

CP MERGING: JOINT LoRA MERGING USING CANONICAL POLYADIC DECOMPOSITION

Anonymous authors

Paper under double-blind review

ABSTRACT

Large language models (LLMs) are often fine-tuned for specific tasks using Low-Rank Adaptation (LoRA), an efficient method that adds small, task-specific modules called LoRA adapters to a pre-trained base model. However, a major challenge arises when merging multiple LoRA adapters trained on different data sources for a specific task: it often leads to task interference, which degrades the model’s performance. While recent SVD-based LoRA merging methods have shown promise by decomposing adapters into orthogonal components and keeping only the most important ones, they have an important limitation: These methods decompose each adapter independently, overlooking potential interactions between different tasks. To address this, we propose a novel LoRA merging method using joint Canonical Polyadic (CP) decomposition (CP merging). We first combine the LoRA adapters into a single third-order tensor. Then, we apply CP decomposition to this tensor to disentangle factors that are unique to each task from those that are shared across tasks. This joint factorization method helps reduce cross-task interference without losing important information. Our extensive experiments on NLP tasks demonstrate that CP merging yields superior performance compared to the existing SVD-based baselines.

1 INTRODUCTION

Parameter-efficient fine-tuning (PEFT) methods have emerged as practical alternatives to full fine-tuning for adapting large language models (LLMs) to downstream tasks (Hu et al., 2021; Houlsby et al., 2019b; Lester et al., 2021). Among these, LoRA adapters (Hu et al., 2021) have gained particular popularity by learning low-rank additive updates that efficiently capture task-specific knowledge. Recent studies (Sheng et al., 2023; Ostapenko et al., 2024; Huang et al., 2023) have shown that LoRA can be fine-tuned for diverse tasks—such as summarization, question answering, and classification—resulting in a set of pre-trained LoRA adapters (LoRA library).

The emergence of model hubs, such as Hugging Face, has further simplified the sharing and reuse of pre-trained and fine-tuned models. These platforms enable collaborative and multi-task learning by allowing users to easily acquire and extend existing models. Within this context, model merging—which aims to combine multiple specialized models into a single model capable of handling diverse tasks—has attracted significant attention (Wortsman et al., 2022; Yadav et al., 2023; Yang et al., 2024; Stoica et al., 2024; Gargiulo et al., 2025). However, merging task-specific models often introduces task interference, since parameter updates optimized for one task may conflict with those optimized for another (Yadav et al., 2023). To mitigate this issue, a variety of model merging methods have been proposed (Davari & Belilovsky, 2024; Ilharco et al., 2022; Matena & Raffel, 2022; Wortsman et al., 2022; Yadav et al., 2023; Yu et al., 2024; Lu et al., 2024; Yang et al., 2024; Deep et al., 2024; Miyano & Arase, 2025; Chen et al., 2025).

Despite these advances, prior work has observed that LoRA adapter weights exhibit weaker alignment compared to fully fine-tuned model weights. Consequently, applying existing model merging techniques directly to LoRAs often yields suboptimal results (Stoica et al., 2024; Panariello et al., 2025). To address this, researchers have explored LoRA-specific merging strategies, such as singular value decomposition (SVD)–based methods. The key idea is that SVD can decompose each LoRA adapter into orthogonal components and retain only the most significant singular vectors, which capture the core task-specific transformations (Lu et al., 2024; Marczak et al., 2025; Gargiulo et al.,

2025). However, current approaches (e.g., TSV-merging, Knots) apply SVD independently to each adapter. As a result, merging LoRAs based on independently derived singular vectors may fail to capture shared and task-specific components in a unified representation.

To address this limitation, a natural idea is to jointly decompose all LoRA adapters so that their shared and task-specific characteristics can be revealed. Concretely, we treat the collection of LoRA adapters as a unified structure by concatenating them into a third-order tensor and then performing tensor decomposition. Formally, given N tasks, each associated with a LoRA adapter $\Delta_i \in \mathbb{R}^{d_{in} \times d_{out}}$, we concatenate all adapters to construct a third-order tensor. Unlike prior approaches that apply independent SVD-based decompositions to mitigate task interference (Stoica et al., 2024; Gargiulo et al., 2025), we employ Canonical Polyadic (CP) decomposition to jointly factorize all task-specific matrices. CP decomposition expresses a tensor as a sum of rank-one components, thereby disentangling shared and task-specific patterns. This joint factorization not only helps reduce task interference but also preserves task-relevant information more effectively. Finally, we merge the LoRA adapters by summing over the task dimension.

To rigorously evaluate our CP-based merging approach, we conduct experiments on 10 held-in tasks and 3 held-out multi-tasks from the Flan datasets (Based on Phi-3 3B) Arnob et al. (2025), as well as on 10 out-of-domain downstream tasks (Based on Phi-3 3B and Mistral-7B) Ostapenko et al. (2024) and skill-composition tasks (Based on Llama 7B) Prabhakar et al. (2024). On the 10 held-in Flan tasks, CP merging outperforms strong SVD-based baselines by +3.19 Rouge-L. For the held-out Flan tasks, CP merging achieves an improvement of +5.69 Rouge-L over the SVD baseline. Interestingly, we observe that uniform merging yields the best performance among all methods on the held-out evaluation. For the 10 out-of-domain (zero-shot) tasks, CP merging surpasses SVD baselines by an average of +1.0%, while uniform merging also shows competitive performance in this setting. Finally, in the skill-composition evaluation, CP merging outperforms SVD-based methods by +2.0% accuracy on the challenging math-hard task.

In summary, our contributions are:

- We propose a CP merging approach that mitigates task interference in LoRA merging by jointly modeling shared and task-specific components in a unified manner.
- Experimental results show that CP merging substantially improves 10 held-in performance (+3.19 Rouge-L), 3 held-out tasks (+5.69 Rouge-L), and skill composition tasks (2% accuracy) over SVD-based methods. This demonstrates that CP decomposition effectively disentangles task-specific factors from shared factors, thereby reducing interference while preserving essential task information.

2 RELATED WORKS

PEFT. Parameter-efficient fine-tuning (PEFT) methods facilitate efficient adaptation of LLMs without updating all the training parameters, thereby reducing the memory and computation cost (Houlsby et al., 2019a; Zhang et al., 2023b; Zaken et al., 2021; Guo et al., 2020; Li & Liang, 2021; Lester et al., 2021). Hu et al. (2021) proposes a method named LoRA that parameterizes incremental weights Δ as a low-rank matrix by the product of the down projector matrix and up projector matrix. Zhang et al. (2023a) proposes Adaptive Low-Rank Adaptation (AdaLoRA), a new method that dynamically allocates the parameter budget among weight matrices during LoRA-like fine-tuning (Zhang et al., 2023a). Liu et al. (2024) introduces a new approach DoRA to investigate the inherent differences between full fine-tuning and LoRA.

LoRA library. In a multi-task scenario, LoRA adapters can be fine-tuned for diverse tasks-resulting in a set of pre-trained LoRA libraries. AdapterSoup Chronopoulou et al. (2023) trains each adapter for each domain and performs weight-space averaging of adapters trained on different domains. Huang et al. (2023) introduces LoRAhub to aggregate the LoRA modules trained on diverse tasks. AdapterFusion (Pfeiffer et al., 2020) proposes a two-stage algorithm that leverages knowledge from multiple tasks. Similar to LoRAhub, a group of task-specific adapters learn to encapsulate the task-specific information, and in the second stage, a fusion layer combines the trained adapters. Ostapenko et al. (2024) propose a model-based clustering library building methods and an Arrow routing function to reuse the LoRA library. Zhao et al. (2024a) propose a retrieval-augmented mix-

ture of LoRA Experts (RAMoLE) that adaptively retrieves and composes multiple LoRAs according to the input prompts.

