

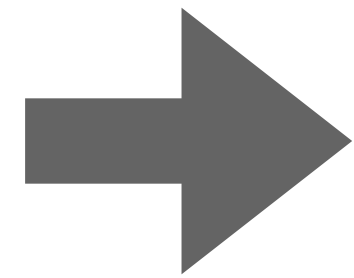
# Large Language Model Basics and MoE Architecture

Yen-Ting Lin 林彥廷

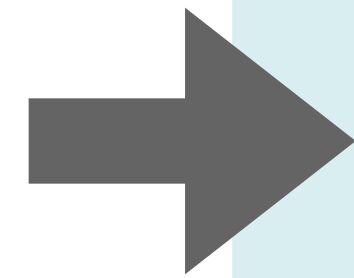
Sep 30, 2024 @ ADL Recitation, NTU

# LLM Development

**Pretraining**



**Instruction  
tuning**

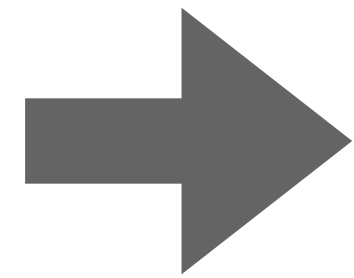


**Learning from  
Feedback**

# LLM Development

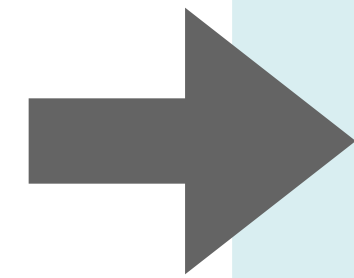
**Scale**

**Pretraining**



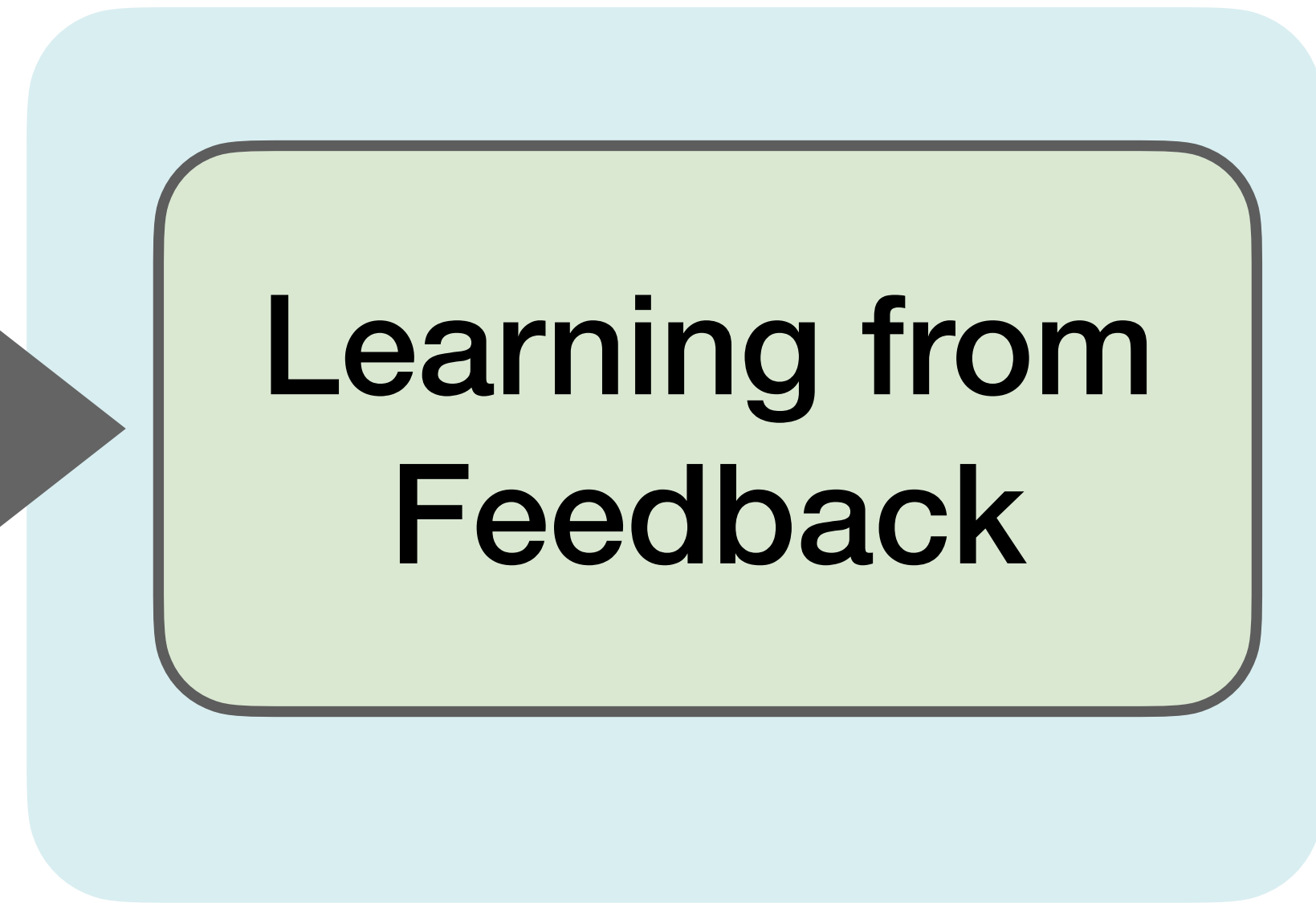
**Adaptation**

**Instruction  
tuning**



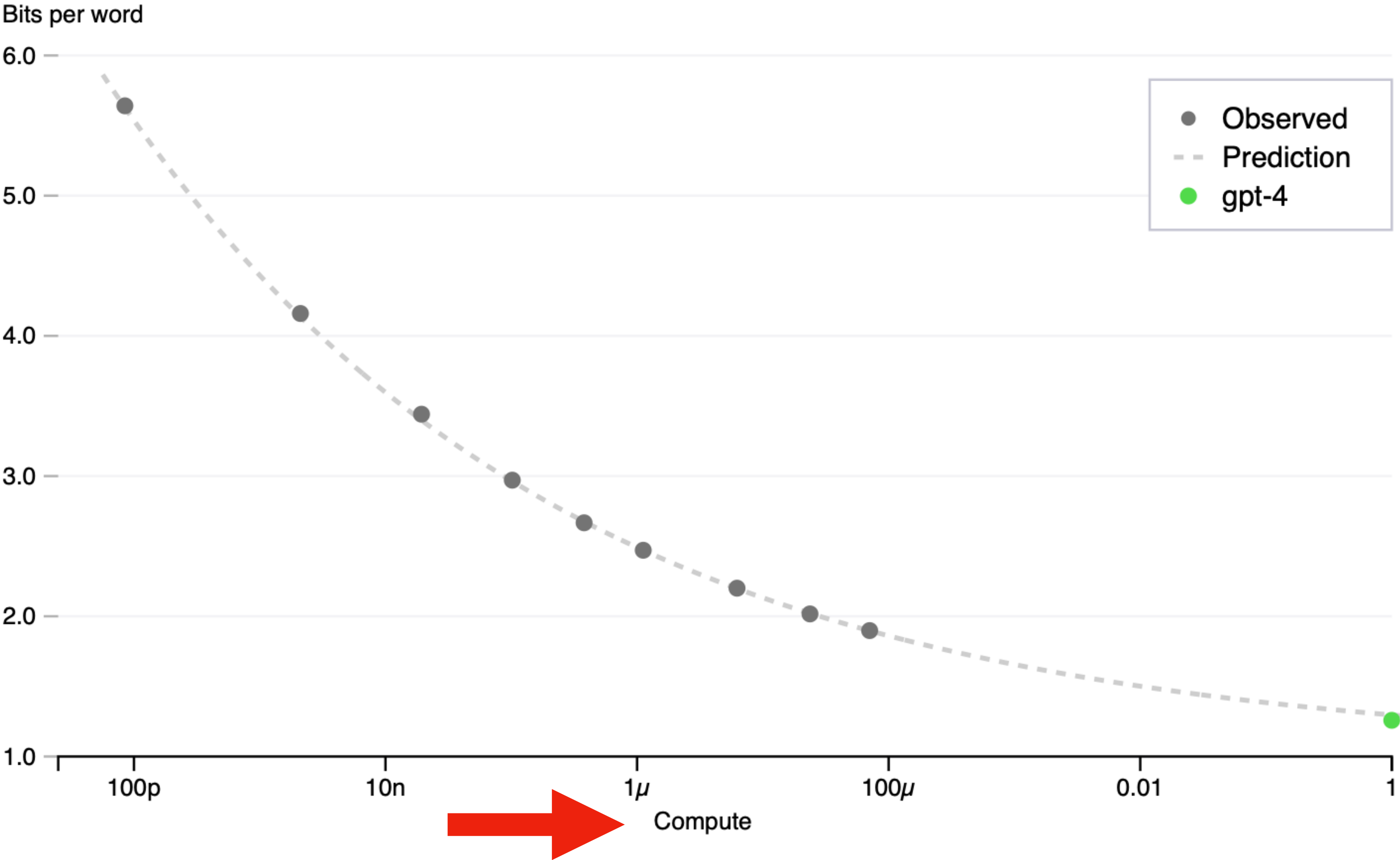
**Alignment**

**Learning from  
Feedback**



# Scaling Law

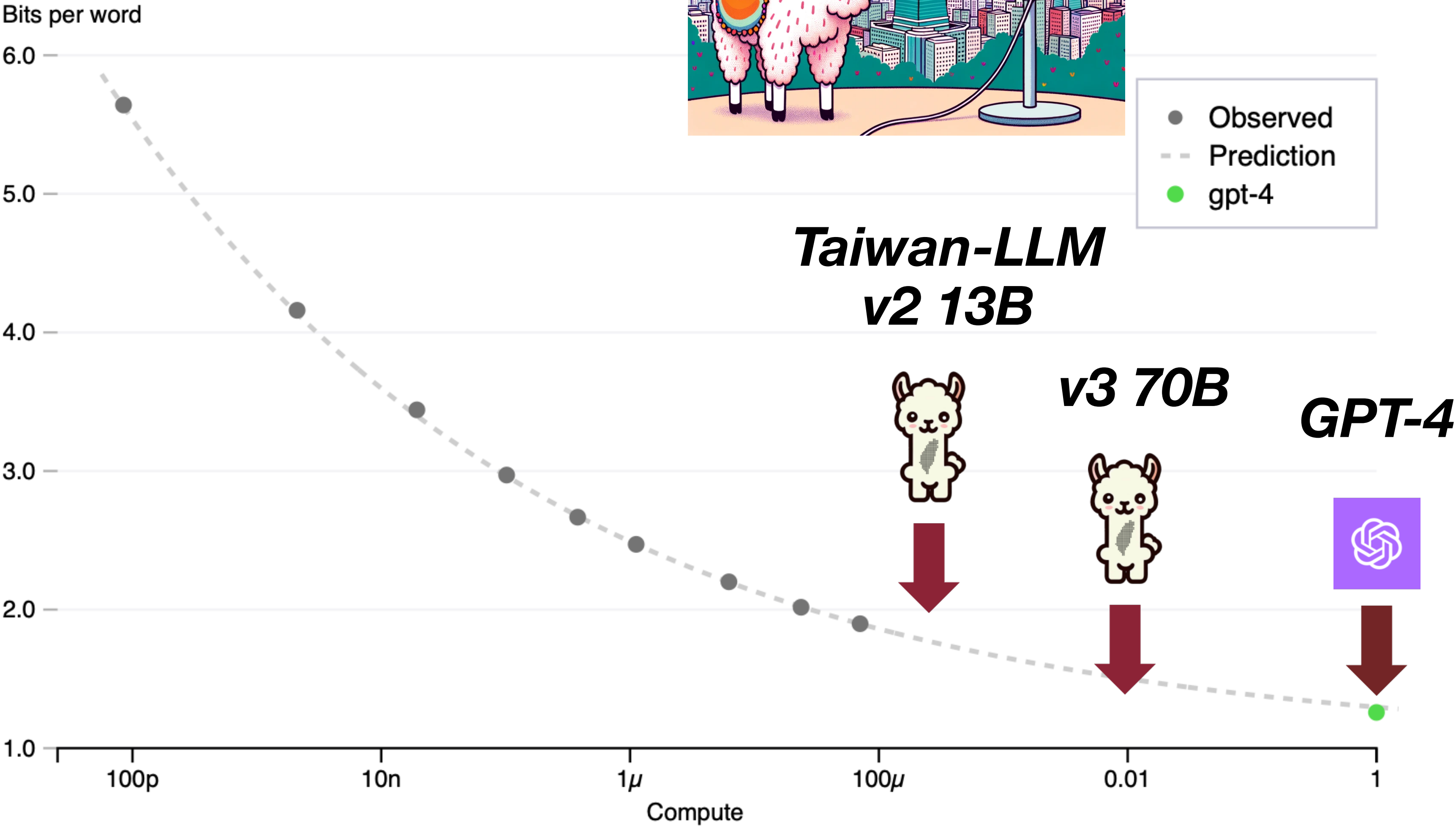
OpenAI codebase next word prediction



# Scaling Law



OpenAI codebase next word prediction



# Emergent Abilities

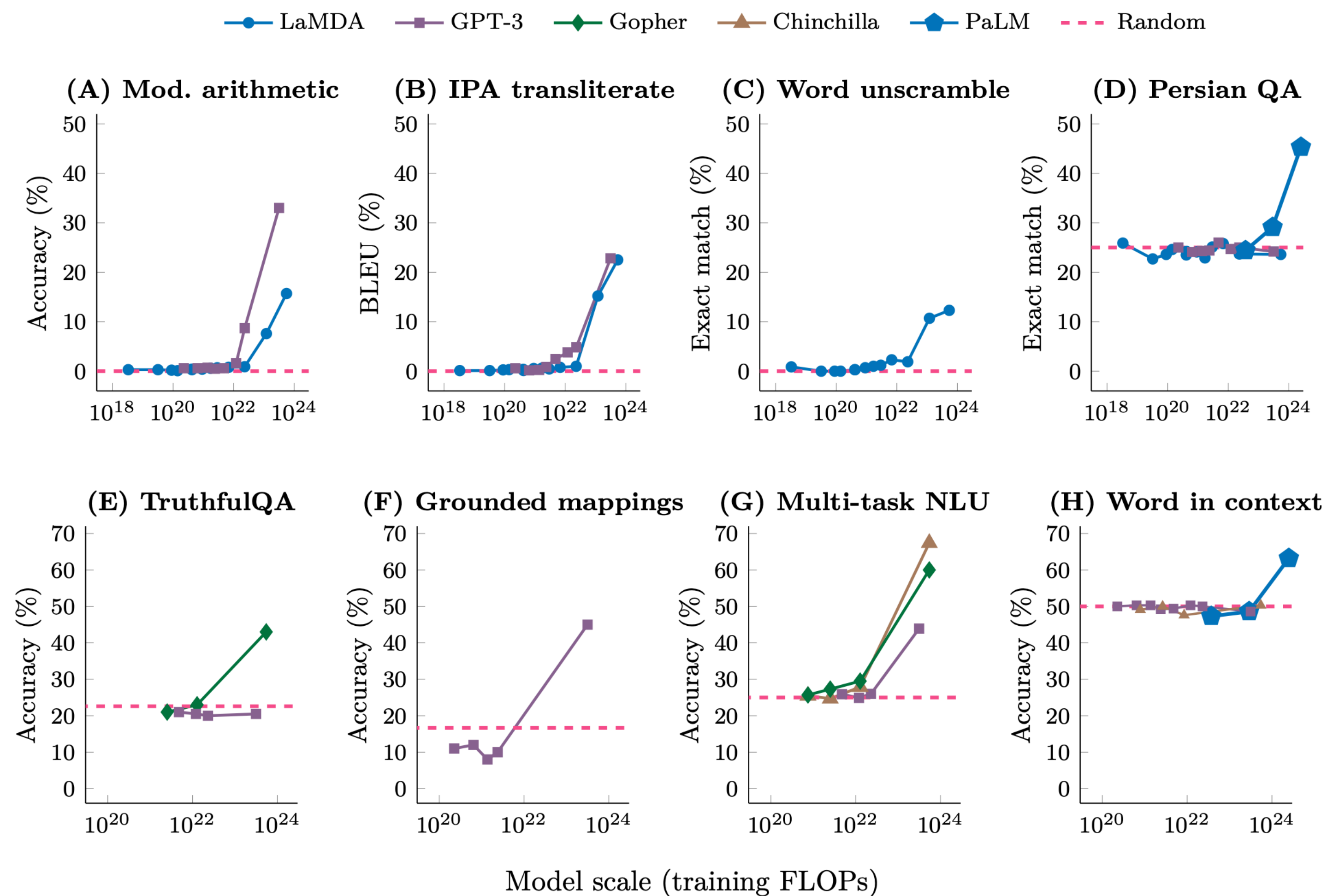
*An ability is emergent if it is not present in smaller models  
but is present in larger models.*

## Emergent Abilities of Large Language Models

Jason Wei<sup>1</sup>

*jasonwei@google.com*

# Emergent Abilities



Emergent Abilities of Large Language Models

# How to Scale?

Model scale (training FLOPs)

