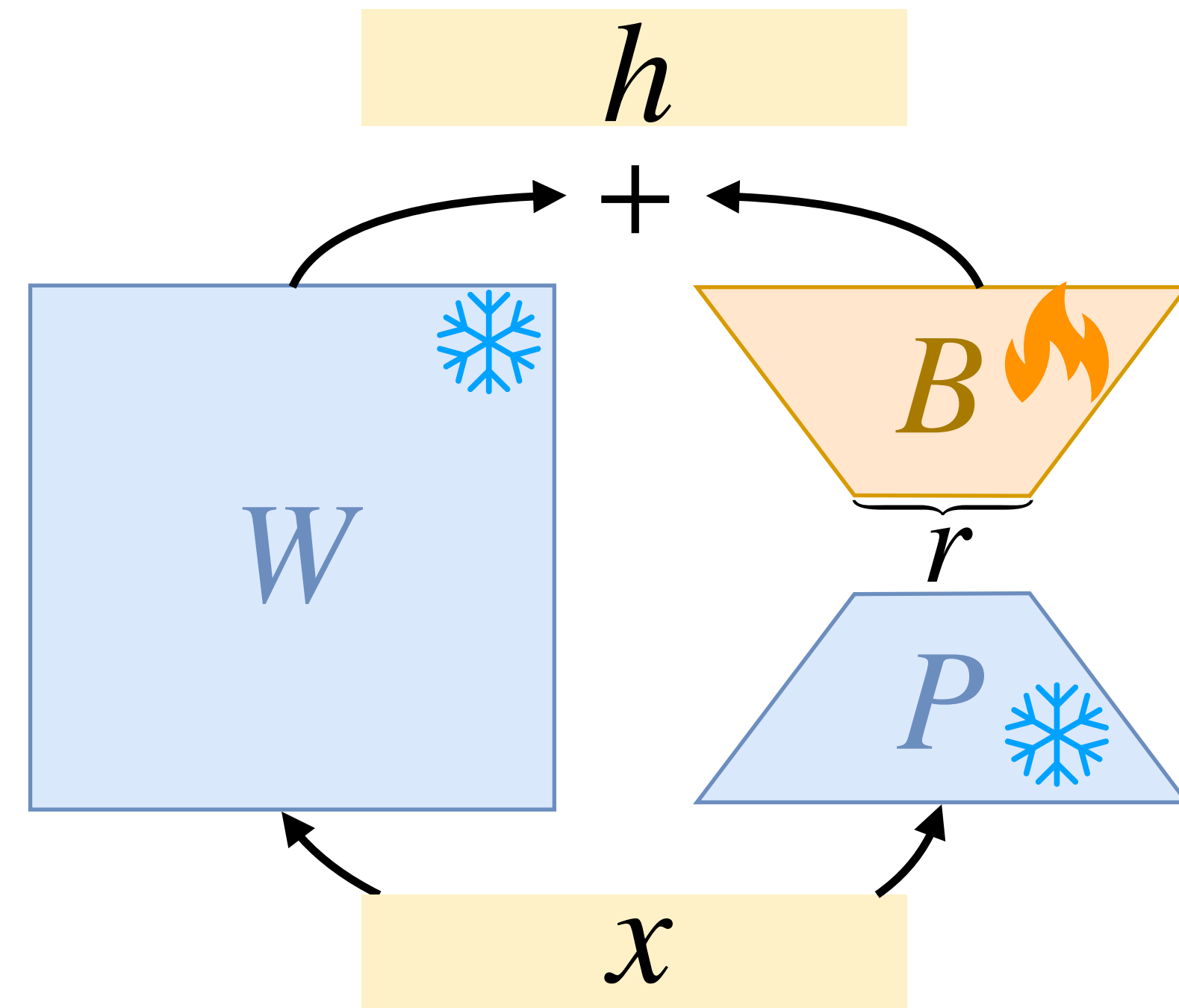
LoQT: Low-Rank Quantized Pre-Training



🔥 Optimize ❄️ Freeze

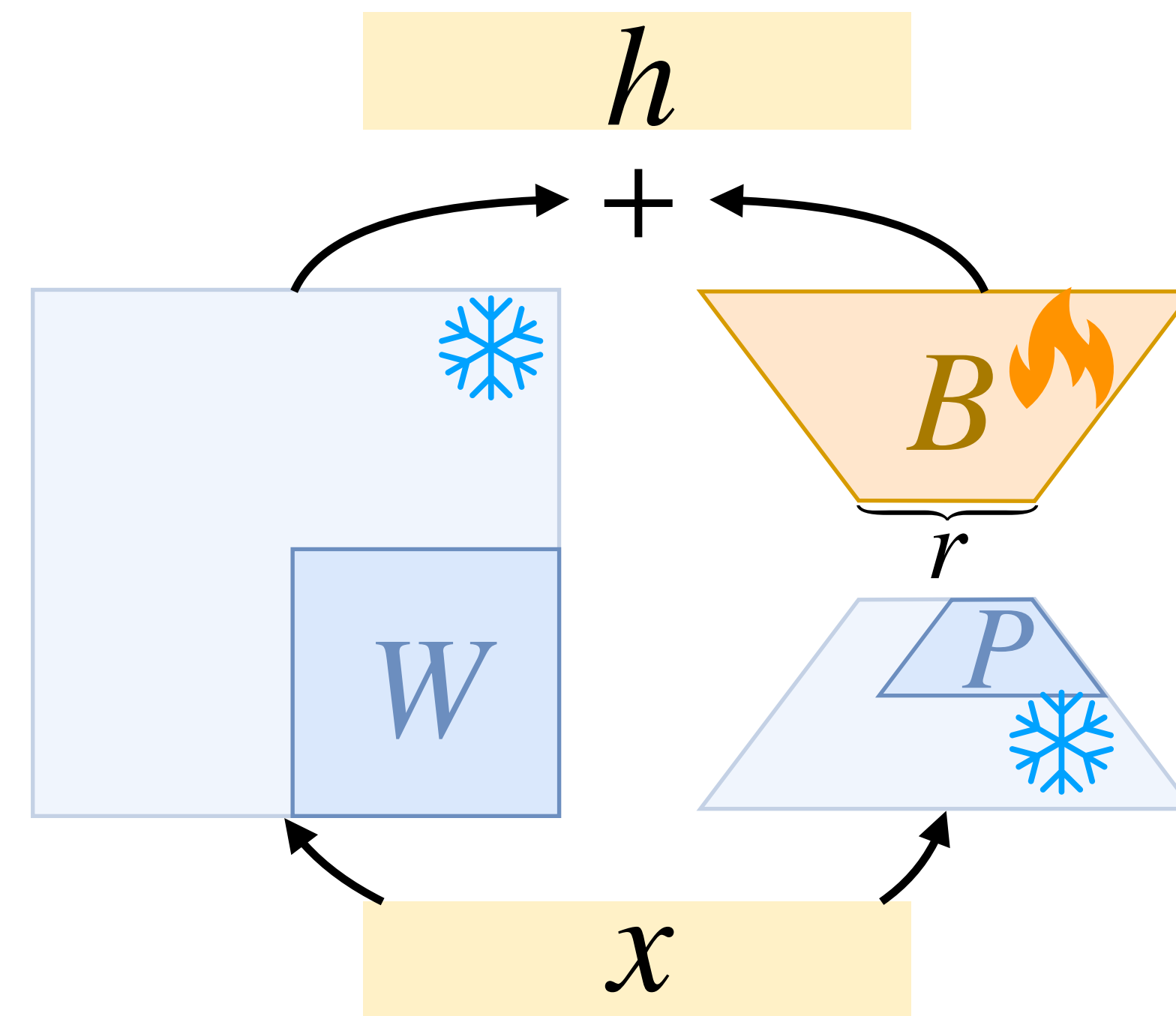
LoQT: Low-Rank Quantized Pre-Training

- **Pre-training from scratch** while only *optimizing a **single low-rank factor** per layer.*



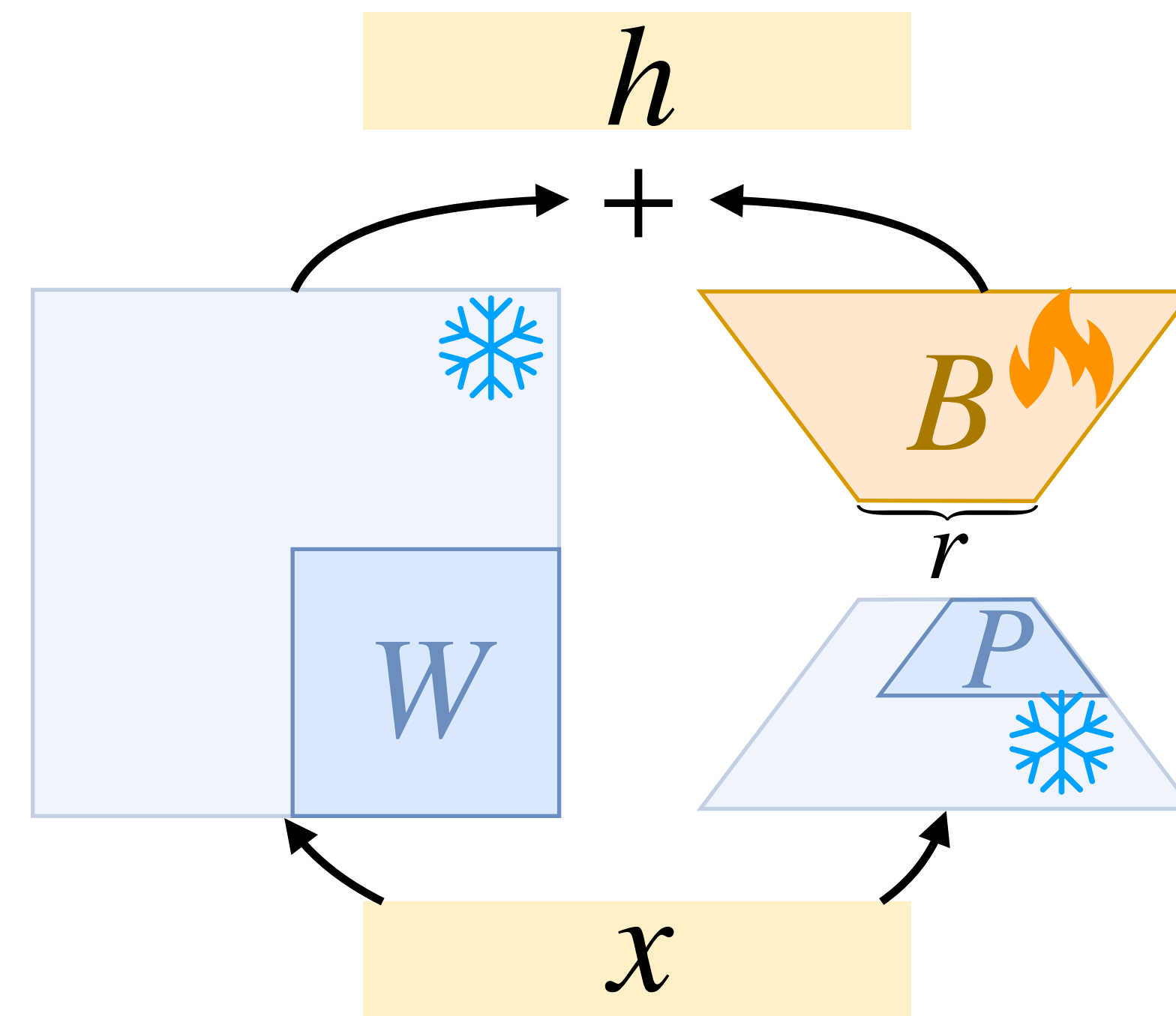
LoQT: Low-Rank Quantized Pre-Training

- **Pre-training from scratch** while only *optimizing a **single low-rank factor** per layer.*
- Keep most *weights in **4-bit precision.***



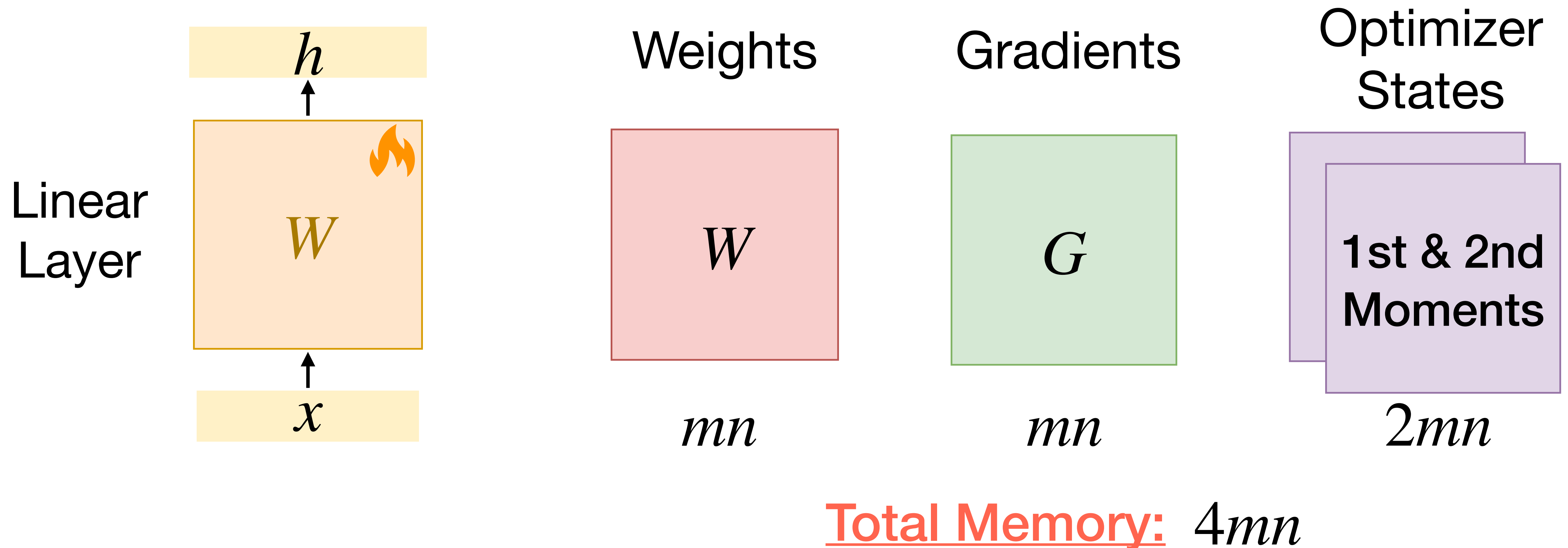
LoQT: Low-Rank Quantized Pre-Training

- **Pre-training from scratch** while only *optimizing a **single low-rank factor per layer***.
- Keep most *weights in **4-bit precision***.
- Matches regular FP16 training

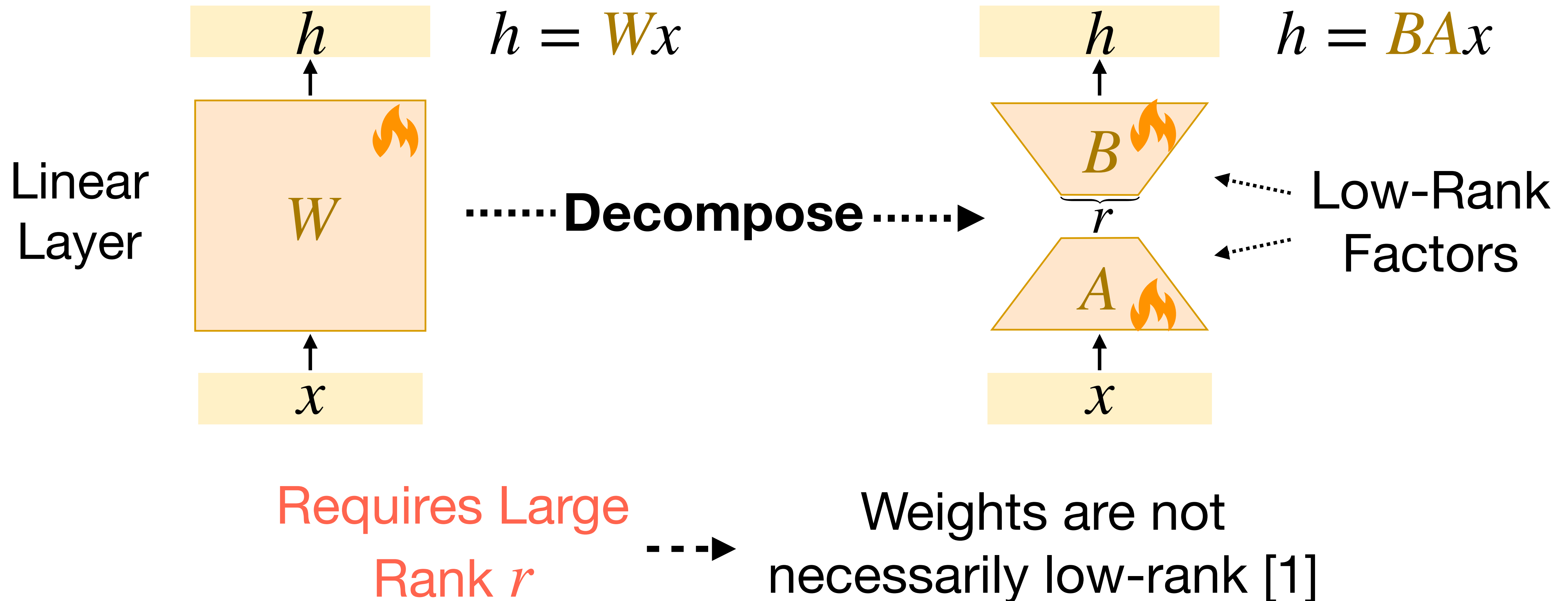


Training with the Adam Optimizer

Memory Requirements



Low-Rank Training



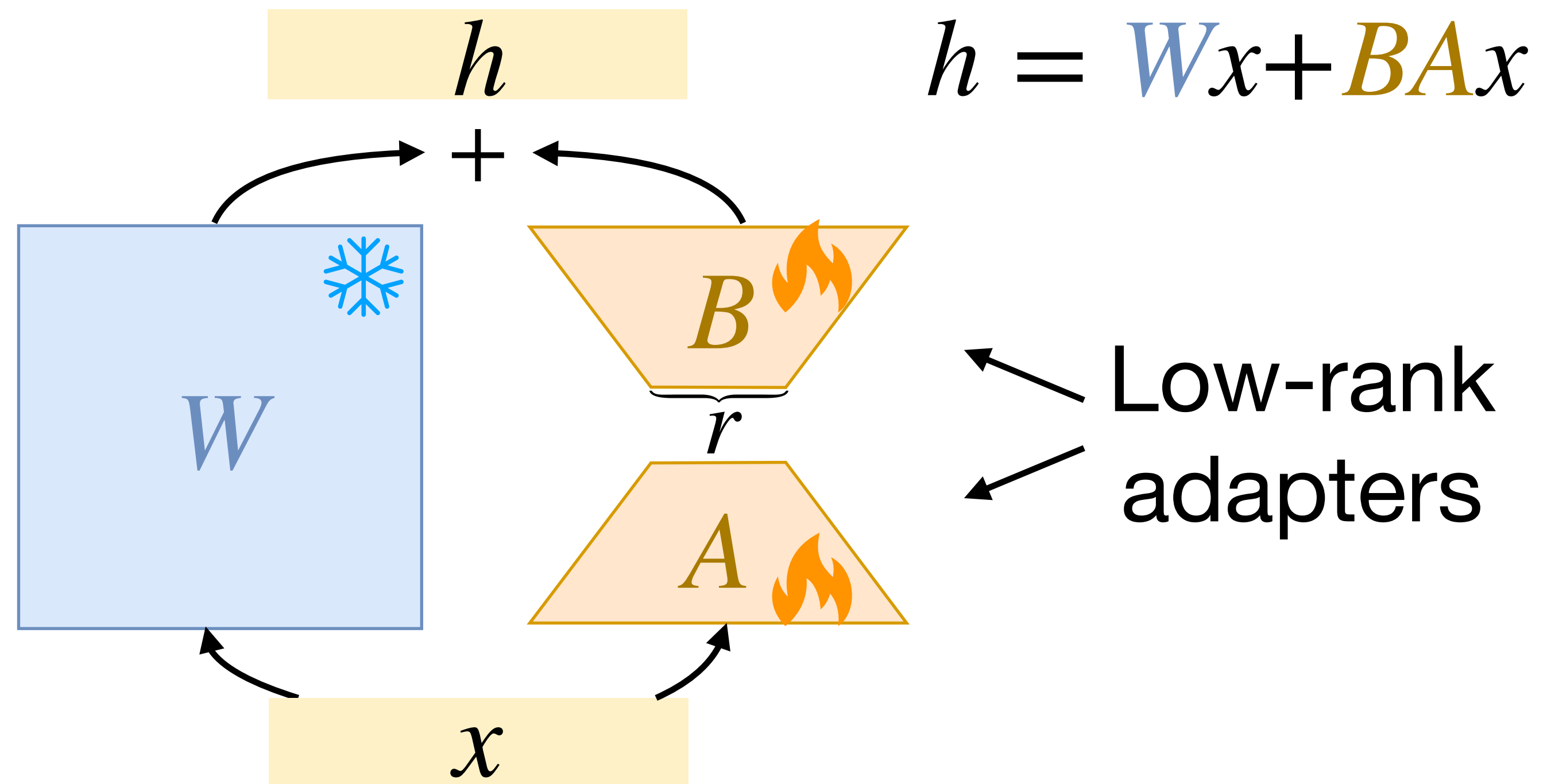
Low-Rank Properties of Gradients

- **GaLore:** “Gradient becomes low-rank during training” [1]
- **InRank:** “Cumulative weight updates follow an incremental low-rank trajectory” [2]

