

Efficient LLMs

Advances in Large Language Models

ELL8299 · AIL861, Semester 1, 2025-26



Yatin Nandwani
Research Scientist, IBM Research

Training Vs Inference in LLMs



$$P_{\theta}(x)$$

$$P_{\theta}(\text{class} | x)$$

$$P_{\theta}(y | x) \quad \forall y \in \mathcal{Y}$$

$$P_{\theta}(X = [\underbrace{x_1}_{50,000}, \dots, x_T]_{\mathcal{V}})$$

$$50,000 = |\mathcal{V}|$$

$$P_{\theta}(x) \uparrow \text{ if } x \text{ is a valid sentence}$$

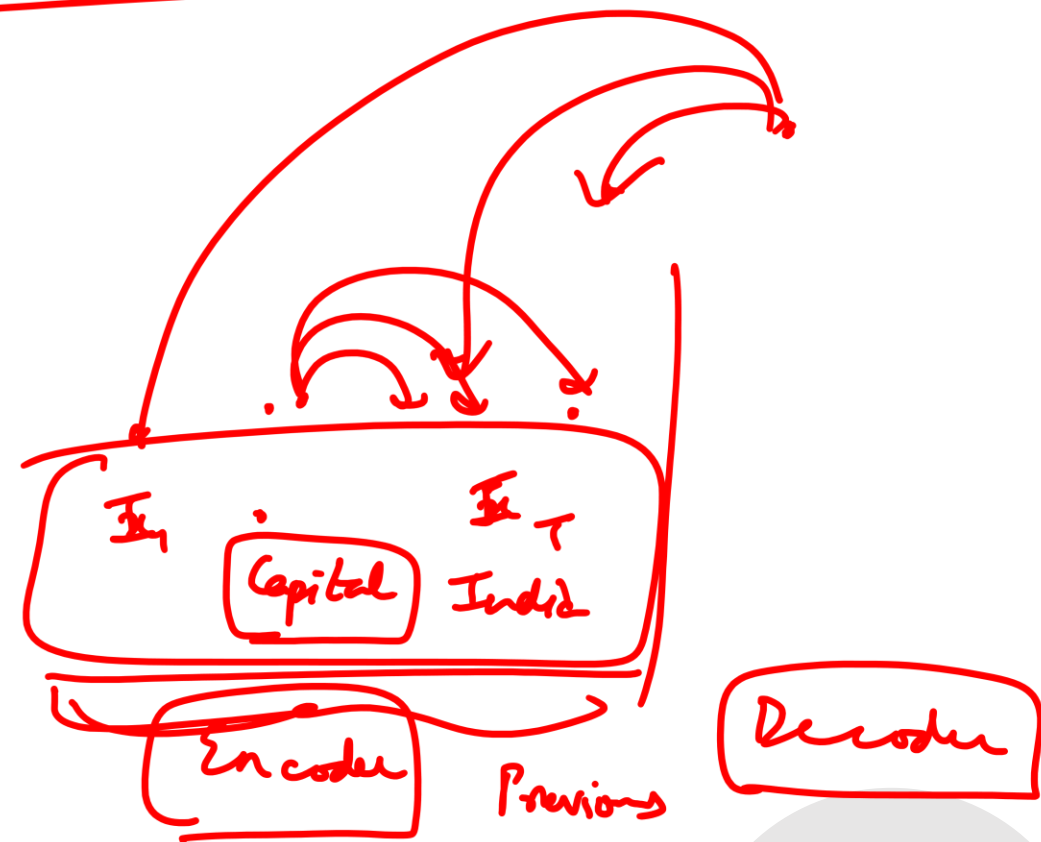
$$P_{\theta}(x) = P_{\theta}(x_1) P_{\theta}(x_2 | x_1) P_{\theta}(x_3 | x_1, x_2) \dots P_{\theta}(x_T | x_1, \dots, x_{T-1})$$

Diagram illustrating the sequence of tokens $x_1, x_2, \dots, x_T$ and the corresponding conditional probabilities. A box labeled x contains a token $*$. A sequence of tokens $x_1, x_2, \dots, x_T$ is shown below, with arrows indicating the flow of information from previous tokens to the current token's probability.



$$\log P_{\theta}(x) = \log P_{\theta}(x_1 | [\text{Bos}]) + \log P_{\theta}(x_2 | [\text{Bos}], x_1) + \dots + \log P_{\theta}(x_T | \dots)$$

$$P_{\theta}(x_T) P_{\theta}(x_H | x_{T-1})$$

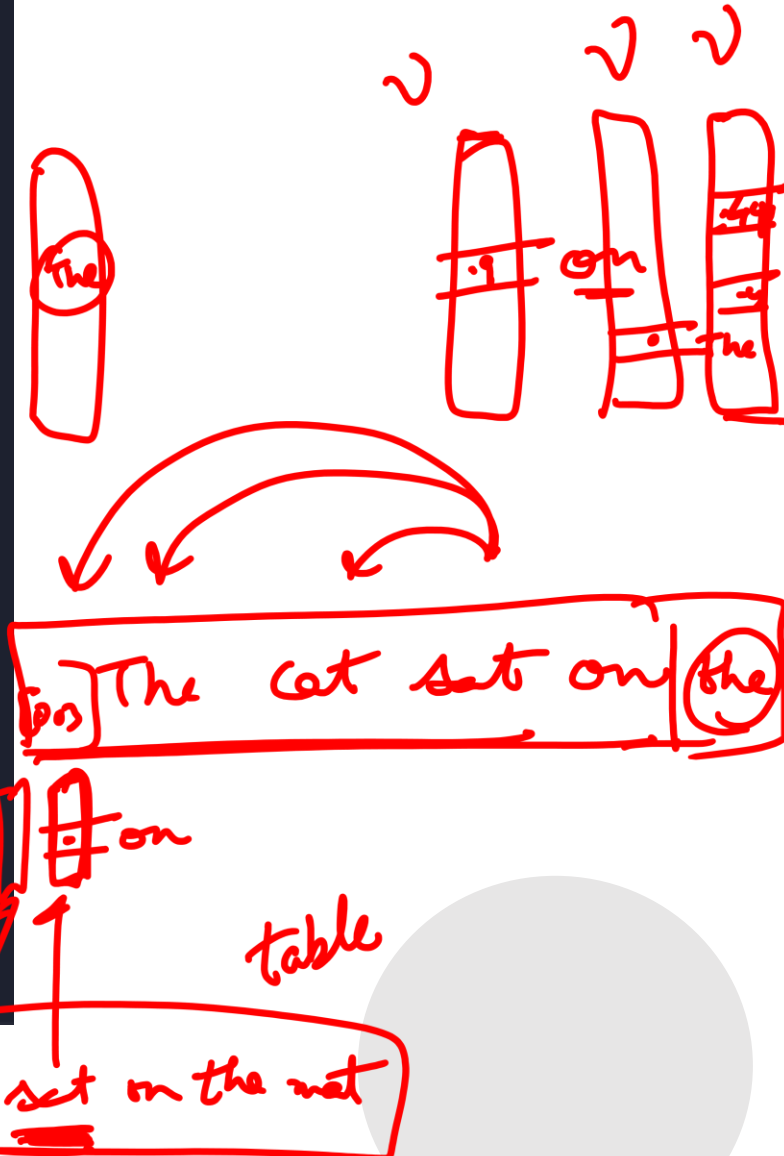


Forward Pass

```

3 class ToyLM(nn.Module):
4     def __init__(self, vocab_size, hidden_dim, num_layers):
5         super().__init__()
6         self.embed = nn.Embedding(vocab_size, hidden_dim) ✓
7         self.layers = nn.ModuleList([
8             TransformerLayer(hidden_dim)
9             for _ in range(num_layers)
10        ])
11        self.lm_head = nn.Linear(hidden_dim, vocab_size)
12
13    def forward(self, input_ids):
14        x = self.embed(input_ids) # input embeddings
15        for layer in self.layers: # stack of transformer layers
16            x = layer(x)
17        logits = self.lm_head(x) # output LM head
18        return logits

```



Training Loop

```
1 import torch
2 model = ToyLM(vocab_size, hidden_dim, num_layers)
3 optimizer = torch.optim.AdamW(model.parameters())
4
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
```

$p_\theta(x_1), p_\theta(x_2|x_1)$
1 -- $p_\theta(x_i|x_{1:i-1})$



Training Loop

```
1 import torch
2 model = ToyLM(vocab_size, hidden_dim, num_layers)
3 optimizer = torch.optim.AdamW(model.parameters())
4
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
```

Input: The cat sat on the mat
↓
Label: cat



Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1)
10    )
11    loss.backward()
12    optimizer.step()
13    optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
11 loss.backward()
12 optimizer.step()
13 optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6     for _ in range(max_len):
7         logits = model(input_ids)
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
11 loss.backward()
12 optimizer.step()
13 optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6
7     logits = model(input_ids)
8     next_token = torch.argmax(logits[:, -1, :],
9                               dim=-1)
10 )
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
11 loss.backward()
12 optimizer.step()
13 optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6
7     logits = model(input_ids)
8     next_token = torch.argmax(logits[:, -1, :],
9                               dim=-1)
10
11     if next_token[0] == "</s>":
12         break
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10
11     loss.backward()
12     optimizer.step()
13     optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6
7     logits = model(input_ids)
8     next_token = torch.argmax(logits[:, -1, :],
9                               dim=-1)
10
11     if next_token[0] == "</s>":
12         break
13     input_ids = torch.cat([
14         input_ids,
15         next_token.unsqueeze(-1)], dim=-1)
16
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10
11     loss.backward()
12     optimizer.step()
13     optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6     for _ in range(max_len):
7         logits = model(input_ids)
8         next_token = torch.argmax(logits[:, -1, :],
9                                   dim=-1)
10    )
11    if next_token[0] == "</s>":
12        break
13    input_ids = torch.cat([
14        input_ids,
15        next_token.unsqueeze(-1)], dim=-1)
16    )
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
11 loss.backward()
12 optimizer.step()
13 optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference

```
1 import torch
2
3 def generate(model, input_string):
4     input_ids = tokenizer.encode(input_string)
5     #input_id_tokens: ["<s>", "The", "cat", "sat"]
6     for _ in range(max_len):
7         logits = model(input_ids)
8         next_token = torch.argmax(logits[:, -1, :],
9                                   dim=-1)
10    )
11    if next_token[0] == "</s>":
12        break
13    input_ids = torch.cat([
14        input_ids,
15        next_token.unsqueeze(-1)], dim=-1)
16    )
17
18 input_string = "The cat sat"
19 model = ToyLM(vocab_size, hidden_dim, num_layers)
20 output_ids = generate(input_string, model)
21 print(tokenizer.decode(input_ids[0]))
```

Training Loop

```
5 for input_ids, labels in dataloader:
6     logits = model(input_ids)
7     loss = F.cross_entropy(
8         logits.view(-1, vocab_size),
9         labels.view(-1))
10 )
11 loss.backward()
12 optimizer.step()
13 optimizer.zero_grad()
```

Forward Pass

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference w/ caching

Forward Pass w/ caching

Forward Pass w/o caching

```
13 def forward(self, input_ids):  
14     x = self.embed(input_ids)  
15     for layer in self.layers:  
16         x = layer(x)  
17     logits = self.lm_head(x)  
18     return logits
```



Inference w/ caching

Forward Pass w/ caching

```
24 def forward(self, input_ids, cache=None):
25     x = self.embed(input_ids)
26
27     new_cache = []
28     for i, layer in enumerate(self.layers):
29         # Each layer consumes (x, past_kv)
30         # and returns (x, new_kv)
31         past_kv = cache[i] if cache else None
32         x, kv = layer(x, past_kv)
33         new_cache.append(kv)
34
35     logits = self.lm_head(x)
36     return logits, new_cache
```

Forward Pass w/o caching

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference w/ caching

Forward Pass w/ caching

```
24 def forward(self, input_ids, cache=None):
25     x = self.embed(input_ids)
26
27     new_cache = []
28     for i, layer in enumerate(self.layers):
29         # Each layer consumes (x, past_kv)
30         # and returns (x, new_kv)
31         past_kv = cache[i] if cache else None
32         x, kv = layer(x, past_kv)
33         new_cache.append(kv)
34
35     logits = self.lm_head(x)
36     return logits, new_cache
```

Forward Pass w/o caching

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference w/ caching

```
38 def generate(model, input_string, max_len, eos_token_id):
39     # 1. Prefill: run full prompt once
40     input_ids = tokenizer.encode(input_string).unsqueeze(0)
41     logits, cache = model(input_ids, cache=None)
42
```

Forward Pass w/ caching

```
24 def forward(self, input_ids, cache=None):
25     x = self.embed(input_ids)
26
27     new_cache = []
28     for i, layer in enumerate(self.layers):
29         # Each layer consumes (x, past_kv)
30         # and returns (x, new_kv)
31         past_kv = cache[i] if cache else None
32         x, kv = layer(x, past_kv)
33         new_cache.append(kv)
34
35     logits = self.lm_head(x)
36     return logits, new_cache
```

Forward Pass w/o caching

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Inference w/ caching

```
38 def generate(model, input_string, max_len, eos_token_id):
39     # 1. Prefill: run full prompt once
40     input_ids = tokenizer.encode(input_string).unsqueeze(0)
41     logits, cache = model(input_ids, cache=None)
42
43     # 2. Iterative decoding
44     for _ in range(max_len):
45         next_token = argmax(logits[:, -1, :])
46         input_ids = cat([input_ids,
47                         next_token.unsqueeze(-1)], dim=-1)
48
49         # --- Early stop if EOS is generated ---
50         if next_token.item() == eos_token_id:
51             break
52
53         # feed only the last token with cache
54         logits, cache = model(input_ids[:, -1:], cache)
55
```

Forward Pass w/ caching

```
24 def forward(self, input_ids, cache=None):
25     x = self.embed(input_ids)
26
27     new_cache = []
28     for i, layer in enumerate(self.layers):
29         # Each layer consumes (x, past_kv)
30         # and returns (x, new_kv)
31         past_kv = cache[i] if cache else None
32         x, kv = layer(x, past_kv)
33         new_cache.append(kv)
34
35     logits = self.lm_head(x)
36     return logits, new_cache
```

Forward Pass w/o caching

```
13 def forward(self, input_ids):
14     x = self.embed(input_ids)
15     for layer in self.layers:
16         x = layer(x)
17     logits = self.lm_head(x)
18     return logits
```



Blockwise Self-Attention Computation



The	key	difference	between	LLM	training	and	inference
K1	K2	K3	K4	K5	K6	K7	K8
							Q8 ✓

V1	The
V2	key
V3	difference
V4	between
V5	LLM
V6	training
V7	and
V8	inference



The	key	difference	between	LLM	training	and	inference
K1	K2	K3	K4	K5	K6	K7	K8
							Q8

V1	The
V2	key
V3	difference
V4	between
V5	LLM
V6	training
V7	and
V8	inference

$$\frac{\exp(q_8 * k_1) * v_1 + \dots + \exp(q_8 * k_4) * v_4 + \exp(q_8 * k_5) * v_5 + \dots + \exp(q_8 * k_8) * v_8}{\exp(q_8 * k_1) + \dots + \exp(q_8 * k_4) + \exp(q_8 * k_5) + \dots + \exp(q_8 * k_8)}$$



The	key	difference	between	LLM	training	and	inference
K1	K2	K3	K4	K5	K6	K7	K8
							Q8

$$\frac{\sum_{i=1..8} \exp(q_8 * k_i) * v_i}{\sum_{i=1..8} \exp(q_8 * k_i)}$$

V1	The
V2	key
V3	difference
V4	between
V5	LLM
V6	training
V7	and
V8	inference



The	key	difference	between	LLM	training	and	inference
K1	K2	K3	K4	K5	K6	K7	K8

$$\frac{\sum_{i=1..8} \exp(q_8 * k_i) * v_i}{\sum_{i=1..8} \exp(q_8 * k_i)}$$

inf.	The	key	diff.	b/w
Q8	K1	K2	K3	K4

s1	s2	s3	s4
----	----	----	----

a1	a2	a3	a4	V1
				V2
				V3
				V4

inf.	LLM	tr.	and	inf.
Q8	K5	K6	K7	K8

s5	s6	s7	s8
----	----	----	----

a5	a6	a7	a8	V5
				V6
				V7
				V8

V1	The
V2	key
V3	difference
V4	between
V5	LLM
V6	training
V7	and
V8	inference

$$l^{(1)} = \sum_{i=1..4} \exp(S_i)$$

$$o^{(1)} = \frac{A^{(1)}V^{(1)}}{l^{(1)}}$$

$$l^{(2)} = \sum_{i=5..8} \exp(S_i)$$

$$o^{(2)} = \frac{A^{(2)}V^{(2)}}{l^{(2)}}$$

$$= \frac{l^{(1)} o^{(1)} + l^{(2)} o^{(2)}}{l^{(1)} + l^{(2)}}$$

$$o = \frac{l^{(1)} o^{(1)} + l^{(2)} o^{(2)}}{l^{(1)} + l^{(2)}}$$

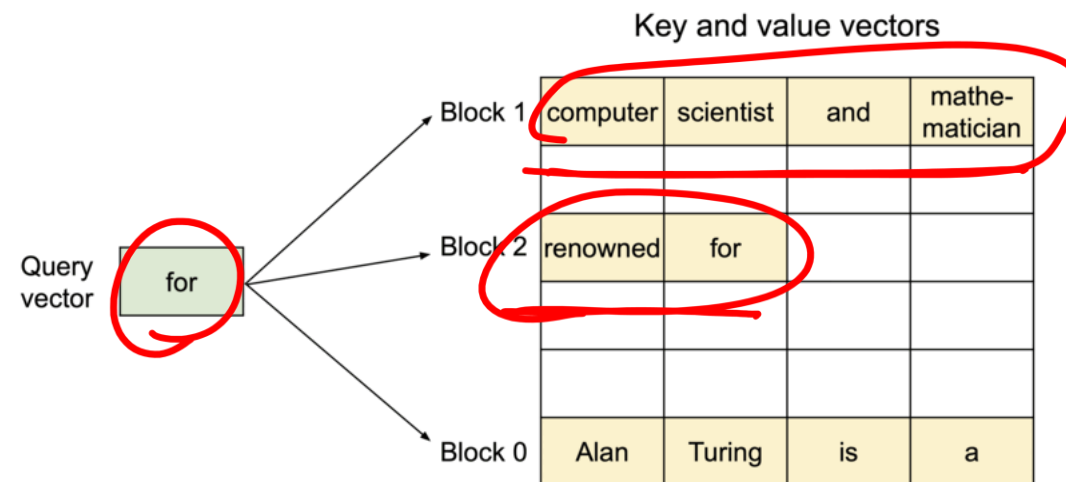
$$l \leftarrow l + l^{(2)}$$

$$o \leftarrow l.o + l^{(2)} o^{(2)}$$



Paged Attention

- Tensor operations require contiguous memory
- How to compute attention *softmax* across fragmented memory?
- Paged Attention!



$$\text{softmax}([A_1, A_2]) = [\alpha \text{softmax}(A_1), \beta \text{softmax}(A_2)]$$

$$\text{softmax}([A_1, A_2]) \begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \alpha \text{softmax}(A_1) * V_1 + \beta \text{softmax}(A_2) * V_2$$



Homework!

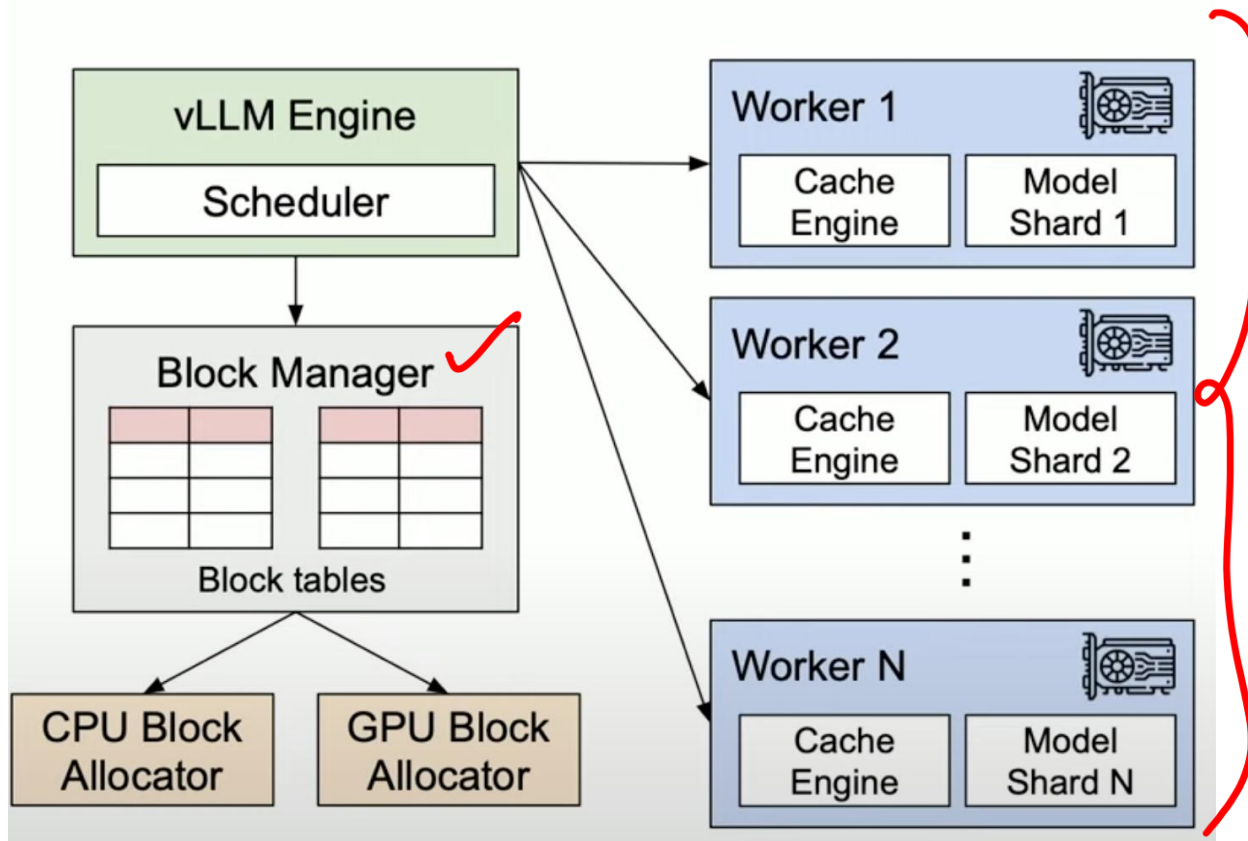
$$\underline{e(10^9)} \quad \underline{10^9} \quad \text{vs} \quad \underline{10^9 + 8}$$

1. Understand the idea of **Safe-Softmax** ✓
2. Read about the concept of **Tiling** used in Matrix Multiplication and Flash Attention

Reference: <https://courses.cs.washington.edu/courses/cse599m/23sp/notes/flashattn.pdf>



System Architecture and Implementation



End to end llm serving engine

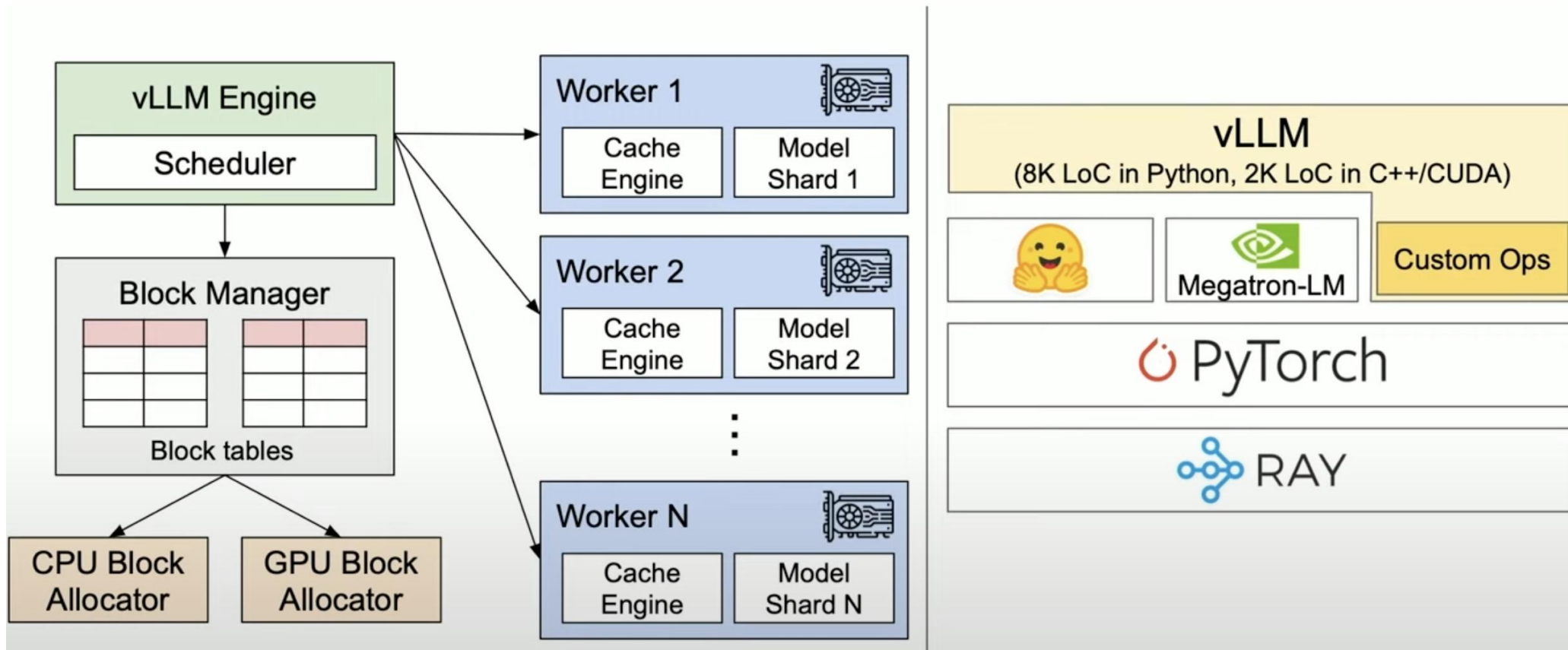
3 components –

- A frontend
- A distributed model executor (TP on single node, PP across nodes)
- A scheduler