Training FLOPs  $\sim$  6 x Model Size x Number of Tokens

# How to Scale?

Model scale (training FLOPs)

Training FLOPs  $\sim 6 \times \text{Model Size} \times \text{Number of Tokens}$

Taiwan-LLM Training FLOPs =  $6 \times 70\text{e}9 \times 500\text{e}9 = 2\text{e}23$

GPT-4 FLOPs =  $2\text{e}25$

# Cost of Frontier Models

 Entrepreneur

## Anthropic CEO: AI Will Cost \$10 Billion to Train By 2025

Anthropic CEO Dario Amodei puts a price tag on AI as it stands today, and hints at the billions of dollars that will be required to make it...

Jul 8, 2024



GPT-4 FLOPs =  $2e^{25}$

~100M USD

2023 Q1

GPT-5 FLOPs =  $2e^{26}$ ?

~1B USD

2024 Q4?

GPT-6 FLOPs =  $2e^{27}$ ?

~10B USD

2025?

# Chinchilla's Law



---

## Training Compute-Optimal Large Language Models

Jordan Hoffmann\*, Sebastian Borgeaud\*, Arthur Mensch\*, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, Tom Hennigan, Eric Noland, Katie Millican, George van den Driessche, Bogdan Damoc, Aurelia Guy, Simon Osindero, Karen Simonyan, Erich Elsen, Jack W. Rae, Oriol Vinyals and Laurent Sifre\*