Model Merging. Model merging integrates the weights of multiple task-specific models into a single multi-task model Davari & Belilovsky (2024); Ilharco et al. (2022); Matena & Raffel (2022); Wortsman et al. (2022); Yadav et al. (2023); Yu et al. (2024); Lu et al. (2024); Yang et al. (2024); Deep et al. (2024); Miyano & Arase (2025); Chen et al. (2025). Ilharco et al. (2022) presents the task vectors, which are the weight differences between pretrained models and the finetuned models. We can merge the task vectors to obtain a merged multi-task model. Ties Yadav et al. (2023) reduces the parameter redundancy by selecting the top- k most significant parameters and then constructing a sign vector based on the majority sign. DARE Yu et al. (2024) randomly drop delta parameters with a ratio p and rescales the remain parameters by $1/(1 - p)$ to reduce the task interference. Fisher Merging Matena & Raffel (2022) and RegMean Jin et al. (2022) merge models by performing weighted averaging, utilizing the Fisher information matrix and inner product of input vectors. Zhao et al. (2024b) proposes a LoRA-LEGO framework to conduct rank-wise parameter clustering by grouping MSUs from different LoRAs into clusters. TSV Gargiulo et al. (2025) shows that SVD decomposition can reduce the parameter redundancy, and merging the singular values can compress the parameters and reduce the task interference. Knots Stoica et al. (2024) uses the SVD to jointly transform the weights of different LoRA models into a shared space. ISO-C Marczak et al. (2025) propose an isotropic merging framework that flattens the isotropic merging framework to flatten singular value spectrum of the task matrix, enhancing alignment and reducing task interference. Instead of decomposing the task matrix separately, we concatenate the task matrix into a third-order tensor and then decompose it using CP decomposition.

3 PRELIMINARIES

3.1 CP DECOMPOSITION

Tensor decomposition aims to approximate a tensor through a set of low-rank factors with diverse applications. The most widely used decompositions are Tucker decomposition (De Lathauwer et al., 2000) and CANDECOMP/PARAFAC (CP) decomposition (Tucker, 1966), both of which can be seen as a generalization of the matrix singular value decomposition (SVD) (Stewart, 1993). We compare the SVD and CP decomposition in Fig. 1. CP decomposition can be seen as a general SVD decomposition for higher-order tensors. In this way, a high-order tensor can be uniquely represented as the sum of a set of rank-one tensors. Given a third-order tensor $\mathcal{T} \in \mathbb{R}^{d \times r \times N}$, the CP decomposition is expressed as:

$$\mathcal{T} = \sum_{i=1}^R \lambda_i \mathbf{u}_i \otimes \mathbf{v}_i \otimes \mathbf{s}_i \quad (1)$$

where R is the rank of CP decomposition. λ_i is the scaling weight. $\mathbf{u}_i$, $\mathbf{v}_i$ and $\mathbf{s}_i$ are the factor vectors corresponding to the three modes of $\mathcal{T}$. The operator $\otimes$ denotes the outer product.

4 METHOD

4.1 BACKGROUND

Multi-task learning is a powerful concept in AI that enables models to adapt and perform new tasks by leveraging previously acquired knowledge. In in-domain multi-task learning, a model f is trained on a set of n tasks $\{T_1, T_2, \dots, T_n\}$, which are considered in-domain tasks. The training process involves optimizing the model’s parameters to minimize the combined loss over all these tasks. For each task T_i , the dataset is denoted as $\mathcal{D}_i = \{(x_1, y_1), \dots, (x_m, y_m)\}$, where x_i are the inputs and y_i are the corresponding outputs.

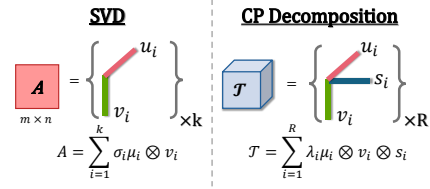


Figure 1: An illustration of SVD decomposition of a matrix $A \in \mathbb{R}^{m \times n}$ and CP decomposition of third-order tensor $\mathcal{T}$.

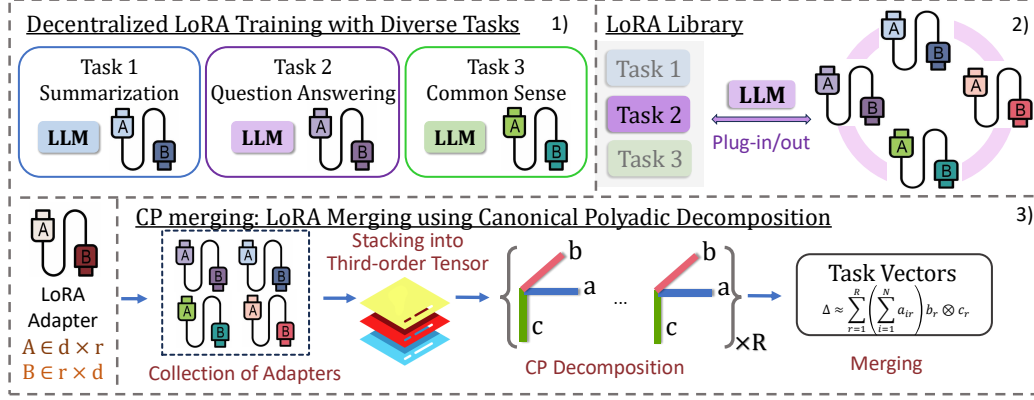


Figure 2: 1) LoRA adapter from various tasks, each LoRA adapter is trained individually. 2) LoRA library aims to leverage the plug-in/out nature of LoRAs to offer the ability to add or remove knowledge from LLMs. 3) CP merging: we stack the LoRA into a third-order tensor and conduct the CP merging.

After this training phase, the model f is evaluated on a new, unseen out-of-domain task, T_{new} , with a dataset $\mathcal{D}_{\text{new}}$. The model must predict outputs $\hat{y}$ for input x without having to train on $\mathcal{D}_{\text{new}}$ explicitly.

With the requirement to handle diverse tasks, the availability of LoRA library makes it possible to reuse the adapters for different tasks without retraining the entire model. Formally, the weights W_{MT} of a multi-task model for N tasks is obtained by aggregating the task-specific weight differences as follows:

$$W_{\text{MT}}^{(l)} = W_0^{(l)} + \alpha \frac{\sum_{i=1}^N \Delta_i^{(l)}}{N} \quad (2)$$

where W_0 is the set of pretrained model weights. α is a scaling factor and $\Delta_i^{(l)} = A_i^{(l)} (B_i^{(l)})^T$ is the LoRA adapter for task i at layer l . For brevity, we omit the layer index and refer to the matrix Δ_i^l at layer l as Δ_i .

4.2 CP MERGING

Task Singular Vectors (Gargiulo et al., 2025) use SVD to decompose the layer task matrices and term the obtained singular vectors *Task Singular Vectors (TSV)*, revealing low-rank properties and deeper insight into the inter-task interactions. Given two tasks i and j , the tasks matrix Δ_i and Δ_j can be written as:

$$\Delta_i = U_i \Sigma_i V_i^T, \quad \Delta_j = U_j \Sigma_j V_j^T \quad (3)$$

where U_i, U_j are the left singular vectors and V_i, V_j are the right singular vectors. Σ_i, Σ_j are diagonal matrices of singular values. As shown in Fig. 1, we can interpret SVD as a sum of rank-one matrices. For each task i , we get the best approximation of each task matrix Δ_i by retaining only the top- k singular values and their corresponding vectors:

$$\hat{\Delta}_i = \sum_{j=1}^k \sigma_j^i \mathbf{u}_j^i \otimes \mathbf{v}_j^i. \quad (4)$$

where $\hat{\Delta}_i$ is the low-rank approximation of the task matrix for task i , obtained using the top- k singular values $\sigma_{i,j}$, left singular vectors $\mathbf{u}_{i,j}$ and right singular vectors $\mathbf{v}_{i,j}$. Knots conduct the SVD

decomposition individually to obtain the right singular matrix $[V_1, V_2, \dots, V_N]$ and then merge the right singular matrix to obtain a merged matrix $V^{(\text{merged})}$. The final task matrix can be computed by: $\Delta W_j^{(\text{merged})} = U \Sigma [V^{(\text{merged})}]^T$. TSV merging decomposes the LoRA adapters individually and then concatenates the left singular matrix $[U_1, U_2, \dots, U_N]$ and right singular matrix $[V_1, V_2, \dots, V_N]$. We notice that SVD decomposes the task adapter separately, leading to independent choices of SVD dimensions. To address this, we propose a more general global decomposition using Canonical Polyadic (CP) decomposition, which takes all the LoRAs together as a tensor input.