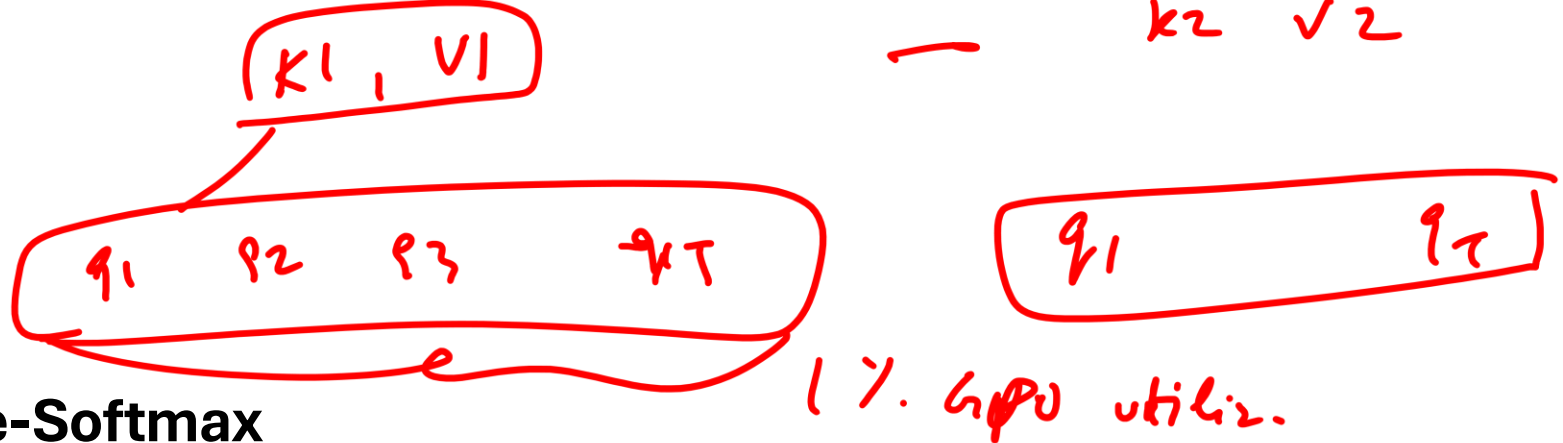
Centralized engine to manage block table

- At each iteration, it sends GPU memory requests to the GPUs;
- Cache engine in the GPU allocates the physical memory blocks





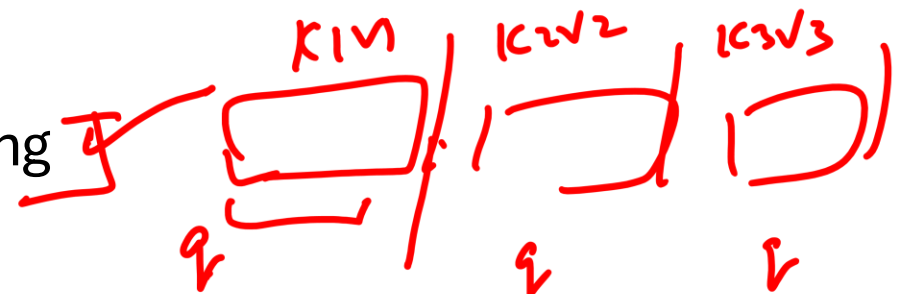
Homework!



1. Understand the idea of **Safe-Softmax**
2. Read about the concept of **Tiling** used in Matrix Multiplication and Flash Attention

Reference: <https://courses.cs.washington.edu/courses/cse599m/23sp/notes/flashattn.pdf>

3. **Flash Decoding** – variant of flash attention for decoding
4. **Speculative decoding** – guess and verify. ✓



Reference: https://www.youtube.com/watch?v=JVUWCv_vMFU&t=1865s&ab_channel=LCS2

o1 o2 o3
v1 v2 v3



So Far...

Efficient Training

1. Parallelism for Scale - ✓
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*
2. Efficient Implementations }
 - *Flash Attention* ✓
 - *Liger Kernels* ✓

Efficient Inference

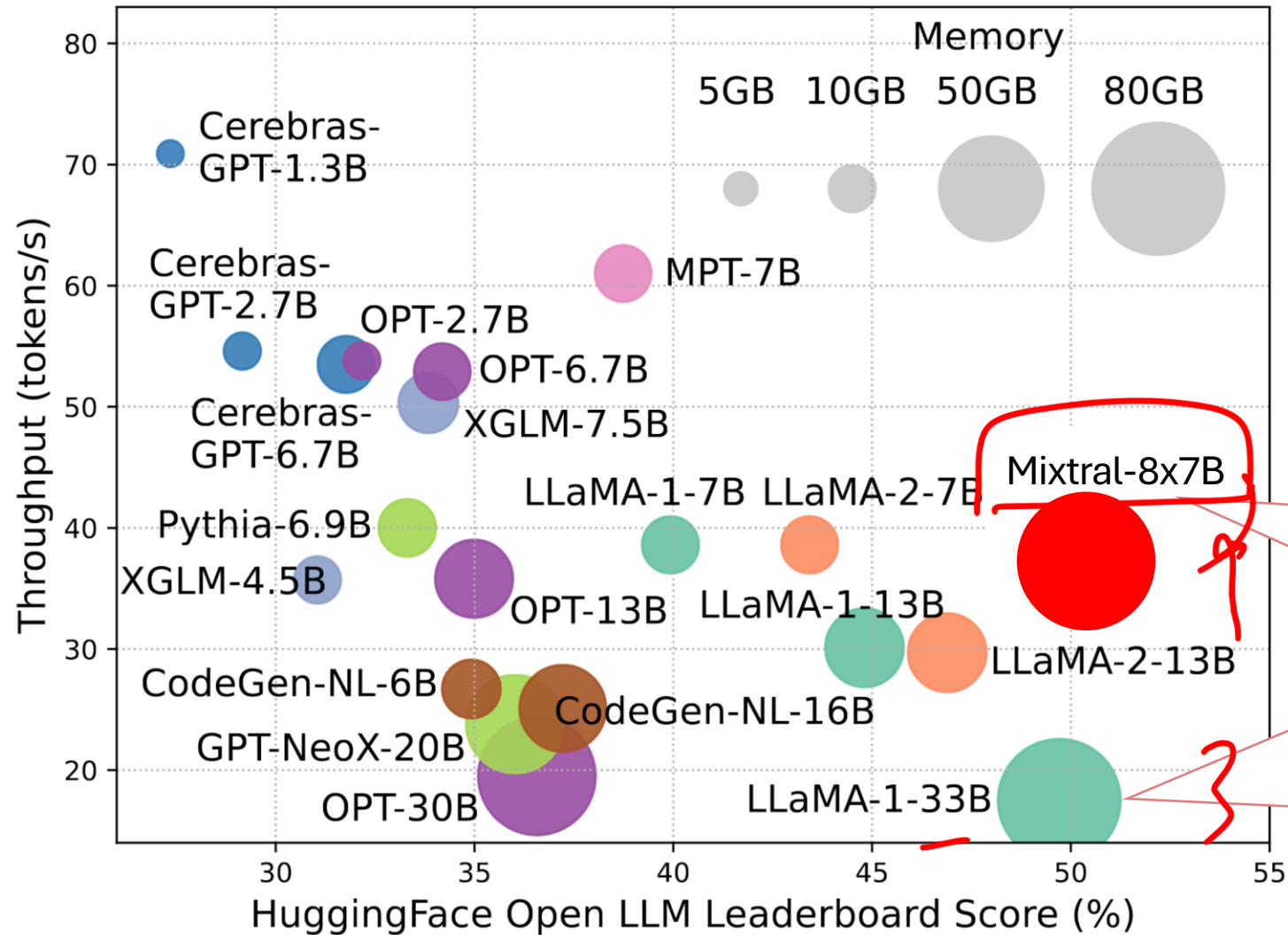
1. KV Caching ✓
 2. vLLMs and paged attention ✓
 3. Flash decoding
 4. Speculative decoding
- Homework!



Inference Throughput vs Performance



Inference Throughput vs Performance



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



So Far...

Efficient Training

1. Parallelism for Scale -
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*
2. Efficient Implementations
 - *Flash Attention*
 - *Liger Kernels*

Efficient Inference

1. KV Caching
2. vLLMs and paged attention
3. Flash decoding
4. Speculative decoding

Homework!

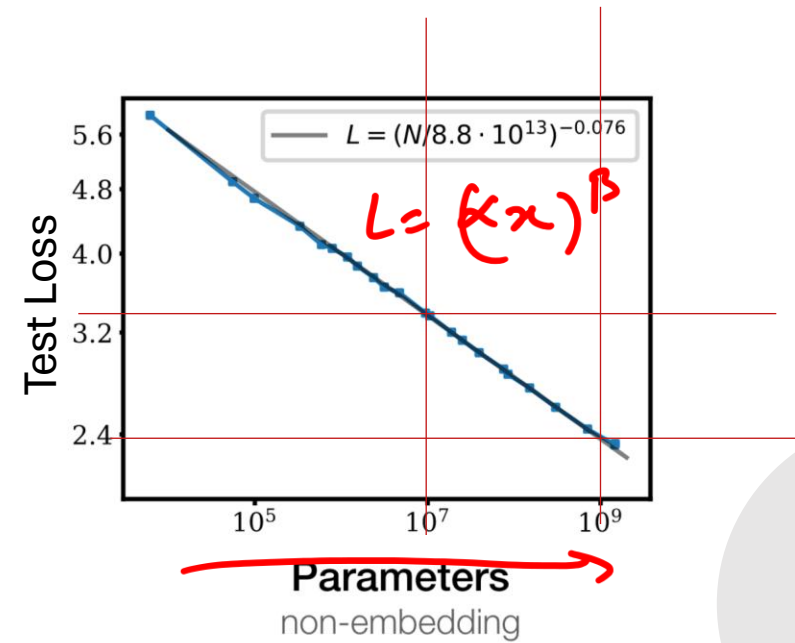
Efficient Design

1. MoE architecture



Neural Scaling Laws Kaplan et al 2020

- Performance improve smoothly as we increase the **model size**
- Larger models → smoother loss curves, better generalization, new emergent abilities

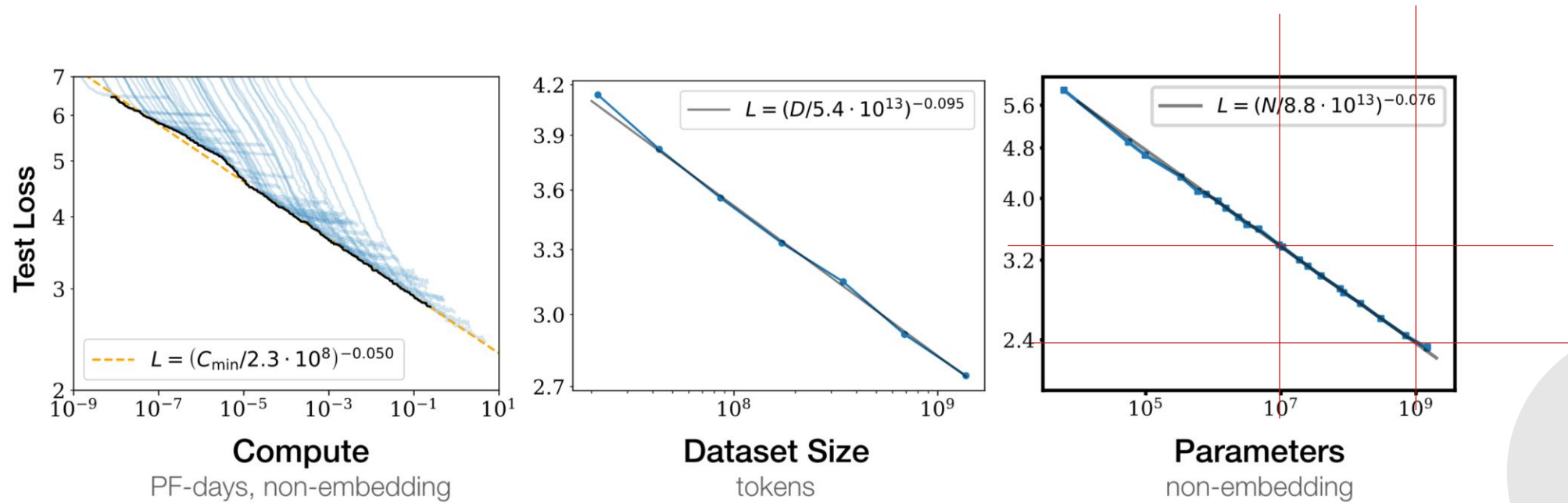


Source: Scaling Laws for Neural Language Models, Kaplan et al . 2020, Open AI



Neural Scaling Laws Kaplan et al 2020

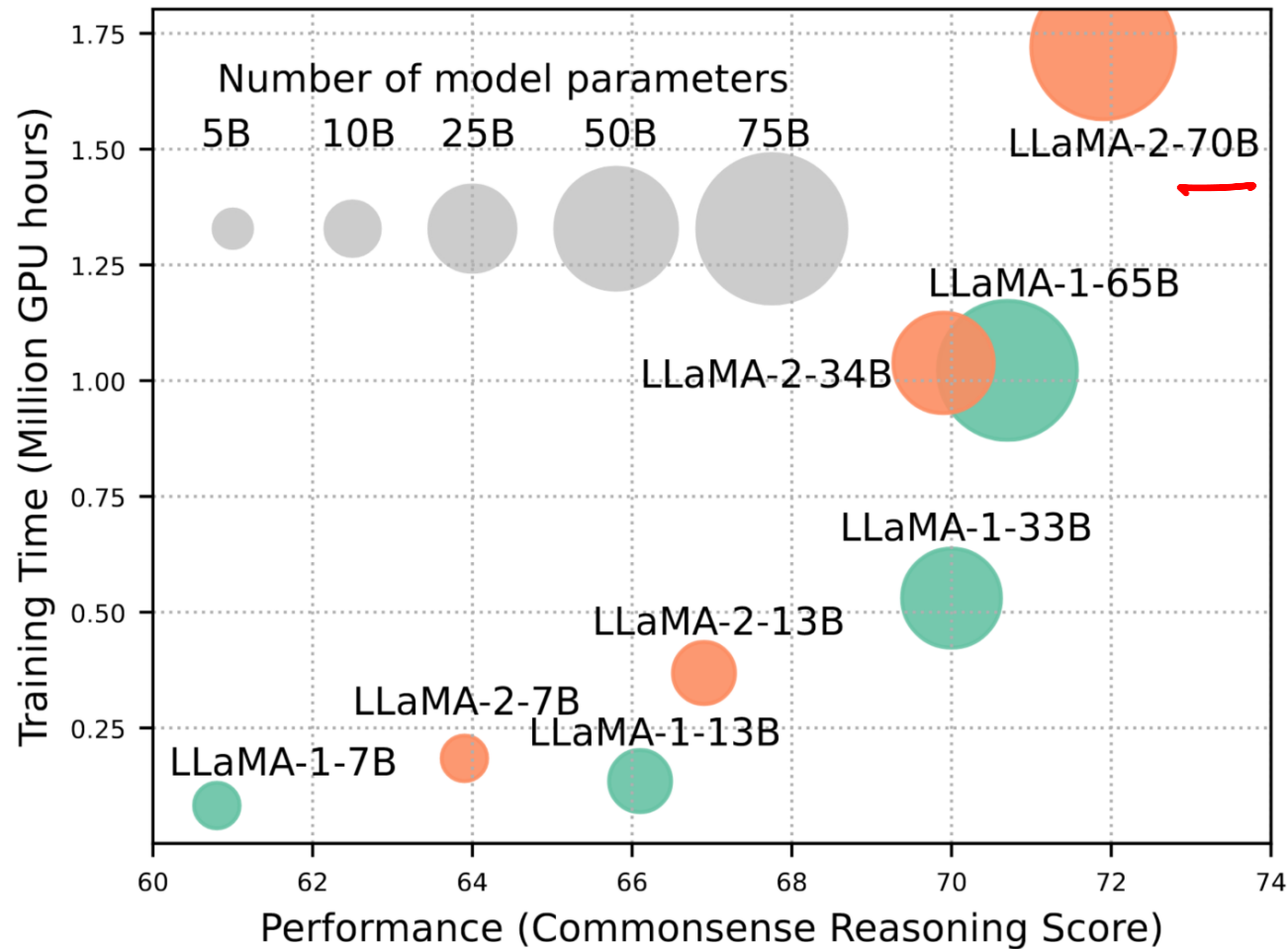
- Performance improve smoothly as we increase the **compute**, **dataset size**, or the model size



Source: Scaling Laws for Neural Language Models, Kaplan et al. 2020, Open AI



Training Resources vs Performance

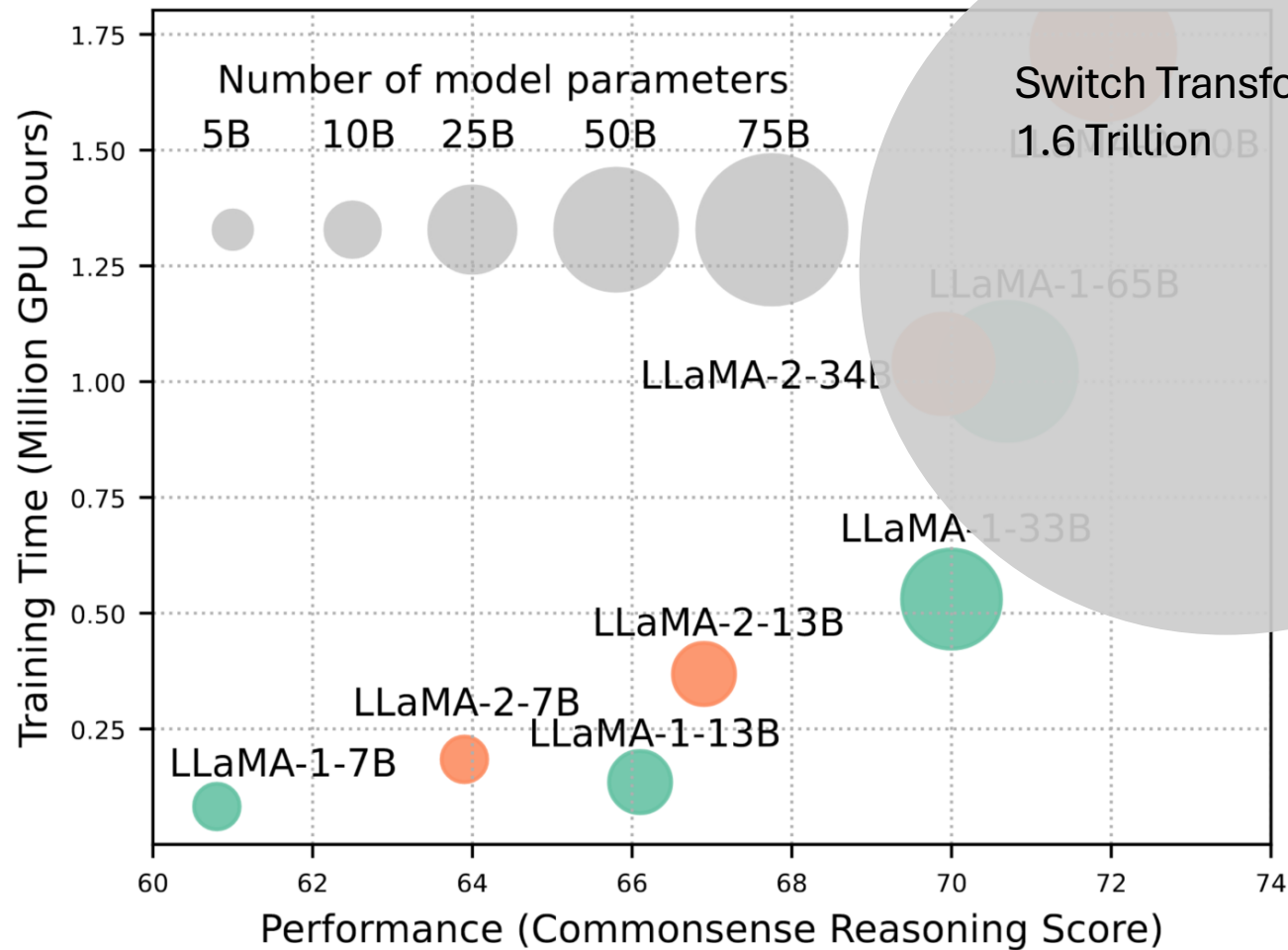


- Based on Nvidia A100 80GB GPU

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



Training Resources vs Performance



- Based on Nvidia A100 80GB GPU

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



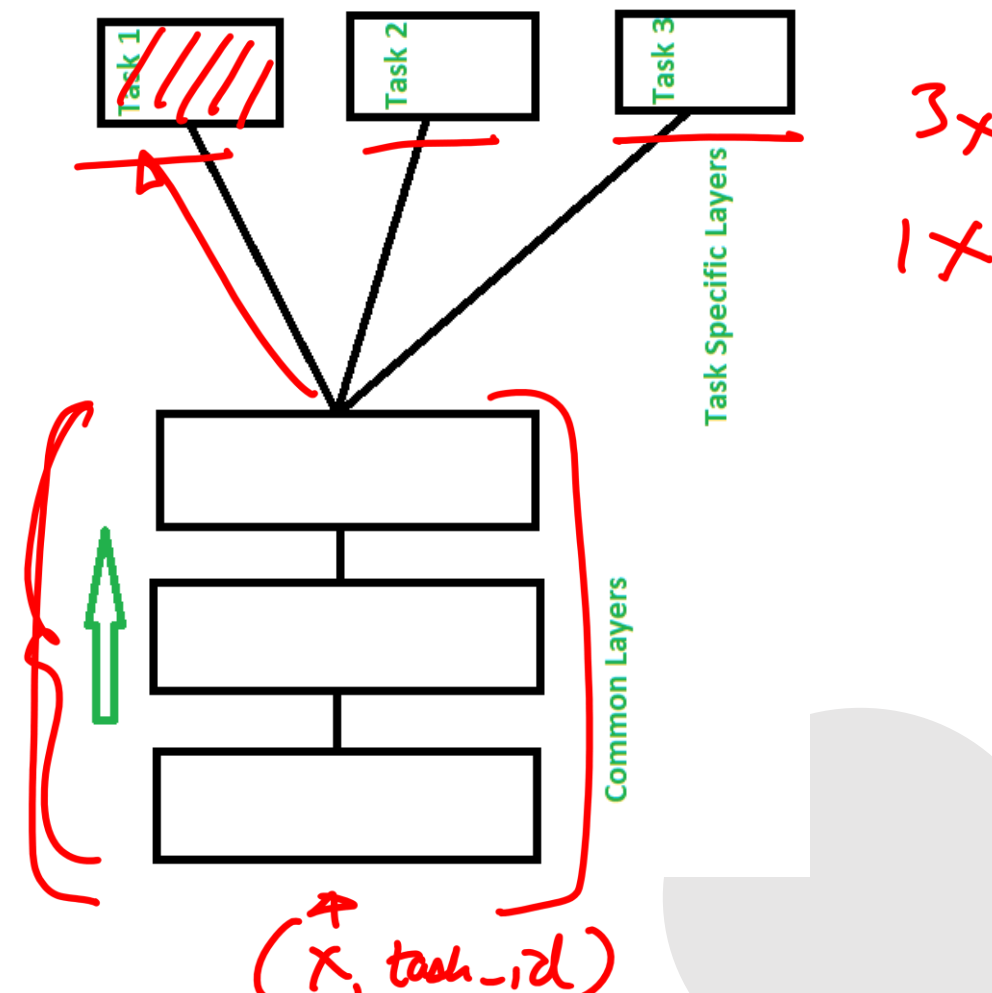
Any efficient way to increase model size?

- **LLMs as generalists** → they implicitly do *multi-task learning* (MTL)
 - summarization, QA, coding, reasoning, etc.
- How does a general MTL model look like?



