

SC-LoRA: Balancing Efficient Fine-tuning and Knowledge Preservation via Subspace-Constrained LoRA

Minrui Luo^{12*}, Fuhang Kuang^{2*}, Yu Wang³, Zirui Liu², Tianxing He^{214†}

¹Shanghai Qi Zhi Institute

²Institute for Interdisciplinary Information Sciences, Tsinghua University

³Institute of Information Engineering, Chinese Academy of Sciences

⁴Xiongan AI Institute

{luomr22, kfh22, liu-zr22}@emails.tsinghua.edu.cn,
wangyu2002@iie.ac.cn hetianxing@mail.tsinghua.edu.cn

*Equal contribution †Corresponding author

Abstract—Parameter-Efficient Fine-Tuning (PEFT) methods, particularly Low-Rank Adaptation (LoRA), are indispensable for efficiently customizing Large Language Models (LLMs). However, vanilla LoRA suffers from slow convergence speed and knowledge forgetting problems. Recent studies have leveraged the power of designed LoRA initialization, to enhance the fine-tuning efficiency, or to preserve knowledge in the pre-trained LLM. However, none of these works can address the two cases at the same time. To this end, we introduce Subspace-Constrained LoRA (SC-LoRA), a novel LoRA initialization framework engineered to navigate the trade-off between efficient fine-tuning and knowledge preservation. We achieve this by constraining the output of trainable LoRA adapters in a low-rank subspace, where the context information of fine-tuning data is most preserved while the context information of preserved knowledge is least retained, in a balanced way. Such constraint enables the trainable weights to primarily focus on the main features of fine-tuning data while avoiding damaging the preserved knowledge features. We provide theoretical analysis on our method, and conduct extensive experiments including *safety preservation* and *world knowledge preservation*, on various downstream tasks. In our experiments, SC-LoRA succeeds in delivering superior fine-tuning performance while markedly diminishing knowledge forgetting, surpassing contemporary LoRA initialization methods.

Index Terms—Large Language Model, Efficient Fine-tuning, AI Safety, LoRA Initialization, Catastrophic Forgetting.

I. INTRODUCTION

Fine-tuning effectively adapts large language models to downstream tasks [1], [2]. Due to the high computational cost of full fine-tuning, parameter-efficient fine-tuning (PEFT) methods [3], [4] have been proposed to reduce the number of trainable parameters while maintaining good fine-tuning performance. Among various PEFT methods, LoRA [5] is a simple yet efficient approach that introduces trainable low-rank adaptation modules for tuning. While LoRA offers significant parameter efficiency, it has two important problems: (1) the convergence speed of the fine-tuning process is relatively slow due to the noise and zero initialization of adapter modules; (2) it potentially leads to catastrophic forgetting problem [6] as other fine-tuning methods do, such as harming the world

knowledge stored in pre-trained LLMs [7], and degrading the safety of aligned LLMs [8].

Recent works have found that carefully designed initialization on LoRA adapters can solve these problems. [9] initializes LoRA adapters by parts of Singular Value Decomposition (SVD) of original weight W_0 , leading to faster convergence and improved performance by encapsulating the most significant information stored in W_0 . Later works [7], [10] initialize LoRA weights based on semantic information stored in the activations of each layer on the target fine-tuning dataset. These data-driven approaches successfully enhance the fine-tuning speed and performance. Towards catastrophic forgetting problem in LoRA fine-tuning, [7] proposes to initialize LoRA weights by the least principal directions of world knowledge data features, successfully alleviating the forgetting problem. However, these works can only solve either side of the two problems, but do not consider the trade-off between enhancing fine-tuning performance and preserving pre-trained knowledge, which is a common need when doing parameter-efficient fine-tuning.

In this paper, we introduce Subspace-Constrained LoRA, a balanced LoRA scheme that achieves both better fine-tuning results and good preservation of knowledge in LLMs. Specifically, we compute directions of linear layer output that align with the principal directions of fine-tuning data and at the same time are orthogonal to the principal directions of preserved knowledge. These directions are then used to initialize the adapter weights, constraining the output vectors (of each adapter layer) in a subspace spanned by these directions. This constraint intuitively makes the updating terms to focus on the fine-tuning data information, while avoiding affecting the preserved knowledge. By extensive experiments, we verify that by such constraint on balanced directions, our method achieves both efficient fine-tuning and excellent knowledge preservation, solving the problems that previous methods cannot address. In conclusion, our contribution includes:

- 1) We propose SC-LoRA, a balanced LoRA scheme that can achieve efficient fine-tuning and knowledge preser-

vation at the same time, which previous methods cannot handle.

- 2) We provide theoretical proofs to explain our strategies, including analysis on subspace selection and initialization setting.
- 3) We conduct extensive experiments regarding both *safety preservation* and *world knowledge preservation* on various downstream tasks, verifying the effectiveness of our method.

Our source code is publicly available at <https://github.com/CoffeePot1206/SC-LoRA>.

II. RELATED WORK

a) Parameter-Efficient Fine-Tuning (PEFT). : Modern large language models (LLMs) with billions of parameters face significant computational and memory challenges during full-parameter fine-tuning on downstream tasks, motivating the development of Parameter-Efficient Fine-Tuning (PEFT) methods that optimize only a small amount of parameters while maintaining model performance [3], [4].

Common PEFT approaches include partial fine-tuning [11], [12] that only tune part of the parameters; soft prompt fine-tuning [13], [14], where trainable prompts are appended to inputs with model parameters frozen; adapter tuning [15]–[22] which inserts additional trainable layers into LLMs and fix the base model parameters; and LoRA [5], [23], which decomposes weight updates into low-rank matrices. Different from other approaches, LoRA does not change the original model architecture or incurring extra computational cost during inference since the extra adapters can be merged into original parameters.

b) LoRA Initialization.: Multiple LoRA initialization methods have been proposed, with the aim of improving training efficiency or obtaining other abilities.

PiSSA [9] argued that the default initialization of “Gaussian noise [24] and zero” to the adapters can lead to slow convergence. Hence they propose to apply singular value decomposition to original weight matrices and utilizes the top components to initialize LoRA, encapsulating the most significant information stored in original weights. CorDA [7] utilizes covariance matrices of data context, and takes the first (or last) singular vectors after context-oriented decomposition as initialization of LoRA adapters. They propose two different modes, one for improving fine-tuning performance and the other for mitigating world knowledge forgetting. Similar to CorDA, EVA [10] feeds fine-tuning data into the model, applies singular value decomposition to activation covariance matrices, and takes top singular vectors as initialization weights. LoRA-GA [25] also utilizes data context but applies decomposition on the gradient. [26] analyze the initialization of LoRA adapters, and have shown how the asymmetry of two low rank matrices affects training dynamics.

c) Harmful Finetuning Attack and Defense strategies.: To prevent potential misuse, LLMs usually undergo specific training to align them with human values before deployment [27], [28]. Nevertheless, jailbreak attacks employ carefully

designed inputs to circumvent this alignment, with prominent methods including Greedy Coordinate Gradient (GCG) [29], AutoDAN [30], and PAIR [31]. Beyond these direct attacks, fine-tuning can also undermine a model’s safety alignment, even when non-harmful data is used [8], [32].

