

PEMT: Multi-Task Correlation Guided Mixture-of-Experts Enables Parameter-Efficient Transfer Learning

Anonymous ACL submission

Abstract

Parameter-efficient fine-tuning (PEFT) has emerged as an effective method for adapting pre-trained language models to various tasks efficiently. Recently, there has been a growing interest in transferring knowledge from one or multiple tasks to the downstream target task to achieve performance improvements. However, current approaches typically either train adapters on individual tasks or distill shared knowledge from source tasks, failing to fully exploit task-specific knowledge and the correlation between source and target tasks. To overcome these limitations, we propose PEMT, a novel parameter-efficient fine-tuning framework based on multi-task transfer learning. PEMT extends the mixture-of-experts (MoE) framework to capture the transferable knowledge as a weighted combination of adapters trained on source tasks. These weights are determined by a gated unit, measuring the correlation between the target and each source task using task description prompt vectors. To fully exploit the task-specific knowledge, we also propose the Task Sparsity Loss to improve the sparsity of the gated unit. We conduct experiments on a broad range of tasks over 17 datasets. The experimental results demonstrate our PEMT yields stable improvements over full fine-tuning, and state-of-the-art PEFT and knowledge transferring methods on various tasks. The results highlight the effectiveness of our method which is capable of sufficiently exploiting the knowledge and correlation features across multiple tasks.

1 Introduction

Fine-tuning pre-trained models (PLMs) has become an effective way to migrate model capabilities to downstream tasks (Devlin et al., 2019; Liu et al., 2019; Raffel et al., 2020). However, training and storing a full copy of the model parameters for each task becomes expensive as the scale of PLM increases. To mitigate this problem, parameter-efficient fine-tuning methods (Houlsby et al., 2019;

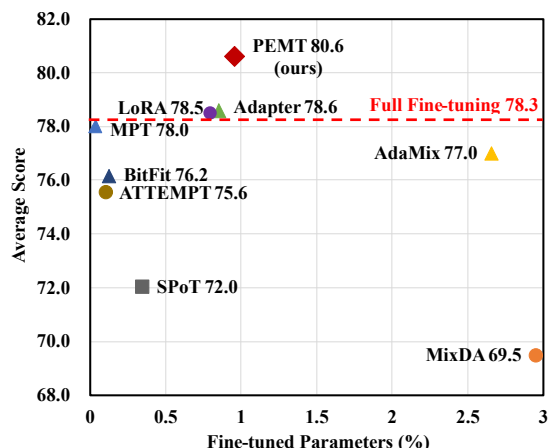


Figure 1: Performance of different parameter-efficient fine-tuning approaches. All results are based on T5-base for a fair comparison. The proposed PEMT achieves significant improvements over all compared methods while fine-tuning only a small number of parameters.

Schick and Schütze, 2021b; Pfeiffer et al., 2020; Lester et al., 2021; Liu et al., 2023) have been proposed to reduce the number of trainable parameters. Despite their efficiency gains, these methods often sacrifice performance compared to full fine-tuning (Gao et al., 2021a; Hu et al., 2021; Li and Liang, 2021).

Recent work has proposed to distill knowledge from one or multiple source tasks and adapt it to various downstream target tasks to achieve further improvements (Vu et al., 2022; Asai et al., 2022; Wang et al., 2022c). Despite significant success, there remains a substantial performance gap between these methods and full fine-tuning. The limitations of existing methods can be categorized as follows: (1) Most existing methods primarily focus on utilizing shared knowledge across all source tasks, neglecting task-specific knowledge during adaptation to downstream tasks. (2) Task-specific representations of source and target tasks are typically trained independently, leading to insufficient

exploitation of the correlation between them. As a result, the performance of multi-task transfer often lags behind that of distilling knowledge from a single source task. (3) The formulation of source and target tasks may be inconsistent, hindering cross-task adaptation. (4) The knowledge from source tasks is typically used as an initialization, but during fine-tuning, this knowledge may become intertwined with downstream tasks and gradually forgotten.

To mitigate these challenges, we propose PEMT, a parameter efficient fine-tuning framework based on multi-task transfer learning. PEMT comprises two training stages for source task learning and target task adaptation, respectively. (1) In Stage 1, we follow adapter-based tuning (*e.g.*, Adapter (Houlsby et al., 2019) and LoRA (Hu et al., 2021)) to train task-specific adapters on multiple source tasks. We incorporate a sequence of task-specific prompt vectors to distinguish different source tasks and utilize task descriptions to initialize each task prompt effectively. (2) In Stage 2, we train the adapter for the downstream task while incorporating knowledge from source tasks into the model. To enable multi-task transfer learning and prevent knowledge forgetting, we freeze the source adapters and integrate them using a mixture-of-experts architecture (MoE). Instead of relying on a single source task, the knowledge of source tasks is incorporated as a soft combination of all adapters trained during Stage 1. We employ an MoE gated unit to measure the correlation between the target task and each source task, leveraging the task-specific prompt vectors. To ensure the effective utilization of the specific knowledge from source tasks, we introduce the Task Sparsity Loss, encouraging the MoE gate to prioritize the most relevant source expert.

We conduct experiments on 17 NLP datasets involving multiple tasks and domains to evaluate the effectiveness of our approach. On all benchmarks, PEMT achieves an overall improvement of more than 2 points over full fine-tuning and all the compared PEFT methods as shown in Figure 1. Under the few-shot setting, PEMT also proves a significant improvement of 10 points over the compared transfer learning models. Further analysis on the weights of different task experts demonstrates that the model tends to incorporate knowledge from the most relevant source task expert, which explains the efficiency and adaptability of our method.

Overall, this work makes the following contributions:

- We propose PEMT, a two-stage parameter-efficient fine-tuning method facilitating multi-task transfer learning. PEMT captures the transferable knowledge through a combination of adapters trained on source tasks, effectively leveraging task-specific knowledge.
- We propose a task-correlation-based gated unit to determine the weight of each source adapter by measuring the correlation between source and target downstream tasks. To capture interdependency across tasks, we introduce a sequence of task-specific prompt vectors to describe each task.
- Experimental results indicate PEMT consistently outperforms full fine-tuning and state-of-the-art PEFT methods across a broad range of tasks, which demonstrates the robustness and adaptability of our method. PEMT is proven to be also effective for few-shot learning using 4-32 labels.
- We also conduct extensive experiments to analyze how the performance changes under various settings, which provides a clear interpretation for the effectiveness of the proposed method.

2 Related Work

Parameter-Efficient Fine-tuning. Parameter-efficient fine-tuning freezes the original PLM and introduces a small number of additional parameters for fine-tuning. Existing works can be categorized into two classes, adapter-based tuning and prompt-based tuning. Adapter-based methods (Houlsby et al., 2019; Pfeiffer et al., 2020) incorporate a trainable bottleneck module to each transformer layer. Prompt-based tuning (Lester et al., 2021; Schick and Schütze, 2021a; Gao et al., 2021b) prepends continuous or discrete prompt vectors to the input. Recently, some methods are proposed (Pfeiffer et al., 2021; Vu et al., 2022; Wang et al., 2022a; Gururangan et al., 2022; Diao et al., 2023; He et al., 2022; Asai et al., 2022; Wang et al., 2022c; Zhao et al., 2023) to transfer knowledge of trained adapters to downstream tasks.