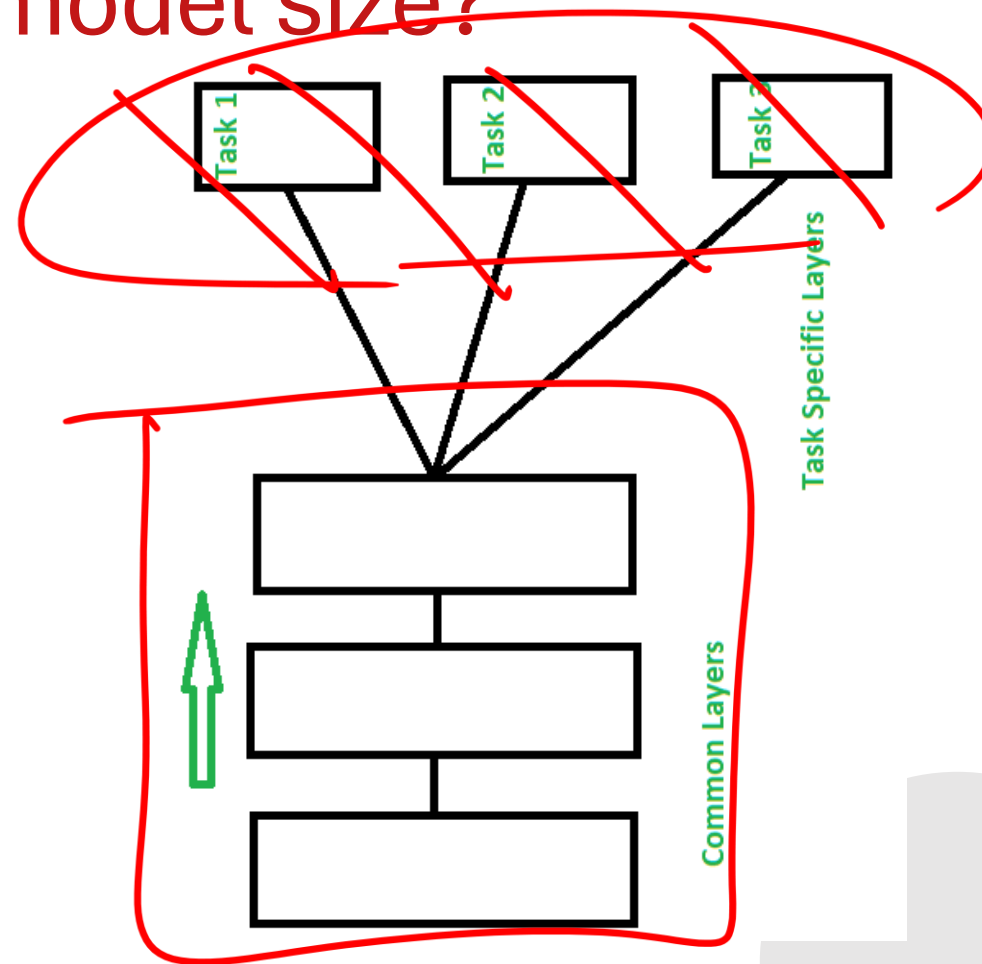
Any efficient way to increase model size?

- **LLMs as generalists** → they implicitly do *multi-task learning* (MTL)
 - summarization, QA, coding, reasoning, etc.
- How does a general MTL model look like?
 - Shared parameters (always active)
 - Task specific parameters (active based on input/desired output)



Any efficient way to increase model size?

- **LLMs as generalists** → they implicitly do *multi-task learning* (MTL)
 - summarization, QA, coding, reasoning, etc.
- How does a general MTL model look like?
 - Shared parameters (always active)
 - Task specific parameters (active based on task)
- In LLMs → all parameters active for all tasks
- Can we create a model where only few parameters are active for a given input?



Mixture of Experts

Adaptive Mixtures of Local Experts

Robert A. Jacobs

Michael I. Jordan

*Department of Brain and Cognitive Sciences, Massachusetts Institute of Technology,
Cambridge, MA 02139 USA*

Steven J. Nowlan

Geoffrey E. Hinton

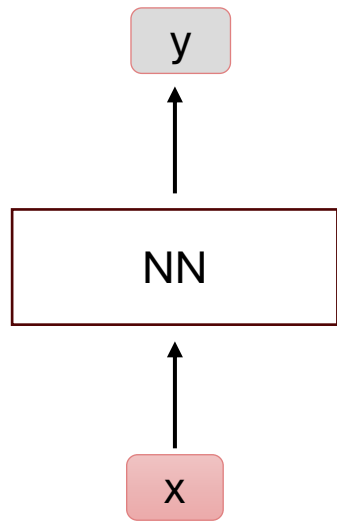
*Department of Computer Science, University of Toronto,
Toronto, Canada M5S 1A4*

Published in 1991!

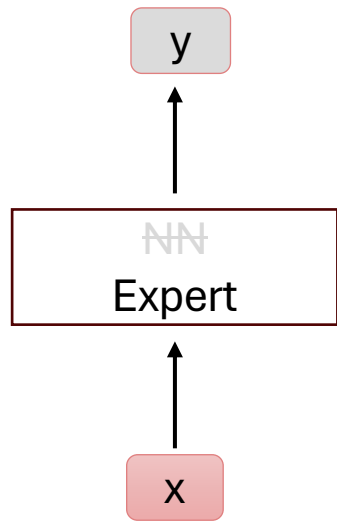
We present a new supervised learning procedure for systems composed of many separate networks, each of which learns to handle a subset of the complete set of training cases. The new procedure can be viewed either as a modular version of a multilayer supervised network, or as an associative version of competitive learning. It therefore provides a new link between these two apparently different approaches. We demonstrate that the learning procedure divides up a vowel discrimination task into appropriate subtasks, each of which can be solved by a very simple expert network.



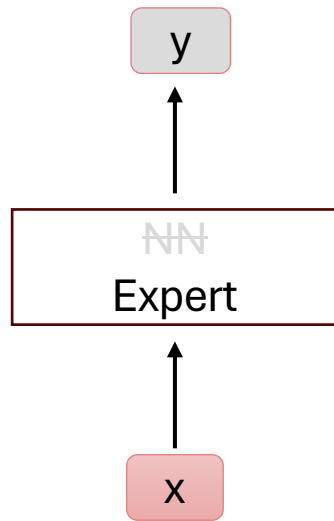
Mixture of Experts



Mixture of Experts



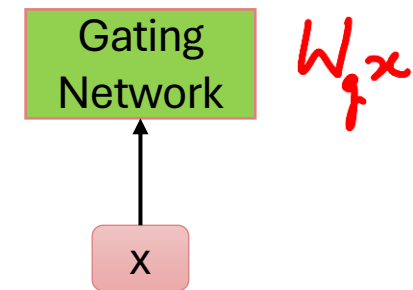
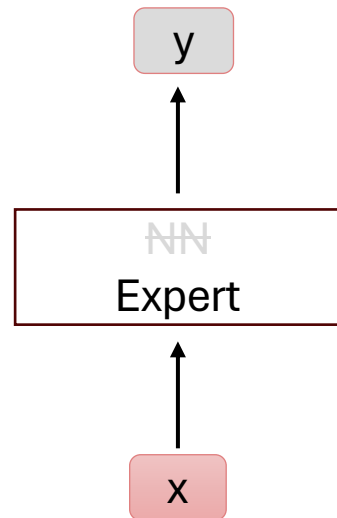
Mixture of Experts



How to decide which expert to use?

Let's have a "Gating Network" that decides (*probabilistically*)

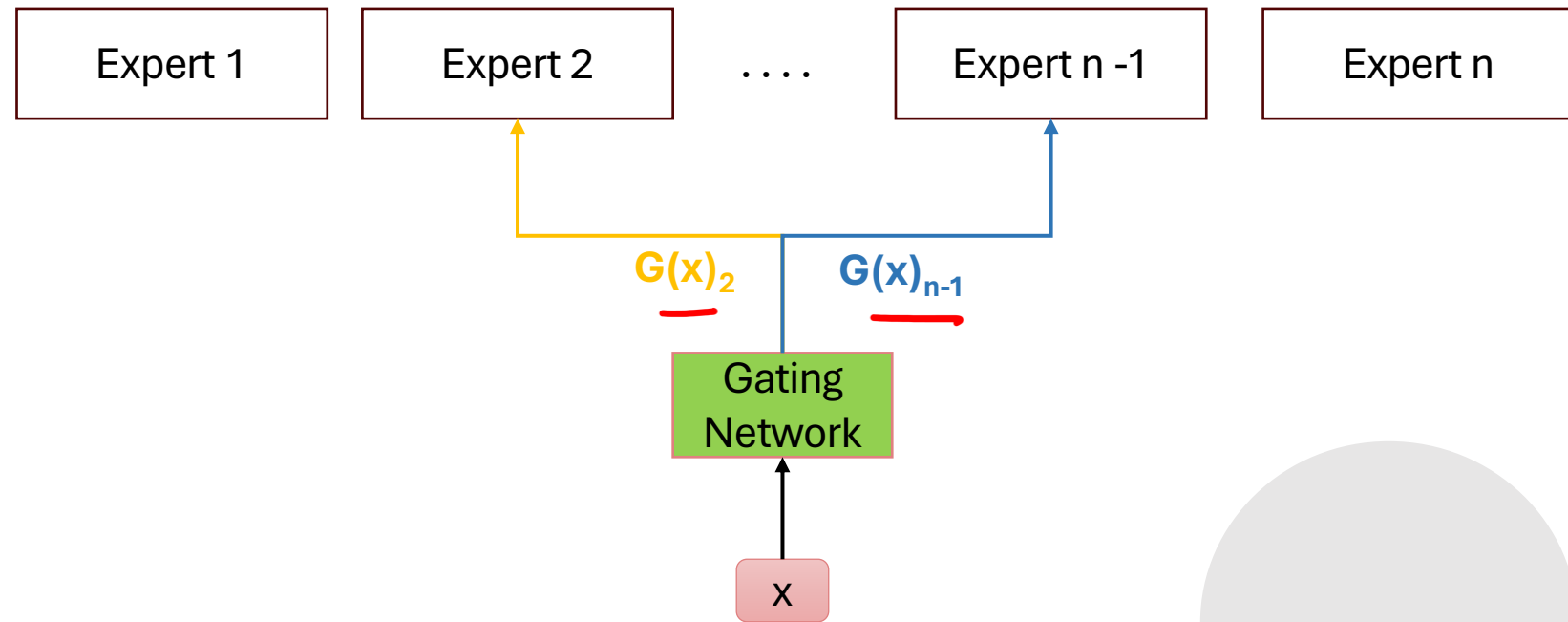
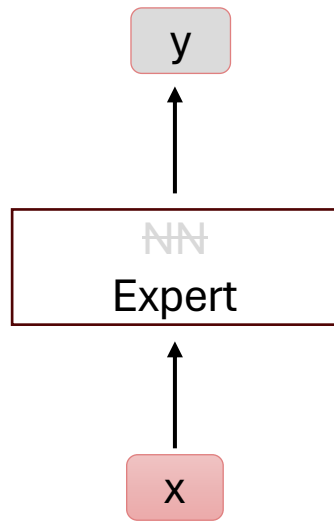
Mixture of Experts



How to decide which expert to use?

Let's have a "Gating Network" that decides (*probabilistically*)

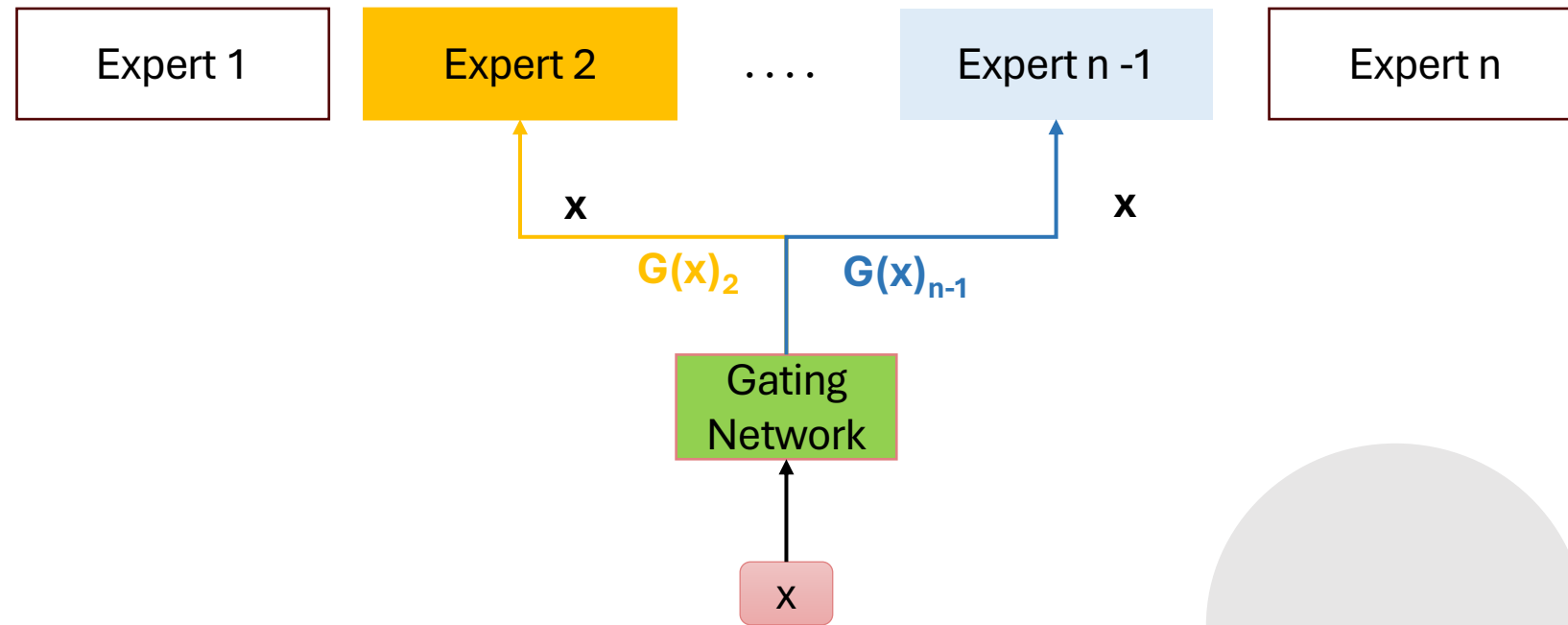
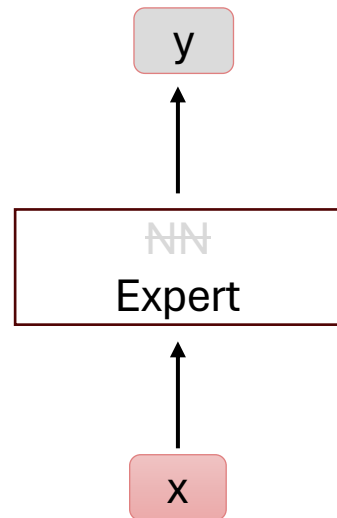
Mixture of Experts



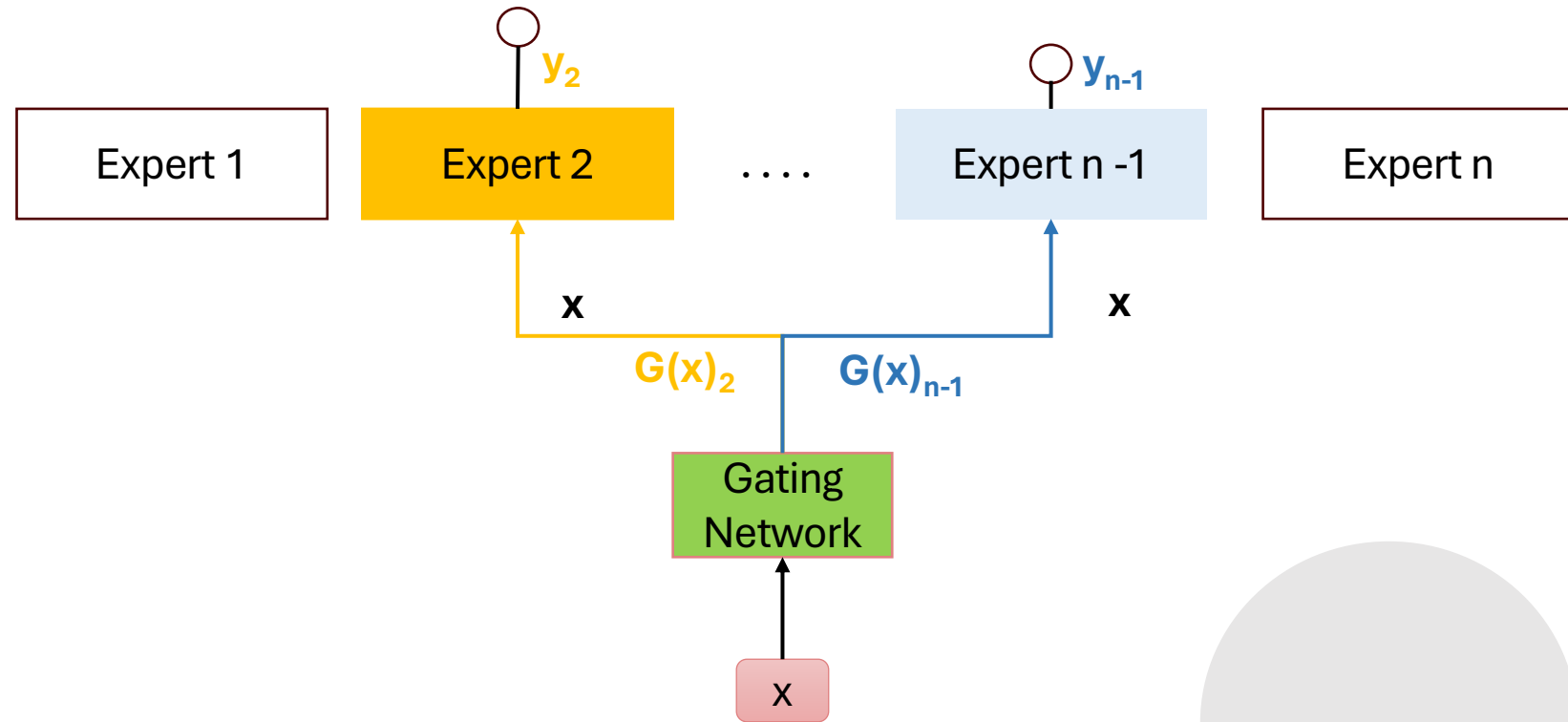
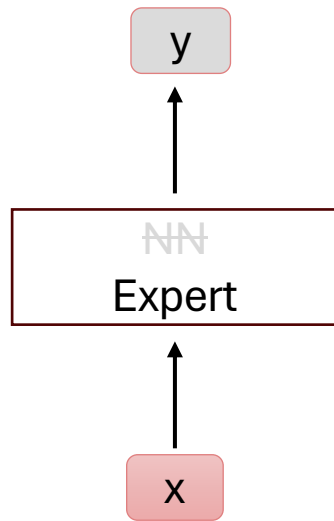
How to decide which expert to use?

Let's have a "Gating Network" that decides (*probabilistically*)

Mixture of Experts



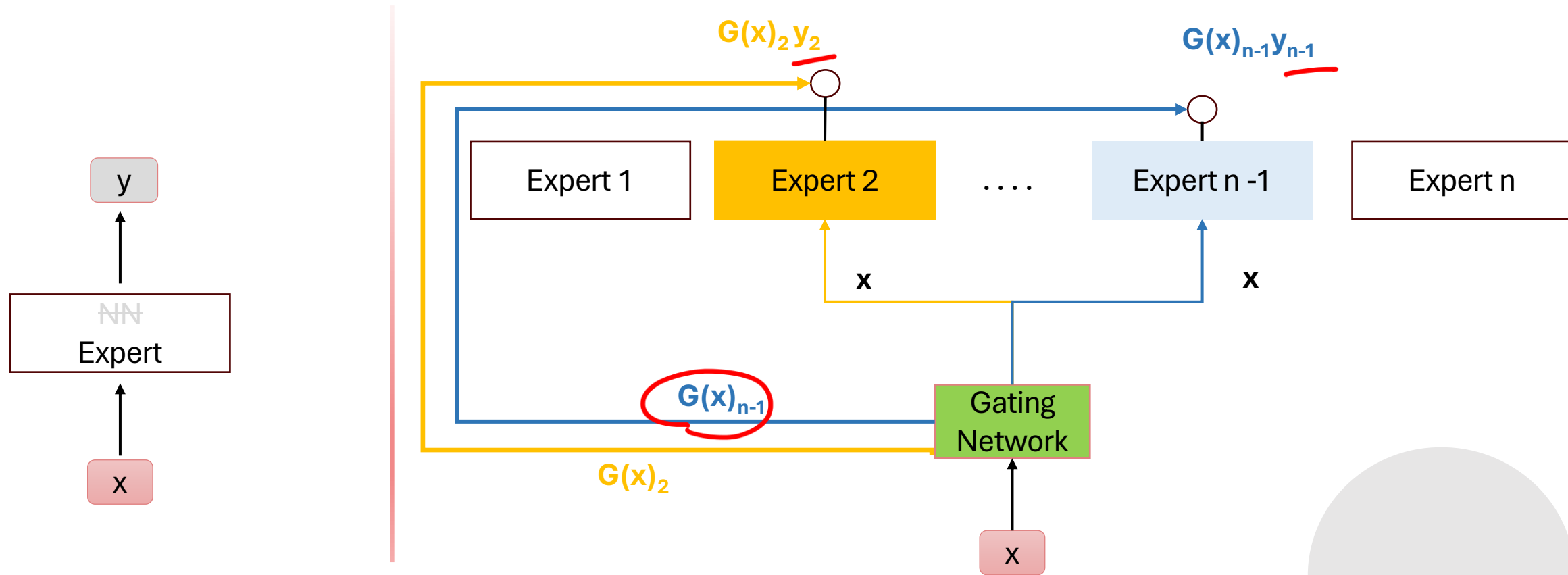
Mixture of Experts



How to decide which expert to use?

Let's have a "Gating Network" that decides (*probabilistically*)

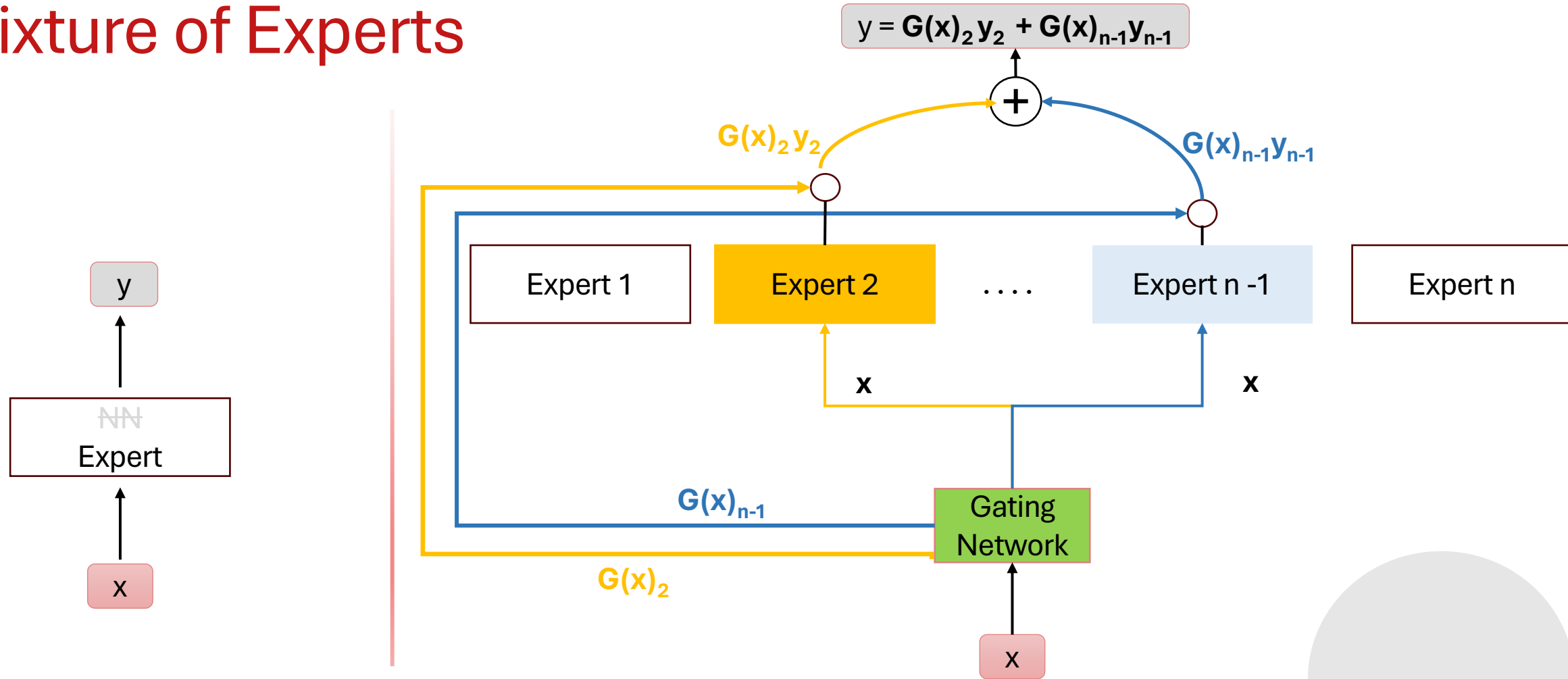
Mixture of Experts



How to decide which expert to use?

