

Model Merging in the Essential Subspace

Longhua Li^{1,2} Lei Qi^{1,2*} Qi Tian³ Xin Geng^{1,2*}

¹School of Computer Science and Engineering, Southeast University, Nanjing, China

²Key Laboratory of New Generation Artificial Intelligence Technology and Its Interdisciplinary Applications (Southeast University), Ministry of Education, China

³Huawei Technologies, Shanghai, China

lhli@seu.edu.cn, qilei@seu.edu.cn, tian.qil@huawei.com, xgeng@seu.edu.cn

Abstract

Model merging aims to integrate multiple task-specific fine-tuned models derived from a shared pre-trained checkpoint into a single multi-task model without additional training. Despite extensive research, task interference remains a major obstacle that often undermines the performance of merged models. In this paper, we propose ESM (Essential Subspace Merging), a robust framework for effective model merging. We begin by performing Principal Component Analysis (PCA) on feature shifts induced by parameter updates. The resulting principal directions span an essential subspace that dominantly influences feature representations. Each task’s parameter update matrix is projected onto its respective essential subspace for low-rank decomposition before merging. This methodology mitigates inter-task interference while preserving core task-specific functionality. Furthermore, we introduce a multi-level polarized scaling strategy that amplifies parameters containing critical knowledge and suppresses redundant ones, preventing essential knowledge from being overwhelmed during fusion. Extensive experiments across multiple task sets and model scales demonstrate that our method achieves state-of-the-art performance in multi-task model merging. Code is available at <https://github.com/kiddo127/ESM>.

1. Introduction

In recent years, the pre-training–fine-tuning paradigm has revolutionized performance across numerous downstream tasks, enabling models to excel by adapting to task-specific data. This success has spurred the creation of many specialized expert models. The growth of public model repositories is now fueling demand for techniques that reuse and integrate these models. Among them, model merging

[18, 39, 46, 53] has emerged as a promising, parameter-efficient approach to combine multiple experts into a single versatile model, thus marking a pivotal advance toward the development of general-purpose artificial intelligence.

However, integrating the capabilities of multiple models remains challenging. A key difficulty is inter-task interference, as fine-tuning on different tasks often drives parameters in conflicting directions. Simple averaging methods such as Model Soup [46] mix these updates, diluting task-specific knowledge and leading to suboptimal performance. To alleviate this problem, subsequent studies analyzed task vectors, defined as the parameter differences between fine-tuned and pre-trained models [18, 19, 29, 52]. More recently, advanced methods have employed Singular Value Decomposition (SVD) on task vectors to identify low-rank subspaces that compactly represent task knowledge and reduce redundancy during merging [11, 28, 38, 57]. However, these SVD-based subspaces are not well aligned with task feature spaces and exhibit limited knowledge concentration.

In this paper, we first theoretically analyze the limitations of SVD-based low-rank decomposition for task-specific parameter updates. Although these methods truncate the smallest singular values, they overlook the feature distribution. Consequently, when input tokens are aligned with the truncated singular vectors, significant truncation errors may still occur, as illustrated in Equation 4. Moreover, when merging a large number of task-specific update parameters, the core knowledge of each task may be overwhelmed. As shown in Figure 2 and 3, we observe that the magnitude (norm) of task parameter updates is closely correlated with their directional importance: large-magnitude updates often correspond to critical adjustments essential for capturing task-specific or shared knowledge across tasks, whereas small-magnitude updates, though potentially negligible for individual task performance, may obscure important task knowledge after model merging.

Building on these insights, our method consists of two components. Instead of directly decomposing the task ma-

*Corresponding authors.

trix (i.e., the parameter update matrix), we first perform Principal Component Analysis (PCA) on the proxy activation shifts induced by the task matrix. The matrix is then decomposed along the resulting principal directions, forming the task’s Essential Subspace, which is directly connected to the output feature space and ensures minimal feature loss compared to other low-rank subspaces of the same rank. Each task matrix is subsequently projected onto its Essential Subspace before fusion. Experiments demonstrate that this subspace aligns more closely with the task’s feature representation, exhibits stronger energy concentration, and isolates task-specific knowledge more effectively than conventional SVD-based subspaces. Next, we introduce a simple yet effective Polarized Scaling mechanism. It adaptively rescales task matrices at three levels: across layers, tasks, and dimensions, amplifying high-norm (strong-signal) components while suppressing low-norm (weak-noise) ones. This polarization effectively reduces inter-task interference during model fusion and ensures the preservation of essential task knowledge in the merged model.

Our main contributions are summarized as follows:

- We propose a decomposition method that builds an essential subspace aware of activation shift distributions, aligning closely with feature distribution and concentrating energy more effectively, thereby better preserving task-specific knowledge and reducing cross-task interference.
- We show that the pre-fusion and post-fusion task matrix norms reflect gradient confidence and inter-task consensus, respectively. Based on this, we introduce a multi-level polarized scaling mechanism that prevents crucial parameters from being overshadowed during merging.
- Extensive experiments show that our method outperforms existing model merging methods and effectively narrows the gap between the merged model and the expert models.

2. Related Work

2.1. Model Merging

Model merging aims to combine multiple task-specific fine-tuned models into a unified multi-task model without retraining. To ensure merging stability, recent studies typically merge models that are fine-tuned from the same pre-trained checkpoint. Model Soup [46] averages fine-tuned weights to improve generalization, while Task Arithmetic [18] introduces task vectors, defined as the parameter differences between fine-tuned and pre-trained models, to enable vector-based knowledge composition. However, direct averaging of task vectors often causes severe task interference due to conflicting updates. To address this, TIES-Merging [52] trims redundant parameters before averaging salient ones, AdaMerging [54] learns adaptive task-wise coefficients, and DARE [55] resets redundant updates while rescaling the rest. Information-weighted methods such as

Fisher Merging [29] and RegMean [19] use Fisher information or input similarity for weighted averaging. Other works refine merging through parameter- or layer-wise strategies [10, 43, 56], or leverage implicit or modular representations to enhance flexibility [3, 16, 58].

Recent advances move beyond raw parameter space to the spectral domain. TSV-M [11] perform Singular Value Decomposition (SVD) on task matrices and merge along the top singular directions that capture dominant functional subspaces. Iso-CTS [28] constructs an isotropic common subspace through singular value normalization followed by task-specific refinements, achieving state-of-the-art performance. However, singular values reflect only the parameter energy rather than their functional impact. To overcome this limitation, we decompose each task matrix within an essential subspace derived from its effect on output activations, which better captures task-specific features. This leads to lower truncation error and improved preservation of task knowledge during merging. Moreover, we introduce a polarized scaling mechanism that amplifies strong signals and suppresses noise, further enhancing fusion performance.

2.2. Model Weight Transformation

Transforming model weights has been extensively studied in areas such as model compression [17, 25, 45] and efficient knowledge transfer [24, 47, 48, 51]. One of the most popular approaches to model weight transformation is based on the low-rank assumption of weight matrices. The LoRA family of methods [8, 9, 15, 27] assumes that fine-tuning updates are inherently low-rank and learns compact matrices to parameterize these updates. Other methods use SVD-based decompositions for parameter-efficient fine-tuning [13, 40] or model compression [26, 35, 44, 45]. More recently, low-rank decompositions have been applied to model merging [11, 28, 38], combining task updates in reduced subspaces to mitigate inter-task interference.

In contrast, our proposed Essential Subspace Decomposition (ESD) constructs the decomposition space not from the weight updates themselves but from the activation shifts induced by these updates. By capturing task-specific principal directions in the activation space, ESD produces sparse yet expressive task representations, reducing cross-task interference while preserving high task fidelity.

3. Methodology

3.1. Preliminaries on Model Merging

Model merging aims to integrate a collection of task-specific models, each fine-tuned from a common pre-trained checkpoint, into a single unified model without additional retraining. Formally, let θ_0 denote the parameters of the pre-trained model, and θ_t be the parameters of the expert model fine-tuned on task t , where $t = 1, \dots, T$. The fundamen-

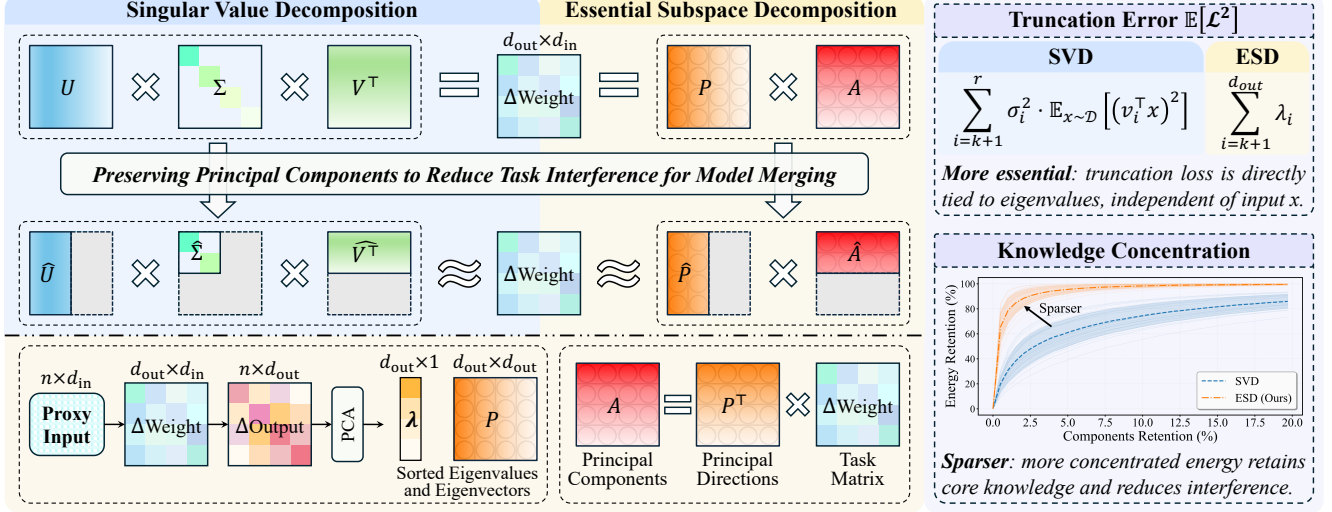


Figure 1. Essential Subspace Decomposition (ESD) versus Singular Value Decomposition (SVD). Unlike SVD, which decomposes the task matrix solely based on weights, ESD decomposes them based on feature shift distributions. When truncating components for merging, ESD’s expected truncation error is directly related to the magnitude of the discarded eigenvalues and yields higher knowledge retention.