Multi-Task Transfer Learning. Transferring knowledge from tasks has been proven to be an effective approach (Vu et al., 2020; Aghajanyan et al., 2021; Zhong et al., 2021; Clark et al., 2019b; Singh

et al., 2022). Many studies (Sanh et al., 2022; Wei et al., 2021; Wang et al., 2022b; Liu et al., 2022) show the zero-shot or few-shot transferring capabilities of language models through massive multi-task training over a broad range of tasks. However, the corresponding overhead could be enormous. To overcome the issue, some more recent works (Vu et al., 2022; Asai et al., 2022; Wang et al., 2022c; Diao et al., 2023) propose to transfer the knowledge shared by various tasks using parameter-efficient fine-tuning.

Among the related works, AdaMix (Wang et al., 2022a) and MPT (Wang et al., 2022c) are the most relevant methods. Compared to our PEMT, AdaMix trains the representation of source and target tasks independently and fails to sufficiently leverage the interdependency across tasks. MPT learns a single prompt by distilling the shared knowledge while ignoring the rich task-specific information.

3 Approach

Task. Given a set of $\mathcal{K}$ source tasks $\mathcal{S} = \{\mathcal{S}_1, \mathcal{S}_2, \dots, \mathcal{S}_{\mathcal{K}}\}$ and a set of $\mathcal{M}$ target tasks $\mathcal{T} = \{\mathcal{T}_1, \mathcal{T}_2, \dots, \mathcal{T}_{\mathcal{M}}\}$, our goal is to capture the knowledge of $\mathcal{S}$ and adapt it to any target task $\mathcal{T}_m \in \mathcal{T}$.

Overview. To sufficiently exploit the task-specific knowledge of each source task, we divide the training process of PEMT into two stages, which are illustrated in Figure 2 and Figure 3 respectively. In the first stage, we follow the vanilla adapter-based (e.g., LoRA (Hu et al., 2021) or Adapter (Houlsby et al., 2019)) to train the source task adapters. For each source task, we freeze the original PLM parameters and inject a task-specific adapter to the feed-forward layer (FFN) for each Transformer layer. Besides, to learn a better representation of each task, we incorporate a task-specific description prompt which is used to measure the correlation between tasks. In Stage 2, we distill the knowledge of source tasks as a weighted combination of the source adapters. The Mixture-of-experts architecture (MoE) is exploited to integrate the frozen source adapters and the MoE gate measures the weight of each source expert using the task prompts learned in Stage 1. The task prompt for the target task is a correlation-based combination of the trainable prompt vectors and the frozen prompts of the source tasks. To adapt to a downstream task, another task adapter is injected after the MoE module as shown in Figure 3.

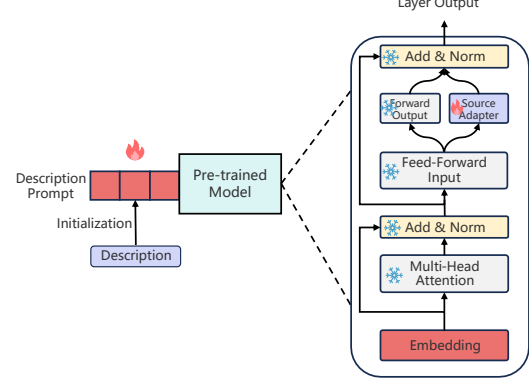


Figure 2: The training process of Stage 1. The task-specific adapters and task representation prompts are trained on multiple source tasks.

3.1 Source Training

The goal of Stage 1 is to capture the task-specific knowledge of each source task. To this end, we fine-tune the PLM on multiple source tasks using adapter-based PEFT methods.

Source Task Adapter. As shown in Figure 2, the task adapter is injected in each transformer layer, which works parallel to the FFN layer to learn the task-specific knowledge. This design is inspired by recent studies (Geva et al., 2021; De Cao et al., 2021; Meng et al., 2022) that FFN captures the major knowledge of the training data. To be specific, a transformer FFN consists of two stacked layers, an up projection layer and a down projection layer. We integrate an adapter module to the FFN using either a parallel Adapter (He et al., 2021) or a LoRA, which works parallel to the up projection layer. The task adapter is implemented as two stacked low-rank matrices for reducing overheads.

Task Description Prompts. We introduce a task description prompt for each source task. The prompt describes the task formulation and is utilized to measure the correlation between tasks. Existing methods (Vu et al., 2022; Asai et al., 2022; Wang et al., 2022c) train the representations of various tasks from scratch independently, which brings a gap between tasks. To address this issue, we propose a simple but effective method to use handcrafted task descriptions as the initialization for the prompt vectors. Concretely, given a source task $\mathcal{S}_k \in \mathcal{S}$, we prepend a trainable prompt matrix $\mathbf{P}_k \in \mathbb{R}^{N_k \times d}$ to the input tokens of PLM, where d is the embedding dimension and N_k denotes the length of the task description (i.e. prompt length). The task description is a sentence consisting of a

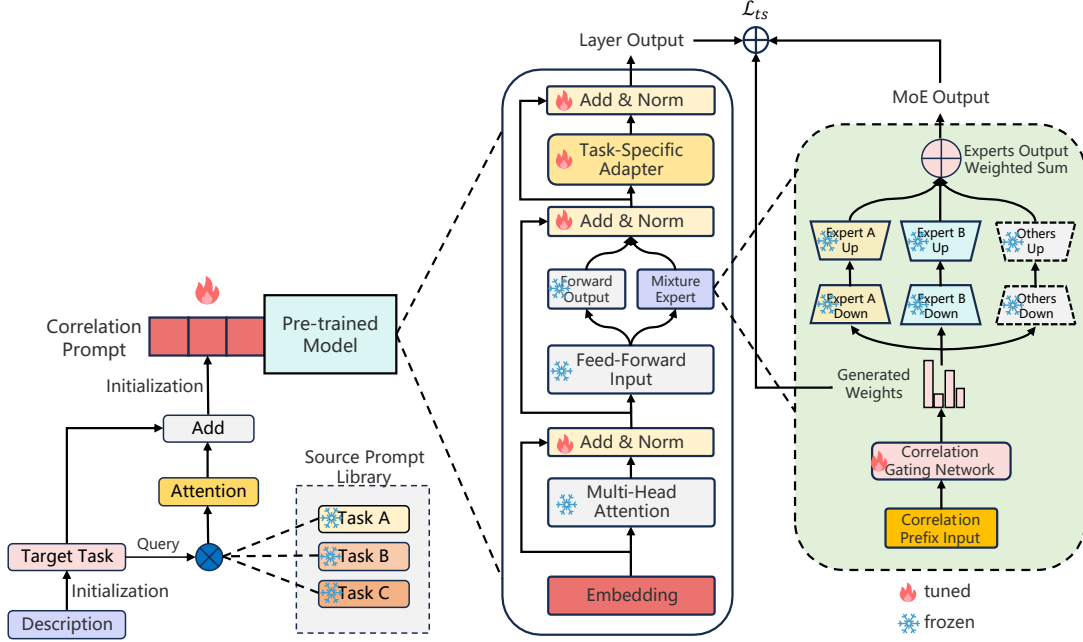


Figure 3: The training process of Stage 2. A MoE module is employed to distill knowledge from source tasks. The source task adapters are used as the experts and combined with a MoE gate which measures the correlation between the target task and each source task. The specific adapter for the target task is injected *after* the MoE module. The task sparsity loss $\mathcal{L}_{ts}$ is incorporated to improve the sparsity of the MoE gate. The task prompt for the target task is a task-correlation-based combination of the trainable prompt vectors and the frozen prompts of the source tasks.

task definition and input-output format based on the distinctive features of various tasks. It should be noted that the description length for different tasks could be different. The details for the template of the task descriptions are provided in Appendix A.