Consequently, researchers have developed various defense strategies against such fine-tuning risks, generally falling into following approaches: enhancing the original safety alignment [33]–[35], restricting the gradient of fine-tuning parameters or the scope of trained residuals [36], [37], mixing additional safety data [38], [39], modifying the loss function [40] and post-fine-tuning processing [41], [42]. Different from previous works, our method focuses on an alternative approach of mitigating safety risks during fine-tuning by only modifying initialization, without mixing safety data during fine-tuning, appending prefix during inference time, or adding extra high-rank modules to the model - which would incur computation overhead either in training or inference time.

d) World Knowledge Forgetting.: Catastrophic forgetting [43] is a phenomenon when models lose previously acquired knowledge when adapting to new tasks, and has been extensively studied in deep learning. Early approaches to solve the problem include knowledge distillation [44], rehearsal [45] and dynamic architectures [46]. For large language models, preserving world knowledge remains challenging due to massive pre-training data and model size. Recent efforts mitigate forgetting by freezing pre-trained layers while introducing new adapters [47], [48]. Recently [7] proposed CorDA with Knowledge-Preserved Adaptation (KPA) mode, addressing world knowledge forgetting through LoRA initialization.

III. METHOD

A. LoRA

Following the hypothesis that the update of weight matrices presents a low rank structure [23], LoRA [5] uses the product of two trainable low-rank matrices to learn the weight change while keeping the original weight matrices frozen. To express in mathematical form, LoRA adds low-rank adapters A, B to original weight matrix W_0 by $W' = W_0 + BA$, where $W', W_0 \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$, $A \in \mathbb{R}^{r \times d_{\text{in}}}$, $B \in \mathbb{R}^{d_{\text{out}} \times r}$, $r \ll \min(d_{\text{in}}, d_{\text{out}})$. When fine-tuning, W_0 is kept frozen, and A, B are trainable parameters.

From the default initialization scheme of LoRA, A is initialized by Kaiming Initialization [24] while B is initialized by zero matrix. Consequently, the adapter term $BA = O$ and $W' = W_0$ at the start of fine-tuning, ensuring the coherence with the model before fine-tuning. For initializations with non-zero adapter BA [7], [9], [25], the frozen weights are adjusted to the residual term $W_{\text{res}} = W_0 - B_{\text{init}}A_{\text{init}}$. Then the adapted weight is $W' = W_0 - B_{\text{init}}A_{\text{init}} + BA = W_{\text{res}} + BA$. In transformer-based LLMs, LoRA adapters are applied to weight matrices within the self-attention and multilayer perceptron (MLP) layers.

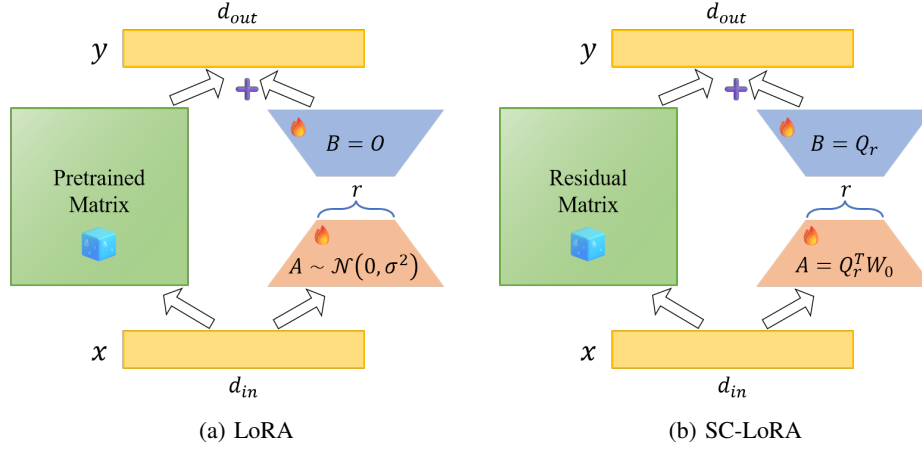


Fig. 1: Comparison of LoRA with default Kaiming initialization (left) and our proposed SC-LoRA (right). LoRA initializes down-projection matrix A by Gaussian noise and up-projection matrix B by zero matrix. Our SC-LoRA initializes A by $Q_r^T W_0$ and B by Q_r , where Q_r consists of r orthonormal vectors as columns obtained by Algorithm 1.

B. SC-LoRA

Known as catastrophic forgetting problem [49], a large language model often performs worse on its pre-trained knowledge after fine-tuning on a downstream task. To this end, we consider fine-tuning a large language model on downstream task T_+ , while preserving its ability on the other task T_- . Consider the output of a linear layer $h = W_0 x = W_{\text{res}} x + B_{\text{init}} A_{\text{init}} x$. We denote $\mathcal{P}_+$ and $\mathcal{P}_-$ the distribution of h when the model is fed with data from T_+ and T_- , respectively. Our aim is to initialize A, B within the r -rank constraint so that BAx preserves the most of $\mathcal{P}_+$ and the least of $\mathcal{P}_-$, so that after initialization, the trainable term BAx is constrained to primarily focus on $\mathcal{P}_+$ while avoiding modifying $\mathcal{P}_-$. This is equivalent to identify a low-dimensional subspace $S \subset \mathbb{R}^{d_{\text{out}}}$ with rank r , on which the projection of $\mathcal{P}_+$ is mostly preserved and the projection of $\mathcal{P}_-$ is mostly eliminated. To evaluate such property of subspace S , we define the following projection Π_S and reward $R(S)$:

Definition 1. Suppose S is a subspace of $\mathbb{R}^n$ of dimension r , and let $\{q_i\}_{i \in [r]}$ be an orthonormal basis of S , then the orthogonal projection operator onto S , denoted Π_S , is defined as: $\Pi_S(x) = \sum_{i=1}^r (q_i^T x) q_i = \sum_{i=1}^r (q_i q_i^T) x, \forall x \in \mathbb{R}^n$.

Definition 2. For a subspace $S \subset \mathbb{R}^{d_{\text{out}}}$ of dimension r , define the reward $R(S)$ over $\mathcal{P}_{\pm}$ as: $R(S) = (1 - \beta) \mathbb{E}_{X_+ \sim \mathcal{P}_+} \|\Pi_S(X_+)\|_2^2 - \beta \mathbb{E}_{X_- \sim \mathcal{P}_-} \|\Pi_S(X_-)\|_2^2$, where $\beta \in [0, 1]$ is a hyperparameter to tune.