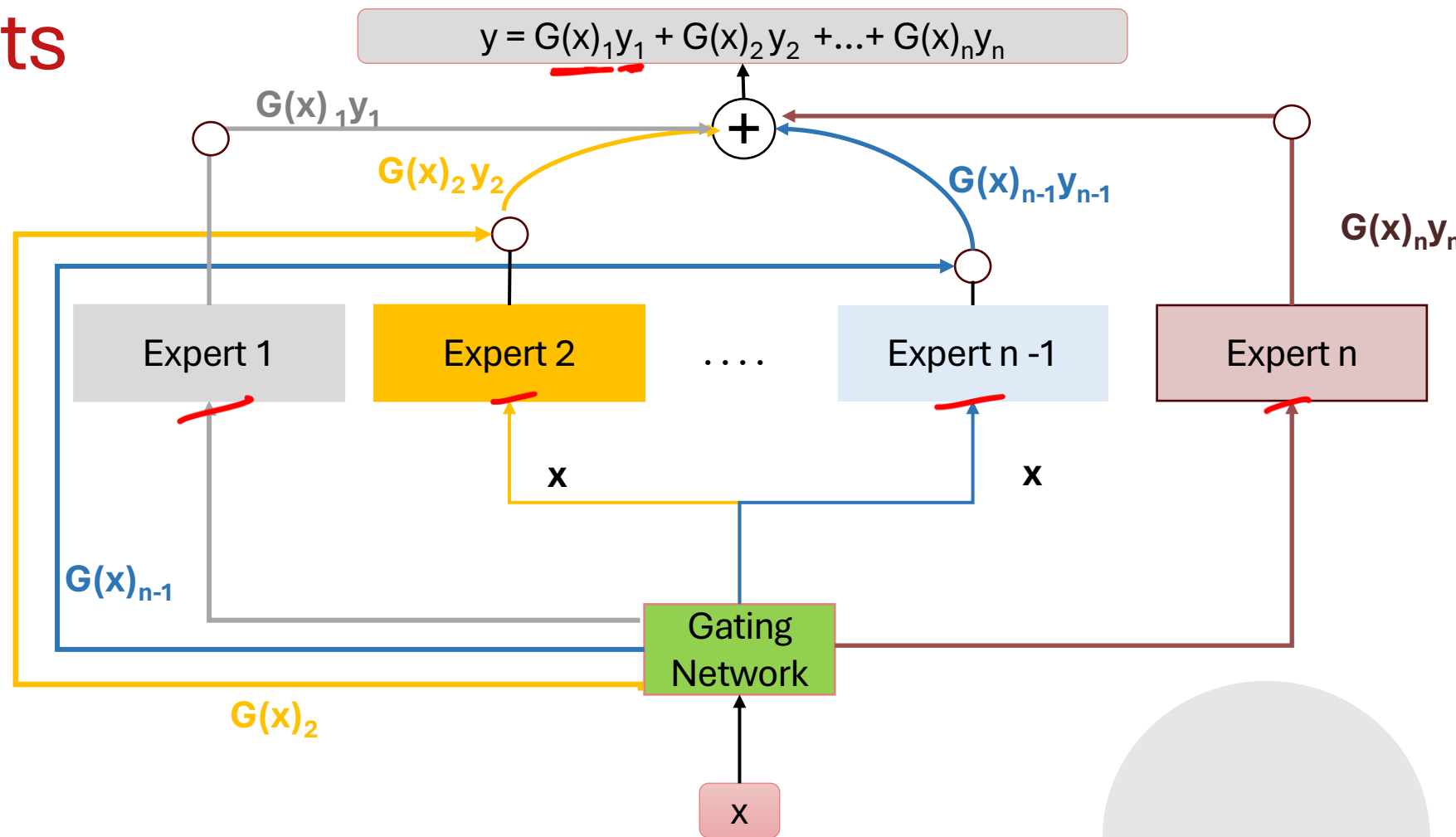
Let's have a "Gating Network" that decides (*probabilistically*)

Mixture of Experts



Mixture of Experts

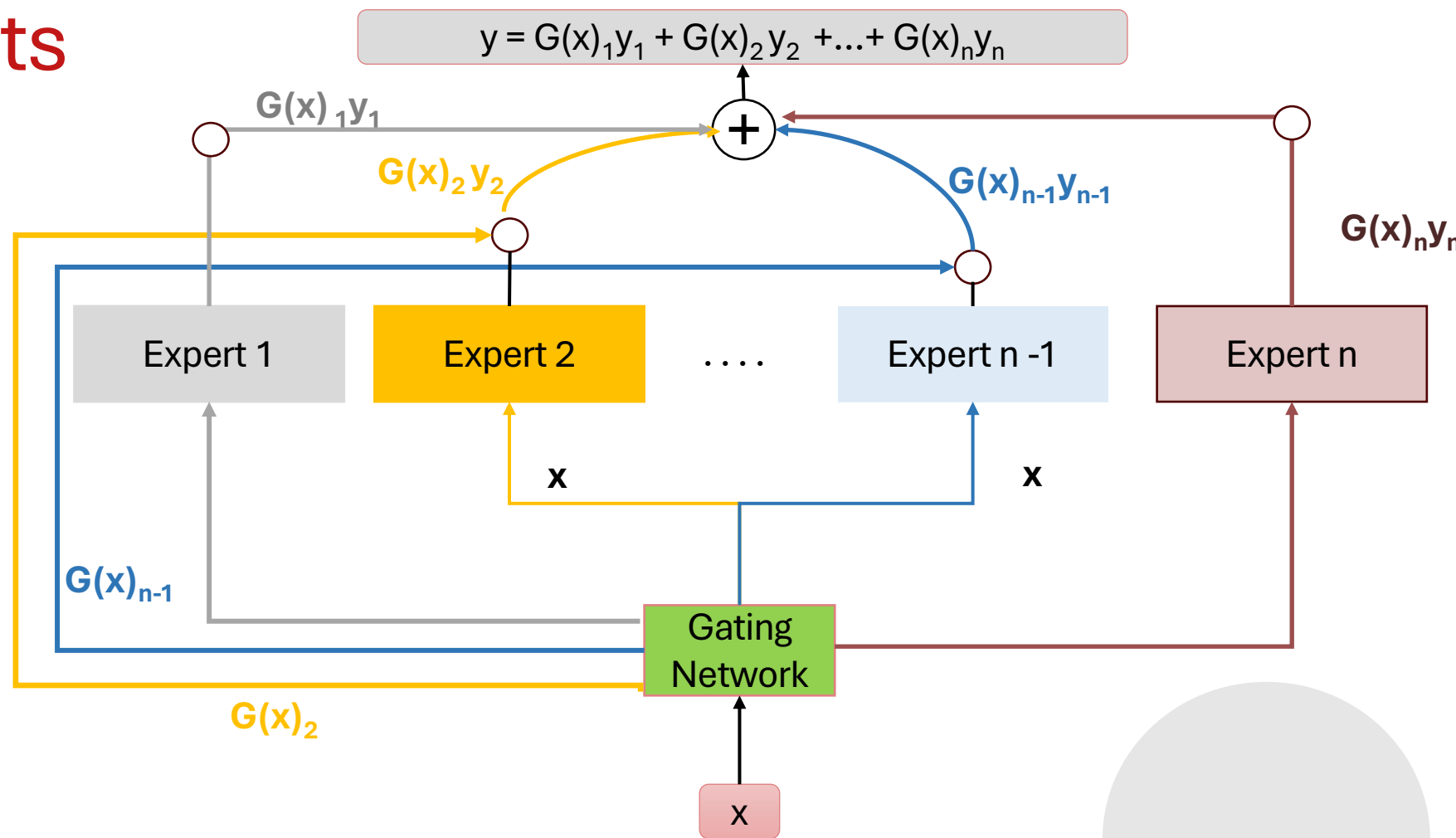
$$\underline{G_\sigma(x)} = \text{Softmax}(x \cdot \underline{W_g})$$



Mixture of Experts

$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

$$y = \sum_{i=1}^n G(x)_i E_i(x)$$



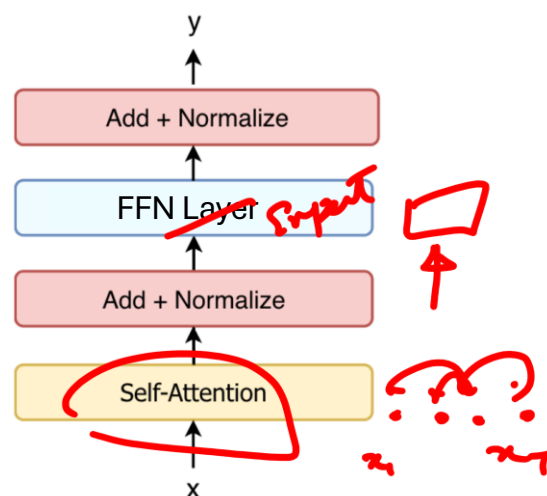
Mixture of Experts - Chronology

- Mixture of Experts Model [Jacobs et al., 1991; Jordan and Jacobs, 1994; Jordan et al., 1997; Tresp, 2001; Collobert et al., 2002;]
- MoE layer in Deep Learning [Eigen et al., 2013; Shazeer et al., 2017; Dauphin et al., 2017; Vaswani et al., 2017]
- MoE layer in Transformer based LLMs [Fedus et al., 2021; Du, Nan, et al., 2021]
- Mixtral-8x7B [Jiang et al. 2024] - Apache 2.0 license; surpassed GPT-3.5 Turbo, Claude-2.1, Gemini Pro, and Llama 2 70B - chat model on human benchmarks
- GPT-4, Deepseek-V3/R1, MiniMax-01

Content credits: <https://www.youtube.com/watch?v=TwHPxUAuqy4>



How/where to use Mixture of Experts in Transformers?

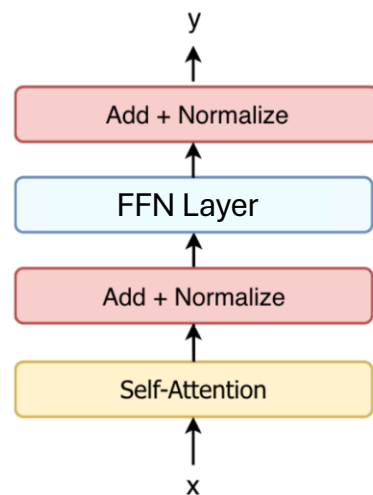


Dense Transformer Block

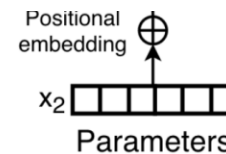
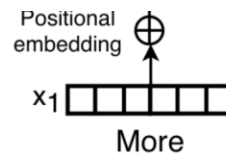
Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Mixture of Experts as a Layer



Dense Transformer Block



Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)

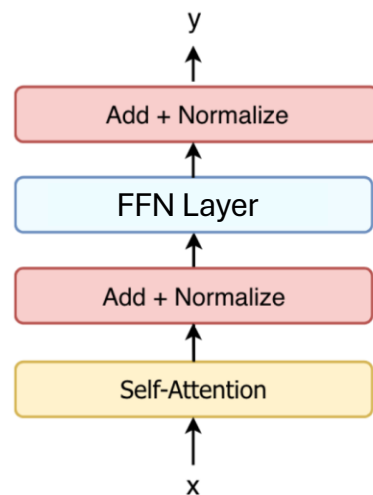


LLMs: Introduction and Recent Advances

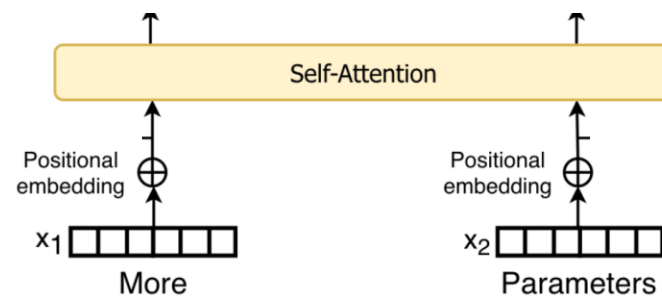


Yatin Nandwani

Mixture of Experts as a Layer



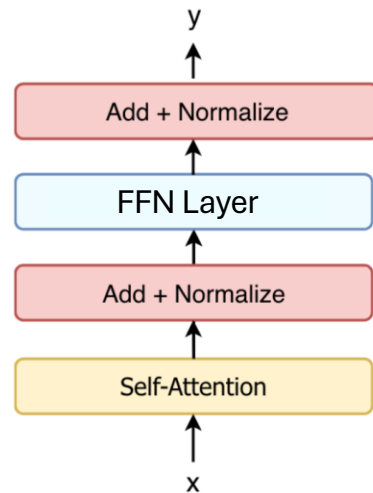
Dense Transformer Block



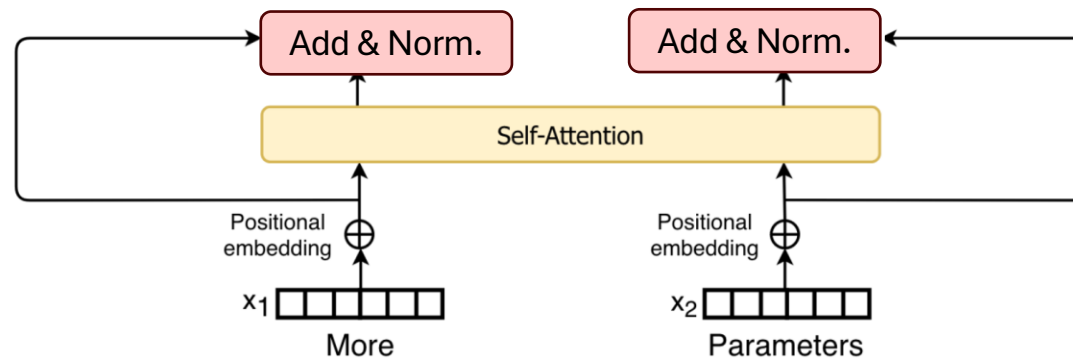
Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Mixture of Experts as a Layer



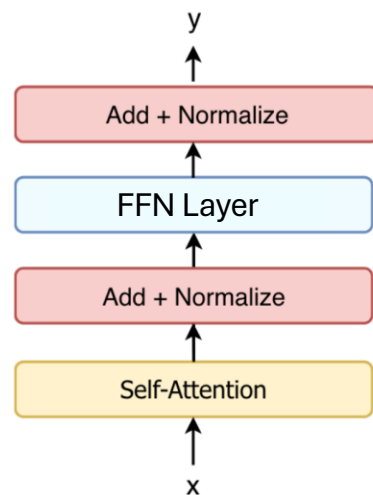
Dense Transformer Block



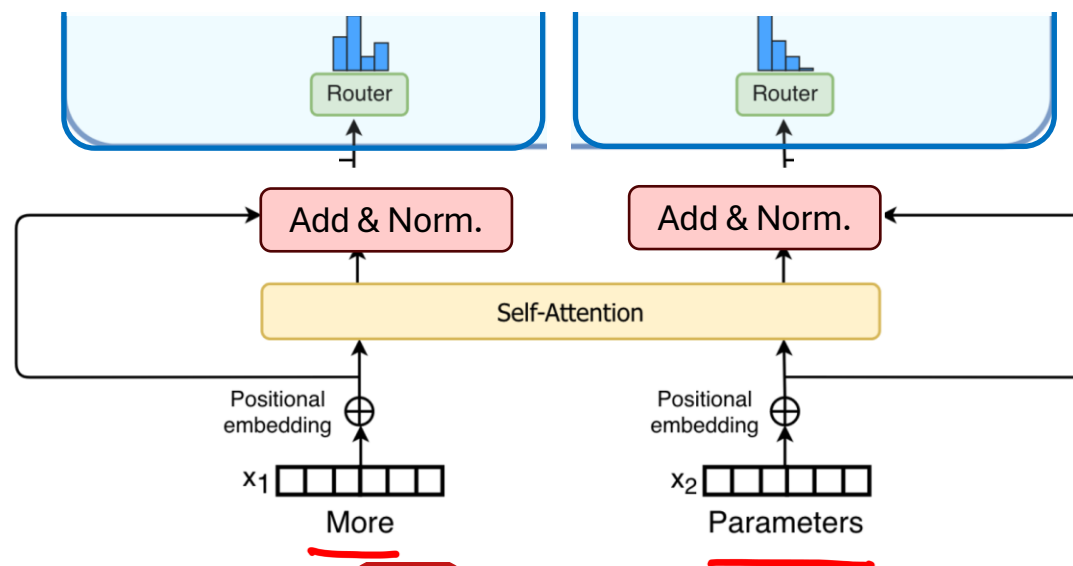
Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Mixture of Experts as a Layer



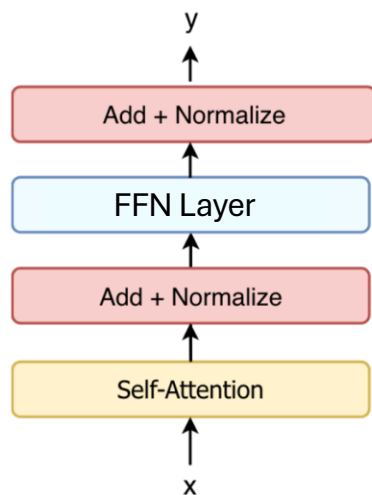
Dense Transformer Block



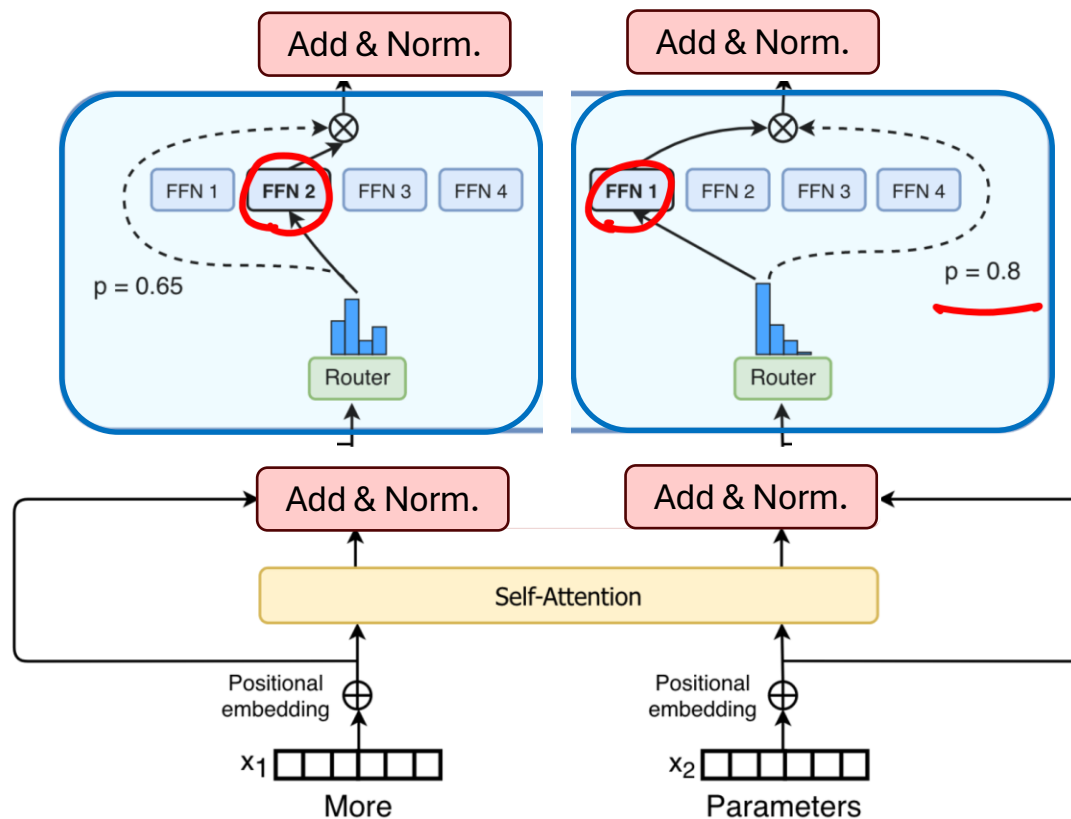
Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Mixture of Experts as a Layer



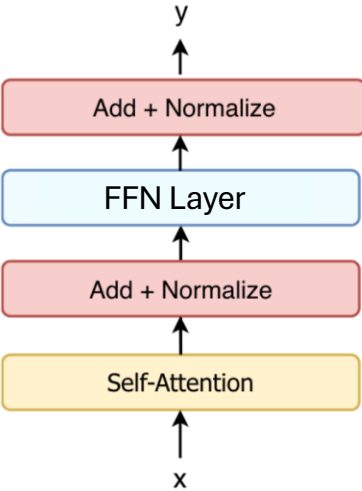
Dense Transformer Block



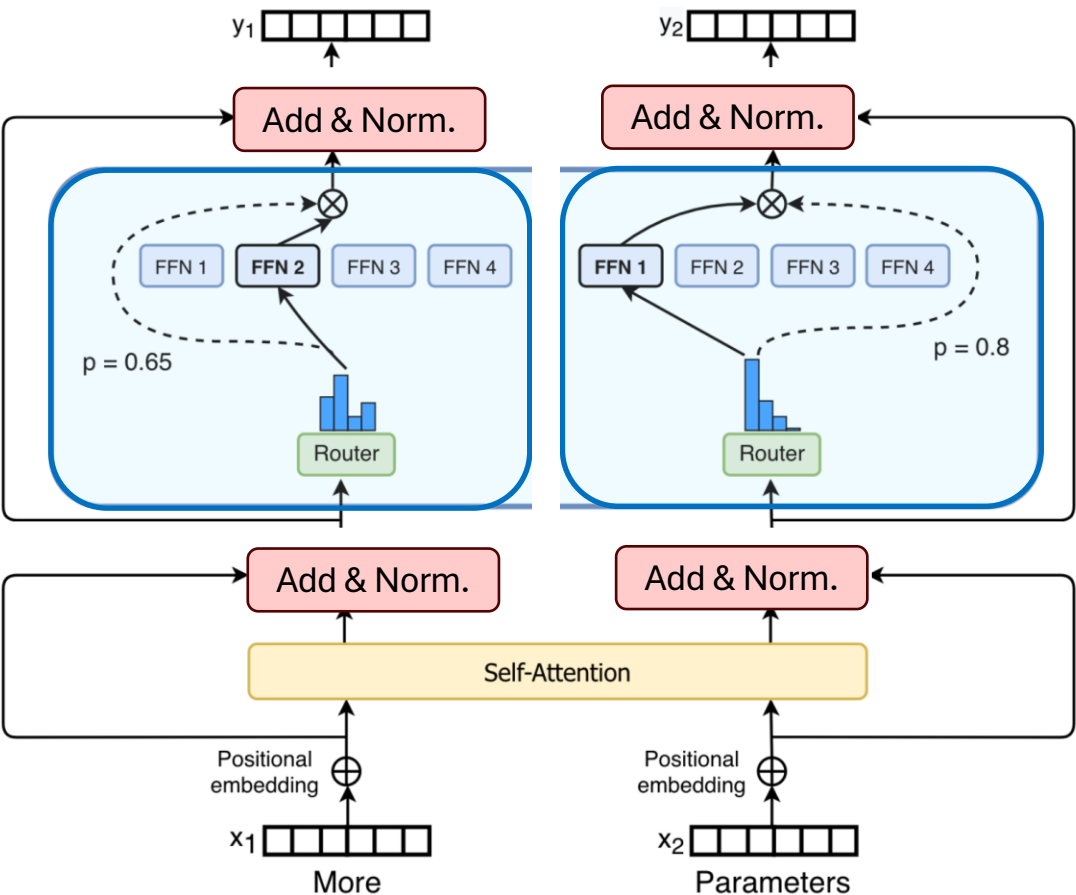
Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Mixture of Experts as a Layer