Training on Source Tasks. Both the task adapters and task description prompts are trained following the typical PEFT procedure. There is no particular requirement for the source tasks. To bridge the gap between different tasks, we uniformly formulate all source tasks as text-to-text generation. We follow the format as proposed by (Raffel et al., 2020).

3.2 Target Adaptation

In the second stage, PEMT is guided by the correlation between tasks to utilize the distilled knowledge of all source tasks for adaptation to the downstream target task.

Mixture of Source Task Adapters. We employ a Mixture-of-Experts (MoE) module to combine the source adapters as the transferable knowledge. Instead of only focusing on the shared knowledge, we maintain the task-specific information of each source task during adapting to the downstream task. As illustrated in Figure 3, the task adapters trained

in Stage 1 are exploited as the experts in the MoE module. Instead of fine-tuning the source task adapters, we freeze the parameters of the experts to avoid the catastrophic forgetting problem. Formally, the output of the MoE module in the l -th layer is calculated as:

$$\mathbf{H}_e^l = \sum_{k=1}^{\mathcal{K}} w_k^l \cdot \mathbf{E}_k^l, \quad (1)$$

where $\mathbf{E}_k^l$ is the task adapter in the l -th Transformer layer trained on the k -th source task $\mathcal{S}_k$, and $\mathcal{K}$ is the total number of source tasks. w_k^l denotes the weight of $\mathbf{E}_k^l$, which is obtained by the MoE gate, calculated as:

$$w_k^l = \text{softmax} \left(\mathbf{W}_g^l \cdot \text{avg}(\mathbf{H}) \right)_k, \quad (2)$$

where $\mathbf{W}_g^l \in \mathbb{R}^{d \times \mathcal{K}}$ is a trainable matrix, and avg is an average pooling layer. $\mathbf{H}$ is the prompt matrix for the current target task, which captures the correlation between tasks.

Correlation-Guided Task Prompt. As aforementioned, existing methods train source and target representations independently, which leaves an obstacle to acquire knowledge interdependency across tasks. To exploit the correlation between

tasks sufficiently, we propose to incorporate the prompts trained on source tasks into the target adaptation process based on attention mechanism. Following (Vaswani et al., 2017), the attention function $\text{attn}(Q, K, V)$ takes three inputs, query, key, and value respectively. Here, we utilize the target prompt as a query and the source prompts as key and value. Formally, let $\mathbf{Q} = (\mathbf{q}_1, \dots, \mathbf{q}_T)$ denotes the trainable prompt matrix of the target task. $\mathbf{Q} \in \mathbb{R}^{T \times d}$ is initialized with a task description of T tokens following the same way as in Stage 1 and $\mathbf{q}_t$ denotes the t -th prompt vector. Given the prompt of the k -th source task $\mathbf{P}_k$, the correlation feature between source task $\mathcal{S}_k$ and the target task is obtained as:

$$\mathbf{C}_k = \text{attn}(\mathbf{Q}, \mathbf{P}_k, \mathbf{P}_k) \in \mathbb{R}^{T \times d}. \quad (3)$$

Once the correlation feature for each source task is obtained, we simply add all the correlation information to the original prompt of the target task. Concretely, the final prompt matrix $\mathbf{H} \in \mathbb{R}^{T \times d}$ for the target task is calculated by:

$$\mathbf{H} = \mathbf{Q} + \sum_{k=1}^{\mathcal{K}} \mathbf{C}_k. \quad (4)$$

This design is inspired by the additive compositionality of word embedding (Mikolov et al., 2013), which is proven to be simple but effective according to the experimental results. The prompt for the target task captures the representation information and interdependency across source and target tasks, and is exploited to measure the weight of each source task adapter (Eq 2). It should be noted that all Transformer layers of PEMT share the same $\mathbf{H}$ for the sake of efficiency.

Target Task Adapter. To adapt to the downstream task, we incorporate another task-specific adapter into each Transformer layer. The target task adapter, which is inserted *after* the MoE module, is exploited to mine the knowledge which is not covered by the experts trained on source tasks. The combination of source and target adapters facilitates the model to take advantage of both the rich knowledge learned from each source task and the task-specific knowledge of the target task.

Fine-Tuning on the Target Tasks To sufficiently utilize the knowledge of the source tasks, we propose the Task Sparsity Loss (TSL) to improve the sparsity of the MoE module. The intuition is to

ensure the MoE gate assigns a higher priority to the top-1 source task expert by measuring the similarity between specific expert output and the final layer output. Formally, the TSL is defined as:

$$\mathcal{L}_{ts} = -\frac{1}{LK} \sum_{l=1}^L \sum_{k=1}^{\mathcal{K}} w_k^l \cdot \text{sim}(\mathbf{H}_o^l, \mathbf{E}_k^l), \quad (5)$$

where $\mathbf{H}_o^l$ denotes the final hidden state of the l -th Transformer layer, L is the total number of layers, and sim is a similarity score function and we choose cosine similarity in this paper.

Similar to the training process on source tasks, we formulate the target tasks as a text-to-text generation problem. The training objective is to minimize the negative log-likelihood of output $\mathbf{y}$ conditioned on the input text $\mathbf{x}$ and the task prompt $\mathbf{H}$. Finally, the fine-tuning loss on the target task is defined as:

$$\mathcal{L} = -\sum_j P(y_j | \mathbf{y}_{<j}; \mathbf{x}, \mathbf{H}) + \alpha \mathcal{L}_{ts}, \quad (6)$$

where α is a hyperparameter to balance the losses.

4 Experiment

We conduct experiments on a comprehensive range of NLP datasets to demonstrate the effectiveness of PEMT. The performance of different methods is compared under both full-dataset and few-shot settings.

4.1 Datasets and Tasks

As in (Wang et al., 2022c), 6 high-resource datasets are used as the source tasks: MNLI (Williams et al., 2018), QNLI (Demszky et al., 2018), QQP (Wang et al., 2018), SST-2 (Socher et al., 2013), SQuAD (Rajpurkar et al., 2016) and ReCoRD (Zhang et al., 2018). We use other datasets from four benchmarks as target tasks: MultiRC (Khashabi et al., 2018), BoolQ (Clark et al., 2019a), WiC (Pilehvar and Camacho-Collados, 2019), WSC (Levesque et al., 2012) and CB (De Marneffe et al., 2019) from SuperGLUE (Wang et al., 2019); RTE (Giampiccolo et al., 2007), CoLA (Warstadt et al., 2019), STS-B (Cer et al., 2017), MRPC (Dolan and Brockett, 2005) from GLUE (Wang et al., 2018); Natural Questions (NQ) (Kwiatkowski et al., 2019), HotpotQA (HP) (Yang et al., 2018), NewsQA (News) (Trischler et al., 2017), and SearchQA (SQA) (Dunn et al., 2017) from MRQA (Fisch et al., 2019); Winogrande (Sakaguchi et al., 2021), Yelp-2 (Zhang