Representing all adapters as a Third-Order Tensor We consider the task adapters Δ_i for all tasks i as slices of a third-order tensor $\mathcal{T}$, where the dimensions could be interpreted as (tasks, input dimension, output dimension). Let us denote this tensor $\mathcal{T} \in \mathbb{R}^{d_{in} \times d_{out} \times N}$ where N is the number of tasks, d_{in}, d_{out} are the dimensions of each task matrix Δ_i . Each Δ_i (a matrix of size $d_{in} \times d_{out}$) is a frontal slice of $\mathcal{T}$ along the task mode. So, we can concatenate the Δ_i matrices along the task dimension to form $\mathcal{T}$.

CP Decomposition Form In CP decomposition, the tensor $\mathcal{T}$ is approximated as a sum of rank-one tensors:

$$\mathcal{T} \approx \sum_{r=1}^R \mathbf{a}_r \otimes \mathbf{b}_r \otimes \mathbf{c}_r \quad (5)$$

where R is the number of components (analogous to k in SVD), $\mathbf{a}_r \in \mathbb{R}^N$ is a vector associated with the task mode, $\mathbf{b}_r \in \mathbb{R}^{d_{in}}$ and $\mathbf{c}_r \in \mathbb{R}^{d_{out}}$ are vectors associated with the row and column modes. For each task i , the i -th slice Δ_i of $\mathcal{T}$ would be:

$$\Delta_i \approx \sum_{r=1}^R a_{ir} \mathbf{b}_r \otimes \mathbf{c}_r \quad (6)$$

where a_{ir} is the i -th element of $\mathbf{a}_r$, selecting the contribution of the r -th component for task i .

CP merging over the task mode To merge the Δ_i into a final Δ , we sum over the task dimension (mode-1). The merged Δ would be:

$$\Delta = \sum_{i=1}^N \Delta_i \approx \sum_{i=1}^N \sum_{r=1}^R a_{ir} \mathbf{b}_r \otimes \mathbf{c}_r \quad (7)$$

Since the outer product $\mathbf{b}_r \otimes \mathbf{c}_r$ is the same for all i , and assuming a_{ir} represents the weight of component r for task i , we can factorize this as:

$$\Delta \approx \sum_{r=1}^R \left(\sum_{i=1}^N a_{ir} \right) \mathbf{b}_r \otimes \mathbf{c}_r \quad (8)$$

where $\sum_{i=1}^N a_{ir}$ is the total contribution of the r -th component across all tasks. The final merged weights can be written as :

$$W_{\text{MT}} = W_0 + \sum_{r=1}^R \left(\sum_{i=1}^N a_{ir} \right) \mathbf{b}_r \otimes \mathbf{c}_r \quad (9)$$

Algorithm 1 Implementation of CP Merging

- 1: **Input:** Pretrained task LoRA adapters $\{\Delta_1, \Delta_2, \dots, \Delta_N\}$, scaling factors α .
 - 2: **Output:** Merged weight W_{MT}
 - 3: **Concatenate** the matrices:
 - 4: $\mathcal{T} \leftarrow [\Delta_1 | \Delta_2 | \dots | \Delta_N]$
 - 5: **CP Decomposition:**
 - 6: $\mathcal{T} \approx \sum_{r=1}^R \mathbf{a}_r \otimes \mathbf{b}_r \otimes \mathbf{c}_r$ (Eq. 5)
 - 7: **Compute the Task Matrix:**
 - 8: $\Delta_i \approx \sum_{r=1}^R a_{ir} \mathbf{b}_r \otimes \mathbf{c}_r$ (Eq. 6)
 - 9: **Reconstruct** the merged matrix:
 - 10: $\Delta \approx \sum_{r=1}^R \left(\sum_{i=1}^N a_{ir} \right) \mathbf{b}_r \otimes \mathbf{c}_r$ (Eq. 8)
 - 11: **Construct** the merged model weights:
 - 12: $W_{\text{MT}} = W_0 + \alpha \sum_{r=1}^R \left(\sum_{i=1}^N a_{ir} \right) \mathbf{b}_r \otimes \mathbf{c}_r$
 - 13: **return** W_{MT}
-

5 EXPERIMENTAL RESULTS

To test the effectiveness of our approach, we evaluate CP Merging on a variety of LoRA merging methods in held-in and held-out multi-tasks based on Flan datasets (Sec. 5.1) Arnob et al. (2025); Ostapenko et al. (2024) and zero-shot task (Sec. 5.2) Ostapenko et al. (2024) and skill composition tasks (Appendix A.2) Prabhakar et al. (2024).

Baselines We compare the following methods in the zero-shot benchmark. **BASE**: The base model without adaptation. **Multi-task**: A single expert LoRA finetuned on a joint training set. **Uniform**: For each task, we train a LoRA adapter, then we merge the LoRA uniformly (Huang et al., 2023; Chronopoulou et al., 2023). **Task Arithmetic (TA)** (Ilharco et al., 2022) subtracts the parameter values of the pre-trained model from those of the fine-tuned models, creating a set of "task-vectors." Then they linearly summed to create a merged model. **Ties-merging** (Yadav et al., 2023): Ties improves TA by resolving the parameter interference between models when merging. It prunes low-magnitude weights and then only averages the weights that share the dominant sign. **DARE** randomly drops fine-tuned weights and rescales the remaining ones to create sparse task vectors Yu et al. (2024). **TSV merging** uses SVD decomposition to compress the LoRA and reduce the task interference based on the singular vectors (Marczak et al., 2025). **Knots** uses the SVD to transform the weights of jointly different LoRA models into a shared space Stoica et al. (2024). **ISO-C** Marczak et al. (2025) proposes an isotropic merging framework that flattens the singular value spectrum of the task matrix.

5.1 HELD-IN AND HELD-OUT MULTI-TASKS

We conduct our experiments using the FLAN dataset (Longpre et al., 2023) and sample 10 held-in and 3 held-out tasks following (Ostapenko et al., 2024; Arnob et al., 2025). Each task is sub-sampled to 10,000 examples. Within these samples, 1,000 are allocated for validation and early stopping.