tal object of interest in model merging is the task update, which captures how fine-tuning shifts the model away from the pre-trained weights. Following Task Arithmetic [18], the task vector for task t is defined as:

$$\tau_t = \text{Flatten}(\theta_t - \theta_0). \quad (1)$$

Given the structured nature of models, it is preferable to retain the matrix form of the update rather than flattening it into a vector. The task matrix of layer ℓ is defined as:

$$\Delta_{W_t}^{(\ell)} = \theta_t^{(\ell)} - \theta_0^{(\ell)}. \quad (2)$$

Each $\Delta_{W_t}^{(\ell)}$ represents the task-specific parameter update at layer ℓ , preserving the row-column structure essential for spectral analysis and subspace alignment. The goal of model merging is to construct merged weights θ_M that support all tasks, typically in the form:

$$\theta_M = \theta_0 + f(\Delta_{W_1}, \dots, \Delta_{W_T}), \quad (3)$$

where $f(\cdot)$ is a merging function, which is the main focus of current model merging research [11, 18, 28, 38].

3.2. Model Merging in Essential Subspace

Unlike previous methods that merge models in the truncated singular vector subspace obtained via SVD, we propose to decompose and merge task matrices within a more essential subspace that is better aligned with the task’s output feature space, as illustrated in Figure 1. For simplicity, unless otherwise specified, we omit the layer index ℓ and task identifier t in the task matrix and denote it simply as Δ_W .

The Limitation of Feature Distribution-Agnostic SVD.

Recent model merging methods often apply SVD to the task matrix $\Delta_W \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$, keeping only the top- k singular components to retain essential knowledge and reduce task interference, yielding the truncated approximation $\widehat{\Delta}_W$. However, this parameter-centric view ignores the input feature distribution. While it minimizes the Frobenius norm reconstruction error of Δ_W , it provides no guarantee about preserving the functional output. For an input $x \sim \mathcal{D}$, the expected output error after discarding the smallest $r - k$ singular components is (see Appendix A.1 for derivation):

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] = \sum_{i=k+1}^r \sigma_i^2 \cdot \mathbb{E}_{x \sim \mathcal{D}} [(v_i^\top x)^2], \quad (4)$$

where r denotes the number of non-zero singular values. As shown, the error depends not only on the discarded singular values σ_i , but also on the alignment between the input distribution and the right singular vectors v_i . A direction with small σ_i may be functionally critical if inputs project strongly onto v_i . By ignoring the input distribution, SVD risks discarding functionally essential information.

Essential Subspace Decomposition (ESD). To address this limitation, we introduce Essential Subspace Decomposition (ESD), an input distribution-aware method that directly connects task matrix decomposition to its functional influence. Instead of relying on the parameter energy of the task matrix, ESD constructs a basis from the principal directions of activation shifts induced by the task matrix. For each task t , we sample a lightweight proxy dataset (32

unlabeled examples in practice). By performing a forward pass through the task-specific fine-tuned model and recording the intermediate activations, we obtain the activation shift. Specifically, given n input tokens of dimension d_{in} , forming $X_{\text{proxy}} \in \mathbb{R}^{n \times d_{\text{in}}}$, the shift is computed as:

$$\Delta_O = X_{\text{proxy}} \Delta_W^\top \in \mathbb{R}^{n \times d_{\text{out}}}, \quad (5)$$

which captures the functional footprint of Δ_W on a representative set of inputs. By performing PCA on Δ_O , we obtain a set of eigenvectors p_i and their corresponding eigenvalues λ_i , sorted by explained variance. These eigenvectors form an orthonormal basis $P = [p_1, p_2, \dots, p_{d_{\text{out}}}] \in \mathbb{R}^{d_{\text{out}} \times d_{\text{out}}}$ for the output space that is inherently aligned with the task matrix’s functional behavior. The original task matrix Δ_W is projected onto the space defined by P , which yields the coordinate matrix $A = P^\top \Delta_W \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$, where A contains the coordinates of Δ_W with respect to the basis P . Thus, Δ_W can be factorized as:

$$\Delta_W = PA = P(P^\top \Delta_W). \quad (6)$$

We truncate to the top- k principal components to form the essential basis $\hat{P} = [p_1, \dots, p_k] \in \mathbb{R}^{d_{\text{out}} \times k}$. The corresponding coordinate matrix is $\hat{A} = \hat{P}^\top \Delta_W \in \mathbb{R}^{k \times d_{\text{in}}}$, leading to the low-rank approximation:

$$\widehat{\Delta_W} = \hat{P} \hat{A} = \hat{P}(\hat{P}^\top \Delta_W). \quad (7)$$

We can theoretically show (see Appendix A.2 for derivation) that the expected output error under this decomposition is now decoupled from the input direction:

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta_W} x\|_2^2] = \sum_{i=k+1}^{d_{\text{out}}} \lambda_i. \quad (8)$$

Unlike SVD (Equation 4), the ESD truncation error (Equation 8) depends only on the sum of discarded eigenvalues. Removing small-eigenvalue components thus discards the least functionally relevant directions, making ESD truly “essential” by optimally preserving task expressiveness and enabling sparser, higher-fidelity model merging.

Essential Subspace Merging (ESM). Building upon the proposed ESD, we propose Essential Subspace Merging (ESM), a strategy that robustly aggregates task matrices by only considering the functionally important components. This method is analogous to TSV-M [11], but is adapted to leverage our distribution-aware low-rank factors. The ESM process follows three sequential steps:

1. **Factorization and Truncation.** For each task $t \in \{1, \dots, T\}$, we first factorize the task matrix Δ_{W_t} into its essential basis P_t and coordinate matrix A_t , such that $\Delta_{W_t} = P_t A_t$, where $P_t \in \mathbb{R}^{d_{\text{out}} \times d_{\text{out}}}$ and $A_t \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$.

Given T tasks, we allocate a rank budget $k = \lfloor d_{\text{out}}/T \rfloor$ to each one. We truncate the task-specific factors to their top- k components, resulting in the sparse factors $\hat{P}_t \in \mathbb{R}^{d_{\text{out}} \times k}$ and $\hat{A}_t \in \mathbb{R}^{k \times d_{\text{in}}}$.

2. **Concatenation.** Next, we form the merged basis and coordinate matrices by horizontally and vertically concatenating the truncated factors across all tasks, respectively:

$$P_{\text{cat}} = [\hat{P}_1 | \hat{P}_2 | \dots | \hat{P}_T] \in \mathbb{R}^{d_{\text{out}} \times (k \cdot T)}, \quad (9)$$

$$A_{\text{cat}} = \begin{bmatrix} \hat{A}_1 \\ \hat{A}_2 \\ \vdots \\ \hat{A}_T \end{bmatrix} \in \mathbb{R}^{(k \cdot T) \times d_{\text{in}}}. \quad (10)$$

The initial merged matrix is simply the product $P_{\text{cat}} A_{\text{cat}}$.

3. **Orthogonalization to Minimize Interference.** The concatenated factors P_{cat} and A_{cat} consist of components derived from different task subspaces, which may not be mutually orthogonal, leading to significant interference in the merged model. To reconstruct the merged matrix with minimized correlation, we must orthogonalize these concatenated factors. Following TSV-M [11], we first compute the SVDs for each concatenated matrix:

$$P_{\text{cat}} = U_P \Sigma_P V_P^\top, \quad A_{\text{cat}} = U_A \Sigma_A V_A^\top. \quad (11)$$

To ensure that more important parameter directions are preferentially preserved during orthogonalization, we apply eigenvalue-based weighting to both the directional vectors of P_{cat} and the coordinate vectors of A_{cat} prior to performing the above SVD. We then perform a whitening operation on the matrices by retaining only the orthogonal components U and V . This transformation effectively rotating the basis to be maximally decorrelated:

$$\tilde{P}_{\text{cat}} = U_P V_P^\top, \quad \tilde{A}_{\text{cat}} = U_A V_A^\top. \quad (12)$$

The final merged task matrix $\Delta_{W_{\text{merged}}}$ is then constructed by multiplying these orthogonalized factors:

$$\Delta_{W_{\text{merged}}} = \tilde{P}_{\text{cat}} \tilde{A}_{\text{cat}}. \quad (13)$$

3.3. Polarized Scaling for Knowledge Enhancement

Although ESD preserves task matrix functionality, intense weight competition during fusion can still overwhelm critical knowledge. We note that the norm of a weight block update reflects its directional confidence and inter-task consensus, serving as an effective scaling factor. We therefore introduce Polarized Scaling (Figure 4) to amplify signals and suppress noise for improved merging performance.

Empirical Evidence: Single-Task Merging. As shown in Figure 2a, we sequentially add task matrices from a fine-tuned model to the zero-shot base model, comparing three

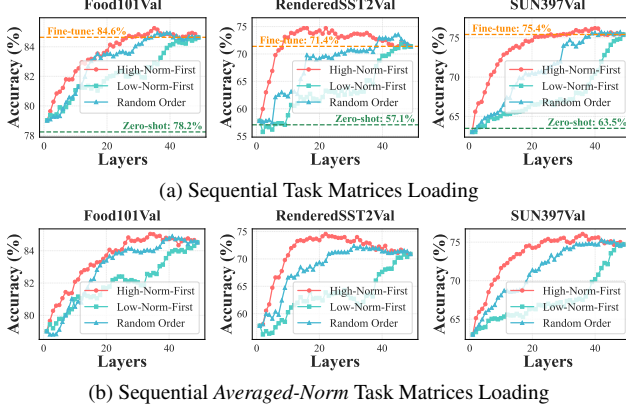


Figure 2. Performance evaluation of layer-wise task matrix loading under different ordering strategies on a pre-trained ViT-B/32 backbone. (a) Direct loading of task matrices. (b) Loading with layer-wise norms averaged to reflect directional importance.