LoRA - Low-Rank Adaption [3]

***Low-Rank
Cumulative Updates***

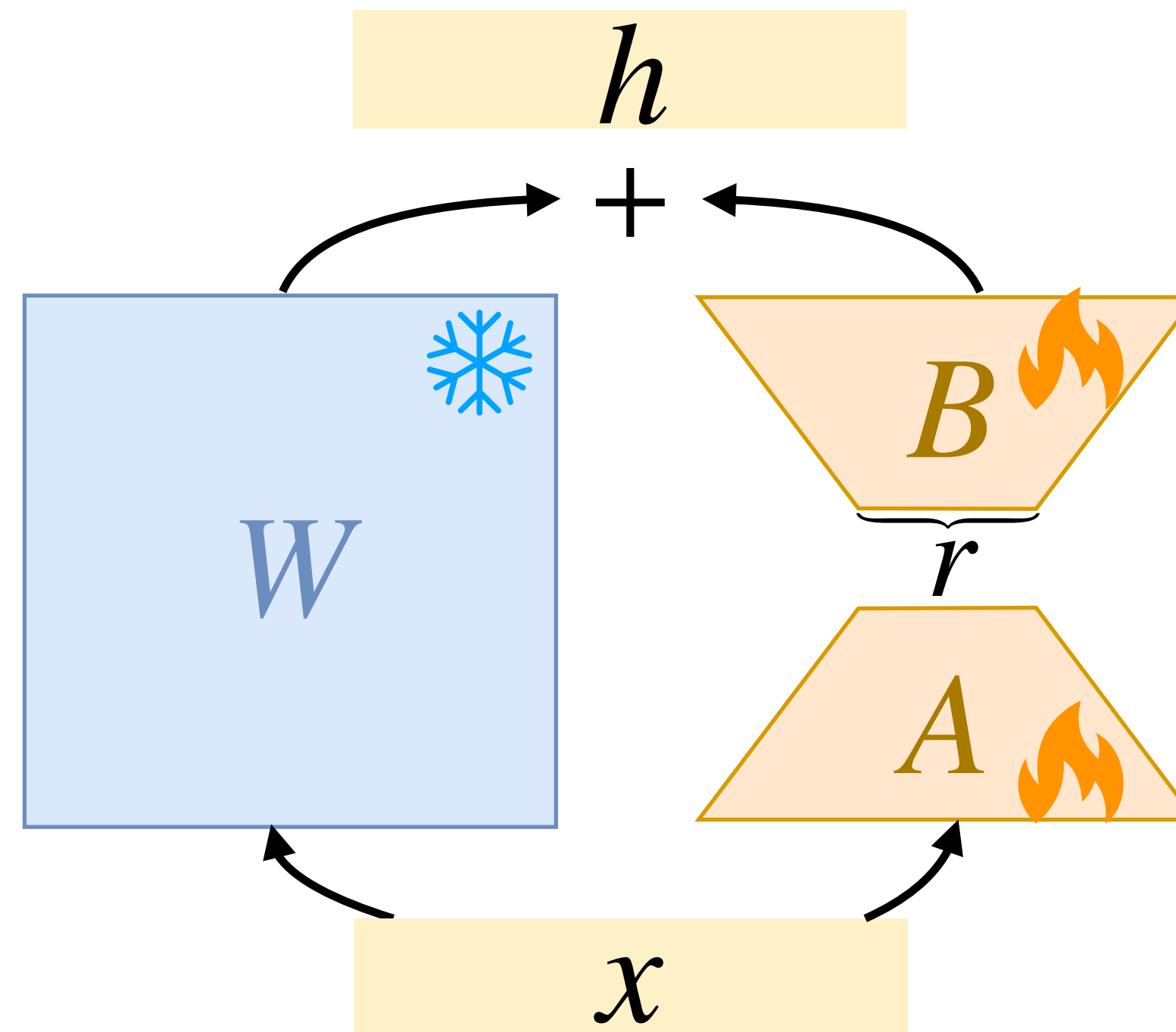
Pre-trained →



LoRA - Low-Rank Adaption [3]

***Low-Rank
Cumulative Updates***

Pre-trained →



$$h = Wx + BAx$$

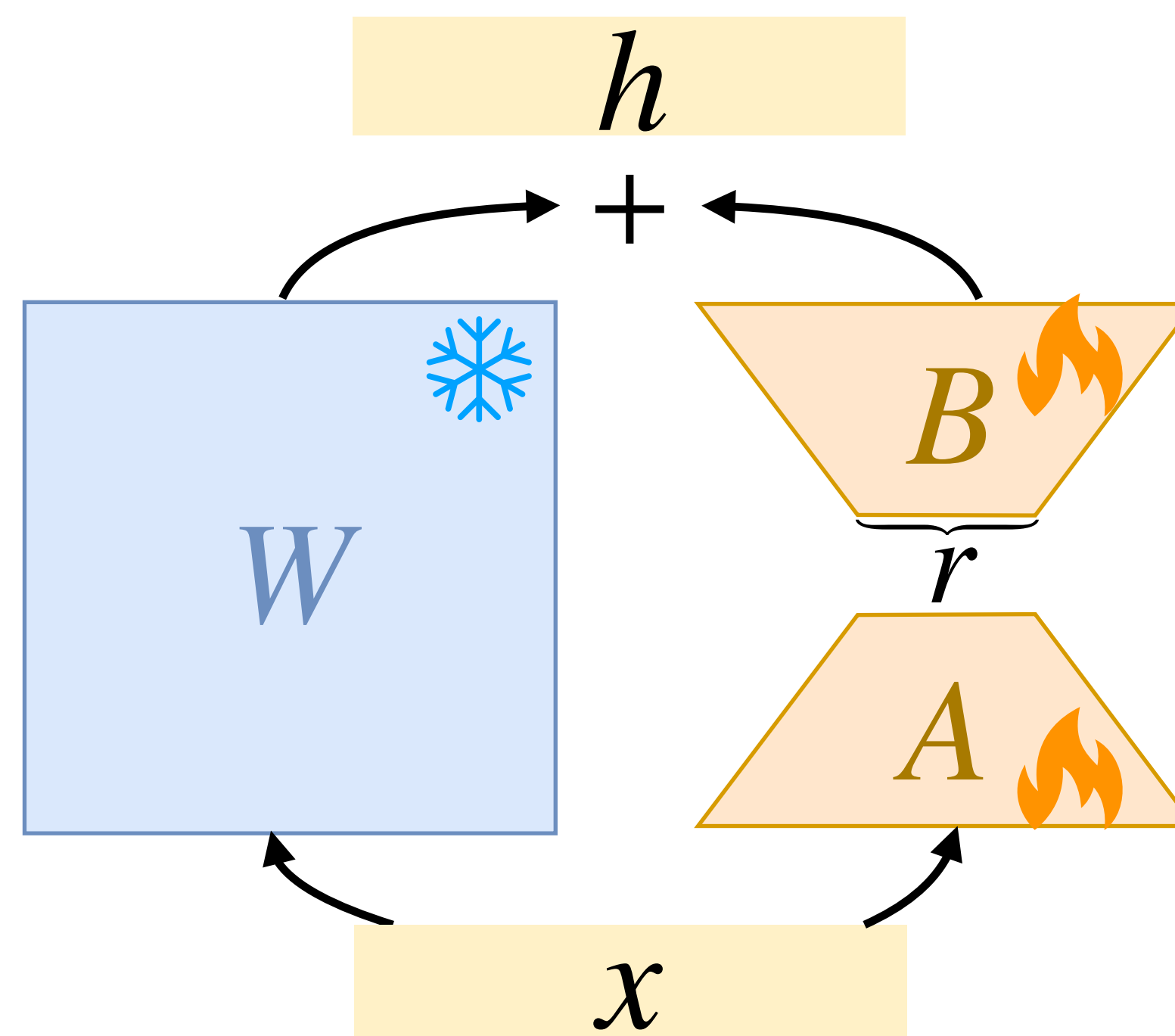
Low-rank
adapters

+ Reduces memory and is
effective for fine-tuning

LoRA - Low-Rank Adaption [3]

**Low-Rank
Cumulative Updates**

Pre-trained →



$$h = Wx + BAx$$

Low-rank
adapters

+ Reduces memory and is effective for fine-tuning

— Does not work well for pre-training

LoQT: Low-Rank Quantized Pre-Training

Key Components

- Pre-train from scratch with low-rank adapters (LoRA)

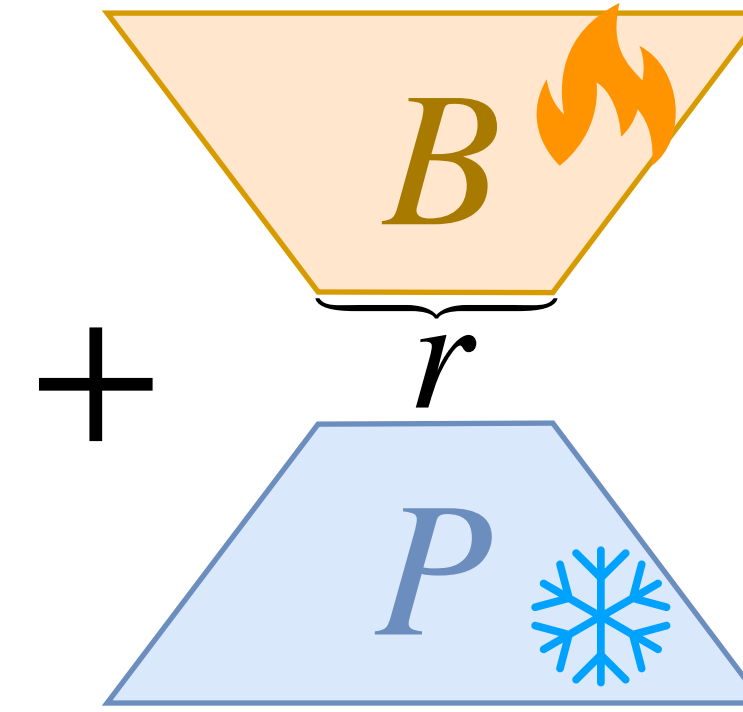
$$\hat{W} = W + PB$$

LoQT: Low-Rank Quantized Pre-Training

Key Components

- Pre-train from scratch with low-rank adapters (LoRA)

Randomly
initialized W



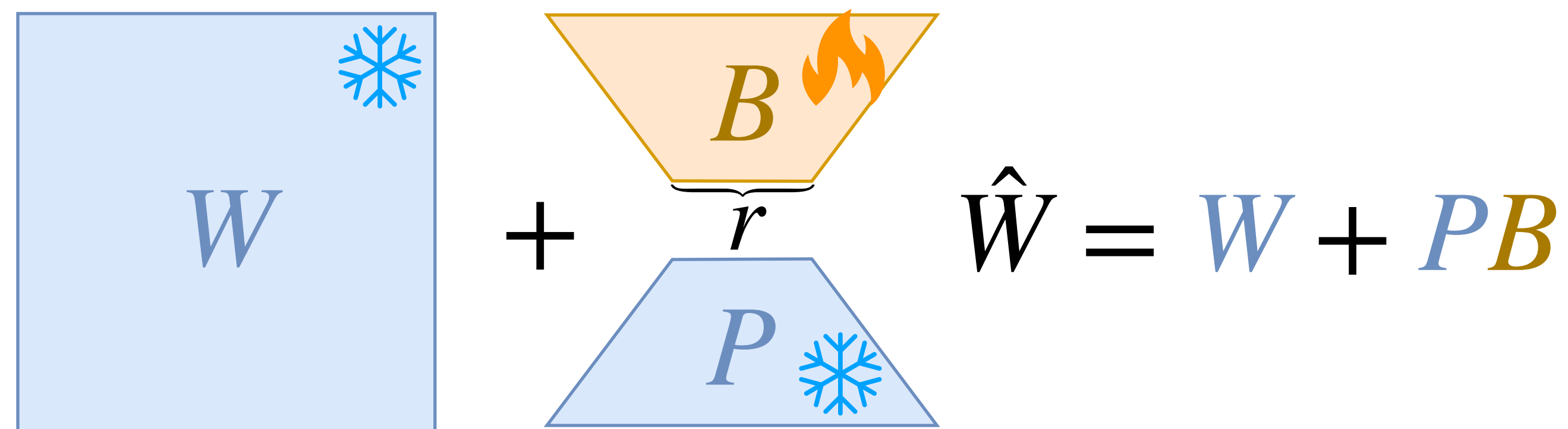
Only optimize a
single adapter

$$\hat{W} = W + PB$$

LoQT: Low-Rank Quantized Pre-Training

Key Components

- Pre-train from scratch with low-rank adapters (LoRA)
- Initialize P using ∇W



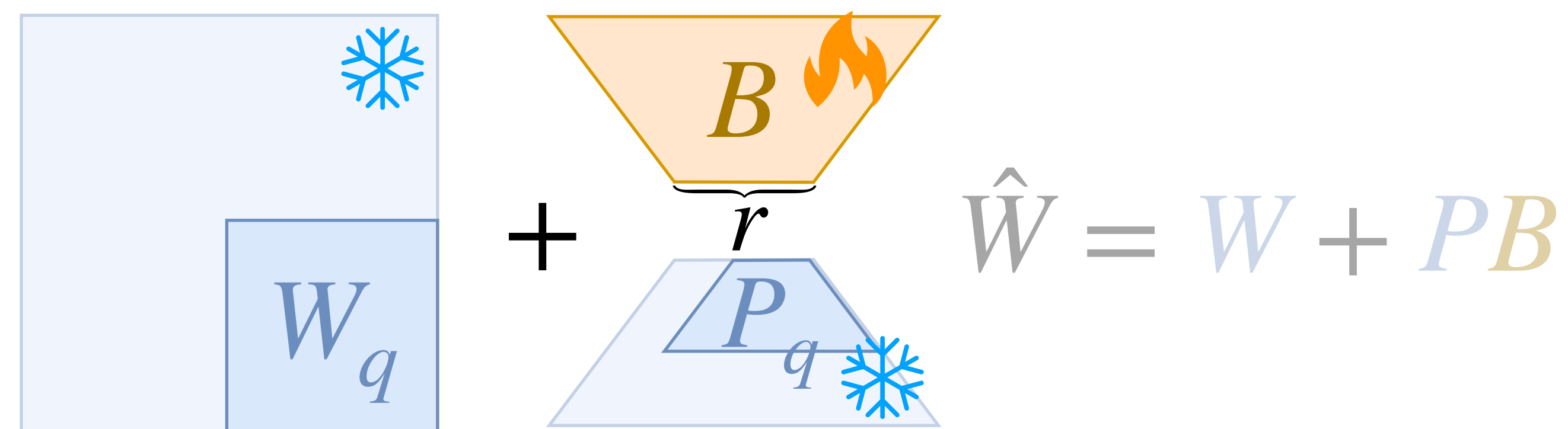
$$\text{SVD}(\nabla W) = U\Sigma V^{\top}$$

$$P = U[:, r]$$

LoQT: Low-Rank Quantized Pre-Training

Key Components

- Pre-train from scratch with low-rank adapters (LoRA)
- Initialize P using ∇W
- Quantize P and W to 4-bit



$$\text{SVD}(\nabla W) = U \Sigma V^T$$

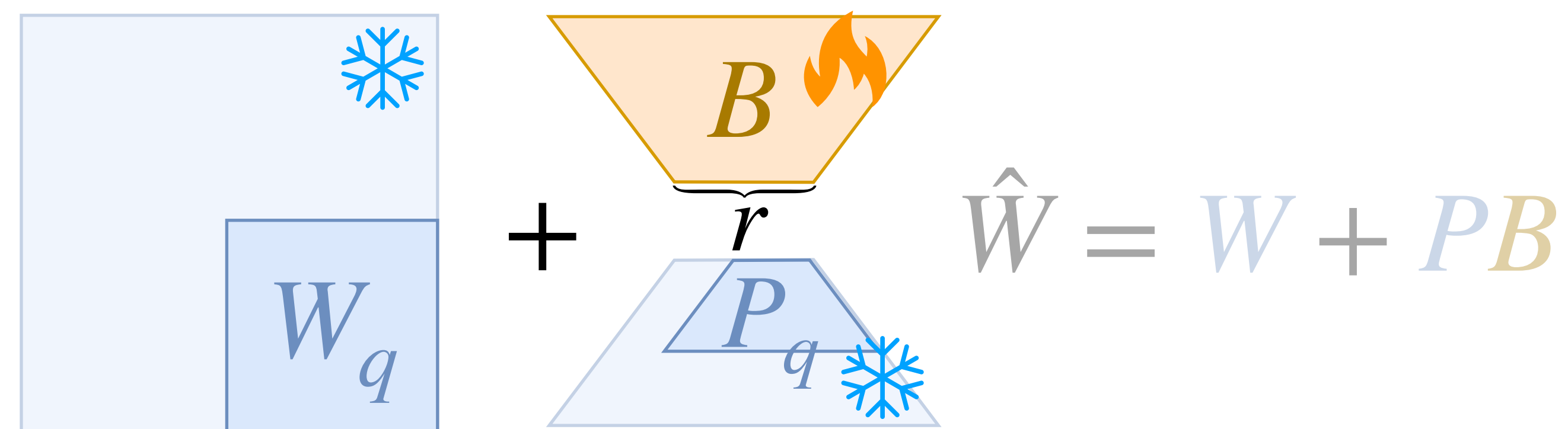
$$P = U[:, r]$$

LoQT: Low-Rank Quantized Pre-Training

Key Components

- Pre-train from scratch with low-rank adapters (LoRA)
- Initialize P using ∇W
- Quantize P and W to 4-bit
- Initialize B using a *quantization error compensation* term.