Model	Wiq	Sci	Adver	Duorc	Cos	Wikiqa	Quail	WikiHop	ParaphQa	Yelp	AVG	3 Held-Out
Phi-3 (3B) with LoRA library												
BASE	14.55	7.99	10.85	7.32	5.00	4.17	11.88	3.51	6.93	8.33	8.05	16.00
Oracle	36.87	79.56	51.58	52.40	42.99	55.94	51.51	46.58	8.325	32.71	45.85	-
Uniform	22.98	61.97	29.61	18.45	28.12	11.56	38.94	7.194	7.18	15.66	24.17	61.13
Ties-merging	23.33	51.52	30.15	18.05	30.12	12.54	35.05	10.03	5.80	16.76	23.34	34.69
Ties w/DARE	25.42	52.76	32.34	19.04	31.20	13.33	36.05	7.05	6.08	16.00	23.93	34.70
ISO-C	14.59	8.19	11.11	7.50	5.38	4.20	12.23	3.53	6.95	8.39	8.20	16.18
TSV Merging	36.75	53.52	44.39	20.20	59.13	23.20	37.25	34.70	4.78	34.95	34.89	45.70
Knots	26.33	90.81	40.15	28.87	50.44	48.35	43.24	13.71	6.63	23.56	<u>37.21</u>	54.46
CP Merging	34.69	89.88	44.32	31.59	66.82	50.34	45.49	10.73	5.90	24.23	40.40	<u>60.15</u>

Table 1: 10 held-in and 3 held-out Flan tasks based on Phi-3 3B model. We report the Rouge-L scores. The oracle here denotes that we use the trained LoRA to test the corresponding tasks, which achieve the best results for held-in results. The tasks details are presented in Appendix A.3.1.

As shown in Tab. 1. The CP Merging method achieves the highest average score of 40.40, indicating it generally performs best across the tested datasets. It also has a strong 3-Held-Out score of 60.15. CP Merging is the standout performer, showing significant gains over other methods, especially in the average score and the 3-Held-Out evaluation. The BASE model provides a baseline performance, while the Oracle provides an upper bound, representing the theoretical maximum performance without any merging techniques. Ties-merging and Ties w/DARE show similar average performance (23-24). ISO-C has a significantly lower average score, similar to the BASE model’s performance. TSV Merging and Knots show moderate performance, with Knots having a higher average score and a much higher 3-Held-Out score. In short, the CP Merging method appears to be the most effective overall, while Uniform merging excels on the specific 3-Held-Out evaluation.

5.2 ZERO-SHOT EXPERIMENTS

To test the out-of-domain generalization, we evaluate CP Merging on a variety of zero-shot task benchmarks. Specifically, we experiment with four broad classes of zero-shot tasks: (1) **Common Sense Reasoning**: WinoGrande (Sakaguchi et al., 2021), HellaSwag (Zellers et al., 2019), and PIQA

	Common Sense			Question Answering				Coding		Reasoning	AVG
	Piqa	Wg	Hswag	Boolq	Obqa	ArcE	ArcC	HE	Mbpp	BBH	AVG
Phi-3 (3B) with LoRA library											
BASE	81.1	70.2	75.7	85.0	49.0	80.0	57.9	53.0	61.1	53.4	66.6
Multi-task	79.1	69.6	75.6	87.1	47.6	82.6	55.6	51.8	52.9	52.1	65.4
Uniform	81.0	70.4	76.0	84.7	48.8	82.4	57.9	52.4	63.0	55.6	67.2
Ties-merging	80.8	70.9	75.4	80.8	46.0	85.8	60.8	54.9	61.5	51.9	66.9
Task Arithmetic	81.0	70.5	75.9	84.7	48.8	82.5	58.2	54.3	61.9	54.2	67.2
Ties w/DARE	80.7	70.8	75.2	82.8	46.3	84.8	60.9	54.9	61.6	52.0	67.0
TA w/DARE	81.2	70.1	75.3	84.9	48.9	82.2	58.3	54.5	61.1	54.3	67.1
TSV Merging	81.1	70.8	75.8	86.0	49.0	84.3	60.4	55.0	61.2	54.0	67.7
C-LoRA	81.5	69.5	76.0	86.6	48.8	86.7	61.8	57.3	65.3	55.9	68.9
Knots	80.8	71.7	74.9	85.0	48.0	84.1	58.2	53.2	62.6	50.4	66.9
TSV Merging	81.3	69.3	75.6	86.4	48.8	86.1	62.3	52.4	63.3	55.1	68.1
CP merging	81.7	70.0	76.0	86.6	48.9	86.7	62.5	57.3	65.4	56.0	69.1

Table 2: 10 downstream Zero-shot results based on Phi-3 (microsoft/Phi-3-mini-4k-instruct). C-LoRA is a LoRA library based on embedding clustering (A.4). We show the results based on Mistral in Appendix A.5. The experimental setting is presented in Appendix A.3.

(Bisk et al., 2020). 2) **Question Answering**: BoolQ Clark et al. (2019), OpenbookQA Mihaylov et al. (2018), ARC-easy Clark et al. (2018), and ARC-challenge Clark et al. (2018). 3) **Coding**: HumanEval Chen et al. (2021), MBPP Austin et al. (2021). 4) **General-Purpose Reasoning**: BBH (Suzgun et al., 2022).

5.2.1 EXPERIMENTAL RESULTS

Tab. 2 presents the average zero-shot accuracy results across 10 downstream tasks. The multi-task based on the Phi-3 model (65.4%) underperforms the base model (66.6%), indicating that the multi-task training may encounter *catastrophic forgetting* issue that newly learned parameters overwrite or interfere with knowledge acquired during prior training (Wang et al., 2024). LoRA library with uniform merging, consisting of 256 experts, achieves accuracies of 67.2% for Phi-3.

Compared to the C-LoRA library, TSV Merging and Knots are comparable in overall performance. TSV Merging has an average score of 68.1, slightly better than Knots’ average of 66.9. While TSV Merging and Knots perform well in some tasks, they do not consistently achieve the top scores that CP merging does. For example, TSV Merging ties for the highest score in Obqa (49.0), but its overall performance is less dominant. CP merging consistently shows superior performance. It achieves the highest average score of 69.1, outperforming all other methods. The results clearly demonstrate the effectiveness of CP merging, which consistently outperforms the other methods, including Knots and TSV Merging. With the highest average score and leading performance in multiple specific tasks, CP merging shows a more robust and effective approach to merging LoRA adapters, successfully mitigating task interference and enhancing overall model performance.

We further analyze performance across different task categories (Fig. 3). CP merging improves over the LoRA library with TSV merging, with the largest gains observed in question answering and coding tasks. In contrast, improvements on common-sense reasoning are minimal, suggesting that the impact of task interference varies by task type.

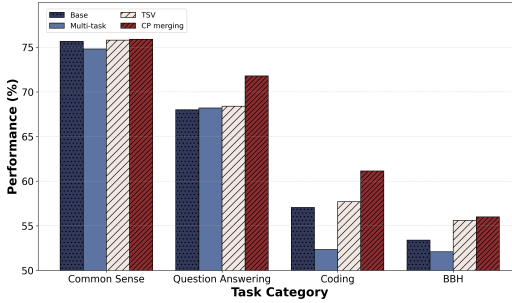


Figure 3: Zero-shot results across different categories. The X-axis represents the task categories.

6 ABLATION STUDIES

6.1 NUMBER OF TASKS ANALYSIS

Overall, the merging methods significantly outperformed the Base model across all task sizes. The Base model consistently showed the lowest ROUGE-L scores, with scores of 11.14, 9.15, and 8.05 for 3, 5, and 10 tasks, respectively. This suggests that LoRA merging is an effective technique for improving performance. CP consistently delivered the best performance on ROUGE-L scores. For 5 tasks, TSV merging and Knots merging performed equally well, both achieving a score of 54.4. The CP merging method surpasses Knots merging with a score of 57.7 for this task size. For 10 tasks, both the Knots and TSV merging methods saw a considerable drop in performance. In contrast, the CP merging method demonstrated superior robustness to an increase in the number of tasks, achieving the highest score of 40.40 in this category. The results indicate that LoRA merging is a highly effective strategy for improving model performance on multiple tasks. Specifically, the CP merging method shows the most promise, as it consistently achieves top-tier performance and exhibits greater stability as the number of tasks increases.