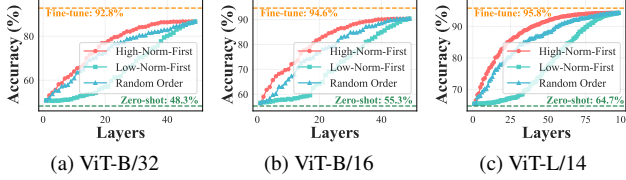


Figure 3. Performance evaluation on the 8-task benchmark when loading merged task matrices layer-by-layer into a pre-trained backbone under different ordering strategies.

orders: high-norm-first, low-norm-first, and random. The results consistently indicate that adding high-norm matrices first yields the best performance. Moreover, incorporating low-norm matrices in later stages leads to clear performance degradation. To isolate the effect of direction rather than scale, we repeat the experiment after normalizing each layer’s task matrix to the same average norm (Figure 2b). The high-norm-first order still performs best, revealing a key insight: *high-norm task matrices matter because they capture consistent, task-aligned directions in the optimization landscape, encoding the task’s essential knowledge.*

Empirical Evidence: Multi-Task Merging. As shown in Figure 3, we load the aggregated task matrices, obtained after Essential Subspace Merging (ESM), into the pre-trained model, again following the three loading orders based on the merged matrix norm. The results confirm that prioritizing high-norm components yields superior merged performance. This suggests that *the high-norm components in the merged matrix are likely the inter-task consensus, representing shared, important knowledge directions.*

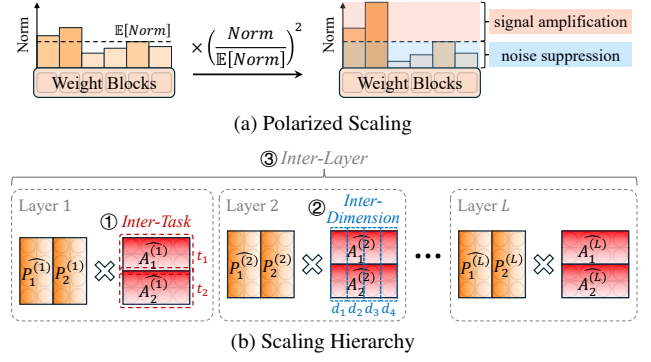


Figure 4. (a) The proposed Polarized Scaling uses the norm of parameter updates as scaling factors to amplify essential parameters submerged by redundant ones. (b) This scaling is applied at three distinct levels: across tasks, dimensions, and layers.

Proposed Polarized Scaling (PS). Based on these observations, we introduce a simple yet effective Polarized Scaling (PS) scheme designed to proactively suppress the influence of noisy, low-confidence parameters and enhance the strength of essential, high-confidence ones. The core idea is to scale a weight block by the square of its relative magnitude. We apply this scaling at three hierarchical levels:

1. **Inter-Task Scaling.** Within the same layer ℓ , we scale each individual truncated coordinate matrix $\hat{A}_t^{(\ell)}$ relative to the average norm across all tasks T . This prevents important task signals from being masked by the accumulation of noise from numerous competing tasks. The scaling coefficient $s_t^{(\ell)}$ for task t at layer ℓ is defined as:

$$s_t^{(\ell)} = \left(\frac{|\hat{A}_t^{(\ell)}|_F}{\mathbb{E}_{i \in \{1, \dots, T\}} [|\hat{A}_i^{(\ell)}|_F]} \right)^2. \quad (14)$$

The scaled coordinate matrix, $s_t^{(\ell)} \cdot \hat{A}_t^{(\ell)}$, is used in place of the original $\hat{A}_t^{(\ell)}$ to construct $A_{\text{cat}}^{(\ell)}$ in Equation 10.

2. **Inter-Dimension Scaling.** After task-wise scaling and concatenation, we further scale the columns of the coordinate matrix $A_{\text{cat}}^{(\ell)}$. This highlights the specific dimensions that exhibit strong consensus across all tasks. Let $\mathbf{a}_j$ be the j -th column vector of $A_{\text{cat}}^{(\ell)}$. The scaling coefficient $c_j^{(\ell)}$ for column j is calculated as follows:

$$c_j^{(\ell)} = \left(\frac{|\mathbf{a}_j^{(\ell)}|_2}{\mathbb{E}_{i \in \{1, \dots, d_{\text{in}}\}} [|\mathbf{a}_i^{(\ell)}|_2]} \right)^2. \quad (15)$$

After applying weighting to each column of $A_{\text{cat}}^{(\ell)}$, the representation becomes more focused on the input dimensions that are effective for every task.

Table 1. Average absolute accuracy on model merging benchmarks, with normalized average accuracy shown as subscripts in parentheses. “Zero-shot” (pre-trained model) and “Fine-tuned” (fine-tuned models) results are presented as the lower and upper bounds, respectively.

Method	ViT-B/32			ViT-B/16			ViT-L/14		
	8 tasks	14 tasks	20 tasks	8 tasks	14 tasks	20 tasks	8 tasks	14 tasks	20 tasks
<i>Zero-shot</i>	48.3	57.2	56.1	55.3	61.3	59.7	64.7	68.2	65.2
<i>Fine-tuned</i>	92.8	90.9	91.3	94.6	92.8	93.2	95.8	94.3	94.7
Weight Averaging [46]	66.3 _(72.1)	64.3 _(71.1)	61.0 _(67.5)	72.2 _(76.6)	69.5 _(74.8)	65.3 _(70.4)	79.6 _(83.2)	76.7 _(81.1)	71.6 _(75.6)
Task Arithmetic [18]	70.8 _(76.5)	65.3 _(72.1)	60.5 _(66.8)	75.4 _(79.6)	70.5 _(75.9)	65.8 _(70.8)	84.9 _(88.7)	79.4 _(84.0)	74.0 _(78.1)
TIES-Merging [52]	75.1 _(81.0)	68.0 _(74.8)	63.4 _(69.9)	79.7 _(84.3)	73.2 _(78.7)	68.2 _(73.3)	86.9 _(90.7)	79.5 _(84.1)	75.7 _(79.8)
Consensus TA [42]	75.0 _(80.8)	70.4 _(77.4)	65.4 _(72.0)	79.4 _(83.9)	74.4 _(79.9)	69.8 _(74.9)	86.3 _(90.1)	82.2 _(86.9)	79.0 _(83.2)
TSV-M [11]	85.9 _(92.3)	80.1 _(87.9)	77.1 _(84.3)	89.0 _(93.9)	84.6 _(91.0)	80.6 _(86.5)	93.0 _(97.0)	89.2 _(94.4)	87.7 _(92.5)
Iso-C [28]	86.3 _(92.9)	80.3 _(88.1)	75.5 _(82.5)	90.6 _(95.6)	84.8 _(91.1)	79.6 _(85.4)	94.2 _(98.3)	89.3 _(94.5)	87.6 _(92.2)
Iso-CTS [28]	86.2 _(92.8)	81.7 _(89.7)	78.1 _(85.5)	91.1 _(96.1)	86.4 _(92.8)	82.4 _(88.4)	94.7 _(98.8)	91.0 _(96.3)	90.1 _(94.9)
ESM (Ours)	88.4_(95.3)	83.7_(92.0)	81.3_(88.9)	91.8_(97.0)	87.4_(94.1)	84.9_(91.1)	94.8_(98.9)	91.3_(96.8)	90.4_(95.3)

3. **Inter-Layer Scaling.** Finally, we apply scaling across the different layers $\ell \in \mathcal{L}$ of the merged task matrix $\Delta_{W_{\text{merged}}}^{(\ell)}$. This is crucial because residual connections introduce inter-layer competition, where less important layers can mask vital information from important layers. We only compare layers of the same nature (e.g., all attention QKV projection layers, or all MLP up-projection layers) as comparing norms across different layer types is meaningless. Let $\mathcal{L}_{\text{type}}$ be the set of layers of the same nature. The scaling factor β_ℓ for layer $\ell \in \mathcal{L}_{\text{type}}$ is:

$$\beta_\ell = \left(\frac{\left| \Delta_{W_{\text{merged}}}^{(\ell)} \right|_F}{\mathbb{E}_{i \in \mathcal{L}_{\text{type}}} \left[\left| \Delta_{W_{\text{merged}}}^{(i)} \right|_F \right]} \right)^2. \quad (16)$$

The scaled merged task matrix is $\beta_\ell \cdot \Delta_{W_{\text{merged}}}^{(\ell)}$. This scaling ensures that layers exhibiting stronger directional consensus (higher norm) retain their relative importance.

The parameter for the ℓ -th layer of the final merged multi-task model is given by:

$$\theta_M^{(\ell)} = \theta_0^{(\ell)} + \alpha \cdot \beta_\ell \cdot \Delta_{W_{\text{merged}}}^{(\ell)}, \quad (17)$$

where α is chosen on a held-out validation set as in [28].