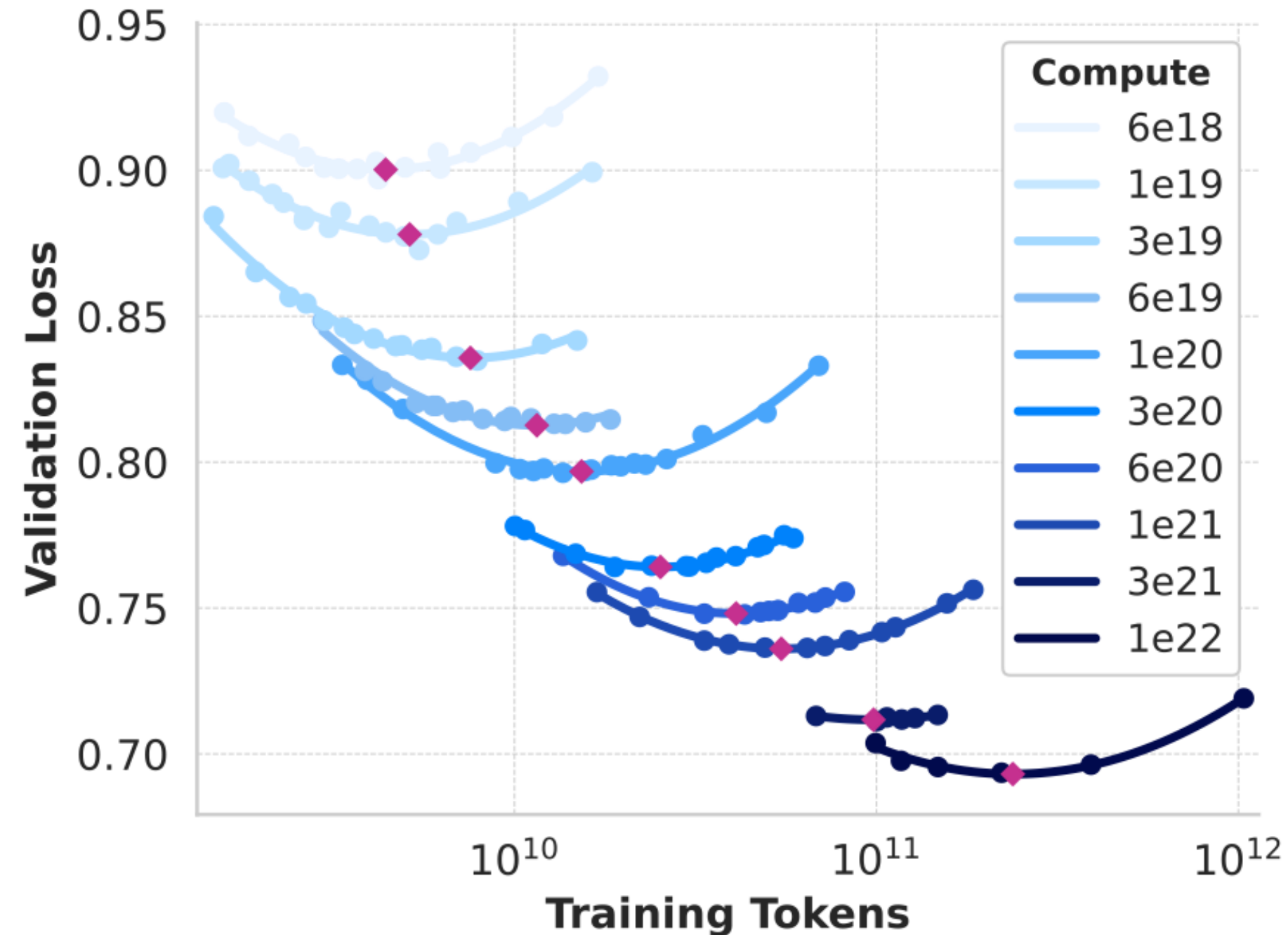
\*Equal contributions

# Research Question

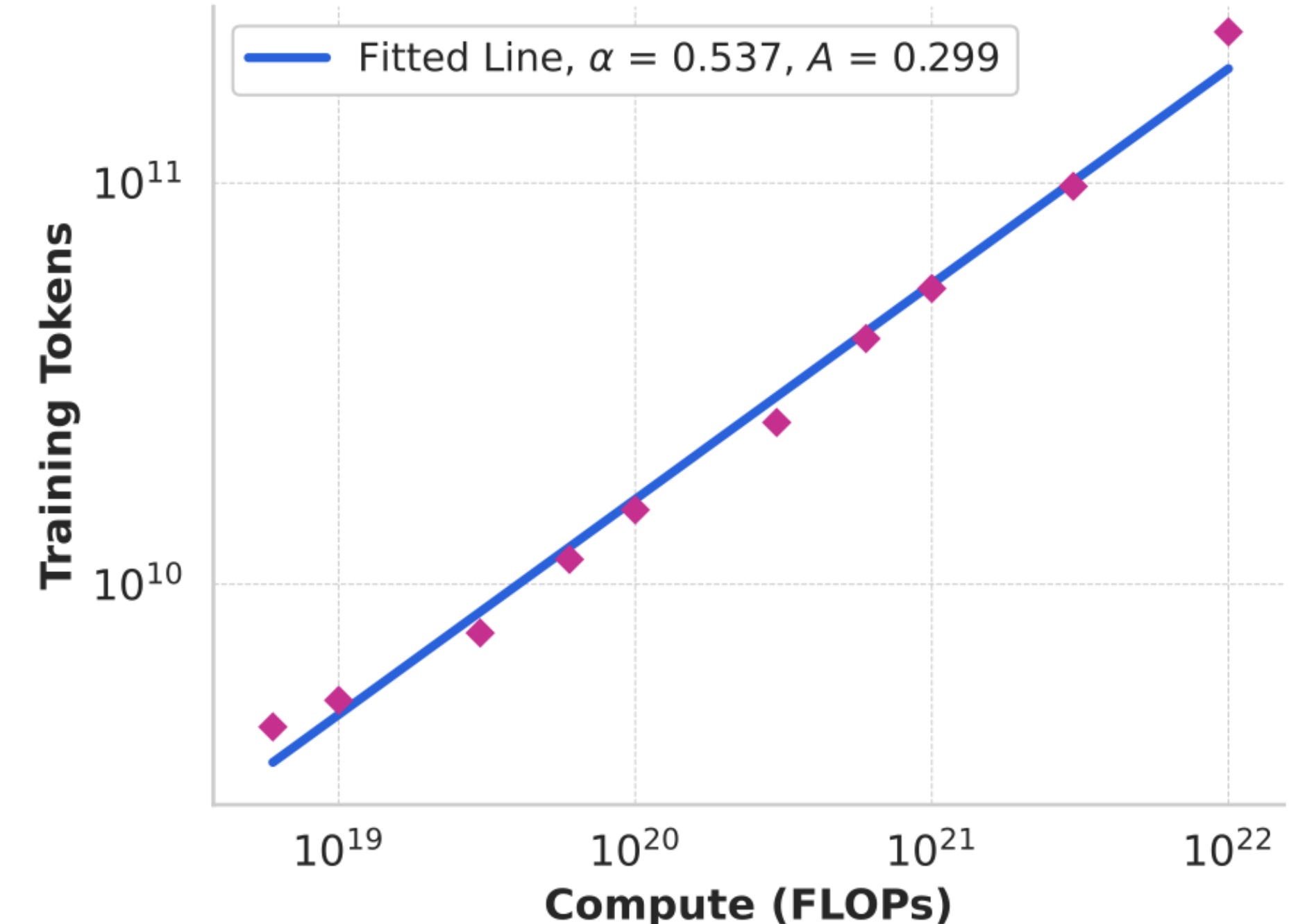
- Given a fixed FLOPs budget,  $C$ , how should one trade-off model size and number of training tokens
- Final pre-training loss as a function of  $L(N, D)$ 
  - $N$ : # parameters
  - $D$ : # training tokens

$$N_{opt}(C), D_{opt}(C) = \underset{N, D \text{ s.t. } \text{FLOPs}(N, D) = C}{\operatorname{argmin}} L(N, D).$$

# Fix FLOPs

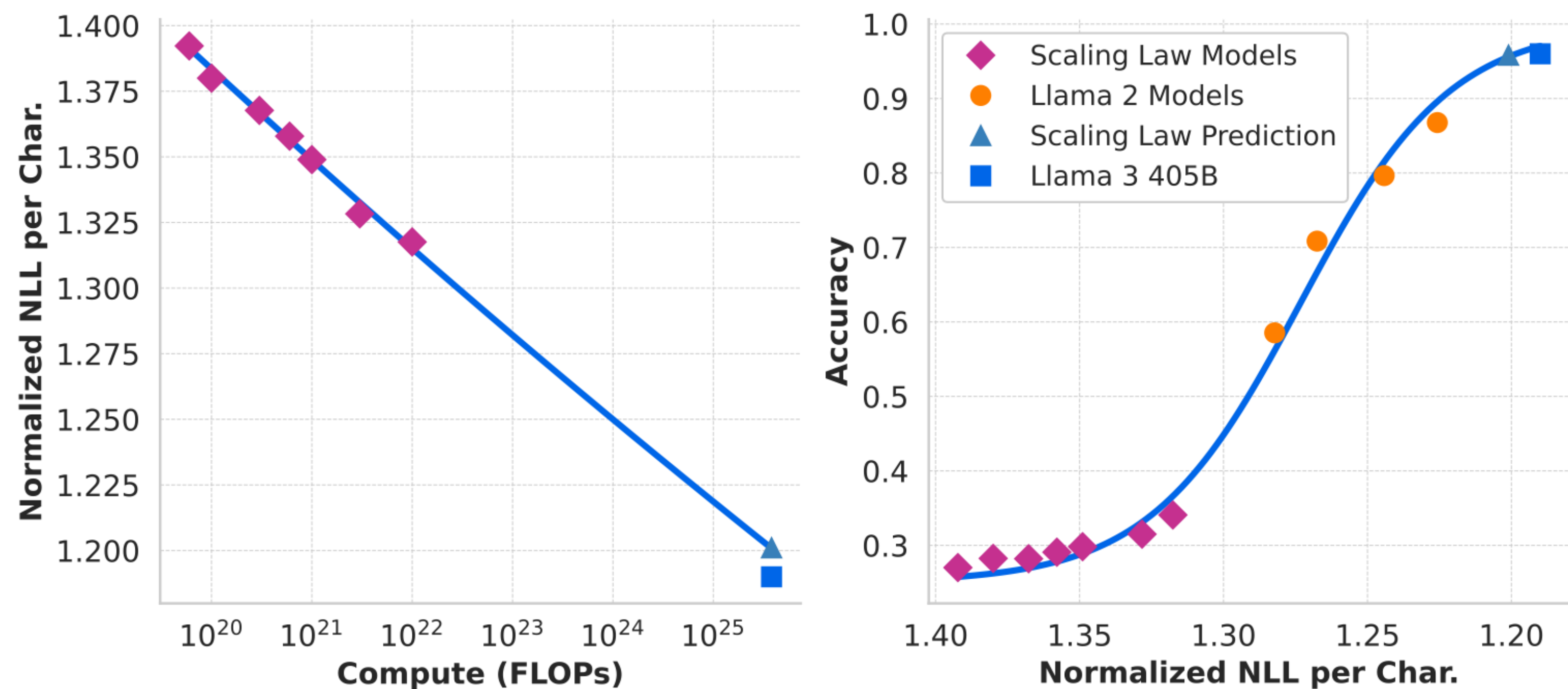


**Figure 2 Scaling law IsoFLOPs curves** between  $6 \times 10^{18}$  and  $10^{22}$  FLOPs. The loss is the negative log-likelihood on a held-out validation set. We approximate measurements at each compute scale using a second degree polynomial.



**Figure 3 Number of training tokens in identified compute-optimal models as a function of pre-training compute budget.** We include the fitted scaling-law prediction as well. The compute-optimal models correspond to the parabola minimums in Figure 2.