Model	3 tasks	5 tasks	10 tasks
Base	11.14	9.15	8.05
TSV merging	44.87	54.4	34.89
Knots	59.04	54.4	37.21
CP merging	59.06	57.7	40.40

Table 3: The influence of the task size for LoRA merging. We select 3, 5, and 10 tasks from the Flan datasets and report the average Rouge-L scores for all the tasks.

6.2 CP MERGING FOR TASK INTERFERENCE

To study the task interference, Gargiulo et al. (2025) introduces a score of task interference (termed as *Singular Task Interference*) based on the interplay of TSVs from different tasks:

$$\text{STI}(\{\Delta_i\}_{i=1}^N) = \|(U^\top U - I) \Sigma (V^\top V - I)\|_1 \quad (10)$$

They assume that higher inner product values for $U^\top U$ and $V^\top V$ imply a higher likelihood (Gargiulo et al., 2025). The intuition is that overlapping singular vectors suggest shared features in the weight space across tasks.

Inspired by this, we can reformulate the STI with CP decomposition. To define a CP-based STI metric, we need to assess the interference caused by the interplay of these factor matrices across tasks. The intuition remains that overlapping components (shared features in the weight space) lead to interference when merged. We can adapt the STI by considering the alignment and orthogonality of the factor matrices a_r , b_r , and c_r across tasks. Then we define the CP-based Singular Task Interference (CP-STI) as:

$$\text{CP-STI}(\{\Delta_i\}_{i=1}^N) = \|(A^\top A - I) \circ (B^\top B - I) \circ (C^\top C - I)\|_1 \quad (11)$$

where $A = [a_1, a_2, \dots, a_R] \in \mathbb{R}^{N \times R}$ is the matrix formed by stacking the task-mode factor vectors a_r , $B = [b_1, b_2, \dots, b_R] \in \mathbb{R}^{d \times R}$ is the matrix of row-mode factors, $C = [c_1, c_2, \dots, c_R] \in \mathbb{R}^{d \times R}$ is the matrix of column-mode factors, $\circ$ denotes the Hadamard (element-wise) product, I is the identity matrix of appropriate dimension (e.g., $R \times R$). The CP-STI metric captures interference by evaluating how the factor matrices A , B , and C deviate from being mutually orthogonal. Overlapping factors (high inner products) suggest shared features across tasks, which can introduce interference when the decomposed representations are merged, potentially degrading individual task performance.

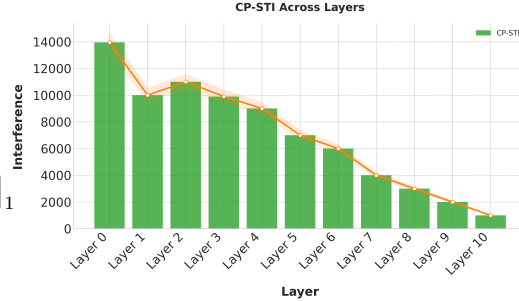


Figure 4: CP-STI across first 10 layers in the Llama 7B for Math and Code LoRA.

We study the task interference on each layer according to Eq. (11). As shown in the Fig. 4, the interference is higher in the lower layer and decreases in the deeper layers. This may indicate that the lower layers capture common features and deeper layers are more specialized for specific tasks.

6.3 RANK R OF THE CP DECOMPOSITION.

CP decomposition factorizes a tensor into a sum of R rank-one components. Wang et al. (2023) observed that increasing R can improve performance on cross-modal tasks. We investigate how the CP rank affects our method’s performance on the GSM8K-hard dataset (Fig.5). The left plot (“CP Rank vs. Accuracy”) shows accuracy rising from 0.12 at rank 5 to a peak of 0.15 at rank 20, with no further gains beyond rank 25. This suggests that higher ranks capture richer patterns, but benefits plateau after rank 20. The right plot (“CP Rank vs. Invalid Code”) shows invalid code outputs dropping from 340 at rank 5 to 200 at rank 20, with little improvement at higher ranks. Overall, increasing the CP rank improves both accuracy and output validity up to an optimal point (around rank 20), after which returns diminish, highlighting the importance of selecting an appropriate rank to balance performance and efficiency in mitigating task interference in skill composition tasks.

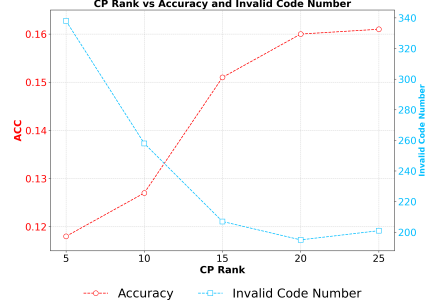


Figure 5: The impact of CP rank on model performance and generation quality for the GSM-8K-hard benchmark.

6.4 ALPHA ANALYSIS

The alpha parameter effectively scales the influence of the adapter’s output before it is added back to the weights of the original model. Table 6 reports the ROUGE-L scores for 10 held-in Flan tasks across seven α values ranging from 0.3 to 0.1. The results highlight the critical importance of the α hyperparameter for LoRA merging methods—including CP merging, knots, TSV merging, and ISO merging. Notably, the optimal α is not universal but depends on the specific task, underscoring the need for careful tuning. On average, moderate values (e.g., $\alpha = 0.15$ – 0.18) yield the best performance. For all experiments, α was selected based on performance over a validation set of 1,000 examples. Finally, we use 0.25 for TSV merging and 0.50 for Knots.

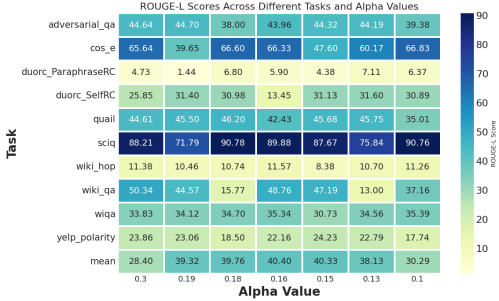


Figure 6: Best performance for different alpha values over 10 held-in Flan datasets.

7 CONCLUSION

We propose a novel LoRA merging method using Canonical Polyadic (CP) decomposition. Unlike existing SVD-based methods that decompose each LoRA adapter independently, our approach factorizes multiple adapters in a unified manner. This allows us to disentangle task-specific components from those that are shared across tasks. Our comprehensive experiments show that CP merging consistently outperforms strong SVD-based baselines across various benchmarks, including in-domain multi-tasks, zero-shot tasks, and skill composition tasks. These results confirm that our method effectively preserves essential task knowledge while significantly reducing cross-task interference. This unified factorization approach is a key advantage over conventional merging techniques. Our findings offer a new perspective on how knowledge is represented and merged within LoRA adapters. The success of CP decomposition suggests that task-specific knowledge and shared knowledge are not merely independent entities but are intricately interwoven. This insight—that a joint, multi-task factorization is superior to a series of independent ones—indicates a promising new direction for future research in parameter-efficient fine-tuning and model merging.