Dense Transformer Block



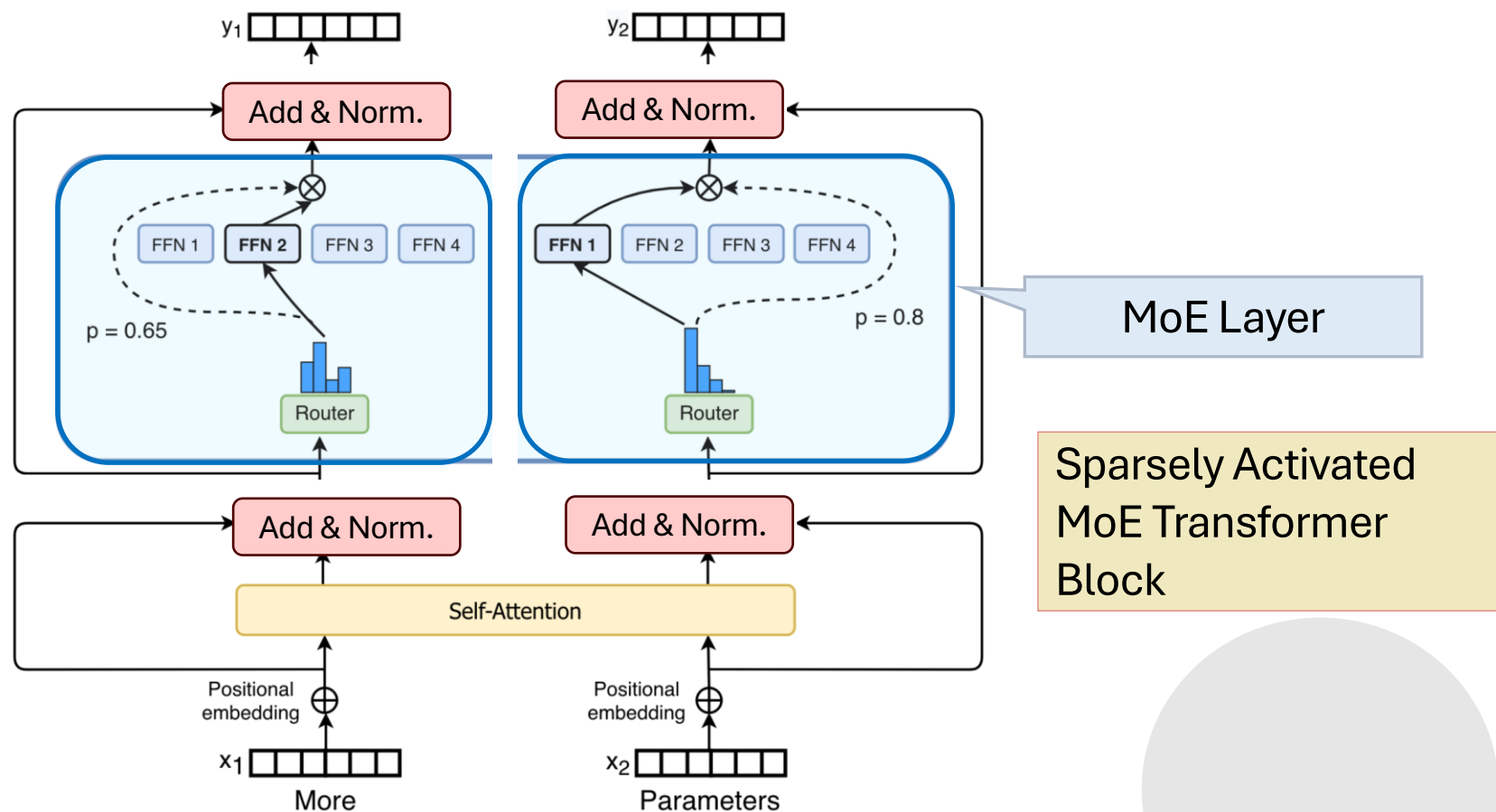
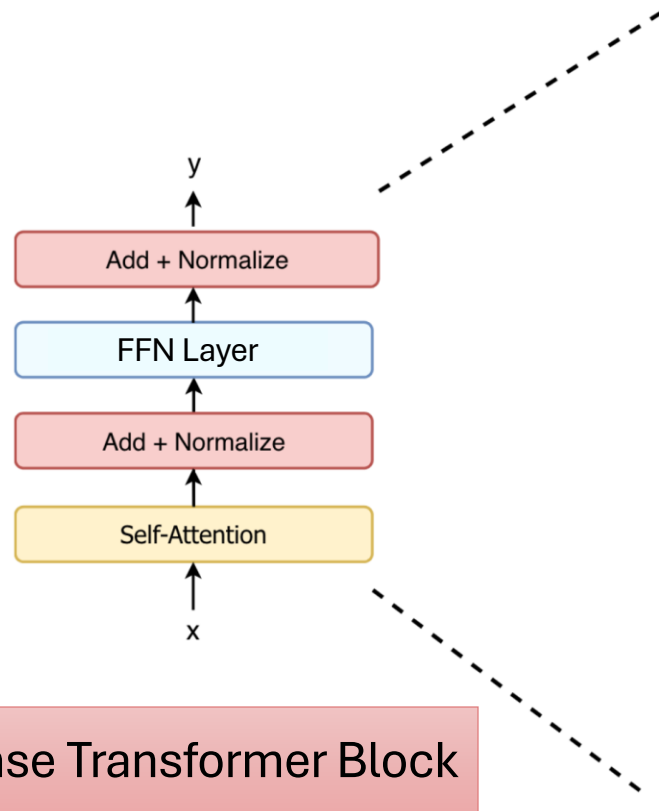
Handwritten notes in red ink:

- $4h^2 + 4h^2 + 4h + h$
- $8h^2 + 5h$ (circled)
- $10 \times 8h^2$
- $3h^2 + 1h^2$
- $4h^2$ (circled)

Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity



Mixture of Experts as a Layer

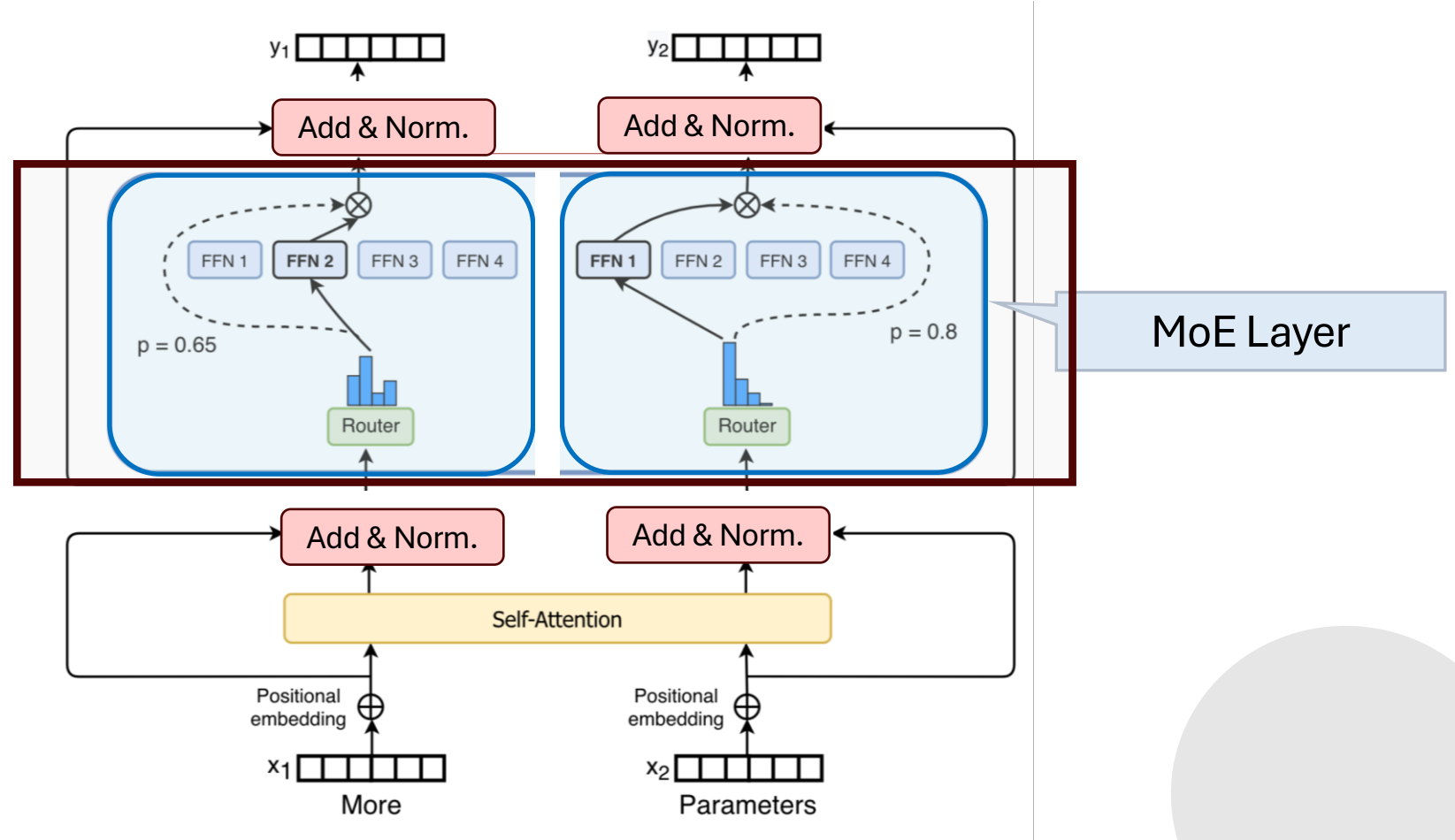


Content credits: [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)



Sparse MoE Layer – Greedy expert selection

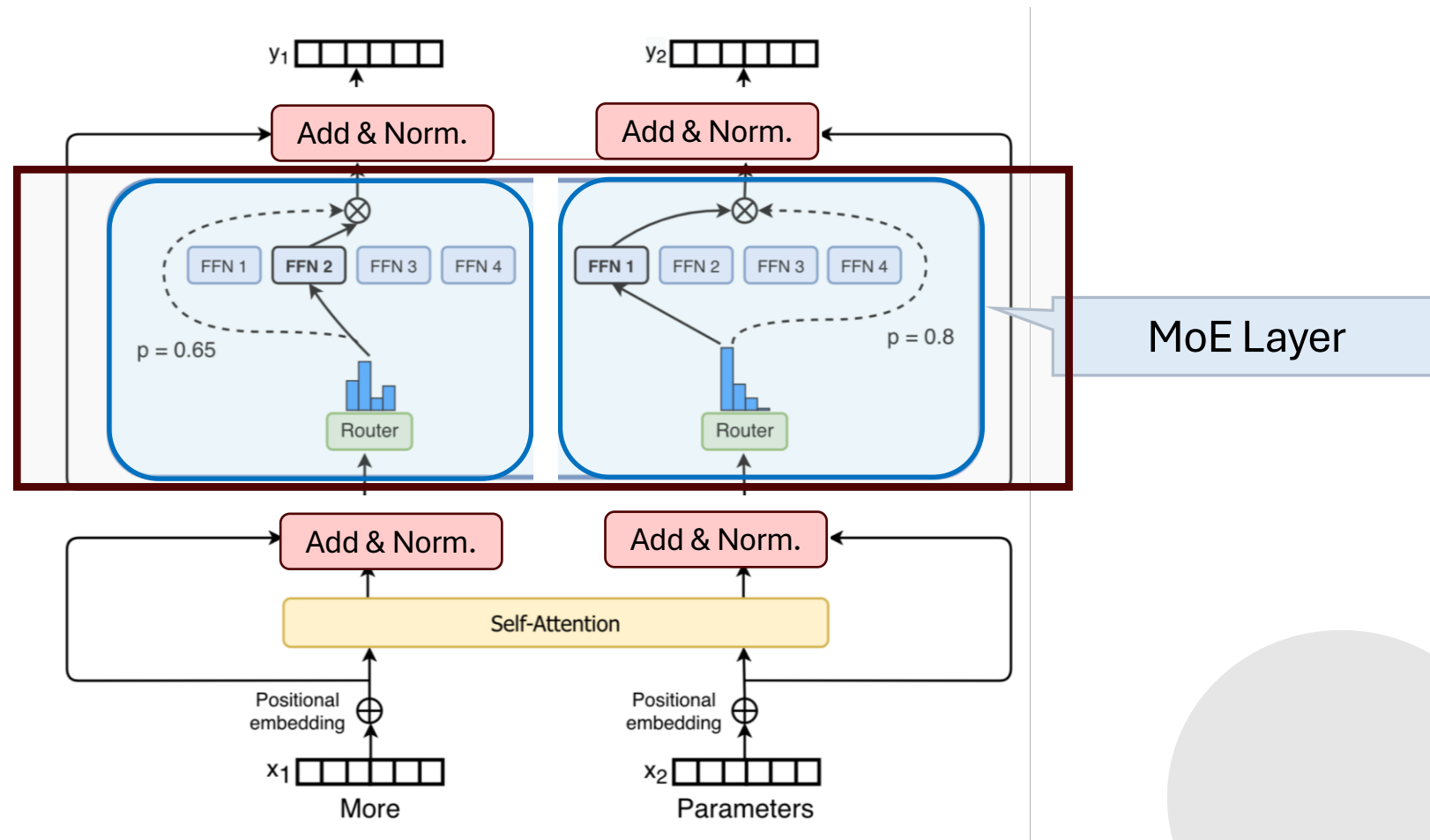
$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$



Sparse MoE Layer – Greedy expert selection

$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

$$i^* = \text{argmax } G_{\sigma}(x)$$

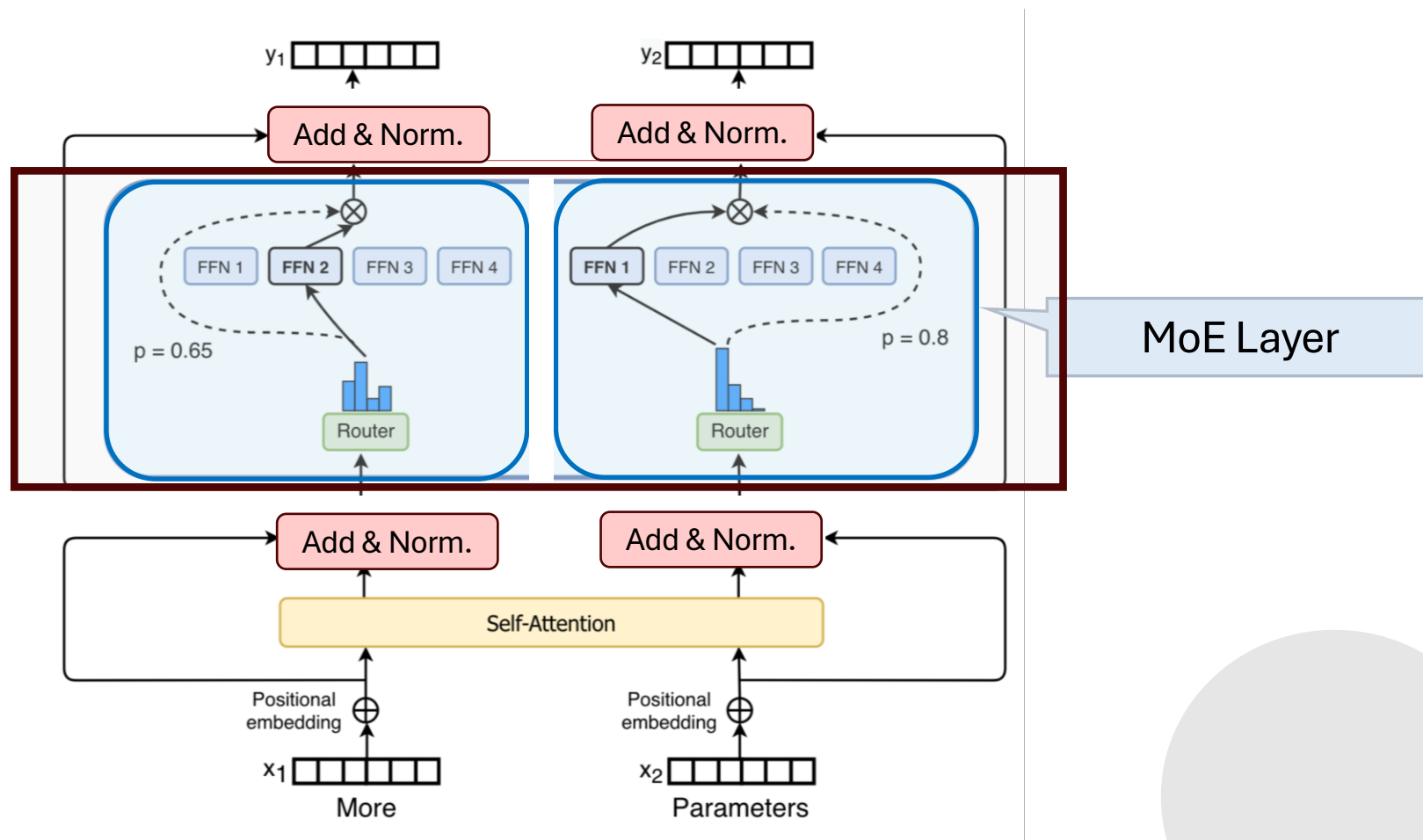


Sparse MoE Layer – Greedy expert selection

$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

$$i^* = \text{argmax } G_{\sigma}(x)$$

$$y = G(x)_{i^*} E_{i^*}(x)$$



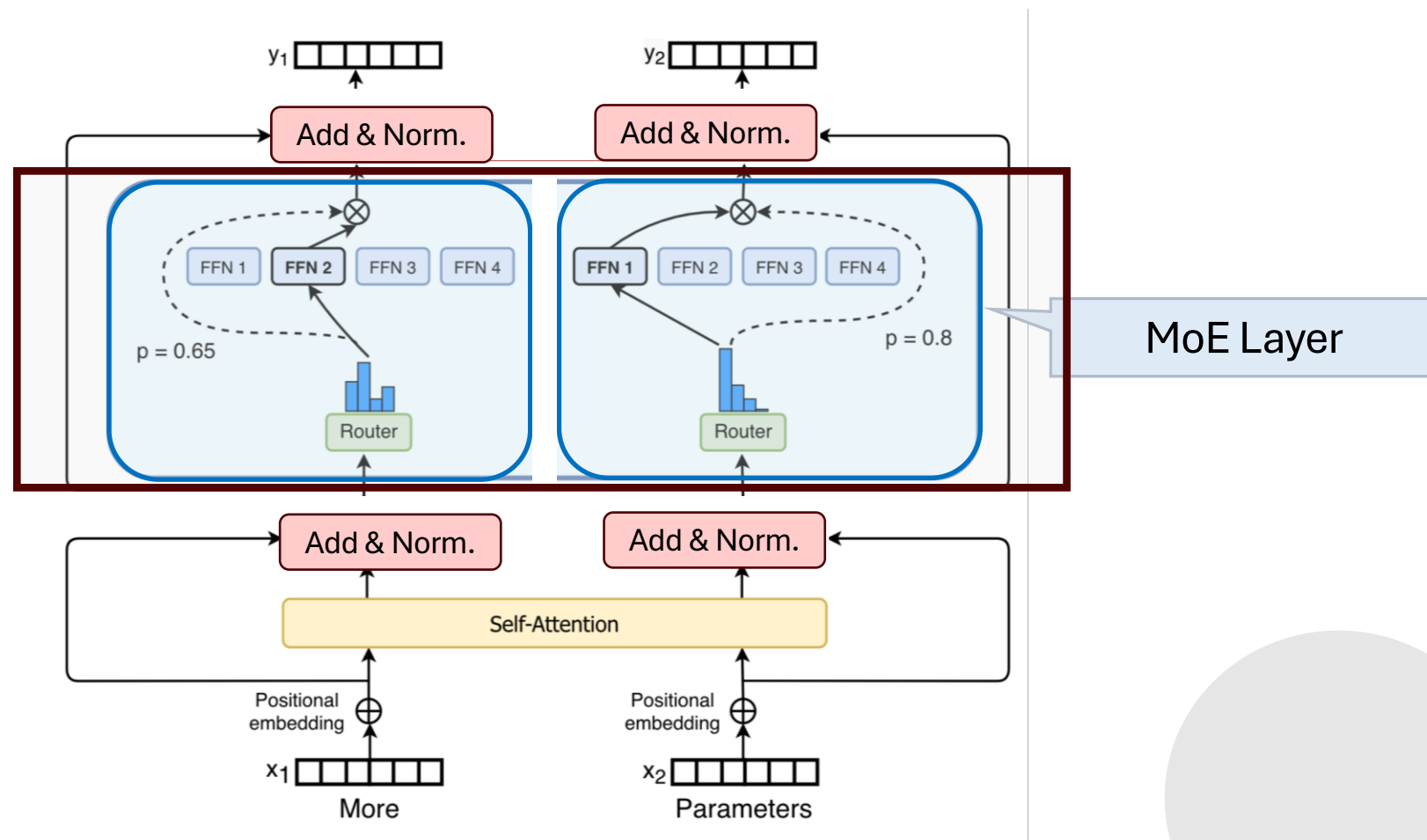
Sparse MoE Layer – Greedy expert selection

$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

$$i^* = \text{argmax } G_{\sigma}(x)$$

$$y = G(x)_{i^*} E_{i^*}(x)$$

$$\frac{\nabla L}{\nabla E_{i^*}(x)} = G(x)_{i^*} \frac{\nabla L}{\nabla y}$$



Sparse MoE Layer – Greedy expert selection

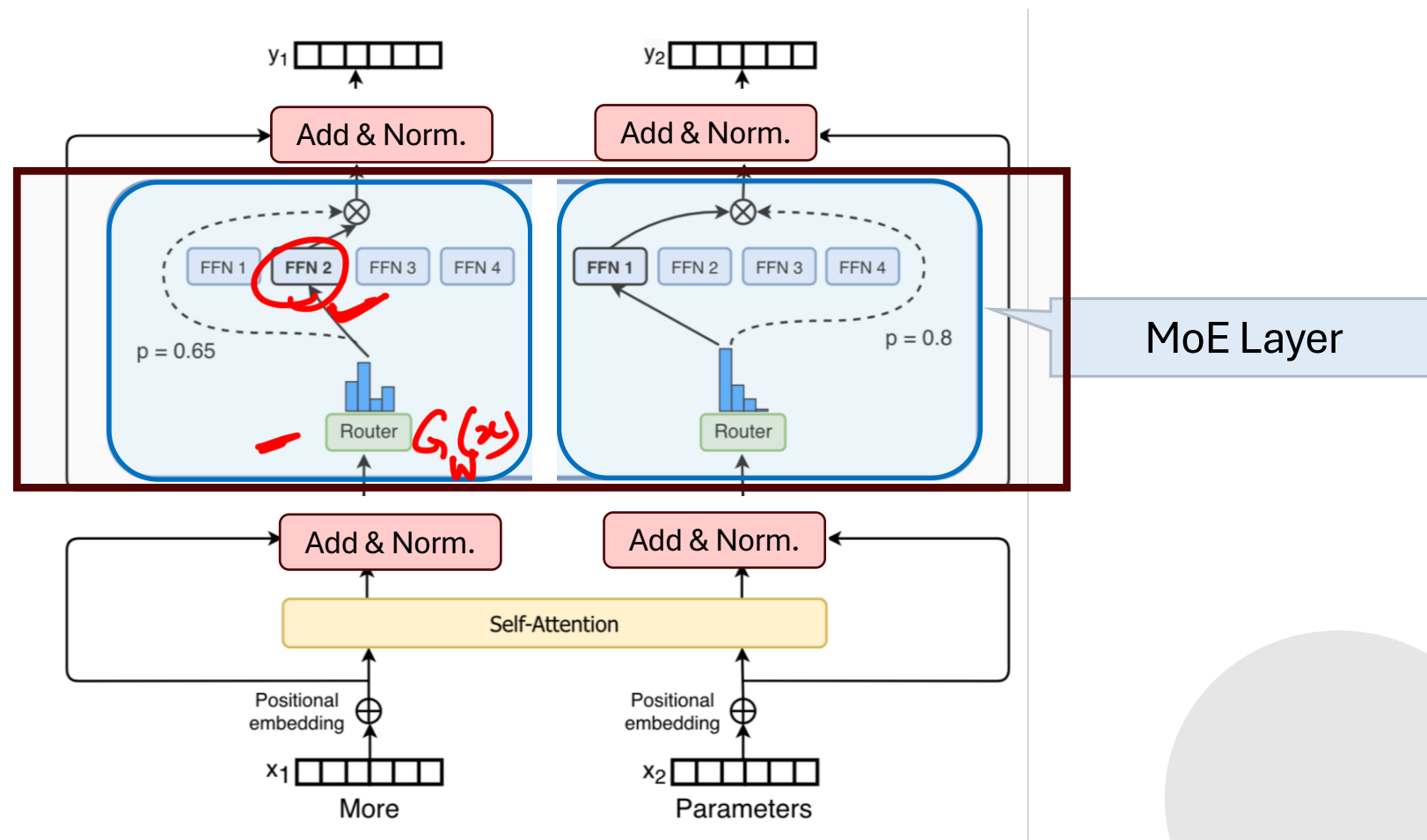
$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

$$i^* = \text{argmax } G_{\sigma}(x)$$

$$y = G(x)_{i^*} E_{i^*}(x)$$

$$\frac{\nabla L}{\nabla E_{i^*}(x)} = G(x)_{i^*} \frac{\nabla L}{\nabla y}$$

$$\frac{\nabla L}{\nabla G(x)_{i^*}} = (E_{i^*}(x))^T \frac{\nabla L}{\nabla y}$$



Sparse MoE Layer – Greedy expert selection

$$G_{\sigma}(x) = \text{Softmax}(x \cdot W_g)$$

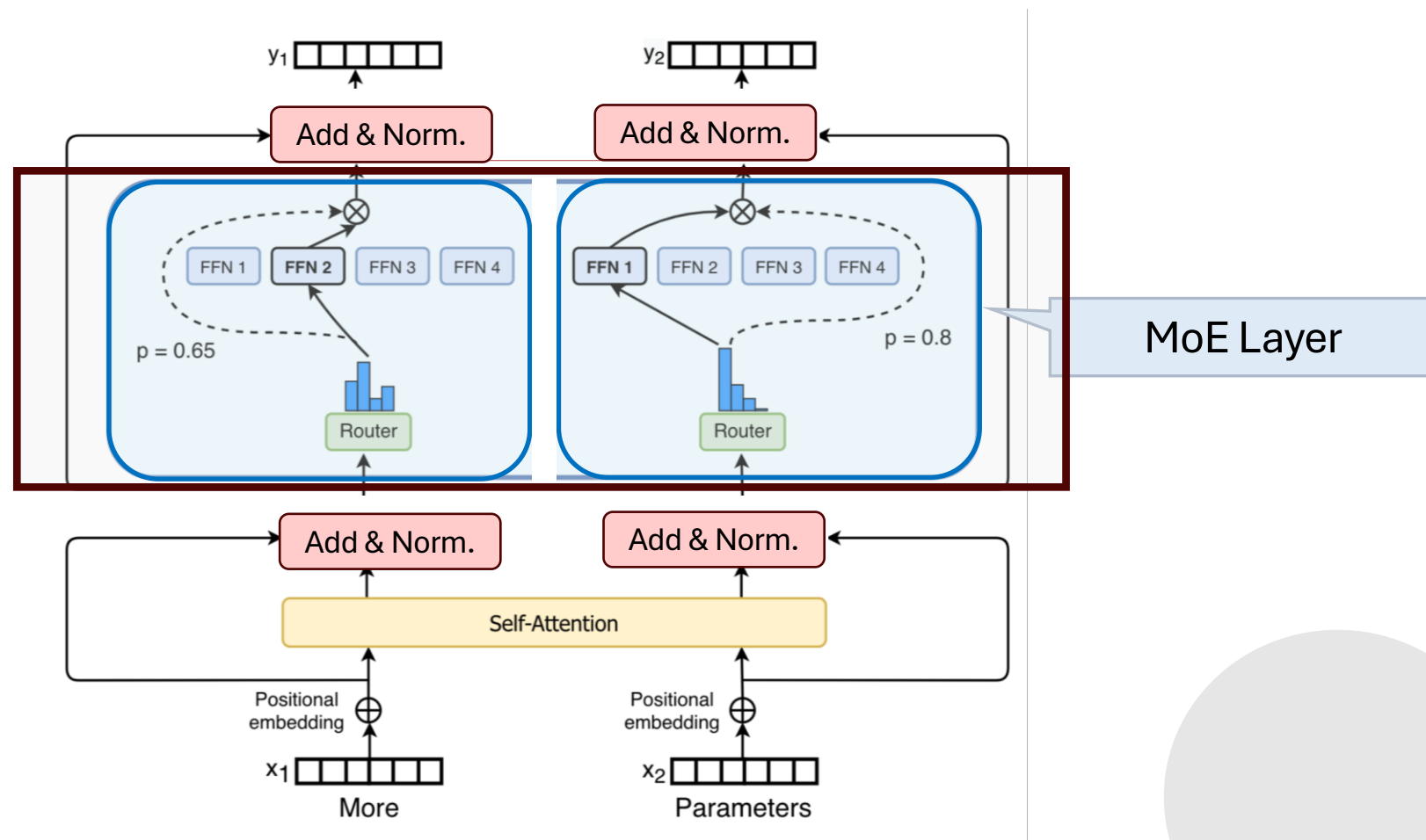
$$i^* = \text{argmax } G_{\sigma}(x)$$

$$y = G(x)_{i^*} E_{i^*}(x)$$

$$\frac{\nabla L}{\nabla E_{i^*}(x)} = G(x)_{i^*} \frac{\nabla L}{\nabla y}$$

$$\frac{\nabla L}{\nabla G(x)_{i^*}} = (E_{i^*}(x))^T \frac{\nabla L}{\nabla y}$$

$$\frac{\nabla L}{\nabla E_i(x)} = \mathbf{0}; \frac{\nabla L}{\nabla G(x)_i} = 0 \text{ for } i \neq i^*$$



Why do experts not collapse in practice?

