

Efficient Mixture of Experts (MoE) in LLMs

: From Routing to Compression

Sehyun Park

June 5, 2026

Seoul National University - IDEA Lab.

- 1 Dense MoE
- 2 Sparse MoE
- 3 Sparse-to-Dense

Dense MoE

► Structure

- In the original MoE formulation (Jacobs et al., 1991; Jordan & Jacobs, 1994), all experts are evaluated for every input: a single network is replaced with N parallel experts and a gating network.
- Each expert is a standard FFN with independent parameters:

$$E_i(\mathbf{x}) = \text{FFN}_i(\mathbf{x}), \quad i = 1, \dots, N \quad (1)$$

- The routing produces a softmax distribution over experts, and the output is a weighted sum of all N experts:

$$\mathbf{g}(\mathbf{x}) = \text{softmax}(\mathbf{x}^\top \mathbf{W}_g) \in \mathbb{R}^N, \quad \mathbf{y} = \sum_{i=1}^N g_i(\mathbf{x}) E_i(\mathbf{x}) \quad (2)$$

※ $\mathbf{x} \in \mathbb{R}^d$: input representation; $\mathbf{W}_g \in \mathbb{R}^{d \times N}$: gating matrix.

► Objective

- Train end-to-end by adding a **balance loss** to the task loss to equalize expert load:

$$\mathcal{L} = \mathcal{L}_{\text{task}} + \lambda \mathcal{L}_{\text{bal}}, \quad \mathcal{L}_{\text{bal}} = \text{CV}(\mathbf{I})^2 = \frac{\text{Var}_i(I_i)}{(\text{Mean}_i(I_i))^2} \quad (3)$$

※ $I_i = \sum_{\mathbf{x} \in \mathcal{B}} g_i(\mathbf{x})$: cumulative gating weight (importance) received by expert i over batch $\mathcal{B}$.

Sparse MoE

Sparse MoE: Token-based Routing

► Motivation

- Dense MoE is computationally expensive at inference: all experts run for every token.
- Token-based routing: the router decides, *per token*, which k experts to activate.

► Top- k Gating

Shazeer et al. (2017); Switch Transformer (2022); Mixtral (2024)

- Only top- k experts ($k \ll N$) are active per token. FLOPs reduce from $O(N \cdot T)$ to $O(k \cdot T)$.

$$G(\mathbf{x}) = \text{softmax}(\text{TopK}(\mathbf{x} \cdot \mathbf{W}_g + \boldsymbol{\epsilon})) , \quad \mathbf{y} = \sum_{i=1}^N G_i(\mathbf{x}) \cdot E_i(\mathbf{x}) \quad (4)$$

- Switch Transformer (Fedus et al., 2022): simplifies to $k = 1$, showing that single-expert routing is sufficient for large-scale pre-training.
- Mixtral (Jiang et al., 2024): uses $k = 2$ with 8 experts; achieves performance comparable to larger dense models at a fraction of the active parameters.

Sparse MoE: Expert-based Routing

► Mixture-of-Experts with Expert Choice Routing

Zhou et al. (NeurIPS 2022)

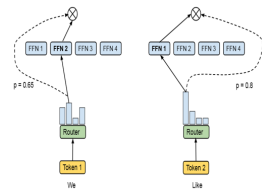
- Instead of tokens selecting experts, **each expert selects the top- k tokens** from the batch.

$$\mathbf{S} = \text{Softmax}(\mathbf{X}\mathbf{W}_g) \in \mathbb{R}^{T \times N}, \mathbf{I}^{(i)} = \text{TopK}(\mathbf{S}_{:,i})$$

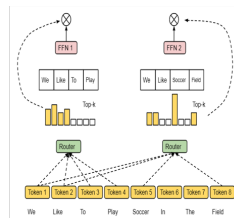
※ $\mathbf{X} \in \mathbb{R}^{T \times d}$: token matrix, $\mathbf{I}^{(i)}$: index set of tokens selected by expert i .

► Experiment Results

- Empirically achieves $> 2\times$ faster training convergence compared to Switch Transformer (top-1) and GShard (top-2) under the same compute budget.