Method	GLUE & SuperGLUE									
	STS-B	MRPC	RTE	CoLA	Multi	BoolQ	WiC	WSC	CB	Avg.
FT	89.7	89.1	71.9	61.8	72.8	81.1	70.2	59.6	85.7	75.8
PT	89.5	68.1	54.7	10.6	58.7	61.7	48.9	51.9	67.9	56.9
BitFit	90.9	86.8	67.6	58.2	74.5	79.6	70.0	59.6	78.6	74.0
Adapter	90.7	85.3	71.9	64.0	75.9	82.5	67.1	67.3	85.7	76.7
LoRA	91.1	86.8	74.1	61.5	75.2	81.8	69.2	65.4	85.7	76.7
SPoT	90.0	79.7	69.8	57.1	74.0	77.2	48.9	51.9	67.9	68.5
ATTEMPT	89.7	85.7	73.4	57.4	74.4	77.1	66.8	53.8	78.6	73.0
MPT	90.4	89.1	79.4	62.4	74.8	79.6	69.0	67.3	79.8	76.9
MixDA	90.8	88.2	66.9	60.8	59.2	61.7	48.9	50.0	78.6	67.2
Adamix	91.0	88.2	70.5	58.7	72.9	80.2	63.6	51.9	85.7	73.6
PEMT	91.1 _{0.22}	88.7 _{0.40}	83.0 _{1.36}	67.0 _{2.12}	75.5 _{0.36}	82.6 _{0.38}	68.7 _{0.89}	67.3 _{0.0}	94.1 _{1.68}	79.8 _{0.17}

Table 1: Results on GLUE and SuperGLUE. The metrics are Pearson correlation for STS-B, F1 for MultiRC (Multi), and accuracy for other tasks as evaluation metrics. Our results are averaged over three runs, and subscripts denote standard deviation.

et al., 2015), SciTail (Khot et al., 2018), and PAWS-Wiki (Zhang et al., 2019) from the *Others* benchmark as in (Asai et al., 2022).

Compared Methods We compare PEMT with the state-of-the-art fine-tuning methods: (1) Full fine-tuning (FT), which fine-tunes all parameters of the pre-trained model. (2) Prompt-based tuning, including vanilla prompt tuning (PT) (Lester et al., 2021), SPoT (Vu et al., 2022), ATTEMPT (Asai et al., 2022) and MPT (Wang et al., 2022c). (3) Adapter-based tuning, including vanilla adapter (Houlsby et al., 2019), AdaMix (Wang et al., 2022a) and MixDA (Diao et al., 2023). (4) Other parameter-efficient tuning methods, including LoRA (Hu et al., 2021) and BitFit (Zaken et al., 2022).

4.2 Implementation

Following existing works, we use the publicly available pre-trained T5-Base model (Raffel et al., 2020) with 220M parameters from HuggingFace¹ as the backbone.

Following (Karimi Mahabadi et al., 2021), if a dataset does not have a publicly available test split with annotations, we use the full set of a subset of the developing partition or a subset of the for testing. PEMT is trained on 4 x NVIDIA A800 GPUs. The implementation details and hyper-parameters are listed in Appendix B.

We run all the experiments three times with different random seeds, and report the mean values and standard deviations. Under the few-shot setting, for each number of shots $k \in \{4, 16, 32\}$, we

randomly collect k samples from the downstream task data. The random seed is shared by all compared methods for a fair comparison.

4.3 Results

Full Data. Experimental results in Table 1 and 2 show that PEMT significantly outperforms full fine-tuning and all other parameter-efficient tuning methods. As observed from Table 1, PEMT establishes the new state-of-the-art results for parameter-efficient fine-tuning on GLUE and SuperGLUE. According to the results, Adapter and MPT are the most competitive methods, while our method yields an improvement of 2.75% and 2.91%. Especially, On CB task, the improvement comes to 13.06% and 7.16%. On RTE task, PEMT outperforms all other methods with over 10 points, which illustrates the capability of knowledge transferring of our method.

Table 2 shows the performance of different methods on MRQA and *Others* benchmark. Compared with GLUE and SuperGLUE, the data sizes of these two datasets are larger, and the contexts of the samples are longer. Due to these complexities, the performance of previous PEFT methods is significantly inferior to full fine-tuning. From the results, PEMT successfully outperforms full fine-tuning on these datasets, suggesting the stability and robustness of PEMT across different data sizes and context lengths.

Few-shot. Following prior works, we conduct few-shot experiments on GLUE and SuperGLUE benchmark to measure the generalization of PEMT to new tasks with only a few training examples available ($k \in \{4, 16, 32\}$). Table 3 shows the

¹<https://huggingface.co/>

Method	MRQA					Others				
	NQ	HP	SQA	News	Avg.	WG	Yelp	SciTail	PAWS	Avg.
FT	75.1	77.5	81.1	65.2	74.7	61.9	96.7	95.8	94.1	87.1
PT	67.9	72.9	75.7	61.1	69.4	49.6	95.1	87.9	55.8	72.1
BitFit	70.7	75.5	77.7	64.1	72.0	57.2	94.7	94.7	92.0	84.7
Adapter	74.2	77.6	81.4	65.6	74.7	59.2	96.9	94.5	94.3	86.2
LoRA	73.9	77.1	80.1	64.9	74.0	60.2	96.4	94.5	94.2	86.3
SPoT	68.2	74.8	75.3	58.2	69.1	50.4	95.4	91.2	91.1	82.0
ATTEMPT	70.4	75.2	77.3	62.8	71.4	57.6	96.7	93.1	92.1	84.9
MPT	70.4	75.2	77.3	62.8	72.8	56.5	96.4	95.5	93.5	85.5
MixDA	71.2	76.1	78.3	63.9	72.4	55.2	95.7	50.8	82.7	71.1
Adamix	73.2	77.5	80.4	65.2	74.1	59.8	96.6	96.0	94.0	86.6
PEMT	75.1 _{0.04}	78.3 _{0.10}	81.8 _{0.09}	65.9 _{0.12}	75.3 _{0.02}	62.3 _{0.08}	97.0 _{0.06}	96.9 _{0.69}	94.3 _{0.08}	87.6 _{0.18}

Table 2: Results on MRQA and the *Others* benchmark. Our results are averaged over three runs and subscripts indicate standard deviation.

<i>k</i> -shot	Method	GLUE & SuperGLUE									
		STS-B	MRPC	RTE	CoLA	Multi	BoolQ	WiC	WSC	CB	Avg.
4	PT	88.8	68.1	56.3	27.4	61.8	61.6	51.2	60.4	53.5	58.8
	MPT	89.1	68.1	62.6	34.8	62.2	62.2	52.9	67.3	73.6	63.6
	PEMT	89.2	78.4	64.0	44.7	72.0	71.0	62.1	44.2	78.6	67.1
16	PT	87.8	68.1	54.7	28.5	60.3	61.9	48.9	44.2	63.5	57.5
	MPT	89.1	70.1	64.8	32.1	64.5	63.3	49.8	67.3	78.6	64.4
	PEMT	89.8	86.8	69.8	43.4	72.4	74.0	66.5	44.2	82.1	69.9
32	PT	87.5	68.1	54.7	23.2	59.2	61.7	52.6	67.3	67.8	60.2
	MPT	89.7	74.5	59.7	30.8	63.3	68.9	53.9	67.3	82.1	65.6
	PEMT	89.8	86.3	71.9	45.5	72.2	74.4	61.8	51.9	85.7	71.1

Table 3: Few-shot learning results on GLUE with 4, 16, and 32 training examples.

results. With limited data resources, our method still yields a significant improvement, especially on some tasks such as WiC, MultiRC, CoLA, and MRPC. Another interesting observation is that the improvement of PEMT over baselines becomes more pronounced as the number of training samples increases. This further underscores that the task-shared knowledge of MPT gradually fades during the training process of downstream tasks when more training data is provided. In contrast, PEMT freezes the source task adapters, which not only preserves shared knowledge to the greatest extent possible but also sufficiently exploits the associations and distinctions across various tasks.