4. Experiments

4.1. Experimental Setup

Following [42], we evaluate multi-task merging on benchmarks of 8, 14, and 20 tasks. The 8-task set includes: Cars [21], DTD [4], EuroSAT [14], GTSRB [37], MNIST [23], RESISC45 [2], SUN397 [50], and SVHN [30]. The 14-task benchmark extends this with CIFAR100 [22], STL10 [6], Flowers102 [31], OxfordIIITPet [32], PCAM [41], and FER2013 [12]. The 20-task set further incorporates EMNIST [7], CIFAR10 [22], Food101 [1], FashionMNIST

[49], RenderedSST2 [36], and KMNIST [5]. We employ three variants of the CLIP model [33] with ViT-B/32, ViT-B/16, and ViT-L/14 visual encoders as our pre-trained base models. For the task-specific parameters, we utilize the fine-tuned checkpoints provided in the TALL-masks [42] repository, corresponding to the tasks listed above. We report performance using both absolute and normalized accuracy, adhering to the standard evaluation practices [42].

4.2. Main Results

As presented in Table 1, we compare our proposed method, ESM, against a comprehensive suite of established model merging methods: Weight Averaging [46], Task Arithmetic [18], TIES-Merging [52], Consensus TA [42], TSV-M [11], and Iso-CTS [28]. We also include the performance of the zero-shot base model and the average performance of single-task fine-tuned expert models as the lower and upper bounds, respectively. Our results demonstrate that ESM achieves state-of-the-art performance across all benchmark settings and base model variants. Notably, the performance gains are particularly pronounced in challenging scenarios with a large number of tasks or when using models with smaller capacity (e.g., ViT-B/32 vs. ViT-L/14). This strongly validates the effectiveness of our method in preserving essential task-specific knowledge while significantly mitigating inter-task interference, leading to a superior and more robust merged model.

4.3. Ablation and Analysis

Ablation of Proposed ESD and PS. We conduct an ablation study of the proposed ESD and PS modules, as shown in Table 2. The results demonstrate that decomposing and merging in the essential subspace leads to a substantial improvement over performing the same operations in the SVD subspace, which is consistent with our theoretical and ex-

Table 2. Ablation study of the proposed Essential Subspace Decomposition (ESD)-based merging and Polarized Scaling. Results are reported in terms of average absolute accuracy, with normalized average accuracy shown as subscripts in parentheses.

Method	Decomposition		Polarized Scaling			ViT-B/16			ViT-L/14		
	SVD	ESD	β_ℓ	$s_t^{(\ell)}$	$c_j^{(\ell)}$	8 tasks	14 tasks	20 tasks	8 tasks	14 tasks	20 tasks
Baseline	✓	✗	✗	✗	✗	89.0 _(93.9)	84.6 _(91.0)	80.6 _(86.5)	93.0 _(97.0)	89.2 _(94.4)	87.7 _(92.5)
	✓	✗	✓	✓	✓	89.6 _(94.6)	85.4 _(91.9)	82.1 _(88.1)	93.4 _(97.4)	89.6 _(95.0)	88.1 _(93.0)
	✗	✓	✗	✗	✗	90.9 _(95.9)	85.8 _(92.2)	82.8 _(88.7)	94.5 _(98.6)	90.7 _(95.9)	89.5 _(94.2)
	✗	✓	✓	✗	✗	91.4 _(96.4)	86.6 _(93.1)	83.7 _(89.6)	94.6 _(98.7)	90.9 _(96.2)	90.0 _(94.8)
ESM (Ours)	✗	✓	✓	✓	✓	91.8 _(97.0)	87.4 _(94.1)	84.9 _(91.1)	94.8 _(98.9)	91.3 _(96.8)	90.4 _(95.3)

Table 3. Time cost of major operations during model merging. “Forward”: forward pass of 32 proxy samples to obtain Δ_O . “PCA”: principal component analysis performed on Δ_O . “Orthogonalization”: orthogonalization of the matrices P_{cat} and A_{cat} .

Model	Forward	PCA	Orthogonalization
ViT-B/32	0.03 s/task	1.31 s/task	13.74 s (once)
ViT-B/16	0.04 s/task	1.39 s/task	13.89 s (once)
ViT-L/14	0.06 s/task	5.20 s/task	70.83 s (once)

perimental analyses. Furthermore, the proposed PS module helps reduce task interference while preserving essential task-specific knowledge, leading to additional performance gains. Notably, PS proves beneficial not only when combined with ESD-based merging, but also significantly improves the performance of the baseline method operating in the SVD subspace, highlighting its general applicability.

Computational Overhead. We report the major computational overhead introduced by our method, which includes the forward pass of a 32-sample proxy set to compute the activation shift Δ_O , the PCA on Δ_O to obtain the principal directions for the essential subspace, and the orthogonalization of the matrices P_{cat} and A_{cat} . As shown in Table 3, the experiments conducted on an RTX 4090 GPU demonstrate that the additional overhead of our method is minimal.

Comparison of ESD and SVD. Figure 5 shows the cumulative energy retained as different proportions of components are preserved, defined using the singular values (SVD) or eigenvalues (ESD) as detailed in Appendix C.7. Our proposed ESD exhibits a highly concentrated energy distribution, indicating its ability to preserve essential task-specific knowledge with fewer components. Figure 6 evaluates feature preservation when retaining only 5% of components per layer, using Centered Kernel Alignment (CKA) similarity [20]. We measure the similarity between the merged model and the fine-tuned expert using the class token from the final layer, based on the feature difference

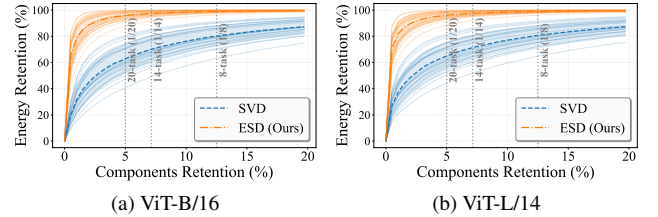


Figure 5. Energy retention as a function of the fraction of principal components retained from the task matrices. Results are reported for each of the 20 tasks as well as the average.

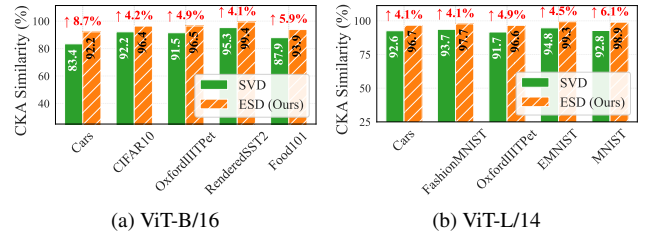


Figure 6. CKA similarity with the fine-tuned model when retaining only 5% of principal components using SVD and ESD.

relative to the zero-shot pre-trained model. ESD more effectively preserves task-specific features than SVD, further confirming its advantage in retaining critical knowledge.

Impact of Proxy Dataset Composition. We analyze how the composition of the proxy dataset affects Essential Subspace Merging (ESM). Table 4 shows the average performance and variance over five runs, with a fixed proxy set size of 32 samples. Our default setup uses unlabeled samples randomly selected from the respective task’s dataset, denoted as “Random (ID)”. We also experiment with extreme scenarios: using samples from only a single class within the task data (“Class Imbalance”) and using samples randomly drawn from an out-of-distribution dataset (ImageNet-1k [34]), denoted as “Random (OOD)”. Results indicate stable merging performance even with large distribution shifts, which we attribute to the consistent sparse pat-

Table 4. Impact of proxy dataset composition on merging performance, evaluated on the 8-task benchmark. “*Random (ID)*”: random sampling from in-distribution task data. “*Class Imbalance*”: sampling only a single class per task. “*Random (OOD)*”: random sampling from out-of-distribution ImageNet-1k [34].

Method	Sampling	ViT-B/32	ViT-B/16	ViT-L/14
SVD	-	86.6	89.6	93.4
ESD	<i>Random (ID)</i>	88.4 $\pm$ 0.06	91.8 $\pm$ 0.04	94.8 $\pm$ 0.07
ESD	<i>Class Imbalance</i>	88.4 $\pm$ 0.10	91.8 $\pm$ 0.06	94.8 $\pm$ 0.04
ESD	<i>Random (OOD)</i>	88.3 $\pm$ 0.04	91.8 $\pm$ 0.04	94.8 $\pm$ 0.01

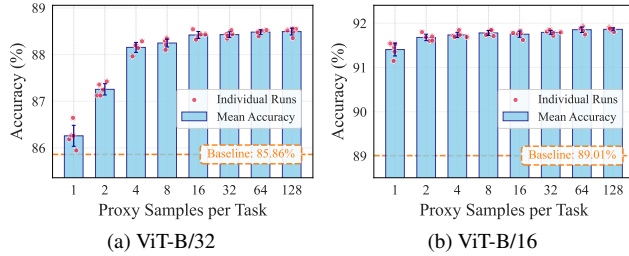


Figure 7. Ablation study on the number of proxy samples used in ESM, reporting the average accuracy on the 8-task benchmark. Each configuration was evaluated over five random runs.

terms in the features extracted by each task-specific model, regardless of the input data distribution. This finding enhances the generality of our approach, demonstrating that it provides stable performance gains over SVD-based merging methods even when task-specific data is unavailable.

Ablation of Proxy Dataset Size. We perform an ablation study on the size of the proxy dataset used in Essential Subspace Merging (ESM), as illustrated in Figure 7. The results show that using only four unlabeled samples per task is sufficient to achieve robust performance, while even a single sample can consistently outperform the baseline.

Detailed Ablation Study of Polarized Scaling. We perform a detailed ablation of Polarized Scaling (PS) in Table 5. We compare three alternatives: (i) “*Reverse*”, which applies the reciprocal of the scaling factors; (ii) “*Noise—*”, which retains only factors < 1 to suppress noisy parameters; and (iii) “*Signal++*”, which retains only factors > 1 to enhance important parameters. Experimental results show that, compared with “*None*” (i.e., without scaling), the “*Reverse*” operation significantly degrades performance because important parameters are overwhelmed by redundant ones. Both “*Noise—*” and “*Signal++*” improve over “*None*” by raising the signal-to-noise ratio of important parameters. Our full PS method, which combines both suppression and amplification, achieves the best performance.