# Pertaining loss as downstream performance indicator?



**Figure 4 Scaling law forecast for ARC Challenge.** *Left:* Normalized negative log-likelihood of the correct answer on the ARC Challenge benchmark as a function of pre-training FLOPs. *Right:* ARC Challenge benchmark accuracy as a function of the normalized negative log-likelihood of the correct answer. This analysis enables us to predict model performance on the ARC Challenge benchmark before pre-training commences. See text for details.

# Transformer Architecture

- Vanilla Transformers vs LLaMa (請見去年實習課影片)
  - RMS Normalization
  - Rotary Positional Encoding
  - KV-Cache
  - Group Query Attention
  - SwiGLU
- **Mixture of Expert**

# Rumor: GPT4 / Gemini Architecrue

## GPT-4 Architecture, Infrastructure, Training Dataset, Costs, Vision, MoE

Demystifying GPT-4: The engineering tradeoffs that led OpenAI to their architecture.



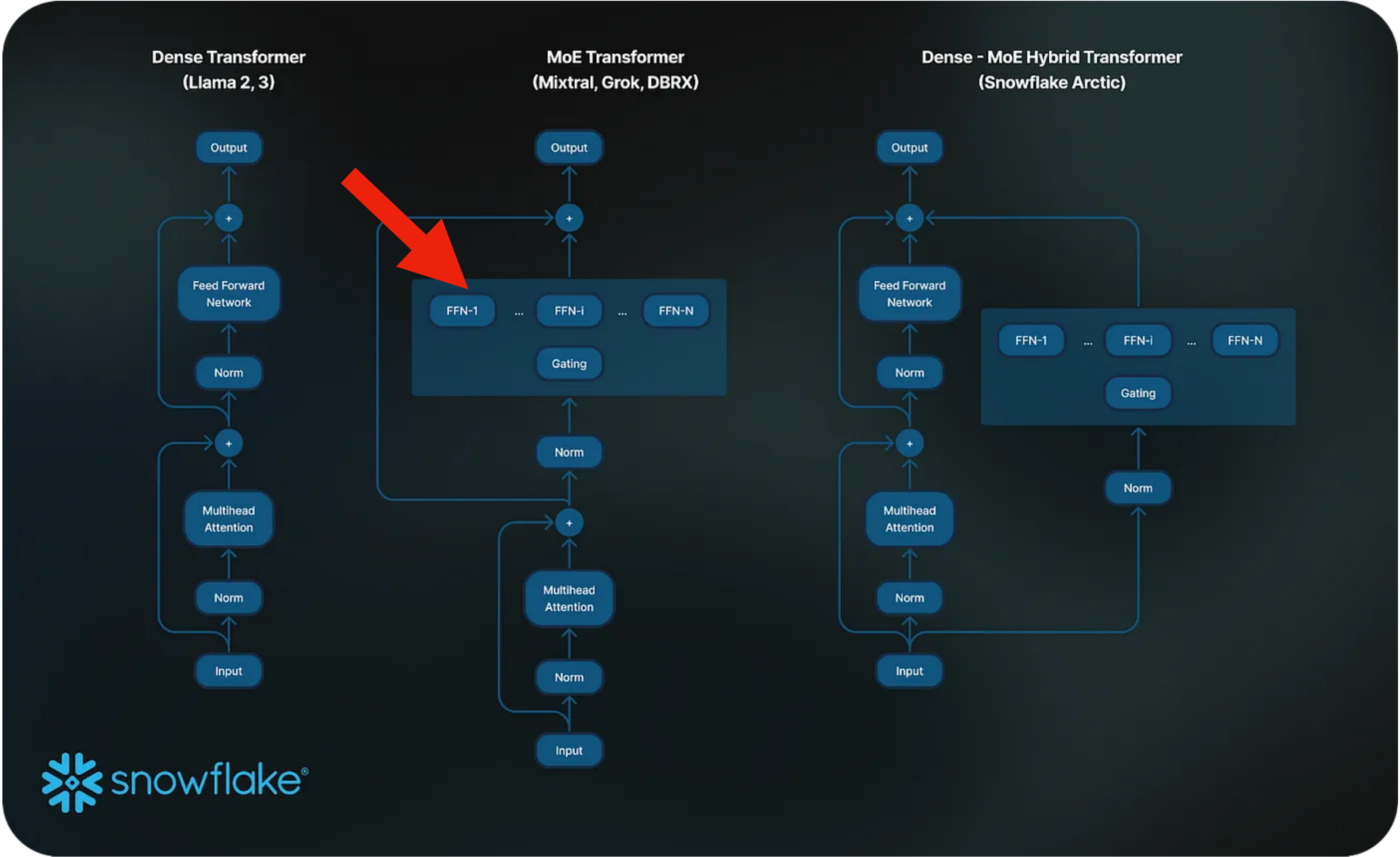
DYLAN PATEL AND GERALD WONG

JUL 11, 2023 • PAID

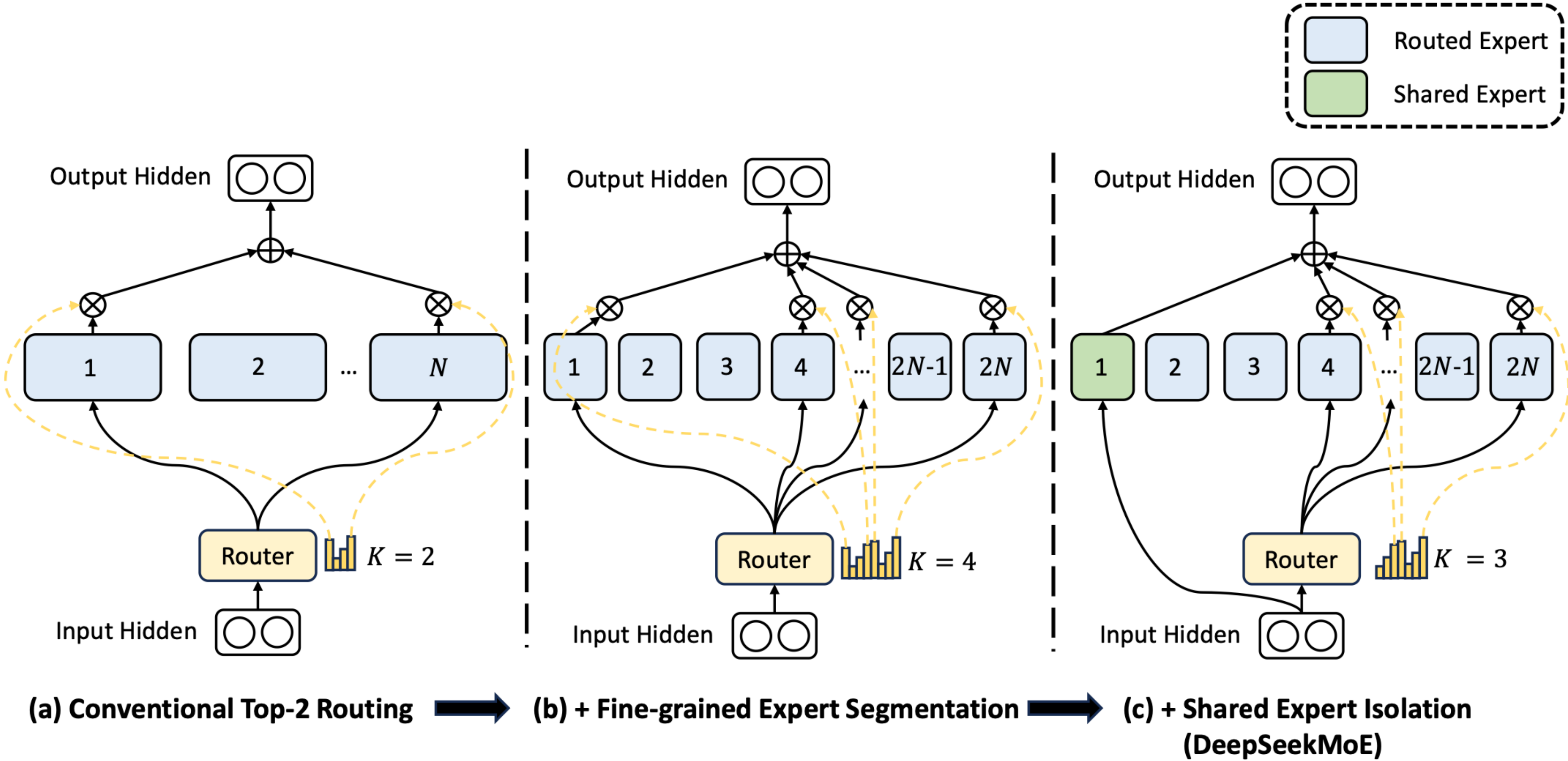
# Open-weight MoE

|              | Org        | #Experts       | Active Params / Total Params |
|--------------|------------|----------------|------------------------------|
| Mixtral 8x7B | Mistral    | 8              | 13B / 56B                    |
| Grok-1       | X.AI       | 8              | 86B / 314B                   |
| DBRX         | Databricks | 16             | 36B / 132B                   |
| Arctic       | Snowflake  | 128            | 17B / 480B                   |
| DeepSeekMoE  | Deepseek   | 162 (2 shared) | 21B / 236B                   |