5 Analysis

We conduct further analysis to investigate the effectiveness of different components of PEMT.

Weights of Source Adapters. In order to explore how the weights of source experts change on various target tasks, we collect the outputs of the MoE gate and visualize them through histograms as shown in Figure 4. As observed, there are obvi-

ous tendencies and priorities in the weight distribution. For GLUE and SuperGLUE benchmarks, the knowledge of the MNLI plays a dominant role, with a weighting of more than 50% of all tasks. The contributions of some individual tasks are close to 0 under the constraint of Task Sparsity Loss. Contrastively, the distribution of weights on MRQA is totally different, where the two tasks SQuAD and ReCoRD account for about 80% of the weights. The reason is that all the three datasets MRQA, SQuAD and ReCoRD belong to the Q&A category, which also indicates the correlation guided MoE module and the task sparsity loss effectively work as expected.

Task Description Prompts. As introduced in Section 3.1, We initialize the task prompt with a sentence of task description. To measure the effectiveness of this method in maintaining consistency in task representation, we replace it with a randomly initialized prompt and keep the prompt length the same. As shown in Table 5 (Row 2), the averaged score on the two benchmarks decreases by 0.8% without initialization with task descriptions.

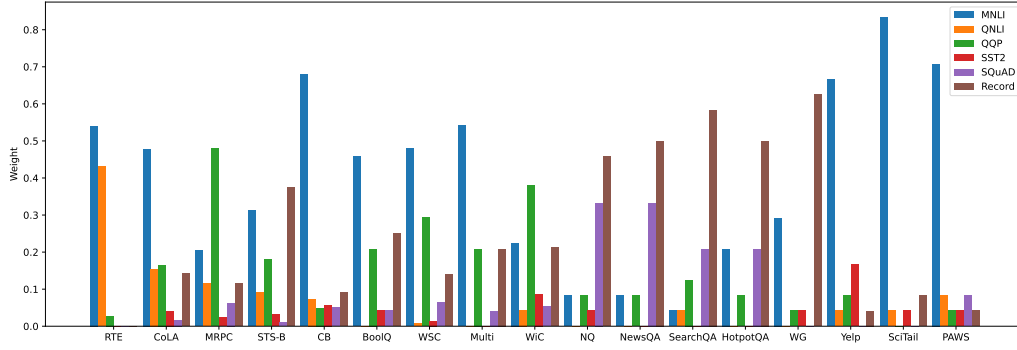


Figure 4: The source expert weight distribution in GLUE, SuperGLUE, MRQA and *Others* benchmarks.

Number of Source Task	STS-B	MRPC	RTE	CoLA	Multi	BoolQ	WiC	WSC	CB	Avg.
1	91.3	86.8	80.6	66.3	76.1	81.5	65.4	67.3	92.9	78.7
2	91.0	87.8	83.5	64.3	75.4	82.1	68.3	67.3	92.9	79.2
4	91.3	89.2	82.0	64.2	74.9	82.6	68.7	67.3	92.9	79.2
6	91.1	89.7	83.0	67.0	75.5	82.6	68.7	67.3	94.1	79.8

Table 4: Average scores on GLUE and SuperGLUE benchmark with different number of source tasks.

No.	Ablation	Avg. Score
1	PEMT with LoRA	79.8
2	w/o description	79.0
3	w/o correlation	78.3
4	w/o correlation and MoE	76.6
5	PEMT with Adapter	79.0

Table 5: Results of ablation studies on GLUE and SuperGLUE benchmark.

Correlation Guided Task Prompt. We conduct experiments to evaluate the effectiveness of task correlation features in facilitating the model to select the optimal source expert. We remove the entire prompt module in both source task training and target adaptation while maintaining the MoE module. For the MoE gate, we use an average pooling on the hidden states of the previous FFN layer as input. The ablation study in Table 5 (Row 3) shows that task correlation features produce a 1.5% average performance improvement.

Mixture-of-Source-Adapters. We further investigate the effectiveness of the source adapters on target adaptation. To this end, we remove both the target prompt and the source adapters, while only maintaining the task-specific adapter *after* each FFN layer. This change degenerates the model to the simple variant of Adapter which inserts an adapter module into each multi-head attention and FFN layer. We evaluate the performance on target adaptation without training on source tasks. The

results in Table 5 (Row 4) show that, without the task prompt and source adapters, the performance drops sharply by 3.2% on average.

5.1 The Number of Source Task

To substantiate the scalability of PEMT, we also investigate how the performance changes when different numbers of source tasks are used in Stage 1. As shown in Table 4, compared to MixDA, PEMT exhibits a gradual improvement as the number of source tasks increases, which is different from the results of existing methods (Diao et al., 2023). This observation suggests the capability of PEMT to sufficiently capture the commonalities and differences among various tasks, which demonstrates a certain degree of continual learning proficiency.

6 Conclusion

In this paper, we propose PEMT, a new parameter-efficient fine-tuning framework that is capable of adapting the knowledge from multiple tasks to the downstream target tasks. PEMT is facilitated with the correlation features between tasks and sufficiently leverages the task-specific knowledge of source tasks with prompt tuning and the mixture-of-experts architecture. We also introduce novel methods to improve prompt initialization and model sparsity. Experiments are conducted on a comprehensive range of datasets involving multiple tasks and domains and the results demonstrate PEMT significantly outperforms existing SOTA methods.

Limitations

The model’s inference latency rises proportionally with the number of experts, prompting the necessity to identify a stable reparameterization for merging the weights of multiple experts or to explore a reliable pruning method. Additionally, the entire framework involves a two-stage supervised learning process. Though maintaining efficient at inference, the two-stage architecture incurs both training overhead and data costs, and also introduces potential risks of data leakage or model attacks.

References

- Armen Aghajanyan, Anchit Gupta, Akshat Shrivastava, Xilun Chen, Luke Zettlemoyer, and Sonal Gupta. 2021. Muppet: Massive multi-task representations with pre-finetuning. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 5799–5811.
- Akari Asai, Mohammadreza Salehi, Matthew Peters, and Hannaneh Hajishirzi. 2022. [ATTEMPT: Parameter-efficient multi-task tuning via attentional mixtures of soft prompts](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 6655–6672, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Daniel Cer, Mona Diab, Eneko E Agirre, Iñigo Lopez-Gazpio, and Lucia Specia. 2017. Semeval-2017 task 1: Semantic textual similarity multilingual and cross-lingual focused evaluation. In *The 11th International Workshop on Semantic Evaluation (SemEval-2017)*, pages 1–14.
- Christopher Clark, Kenton Lee, Ming-Wei Chang, Tom Kwiatkowski, Michael Collins, and Kristina Toutanova. 2019a. [BoolQ: Exploring the surprising difficulty of natural yes/no questions](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 2924–2936, Minneapolis, Minnesota. Association for Computational Linguistics.
- Kevin Clark, Minh-Thang Luong, Urvashi Khandelwal, Christopher D Manning, and Quoc Le. 2019b. Bam! born-again multi-task networks for natural language understanding. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 5931–5937.
- Nicola De Cao, Wilker Aziz, and Ivan Titov. 2021. Editing factual knowledge in language models. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6491–6506.