Table 5. Detailed ablation study of the proposed Polarized Scaling. The symbol γ denotes the scaling factor at three different levels: s , c , or β . The following variants are compared: (i) “*Reverse*”: taking the reciprocal of the scaling factors; (ii) “*Noise—*”: retaining only factors < 1 to suppress noisy parameters; (iii) “*Signal++*”: retaining only factors > 1 to enhance important parameters.

Method	Scaling	ViT-B/32		
		8 tasks	14 tasks	20 tasks
<i>None</i>	-	86.7 _(93.4)	81.1 _(89.0)	78.1 _(85.5)
<i>Reverse</i>	$1/\gamma$	82.9 _(89.2)	76.3 _(83.7)	72.6 _(79.4)
<i>Noise—</i>	$\min(\gamma, 1)$	87.4 _(94.1)	83.0 _(91.2)	80.5 _(88.1)
<i>Signal++</i>	$\max(\gamma, 1)$	87.7 _(94.6)	82.8 _(91.0)	80.2 _(87.8)
PS (Ours)	γ	88.4_(95.3)	83.7_(92.0)	81.3_(88.9)

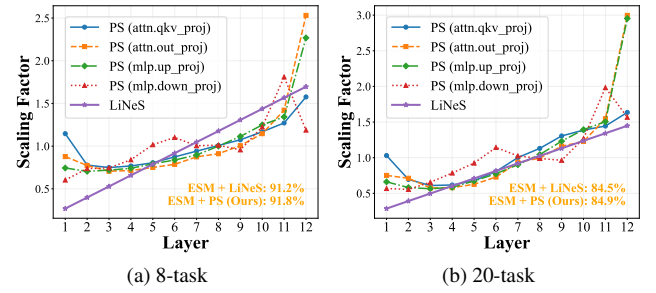


Figure 8. Layer-wise scaling coefficients for merged ViT-B/16.

Analysis of Scaling Coefficients. Figure 8 compares the layer-wise coefficients of our Polarized Scaling (PS) with LiNeS [43]. Based on the functional roles of shallow and deep blocks, LiNeS applies coefficients that increase linearly with depth, requiring a validation set to determine the linear multiplier. In contrast, our PS method provides a more refined measure of parameter importance by leveraging layer norms. It performs separate scaling for four key layer types in a ViT model: the QKV and Output projections in the attention module, and the Up and Down projections in the MLP. This approach not only achieves superior merging performance but also removes the need for a validation set by computing coefficients directly from parameter norms.

5. Conclusion

In this paper, we propose Essential Subspace Merging (ESM), a novel model merging framework that narrows the gap to expert models. It introduces the ESD, a feature shift distribution-aware decomposition method that identifies a sparser and more functionally critical subspace than SVD, thereby reducing inter-task interference. Observing that high-norm updates indicate consensus-driven directions, we further propose Polarized Scaling (PS), a three-level mechanism that amplifies high-confidence signals while suppress-

ing noise. ESM achieves state-of-the-art performance and provides a robust and effective solution for model merging.

Acknowledgment

This research was supported by the Jiangsu Science Foundation (BG2024036, BK20243012), the National Science Foundation of China (62125602, U24A20324, 92464301), the New Cornerstone Science Foundation through the XPLOER PRIZE, the Fundamental Research Funds for the Central Universities (2242025K30024), and the Big Data Computing Center of Southeast University.

References

- [1] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101—mining discriminative components with random forests. In *ECCV*, 2014. 6
- [2] Gong Cheng, Junwei Han, and Xiaoqiang Lu. Remote sensing image scene classification: Benchmark and state of the art. *Proceedings of the IEEE*, 2017. 6
- [3] Runxi Cheng, Feng Xiong, Yongxian Wei, Wanyun Zhu, and Chun Yuan. Whoever started the interference should end it: Guiding data-free model merging via task vectors. In *ICML*, 2025. 2
- [4] Mircea Cimpoi, Subhransu Maji, Iasonas Kokkinos, Sammy Mohamed, and Andrea Vedaldi. Describing textures in the wild. In *CVPR*, 2014. 6
- [5] Tarin Clanuwat, Mikel Bober-Irizar, Asanobu Kitamoto, Alex Lamb, Kazuaki Yamamoto, and David Ha. Deep learning for classical japanese literature. *arXiv preprint arXiv:1812.01718*, 2018. 6
- [6] Adam Coates, Andrew Ng, and Honglak Lee. An analysis of single-layer networks in unsupervised feature learning. In *AISTATS*, 2011. 6
- [7] Gregory Cohen, Saeed Afshar, Jonathan Tapson, and Andre Van Schaik. Emnist: Extending mnist to handwritten letters. In *IJCNN*, 2017. 6
- [8] Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. Qlora: Efficient finetuning of quantized llms. In *NeurIPS*, 2023. 2
- [9] Ning Ding, Yujia Qin, Guang Yang, Fuchao Wei, Zonghan Yang, Yusheng Su, Shengding Hu, Yulin Chen, Chi-Min Chan, Weize Chen, et al. Parameter-efficient fine-tuning of large-scale pre-trained language models. *NMI*, 2023. 2
- [10] Guodong Du, Junlin Lee, Jing Li, Runhua Jiang, Yifei Guo, Shuyang Yu, Hanting Liu, Sim K Goh, Ho-Kin Tang, Daojing He, et al. Parameter competition balancing for model merging. In *NeurIPS*, 2024. 2
- [11] Antonio Andrea Gargiulo, Donato Crisostomi, Maria Sofia Bucarelli, Simone Scardapane, Fabrizio Silvestri, and Emanuele Rodola. Task singular vectors: Reducing task interference in model merging. In *CVPR*, 2025. 1, 2, 3, 4, 6
- [12] Ian J Goodfellow, Dumitru Erhan, Pierre Luc Carrier, Aaron Courville, Mehdi Mirza, Ben Hamner, Will Cukierski, Yichuan Tang, David Thaler, Dong-Hyun Lee, et al. Challenges in representation learning: A report on three machine learning contests. In *ICONIP*, 2013. 6
- [13] Ligong Han, Yinxiao Li, Han Zhang, Peyman Milanfar, Dimitris Metaxas, and Feng Yang. Svdif: Compact parameter space for diffusion fine-tuning. In *ICCV*, 2023. 2
- [14] Patrick Helber, Benjamin Bischke, Andreas Dengel, and Damian Borth. Eurosat: A novel dataset and deep learning benchmark for land use and land cover classification. *JS-TARS*, 2019. 6
- [15] Edward J Hu, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. Lora: Low-rank adaptation of large language models. In *ICLR*, 2022. 2
- [16] Chenyu Huang, Peng Ye, Tao Chen, Tong He, Xiangyu Yue, and Wanli Ouyang. Emr-merging: Tuning-free high-performance model merging. In *NeurIPS*, 2024. 2
- [17] Xin Huang, Longhua Li, Lei Qi, and Xin Geng. Few-shot network thinning: A unified dense compression framework for cnns and vits. *TMC*, 2026. 2
- [18] Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Ludwig Schmidt, Hannaneh Hajishirzi, and Ali Farhadi. Editing models with task arithmetic. In *ICLR*, 2023. 1, 2, 3, 6
- [19] Xisen Jin, Xiang Ren, Daniel Preotiuc-Pietro, and Pengxiang Cheng. Dataless knowledge fusion by merging weights of language models. In *ICLR*, 2023. 1, 2
- [20] Simon Kornblith, Mohammad Norouzi, Honglak Lee, and Geoffrey Hinton. Similarity of neural network representations revisited. In *ICML*, 2019. 7
- [21] Jonathan Krause, Michael Stark, Jia Deng, and Li Fei-Fei. 3d object representations for fine-grained categorization. In *ICCV workshops*, 2013. 6
- [22] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. Technical report, Toronto, ON, Canada, 2009. 6
- [23] Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 2002. 6
- [24] Longhua Li, Lei Qi, and Xin Geng. One-shot knowledge transfer for scalable person re-identification. In *ICCV*, 2025. 2
- [25] Longhua Li, Lei Qi, and Xin Geng. Stratified knowledge-density super-network for scalable vision transformers. *arXiv preprint arXiv:2511.11683*, 2025. 2
- [26] Yixiao Li, Yifan Yu, Qingru Zhang, Chen Liang, Pengcheng He, Weizhu Chen, and Tuo Zhao. Lospars: Structured compression of large language models based on low-rank and sparse approximation. In *ICML*, 2023. 2
- [27] Xin Lin, Yajiao Wang, Zhixiong Zhang, and Mengting Zhang. Scientific claim recognition via staged fine-tuning with lora. *DATA INTELLIGENCE*, 2025. 2
- [28] Daniel Marczak, Simone Magistri, Sebastian Cygert, Bartłomiej Twardowski, Andrew D Bagdanov, and Joost van de Weijer. No task left behind: Isotropic model merging with common and task-specific subspaces. In *ICML*, 2025. 1, 2, 3, 6