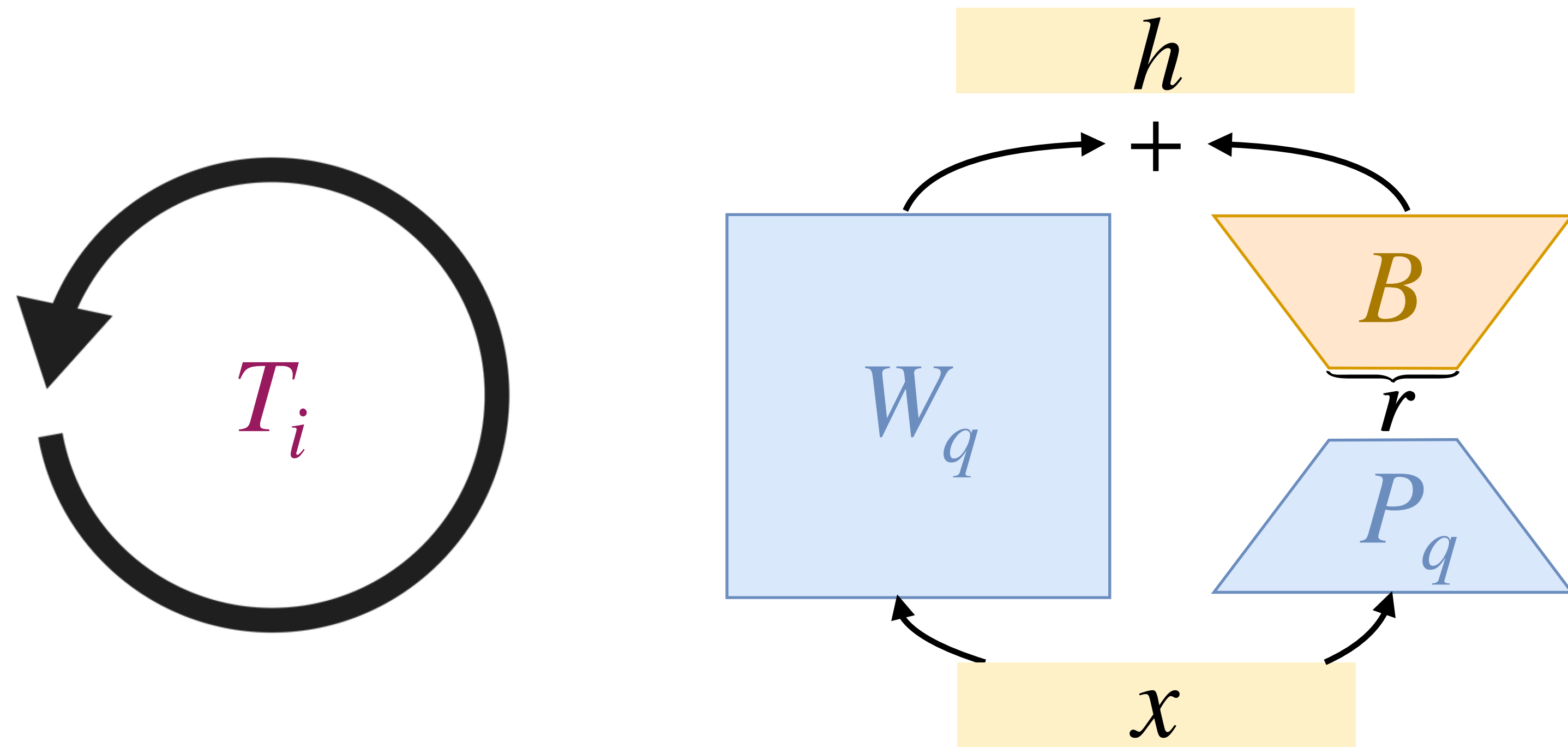
$$B = P^{-1}(W_q - W)$$



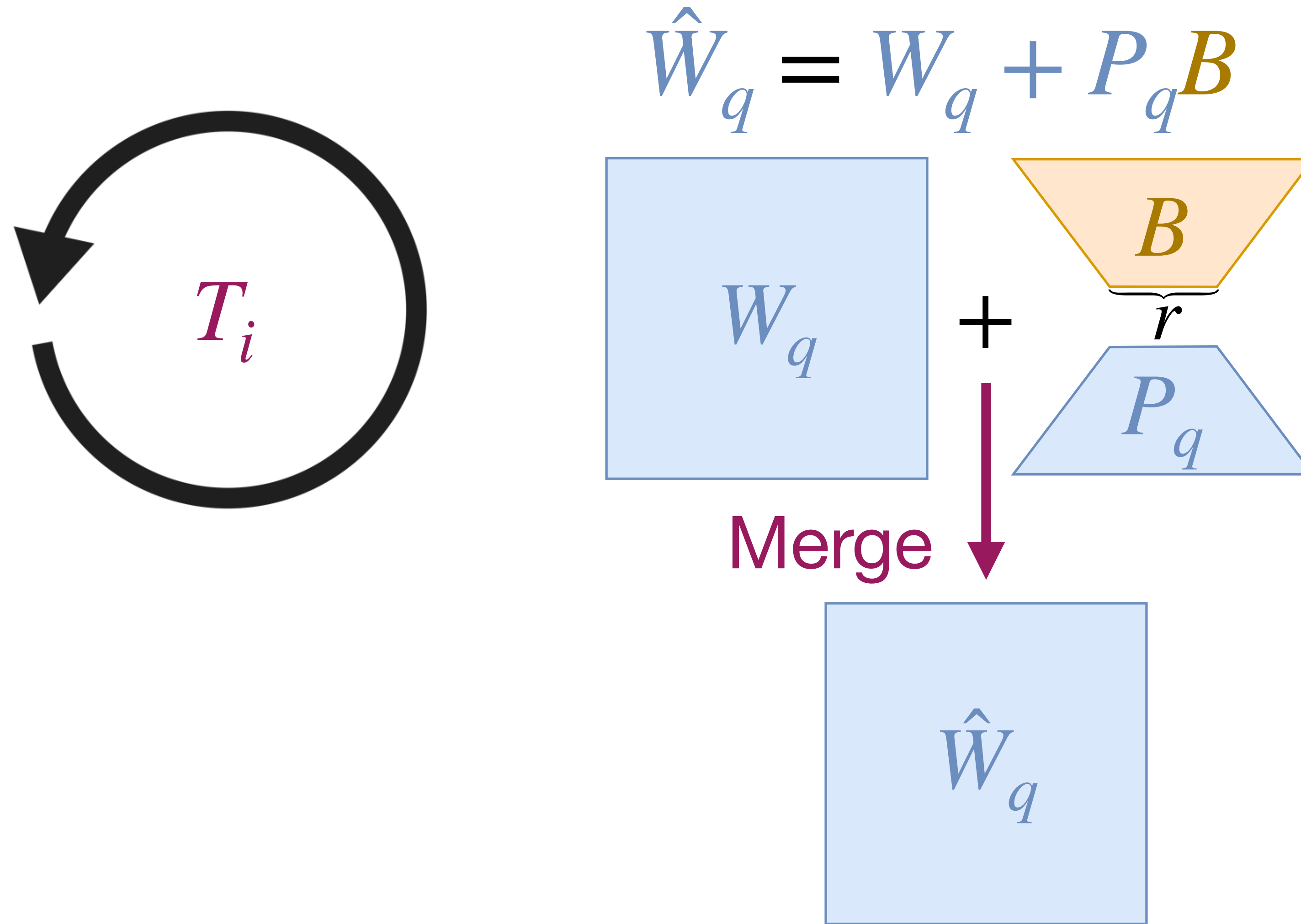
$$\text{SVD}(\nabla W) = U \Sigma V^T$$

$$P = U[:, r]$$

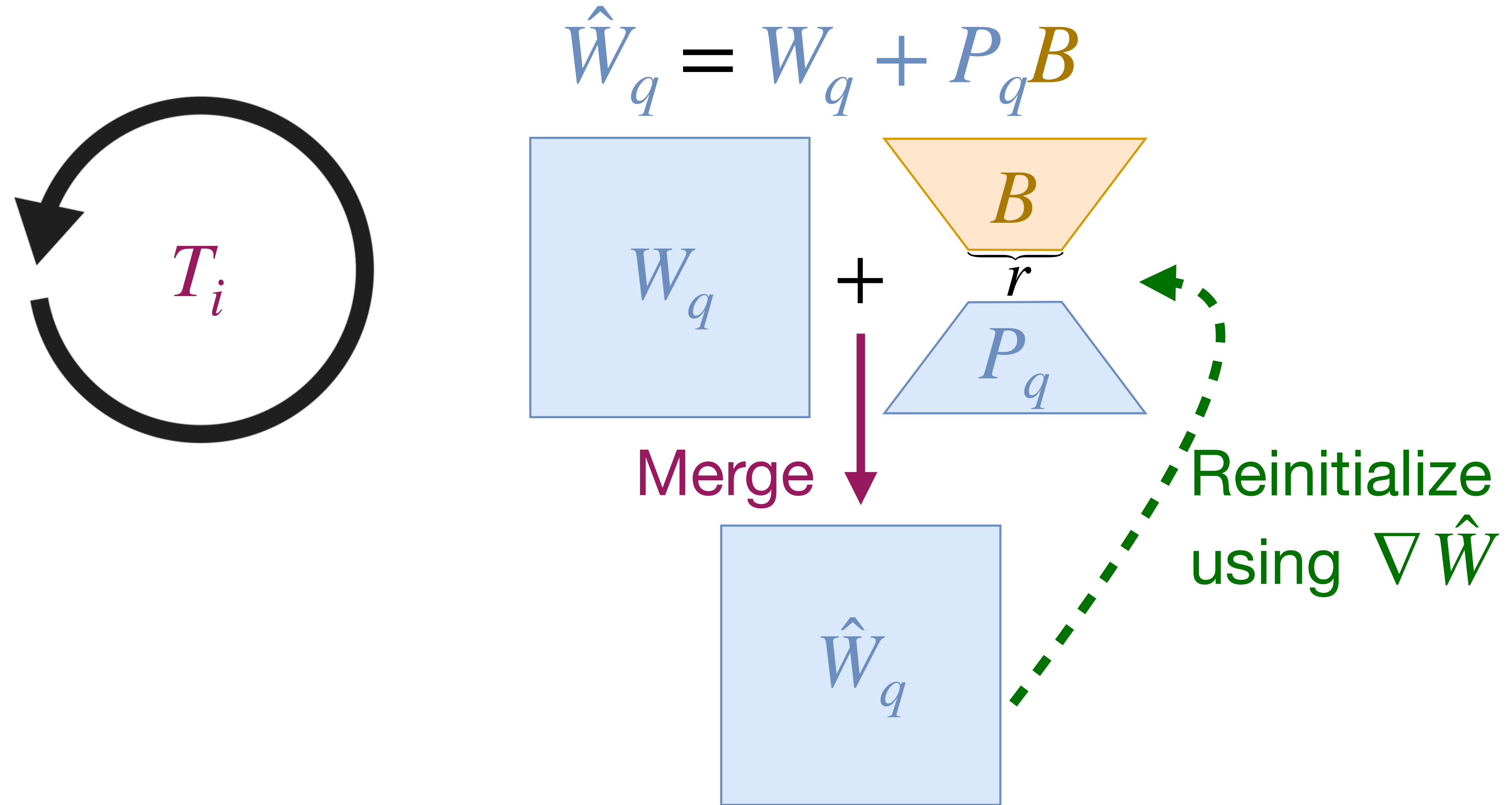
LoQT - Training



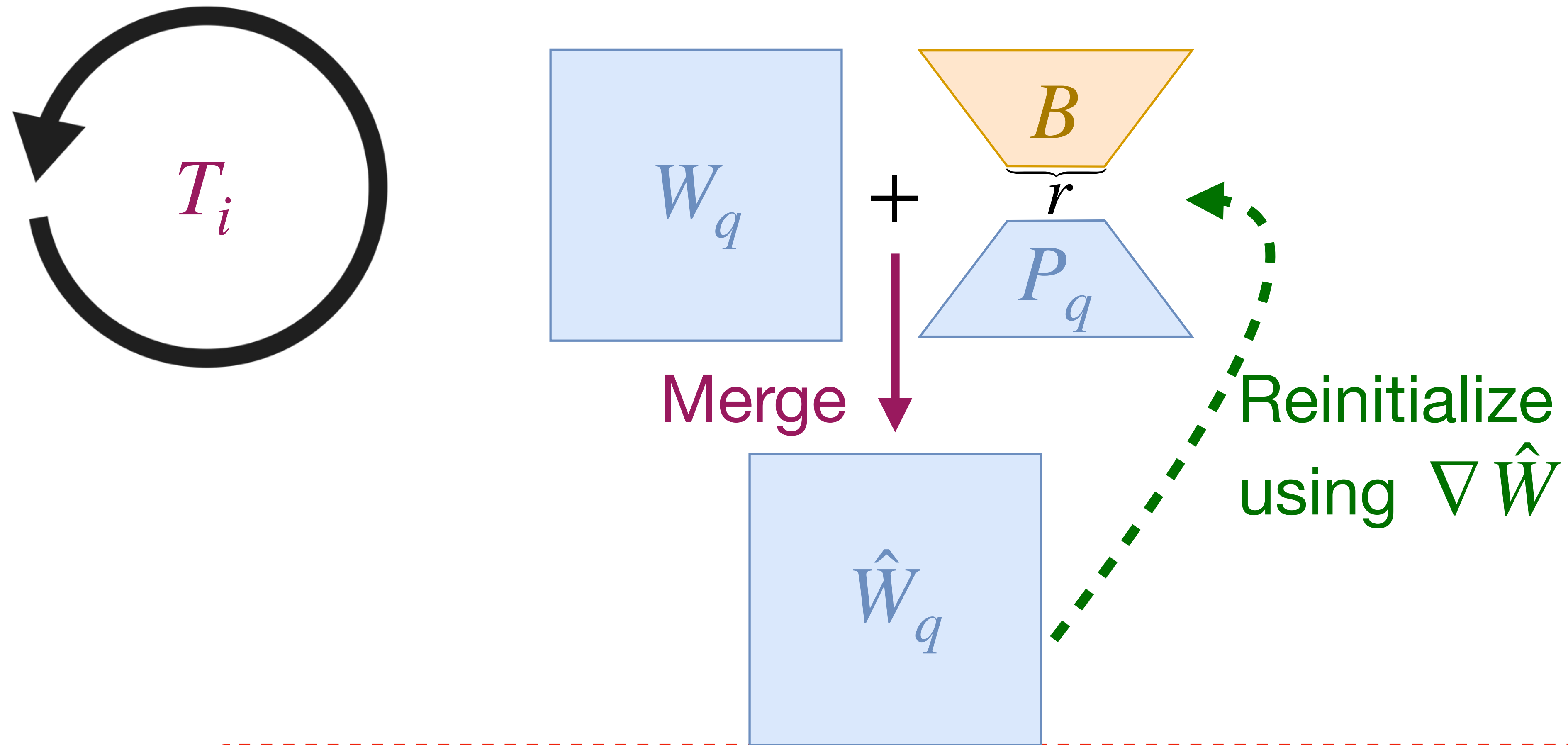
LoQT - Merge and Reinitialize



LoQT - Merge and Reinitialize

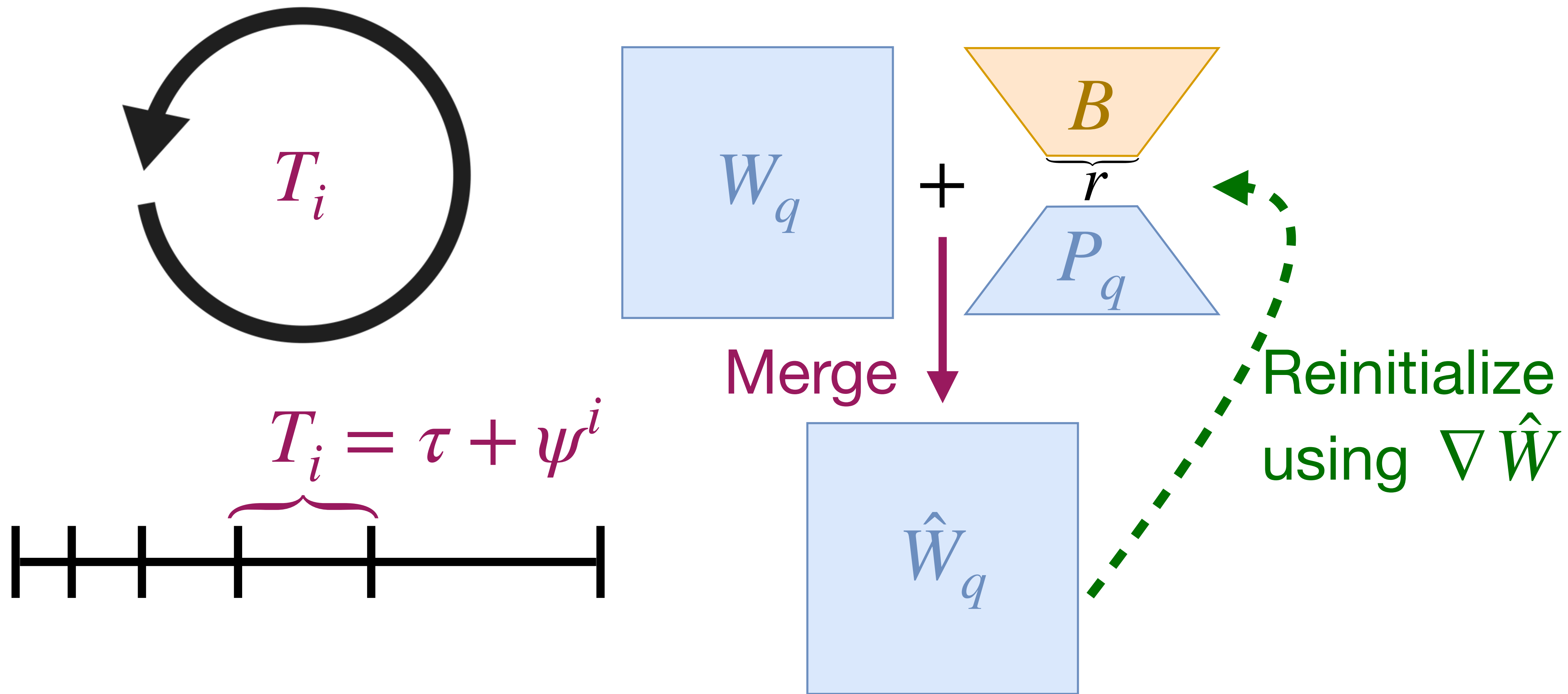


LoQT - Merge and Reinitialize



$$W_N = W_0 + P_0 B_0 + P_1 B_1 + \dots + P_N B_N$$

LoQT - Merge and Reinitialize



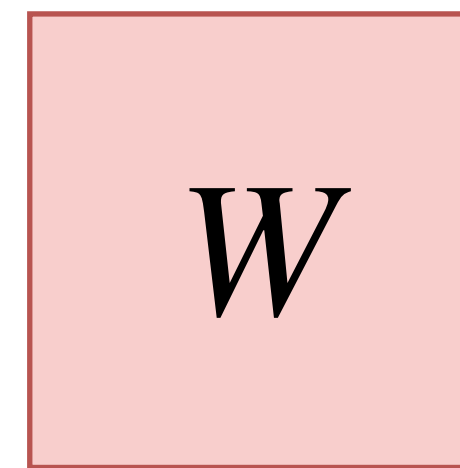
$$W_N = W_0 + P_0 B_0 + P_1 B_1 + \dots + P_N B_N$$

Memory Requirements with Adam

Regular Training

$$h = Wx$$

Weights



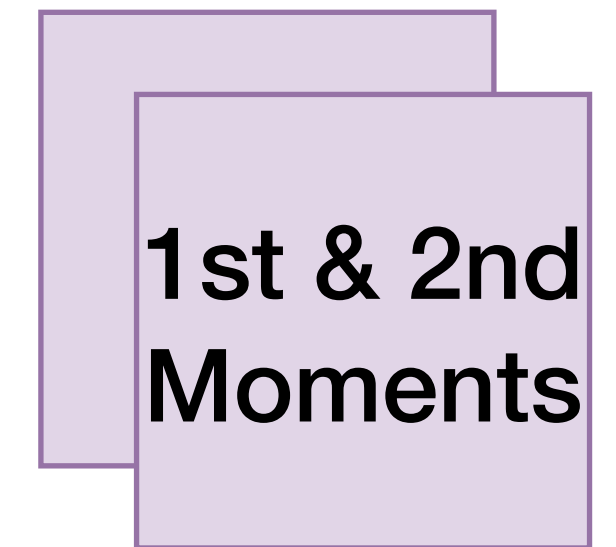
mn

Gradients



mn

Optimizer States



1st & 2nd Moments

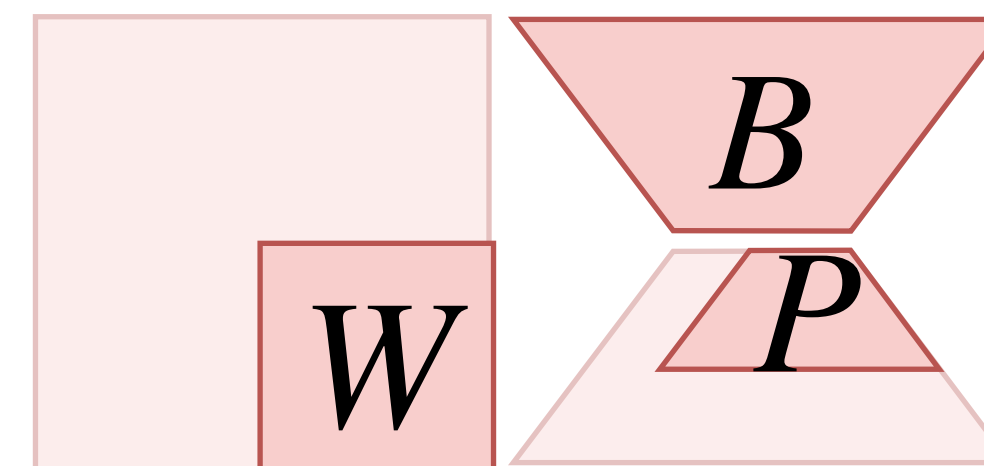
$2mn$

LoQT

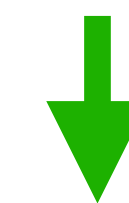
$$h = Wx + PBx$$

Good

if $r \ll \min(m, n)$

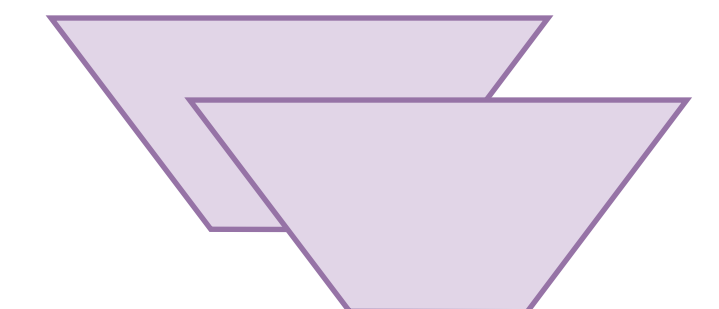
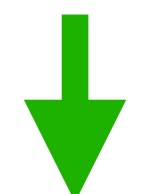