# Dense vs MoE



# 共享 Expert



# Attention 也可以 MoE?

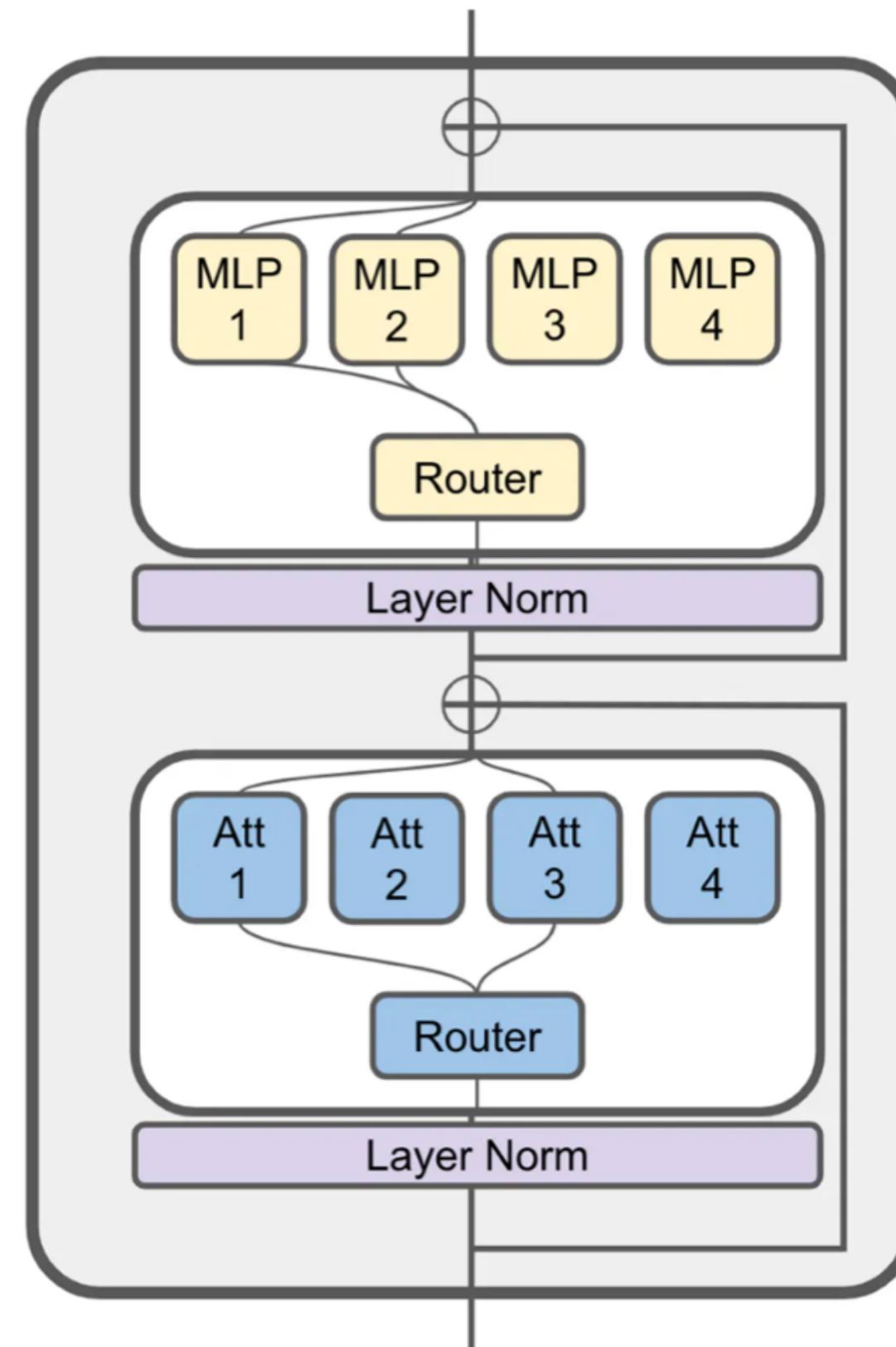


Figure 1: JetMoE architecture

# Attention 也可以 MoE?

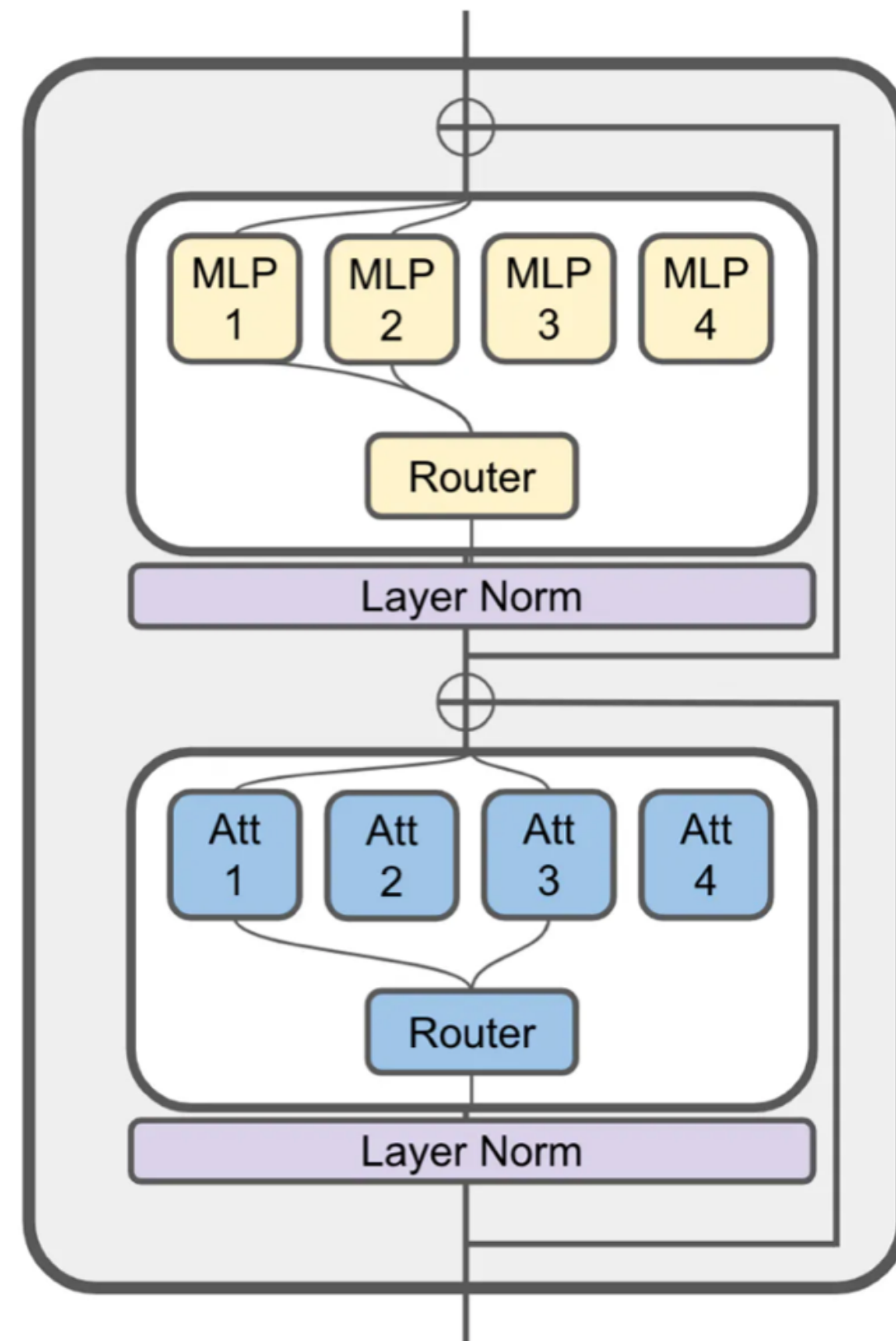


Figure 1: JetMoE architecture

## Llama 3.1 405B 參數分佈

| Tensors                                        | Shape             | Precision |
|------------------------------------------------|-------------------|-----------|
| model                                          |                   |           |
| model.embed_tokens.weight                      | [128 256, 16 384] | BF16      |
| model.layers.0                                 |                   |           |
| model.layers.0.input_layernorm.weight          | [16 384]          | BF16      |
| model.layers.0.mlp.down_proj.weight            | [16 384, 53 248]  | BF16      |
| model.layers.0.mlp.gate_proj.weight            | [53 248, 16 384]  | BF16      |
| model.layers.0.mlp.up_proj.weight              | [53 248, 16 384]  | BF16      |
| model.layers.0.post_attention_layernorm.weight | [16 384]          | BF16      |
| model.layers.0.self_attn.k_proj.weight         | [1 024, 16 384]   | BF16      |
| model.layers.0.self_attn.o_proj.weight         | [16 384, 16 384]  | BF16      |
| model.layers.0.self_attn.q_proj.weight         | [16 384, 16 384]  | BF16      |
| model.layers.0.self_attn.v_proj.weight         | [1 024, 16 384]   | BF16      |