- [29] Michael S Matena and Colin A Raffel. Merging models with fisher-weighted averaging. In *NeurIPS*, 2022. 1, 2
- [30] Yuval Netzer, Tao Wang, Adam Coates, Alessandro Bis-sacco, Baolin Wu, Andrew Y Ng, et al. Reading digits in nat-ural images with unsupervised feature learning. In *NeurIPS workshops*, 2011. 6
- [31] Maria-Elena Nilsback and Andrew Zisserman. Automated flower classification over a large number of classes. In *ICVGIP*, 2008. 6
- [32] Omkar M Parkhi, Andrea Vedaldi, Andrew Zisserman, and CV Jawahar. Cats and dogs. In *CVPR*, 2012. 6
- [33] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Aspell, Pamela Mishkin, Jack Clark, et al. Learn-ing transferable visual models from natural language super-vision. In *ICML*, 2021. 6
- [34] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, San-jeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, et al. Imagenet large scale visual recognition challenge. *IJCV*, 2015. 7, 8
- [35] Rajarshi Saha, Varun Srivastava, and Mert Pilanci. Matrix compression via randomized low rank and low precision fac-torization. In *NeurIPS*, 2023. 2
- [36] Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D Manning, Andrew Y Ng, and Christopher Potts. Recursive deep models for semantic compositional-ity over a sentiment treebank. In *EMNLP*, 2013. 6
- [37] Johannes Stallkamp, Marc Schlipsing, Jan Salmen, and Christian Igel. The german traffic sign recognition bench-mark: a multi-class classification competition. In *IJCNN*, 2011. 6
- [38] George Stoica, Pratik Ramesh, Boglarka Ecsedi, Leshem Choshen, and Judy Hoffman. Model merging with svd to tie the knots. In *ICLR*, 2025. 1, 2, 3
- [39] Wenju Sun, Qingyong Li, Wen Wang, Yang Liu, Yangliao Geng, and Boyang Li. Towards minimizing feature drift in model merging: Layer-wise task vector fusion for adaptive knowledge integration. In *NeurIPS*, 2025. 1
- [40] Yanpeng Sun, Qiang Chen, Xiangyu He, Jian Wang, Haocheng Feng, Junyu Han, Errui Ding, Jian Cheng, Zechao Li, and Jingdong Wang. Singular value fine-tuning: Few-shot segmentation requires few-parameters fine-tuning. In *NeurIPS*, 2022. 2
- [41] Bastiaan S Veeling, Jasper Linmans, Jim Winkens, Taco Co-hen, and Max Welling. Rotation equivariant cnns for digital pathology. In *MICCAI*, 2018. 6
- [42] Ke Wang, Nikolaos Dimitriadis, Guillermo Ortiz-Jimenez, François Fleuret, and Pascal Frossard. Localizing task in-formation for improved model merging and compression. In *ICML*, 2024. 6
- [43] Ke Wang, Nikolaos Dimitriadis, Alessandro Favero, Guillermo Ortiz-Jimenez, François Fleuret, and Pascal Frossard. Lines: Post-training layer scaling prevents forget-ting and enhances model merging. In *ICLR*, 2025. 2, 8
- [44] Xin Wang, Samiul Alam, Zhongwei Wan, Hui Shen, and Mi Zhang. Svd-llm v2: Optimizing singular value truncation for large language model compression. In *NAACL*, 2025. 2
- [45] Xin Wang, Yu Zheng, Zhongwei Wan, and Mi Zhang. Svd-llm: Truncation-aware singular value decomposition for large language model compression. In *ICLR*, 2025. 2
- [46] Mitchell Wortsman, Gabriel Ilharco, Samir Ya Gadre, Re-becca Roelofs, Raphael Gontijo-Lopes, Ari S Morcos, Hongseok Namkoong, Ali Farhadi, Yair Carmon, Simon Kor-nblith, et al. Model soups: averaging weights of multiple fine-tuned models improves accuracy without increasing in-ference time. In *ICML*, 2022. 1, 2, 6
- [47] Shiyu Xia, Miaosen Zhang, Xu Yang, Ruiming Chen, Haokun Chen, and Xin Geng. Transformer as linear expan-sion of learngene. In *AAAI*, 2024. 2
- [48] Shi-Yu Xia, Wenxuan Zhu, Xu Yang, and Xin Geng. Explor-ing learngene via stage-wise weight sharing for initializing variable-sized models. In *IJCAI*, 2024. 2
- [49] Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017. 6
- [50] Jianxiong Xiao, Krista A Ehinger, James Hays, Antonio Tor-ralba, and Aude Oliva. Sun database: Exploring a large col-lection of scene categories. *IJCV*, 2016. 6
- [51] Jiase Xu, Shiyu Xia, Xu Yang, Jiaqi Lv, and Xin Geng. Learngene tells you how to customize: Task-aware param-eter initialization at flexible scales. In *ICML*, 2025. 2
- [52] Prateek Yadav, Derek Tam, Leshem Choshen, Colin A Raf-fel, and Mohit Bansal. Ties-merging: Resolving interference when merging models. In *NeurIPS*, 2023. 1, 2, 6
- [53] Kunda Yan, Min Zhang, Sen Cui, Qu Zikun, Bo Jiang, Feng Liu, and Changshui Zhang. Calm: Consensus-aware local-ized merging for multi-task learning. In *ICML*, 2025. 1
- [54] Enneng Yang, Zhenyi Wang, Li Shen, Shiwei Liu, Guibing Guo, Xingwei Wang, and Dacheng Tao. Adamerging: Adap-tive model merging for multi-task learning. In *ICLR*, 2024. 2
- [55] Le Yu, Bowen Yu, Haiyang Yu, Fei Huang, and Yongbin Li. Language models are super mario: Absorbing abilities from homologous models as a free lunch. In *ICML*, 2024. 2
- [56] Frederic Z Zhang, Paul Albert, Cristian Rodriguez-Opazo, Anton van den Hengel, and Ehsan Abbasnejad. Knowledge composition using task vectors with learned anisotropic scal-ing. In *NeurIPS*, 2024. 2
- [57] Han-Chen Zhang, Zi-Hao Zhou, Mao-Lin Luo, Shimin Di, Min-Ling Zhang, and Tong Wei. Dc-merge: Improving model merging with directional consistency. *arXiv preprint arXiv:2603.06242*, 2026. 1
- [58] Shenghe Zheng and Hongzhi Wang. Free-merging: Fourier transform for efficient model merging. In *ICCV*, 2025. 2

Model Merging in the Essential Subspace

Supplementary Material

A. Proofs

This section provides the proofs for the expected output error after truncation for both the standard SVD and the proposed Essential Subspace Decomposition (ESD), as well as the equivalence of whitening and Orthogonal Procrustes.

A.1. Proof for SVD Truncation Error

Theorem 1. Given a task matrix $\Delta_W \in \mathbb{R}^{d_{out} \times d_{in}}$ with its singular value decomposition $\Delta_W = U\Sigma V^\top = \sum_{i=1}^r \sigma_i u_i v_i^\top$. Let $\widehat{\Delta}_W = \sum_{i=1}^k \sigma_i u_i v_i^\top$ be its top- k rank approximation. For an input x drawn from a distribution $\mathcal{D}$, the expected squared L_2 error on the output activation is:

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] = \sum_{i=k+1}^r \sigma_i^2 \cdot \mathbb{E}_{x \sim \mathcal{D}} [(v_i^\top x)^2].$$

Proof. The error matrix resulting from the truncation is the sum of the discarded components:

$$\Delta_W - \widehat{\Delta}_W = \sum_{i=k+1}^r \sigma_i u_i v_i^\top. \quad (18)$$

The error on the output activation for a given input x is:

$$(\Delta_W - \widehat{\Delta}_W)x = \left(\sum_{i=k+1}^r \sigma_i u_i v_i^\top \right) x = \sum_{i=k+1}^r \sigma_i u_i (v_i^\top x). \quad (19)$$

Since $v_i^\top x$ is a scalar, we can rewrite this as a linear combination of the orthonormal vectors u_i . The squared L_2 norm of this error vector is:

$$\|(\Delta_W - \widehat{\Delta}_W)x\|_2^2 = \left\| \sum_{i=k+1}^r (\sigma_i v_i^\top x) u_i \right\|_2^2. \quad (20)$$

Because the left singular vectors $\{u_i\}$ form an orthonormal set, the squared norm of their weighted sum is the sum of the squares of the weights:

$$\|(\Delta_W - \widehat{\Delta}_W)x\|_2^2 = \sum_{i=k+1}^r (\sigma_i v_i^\top x)^2 = \sum_{i=k+1}^r \sigma_i^2 (v_i^\top x)^2. \quad (21)$$

By taking the expectation over the input distribution $\mathcal{D}$ and applying the linearity of expectation, we arrive at the final expression:

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] = \sum_{i=k+1}^r \sigma_i^2 \cdot \mathbb{E}_{x \sim \mathcal{D}} [(v_i^\top x)^2]. \quad (22)$$

This completes the proof. $\square$

A.2. Proof for ESD Truncation Error

Theorem 2. Given a task matrix $\Delta_W \in \mathbb{R}^{d_{out} \times d_{in}}$, let $\hat{P} \in \mathbb{R}^{d_{out} \times k}$ be the matrix whose columns are the top- k principal components (eigenvectors) derived from the activation shift matrix $\Delta_O = X_{\text{proxy}} \Delta_W^\top$. Let $\widehat{\Delta}_W = \hat{P} \hat{A} = \hat{P}(\hat{P}^\top \Delta_W)$ be the ESD reconstruction. For an input $x \sim \mathcal{D}$, the expected squared L_2 error on the output is proportional to the sum of the discarded eigenvalues:

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] = \sum_{i=k+1}^{d_{out}} \lambda_i.$$

Proof. The error on the output activation for an input x is:

$$\Delta_W x - \widehat{\Delta}_W x = \Delta_W x - \hat{P} \hat{P}^\top \Delta_W x = (I - \hat{P} \hat{P}^\top) \Delta_W x. \quad (23)$$

The matrix $(I - \hat{P} \hat{P}^\top)$ is the projection matrix onto the subspace spanned by the discarded eigenvectors $\{p_{k+1}, \dots, p_{d_{out}}\}$. Let $y = \Delta_W x$ be the activation shift for input x . The error vector can be expressed as the projection of y onto this orthogonal subspace:

$$(I - \hat{P} \hat{P}^\top) y = \sum_{i=k+1}^{d_{out}} (p_i^\top y) p_i. \quad (24)$$

Since $\{p_i\}$ form an orthonormal basis, the squared L_2 norm is the sum of the squares of the projection coefficients:

$$\begin{aligned} \|\Delta_W x - \widehat{\Delta}_W x\|_2^2 &= \left\| \sum_{i=k+1}^{d_{out}} (p_i^\top y) p_i \right\|_2^2 \\ &= \sum_{i=k+1}^{d_{out}} (p_i^\top y)^2 \\ &= \sum_{i=k+1}^{d_{out}} (p_i^\top \Delta_W x)^2. \end{aligned} \quad (25)$$