REFERENCES

- Marah Abidin, Sam Ade Jacobs, Ammar Ahmad Awan, Jyoti Aneja, Ahmed Awadallah, Hany Awadalla, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Harkirat Behl, et al. Phi-3 technical report: A highly capable language model locally on your phone. *arXiv preprint arXiv:2404.14219*, 2024.
- Samin Yeasar Arnob, Zhan Su, Minseon Kim, Oleksiy Ostapenko, Riyasat Ohib, Esra’ Saleh, Doina Precup, Lucas Caccia, and Alessandro Sordoni. Exploring sparse adapters for scalable merging of parameter efficient experts. *arXiv preprint arXiv:2507.07140*, 2025.
- Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*, 2021.
- Stephen H Bach, Victor Sanh, Zheng-Xin Yong, Albert Webson, Colin Raffel, Nihal V Nayak, Abheesht Sharma, Taewoon Kim, M Saiful Bari, Thibault Fevry, et al. Promptsources: An integrated development environment and repository for natural language prompts. *arXiv preprint arXiv:2202.01279*, 2022.
- Yonatan Bisk, Rowan Zellers, Jianfeng Gao, Yejin Choi, et al. Piqa: Reasoning about physical commonsense in natural language. In *Proceedings of the AAAI conference on artificial intelligence*, pp. 7432–7439, 2020.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, et al. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*, 2021.
- Shilian Chen, Jie Zhou, Tianyu Huai, Yujiang Lu, Junsong Li, Bihao Zhan, Qianjun Pan, Yutao Yang, Xin Li, Qin Chen, et al. Black-box model merging for language-model-as-a-service with massive model repositories. *arXiv preprint arXiv:2509.12951*, 2025.
- Alexandra Chronopoulou, Matthew E Peters, Alexander Fraser, and Jesse Dodge. Adaptersoup: Weight averaging to improve generalization of pretrained language models. *arXiv preprint arXiv:2302.07027*, 2023.
- Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Yunxuan Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, et al. Scaling instruction-finetuned language models. *Journal of Machine Learning Research*, 25(70):1–53, 2024.
- Christopher Clark, Kenton Lee, Ming-Wei Chang, Tom Kwiatkowski, Michael Collins, and Kristina Toutanova. Boolq: Exploring the surprising difficulty of natural yes/no questions. In *NAACL*, 2019.
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. Think you have solved question answering? try arc, the ai2 reasoning challenge. *arXiv preprint arXiv:1803.05457*, 2018.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021.
- MohammadReza Davari and Eugene Belilovsky. Model breadcrumbs: Scaling multi-task model merging with sparse masks. In *European Conference on Computer Vision*, pp. 270–287. Springer, 2024.
- Lieven De Lathauwer, Bart De Moor, and Joos Vandewalle. A multilinear singular value decomposition. *SIAM journal on Matrix Analysis and Applications*, 21(4):1253–1278, 2000.
- Pala Tej Deep, Rishabh Bhardwaj, and Soujanya Poria. Della-merging: Reducing interference in model merging through magnitude-based sampling. *arXiv preprint arXiv:2406.11617*, 2024.
- Luyu Gao, Aman Madaan, Shuyan Zhou, Uri Alon, Pengfei Liu, Yiming Yang, Jamie Callan, and Graham Neubig. Pal: Program-aided language models. In *International Conference on Machine Learning*, pp. 10764–10799. PMLR, 2023.

- 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 *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 18695–18705, 2025.
- Demi Guo, Alexander M Rush, and Yoon Kim. Parameter-efficient transfer learning with diff pruning. *arXiv preprint arXiv:2012.07463*, 2020.
- Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. Parameter-efficient transfer learning for nlp. In *International Conference on Machine Learning*, pp. 2790–2799. PMLR, 2019a.
- Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. Parameter-efficient transfer learning for nlp. In *International conference on machine learning*, pp. 2790–2799. PMLR, 2019b.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*, 2021.
- Chengsong Huang, Qian Liu, Bill Yuchen Lin, Tianyu Pang, Chao Du, and Min Lin. Lorahub: Efficient cross-task generalization via dynamic lora composition. *arXiv preprint arXiv:2307.13269*, 2023.
- Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Suchin Gururangan, Ludwig Schmidt, Hannaneh Hajishirzi, and Ali Farhadi. Editing models with task arithmetic. *arXiv preprint arXiv:2212.04089*, 2022.
- AQ Jiang, A Sablayrolles, A Mensch, C Bamford, DS Chaplot, D de las Casas, F Bressand, G Lengyel, G Lample, L Saulnier, et al. Mistral 7b (2023). *arXiv preprint arXiv:2310.06825*, 2023.
- Xisen Jin, Xiang Ren, Daniel Preotiuc-Pietro, and Pengxiang Cheng. Dataless knowledge fusion by merging weights of language models. *arXiv preprint arXiv:2212.09849*, 2022.
- Brian Lester, Rami Al-Rfou, and Noah Constant. The power of scale for parameter-efficient prompt tuning. *arXiv preprint arXiv:2104.08691*, 2021.
- Xiang Lisa Li and Percy Liang. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*, 2021.
- Shih-Yang Liu, Chien-Yi Wang, Hongxu Yin, Pavlo Molchanov, Yu-Chiang Frank Wang, Kwang-Ting Cheng, and Min-Hung Chen. Dora: Weight-decomposed low-rank adaptation. *arXiv preprint arXiv:2402.09353*, 2024.
- Shayne Longpre, Le Hou, Tu Vu, Albert Webson, Hyung Won Chung, Yi Tay, Denny Zhou, Quoc V Le, Barret Zoph, Jason Wei, et al. The flan collection: Designing data and methods for effective instruction tuning. In *International Conference on Machine Learning*, pp. 22631–22648. PMLR, 2023.
- Zhenyi Lu, Chenghao Fan, Wei Wei, Xiaoye Qu, Danyang Chen, and Yu Cheng. Twin-merging: Dynamic integration of modular expertise in model merging. *Advances in Neural Information Processing Systems*, 37:78905–78935, 2024.
- 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. *arXiv preprint arXiv:2502.04959*, 2025.
- Michael S Matena and Colin A Raffel. Merging models with fisher-weighted averaging. *Advances in Neural Information Processing Systems*, 35:17703–17716, 2022.
- Todor Mihaylov, Peter Clark, Tushar Khot, and Ashish Sabharwal. Can a suit of armor conduct electricity? a new dataset for open book question answering. *arXiv preprint arXiv:1809.02789*, 2018.

- Ryota Miyano and Yuki Arase. Adaptive lora merge with parameter pruning for low-resource generation. *arXiv preprint arXiv:2505.24174*, 2025.
- Oleksiy Ostapenko, Zhan Su, Edoardo Maria Ponti, Laurent Charlin, Nicolas Le Roux, Matheus Pereira, Lucas Caccia, and Alessandro Sordani. Towards modular llms by building and reusing a library of loras. *arXiv preprint arXiv:2405.11157*, 2024.
- Aniello Panariello, Daniel Marczak, Simone Magistri, Angelo Porrello, Bartłomiej Twardowski, Andrew D Bagdanov, Simone Calderara, and Joost van de Weijer. Accurate and efficient low-rank model merging in core space. *arXiv preprint arXiv:2509.17786*, 2025.
- Jonas Pfeiffer, Aishwarya Kamath, Andreas Rücklé, Kyunghyun Cho, and Iryna Gurevych. Adapterfusion: Non-destructive task composition for transfer learning. *arXiv preprint arXiv:2005.00247*, 2020.
- Akshara Prabhakar, Yuanzhi Li, Karthik Narasimhan, Sham Kakade, Eran Malach, and Samy Jelassi. Lora soups: Merging loras for practical skill composition tasks. *arXiv preprint arXiv:2410.13025*, 2024.
- Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhagavatula, and Yejin Choi. Winogrande: An adversarial winograd schema challenge at scale. *Communications of the ACM*, 64(9):99–106, 2021.
- Ying Sheng, Shiyi Cao, Dacheng Li, Coleman Hooper, Nicholas Lee, Shuo Yang, Christopher Chou, Banghua Zhu, Lianmin Zheng, Kurt Keutzer, et al. S-lora: Serving thousands of concurrent lora adapters. *arXiv preprint arXiv:2311.03285*, 2023.
- Gilbert W Stewart. On the early history of the singular value decomposition. *SIAM review*, 35(4):551–566, 1993.
- George Stoica, Pratik Ramesh, Boglarka Ecsedi, Leshem Choshen, and Judy Hoffman. Model merging with svd to tie the knots. *arXiv preprint arXiv:2410.19735*, 2024.
- Mirac Suzgun, Nathan Scales, Nathanael Schärli, Sebastian Gehrmann, Yi Tay, Hyung Won Chung, Aakanksha Chowdhery, Quoc V Le, Ed H Chi, Denny Zhou, et al. Challenging big-bench tasks and whether chain-of-thought can solve them. *arXiv preprint arXiv:2210.09261*, 2022.
- Ledyard R Tucker. Some mathematical notes on three-mode factor analysis. *Psychometrika*, 31(3):279–311, 1966.
- Haixin Wang, Xinlong Yang, Jianlong Chang, Dian Jin, Jinan Sun, Shikun Zhang, Xiao Luo, and Qi Tian. Parameter-efficient tuning of large-scale multimodal foundation model. *Advances in Neural Information Processing Systems*, 36:15752–15774, 2023.
- Liyuan Wang, Xingxing Zhang, Hang Su, and Jun Zhu. A comprehensive survey of continual learning: theory, method and application. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M Dai, and Quoc V Le. Finetuned language models are zero-shot learners. *arXiv preprint arXiv:2109.01652*, 2021.
- Mitchell Wortsman, Gabriel Ilharco, Samir Ya Gadre, Rebecca Roelofs, Raphael Gontijo-Lopes, Ari S Morcos, Hongseok Namkoong, Ali Farhadi, Yair Carmon, Simon Kornblith, et al. Model soups: averaging weights of multiple fine-tuned models improves accuracy without increasing inference time. In *International conference on machine learning*, pp. 23965–23998. PMLR, 2022.
- Prateek Yadav, Derek Tam, Leshem Choshen, Colin A Raffel, and Mohit Bansal. Ties-merging: Resolving interference when merging models. *Advances in Neural Information Processing Systems*, 36:7093–7115, 2023.
- Enneng Yang, Li Shen, Guibing Guo, Xingwei Wang, Xiaochun Cao, Jie Zhang, and Dacheng Tao. Model merging in llms, mllms, and beyond: Methods, theories, applications and opportunities. *arXiv preprint arXiv:2408.07666*, 2024.