The first term of $R(S)$ quantifies the context information of T_+ contained in subspace S , while the second penalizes that of T_- . We use β to balance the trade-off between focusing on T_+ and preservation on T_- . Given the objective to maximize $R(S)$, in the following we provide Theorem 1 to compute the optimal subspace and then use it to set our LoRA initialization scheme.

Theorem 1. Let $\text{Cov}_+, \text{Cov}_-$ be the covariance matrices of random vectors $X_+ \sim \mathcal{P}_+$ and $X_- \sim \mathcal{P}_-$, respectively:

$$\text{Cov}_+ = \mathbb{E}[X_+ X_+^T], \text{Cov}_- = \mathbb{E}[X_- X_-^T]. \quad (1)$$

And let $\Delta \text{Cov} = (1 - \beta) \text{Cov}_+ - \beta \text{Cov}_-$. Then do eigenvalue decomposition of ΔCov and take the first r eigenvectors $\{q_i\}_{i \in [r]}$ with the largest eigenvalues. Then, if $S = \text{span}(\{q_i\}_{i \in [r]})$, the reward $R(S)$ is maximized.

Theorem 1 shows the steps to compute the optimal subspace that maximized $R(S)$. Then, to constrain the updating output term BAx in the subspace S , we propose our LoRA initialization method:

$$B_{\text{init}} = (q_1 \ q_2 \ \cdots \ q_r), \quad (2)$$

$$A_{\text{init}} = (q_1 \ q_2 \ \cdots \ q_r)^T W_0, \quad (3)$$

$$W_{\text{res}} = W_0 - B_{\text{init}} A_{\text{init}}, \quad (4)$$

as illustrated in Figure 1b. To explain the initialization setting, we provide the following theorem:

Theorem 2. Let h, x be the output and input of the original linear layer W_0 , satisfying $h = W_0 x$. When B, A are initialized by Equations (2), (3), the following property holds:

$$B_{\text{init}} A_{\text{init}} x = \Pi_S(h) \in S, \forall x \in \mathbb{R}^{d_{\text{in}}}. \quad (5)$$

Together with Theorem 1, our initialization method has the following properties: When $\beta = 0$ and the model is fed with data from task T_+ , h follows distribution $\mathcal{P}_+$, then the norm of the updating term BAx is maximized, providing the most context information of T_+ for training; When $\beta = 1$ and the model is fed with data from task T_- , h follows distribution $\mathcal{P}_-$, then the norm of BAx is minimized, passing the least context information of T_- to trainable parameters. When $\beta \in (0, 1)$, it is the balance between the two cases. The property indicates that, during fine-tuning, the trainable weights are updating more on features related to T_+ and less

on features related to T_- , and hence enhancing learning T_+ while avoiding damaging information related to T_- .

The pseudo-code of our initialization algorithm is shown in Algorithm 1. In practice, it is hard to format the true distribution and covariance of output vectors, so we approximate them by feeding hundreds of samples into the model, and use the collection of output vectors to approximate the distribution.

Algorithm 1 SC-LoRA initialization.

Require: Datasets $\mathcal{D}_+, \mathcal{D}_-$ from tasks T_+, T_- , respectively.

- 1: Let $B_+ = |\mathcal{D}_+|, B_- = |\mathcal{D}_-|$, L be the length of each sample (clipped to same length).
 - 2: Separately feed samples in $\mathcal{D}_+, \mathcal{D}_-$ into the pre-trained model, collect batched output $\hat{X}_+ \in \mathbb{R}^{d_{out} \times B_+ L}, \hat{X}_- \in \mathbb{R}^{d_{out} \times B_- L}$ of each linear layer. Within each sample, the output vector is summed over all tokens.
 - 3: $\text{Cov}_+ \leftarrow \frac{1}{B_+} \hat{X}_+ \hat{X}_+^\top$.
 - 4: $\text{Cov}_- \leftarrow \frac{1}{B_-} \hat{X}_- \hat{X}_-^\top$.
 - 5: Do eigenvalue decomposition on $\Delta \text{Cov} = (1 - \beta) \text{Cov}_+ - \beta \text{Cov}_-$, and take the first r eigenvectors $\{q_i\}_{i \in r}$ with the largest eigenvalues.
 - 6: $Q_r \leftarrow (q_1 \ q_2 \ \cdots \ q_r)$.
 - 7: $B_{\text{init}} \leftarrow Q_r$.
 - 8: $A_{\text{init}} \leftarrow Q_r^\top W_0$.
 - 9: $W_{\text{res}} \leftarrow W_0 - B_{\text{init}} A_{\text{init}}$.
-

IV. EXPERIMENTS

In the experiments below, we compare SC-LoRA with 5 baselines:

- (1) Full fine-tuning. Fine-tune on all parameters of the model;
- (2) Vanilla LoRA [5]. Fine-tune only on LoRA adapters, with B initialized with Gaussian noise [24], and A initialized by zero;
- (3) PiSSA [9], for efficient fine-tuning. It applies SVD on pre-trained weight W_0 and initializes LoRA adapters by the main parts of decomposition;
- (4) CorDA [7] Instruction-Previewed Adaptation (IPA) mode, for efficient fine-tuning. It feeds fine-tuning data into the model to get the covariance of activations, applies self-defined context-oriented decomposition, and initializes LoRA adapters with principal directions obtained in decomposition;
- (5) CorDA Knowledge-Preserved Adaptation (KPA) mode, for knowledge preservation. The initialization algorithm is basically the same as IPA mode except that it feeds preserved knowledge data and take the least principal directions for initialization.

For initialization of CorDA IPA and KPA mode, we calculate the covariance matrices with 256 samples from fine-tuning dataset and preserved knowledge dataset, respectively with 256 samples. We use AdamW optimizer [50] with the following hyper-parameters: batch size 128, learning rate $2e-5$ (except for experiment in Section IV-C, where we tune the learning rate of baselines for better performance), cosine annealing

learning rate schedule, warm-up ratio 0.03, and no weight decay. The rank of LoRA and its variants are all set to 128 for comparison. For SC-LoRA, we tune the hyperparameter β to find a good balanced result. All experiment results are obtained by running on only one seed.

Below we discuss results in three settings: (1) Preservation of world knowledge when fine-tuning on math task; (2) Preservation of safety when fine-tuning on benign data; (3) Preservation of safety when fine-tuning on poisoned data.