Now, we take the expectation over the input distribution $\mathcal{D}$:

$$\begin{aligned} \mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] &= \mathbb{E}_{x \sim \mathcal{D}} \left[\sum_{i=k+1}^{d_{out}} (p_i^\top \Delta_W x)^2 \right] \\ &= \sum_{i=k+1}^{d_{out}} \mathbb{E}_{x \sim \mathcal{D}} [(p_i^\top \Delta_W x)^2]. \end{aligned} \quad (26)$$

By the definition of Principal Component Analysis (PCA), the eigenvalue λ_i of the covariance matrix of activation

shifts corresponds to the variance of the activation shifts projected onto the i -th principal component p_i . Assuming the activation shifts are centered (or mean-subtracted during PCA), this variance is:

$$\begin{aligned}\lambda_i &= \text{Var}(p_i^\top \Delta_W x) \\ &= \mathbb{E}_{x \sim \mathcal{D}} [(p_i^\top \Delta_W x)^2] - (\mathbb{E}_{x \sim \mathcal{D}} [p_i^\top \Delta_W x])^2 \\ &= \mathbb{E}_{x \sim \mathcal{D}} [(p_i^\top \Delta_W x)^2].\end{aligned}\quad (27)$$

Substituting this result back into our error expression, we get:

$$\mathbb{E}_{x \sim \mathcal{D}} [\|\Delta_W x - \widehat{\Delta}_W x\|_2^2] = \sum_{i=k+1}^{d_{\text{out}}} \lambda_i. \quad (28)$$

This completes the proof, showing that the expected error is solely dependent on the magnitude of the discarded eigenvalues, which represent the functional variance captured by those directions. $\square$

A.3. Equivalence of Whitening and Orthogonal Procrustes

Theorem 3. The transformations $X \mapsto X(X^\top X)^{-1/2}$ (whitening) and $X \mapsto UV^\top$ (Orthogonal Procrustes), where $X = U\Sigma V^\top$ is the singular value decomposition (SVD) of X , are equivalent.

Proof. Let $X = U\Sigma V^\top$ be the SVD of X , where Σ is a diagonal matrix of singular values, we have:

$$X^\top X = V\Sigma^2 V^\top. \quad (29)$$

It follows that:

$$(X^\top X)^{-1/2} = V\Sigma^{-1} V^\top. \quad (30)$$

Substituting this into the whitening transformation gives:

$$\begin{aligned}X(X^\top X)^{-1/2} &= (U\Sigma V^\top)(V\Sigma^{-1} V^\top) \\ &= U\Sigma\Sigma^{-1} V^\top \\ &= UV^\top.\end{aligned}\quad (31)$$

The whitening operation is equivalent to solving the Orthogonal Procrustes problem, completing the proof. $\square$

B. Method Details

B.1. Methodology Pseudocode

The pseudocode of our proposed model merging approach is presented in Algorithm 1. Steps highlighted in *orange* correspond to the essential subspace merging, while *blue* annotations indicate the polarized scaling. Our method decomposes each task matrix within its essential subspace and performs truncation, then concatenates the components from all task matrices to reconstruct a new merged matrix. Through three-level scaling, we ultimately obtain a fused model that effectively preserves essential task knowledge while maintaining minimal inter-task interference.

Algorithm 1 Essential Subspace Merging

Require: Task matrices $\{\Delta_{W_t}^{(\ell)}\}_{t=1}^T$ for all layers $\ell \in \mathcal{L}$, pre-trained weights θ_0 , validation set $\mathcal{D}_{\text{val}}$

Ensure: Merged model parameters θ_M

```

1: Decomposition and Truncation
   for each task  $t = 1$  to  $T$  and each layer  $\ell$  do
2:   Obtain essential basis  $P_t^{(\ell)}$ 
3:   Compute coordinate matrix:  $A_t^{(\ell)} \leftarrow (P_t^{(\ell)})^\top \Delta_{W_t}^{(\ell)}$ 
4:   Truncate to top- $k$  components:  $\hat{P}_t^{(\ell)} \leftarrow P_t^{(\ell)}[:, 1 : k]$ ,  $\hat{A}_t^{(\ell)} \leftarrow A_t^{(\ell)}[1 : k, :]$ , where  $k = \lfloor d_{\text{out}}/T \rfloor$ 
5: end for
6: Concatenation
   for each layer  $\ell$  do
7:    $P_{\text{cat}}^{(\ell)} \leftarrow [\hat{P}_1^{(\ell)}, \hat{P}_2^{(\ell)}, \dots, \hat{P}_T^{(\ell)}]$ 
8:    $A_{\text{cat}}^{(\ell)} \leftarrow [\hat{A}_1^{(\ell)}; \hat{A}_2^{(\ell)}; \dots; \hat{A}_T^{(\ell)}]$ 
9:   Inter-Task Scaling
   for each task  $t = 1$  to  $T$  do
10:     $\hat{A}_t^{(\ell)} \leftarrow s_t^{(\ell)} \cdot \hat{A}_t^{(\ell)}$ 
11:   end for
12:   Inter-Dimension Scaling
   for each column  $j = 1$  to  $d_{\text{in}}$  do
13:     $a_j^{(\ell)} \leftarrow c_j^{(\ell)} \cdot a_j^{(\ell)}$ 
14:   end for
15: end for
16: Orthogonalization and Reconstruction
   for each layer  $\ell$  do
17:   Compute SVD:  $P_{\text{cat}}^{(\ell)} = U_P^{(\ell)} \Sigma_P^{(\ell)} (V_P^{(\ell)})^\top$ 
18:   Compute SVD:  $A_{\text{cat}}^{(\ell)} = U_A^{(\ell)} \Sigma_A^{(\ell)} (V_A^{(\ell)})^\top$ 
19:   Orthogonalize:  $\tilde{P}^{(\ell)} \leftarrow U_P^{(\ell)} (V_P^{(\ell)})^\top$ ,  $\tilde{A}^{(\ell)} \leftarrow U_A^{(\ell)} (V_A^{(\ell)})^\top$ 
20:   Reconstruct:  $\Delta_{\text{merged}}^{(\ell)} \leftarrow \tilde{P}^{(\ell)} \tilde{A}^{(\ell)}$ 
21: end for
22: Inter-Layer Scaling
   for each layer  $\ell \in \mathcal{L}$  do
23:     $\Delta_{\text{merged}}^{(\ell)} \leftarrow \beta_\ell \cdot \Delta_{\text{merged}}^{(\ell)}$ 
24:   end for
25: Find optimal  $\alpha^*$  on  $\mathcal{D}_{\text{val}}$  that maximizes performance
26:  $\theta_M \leftarrow \theta_0 + \alpha^* \cdot \{\Delta_{\text{merged}}^{(\ell)}\}_{\ell=1}^L$ 
27: return  $\theta_M$ 

```

B.2. Details on Polarized Scaling

In the empirical evidence presented in Section 3.3, individual or merged task matrices are loaded in different sequences. To isolate the phenomenon from the distinct characteristics of different layer types, our analysis focuses solely on the four principal layer types within the ViT: namely, the QKV projection and the Output projection in the attention module, along with the Up- and Down-projections in the MLP. Furthermore, when merging layers

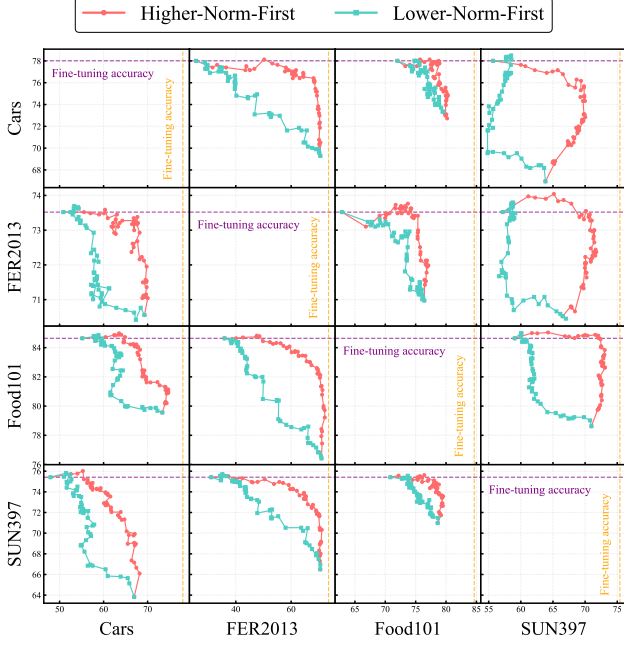


Figure 9. Illustration of task invasion: using the task matrices from one task to invade the fine-tuned model of another, performed layer-by-layer based on the norm order. Rows represent the invaded task, and columns represent the invading task.

sequentially, we group the layers by these four layer types. The merging process then proceeds by cyclically selecting task matrices from each group in a fixed order. This design ensures that the observed effects are attributable to the merging strategy rather than to variations across layer types.

Empirical Evidence: Pairwise Task Interaction. We further analyze the pairwise influence between task matrices. As shown in Figure 9, each column represents how the performance of two tasks changes when the task matrix of the column’s task is added in different orders to the fine-tuned model of the row’s task. The results show that adding task matrices in descending order of their norms yields a better Pareto optimal trade-off than adding them in ascending order. Introducing a large number of low-norm updates brings little improvement to the target task but significantly harms the performance of others. This highlights the importance of suppressing less critical or noisy updates while emphasizing the most essential ones in model merging.