◀Token-based MoE▶



◀Expert-based MoE▶

Sparse-to-Dense

Sparse-to-Dense: Motivation

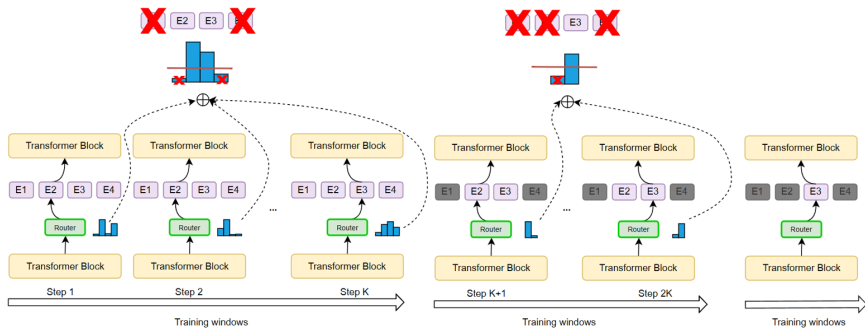
- Sparse MoE reduces *compute* per token, but all expert parameters must reside in memory.
- A Mixtral-8×7B model stores $\approx 46.7\text{B}$ parameters; deploying it requires $\sim 87\text{GB}$ GPU memory.
- Sparse-to-Dense research compresses a trained Sparse MoE into one or a small number of dense models:

Method	Key idea
Pruning	Remove low-importance experts
Knowledge Distillation	Train a single FFN student
Expert Merging (Sub-MoE)	Merge experts

Sparse-to-Dense: Pruning

► TSEP: Task-Specific Expert Pruning

Chen et al. (2022)



〈The pipeline of Task-Specific Expert Pruning〉

Sparse-to-Dense: Pruning

► TSEP: Task-Specific Expert Pruning

Chen et al. (arxiv, 2022); Sarkar et al. (NeurIPS, 2024)

- Observation: for a specific downstream task, only a subset of experts are consistently activated; the rest contribute little.
- An *expert proficiency score* is defined per expert i on task dataset $\mathcal{D}$:

$$\text{score}(E_i) = \sum_{\mathbf{x} \in \mathcal{D}} G_i(\mathbf{x}) \quad (5)$$

※ $G_i(\mathbf{x})$ is indicator that expert i is selected for input $\mathbf{x}$.

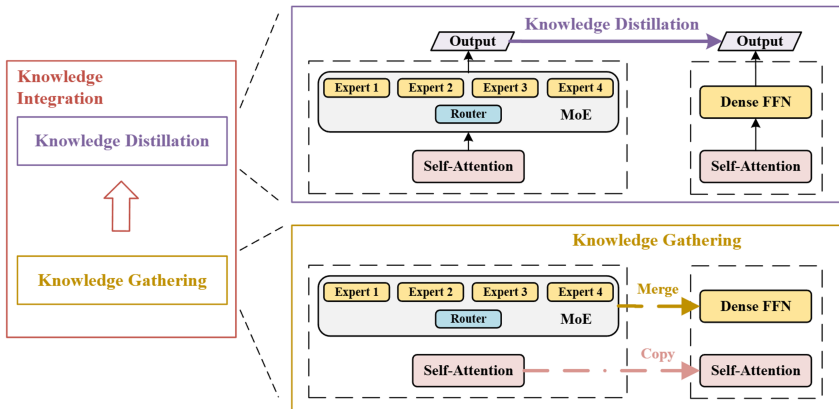
► Results

- The resulting single-expert model preserves 99.3% of MoE performance across six downstream tasks while achieving $2\times$ inference speedup by eliminating expert parallelism overhead.
- However, Sarkar et al. (2024) show that naive activation-based pruning fails to recover the performance of a scratch-trained SMoE of equivalent size.

Sparse-to-Dense: Knowledge Distillation

► OneS: One Student Knows All Experts Know

Xue et al. (arxiv, 2022)



〈An overview of OneS〉

Sparse-to-Dense: Knowledge Distillation

► OneS: One Student Knows All Experts Know

Xue et al. (arxiv, 2022)

- Goal: obtain a **single dense student model** that matches the knowledge of a sparse MoE teacher.
- Two-stage framework:
 - ① **Knowledge Gathering**: aggregate N expert weights $\{\mathbf{W}_i\}_{i=1}^N$ into one set of parameters:

- SVD-KG:
$$\mathbf{W}_{\text{student}} = \frac{1}{N} \sum_{i=1}^N \mathbf{U}_{r,i} \mathbf{\Sigma}_{r,i} \mathbf{V}_{r,i}^\top,$$

where $\mathbf{W}_i \approx \mathbf{U}_{r,i} \mathbf{\Sigma}_{r,i} \mathbf{V}_{r,i}^\top$ is the rank- r SVD approximation of the i -th expert