❖ All experts have:

- Same architecture
- Same initialization scheme
- Trained with same optimizer

❖ So why don't they collapse? I.e.



Why do experts not collapse in practice?

Towards Understanding the Mixture-of-Experts Layer in Deep Learning

Zixiang Chen

Department of Computer Science
University of California, Los Angeles
Los Angeles, CA 90095, USA
chenzx19@cs.ucla.edu

Yihe Deng

Department of Computer Science
University of California, Los Angeles
Los Angeles, CA 90095, USA
yihedeng@cs.ucla.edu

Yue Wu

Department of Computer Science
University of California, Los Angeles
Los Angeles, CA 90095, USA
ywu@cs.ucla.edu

Quanquan Gu

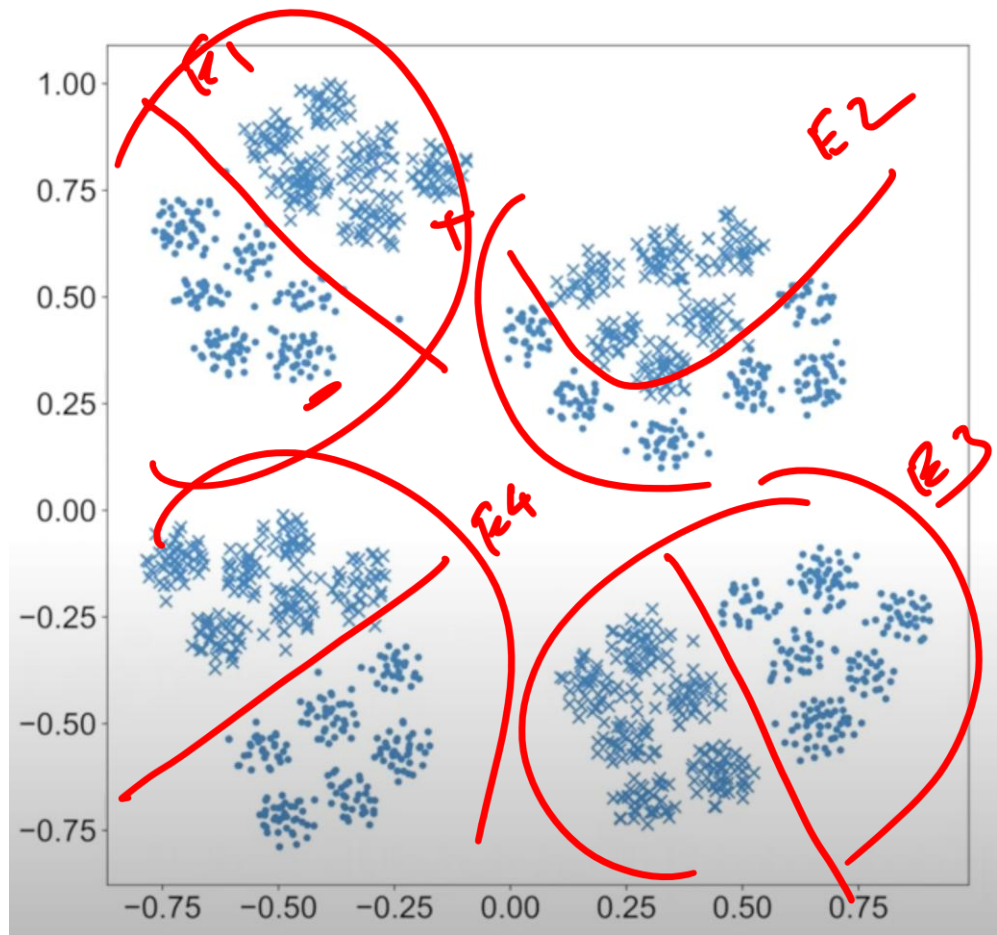
Department of Computer Science
University of California, Los Angeles
Los Angeles, CA 90095, USA
qgu@cs.ucla.edu

Yuanzhi Li

Machine Learning Department
Carnegie Mellon University
Pittsburgh, PA 15213, USA
yuanzhil@andrew.cmu.edu



Why do experts not collapse in practice?



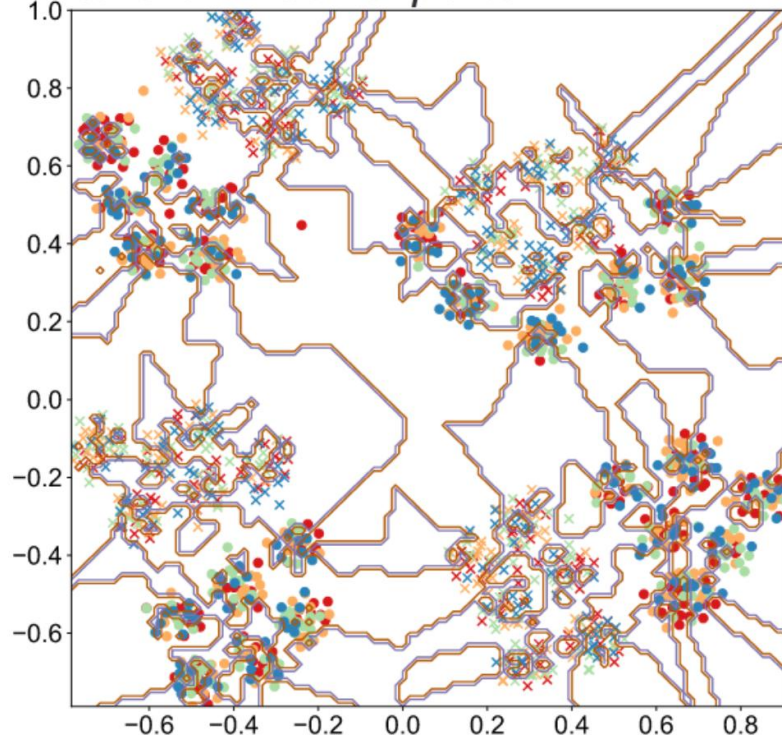
- ✓ Synthetically generated 50 dim. data
- ✓ Visualization in 2-D space
- ✓ 4 clusters
- ✓ Each cluster is linearly separable
- ✓ Ideally, a mixture of 4 linear experts sufficient for classification



$$E \rightarrow W_{ei} x$$

Why do experts not collapse in practice?

Mixture of linear experts

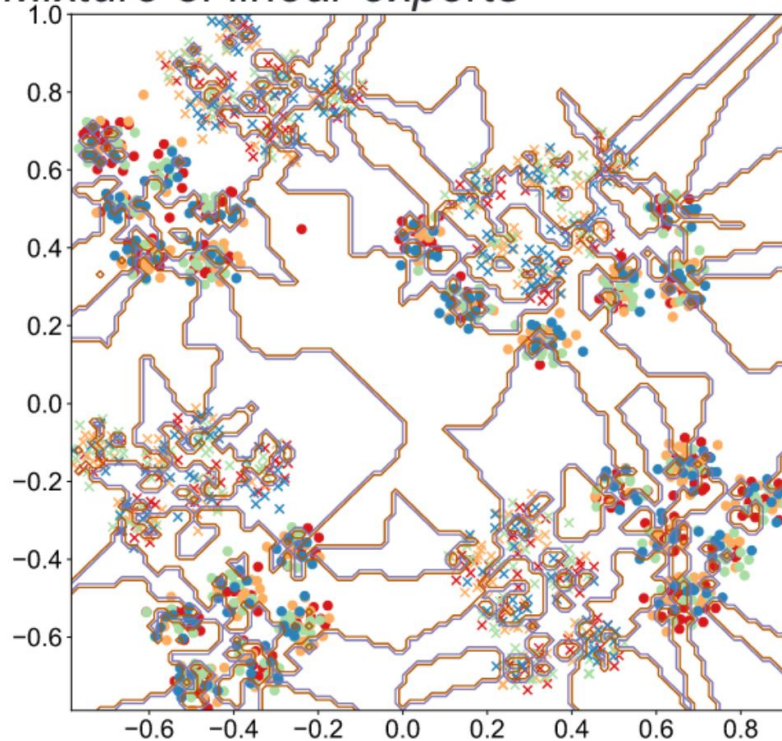


Initialization

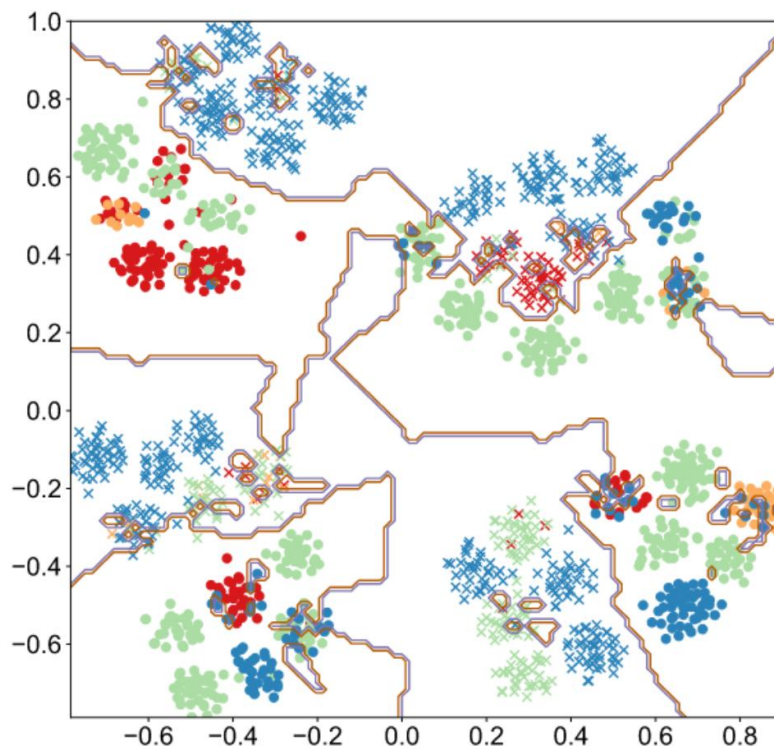


Why do experts not collapse in practice?

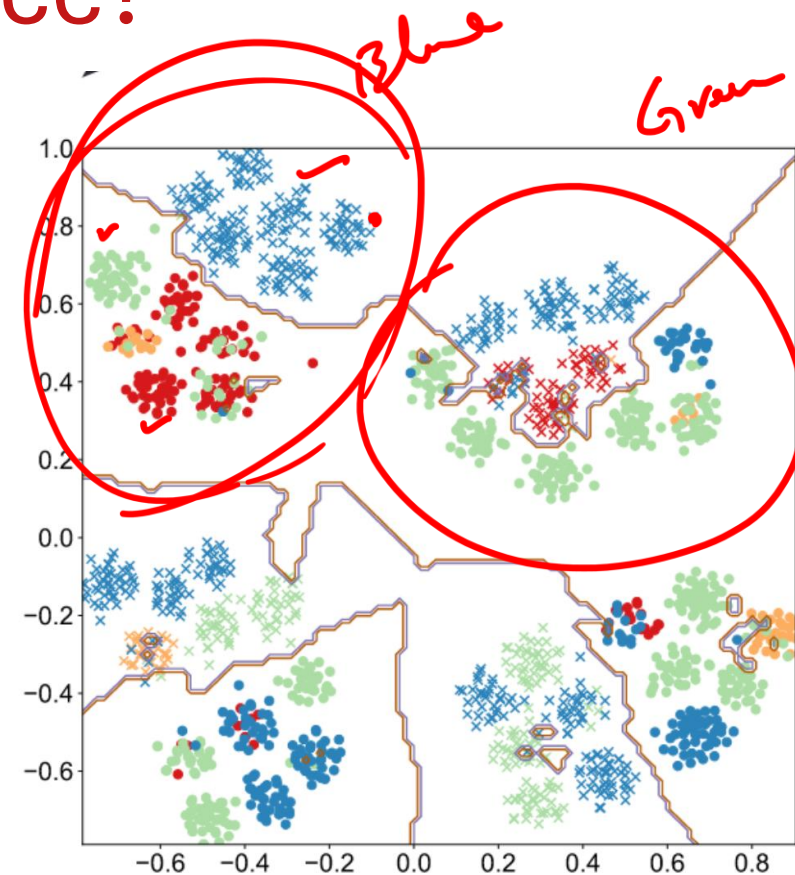
Mixture of linear experts



Initialization



Training

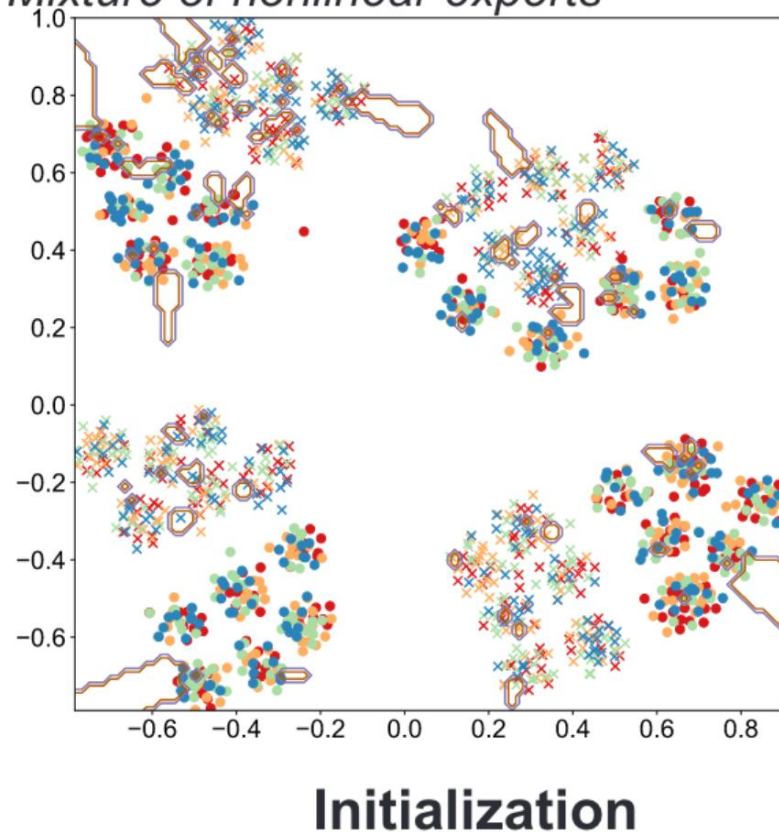


Finished



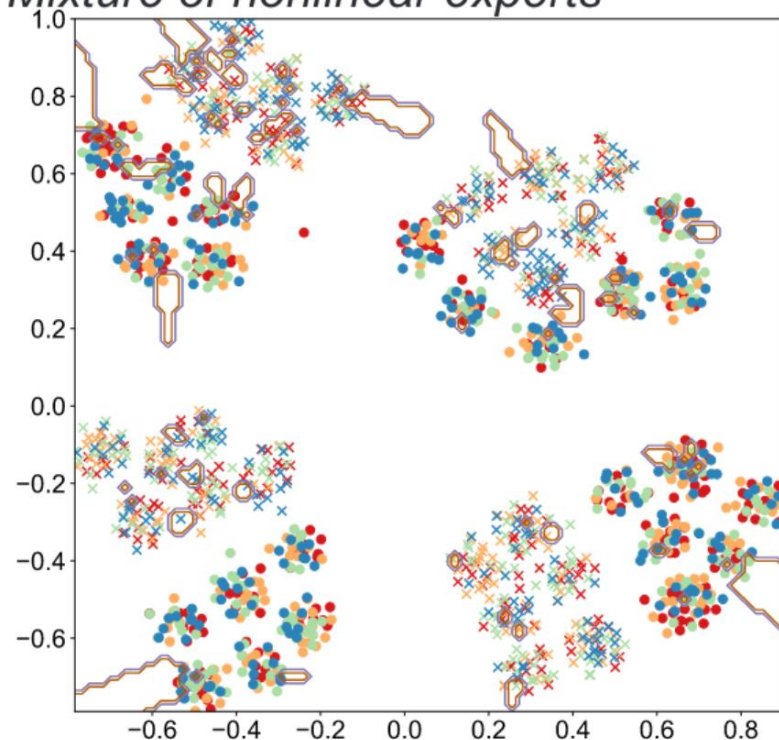
Why do experts not collapse in practice?

Mixture of nonlinear experts

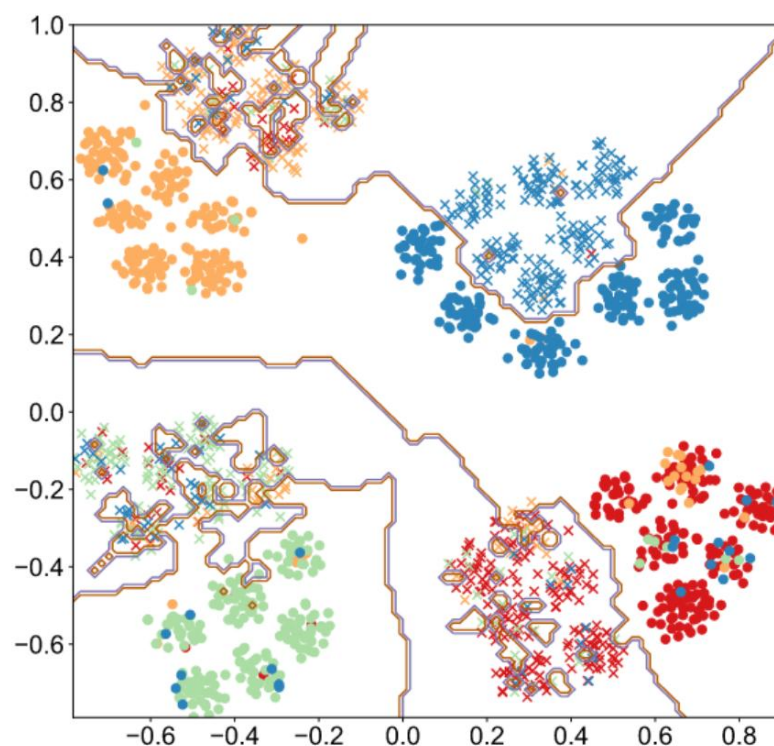


Why do experts not collapse in practice?