B.3. Merging Non-Matrix Parameters

While most parameters in the ViT are 2D matrices merged using our proposed ESM within the Essential Subspace, the network also includes other parameter types. For non-matrix parameters such as bias vectors, layer normalization parameters, and the convolutional stem, we follow the stan-

Table 6. Performance on the 8-task benchmark. “ID”: the average accuracy of the merged model across all tasks. “OOD”: the result of merging fine-tuned models from 7 tasks and testing on the remaining unseen task. Each task is alternately used as the test task, and the reported OOD score is the average across all such settings.

Method	ViT-B/32		ViT-B/16		ViT-L/14	
	ID	OOD	ID	OOD	ID	OOD
Baseline	85.9	49.9	89.0	54.7	93.0	65.9
ESM (Ours)	88.4	51.3	91.8	55.5	94.8	66.4

Table 7. Scaling coefficient α selected using the validation set. T represents the total number of merged task-specific expert models.

ViT-B/32			ViT-B/16			ViT-L/14		
$T=8$	$T=14$	$T=20$	$T=8$	$T=14$	$T=20$	$T=8$	$T=14$	$T=20$
0.88	0.76	0.69	0.82	0.76	0.67	0.91	0.72	0.64

dard practice in [11] and apply simple averaging.

C. Experiment Details

C.1. OOD Performance

We evaluate the out-of-domain (OOD) generalization of our model merging method and that of the baseline [11], as presented in Table 6. Based on the 8-task benchmark, we treat each task in turn as the test task, merge the fine-tuned models from the remaining seven tasks, and evaluate the merged model on the held-out task. The final OOD score is the average performance across all test tasks. The results show that our method not only achieves substantial improvements in conventional in-domain evaluations but also demonstrates superior generalization in the out-of-domain setting.

C.2. Selection of the Global Scaling Coefficient α

We report the global scaling coefficient α selected on the validation set, as shown in Table 7. Based on the empirical ranges used in previous model merging studies [11, 28], we set the search interval for α between 0.0 and 2.0 and perform binary search to determine the optimal value. The results show that the optimal α decreases as the number of tasks increases, likely because merging more tasks amplifies the norm of the combined updates. Consequently, a smaller scaling factor helps balance the output feature norm and achieves better validation performance.

C.3. Performance on Individual Tasks

We present detailed model fusion results across multiple models and datasets, as shown in Figure 10. Our proposed ESM framework demonstrates consistent performance improvements across nearly all datasets and models.

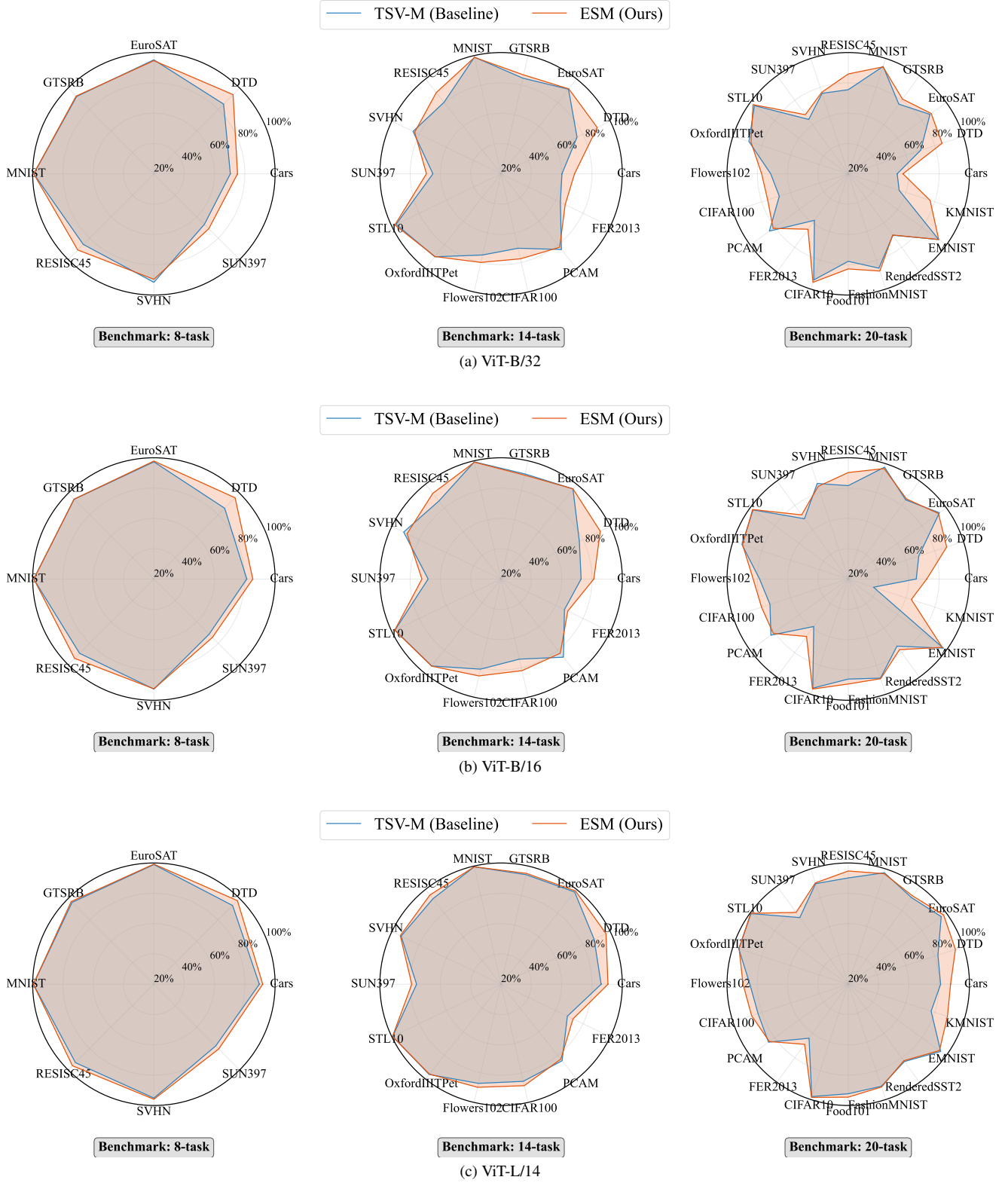
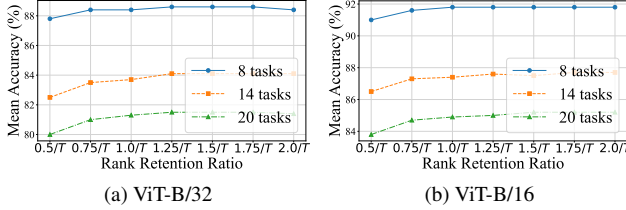


Figure 10. Performance comparison between the baseline method TSV-M [11] and our proposed ESM framework.

Table 8. Ablation study on the order of the three-level scaling in Polarized Scaling.

Scaling Order	ViT-B/32			ViT-B/16		
	8 tasks	14 tasks	20 tasks	8 tasks	14 tasks	20 tasks
(1) Inter-Dimension, (2) Inter-Task, (3) Inter-layer	88.0 _(94.1)	83.2 _(91.4)	80.8 _(88.3)	91.6 _(96.7)	87.1 _(93.7)	84.7 _(90.9)
(1) Inter-Task, (2) Inter-Dimension, (3) Inter-layer	88.4 _(95.3)	83.7 _(92.0)	81.3 _(88.9)	91.8 _(97.0)	87.4 _(94.1)	84.9 _(91.1)

Figure 11. Ablation study on the impact of rank retention ratio on merged model performance. T denotes the number of tasks.

C.4. Order of Scaling Across Tasks, Dimensions, and Layers

As shown in Algorithm 1, the default order of Polarized Scaling in our experiments is: first across tasks, then across dimensions, and finally across layers. Since inter-layer scaling is designed to emphasize task consensus, it is applied by default to the final fused task matrices. The order of inter-task and inter-dimension scaling, however, can be interchanged. Therefore, we ablate the sequence of these two operations in Table 8. The results indicate that our default order, which applies inter-task scaling first followed by inter-dimension scaling, yields better performance.

C.5. Ablation Study on Rank k Selection

In our method and experiments, the default setting uses $k = \lfloor d_{\text{out}}/T \rfloor$ as the rank budget for low-rank decomposition of each task matrix, where T denotes the number of tasks and d_{out} represents the original output dimension. We conduct an ablation study on the selection of rank k , as shown in Figure 11. The results demonstrate that the merged model exhibits robustness to the choice of rank k , maintaining comparable performance across a wide range of values ($\lfloor 0.5 \cdot d_{\text{out}}/T \rfloor \sim \lfloor 2.0 \cdot d_{\text{out}}/T \rfloor$).

C.6. Ablation Study on the Exponent of the Scaling Factor

The default configuration of our method employs a power of 2 in the polarized scaling coefficient, i.e., $(\frac{\text{norm}}{\mathbb{E}[\text{norm}]})^2$. The rationale for this choice is to amplify significant parameters while suppressing redundant ones. To validate the sensitivity of our approach to this hyperparameter, we conducted an ablation study. As shown in Figure 12, the results indicate that model merging is robust across a range of exponents.

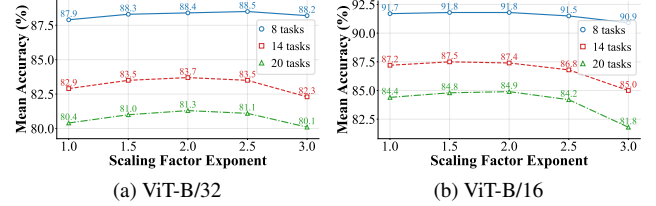


Figure 12. Impact of the scaling factor exponent on merged model performance.

The value of 2 was chosen as the default because it achieves optimal performance.

C.7. Calculation of Energy Retention

For the SVD-based method, energy retention is calculated as the ratio of the sum of squares of the retained singular values to the sum of squares of all singular values. For our ESD method, it is defined as the ratio of the sum of the retained eigenvalues to the sum of all eigenvalues. This is because the square of a singular value and an eigenvalue both correspond to the explained variance.