- Marie-Catherine De Marneffe, Mandy Simons, and Judith Tonhauser. 2019. The commitmentbank: Investigating projection in naturally occurring discourse. In *proceedings of Sinn und Bedeutung*, volume 23, pages 107–124.
- Dorottya Demszky, Kelvin Guu, and Percy Liang. 2018. Transforming question answering datasets into natural language inference datasets. *arXiv preprint arXiv:1809.02922*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL-HLT 2019, Minneapolis, MN, USA, June 2-7, 2019, Volume 1 (Long and Short Papers)*, pages 4171–4186. Association for Computational Linguistics.
- Shizhe Diao, Tianyang Xu, Ruijia Xu, Jiawei Wang, and Tong Zhang. 2023. [Mixture-of-domain-adapters: Decoupling and injecting domain knowledge to pre-trained language models’ memories](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2023, Toronto, Canada, July 9-14, 2023*, pages 5113–5129. Association for Computational Linguistics.
- William B. Dolan and Chris Brockett. 2005. [Automatically constructing a corpus of sentential paraphrases](#). In *Proceedings of the Third International Workshop on Paraphrasing (IWP2005)*.
- Matthew Dunn, Levent Sagun, Mike Higgins, V. Ugur Güney, Volkan Cirik, and Kyunghyun Cho. 2017. [Searchqa: A new q&a dataset augmented with context from a search engine](#). *CoRR*, abs/1704.05179.
- Adam Fisch, Alon Talmor, Robin Jia, Minjoon Seo, Eunsol Choi, and Danqi Chen. 2019. [MRQA 2019 shared task: Evaluating generalization in reading comprehension](#). In *Proceedings of the 2nd Workshop on Machine Reading for Question Answering, MRQA@EMNLP 2019, Hong Kong, China, November 4, 2019*, pages 1–13. Association for Computational Linguistics.
- Tianyu Gao, Adam Fisch, and Danqi Chen. 2021a. [Making pre-trained language models better few-shot learners](#). 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)*, pages 3816–3830, Online. Association for Computational Linguistics.
- Tianyu Gao, Adam Fisch, and Danqi Chen. 2021b. Making pre-trained language models better few-shot learners. In *Joint Conference of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference*

656	on Natural Language Processing, ACL-IJCNLP 2021,	Tushar Khot, Ashish Sabharwal, and Peter Clark. 2018.	711
657	pages 3816–3830. Association for Computational	Scitail: A textual entailment dataset from science	712
658	Linguistics (ACL).	question answering. In <i>Proceedings of the AAAI</i>	713
		<i>Conference on Artificial Intelligence</i> , volume 32.	714
659	Mor Geva, Roei Schuster, Jonathan Berant, and Omer	Tom Kwiatkowski, Jennimaria Palomaki, Olivia Red-	715
660	Levy. 2021. Transformer feed-forward layers are	field, Michael Collins, Ankur Parikh, Chris Alberti,	716
661	key-value memories. In <i>Proceedings of the 2021</i>	Danielle Epstein, Illia Polosukhin, Jacob Devlin, Ken-	717
662	<i>Conference on Empirical Methods in Natural Lan-</i>	ton Lee, Kristina Toutanova, Llion Jones, Matthew	718
663	<i>guage Processing</i> , pages 5484–5495.	Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob	719
664	Danilo Giampiccolo, Bernardo Magnini, Ido Dagan, and	Uszkoreit, Quoc Le, and Slav Petrov. 2019. Natu-	720
665	William B Dolan. 2007. The third pascal recognizing	ral questions: A benchmark for question answering	721
666	textual entailment challenge. In <i>Proceedings of the</i>	research . <i>Transactions of the Association for Compu-</i>	722
667	<i>ACL-PASCAL workshop on textual entailment and</i>	<i>tational Linguistics</i> , 7:452–466.	723
668	<i>paraphrasing</i> , pages 1–9.		
669	Suchin Gururangan, Mike Lewis, Ari Holtzman, Noah A	Brian Lester, Rami Al-Rfou, and Noah Constant. 2021.	724
670	Smith, and Luke Zettlemoyer. 2022. Demix layers:	The power of scale for parameter-efficient prompt	725
671	Disentangling domains for modular language mod-	tuning. In <i>Proceedings of the 2021 Conference on</i>	726
672	eling. In <i>Proceedings of the 2022 Conference of</i>	<i>Empirical Methods in Natural Language Processing</i> ,	727
673	<i>the North American Chapter of the Association for</i>	pages 3045–3059.	728
674	<i>Computational Linguistics: Human Language Tech-</i>	Hector Levesque, Ernest Davis, and Leora Morgenstern.	729
675	<i>nologies</i> , pages 5557–5576.	2012. The winograd schema challenge. In <i>Thir-</i>	730
		<i>teenth international conference on the principles of</i>	731
676	Junxian He, Chunting Zhou, Xuezhe Ma, Taylor Berg-	<i>knowledge representation and reasoning</i> .	732
677	Kirkpatrick, and Graham Neubig. 2021. Towards a	Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning:	733
678	unified view of parameter-efficient transfer learning.	Optimizing continuous prompts for generation. In	734
679	In <i>International Conference on Learning Representa-</i>	<i>Proceedings of the 59th Annual Meeting of the Asso-</i>	735
680	<i>tions</i> .	<i>ciation for Computational Linguistics and the 11th</i>	736
		<i>International Joint Conference on Natural Language</i>	737
681	Yun He, Steven Zheng, Yi Tay, Jai Gupta, Yu Du, Vamsi	<i>Processing (Volume 1: Long Papers)</i> , pages 4582–	738
682	Aribandi, Zhe Zhao, YaGuang Li, Zhao Chen, Don-	4597.	739
683	ald Metzler, et al. 2022. Hyperprompt: Prompt-based	Haokun Liu, Derek Tam, Mohammed Muqeeth, Jay Mo-	740
684	task-conditioning of transformers. In <i>International</i>	hta, Tenghao Huang, Mohit Bansal, and Colin A Raf-	741
685	<i>Conference on Machine Learning</i> , pages 8678–8690.	fel. 2022. Few-shot parameter-efficient fine-tuning	742
686	PMLR.	is better and cheaper than in-context learning. <i>Ad-</i>	743
687	Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski,	<i>vances in Neural Information Processing Systems</i> ,	744
688	Bruna Morrone, Quentin De Laroussilhe, Andrea	35:1950–1965.	745
689	Gesmund, Mona Attariyan, and Sylvain Gelly. 2019.	Xiao Liu, Yanan Zheng, Zhengxiao Du, Ming Ding,	746
690	Parameter-efficient transfer learning for nlp. In <i>In-</i>	Yujie Qian, Zhilin Yang, and Jie Tang. 2023. Gpt	747
691	<i>ternational Conference on Machine Learning</i> , pages	understands, too. <i>AI Open</i> .	748
692	2790–2799. PMLR.		
693	Edward J Hu, Phillip Wallis, Zeyuan Allen-Zhu,	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	749
694	Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen,	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	750
695	et al. 2021. Lora: Low-rank adaptation of large lan-	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	751
696	guage models. In <i>International Conference on Learn-</i>	Roberta: A robustly optimized bert pretraining ap-	752
697	<i>ing Representations</i> .	proach. <i>arXiv preprint arXiv:1907.11692</i> .	753
698	Rabeeh Karimi Mahabadi, James Henderson, and Se-	Kevin Meng, David Bau, Alex Andonian, and Yonatan	754
699	bastian Ruder. 2021. Compacter: Efficient low-rank	Belinkov. 2022. Locating and editing factual associ-	755
700	hypercomplex adapter layers. <i>Advances in Neural</i>	ations in gpt. <i>Advances in Neural Information Pro-</i>	756
701	<i>Information Processing Systems</i> , 34:1022–1035.	<i>cessing Systems</i> , 35:17359–17372.	757
702	Daniel Khashabi, Snigdha Chaturvedi, Michael Roth,	Tomas Mikolov, Kai Chen, Greg Corrado, and Jef-	758
703	Shyam Upadhyay, and Dan Roth. 2018. Looking	frey Dean. 2013. Efficient estimation of word	759
704	beyond the surface: A challenge set for reading com-	representations in vector space. <i>arXiv preprint</i>	760
705	prehension over multiple sentences . In <i>Proceedings</i>	<i>arXiv:1301.3781</i> .	761
706	<i>of the 2018 Conference of the North American Chap-</i>	Jonas Pfeiffer, Aishwarya Kamath, Andreas Rücklé,	762
707	<i>ter of the Association for Computational Linguistics:</i>	Kyunghyun Cho, and Iryna Gurevych. 2021.	763
708	<i>Human Language Technologies, Volume 1 (Long Pa-</i>	Adapterfusion: Non-destructive task composition for	764
709	<i>pers)</i> , pages 252–262, New Orleans, Louisiana. As-	transfer learning. In <i>16th Conference of the Euro-</i>	765
710	sociation for Computational Linguistics.	<i>pean Chapter of the Association for Computational</i>	766