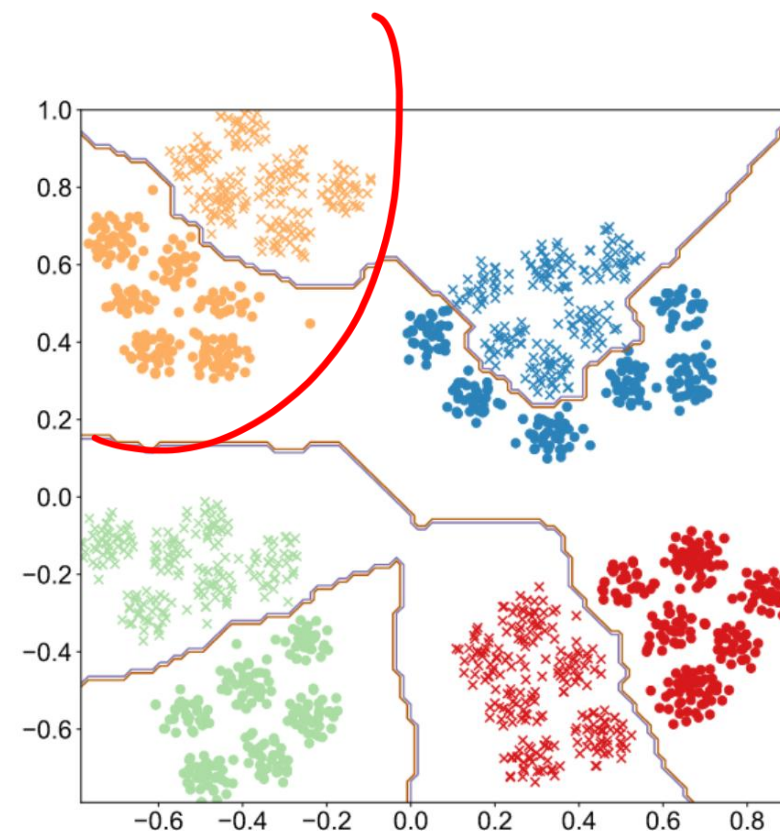
Mixture of nonlinear experts



Initialization



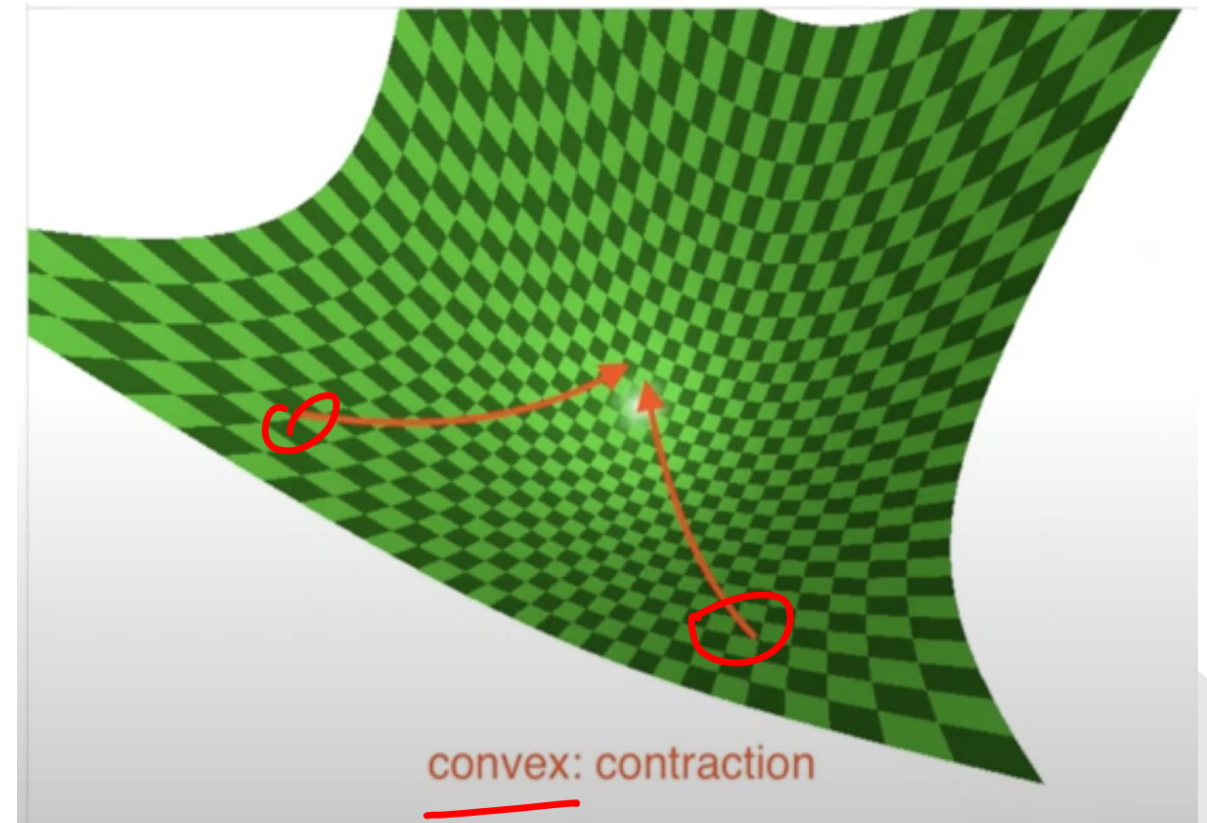
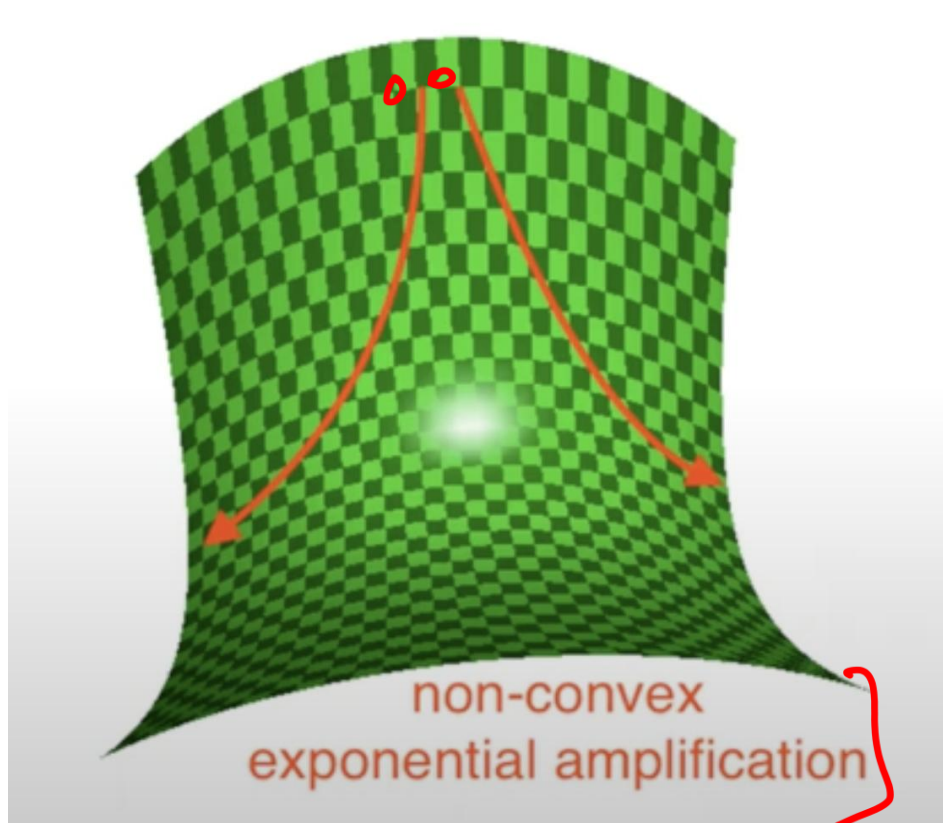
Training



Finished



Why do experts not collapse in practice?



Pros and Cons of Sparse MoE Layer

Pros

- 👍 Increased model parameters
- 👍 Efficient pretraining due to conditional (sparse) computation
- 👍 Faster inference

Cons

- 👎 High memory requirement - all parameters need to be loaded in vRAM (GPU memory)



Pros and Cons of Sparse MoE Layer

Pros

- 👍 Increased model parameters
- 👍 Efficient pretraining due to conditional (sparse) computation
- 👍 Faster inference

Cons

- 👎 High memory requirement - all parameters need to be loaded in vRAM (GPU memory)

Can we exploit Parallelism?



Parallelism ...

- Data Parallelism / Fully Sharded Data Parallelism
- Tensor Parallelism (w/ sequence parallelism)
- Context Parallelism
- Pipeline Parallelism



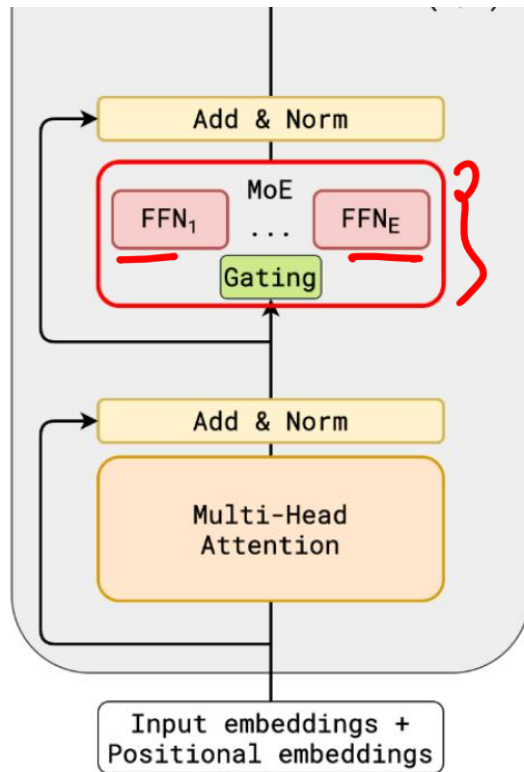
Parallelism ...

- Data Parallelism / Fully Sharded Data Parallelism
- Tensor Parallelism (w/ sequence parallelism)
- Context Parallelism
- Pipeline Parallelism

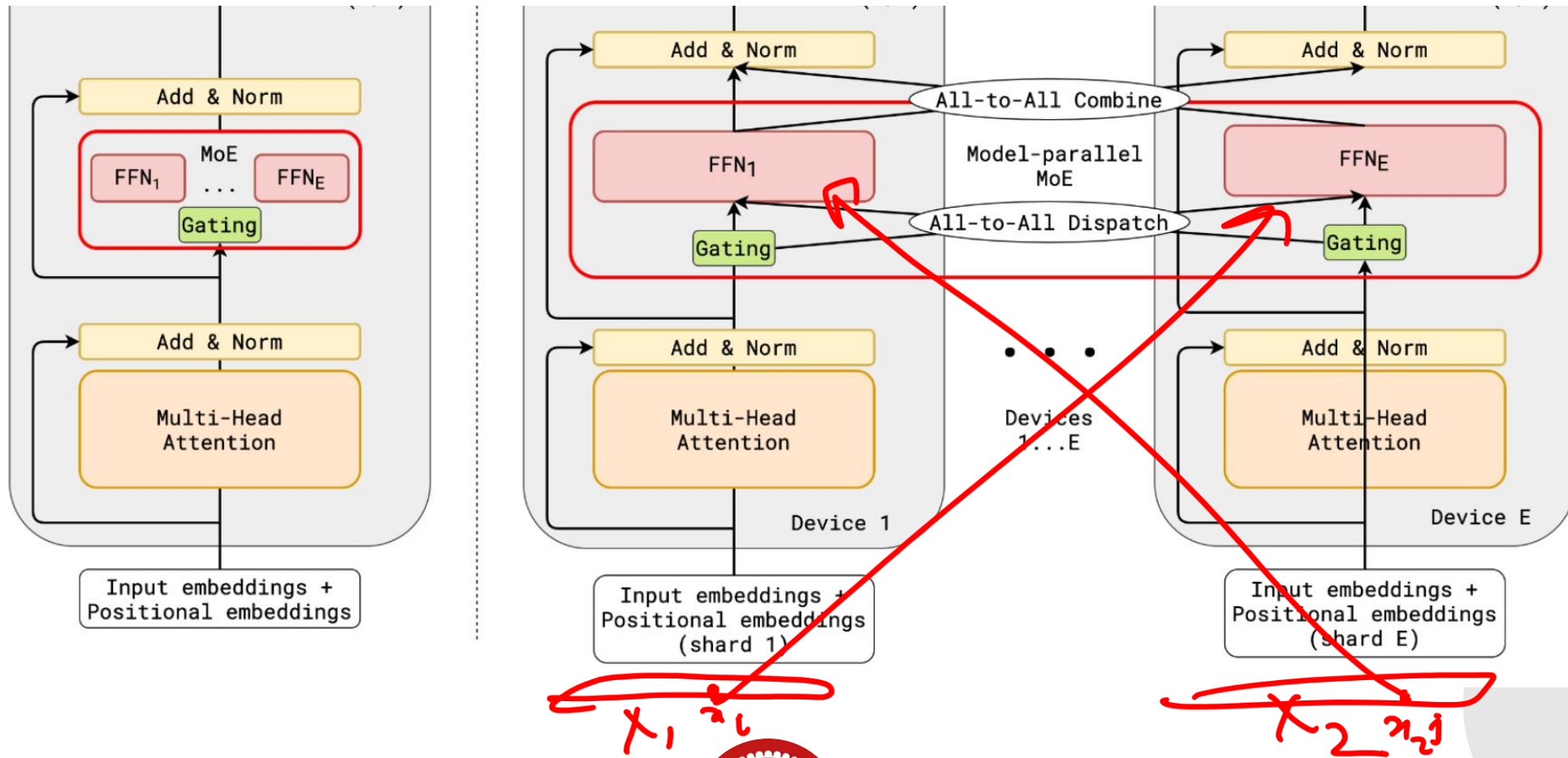
✓ • **Expert Parallelism?**



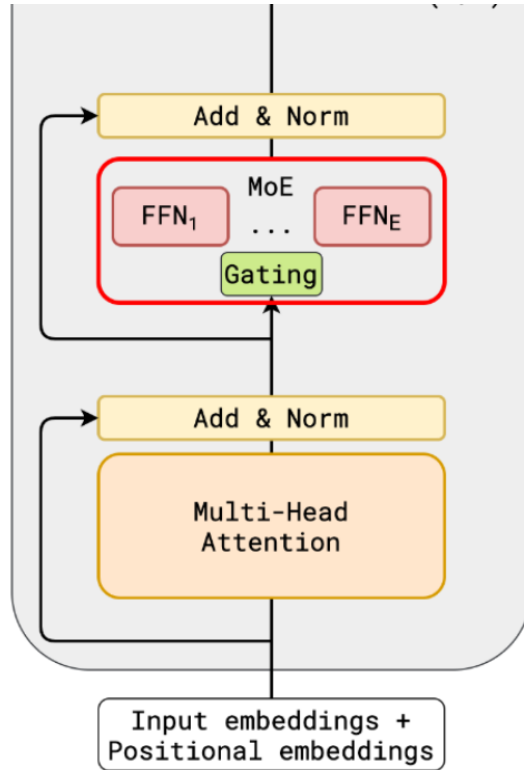
Expert Parallel for Sparse MoEs



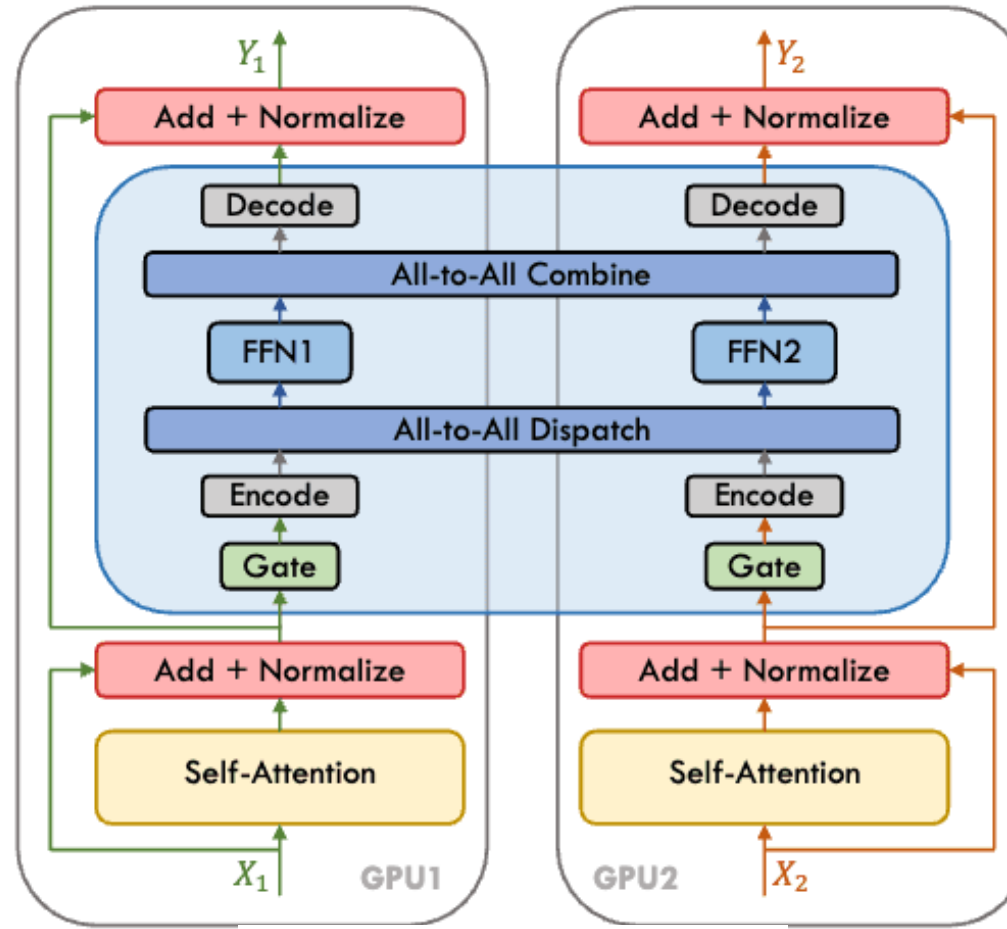
Expert Parallel for Sparse MoEs



Expert Parallel for Sparse MoEs



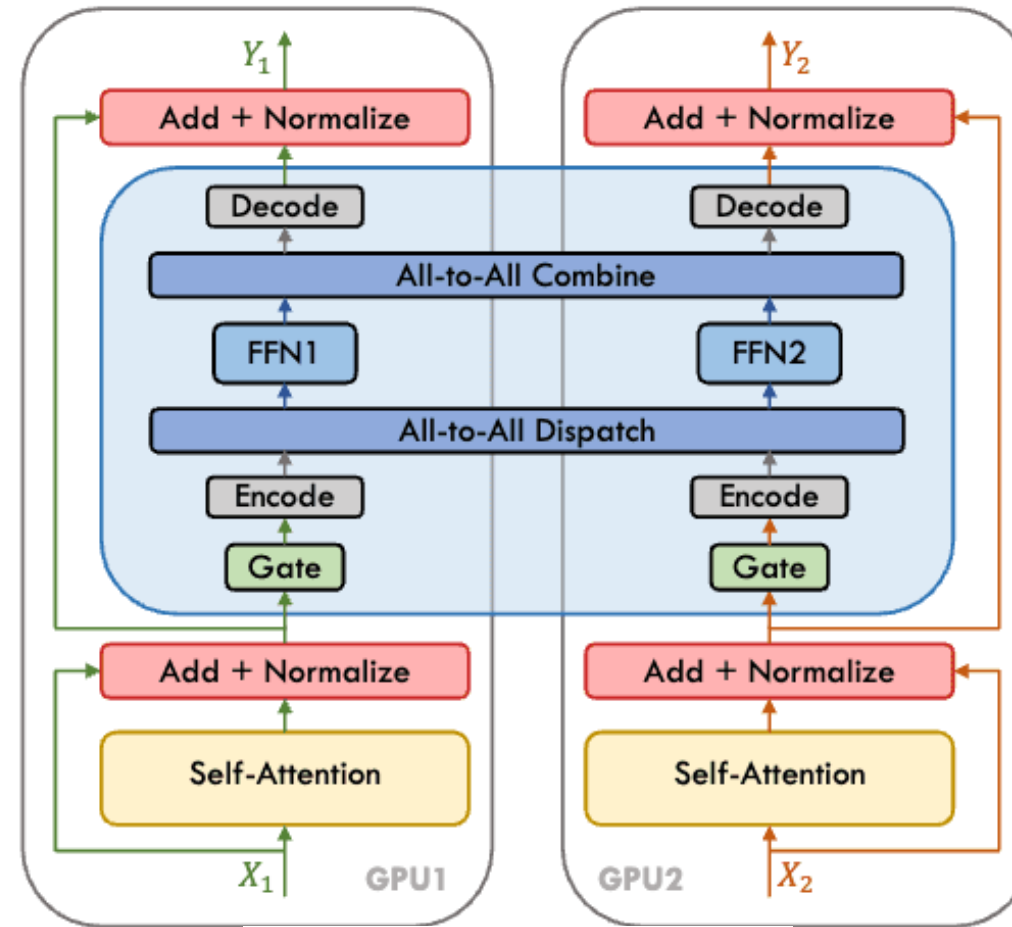
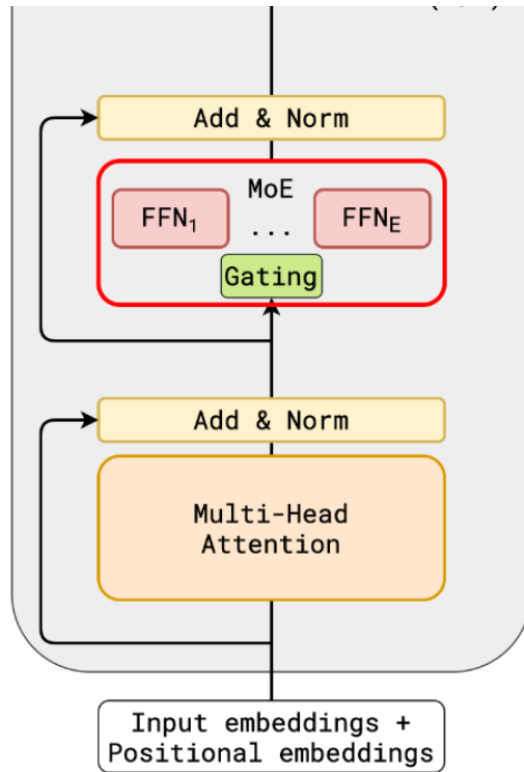
Used in conjunction with other parallelism strategies



Data + Expert Parallelism



Expert Parallel for Sparse MoEs



Data + Expert Parallelism

But what if all tokens are routed to single expert?



Pros and Cons of Sparse MoE Layer

Pros

- 👍 Increased model parameters
- 👍 Efficient pretraining due to conditional (sparse) computation
- 👍 Faster inference

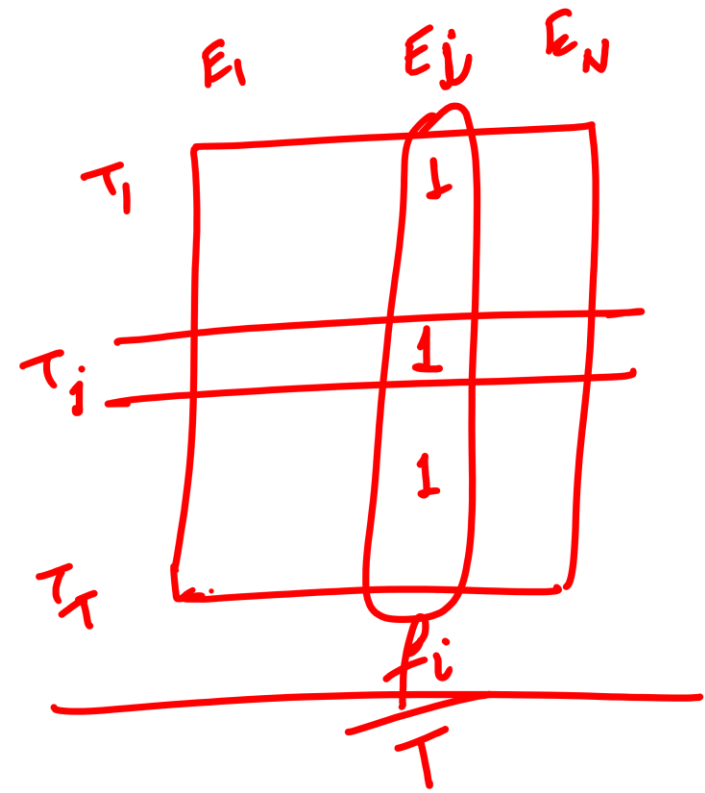
Cons

- 👎 High memory requirement - all parameters need to be loaded in vRAM (GPU memory)
- 👎 Unstable training
 - 😞 Router collapse—router sends all tokens to the same expert
 - 😞 May diverge



Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i



Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity
<https://www.youtube.com/watch?v=U8J3Z3qV8s&t=2816s>

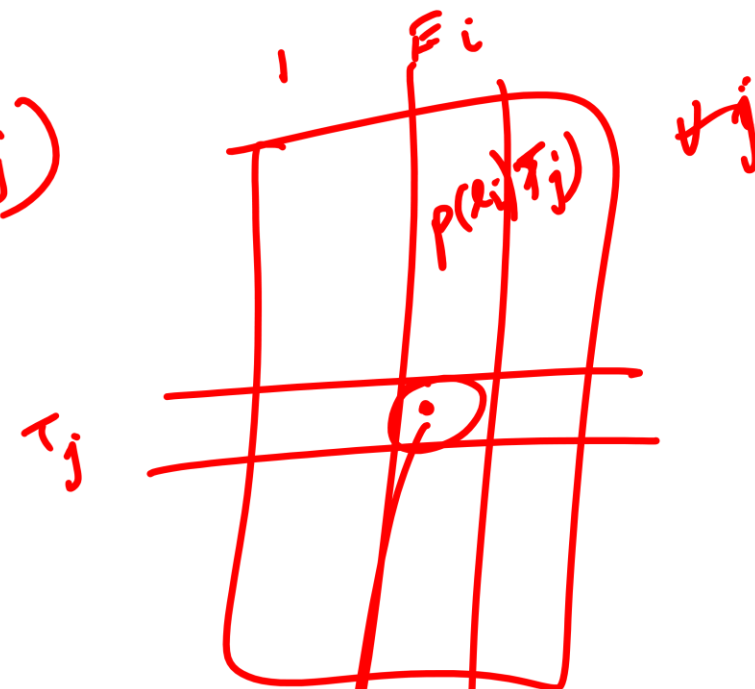


Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

$$\frac{1}{|T|} \sum_j p(e_i | \tau_j)$$



$$p(D_1) \leftarrow \begin{matrix} 5 \rightarrow 10 \\ 15 \rightarrow 20 \end{matrix}$$

$$p(D_1) \leftarrow \left(\frac{N_1 \beta_1}{N_1 + N_2} \right) \left(\frac{N_2 \beta_2}{N_1 + N_2} \right)$$

$$p(e_i) = \sum_j p(\tau_j) p(e_i | \tau_j) = \mathbb{E}(p(e_i | \tau_j))$$

Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity
<https://www.youtube.com/watch?v=U8J32Z3qV8s&t=2816s>



Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

- P_i : Expected Probability of selecting expert i

Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity
<https://www.youtube.com/watch?v=U8J3Z3qV8s&t=2816s>



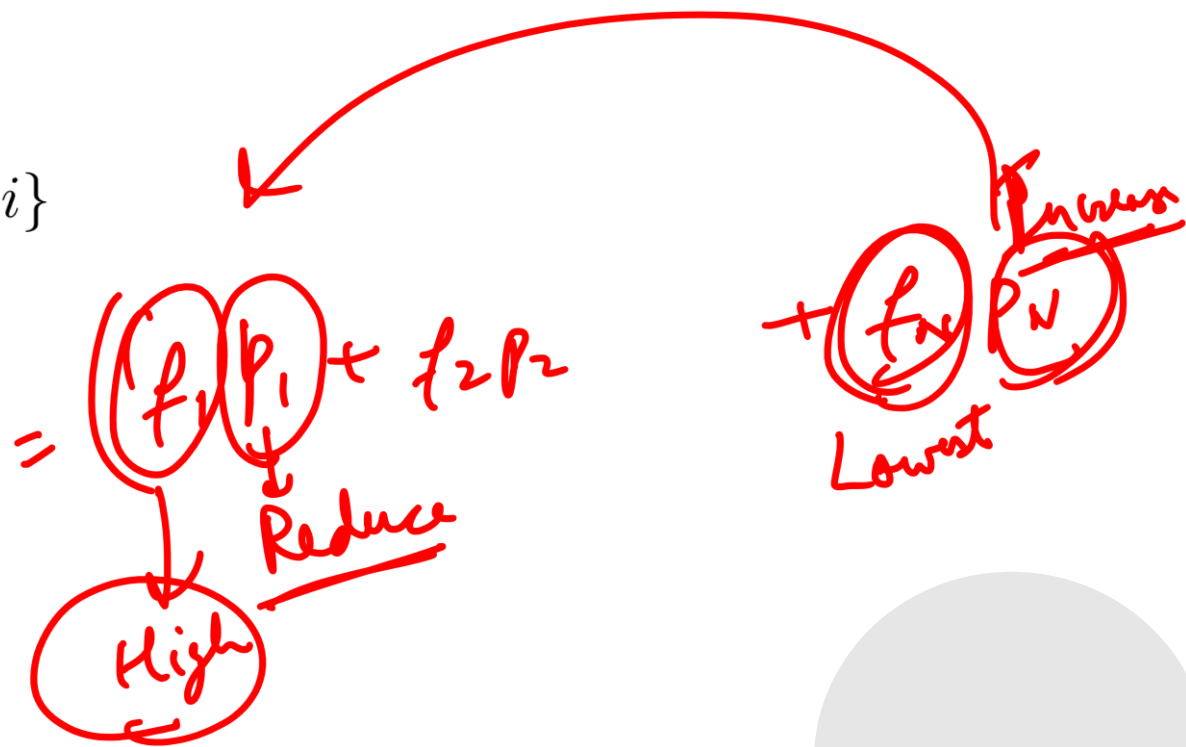
Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

- P_i : Expected Probability of selecting expert i

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



Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity
<https://www.youtube.com/watch?v=U8J32Z3qV8s&t=2816s>



Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

- P_i : Expected Probability of selecting expert i

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

Using sample mean as an empirical estimate

Content credits: Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity
<https://www.youtube.com/watch?v=U8J3Z3qV8s&t=2816s>



Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

- P_i : Expected Probability of selecting expert i

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

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



Load Balancing Loss

- N experts; T tokens in a batch $\mathcal{B}$
- f_i : Fraction of tokens dispatched to expert i

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

- P_i : Expected Probability of selecting expert i

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

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

👍 Prevents router collapse

👍 Improves training efficiency by using all the devices equally (remember that each expert is on a separate device)



Homework!