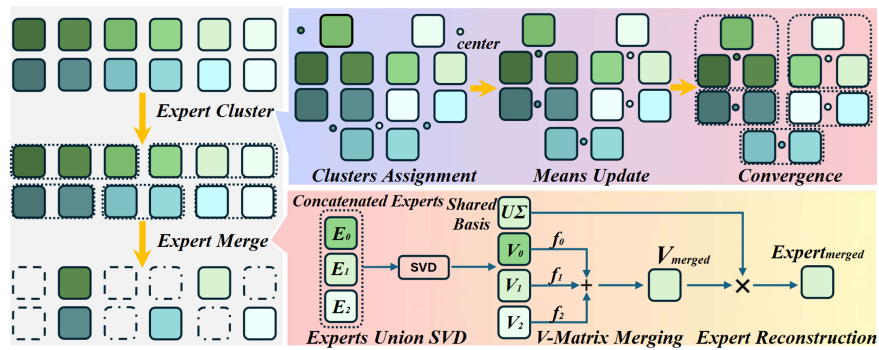
- ② **Knowledge Distillation**: train the student to minimize the following objective:

$$\mathcal{L} = (1 - \lambda) \mathcal{L}_{\text{CE}} + \lambda \text{KL}(p_{\text{teacher}}(\mathbf{x}; T) \parallel p_{\text{student}}(\mathbf{x}; T)) \quad (6)$$

Sparse-to-Dense: Expert Merging (Sub-MoE)

► Sub-MoE: Efficient Mixture-of-Expert LLMs Compression

Li et al. (AAAI 2026)



◁An overview of Sub-MoE▷

Sparse-to-Dense: Expert Merging (Sub-MoE)

► Sub-MoE: Efficient Mixture-of-Expert LLMs Compression

Li et al. (AAAI 2026)

- Rather than producing a single dense model, Sub-MoE merges groups of experts into fewer experts within the MoE structure, preserving the routing mechanism.
- Core problem: direct weight averaging causes *parameter conflicts* because experts operate in distinct representation spaces (inter-expert similarity ≈ 0.1 – 0.3 in Mixtral-8 $\times$ 7B).

► Stage 1: Adaptive Expert Clustering

- Experts that respond similarly to the same inputs are better merge candidates. Functional similarity is measured as:

$$\text{Sim}(E_i, E_j) = \frac{1}{|\mathcal{X}|} \sum_{\mathbf{X} \in \mathcal{X}} \frac{1}{|\mathbf{X}|} \sum_{\mathbf{x} \in \mathbf{X}} \frac{E_i(\mathbf{x}) \cdot E_j(\mathbf{x})}{\|E_i(\mathbf{x})\| \|E_j(\mathbf{x})\|} \quad (7)$$

- Experts are grouped via **K-means clustering** on this similarity.

Sparse-to-Dense: Expert Merging (Sub-MoE)

► Stage 2: Subspace Expert Merging

- **Experts Union Decomposition:** for each cluster, concatenate expert weights and apply joint SVD to find a shared subspace:

$$\text{SVD}([\mathbf{W}_1; \dots; \mathbf{W}_n]) = \mathbf{U}\mathbf{\Sigma}[\mathbf{V}^{(1)}; \dots; \mathbf{V}^{(n)}]^\top \quad (8)$$

The shared $\mathbf{U}$ aligns all experts; conflicts are isolated to the expert-specific $\mathbf{V}^{(i)}$.

- **Frequency-based V-Matrix Merging:** merge expert-specific components weighted by router activation frequency:

$$f_i(\mathbf{X}) = \frac{\sum_{\mathbf{x} \in \mathbf{X}} \mathbb{I}[i \in \text{TopK}(G(\mathbf{x}))]}{|\mathbf{X}|}, \quad \mathbf{V}_{\text{merged}} = \frac{\sum_{i=1}^N f_i(\mathbf{X}) \cdot \mathbf{V}^{(i)}}{\sum_{i=1}^N f_i(\mathbf{X})} \quad (9)$$

- **Expert Reconstruction:**

$$\mathbf{W}_{\text{merged}} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}_{\text{merged}}^\top \quad (10)$$

- On Mixtral-8×7B: 50% expert reduction with 86% performance retained.

End