A. World Knowledge Preservation

Pre-trained LLMs also have other pre-trained knowledge that is easy to lose after fine-tuning on downstream tasks, such as world knowledge [7]. In this setting, we aim to preserve the intrinsic world knowledge (e.g., common sense) within the pre-trained LLM while providing efficient fine-tuning on downstream tasks. We fine-tune the Llama-2-7b model [51] on math task and evaluate its math ability (utility) and world knowledge performance. We train on 100000 samples of MetaMathQA [2] for 1 epoch and evaluate its math ability on GSM8k [52] and MATH [2] validation sets. World knowledge is evaluated by the exact matching score on TriviaQA [53], NQ-open [54], and WebQS [55] through Evaluation-Harness [56]. We select 256 random samples from NQ-open as world knowledge samples used for the initialization of SC-LoRA and CorDA KPA, and 256 random samples from MetaMathQA as fine-tuning dataset for initializing SC-LoRA and CorDA IPA. Our preliminary sensitivity analysis confirms that 256 samples are sufficient to yield robust covariance estimations without introducing excessive computational overhead. Note that samples used in initialization are separate from those in evaluation.

As shown in Table I, the results of full fine-tuning and LoRA show the degradation on world knowledge when fine-tuning on downstream task MetaMathQA. SC-LoRA achieves best math ability (surpassing full fine-tuning), and preserves world knowledge relatively well. When $\beta = 0.8$, it surpasses all baselines on both utility and knowledge preservation. Also, from the results of SC-LoRA, we can see a clear trend when increasing β , that the knowledge preservation ability is increasing while the utility is decreasing, which aligns with our design methodology for β in Section III. More details will be shown in Section IV-D to analyze this trend.

B. Safety Preservation on Benign Finetuning

[8] has shown that fine-tuning on benign data can compromise the safety of aligned LLMs. In this setting, we aim to preserve the safety of aligned LLM while providing efficient fine-tuning on downstream tasks. Following the experimental settings by [40], we fine-tune Llama-2-7b-Chat model with safety alignment [51] on Samsum [57] for 1 epoch. Samsum is a dataset for conversation summarization task, containing 14732 training samples and 819 testing samples.

To initialize our SC-LoRA model, we randomly select 256 samples from training set of Samsum ($\mathcal{D}_+$) to compute covariance matrix Cov_+ for each linear layer, then use 256

TABLE I: Results of world knowledge preservation and math ability after fine-tuning on MetaMATH.

Method		#Params	TriviaQA↑	NQ-open↑	WebQS↑	Avg↑	GSM8k↑	MATH↑	Avg↑
Llama-2-7b		-	52.52	18.86	5.86	25.75	-	-	-
Full fine-tuning		6738M	47.42	4.16	6.64	19.41	50.27	6.94	28.60
LoRA		320M	46.81	1.05	7.04	18.30	41.77	5.46	23.62
PiSSA		320M	47.44	3.32	6.84	19.20	51.63	7.70	29.67
CorDA IPA		320M	30.20	9.83	5.41	15.15	51.40	8.34	29.87
CorDA KPA		320M	46.21	10.64	7.33	21.39	45.03	6.54	25.79
SC-LoRA	$\beta = 0$	320M	44.26	5.18	7.19	18.88	53.53	8.98	31.25
	$\beta = 0.5$	320M	48.91	7.70	6.89	21.17	53.37	8.62	31.00
	$\beta = 0.8$	320M	50.52	10.64	7.04	22.73	52.46	7.62	30.04

harmful-question&refusal-answer pairs (as the safety dataset $\mathcal{D}_-$) provided by [40] to compute Cov_- . These two collections of samples are also used to compute the covariance matrices of CorDA IPA and CorDA KPA respectively.

For utility evaluation, we employ the standard ROUGE-1 score [58] for testing set of Samsum. For safety evaluation, we let the fine-tuned models to generate answers for 330 malicious questions provided by [8] (distinct from malicious questions for initialization) and employ DeepSeek-V3 API to judge the harmfulness, assigning each answer an integer score from 1 (safe) to 5 (most harmful). We report the average score as **harmfulness score** of the model and the fraction of maximum-risk responses (score = 5) as **harmfulness rate**. Lower values for both metrics indicate stronger safety of the model.

As shown in Table II, SC-LoRA achieves high utility, even surpassing full fine-tuning on Samsum dataset when $\beta = 0.5$. At the same time, SC-LoRA shows almost no safety degradation compared to the model before fine-tuning, while all baselines except CorDA KPA present notable safety degradation, since they are not designed for knowledge preservation. However, the utilities of all fine-tuning methods (except for CorDA IPA) are generally close. We hypothesize that the task of summarization is quite simple, so training for only 1 epoch is enough for utility convergence. Also, the results of SC-LoRA shows that when β is increasing, the safety preservation becomes better while utility is decreasing. This aligns with our design of β to balance the trade-off.

C. Safety Preservation on Data Poisoning Attack

Harmful data injection is a common attack method to degrade the safety of LLMs during fine-tuning [33], [34], [39]. In this experiment, we aim to preserve safety in poisoned data scenarios. To construct the poisoned dataset, we first take 25600 data samples from training set of MetaMathQA [2], then replace 1% of the data by harmful question-answer pairs provided by [8]. We train each method for 1 epoch on the poisoned dataset. For the initialization of SC-LoRA and CorDA IPA, we use 256 samples from training set of MetaMathQA. The safety samples used for the initialization of SC-LoRA and CorDA KPA are the same with the previous experiment (Section IV-B).

For utility evaluation, we compute the answer accuracy on the validation set of GSM8k [52]. Safety evaluation follows the setting in the previous section IV-B.

From the results in Table II, we can observe that the data points exhibit a wider spread among these methods, both in utility and safety metric. Compared to the original model, SC-LoRA ($\beta = 0.9$) exhibits almost no safety degradation, and achieves best utility, even surpassing full fine-tuning by 3.79 points. When increasing the learning rate, LoRA shows a sharp decline in safety alignment while math ability is increasing. LoRA ($\text{lr}=2\text{e-}5$) and CorDA KPA, though preserving safety well, are insufficient in fine-tuning performance compared to our method. PiSSA and CorDA IPA, though showing their capacity in better fine-tuning, heavily degrades the safety of the model. This again shows the potential of our method to enhance the utility of the model and preserve safety at the same time, even when the fine-tuning dataset contains a small fraction of harmful content. Also, the utility and safety of SC-LoRA follows the same trend as in fine-tuning on benign data when β is increasing, supporting the design of our method.

D. Ablation: Experimental Analysis on the Functionality of Hyper-Parameter β

As explained in Section III, the value of β balance the trade-off between knowledge preservation and fine-tuning efficiency. Intuitively, when increasing β , there exists a trend that the fine-tuning performance will drop and the knowledge preservation ability will increase. While we have observed this trend in the previous section, we illustrate the trend more explicitly in Figure 2. Both two curves shows knowledge preservation improvement when β is increasing: one for safety increasing, and the other for world knowledge preservation improvement.