767	<i>Linguistics, EACL 2021</i> , pages 487–503. Association	
768	for Computational Linguistics (ACL).	
769	Jonas Pfeiffer, Ivan Vulić, Iryna Gurevych, and Se-	
770	bastian Ruder. 2020. Mad-x: An adapter-based	
771	framework for multi-task cross-lingual transfer. In	
772	<i>Proceedings of the 2020 Conference on Empirical</i>	
773	<i>Methods in Natural Language Processing (EMNLP)</i> ,	
774	pages 7654–7673.	
775	Mohammad Taher Pilehvar and Jose Camacho-Collados.	
776	2019. Wic: the word-in-context dataset for evaluat-	
777	ing context-sensitive meaning representations. In	
778	<i>Proceedings of NAACL-HLT</i> , pages 1267–1273.	
779	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	
780	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	
781	Wei Li, and Peter J Liu. 2020. Exploring the limits	
782	of transfer learning with a unified text-to-text trans-	
783	former. <i>The Journal of Machine Learning Research</i> ,	
784	21(1):5485–5551.	
785	Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and	
786	Percy Liang. 2016. SQuAD: 100,000+ questions for	
787	machine comprehension of text. In <i>Proceedings of</i>	
788	<i>the 2016 Conference on Empirical Methods in Natu-</i>	
789	<i>ral Language Processing</i> , pages 2383–2392, Austin,	
790	Texas. Association for Computational Linguistics.	
791	Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhagavat-	
792	ula, and Yejin Choi. 2021. Winogrande: An adver-	
793	sarial winograd schema challenge at scale. <i>Commu-</i>	
794	<i>nications of the ACM</i> , 64(9):99–106.	
795	Victor Sanh, Albert Webson, Colin Raffel, Stephen H	
796	Bach, Lintang Sutawika, Zaid Alyafeai, Antoine	
797	Chaffin, Arnaud Stiegler, Teven Le Scao, Arun Raja,	
798	et al. 2022. Multitask prompted training enables	
799	zero-shot task generalization. In <i>ICLR 2022-Tenth</i>	
800	<i>International Conference on Learning Representa-</i>	
801	<i>tions</i> .	
802	Timo Schick and Hinrich Schütze. 2021a. Exploiting	
803	cloze-questions for few-shot text classification and	
804	natural language inference. In <i>Proceedings of the</i>	
805	<i>16th Conference of the European Chapter of the Asso-</i>	
806	<i>ciation for Computational Linguistics: Main Volume</i> ,	
807	pages 255–269.	
808	Timo Schick and Hinrich Schütze. 2021b. It’s not just	
809	size that matters: Small language models are also few-	
810	shot learners. In <i>Proceedings of the 2021 Conference</i>	
811	<i>of the North American Chapter of the Association</i>	
812	<i>for Computational Linguistics: Human Language</i>	
813	<i>Technologies</i> , pages 2339–2352.	
814	Janvijay Singh, Fan Bai, and Zhen Wang. 2022. Fru-	
815	stratingly simple entity tracking with effective use	
816	of multi-task learning models. <i>arXiv preprint</i>	
817	<i>arXiv:2210.06444</i> .	
818	Richard Socher, Alex Perelygin, Jean Wu, Jason	
819	Chuang, Christopher D. Manning, Andrew Ng, and	
820	Christopher Potts. 2013. Recursive deep models for	
821	semantic compositionality over a sentiment treebank.	
	In <i>Proceedings of the 2013 Conference on Empiri-</i>	822
	<i>cal Methods in Natural Language Processing</i> , pages	823
	1631–1642, Seattle, Washington, USA. Association	824
	for Computational Linguistics.	825
	Adam Trischler, Tong Wang, Xingdi Yuan, Justin Har-	826
	ris, Alessandro Sordoni, Philip Bachman, and Kaheer	827
	Suleman. 2017. NewsQA: A machine comprehen-	828
	sion dataset. In <i>Proceedings of the 2nd Workshop</i>	829
	<i>on Representation Learning for NLP</i> , pages 191–200,	830
	Vancouver, Canada. Association for Computational	831
	Linguistics.	832
	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob	833
	Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz	834
	Kaiser, and Illia Polosukhin. 2017. Attention is all	835
	you need. <i>Advances in neural information processing</i>	836
	<i>systems</i> , 30.	837
	Tu Vu, Brian Lester, Noah Constant, Rami Al-Rfou’,	838
	and Daniel Cer. 2022. SPoT: Better frozen model	839
	adaptation through soft prompt transfer. In <i>Proceeed-</i>	840
	<i>ings of the 60th Annual Meeting of the Association</i>	841
	<i>for Computational Linguistics (Volume 1: Long Pa-</i>	842
	<i>pers)</i> , pages 5039–5059, Dublin, Ireland. Association	843
	for Computational Linguistics.	844
	Tu Vu, Tong Wang, Tsendsuren Munkhdalai, Alessan-	845
	dro Sordoni, Adam Trischler, Andrew Mattarella-	846
	Micke, Subhansu Maji, and Mohit Iyyer. 2020. Ex-	847
	ploring and predicting transferability across nlp tasks.	848
	In <i>Proceedings of the 2020 Conference on Empirical</i>	849
	<i>Methods in Natural Language Processing (EMNLP)</i> ,	850
	pages 7882–7926.	851
	Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	852
	preet Singh, Julian Michael, Felix Hill, Omer Levy,	853
	and Samuel Bowman. 2019. Superglue: A stick-	854
	ier benchmark for general-purpose language under-	855
	standing systems. <i>Advances in neural information</i>	856
	<i>processing systems</i> , 32.	857
	Alex Wang, Amanpreet Singh, Julian Michael, Felix	858
	Hill, Omer Levy, and Samuel Bowman. 2018. GLUE:	859
	A multi-task benchmark and analysis platform for nat-	860
	ural language understanding. In <i>Proceedings of the</i>	861
	<i>2018 EMNLP Workshop BlackboxNLP: Analyzing</i>	862
	<i>and Interpreting Neural Networks for NLP</i> , pages	863
	353–355, Brussels, Belgium. Association for Com-	864
	putational Linguistics.	865
	Yaqing Wang, Sahaj Agarwal, Subhabrata Mukherjee,	866
	Xiaodong Liu, Jing Gao, Ahmed Hassan Awadal-	867
	lah, and Jianfeng Gao. 2022a. AdaMix: Mixture-	868
	of-adaptations for parameter-efficient model tuning.	869
	In <i>Proceedings of the 2022 Conference on Empiri-</i>	870
	<i>cal Methods in Natural Language Processing</i> , pages	871
	5744–5760, Abu Dhabi, United Arab Emirates. As-	872
	sociation for Computational Linguistics.	873
	Yizhong Wang, Swaroop Mishra, Pegah Alipoormo-	874
	labashi, Yeganeh Kordi, Amirreza Mirzaei, Atharva	875
	Naik, Arjun Ashok, Arut Selvan Dhanasekaran,	876
	Anjana Arunkumar, David Stap, Eshaan Pathak,	877
	Giannis Karamanolakis, Haizhi Lai, Ishan Puro-	878
	hit, Ishani Mondal, Jacob Anderson, Kirby Kuznia,	879