$m_q n_q + nr + m_q r$



nr

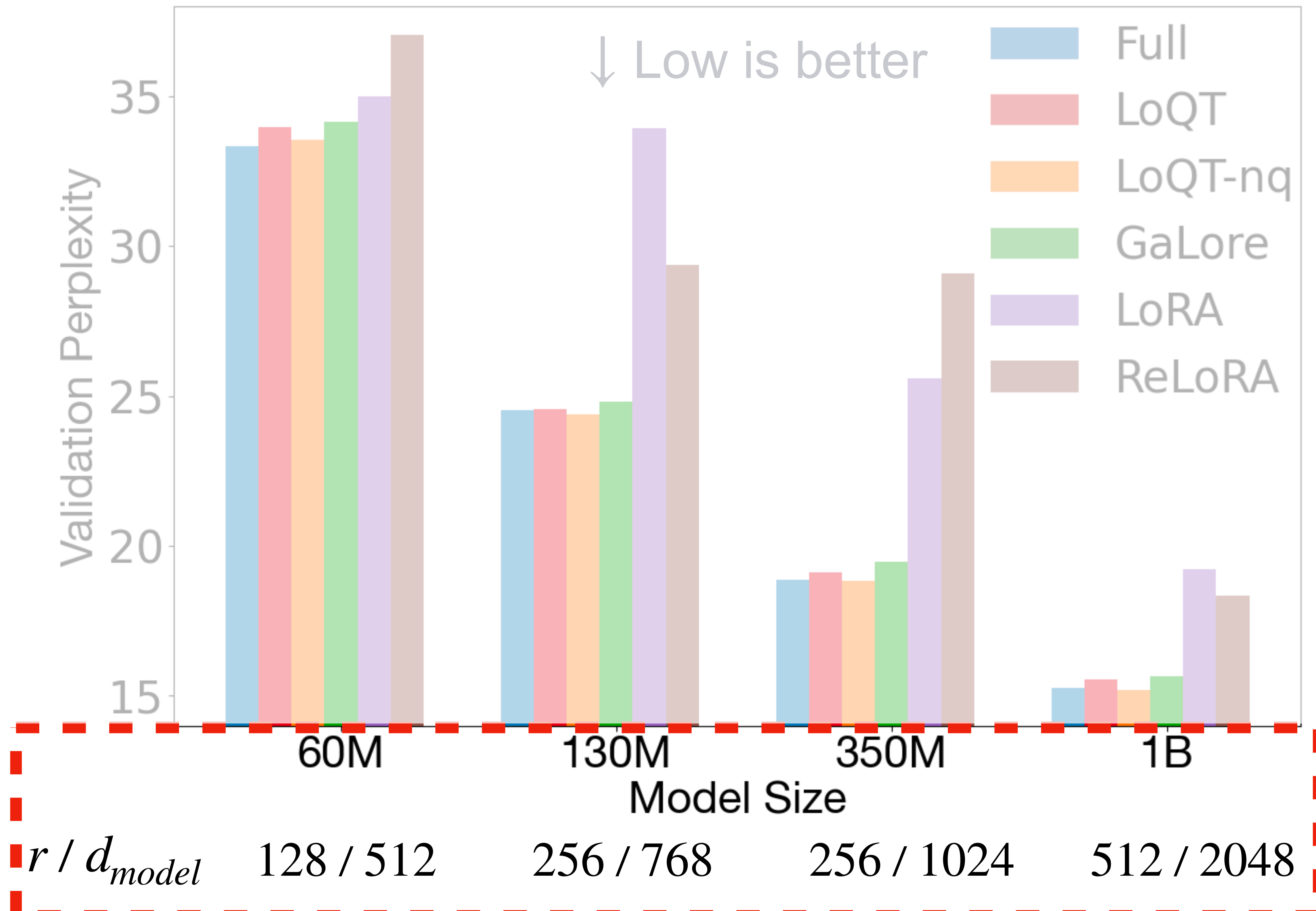
$\frac{m}{r}$



$2nr$

Results

Pre-training LLaMA2 Models on the C4 Dataset



Pre-training LLaMA2 Models on the C4 Dataset

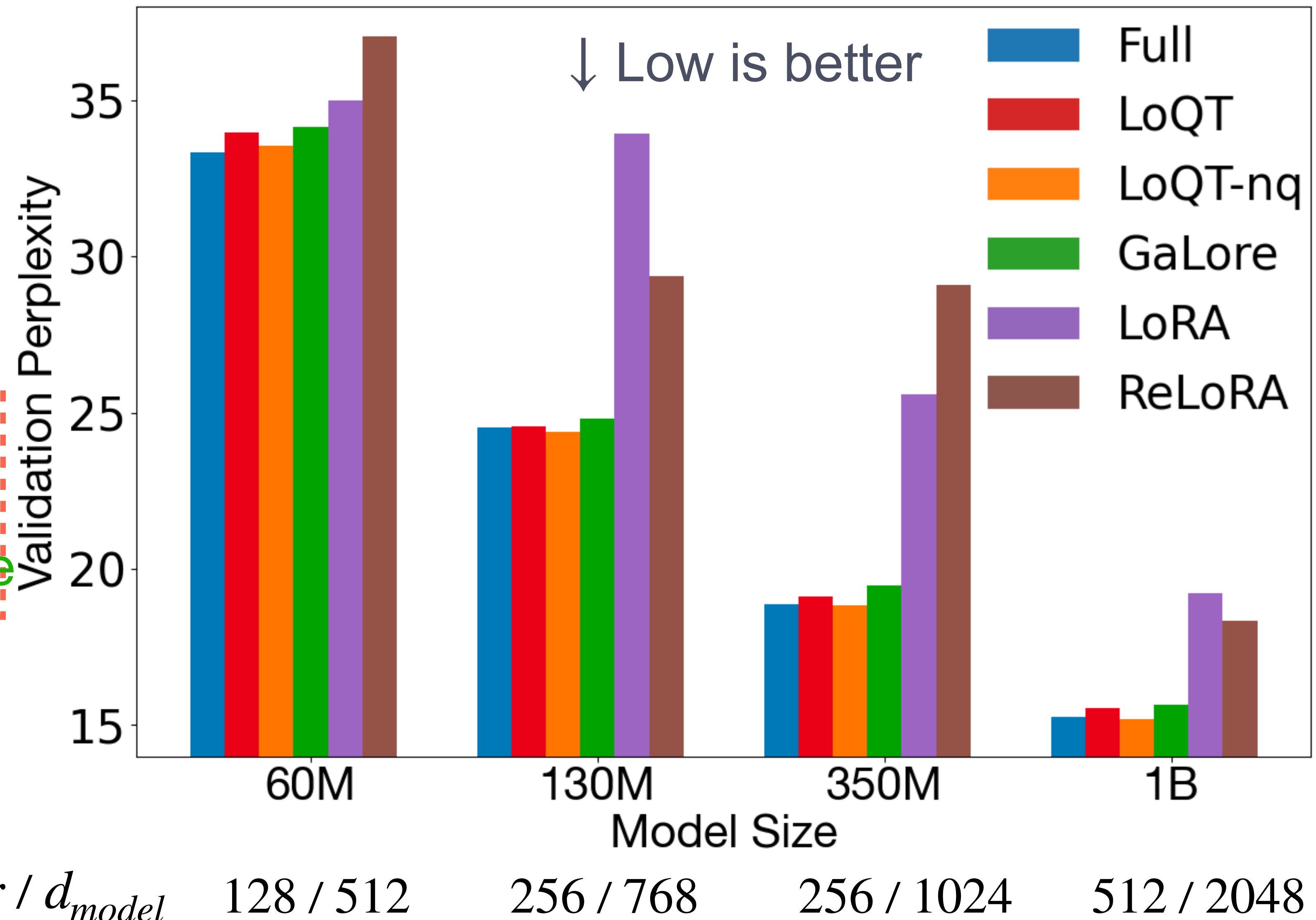
Takeaways

- **LoQT** is very comparable to **Full** training and **GaLore**

LoQT - 1B model

61% less memory than **Full**

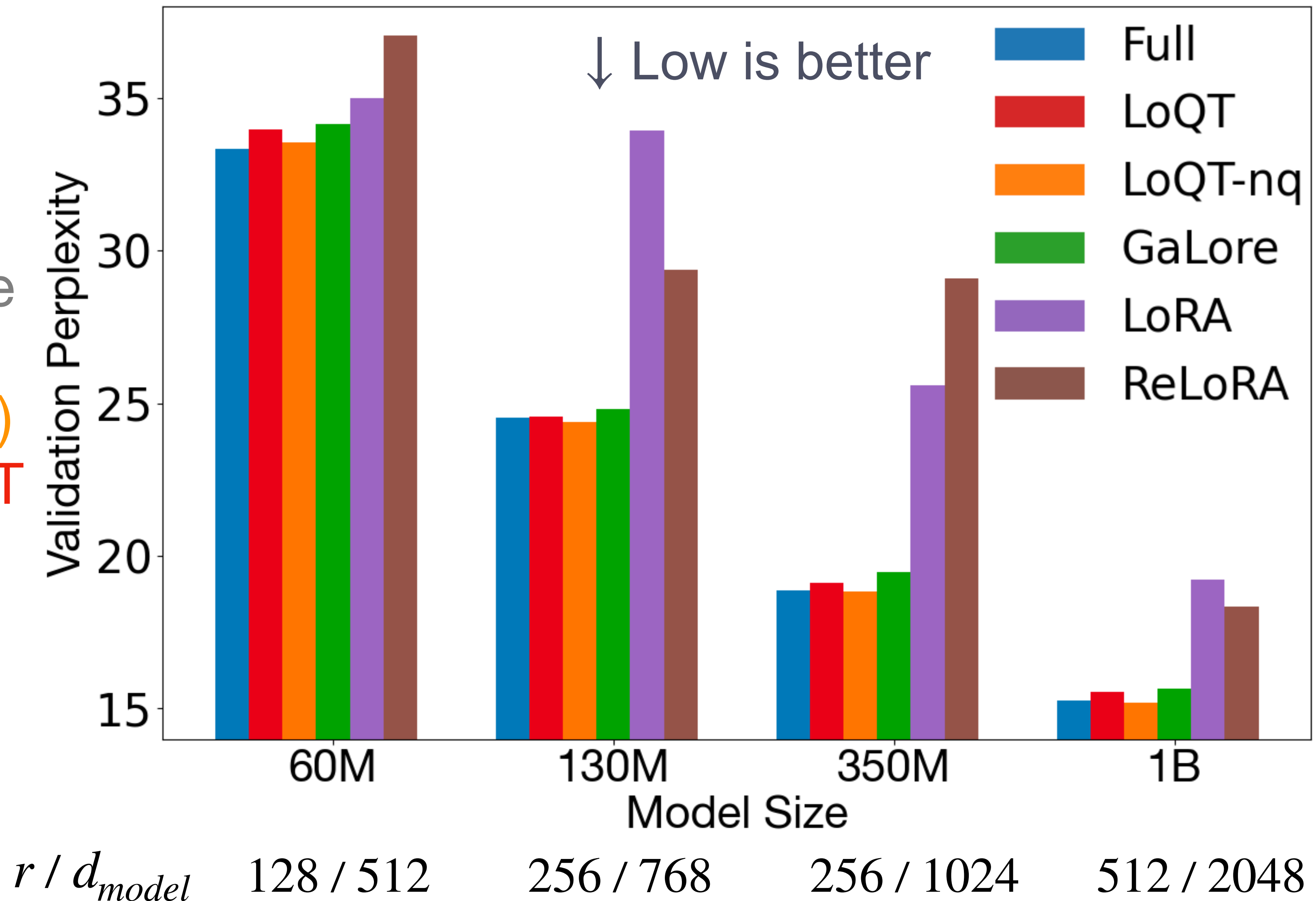
41% less memory than **GaLore**



Pre-training LLaMA2 Models on the C4 Dataset

Takeaways

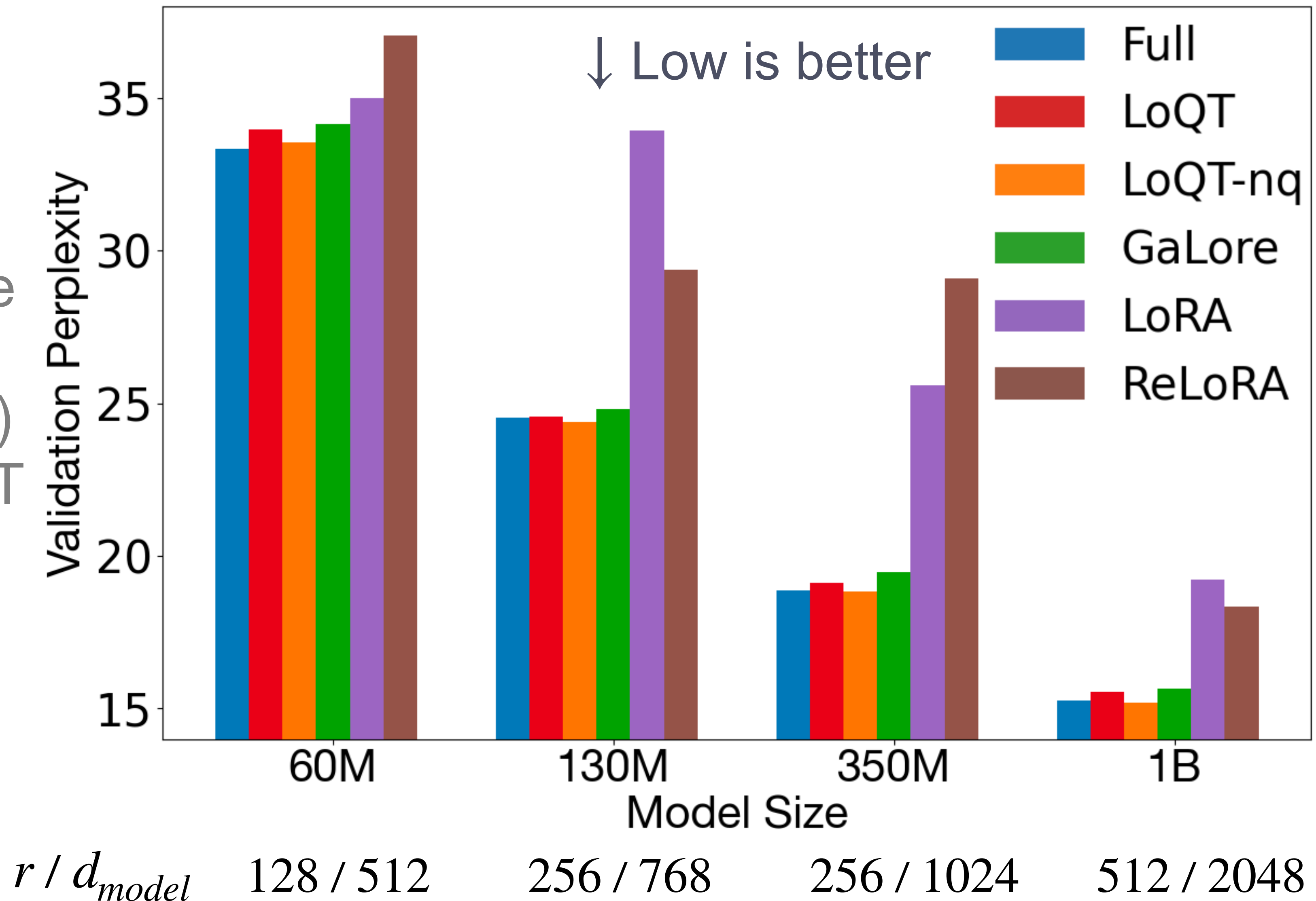
- LoQT is very comparable to Full training and GaLore
- No-quantization (LoQT-nq) is slightly better than LoQT



Pre-training LLaMA2 Models on the C4 Dataset

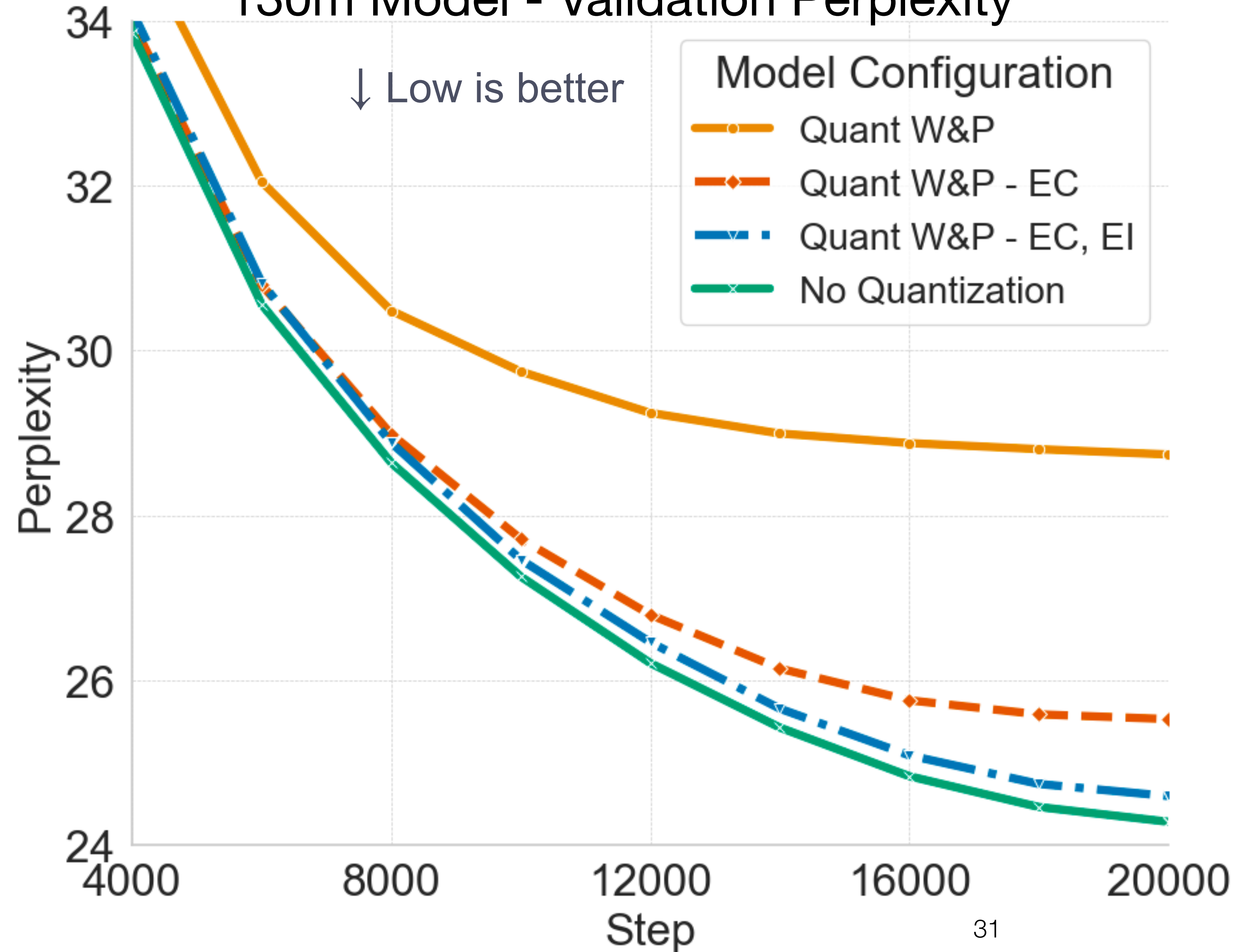
Takeaways

- LoQT is very comparable to Full training and GaLore
- No-quantization (LoQT-nq) is slightly better than LoQT
- **LoQT** is much better than **LoRA** and **ReLoRA**