Attn 做 MoE 好處不大

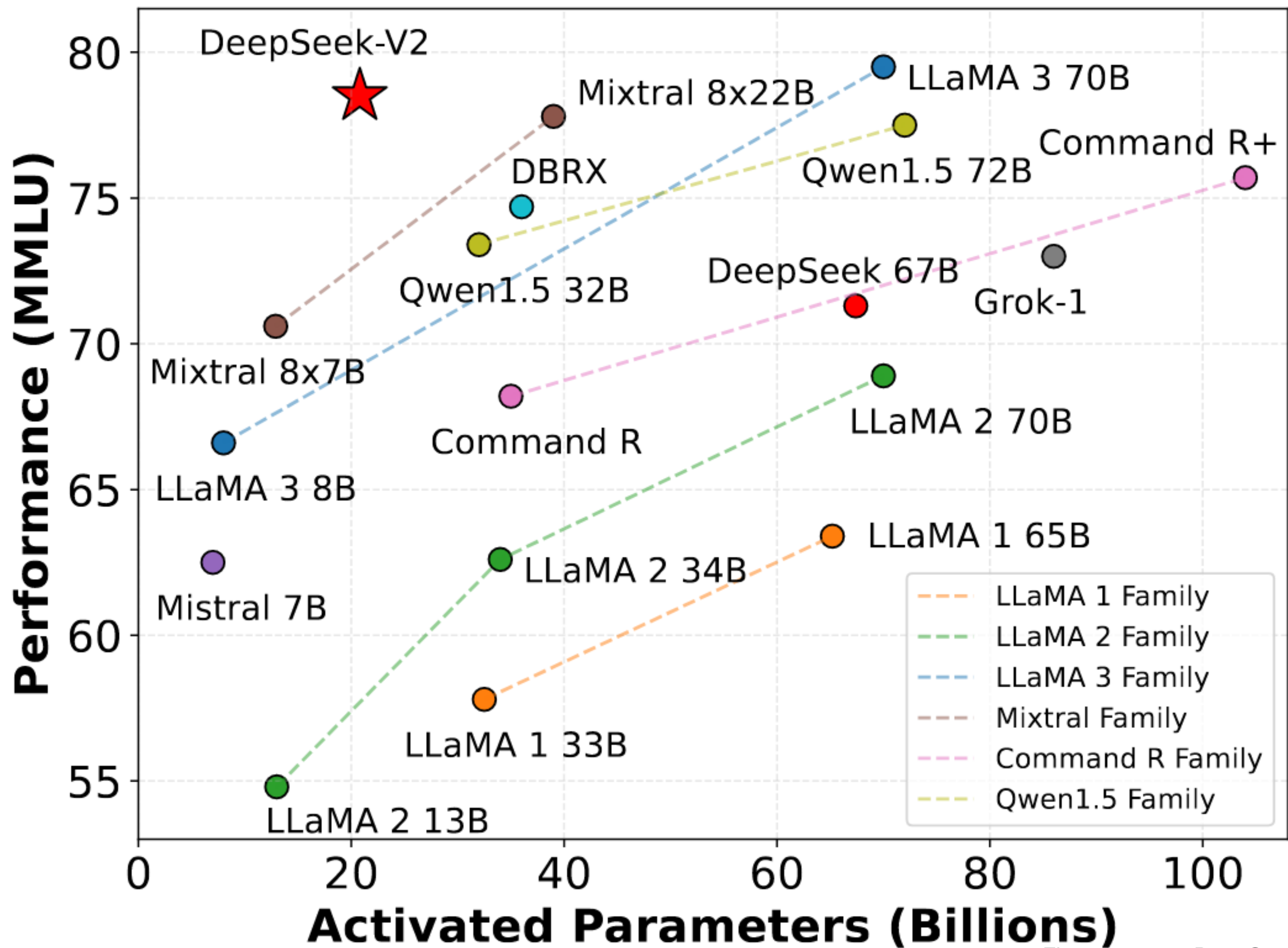


Figure source: DeepSeek

William Fedus\*  
LIAMFEDUS@GOOGLE.COM

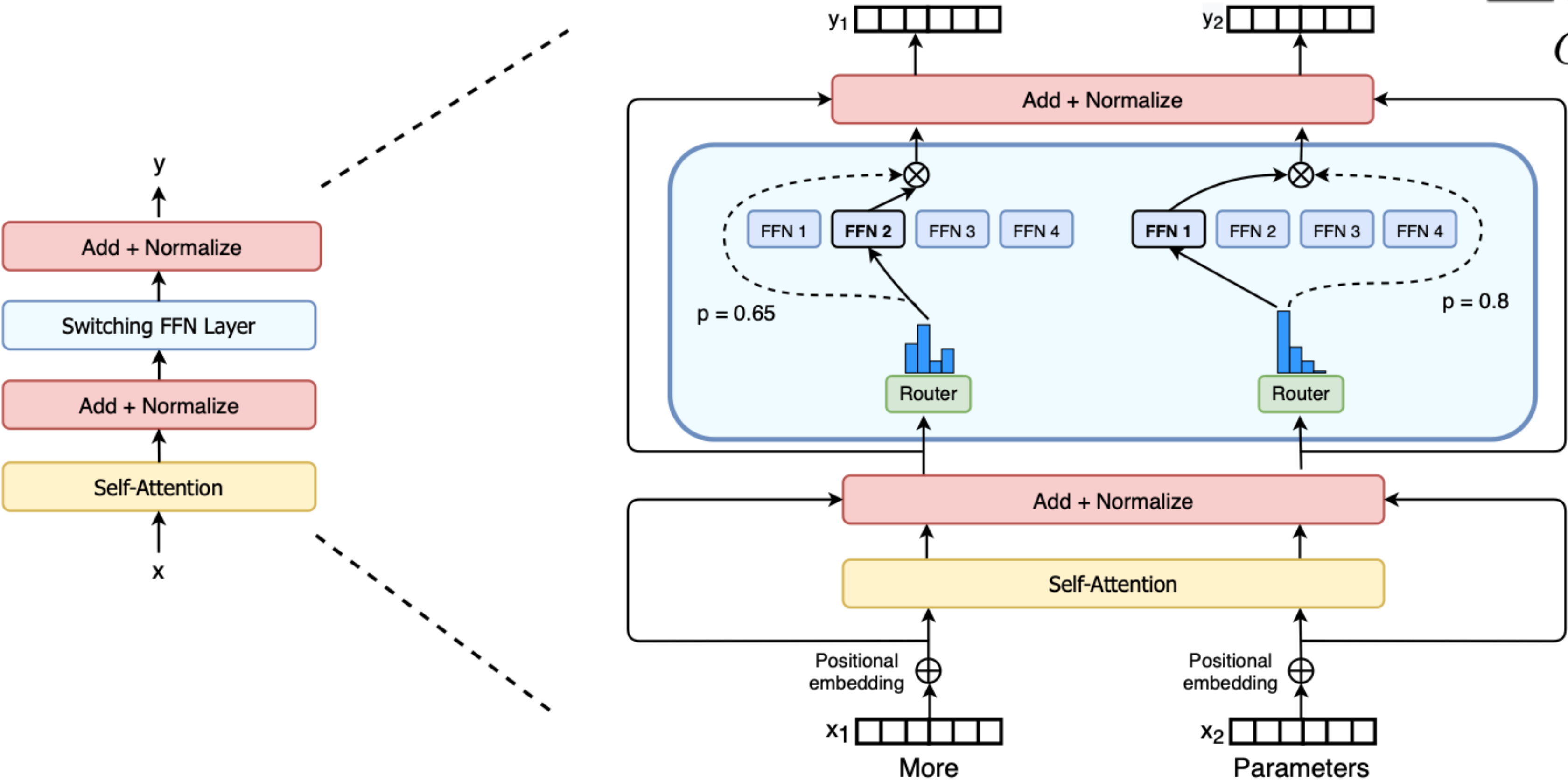
Barret Zoph\*  
BARRETZOPH@GOOGLE.COM

Noam Shazeer  
NOAM@GOOGLE.COM  
*Google, Mountain View, CA 94043, USA*

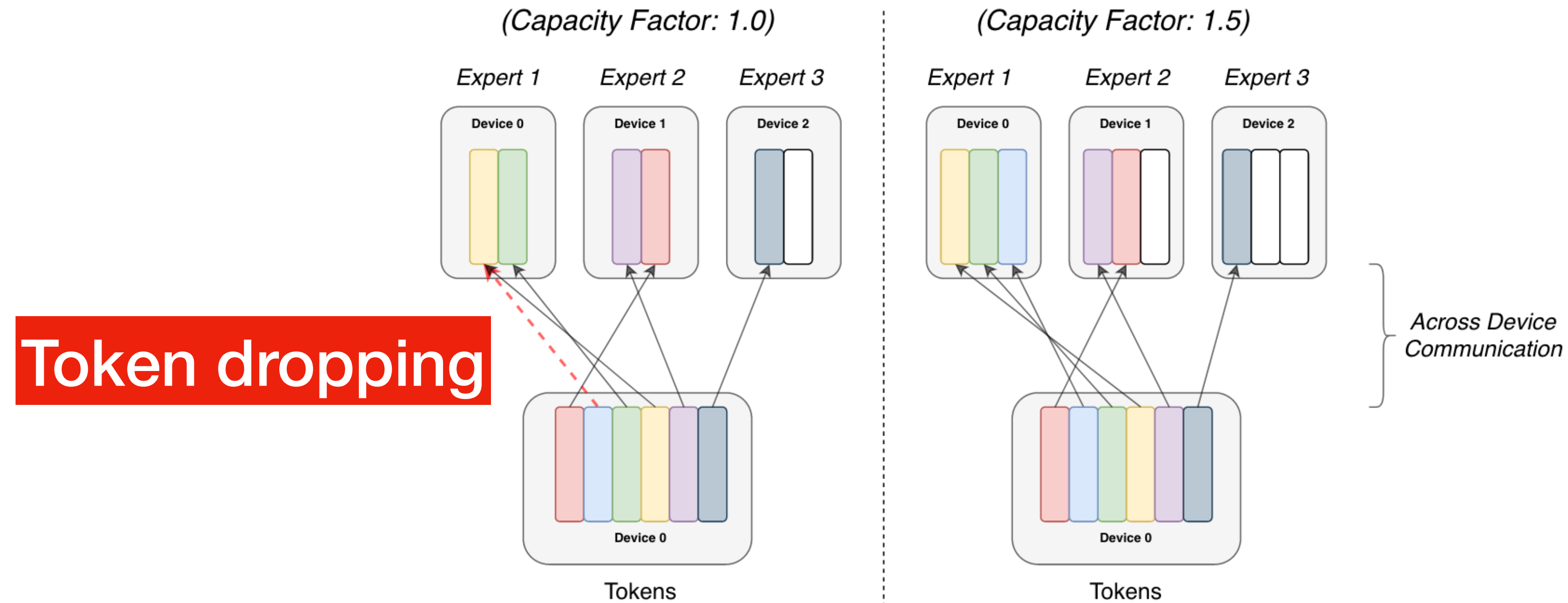
# Token 先決

## SWITCH TRANSFORMERS

$$\sum_{e=1}^E \mathcal{G}_{s,e} \cdot \text{FFN}_e(x_s)$$
$$G(x) := \text{Softmax}(\text{TopK}(x \cdot W_g)),$$



# Token 塞車



$$expert\_capacity = \frac{num\_tokens}{num\_experts} \times capacity\_factor$$

# 負載平衡

$$\text{loss} = \alpha \cdot N \cdot \sum_{i=1}^N f_i \cdot P_i \quad (4)$$

where  $f_i$  is the fraction of tokens dispatched to expert  $i$ ,

$$f_i = \frac{1}{T} \sum_{x \in \mathcal{B}} \mathbb{1}\{\text{argmax } p(x) = i\} \quad (5)$$

and  $P_i$  is the fraction of the router probability allocated for expert  $i$ ,<sup>2</sup>

$$P_i = \frac{1}{T} \sum_{x \in \mathcal{B}} p_i(x). \quad (6)$$

# 負載平衡

|         | Expert 1 | Expert 2 | $f_1 = 2/3$<br>$P_1 = (0.5+0.5+0.5) / 3$<br>$f_2 = 1/3$<br>$P_2 = (0.5+0.5+0.5) / 3$<br><br>$Loss = 2/3 \times 0.5 + 1/3 \times 0.5 = 0.5$ |
|---------|----------|----------|--------------------------------------------------------------------------------------------------------------------------------------------|
| Token 1 | 0.5      | 0.5      |                                                                                                                                            |
| Token 2 | 0.5      | 0.5      |                                                                                                                                            |
| Token 3 | 0.5      | 0.5      |                                                                                                                                            |

|         | Expert 1 | Expert 2 | $f_1 = 2/3$<br>$P_1 = (0.51+0.51) / 3$<br>$f_2 = 1/3$<br>$P_2 = (0.49+0.49+1) / 3$<br><br>$Loss = 0.447$ |
|---------|----------|----------|----------------------------------------------------------------------------------------------------------|
| Token 1 | 0.51     | 0.49     |                                                                                                          |
| Token 2 | 0.51     | 0.49     |                                                                                                          |
| Token 3 | 0        | 1        |                                                                                                          |