- 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 *Forty-first International Conference on Machine Learning*, 2024.
- Elad Ben Zaken, Shauli Ravfogel, and Yoav Goldberg. Bitfit: Simple parameter-efficient fine-tuning for transformer-based masked language-models. *arXiv preprint arXiv:2106.10199*, 2021.
- Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. Hellaswag: Can a machine really finish your sentence? *arXiv preprint arXiv:1905.07830*, 2019.
- Qingru Zhang, Minshuo Chen, Alexander Bukharin, Pengcheng He, Yu Cheng, Weizhu Chen, and Tuo Zhao. Adaptive budget allocation for parameter-efficient fine-tuning. *arXiv preprint arXiv:2303.10512*, 2023a.
- Renrui Zhang, Jiaming Han, Aojun Zhou, Xiangfei Hu, Shilin Yan, Pan Lu, Hongsheng Li, Peng Gao, and Yu Qiao. Llama-adapter: Efficient fine-tuning of language models with zero-init attention. *arXiv preprint arXiv:2303.16199*, 2023b.
- Ziyu Zhao, Leilei Gan, Guoyin Wang, Yuwei Hu, Tao Shen, Hongxia Yang, Kun Kuang, and Fei Wu. Retrieval-augmented mixture of lora experts for uploadable machine learning. *arXiv preprint arXiv:2406.16989*, 2024a.
- Ziyu Zhao, Tao Shen, Didi Zhu, Zexi Li, Jing Su, Xuwu Wang, Kun Kuang, and Fei Wu. Merging loras like playing lego: Pushing the modularity of lora to extremes through rank-wise clustering. *arXiv preprint arXiv:2409.16167*, 2024b.

A APPENDIX

A.1 LIMITATIONS

In this work, we evaluate our approach on natural language tasks using models such as Phi-3B, Mistral-7B, and LLaMA-7B. Extending the analysis to larger-scale models remains an important direction for future research. Our experiments show that CP merging achieves superior performance compared to SVD-based LoRA merging methods. However, on held-out and zero-shot tasks, uniform merging achieves competitive—and in some cases even better—results, suggesting that improving the generalization ability of LoRA merging methods remains an open challenge. Moreover, our current evaluation is limited to at most 10 tasks. Scaling to a larger number of tasks presents additional challenges for LoRA merging, which we plan to investigate in future work.

Use of large language models statement We use the LLM to polish the writing. All other parts, including experimental results, analyses were written by the authors and carefully verified for accuracy before and after any LLM-assisted editing.

A.2 SKILL COMPOSITION TASKS

GSM8k vs GSM8k-hard

GSM8k: A robe takes 2 bolts of blue fiber and half that much white fiber. How many bolts in total does it take?

Reference Answer: 18

GSM8k-hard: A robe takes 2287720 bolts of blue fiber and half that much white fiber. How many bolts in total does it take?

Reference Answer: 3,431,580

Figure 7: The comparison between GSM8k and GSM8k-hard. The GSM8K uses a larger number to replace the original number.

Prabhakar et al. (2024) have used LoRA merging to achieve the *skill composition*. Skill here refers to specific capabilities that the LLM needs for customization to downstream use cases. For example, achieving a high score in GSM8k benchmark Cobbe et al. (2021) needs good *commonsense* and *arithmetic* skills.

To explicitly examine how reducing task interference, we evaluate CP merging on hard math word problems. Compared to GSM-8K (Cobbe et al., 2021), GSM8K-hard (Gao et al., 2023) contains problems of similar types but with more complex arithmetic operations. Gao et al. (2023) propose a program-aided method in which an LLM first generates a program as an intermediate reasoning step and then executes it using a Python interpreter to obtain the final answer. This setting requires the LLM to be proficient in both mathematical reasoning (to analyze the problem) and coding (to translate the reasoning into executable code). Following the same experimental setting in Prabhakar et al. (2024), we train separate LoRA adapters for math and code skills.

A.2.1 GSM8K-HARD RESULTS

As depicted in Tab. 4, we evaluate the accuracy scores on the GSM-8k hard benchmark across various models. The base model achieves a modest accuracy of 0.043, highlighting its limitations in tackling GSM-hard tasks effectively. By training a LoRA adapter with the MathQA instruction dataset, we enhance the model’s mathematical capabilities. Similarly, training with the Alpaca-code dataset boosts the model’s coding proficiency. Combining both MathQA and Alpaca-code datasets in a single adapter yields an accuracy of 0.1349. Merging the math and coding skills through LoRA merging results in an accuracy of 0.1296, demonstrating the potential to integrate these abilities. TSV Merging and Knots perform similarly in terms of accuracy, with scores of 0.1349 and 0.1344, respectively. However, TSV Merging generates fewer invalid codes (224) compared to Knots (263), suggesting it is slightly more stable in this regard. Among the competing baselines, our proposed CP merging achieves the highest accuracy of 0.1569. Additionally, we assess the “invalid code”

Model	ACC	Invalid Code
BASE	0.0430	524
LoRA(math)	0.1175	562
LoRA(code)	0.0796	159
Multi-task	0.1349	413
Uniform	0.1296	217
Ties-merging	0.1060	223
Task Arithmetic	0.1190	232
Ties w/DARE	0.1162	227
TA w/DARE	0.1200	247
ISO Merging	0.0840	394
TSV Merging	0.1349	224
Knots	0.1344	263
CP merging	0.1569	201

Table 4: Evaluation results on MATH-hard tasks. "Invalid" is the number of invalid codes that can not be executed by Python.

metrics across methods. Notably, the LoRA code approach reduces invalid code instances from 524 in the base model to 159, indicating improved code quality.

The results demonstrate that our CP merging effectively reduces task interference by integrating math and coding skills, as evidenced by the lower number of invalid codes (201) compared to the base model (524) and other methods. In conclusion, CP merging not only achieves the highest accuracy but also maintains a low rate of invalid code generation, clearly demonstrating its superior performance in handling complex math-hard tasks.