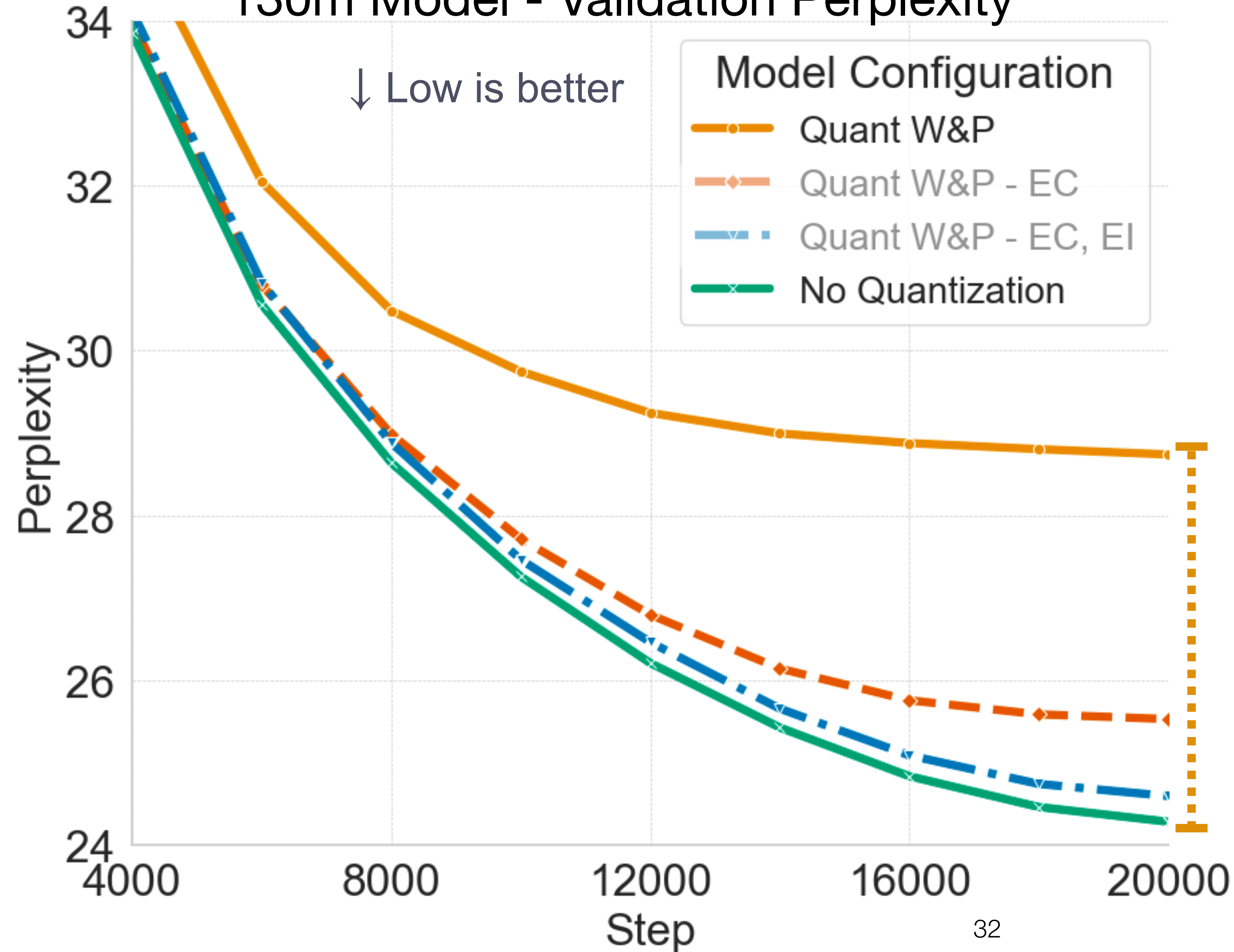
Component Ablations

130m Model - Validation Perplexity



Component Ablations

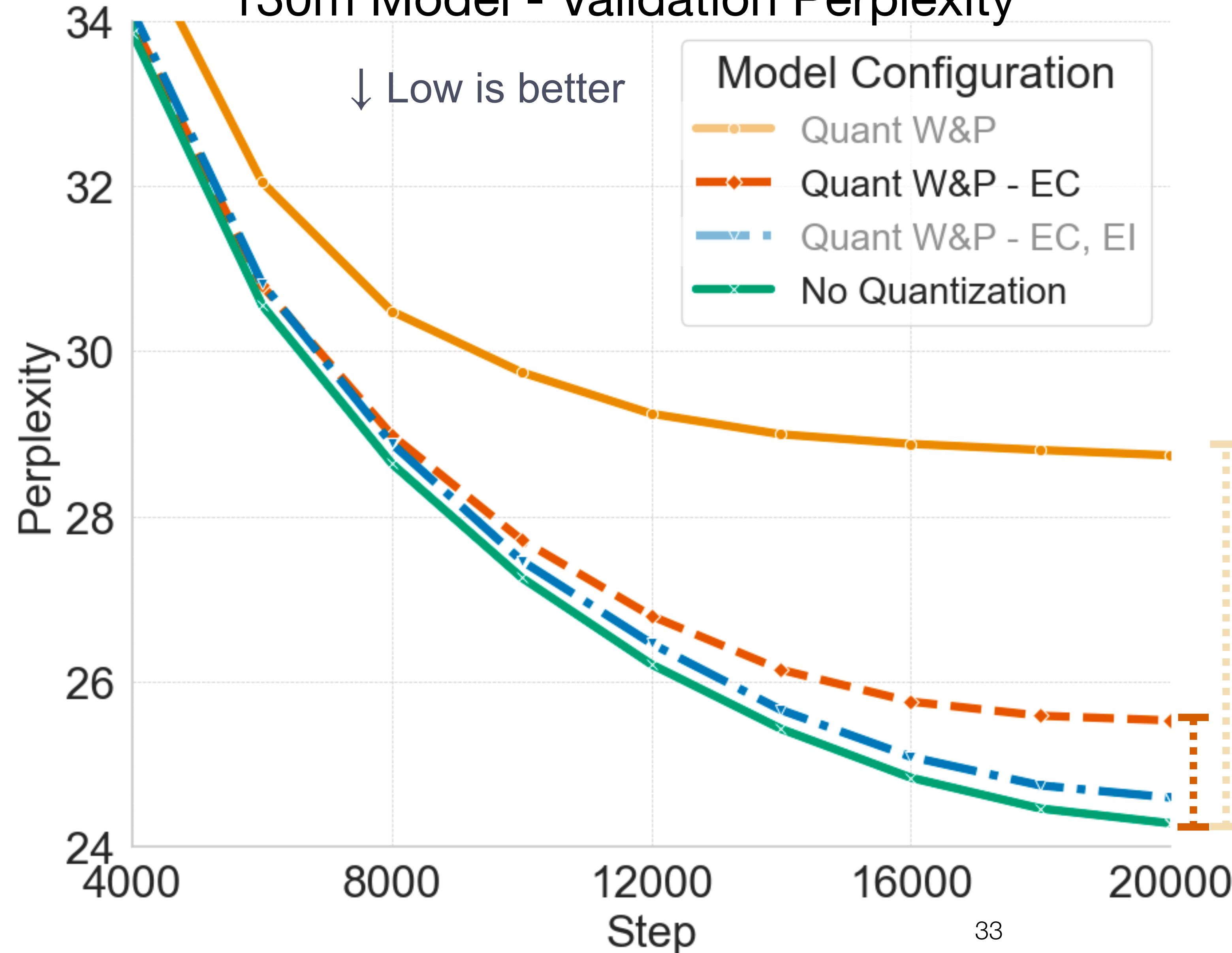
130m Model - Validation Perplexity



- **Simple 4-bit Quantization of W and P :** Flattens early.

Component Ablations

130m Model - Validation Perplexity

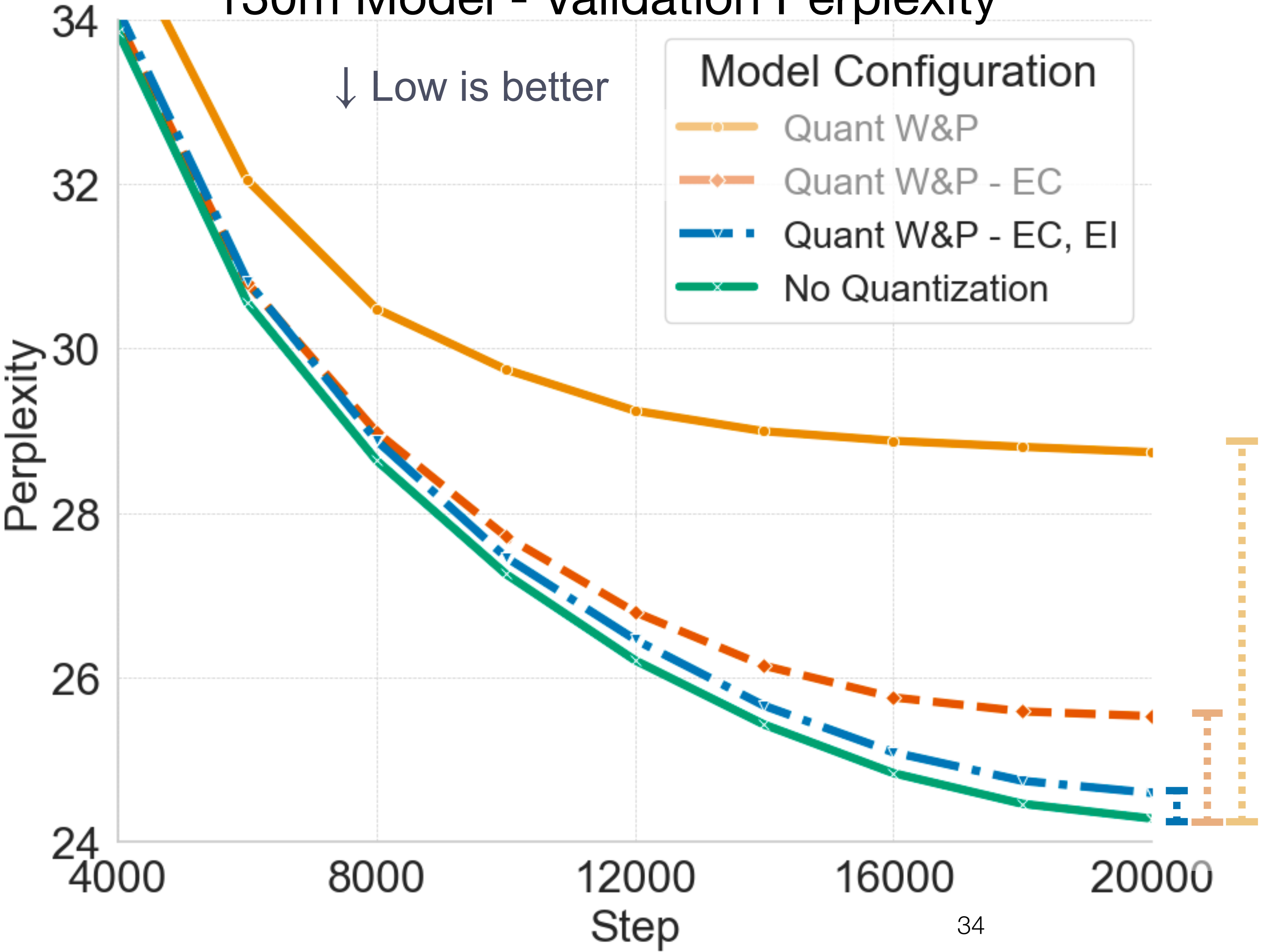


$$B = P^{-1}(W_q - W)$$

- **Simple 4-bit Quantization of W and P :** Flattens early.
- **Error Compensation (EC):** This gap is reduced by 72%.

Component Ablations

130m Model - Validation Perplexity



$$B = P^{-1}(W_q - W)$$

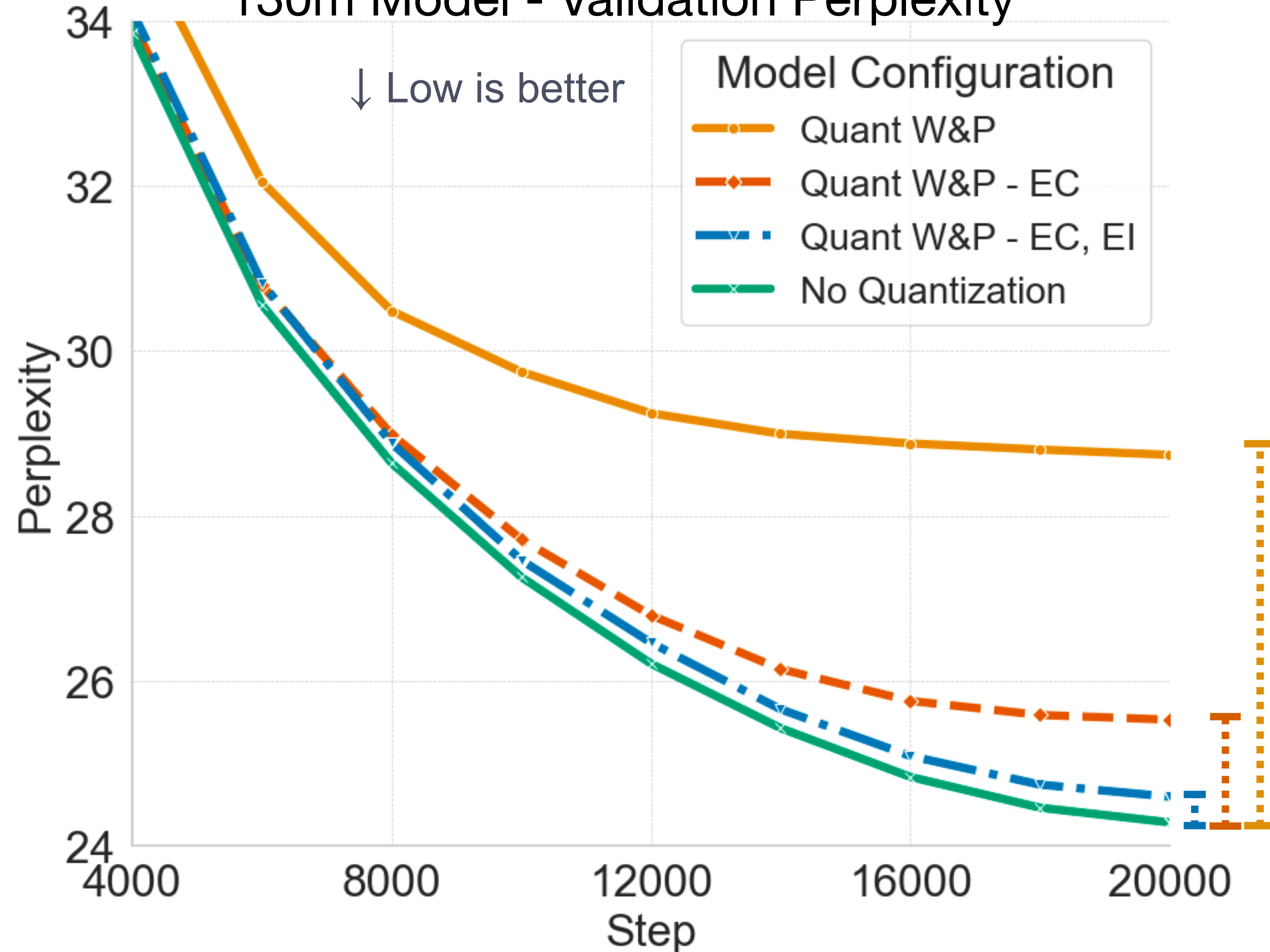
$$T_i = \tau + \psi^i$$



- **Simple 4-bit Quantization of W and P :** Flattens early.
- **Error Compensation (EC):** This gap is reduced by 72%.
- **Exponentially increasing update intervals (EI) and EC:** the gap is reduced by 93%.

Component Ablations

130m Model - Validation Perplexity



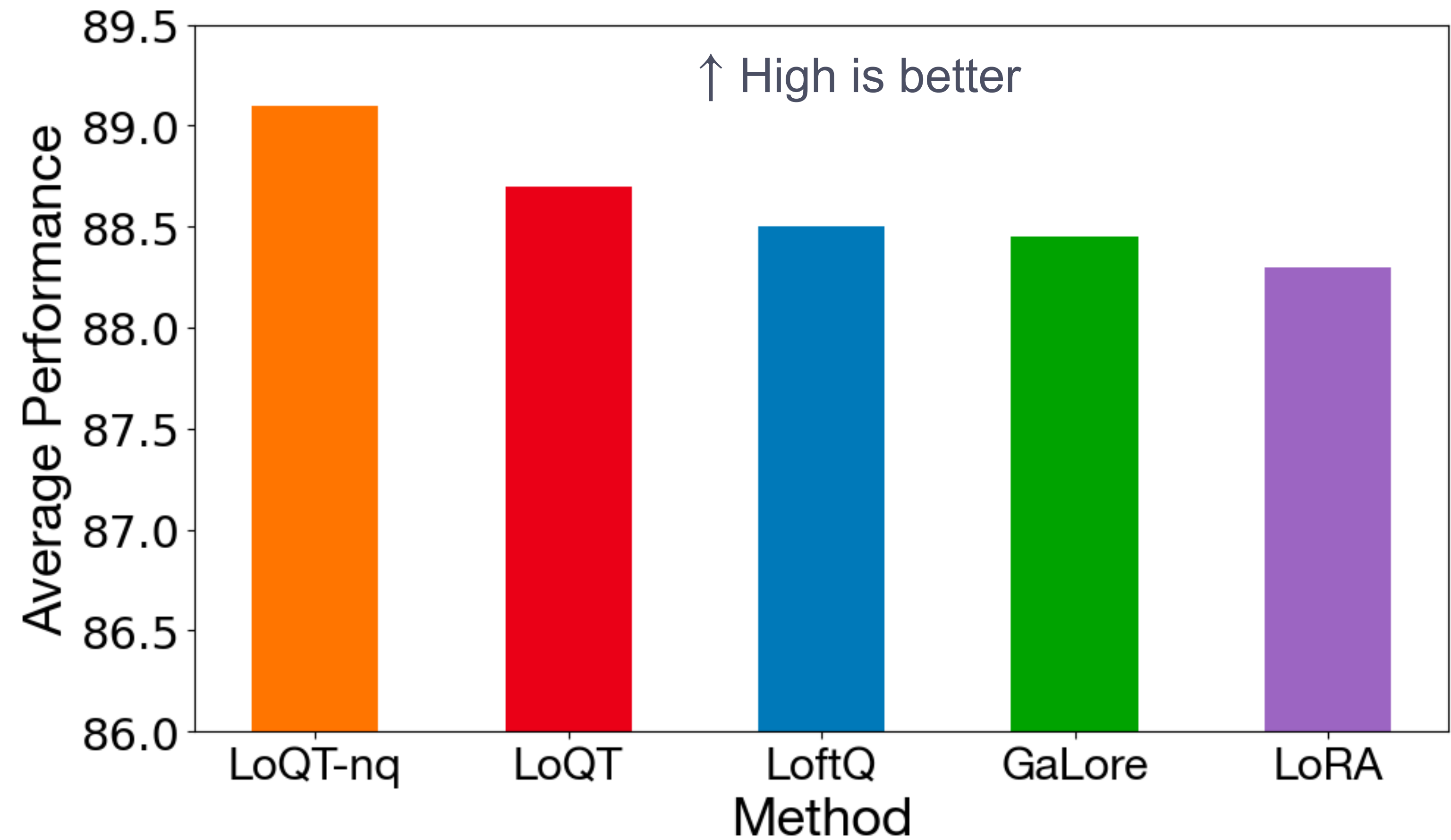
$$B = P^{-1}(W_q - W)$$

$$T_i = \tau + \psi^i$$



- **Simple 4-bit Quantization of W and P :** Flattens early.
- **Error Compensation (EC):** This gap is reduced by 72%.
- **Exponentially increasing update intervals (EI) and EC:** the gap is reduced by 93%.