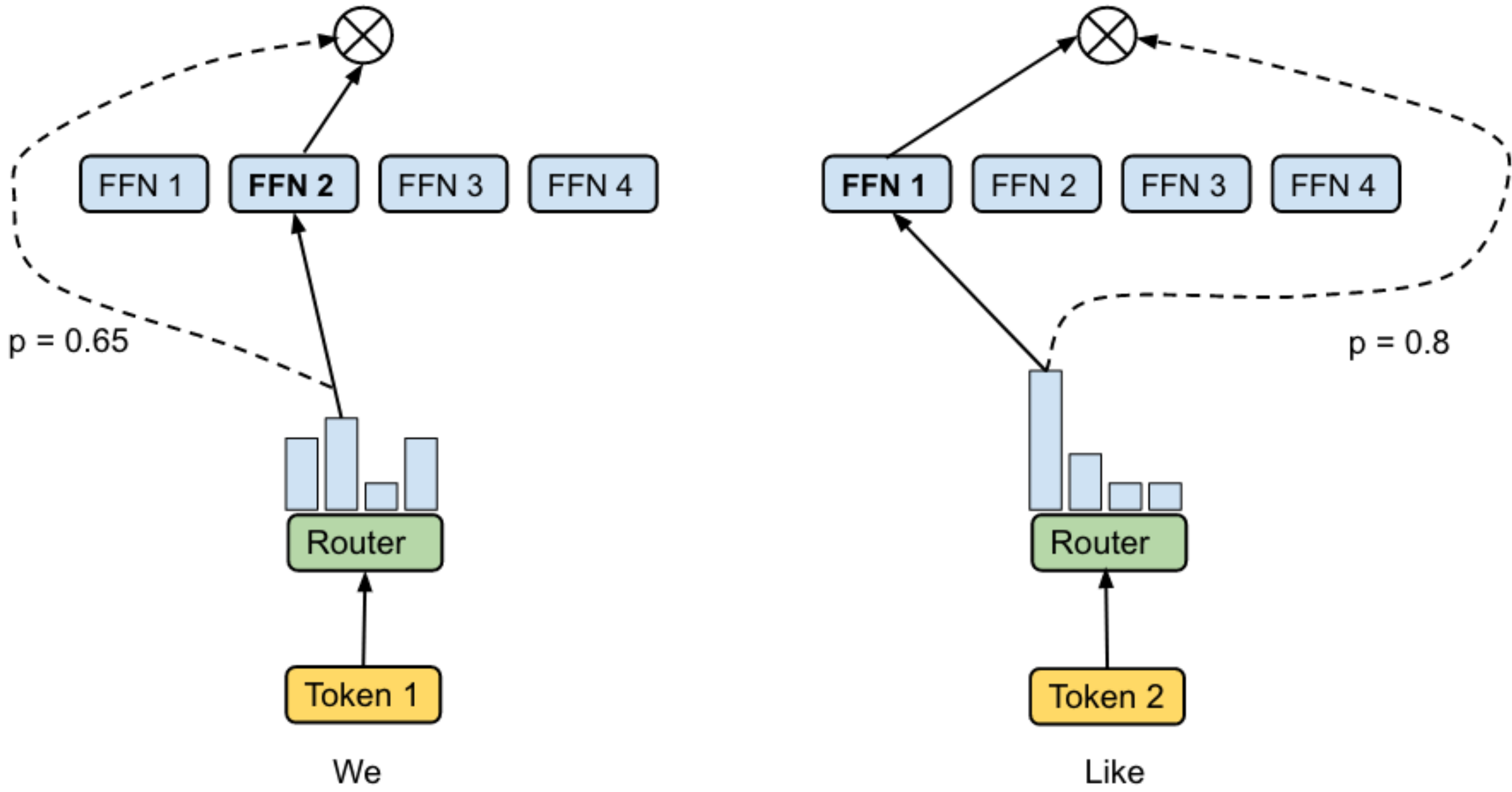
# Token 先決的問題

- 負載不平衡 / Balancing Loss 不穩定
- 每一個 token 用相同的 FLOPs

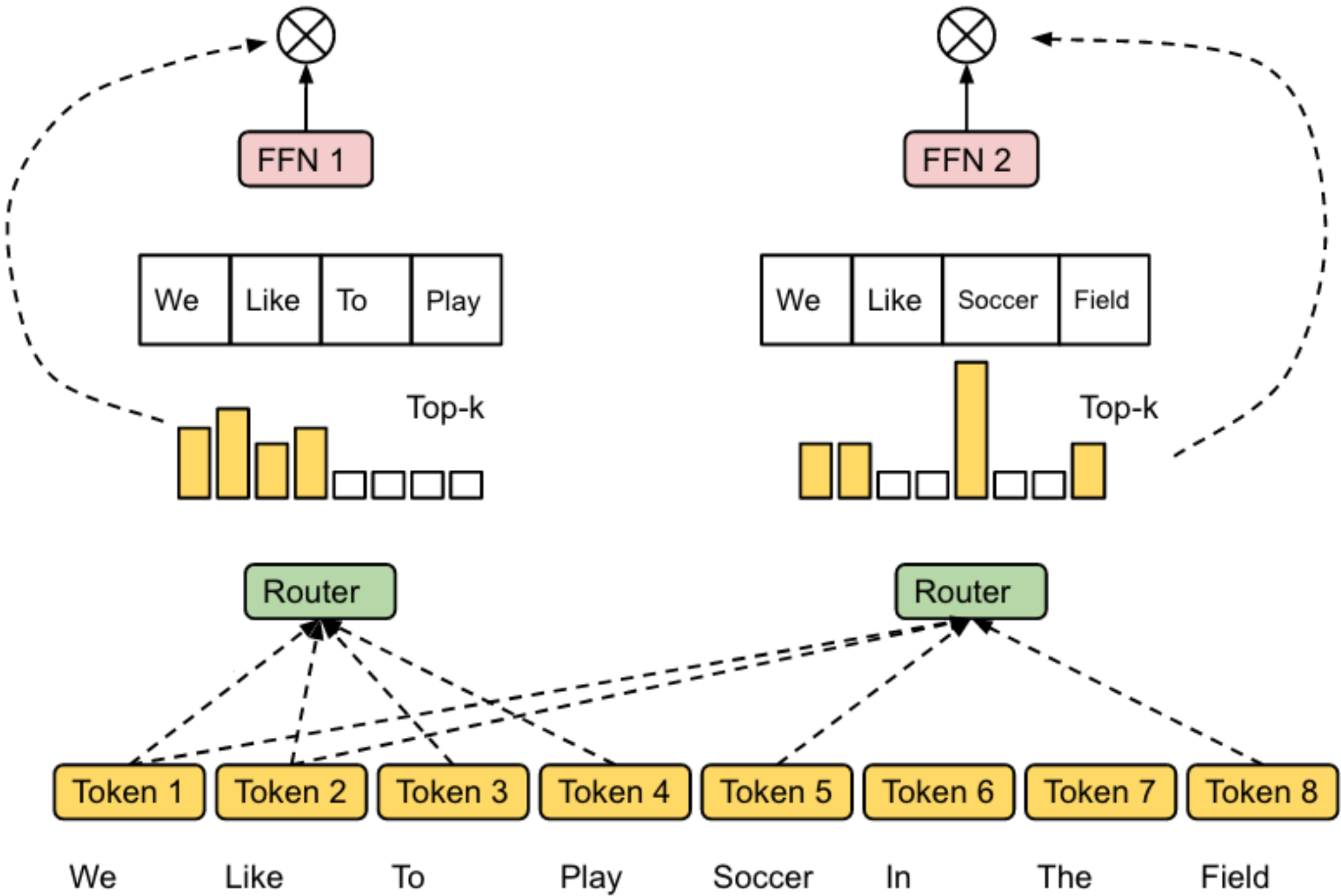
(Expert-choice)

(Mixture of Depth)