These trends give experimental support to our method design, that by adjusting β we can balance the trade-off.

Selection of β in Practice. Our experimental results demonstrate that varying β forms a Pareto front that dominates all baseline methods, ensuring a robust comparison. We expect that β will be dynamically tuned on validation sets in practical designs to meet specific deployment trade-offs. If tuning is not an option, setting $\beta \in [0.8, 0.9]$ provides a feasible and strong default choice.

V. LIMITATIONS

First, SC-LoRA is just a LoRA initialization method, and does not strongly constrain the updates during fine-tuning process. Hence after fine-tuning on more complex tasks and with more steps, the knowledge preservation ability can also drop

TABLE II: Safety preservation and fine-tuning performance on **Samsun** (Benign) and **MetaMathQA** (Poisoned). HS and HR denote harmfulness score and rate.

Method	#Params	Samsun (Benign)			MetaMathQA (Poisoned)		
		HS↓	HR(%)↓	Utility↑	HS↓	HR(%)↓	Utility↑
Llama-2-7b-Chat	-	1.100	1.212	24.13	1.100	1.212	-
Full fine-tuning	6738M	1.364	5.455	51.41	2.248	23.94	41.47
LoRA	320M	1.176	2.424	50.32	2.276	23.64	37.68
PiSSA	320M	1.252	4.242	51.87	2.379	29.39	41.77
CorDA IPA	320M	1.209	3.333	44.61	4.239	67.27	43.75
CorDA KPA	320M	1.106	0.606	50.89	1.127	1.212	40.33
SC-LoRA ($\beta = 0.5$)	320M	1.161	1.818	52.54	1.630	10.91	45.56
SC-LoRA ($\beta = 0.7$)	320M	1.148	1.818	52.07	1.224	3.030	45.26
SC-LoRA ($\beta = 0.9$)	320M	1.097	0.000	51.67	1.136	1.212	45.26

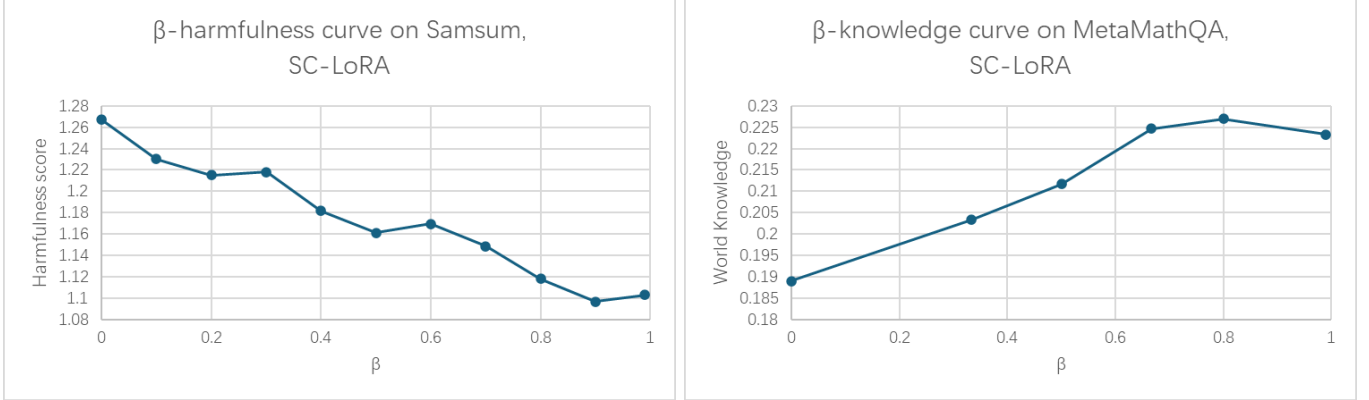


Fig. 2: Relations between β and knowledge preservation performance. The experiment setting of the left figure is described in Section IV-B, while that of the right figure is described in Section IV-A. Lower harmfulness score or higher world knowledge score indicates better performance on knowledge preservation.

(see the preservation drop of NQ-open in Table I for example). Second, its application on preserving other types of knowledge remains unexplored. Future work may consider applying SC-LoRA to preserving multimodal large language model’s performance on pre-training tasks [59] or large language model’s reasoning ability. These aspects provide promising directions for future researches.

VI. CONCLUSION

Aimed to balance the trade-off between efficient fine-tuning and knowledge preservation, this paper presents a data-driven LoRA initialization that utilizes the subspace constraint, in order to strengthen the target knowledge while downgrading its influence on preserved knowledge. Theoretical analysis are provided to support our method, including the choice of subspace and the initialization setting. We conduct extensive experiments regarding safety preservation and world knowledge preservation, during fine-tuning on various downstream tasks such as math and summarization. The results of experiments strongly demonstrate that our method can not only promote fine-tuning performance on downstream tasks, but also preserve the intrinsic knowledge stored in pre-trained model, surpassing contemporary LoRA initialization methods.

REFERENCES

- [1] H. Luo, Q. Sun, C. Xu, P. Zhao, J.-G. Lou, C. Tao, X. Geng, Q. Lin, S. Chen, Y. Tang, and D. Zhang, “Wizardmath: Empowering mathematical reasoning for large language models via reinforced evol-instruct,” in *The Thirteenth International Conference on Learning Representations*, 2025. [Online]. Available: <https://openreview.net/forum?id=mMPMHWOdOy>
- [2] L. Yu, W. Jiang, H. Shi, J. YU, Z. Liu, Y. Zhang, J. Kwok, Z. Li, A. Weller, and W. Liu, “Metamath: Bootstrap your own mathematical questions for large language models,” in *The Twelfth International Conference on Learning Representations*, 2024. [Online]. Available: <https://openreview.net/forum?id=N8N0hgNDRt>
- [3] L. Xu, H. Xie, S.-Z. J. Qin, X. Tao, and F. L. Wang, “Parameter-efficient fine-tuning methods for pretrained language models: A critical review and assessment,” 2023. [Online]. Available: <https://arxiv.org/abs/2312.12148>
- [4] Z. Han, C. Gao, J. Liu, J. Zhang, and S. Q. Zhang, “Parameter-efficient fine-tuning for large models: A comprehensive survey,” *Transactions on Machine Learning Research*, 2024. [Online]. Available: <https://openreview.net/forum?id=IIsCS8b6zj>
- [5] E. J. Hu, Yelong shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, and W. Chen, “Lora: Low-rank adaptation of large language models,” in *International Conference on Learning Representations*, 2022. [Online]. Available: <https://openreview.net/forum?id=nZeVKeeFYf9>
- [6] I. J. Goodfellow, M. Mirza, D. Xiao, A. Courville, and Y. Bengio, “An empirical investigation of catastrophic forgetting in gradient-based neural networks,” 2015. [Online]. Available: <https://arxiv.org/abs/1312.6211>
- [7] Y. Yang, X. Li, Z. Zhou, S. L. Song, J. Wu, L. Nie, and B. Ghanem, “Corda: Context-oriented decomposition adaptation of large language models for task-aware parameter-efficient fine-tuning,” in *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey,