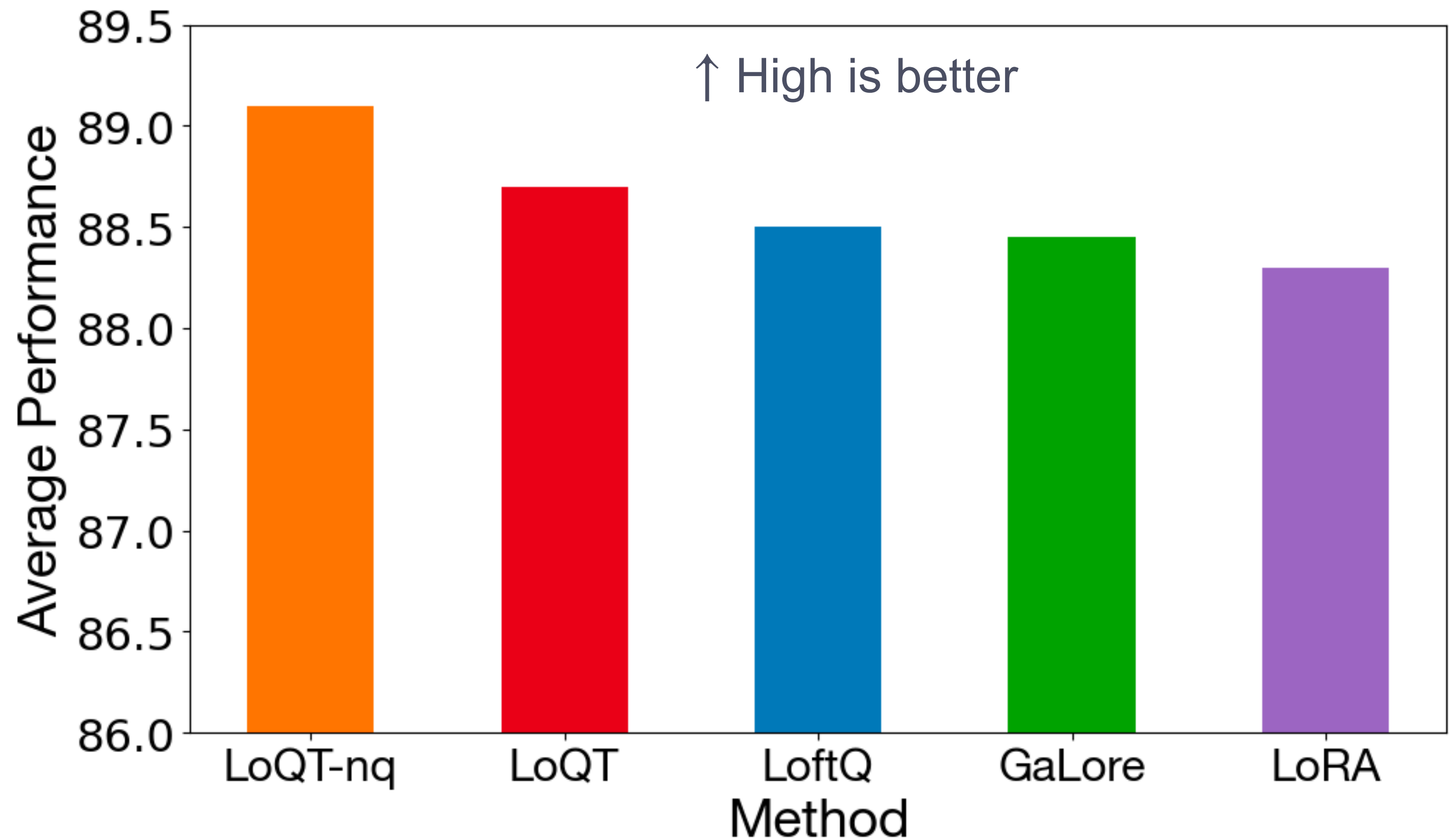
Fine-tuning DeBERTaV3-base Models on GLUE



Fine-tuning DeBERTaV3-base Models on GLUE

Takeaways

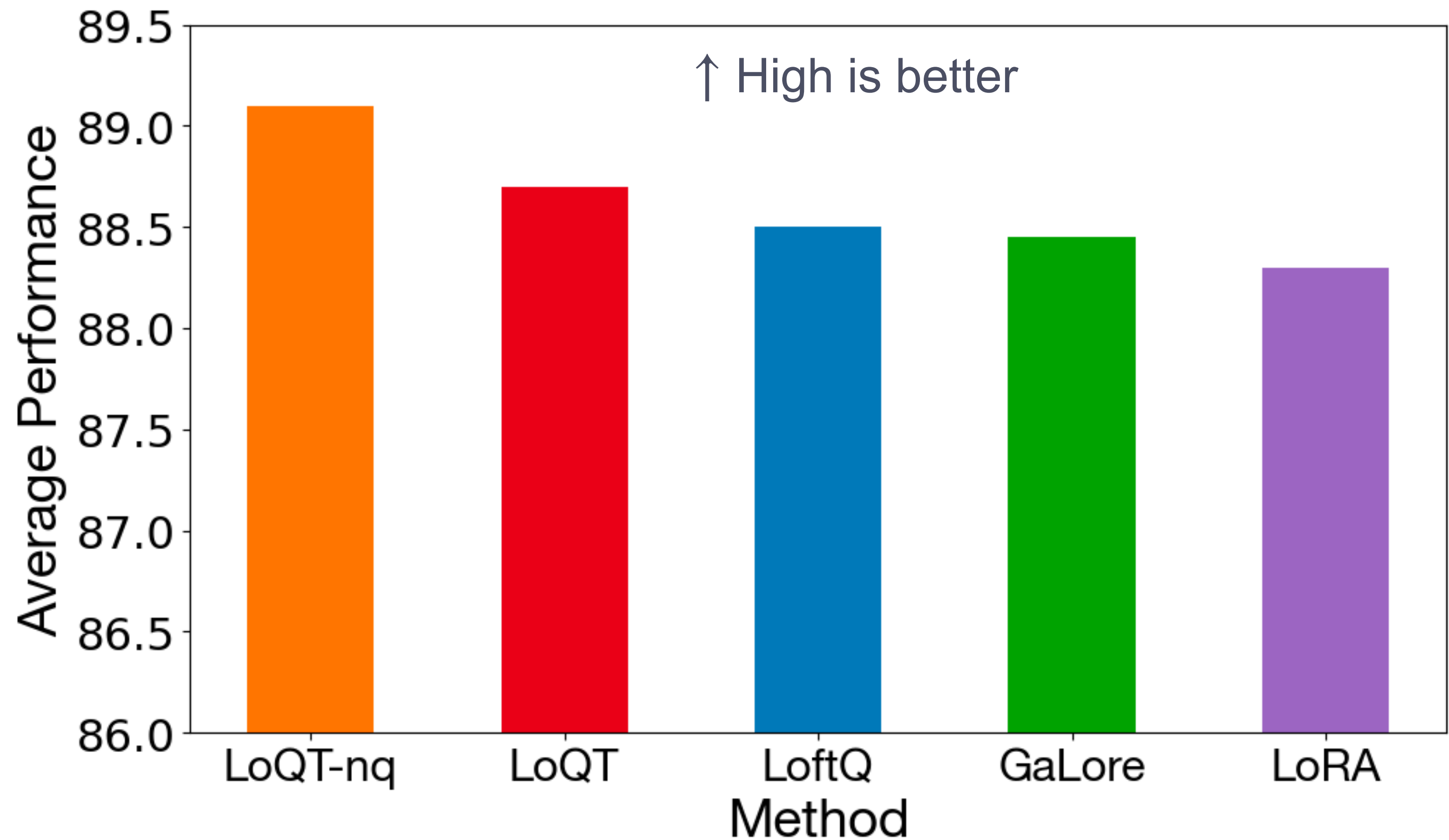
- **LoQT** and **LoQT-*nq*** show promising results.



Fine-tuning DeBERTaV3-base Models on GLUE

Takeaways

- **LoQT** and **LoQT-*nq*** show promising results.
- Gradient-initialized LoRA factors can benefit fine-tuning.

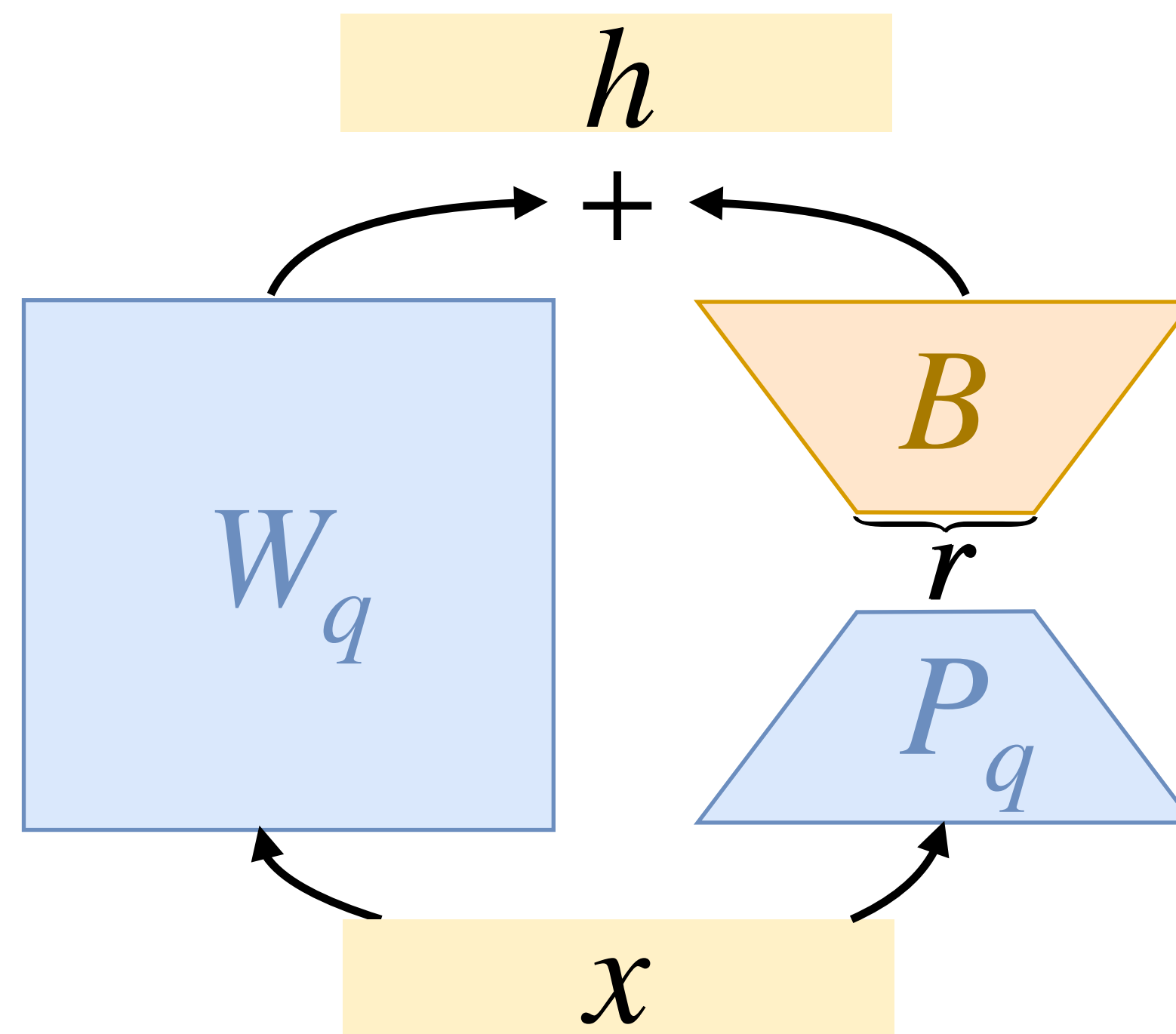


GSM8k

Continued Pretraining Results

Limitations and Future Work

- **Speed:** LoQT is 16% slower than FP16 training.



Limitations and Future Work

- **Speed:** LoQT is 16% slower than FP16 training.
- **Dynamic Scheduling** for update intervals.

$$T_i = \tau + \psi^i$$


Limitations and Future Work

- **Speed:** LoQT is 16% slower than FP16 training.
- **Dynamic Scheduling** for update intervals.
- **Alternative Architectures and Quantization.**

Key Takeaways

- **Pretraining** from scratch with **low-rank updates** and **4-bit quantization** can match FP16 training.

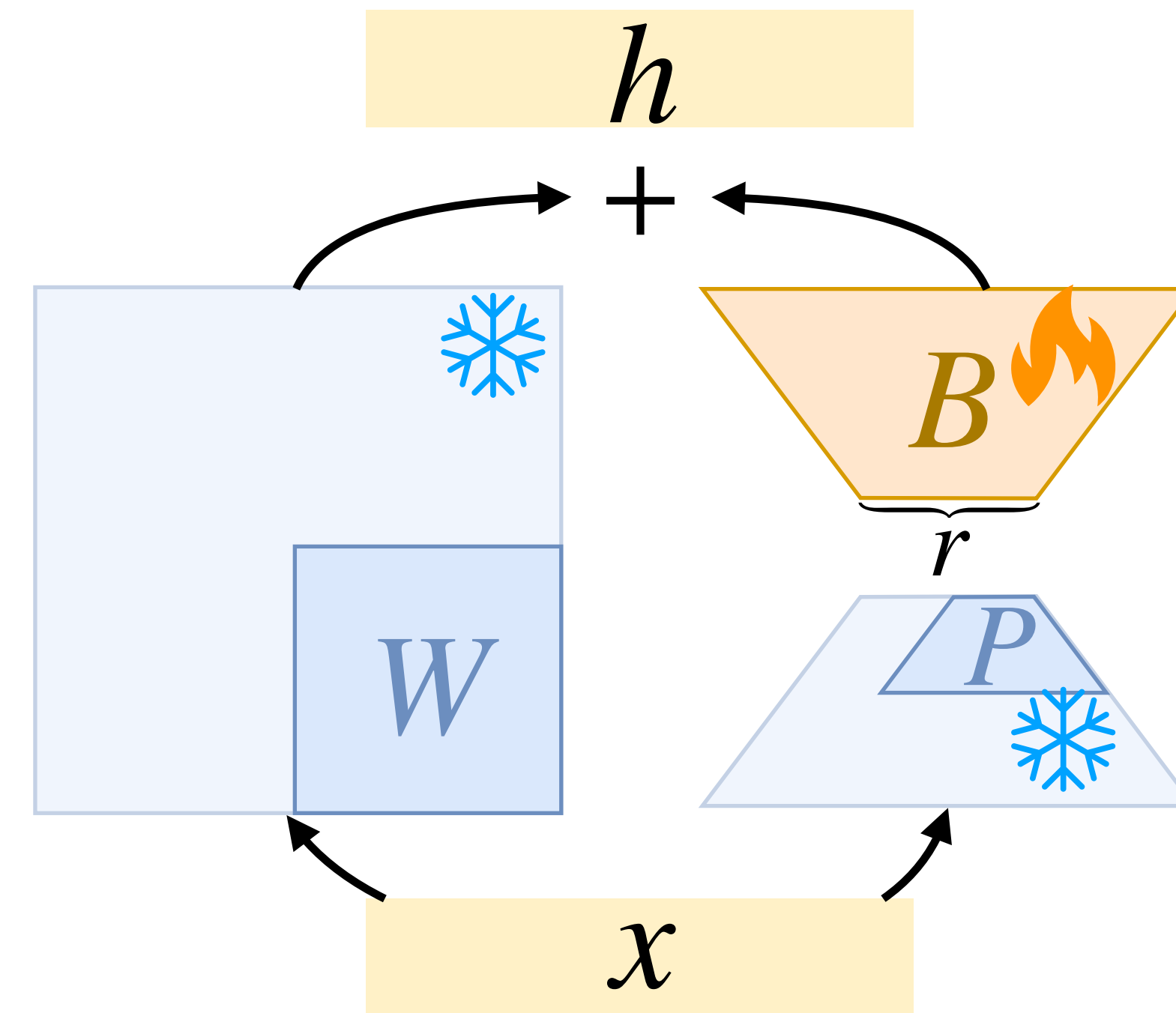
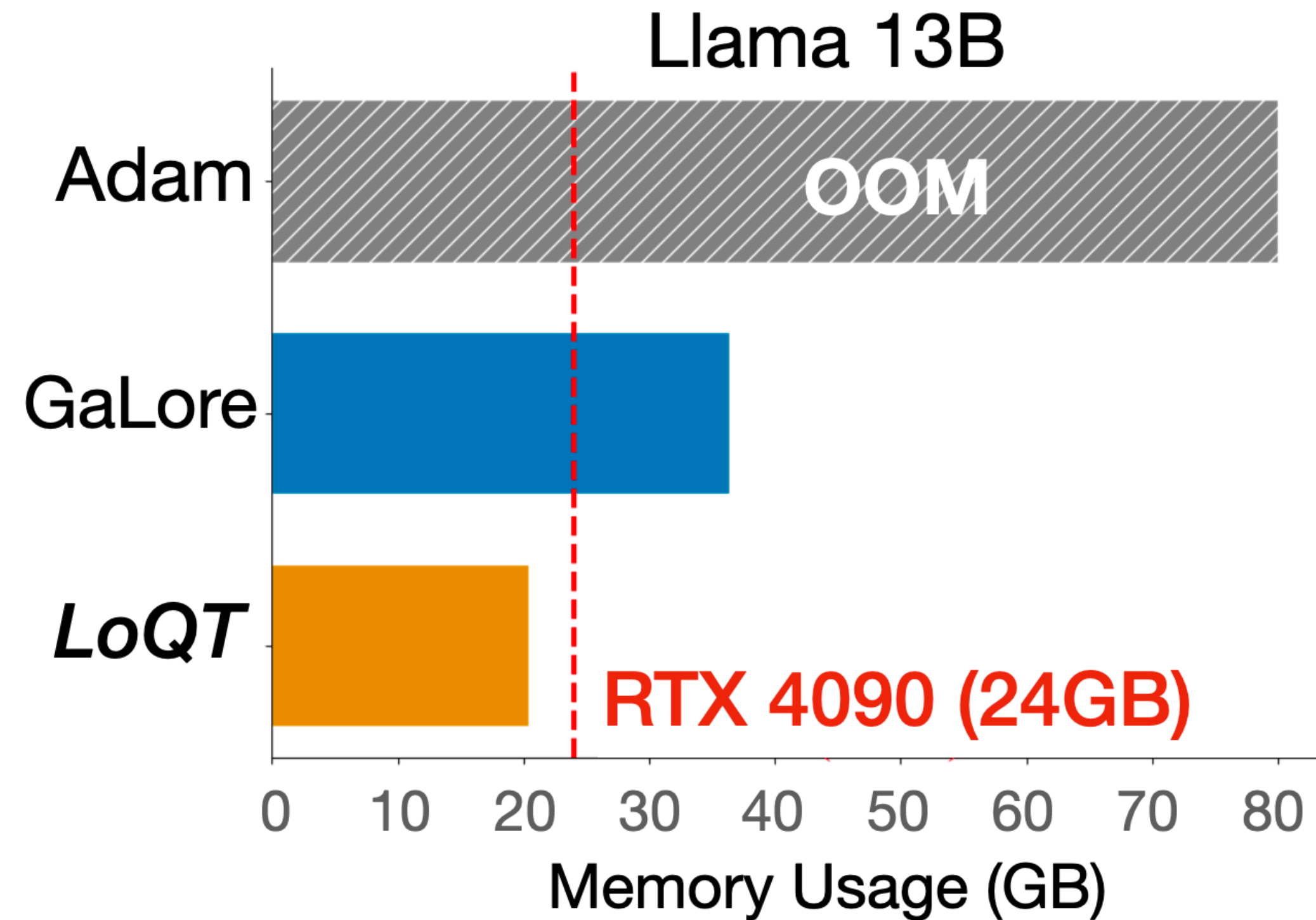
Key Takeaways

- **Pretraining** from scratch with **low-rank updates** and **4-bit quantization** can match FP16 training.
- **Low Memory:** LoQT enables training on consumer hardware.

Key Takeaways

- **Pretraining** from scratch with **low-rank updates** and **4-bit quantization** can match FP16 training.
- **Low Memory:** LoQT enables training on consumer hardware.
- **Fine-tuning** may benefit from *gradient-based initialization*.

LoQT: Low Rank Adapters for Quantized Pre-Training



Sebastian Loeschcke*, Mads Tofttrup*, Michael J. Kastoryano,
Serge Belongie, and Vésteinn Snæbjarnarson