A.3 EXPERIMENTAL SETTING FOR FLAN TASKS AND ZERO-SHOT DOWNSTREAM TASKS.

To enhance the LLM with various abilities for the above tasks, we construct a collection of LoRA adapters on diverse tasks. We use the FlanV2, an instruction-tuning dataset designed to scale both task diversity and model size, which has been shown to improve model performance significantly (Chung et al., 2024). We train each LoRA adapter independently, enabling it to specialize in a specific task. Finally, we select 256 tasks (Longpre et al., 2023), with the majority derived from P3 (Bach et al., 2022) and Flan2021 (Wei et al., 2021).

To evaluate the effectiveness of our approach across different backbone models, we conduct experiments using Phi-3 (Abdin et al., 2024), and Mistral-7B(mistralai/Mistral-7B-v0.1) (Jiang et al., 2023). In all cases, we apply LoRA adapters exclusively to the attention layers. Unless stated otherwise, our multi-task training and single-task adaptation scenarios employ a LoRA rank of 4, a dropout rate of 0.05, and a learning rate of 1e-4. Each LoRA adapter is trained for 5 epochs on either an H100 or A100 GPU. For the clustering step, we use the `sentence-transformers/sentence-t5-xxl` model for encoding. Given the typically large size of the dataset, we sample 20% of the data to train the k-means algorithm and then use the resulting clusters to predict labels for all samples.

A.3.1 10 HELD-IN AND 3 HELD OUT TASK NAMES

There are 256 pre-trained LoRA adapters used in Ostapenko et al. (2024). We randomly select 10 tasks as the held-in tasks and 3 tasks as the held-out tasks, followed by Arnob et al. (2025).

Held-in tasks

- `wiqa_what_is_the_final_step_of_the_following_process`
- `sciq_Multiple_Choice`
- `adversarial_qa_droberta_answer_the_following_q`

- duorc_SelfRC_question_answering
- cos_e_v1_l1_description_question_option_id
- wiki_qa_Is_This_True_
- quail_description_context_question_text
- wiki_hop_original_explain_relation
- duorc_ParaphraseRC_build_story_around_qa
- yelp_polarity_reviews_0.2_0

3 held-out tasks

- glue_sst2_2_0_0
- dream_read_the_following_conversation_and_answer_the_question
- race_middle_Read_the_article_and_answer_the_question_no_option_

A.4 CLUSTERIZED LORA LIBRARY (C-LoRA)

As noted in Ostapenko et al. (2024), LoRA similarity can contribute to positive transfer. Their approach clusters tasks based on LoRA similarity, requiring the *task ID* for each example and the pre-training of a collection of LoRA adapters before clustering. In contrast, we argue that text-level clustering is more convenient while achieving competitive results. As shown in Fig. 8, we select 100 examples per cluster and visualize their embeddings. The embeddings reveal that similar texts are grouped into the same clusters, demonstrating the quality of our embedding representations.

To mitigate task conflicts at the *text-level*, we train each *expert* on a collection of similar tasks. This is achieved by clustering instructions and partitioning the training data into multiple clusters without requiring manual intervention. Clustering input texts into several groups can reduce task interference by allowing the model to adapt its task adapters or representations to distinct subsets of data with similar characteristics. Formally, let $E(\cdot)$ denote a pretrained sentence encoder. For each data sample, the sentence representation of an instruction I_i is computed as $e_i = E(I_i)$. We then apply the k-means algorithm to cluster all instruction embeddings $\{e_i\}$ in the training dataset into K clusters. We train each LoRA adapter based on the instructions from each cluster. For each pre-trained weight W_0 , we obtain corresponding LoRA adapters $\{A_i B_i, \dots, A_K B_K\}$.

A.4.1 CLUSTERING TO REDUCE THE TASK INTERFERENCE AT THE TEXT-LEVEL

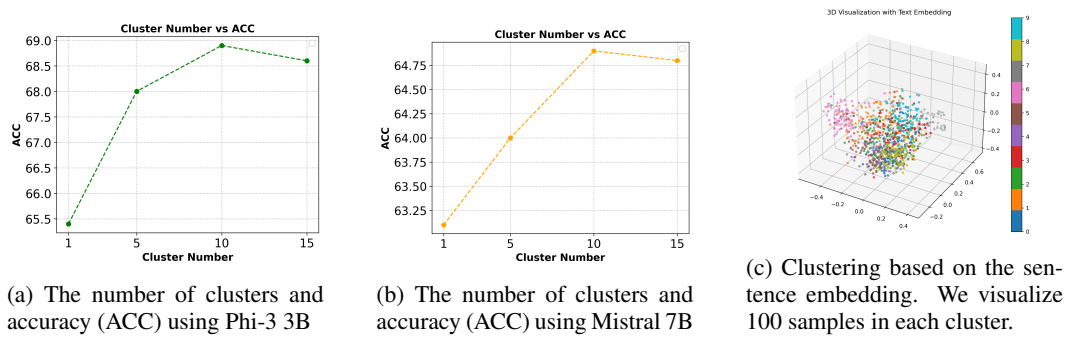


Figure 8: Cluster analysis. The number of clusters and accuracy (ACC) using Phi-3 3B and Mistral 7B.

The Fig. 8 shows the relationship between the number of clusters and accuracy (ACC) using two different model backbones: Phi-3 3B (left) and Mistral 7B (right). For the Phi-3 3B backbone, accuracy increases from approximately 65.4% at 1 cluster to a peak of 68.9% at 10 clusters, stabilizing thereafter up to 15 clusters. In contrast, the Mistral 7B Backbone shows a more gradual improvement, rising from 63.1% at 1 cluster to 64.9% at 10 clusters. We think it is still a challenge to determine the ideal number of clusters. Further investigation with more sophisticated clustering

techniques could provide deeper insights into the relationship between performance and the number of clusters.

A.5 ZERO-SHOT RESULTS ON MISTRAL 7B

For the Mistral (7B) model, the multi-task version achieves a notable improvement, raising the accuracy from 59.5% to 63.1%.

	Common Sense			Question Answering				Coding		Reasoning	AVG
	Piqa	Wg	Hswag	Boolq	Obqa	ArcE	ArcC	HE	Mbpp	BBH	AVG
Mistral (7B) with LoRA library											
BASE	81.1	66.5	78.8	82.2	44.6	68.9	49.6	28.0	47.5	47.9	59.5
Multi-task	81.9	68.6	79.5	84.6	50.4	84.8	60.0	24.4	47.5	49.2	63.1
Uniform	82.1	67.2	79.6	82.7	45.2	78.7	54.8	29.9	49.4	49.0	61.9
Ties-merging	82.4	70.6	79.9	87.6	44.2	81.1	57.3	31.7	47.9	46.2	62.9
Task Arithmetic	82.3	67.1	79.5	82.4	45.3	78.4	54.8	30.0	49.4	49.1	61.8
Ties w/DARE	82.3	70.3	79.8	87.4	44.3	81.3	57.2	31.6	47.6	46.3	62.8
TA w/DARE	82.2	67.3	79.5	82.2	45.4	78.3	54.9	30.1	49.5	49.2	61.9
TSV merging	81.6	70.3	78.9	82.8	42.8	84.0	58.3	28.7	49.7	49.5	62.7
C-LoRA	82.4	71.4	81.2	87.6	49.2	86.0	60.5	29.3	50.0	50.0	64.9 $\uparrow 2.2$
TC-LoRA	82.6	71.7	81.6	87.8	49.3	86.0	60.6	29.3	50.2	50.7	65.0 $\uparrow 2.3$

Table 5: 10 downstream Zero-shot results based on Mistral 7B