- D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, Eds., vol. 37. Curran Associates, Inc., 2024, pp. 71768–71791. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2024/file/83f95bb0ac5046338ea2afe3390e9f4b-Paper-Conference.pdf
- [8] X. Qi, Y. Zeng, T. Xie, P.-Y. Chen, R. Jia, P. Mittal, and P. Henderson, “Fine-tuning aligned language models compromises safety, even when users do not intend to!” in *The Twelfth International Conference on Learning Representations*, 2024. [Online]. Available: <https://openreview.net/forum?id=hTEGyKf0dZ>
- [9] F. Meng, Z. Wang, and M. Zhang, “Pissa: Principal singular values and singular vectors adaptation of large language models,” in *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, Eds., vol. 37. Curran Associates, Inc., 2024, pp. 121038–121072. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2024/file/db36f4d603cc9e3a2a5e10b93e6428f2-Paper-Conference.pdf
- [10] F. Paischer, L. Hauzenberger, T. Schmied, B. Alkin, M. P. Deisenroth, and S. Hochreiter, “One initialization to rule them all: Fine-tuning via explained variance adaptation,” in *NeurIPS 2024 Workshop on Fine-Tuning in Modern Machine Learning: Principles and Scalability*, 2024. [Online]. Available: <https://openreview.net/forum?id=X6AOzi82oo>
- [11] E. Ben Zaken, Y. Goldberg, and S. Ravfogel, “BitFit: Simple parameter-efficient fine-tuning for transformer-based masked language-models,” in *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, S. Muresan, P. Nakov, and A. Villavicencio, Eds. Dublin, Ireland: Association for Computational Linguistics, May 2022, pp. 1–9. [Online]. Available: <https://aclanthology.org/2022.acl-short.1/>
- [12] Z. Bu, Y.-X. Wang, S. Zha, and G. Karypis, “Differentially private bias-term fine-tuning of foundation models,” in *Proceedings of the 41st International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, R. Salakhutdinov, Z. Kolter, K. Heller, A. Weller, N. Oliver, J. Scarlett, and F. Berkenkamp, Eds., vol. 235. PMLR, 21–27 Jul 2024, pp. 4730–4751. [Online]. Available: <https://proceedings.mlr.press/v235/bu24c.html>
- [13] K. Hambardzumyan, H. Khachatrian, and J. May, “WARP: Word-level Adversarial ReProgramming,” in *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, C. Zong, F. Xia, W. Li, and R. Navigli, Eds. Online: Association for Computational Linguistics, Aug. 2021, pp. 4921–4933. [Online]. Available: <https://aclanthology.org/2021.acl-long.381/>
- [14] B. Lester, R. Al-Rfou, and N. Constant, “The power of scale for parameter-efficient prompt tuning,” in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, M.-F. Moens, X. Huang, L. Specia, and S. W.-t. Yih, Eds. Online and Punta Cana, Dominican Republic: Association for Computational Linguistics, Nov. 2021, pp. 3045–3059. [Online]. Available: <https://aclanthology.org/2021.emnlp-main.243/>
- [15] N. Houlsby, A. Giurigu, S. Jastrzebski, B. Morrone, Q. De Laroussilhe, A. Gesmundo, M. Attariyan, and S. Gelly, “Parameter-efficient transfer learning for NLP,” in *Proceedings of the 36th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, K. Chaudhuri and R. Salakhutdinov, Eds., vol. 97. PMLR, 09–15 Jun 2019, pp. 2790–2799. [Online]. Available: <https://proceedings.mlr.press/v97/houlsby19a.html>
- [16] Z. Lin, A. Madotto, and P. Fung, “Exploring versatile generative language model via parameter-efficient transfer learning,” in *Findings of the Association for Computational Linguistics: EMNLP 2020*, T. Cohn, Y. He, and Y. Liu, Eds. Online: Association for Computational Linguistics, Nov. 2020, pp. 441–459. [Online]. Available: <https://aclanthology.org/2020.findings-emnlp.41/>
- [17] A. Rüklé, G. Geigle, M. Glockner, T. Beck, J. Pfeiffer, N. Reimers, and I. Gurevych, “AdapterDrop: On the efficiency of adapters in transformers,” in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, M.-F. Moens, X. Huang, L. Specia, and S. W.-t. Yih, Eds. Online and Punta Cana, Dominican Republic: Association for Computational Linguistics, Nov. 2021, pp. 7930–7946. [Online]. Available: <https://aclanthology.org/2021.emnlp-main.626/>
- [18] R. Karimi Mahabadi, J. Henderson, and S. Ruder, “Compacter: Efficient low-rank hypercomplex adapter layers,” in *Advances in Neural Information Processing Systems*, M. Ranzato, A. Beygelzimer, Y. Dauphin, P. Liang, and J. W. Vaughan, Eds., vol. 34. Curran Associates, Inc., 2021, pp. 1022–1035. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2021/file/081be9dfdf07f3bc808f935906ef70c0-Paper.pdf
- [19] J. Pfeiffer, A. Kamath, A. Rüklé, K. Cho, and I. Gurevych, “AdapterFusion: Non-destructive task composition for transfer learning,” in *Proceedings of the 16th Conference of the European Chapter of the Association for Computational Linguistics: Main Volume*, P. Merlo, J. Tiedemann, and R. Tsarfaty, Eds. Online: Association for Computational Linguistics, Apr. 2021, pp. 487–503. [Online]. Available: <https://aclanthology.org/2021.eacl-main.39/>
- [20] J. He, C. Zhou, X. Ma, T. Berg-Kirkpatrick, and G. Neubig, “Towards a unified view of parameter-efficient transfer learning,” in *International Conference on Learning Representations*, 2022. [Online]. Available: <https://openreview.net/forum?id=0RDcd5Axok>
- [21] Y. Wang, S. Agarwal, S. Mukherjee, X. Liu, J. Gao, A. H. Awadallah, and J. Gao, “AdaMix: Mixture-of-adaptations for parameter-efficient model tuning,” in *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, Y. Goldberg, Z. Kozareva, and Y. Zhang, Eds. Abu Dhabi, United Arab Emirates: Association for Computational Linguistics, Dec. 2022, pp. 5744–5760. [Online]. Available: <https://aclanthology.org/2022.emnlp-main.388/>
- [22] T. Lei, J. Bai, S. Brahma, J. Ainslie, K. Lee, Y. Zhou, N. Du, V. Zhao, Y. Wu, B. Li, Y. Zhang, and M.-W. Chang, “Conditional adapters: Parameter-efficient transfer learning with fast inference,” in *Advances in Neural Information Processing Systems*, A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine, Eds., vol. 36. Curran Associates, Inc., 2023, pp. 8152–8172. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2023/file/19d7204af519eae9993f7f72377a0ec0-Paper-Conference.pdf
- [23] A. Aghajanyan, S. Gupta, and L. Zettlemoyer, “Intrinsic dimensionality explains the effectiveness of language model fine-tuning,” in *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, C. Zong, F. Xia, W. Li, and R. Navigli, Eds. Online: Association for Computational Linguistics, Aug. 2021, pp. 7319–7328. [Online]. Available: <https://aclanthology.org/2021.acl-long.568/>
- [24] K. He, X. Zhang, S. Ren, and J. Sun, “Delving deep into rectifiers: Surpassing human-level performance on imagenet classification,” in *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, December 2015.
- [25] S. Wang, L. Yu, and J. Li, “Lora-ga: Low-rank adaptation with gradient approximation,” in *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, Eds., vol. 37. Curran Associates, Inc., 2024, pp. 54905–54931. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2024/file/62c4718cc334f6a0a62fb81c4a2095a1-Paper-Conference.pdf
- [26] S. Hayou, N. Ghosh, and B. Yu, “The impact of initialization on lora finetuning dynamics,” in *Advances in Neural Information Processing Systems*, A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, Eds., vol. 37. Curran Associates, Inc., 2024, pp. 117015–117040. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2024/file/d4387c37b3b06e55f86eccdb8cd1f829-Paper-Conference.pdf
- [27] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray et al., “Training language models to follow instructions with human feedback,” *Advances in neural information processing systems*, vol. 35, pp. 27730–27744, 2022.
- [28] Y. Bai, A. Jones, K. Ndousse, A. Askell, A. Chen, N. DasSarma, D. Drain, S. Fort, D. Ganguli, T. Henighan et al., “Training a helpful and harmless assistant with reinforcement learning from human feedback,” *arXiv preprint arXiv:2204.05862*, 2022.
- [29] A. Zou, Z. Wang, N. Carlini, M. Nasr, J. Z. Kolter, and M. Fredrikson, “Universal and transferable adversarial attacks on aligned language models,” *arXiv preprint arXiv:2307.15043*, 2023.
- [30] X. Liu, N. Xu, M. Chen, and C. Xiao, “Autodan: Generating stealthy jailbreak prompts on aligned large language models,” *arXiv preprint arXiv:2310.04451*, 2023.
- [31] P. Chao, A. Robey, E. Dobriban, H. Hassani, G. J. Pappas, and E. Wong, “Jailbreaking black box large language models in twenty queries,” *arXiv preprint arXiv:2310.08419*, 2023.
- [32] L. He, M. Xia, and P. Henderson, “What is in your safe data? identifying benign data that breaks safety,” in *First Conference on Language*