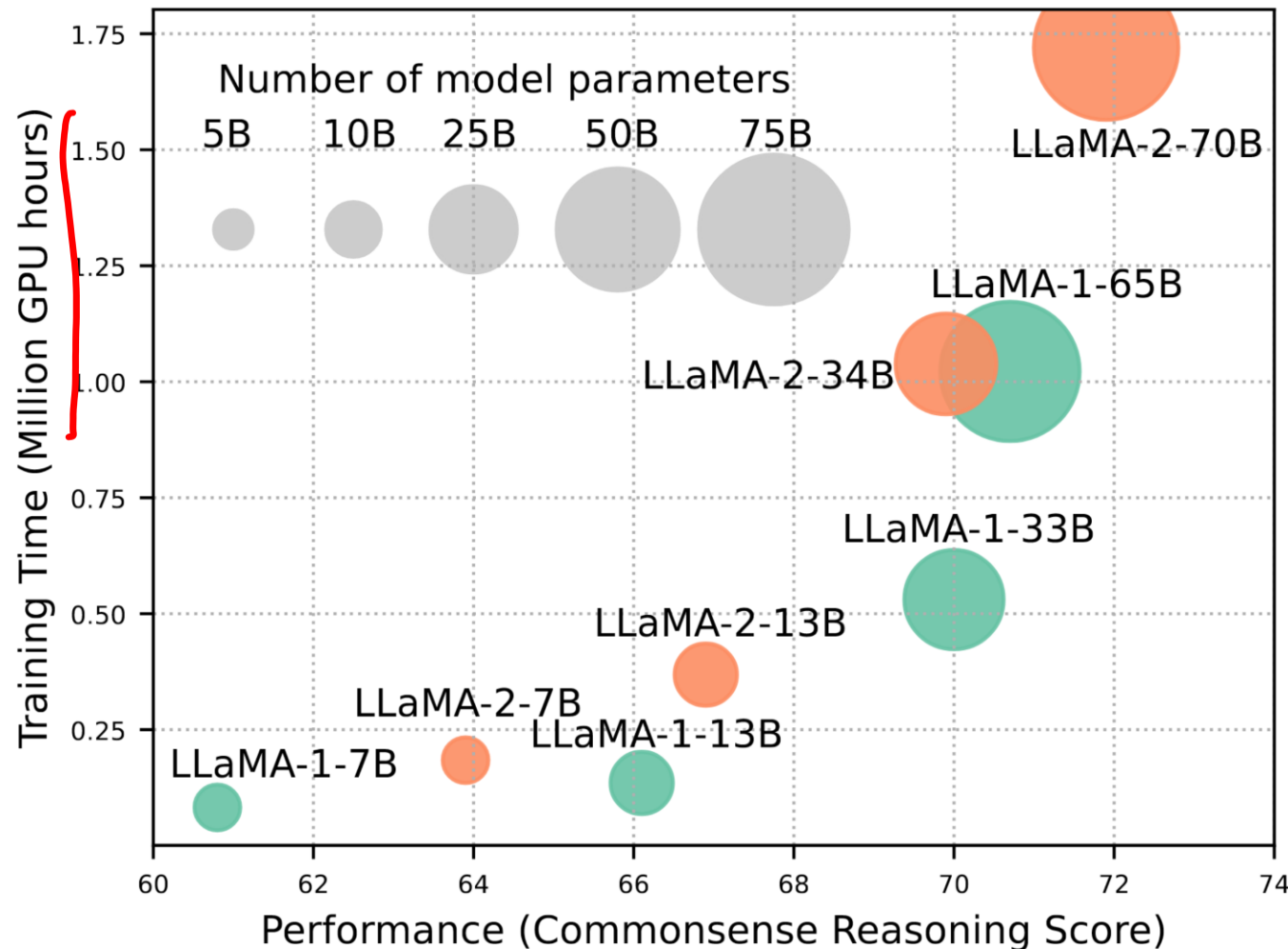
- Read the following papers –
 - [Switch Transformers: Scaling to Trillion Parameter Models with Simple and Efficient Sparsity](#)
 - [Mixtral of Experts](#)
 - [MiniMax-01: Scaling Foundation Models with Lightning Attention](#)



Summary



Training Resources vs Performance



- Based on Nvidia A100 80GB GPU

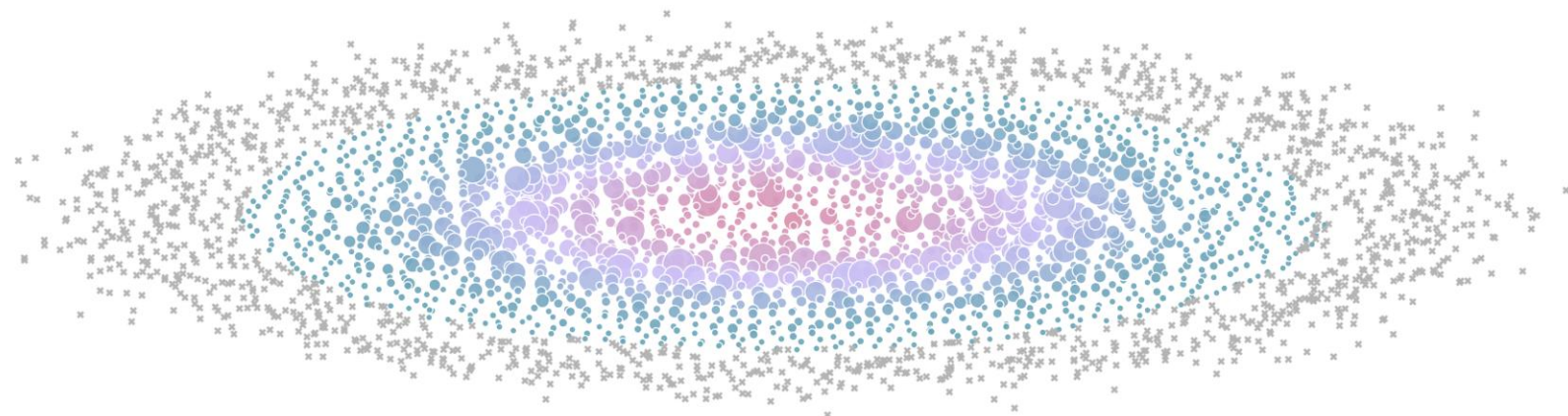
<https://huggingface.co/spaces/optimum/llm-perf-leaderboard>



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.



Summary

Efficient Training

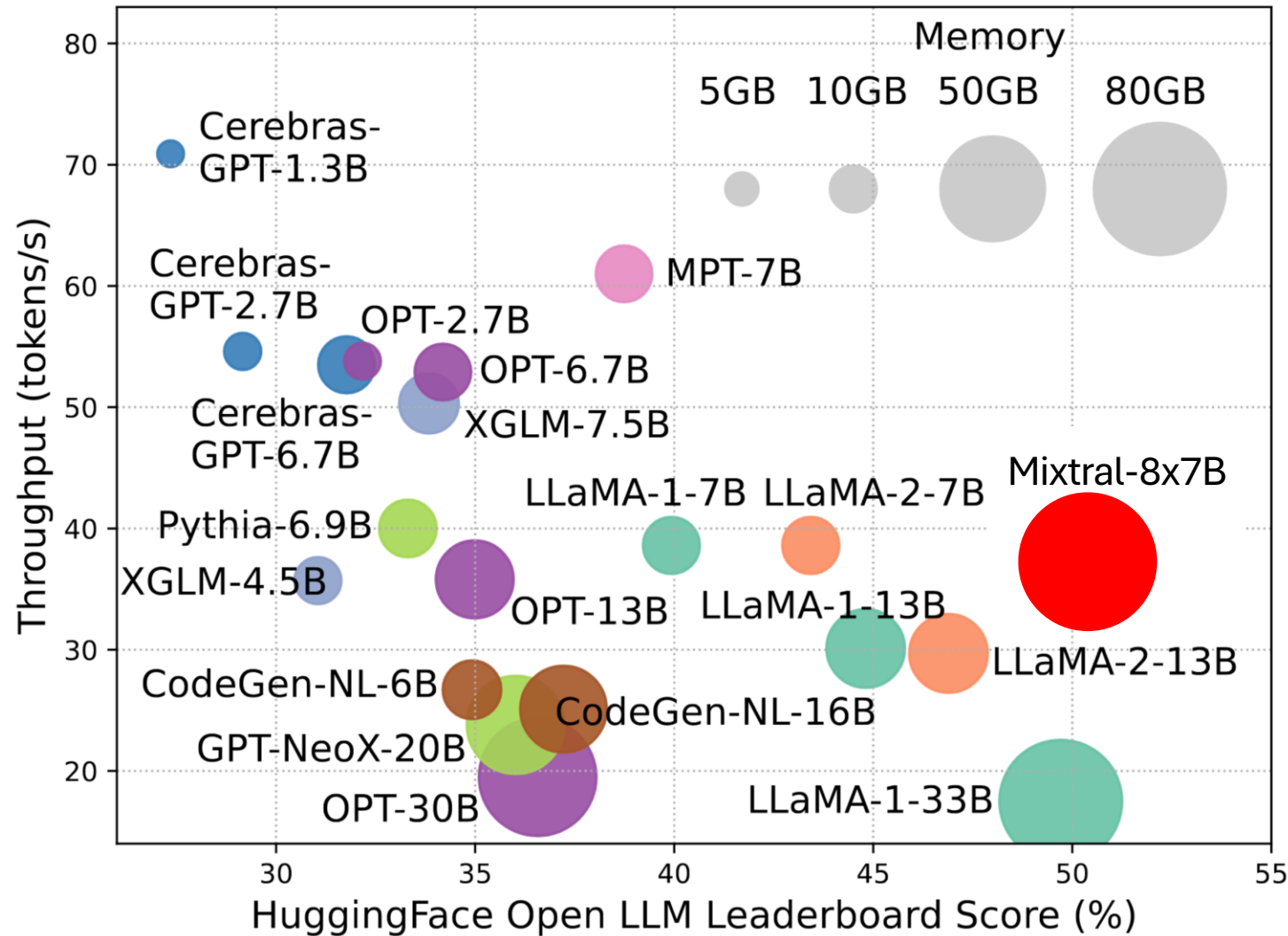
1. Parallelism for Scale -
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*



Inference Throughput vs Performance



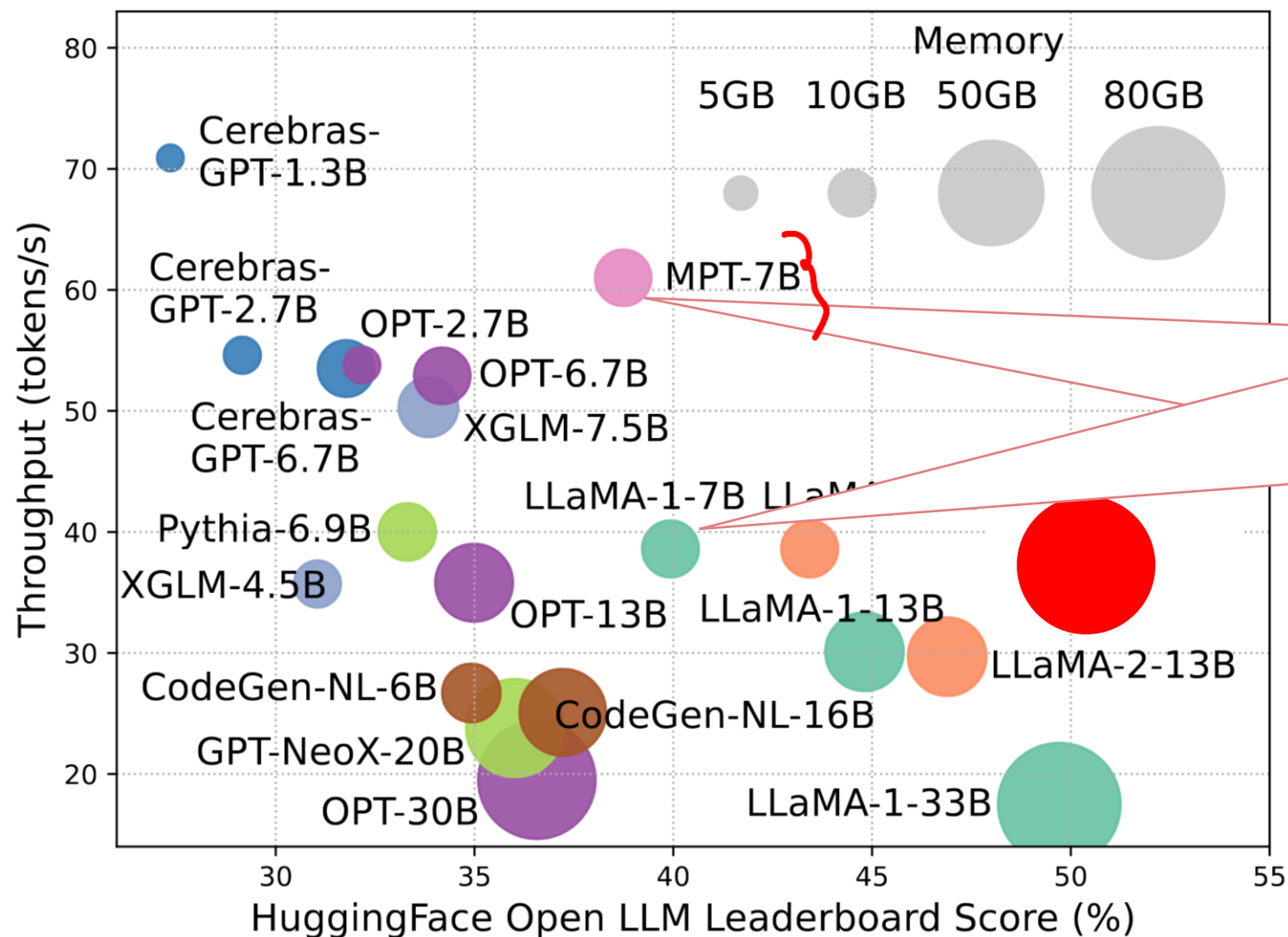
Inference Throughput vs Performance



- 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



Summary

Efficient Training

1. Parallelism for Scale -
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*
2. Efficient Implementations
 - *Flash Attention*
 - *Liger Kernels*



Summary

Efficient Training

1. Parallelism for Scale -
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*
2. Efficient Implementations
 - *Flash Attention*
 - *Liger Kernels*

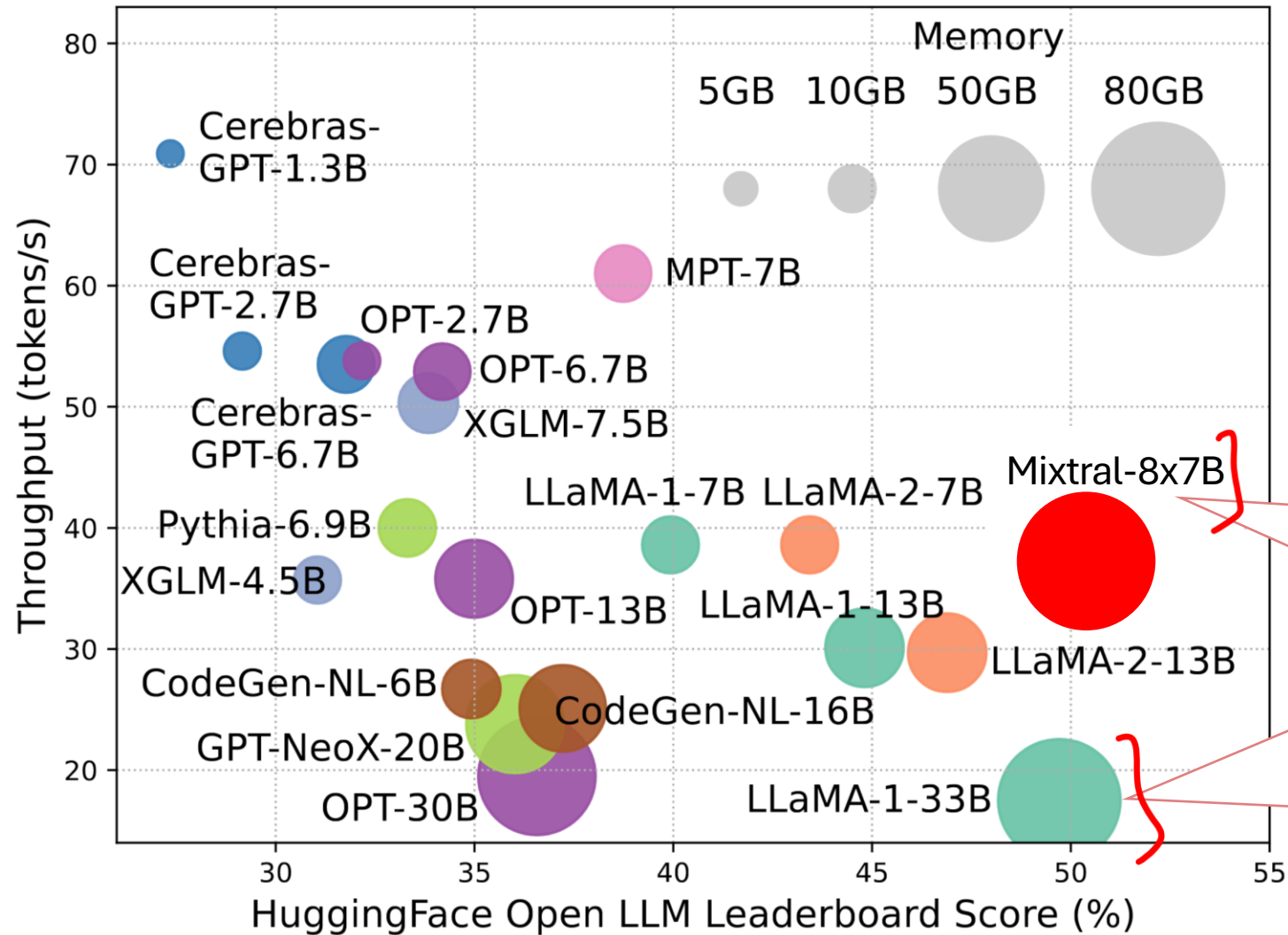
Efficient Inference

1. KV Caching
2. vLLMs and paged attention
3. Flash decoding
4. Speculative decoding

Homework!



Inference Throughput vs Performance



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



Summary

Efficient Training

1. Parallelism for Scale -
 - *Distribute model & data to fit massive training*
 - *DP, ZeRO-1/2/3, TP (w/ SP), CP, PP*
2. Efficient Implementations
 - *Flash Attention*
 - *Liger Kernels*

Efficient Inference

1. KV Caching
2. vLLMs and paged attention
3. Flash decoding
4. Speculative decoding

Homework!

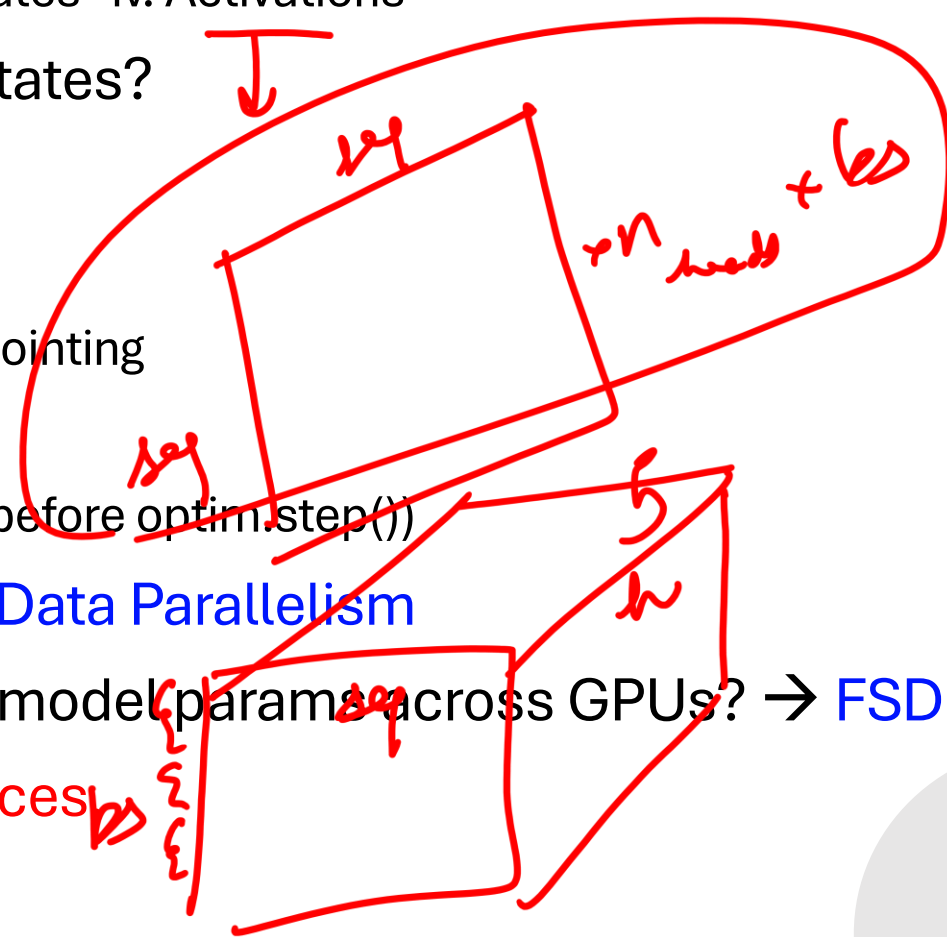
Efficient Design

1. MoE as a Layer
 1. Sparsification
 2. Expert Parallelism
 3. Load Balancing Loss
2. MoE Models
 1. Switch Transformers
 2. Mixtral of Experts
 3. Minimax

Homework!



- How to train big models on big data?
- What's in the GPU memory during training?
 - i. Model weights ; ii. Param gradients iii. Optim states iv. Activations
- What is the size of params / grads / optim states?
- What is the size of activations?
- How to reduce activation memory?
 - Activation re-computation aka Gradient checkpointing
- How to increase batch size?
 - Gradient accumulation (run fwd / bwd k times before optim.step())
- Can we parallelize grad. Accumulation? → Data Parallelism
- Can we shard the optim. states, grads, and model params across GPUs? → FSDP
- Still not good for big models & large sequences



Recap

- Can we split activations of one input across GPUs?
- Split along the hidden dimension -- let's exploit the distributive property of matrix multiplication! → **Tensor Parallelism**
- **TP: both model weights & activations are split across GPUs!**
- **TP: LayerNorm & Dropout – same computation on same data - wastage of resources :-)**
- Combine TP with **Sequence Parallel – reduce & scatter** along seq. length dimension
- **How to handle very very long sequences?**
- **Context Parallel** → split sequence into chunks & process each chunk on a different GPU (same weights but different activations on each GPU)
- How to apply attention on a sequence split on multiple GPUs? → **Ring Attention**
- **TP: Communication overhead beyond a node is prohibitive.**
- **How to handle very large models? → Pipeline Parallelism**