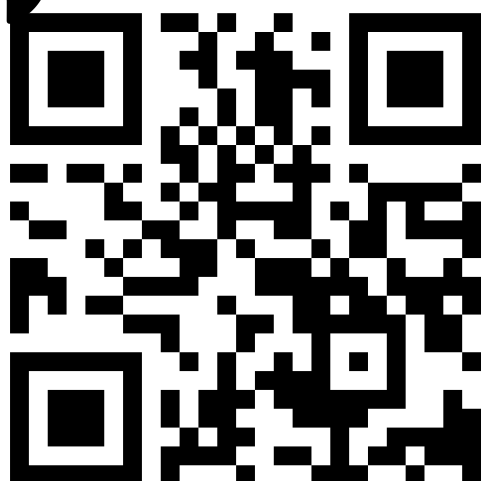
UNIVERSITY OF
COPENHAGEN



AARHUS UNIVERSITY

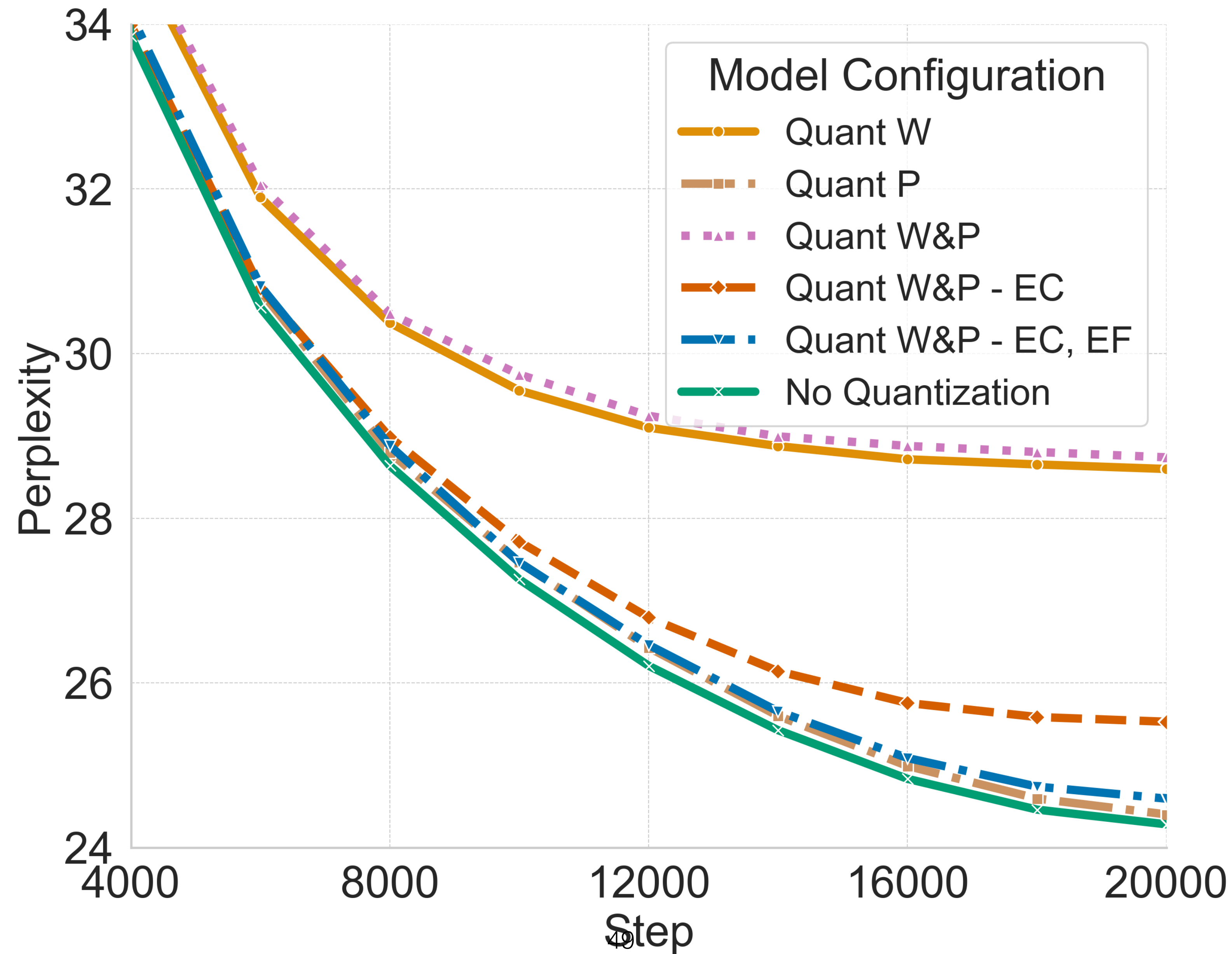


PIONEER CENTRE FOR
ARTIFICIAL INTELLIGENCE

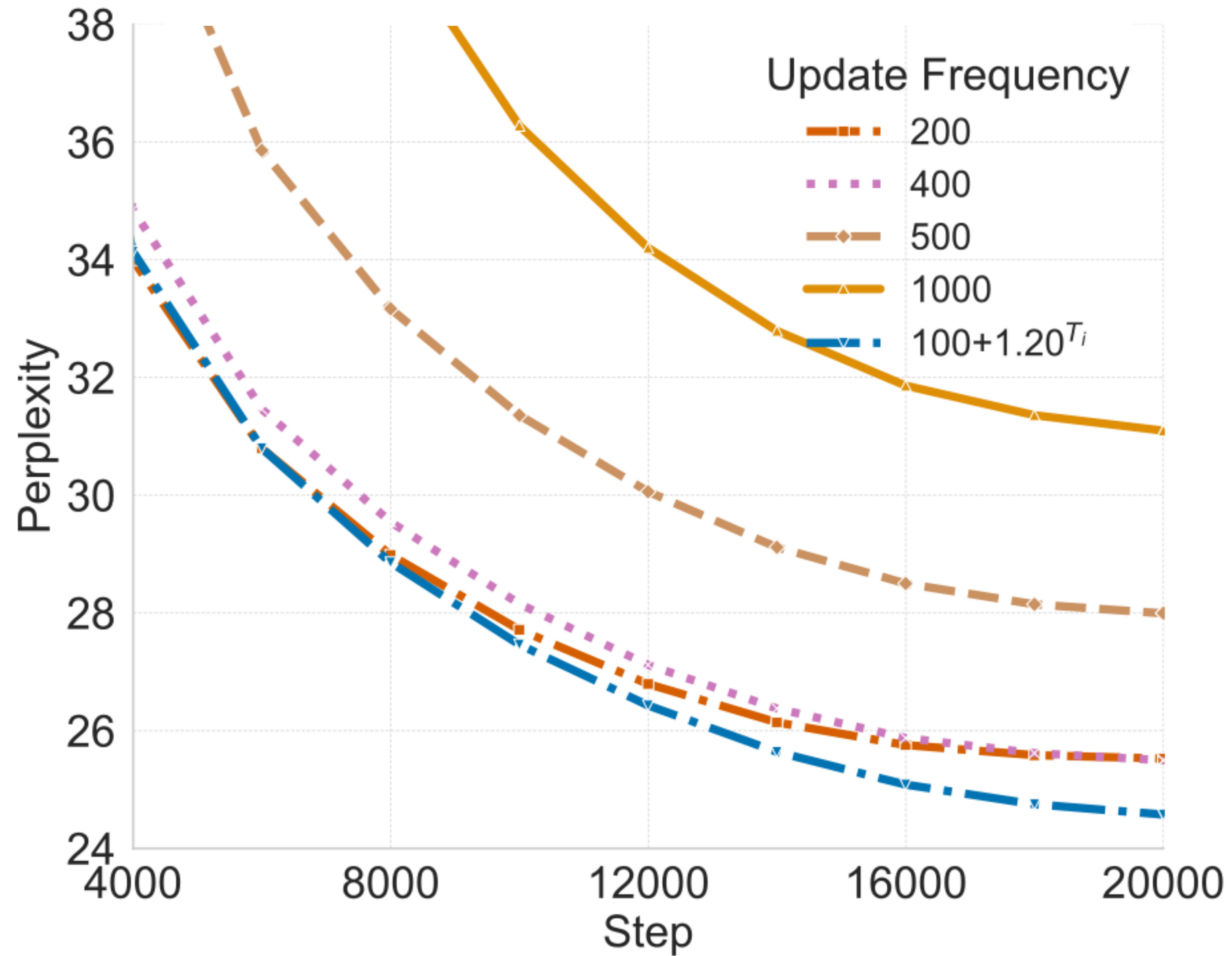


Ekstra

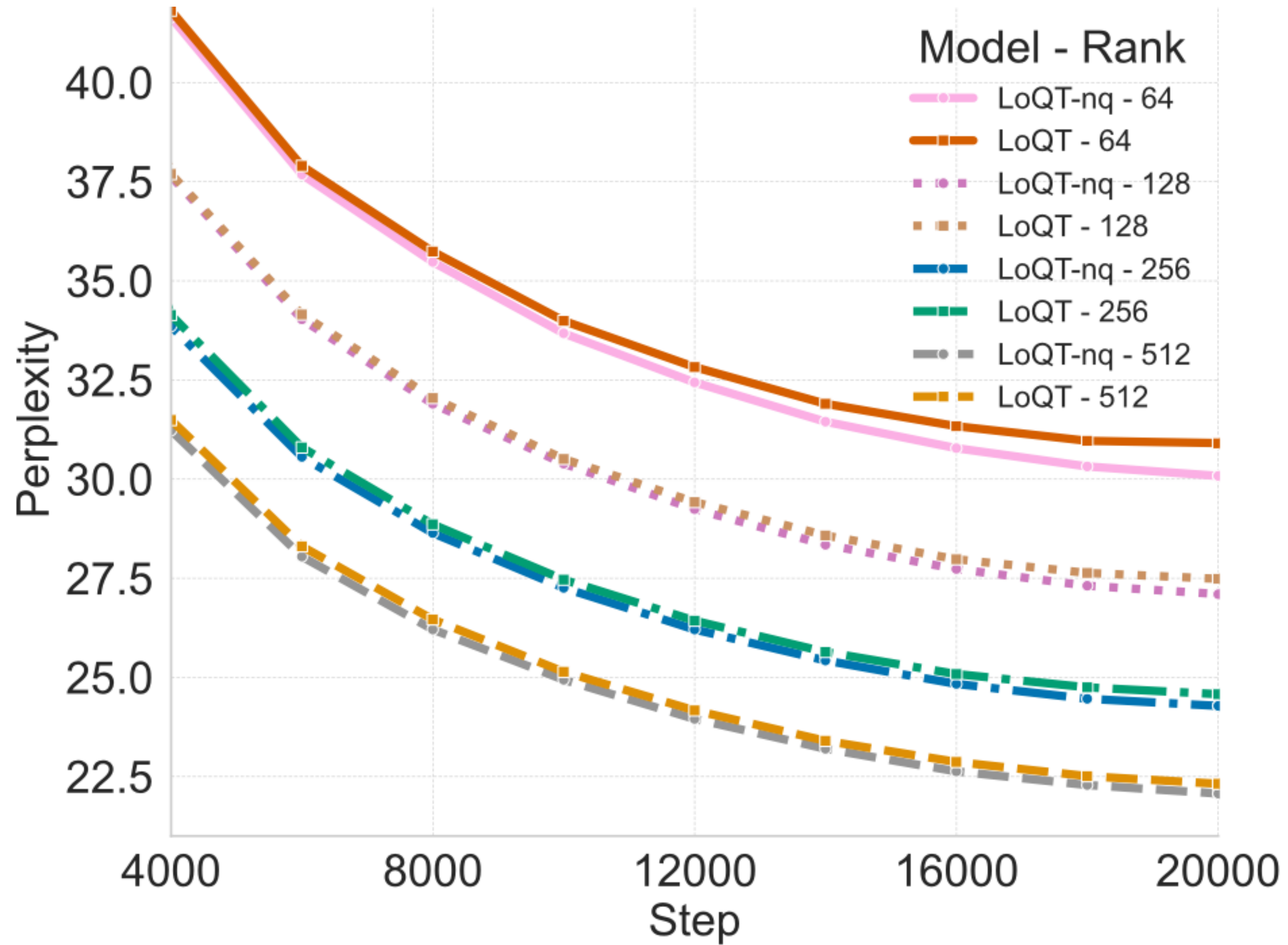
Component Ablations



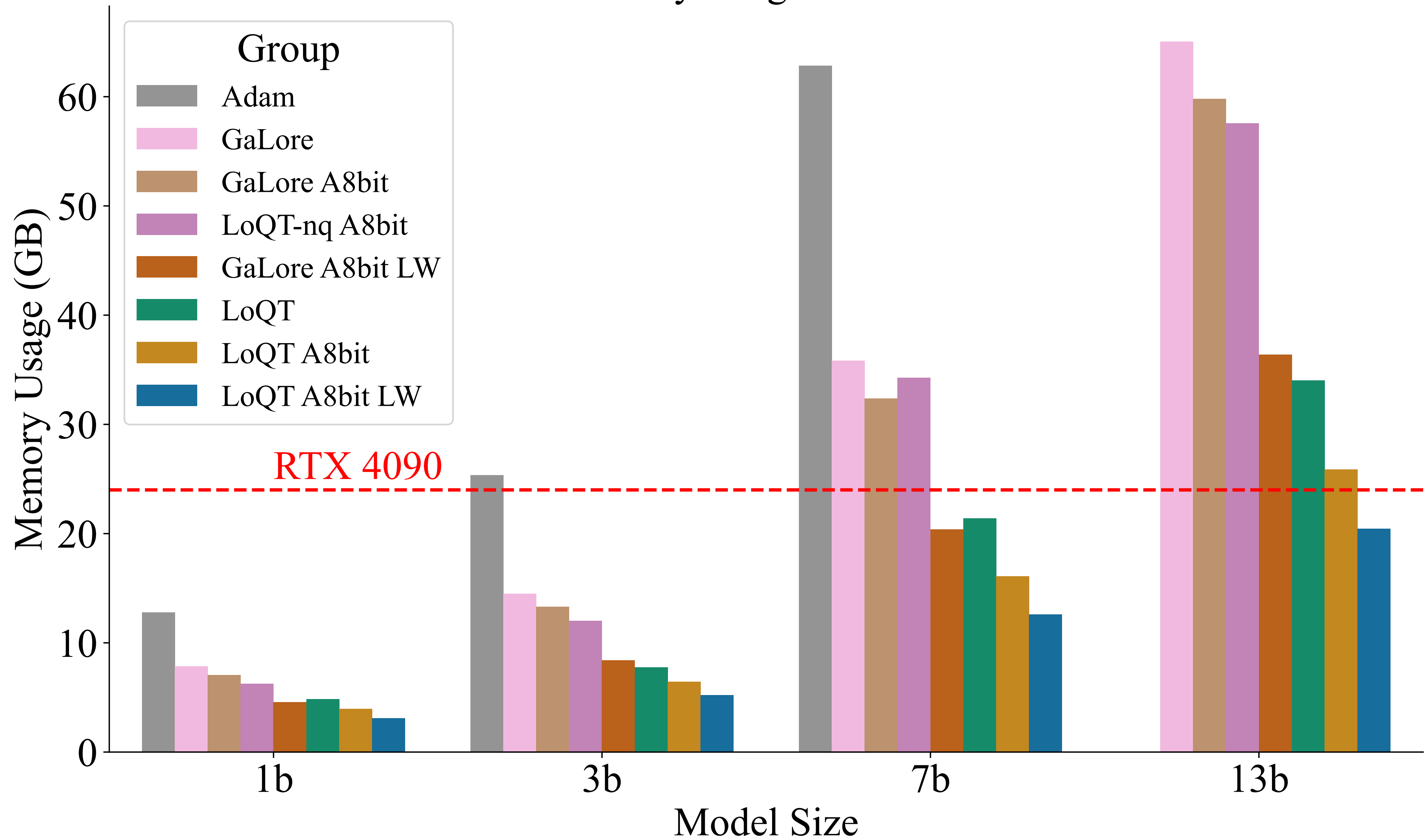
Update Frequency Ablations



Rank Ablations

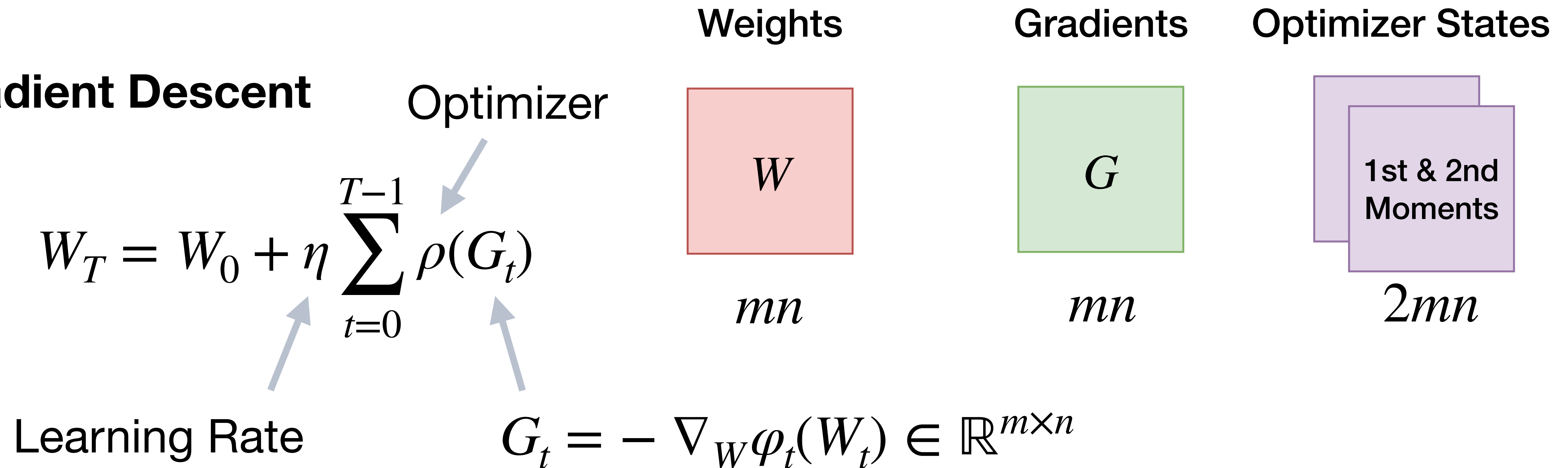


Memory Usage vs. Model



Memory requirements by LLMs

- **Gradient Descent**



LoRA - Low-Rank Adaption

- Gradient Descent

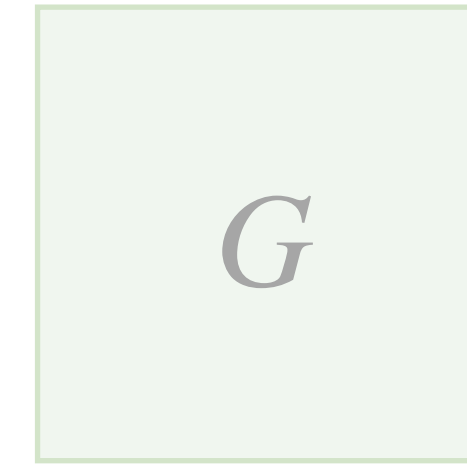
$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(G_t)$$

Weights



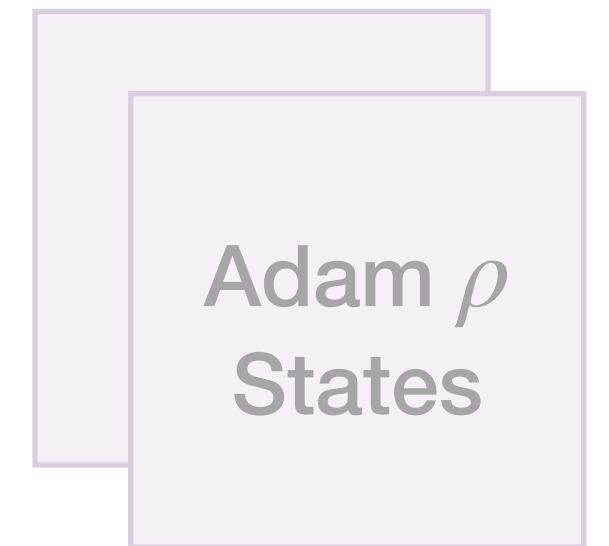
mn

Gradients



mn

Optimizer States

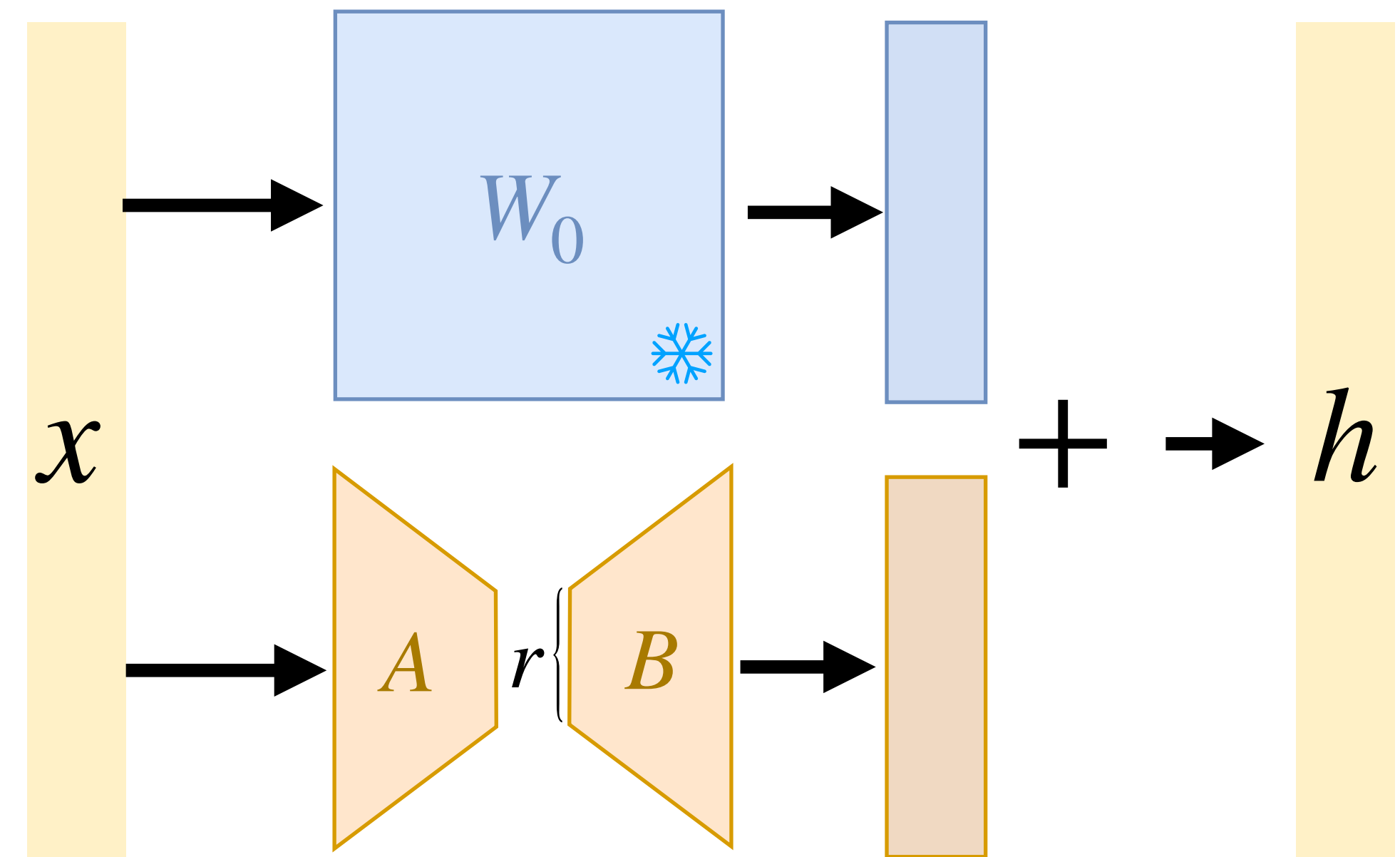


Adam ρ
States

$2mn$

- LoRA (Hu et. al. 2021)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(A)\rho(B)$$