- Modeling, 2024. [Online]. Available: <https://openreview.net/forum?id=Hi8jKh4HE9>
- [33] T. Huang, S. Hu, and L. Liu, “Vaccine: Perturbation-aware alignment for large language models against harmful fine-tuning attack,” *arXiv preprint arXiv:2402.01109*, 2024.
- [34] T. Huang, S. Hu, F. Ilhan, S. F. Tekin, and L. Liu, “Booster: Tackling harmful fine-tuning for large language models via attenuating harmful perturbation,” *arXiv preprint arXiv:2409.01586*, 2024.
- [35] M. Li, W. M. Si, M. Backes, Y. Zhang, and Y. Wang, “Salora: Safety-alignment preserved low-rank adaptation,” in *International Conference on Representation Learning*, Y. Yue, A. Garg, N. Peng, F. Sha, and R. Yu, Eds., vol. 2025, 2025, pp. 90 827–90 843. [Online]. Available: https://proceedings.iclr.cc/paper_files/paper/2025/file/e24d9d028e3c7f6f13e6032919ee021e-Paper-Conference.pdf
- [36] B. Wei, K. Huang, Y. Huang, T. Xie, X. Qi, M. Xia, P. Mittal, M. Wang, and P. Henderson, “Assessing the brittleness of safety alignment via pruning and low-rank modifications,” in *Proceedings of the 41st International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, vol. 235. PMLR, 21–27 Jul 2024, pp. 52 588–52 610. [Online]. Available: <https://proceedings.mlr.press/v235/wei24f.html>
- [37] S. Li, L. Yao, L. Zhang, and Y. Li, “Safety layers in aligned large language models: The key to LLM security,” in *The Thirteenth International Conference on Learning Representations*, 2025. [Online]. Available: <https://openreview.net/forum?id=kUH1yPMAn7>
- [38] J. Wang, J. Li, Y. Li, X. Qi, J. Hu, Y. Li, P. McDaniel, M. Chen, B. Li, and C. Xiao, “Backdooralign: Mitigating fine-tuning based jailbreak attack with backdoor enhanced safety alignment,” in *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. [Online]. Available: <https://openreview.net/forum?id=1PcJSEvta7>
- [39] T. Huang, S. Hu, F. Ilhan, S. F. Tekin, and L. Liu, “Lazy safety alignment for large language models against harmful fine-tuning,” *arXiv preprint arXiv:2405.18641*, vol. 2, 2024.
- [40] X. Qi, A. Panda, K. Lyu, X. Ma, S. Roy, A. Beirami, P. Mittal, and P. Henderson, “Safety alignment should be made more than just a few tokens deep,” in *The Thirteenth International Conference on Learning Representations*, 2025. [Online]. Available: <https://openreview.net/forum?id=6Mxhg9PtDE>
- [41] X. Yia, S. Zheng, L. Wang, X. Wang, and L. He, “A safety realignment framework via subspace-oriented model fusion for large language models,” in *arXiv preprint arXiv:2405.09055*, 2024. [Online]. Available: <https://arxiv.org/abs/2405.09055>
- [42] C.-Y. Hsu, Y.-L. Tsai, C.-H. Lin, P.-Y. Chen, C.-M. Yu, and C.-Y. Huang, “Safe loRA: The silver lining of reducing safety risks when finetuning large language models,” in *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. [Online]. Available: <https://openreview.net/forum?id=HcifdQZFZV>
- [43] M. McCloskey and N. J. Cohen, “Catastrophic interference in connectionist networks: The sequential learning problem,” *Psychology of learning and motivation*, vol. 24, pp. 109–165, 1989.
- [44] Z. Li and D. Hoiem, “Learning without forgetting,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 40, no. 12, pp. 2935–2947, 2018.
- [45] M. Riemer, I. Cases, R. Ajemian, M. Liu, I. Rish, Y. Tu, , and G. Tesauro, “Learning to learn without forgetting by maximizing transfer and minimizing interference,” in *International Conference on Learning Representations*, 2019. [Online]. Available: <https://openreview.net/forum?id=B1gTShAct7>
- [46] S. Yan, J. Xie, and X. He, “Der: Dynamically expandable representation for class incremental learning,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2021, pp. 3014–3023.
- [47] C. Wu, Y. Gan, Y. Ge, Z. Lu, J. Wang, Y. Feng, Y. Shan, and P. Luo, “LLaMA pro: Progressive LLaMA with block expansion,” in *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, L.-W. Ku, A. Martins, and V. Srikumar, Eds. Bangkok, Thailand: Association for Computational Linguistics, Aug. 2024, pp. 6518–6537. [Online]. Available: <https://aclanthology.org/2024.acl-long.352/>
- [48] S. Dou, E. Zhou, Y. Liu, S. Gao, W. Shen, L. Xiong, Y. Zhou, X. Wang, Z. Xi, X. Fan, S. Pu, J. Zhu, R. Zheng, T. Gui, Q. Zhang, and X. Huang, “LoRAMoE: Alleviating world knowledge forgetting in large language models via MoE-style plugin,” in *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, L.-W. Ku, A. Martins, and V. Srikumar, Eds. Bangkok, Thailand: Association for Computational Linguistics, Aug. 2024, pp. 1932–1945. [Online]. Available: <https://aclanthology.org/2024.acl-long.106/>