# Token 先決



# Expert 先決



# 效果顯著...嗎？

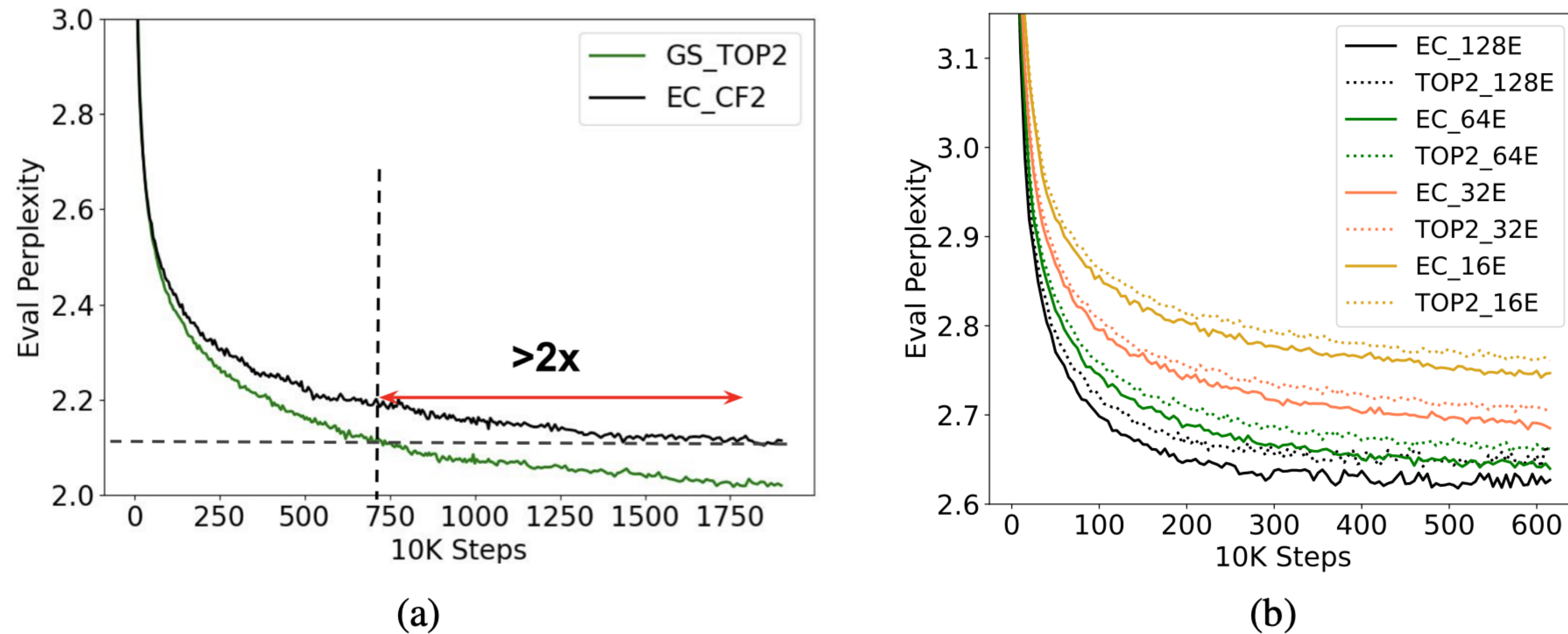


Figure 2: (a) Training convergence is more than 2x faster using our method compared to GShard top-2 gating. (b) Training perplexity scales strongly with the number of experts while keeping the expert size fixed. EC consistently outperforms GShard top-2 gating.

# 效果顯著...嗎？

## 6 Limitations

The expert choice method might not immediately apply to auto-regressive text generation as our current implementation takes in the past and future tokens to perform the top- $k$  selection. One possible solution is to collect a large batch of input sequences, dispatch tokens of the same sequence into separate groups, and perform expert choice routing for each group. Another scenario where the expert choice method does not immediately apply is when the batch size becomes very small during serving or inference. A global top- $k$  can be selected instead and we can cap the number of times each expert or token gets selected. We leave these possible improvements for future work.

**需要多解決一個問題：訓練與推論時 router 的輸入不同**

# Token 先決的問題

- 負載不平衡 / Balancing Loss 不穩定
- 每一個 token 用相同的 FLOPs

(Expert-choice)

(Mixture of Depth)

# 跳層

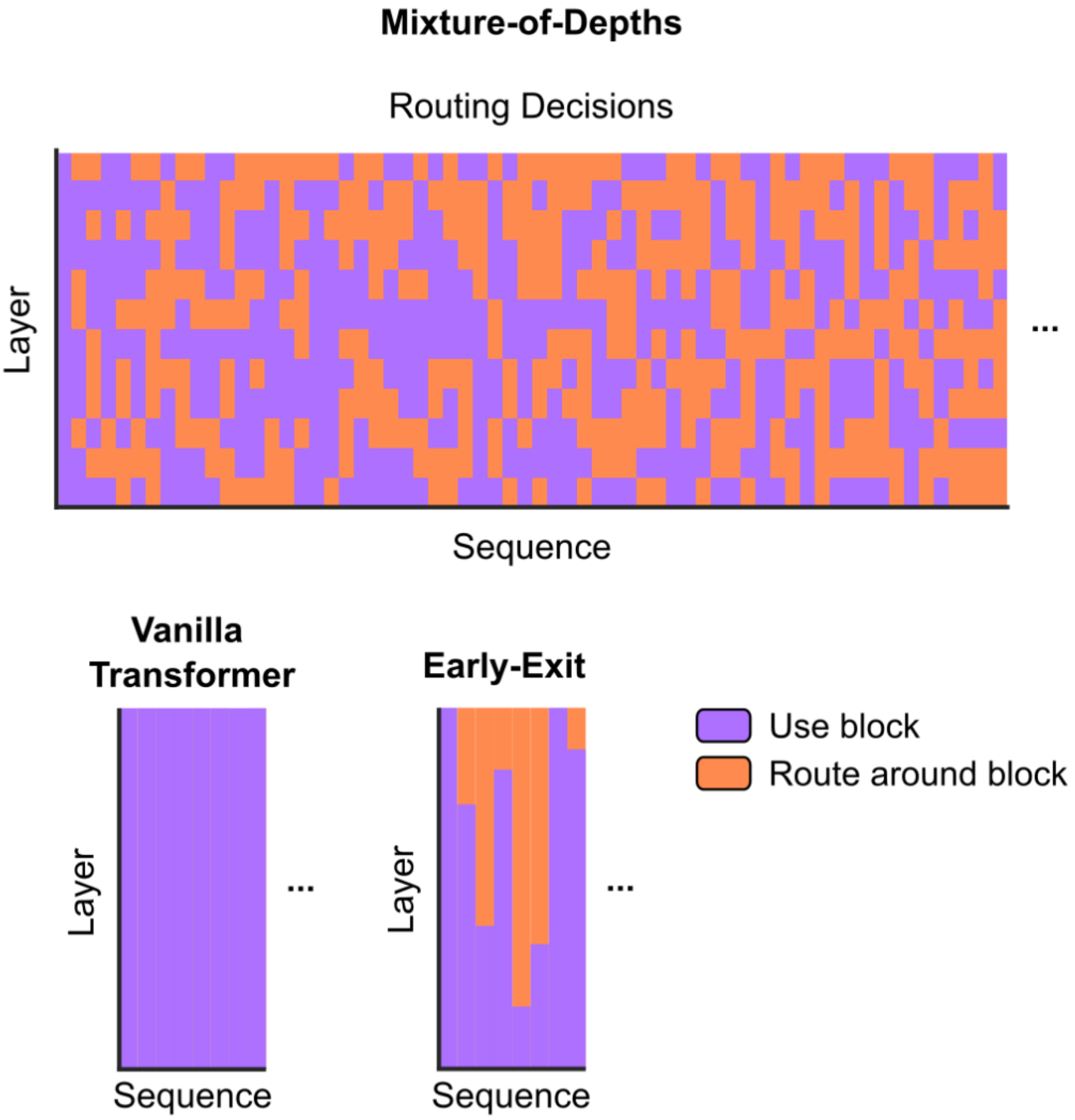
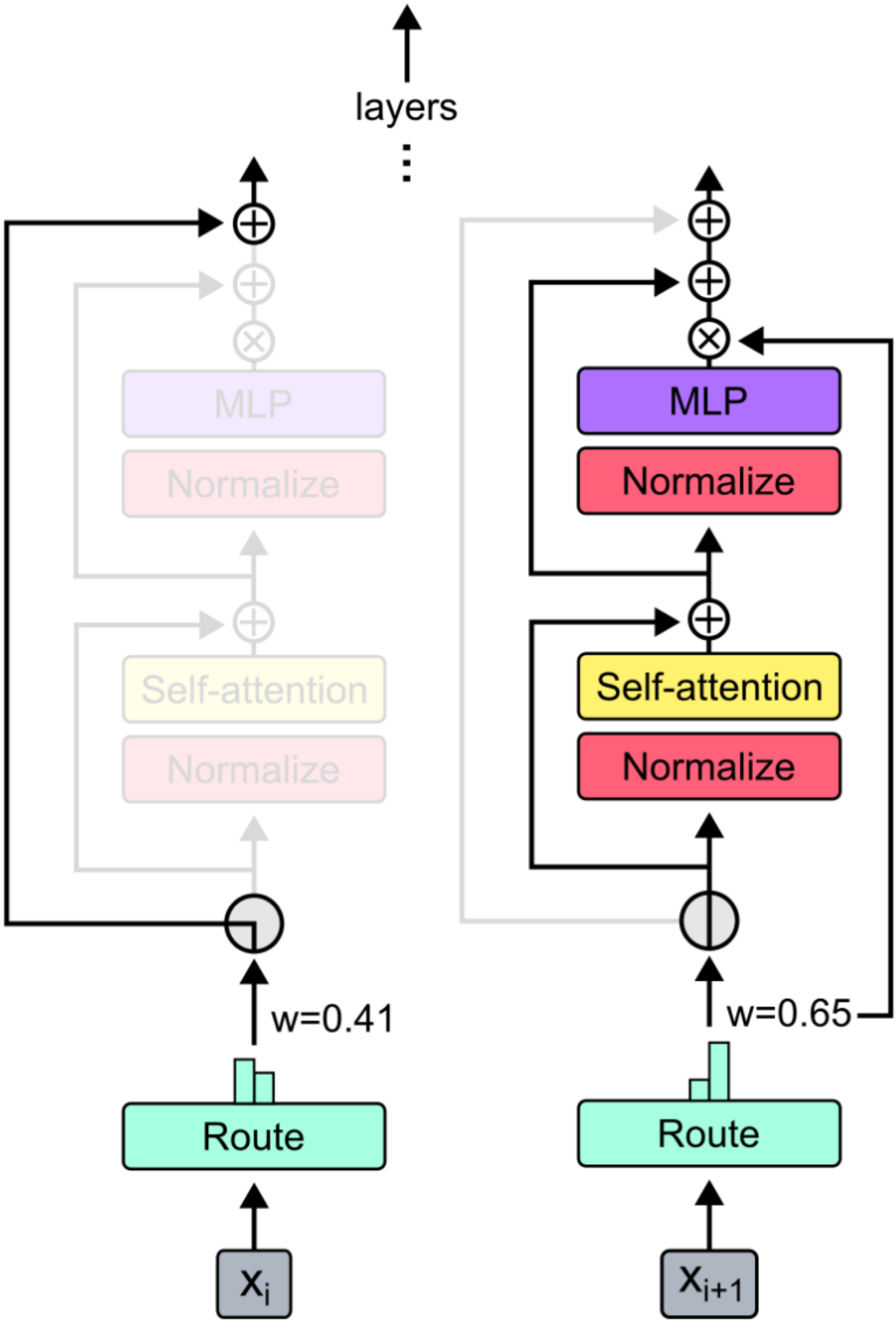


Figure source: Mixture-of-Depths: Dynamically allocating compute in transformer-based language models