LoRA - Low-Rank Adaption

- Gradient Descent

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(G_t)$$

Weights



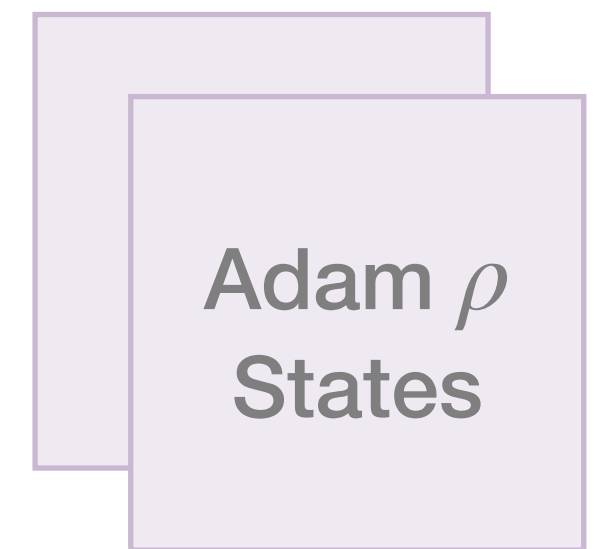
mn

Gradients



mn

Optimizer States



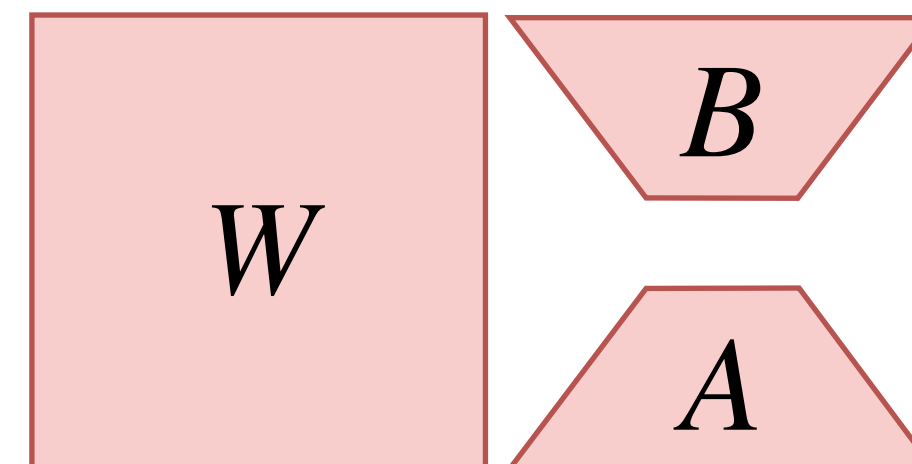
Adam ρ
States

$2mn$

- LoRA (Hu et. al. 2021)

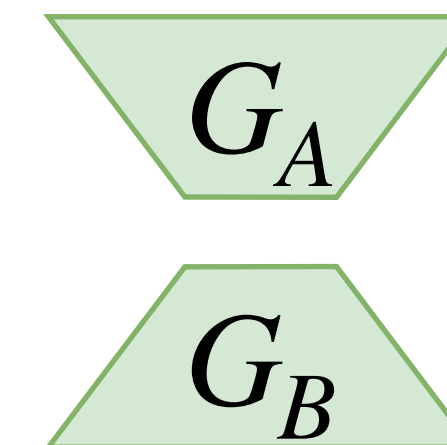
$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(A)\rho(B)$$

↓ Increase



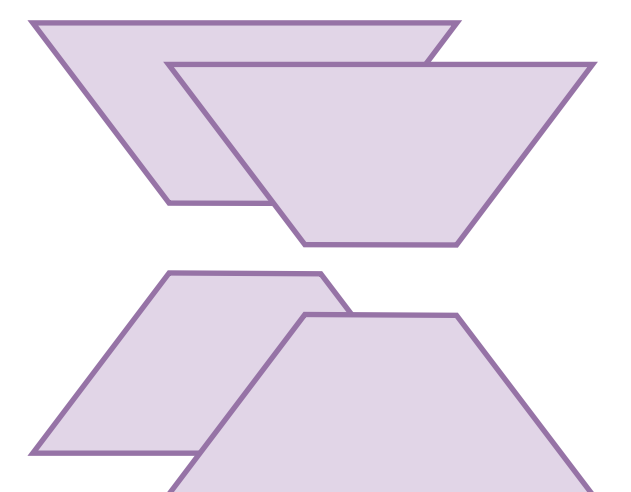
$mn + nr + mr$

↓ Decrease



$nr + mr$

↓ Decrease

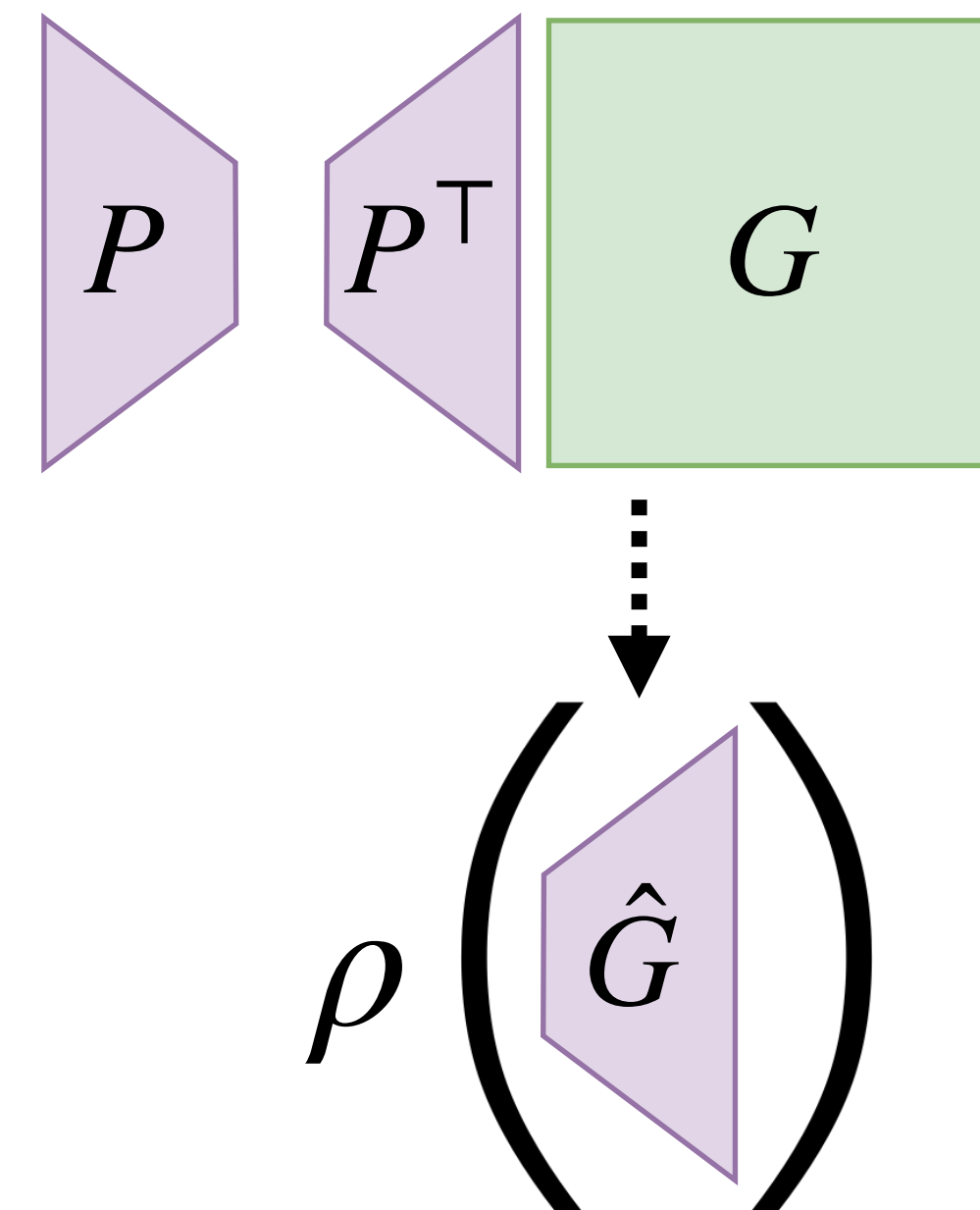


$2nr + 2mr$

Low-Rank Projection of Gradient

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times n}}$$



**Gradient
Descent**

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(G_t)$$

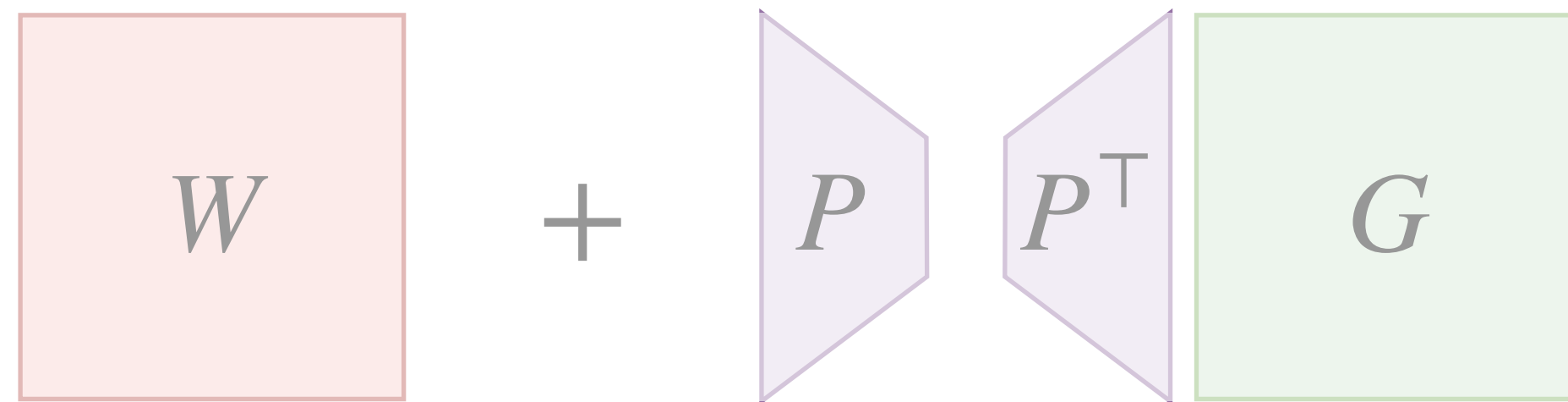
LoRA

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(A)\rho(B)$$

Low-Rank Projection of Gradient

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times m}}$$



$$SVD(G_t) = U \Sigma V^\top$$

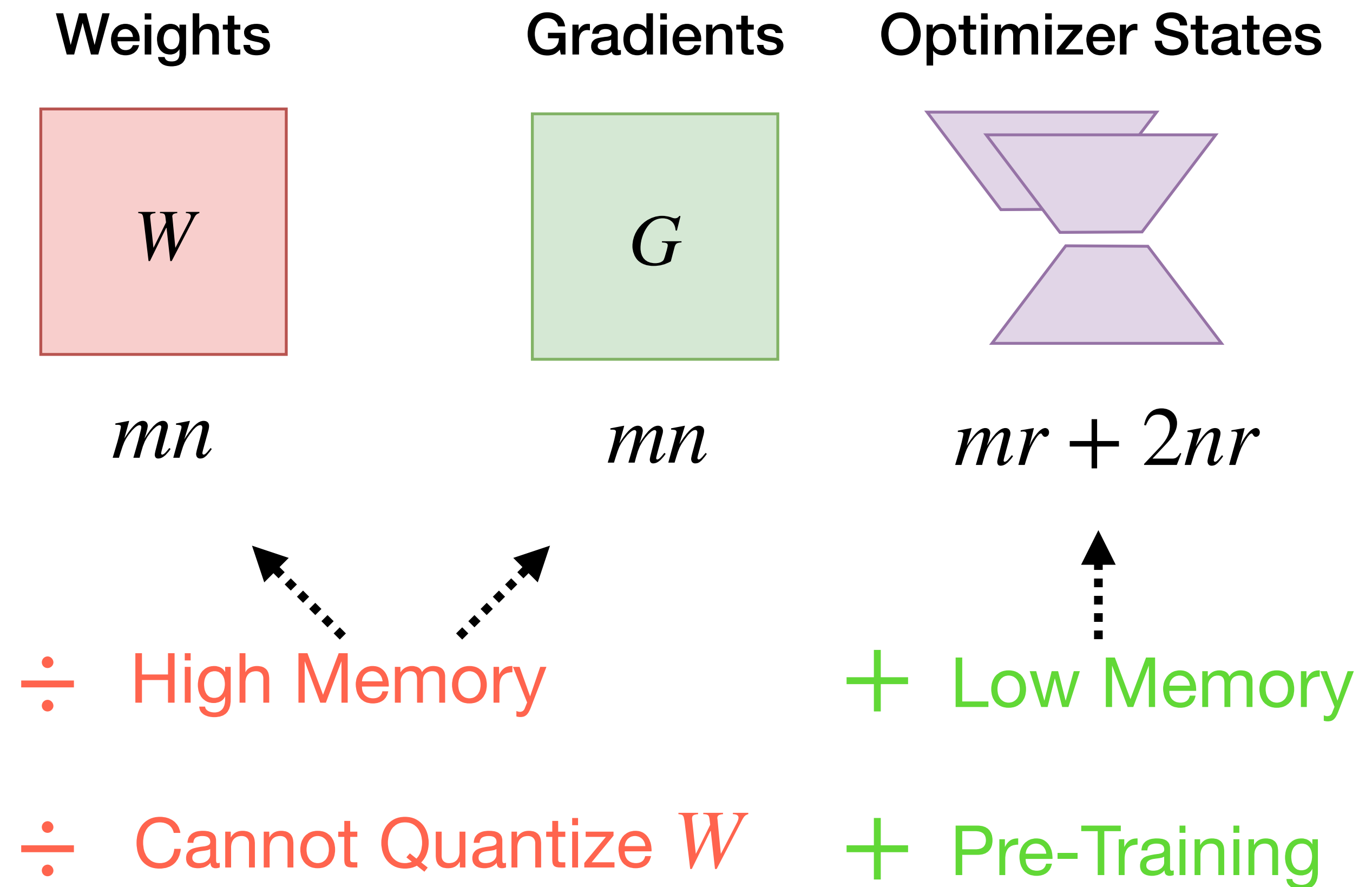
$$P_t = U[:, :r]$$



Low-Rank Projection of Gradient

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times n}}$$



LoQT - Low Rank Adapters for Quantized Pre-Training

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times n}}$$

Weights



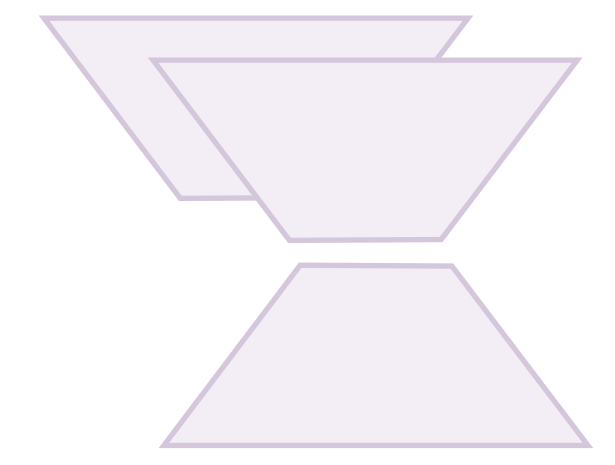
mn

Gradients



mn

Optimizer States



$mr + 2nr$

- **LoQT**

LoRA factors:

A

B

LoRA: $W_T = W_0 + \eta \sum_{t=0}^{T-1} \rho(A) \rho(B)$

LoQT - Low Rank Adapters for Quantized Pre-Training

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times n}}$$

- **LoQT**

LoRA factors:

$$W_T = \underbrace{W_0}_{\text{Frozen}} + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\text{Optimize}} \rho(\underbrace{B_t}_{\text{Optimize}})$$

Weights



mn

↓ Increase

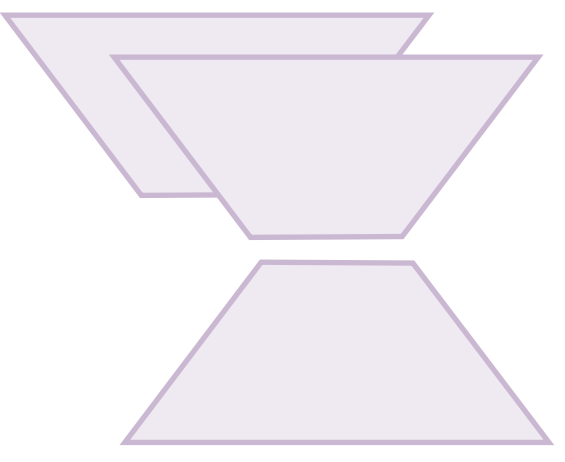
Gradients



mn

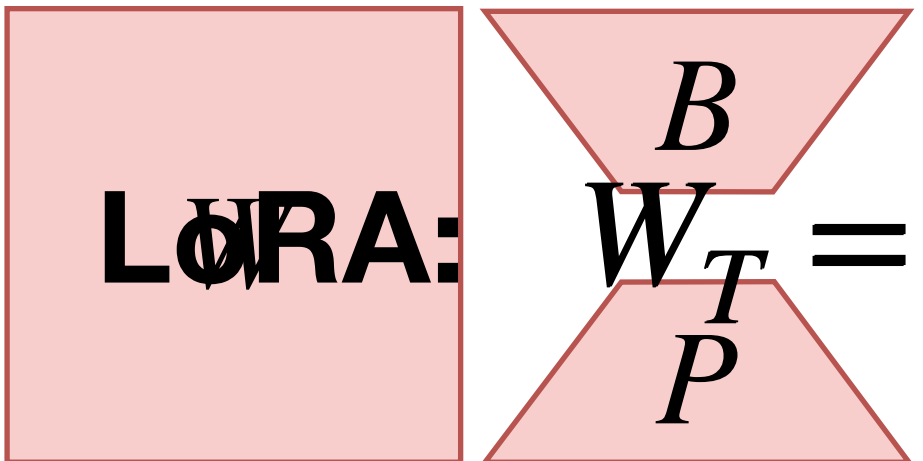
↓ Decrease

Optimizer States



$mr + 2nr$

↓ Decrease



$mn + nr + mr$

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{G_B}_{\text{Frozen}} \rho(\underbrace{A}_{\text{Optimize}}) \rho(\underbrace{B}_{\text{Optimize}})$$

nr

$2nr$

❄️ Frozen

🔥 Optimize

LoQT and GaLore have equivalent Gradient Updates

- **Forward pass for LoQT:** $y = xW + xPB$
- The gradient of the loss w.r.t. B using the Chain Rule:

$$\frac{\partial B}{\partial \mathcal{L}} = (xP)^\top \frac{\partial \mathcal{L}}{\partial y} = P^\top \underbrace{x^\top \frac{\partial \mathcal{L}}{\partial y}}_{G^W} = P^\top G^W$$


$$\text{GaLore } W_T = W_0 + \eta \sum_{t=0}^{T-1} P_t \rho(\overbrace{P_t^\top G_t})$$

LoQT - Low Rank Adapters for Quantized Pre-Training

- **GaLore** (Zhao et. al., 2024)

$$W_T = W_0 + \eta \sum_{t=0}^{T-1} \underbrace{P_t}_{\in \mathbb{R}^{m \times r}} \underbrace{\rho(P_t^\top G_t)}_{\in \mathbb{R}^{r \times n}}$$

Weights



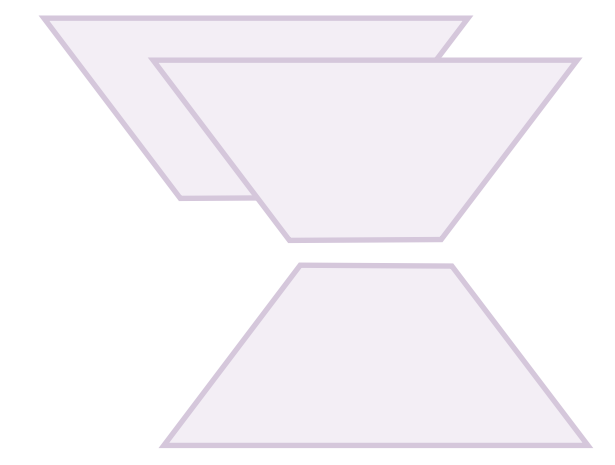
mn

Gradients



mn

Optimizer States



$mr + 2nr$

- **LoQT**

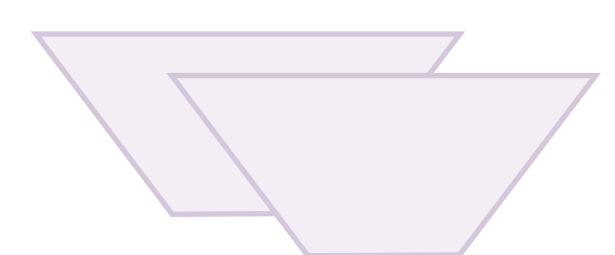
$$W_T = W_0 + \eta \sum_{t=0}^{T-1} P_t \rho(B_t)$$



$mn + nr + mr$



nr



$2nr$

Note: The gradient of the loss w.r.t. B_t : $\frac{\partial \mathcal{L}}{\partial B_t} = P_t^\top G_t \longrightarrow$ Same Gradient Updates