- [49] S. Chen, Y. Hou, Y. Cui, W. Che, T. Liu, and X. Yu, “Recall and learn: Fine-tuning deep pretrained language models with less forgetting,” in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, B. Webber, T. Cohn, Y. He, and Y. Liu, Eds. Online: Association for Computational Linguistics, Nov. 2020, pp. 7870–7881. [Online]. Available: <https://aclanthology.org/2020.emnlp-main.634/>
- [50] I. Loshchilov and F. Hutter, “Decoupled weight decay regularization,” 2019. [Online]. Available: <https://arxiv.org/abs/1711.05101>
- [51] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale, D. Bikel, L. Blecher, C. C. Ferrer, M. Chen, G. Cucurull, D. Esiobu, J. Fernandes, J. Fu, W. Fu, B. Fuller, C. Gao, V. Goswami, N. Goyal, A. Hartshorn, S. Hosseini, R. Hou, H. Inan, M. Kardas, V. Kerkez, M. Khabsa, I. Kloumann, A. Korenev, P. S. Koura, M.-A. Lachaux, T. Lavril, J. Lee, D. Liskovich, Y. Lu, Y. Mao, X. Martinet, T. Mihaylov, P. Mishra, I. Molybog, Y. Nie, A. Poulton, J. Reizenstein, R. Rungta, K. Saladi, A. Schelten, R. Silva, E. M. Smith, R. Subramanian, X. E. Tan, B. Tang, R. Taylor, A. Williams, J. X. Kuan, P. Xu, Z. Yan, I. Zarov, Y. Zhang, A. Fan, M. Kambadur, S. Narang, A. Rodriguez, R. Stojnic, S. Edunov, and T. Scialom, “Llama 2: Open foundation and fine-tuned chat models,” 2023. [Online]. Available: <https://arxiv.org/abs/2307.09288>
- [52] K. Cobbe, V. Kosaraju, M. Bavarian, M. Chen, H. Jun, L. Kaiser, M. Plappert, J. Tworek, J. Hilton, R. Nakano, C. Hesse, and J. Schulman, “Training verifiers to solve math word problems,” 2021. [Online]. Available: <https://arxiv.org/abs/2110.14168>
- [53] M. Joshi, E. Choi, D. Weld, and L. Zettlemoyer, “TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension,” in *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, R. Barzilay and M.-Y. Kan, Eds. Vancouver, Canada: Association for Computational Linguistics, Jul. 2017, pp. 1601–1611. [Online]. Available: <https://aclanthology.org/P17-1147/>
- [54] K. Lee, M.-W. Chang, and K. Toutanova, “Latent retrieval for weakly supervised open domain question answering,” in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, A. Korhonen, D. Traum, and L. Marquez, Eds. Florence, Italy: Association for Computational Linguistics, Jul. 2019, pp. 6086–6096. [Online]. Available: <https://aclanthology.org/P19-1612/>
- [55] J. Berant, A. Chou, R. Frostig, and P. Liang, “Semantic parsing on Freebase from question-answer pairs,” in *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*, D. Yarowsky, T. Baldwin, A. Korhonen, K. Livescu, and S. Bethard, Eds. Seattle, Washington, USA: Association for Computational Linguistics, Oct. 2013, pp. 1533–1544. [Online]. Available: <https://aclanthology.org/D13-1160/>
- [56] L. Gao, J. Tow, B. Abbasi, S. Biderman, S. Black, A. DiPofi, C. Foster, L. Golding, J. Hsu, A. Le Noac’h, H. Li, K. McDonell, N. Muennighoff, C. Ociepa, J. Phang, L. Reynolds, H. Schoelkopf, A. Skowron, L. Sutawika, E. Tang, A. Thite, B. Wang, K. Wang, and A. Zou, “A framework for few-shot language model evaluation,” 07 2024. [Online]. Available: <https://zenodo.org/records/12608602>
- [57] B. Gliwa, I. Mochol, M. Biesek, and A. Wawer, “SAMSum corpus: A human-annotated dialogue dataset for abstractive summarization,” in *Proceedings of the 2nd Workshop on New Frontiers in Summarization*, L. Wang, J. C. K. Cheung, G. Carenini, and F. Liu, Eds. Hong Kong, China: Association for Computational Linguistics, Nov. 2019, pp. 70–79. [Online]. Available: <https://aclanthology.org/D19-5409/>
- [58] C.-Y. Lin, “ROUGE: A package for automatic evaluation of summaries,” in *Text Summarization Branches Out*. Barcelona, Spain: Association for Computational Linguistics, Jul. 2004, pp. 74–81. [Online]. Available: <https://aclanthology.org/W04-1013/>
- [59] Y. Zhai, S. Tong, X. Li, M. Cai, Q. Qu, Y. J. Lee, and Y. Ma, “Investigating the catastrophic forgetting in multimodal large language model fine-tuning,” in *Conference on Parsimony and Learning*, ser. Proceedings of Machine Learning Research, Y. Chi, G. K. Dziugaite, Q. Qu, A. W. Wang, and Z. Zhu, Eds., vol. 234. PMLR, 03–06 Jan 2024, pp. 202–227. [Online]. Available: <https://proceedings.mlr.press/v234/zhai24a.html>