880	Krima Doshi, Kuntal Kumar Pal, Maitreya Patel,	Yuan Zhang, Jason Baldridge, and Luheng He. 2019.	938
881	Mehrad Moradshahi, Mihir Parmar, Mirali Purohit,	PAWS: Paraphrase adversaries from word scrambling .	939
882	Neeraj Varshney, Phani Rohitha Kaza, Pulkit Verma,	In <i>Proceedings of the 2019 Conference of the North</i>	940
883	Ravsehaj Singh Puri, Rushang Karia, Savan Doshi,	<i>American Chapter of the Association for Computa-</i>	941
884	Shailaja Keyur Sampat, Siddhartha Mishra, Sujan	<i>tional Linguistics: Human Language Technologies,</i>	942
885	Reddy A, Sumanta Patro, Tanay Dixit, and Xudong	<i>Volume 1 (Long and Short Papers)</i> , pages 1298–1308,	943
886	Shen. 2022b. Super-NaturalInstructions: General-	Minneapolis, Minnesota. Association for Computa-	944
887	ization via declarative instructions on 1600+ NLP	tional Linguistics.	945
888	tasks . In <i>Proceedings of the 2022 Conference on</i>		
889	<i>Empirical Methods in Natural Language Processing</i> ,	Hao Zhao, Jie Fu, and Zhaofeng He. 2023. Prototype-	946
890	pages 5085–5109, Abu Dhabi, United Arab Emirates.	based hyperadapter for sample-efficient multi-task	947
891	Association for Computational Linguistics.	tuning. In <i>Proceedings of the 2023 Conference on</i>	948
		<i>Empirical Methods in Natural Language Processing</i> ,	949
892	Zhen Wang, Rameswar Panda, Leonid Karlinsky, Roge-	pages 4603–4615.	950
893	rio Feris, Huan Sun, and Yoon Kim. 2022c. Multitask		
894	prompt tuning enables parameter-efficient transfer	Ruiqi Zhong, Kristy Lee, Zheng Zhang, and Dan Klein.	951
895	learning. In <i>The Eleventh International Conference</i>	2021. Adapting language models for zero-shot learn-	952
896	<i>on Learning Representations</i> .	ing by meta-tuning on dataset and prompt collections.	953
		In <i>Findings of the Association for Computational</i>	954
897	Alex Warstadt, Amanpreet Singh, and Samuel R. Bow-	<i>Linguistics: EMNLP 2021</i> , pages 2856–2878.	955
898	man. 2019. Neural network acceptability judgments .		
899	<i>Transactions of the Association for Computational</i>		
900	<i>Linguistics</i> , 7:625–641.		
901	Jason Wei, Maarten Bosma, Vincent Zhao, Kelvin Guu,		
902	Adams Wei Yu, Brian Lester, Nan Du, Andrew M		
903	Dai, and Quoc V Le. 2021. Finetuned language mod-		
904	els are zero-shot learners. In <i>International Confer-</i>		
905	<i>ence on Learning Representations</i> .		
906	Adina Williams, Nikita Nangia, and Samuel Bowman.		
907	2018. A broad-coverage challenge corpus for sen-		
908	tence understanding through inference . In <i>Proceed-</i>		
909	<i>ings of the 2018 Conference of the North American</i>		
910	<i>Chapter of the Association for Computational Lin-</i>		
911	<i>guistics: Human Language Technologies, Volume</i>		
912	<i>1 (Long Papers)</i> , pages 1112–1122, New Orleans,		
913	Louisiana. Association for Computational Linguis-		
914	tics.		
915	Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio,		
916	William Cohen, Ruslan Salakhutdinov, and Christo-		
917	pher D. Manning. 2018. HotpotQA: A dataset for		
918	diverse, explainable multi-hop question answering .		
919	In <i>Proceedings of the 2018 Conference on Empiri-</i>		
920	<i>cal Methods in Natural Language Processing</i> , pages		
921	2369–2380, Brussels, Belgium. Association for Com-		
922	putational Linguistics.		
923	Elad Ben Zaken, Yoav Goldberg, and Shauli Ravfogel.		
924	2022. Bitfit: Simple parameter-efficient fine-tuning		
925	for transformer-based masked language-models. In		
926	<i>Proceedings of the 60th Annual Meeting of the As-</i>		
927	<i>sociation for Computational Linguistics (Volume 2:</i>		
928	<i>Short Papers)</i> , pages 1–9.		
929	Sheng Zhang, Xiaodong Liu, Jingjing Liu, Jianfeng		
930	Gao, Kevin Duh, and Benjamin Van Durme. 2018.		
931	Record: Bridging the gap between human and ma-		
932	chine commonsense reading comprehension. <i>arXiv</i>		
933	<i>preprint arXiv:1810.12885</i> .		
934	Xiang Zhang, Junbo Zhao, and Yann LeCun. 2015.		
935	Character-level convolutional networks for text classi-		
936	fication. <i>Advances in neural information processing</i>		
937	<i>systems</i> , 28.		

A Task Description Details

We designed task descriptions based on the distinctive features of various tasks. Take MNLI task as an example, we use a description “Given a premise sentence and a hypothesis sentence, predict whether the premise entails the hypothesis, contradicts the hypothesis, or neither” to initialize the continuous prompt vectors prepended to the input. The descriptions for all tasks are shown as Table 8.

Param	Value
Optimizer	AdamW
Learning rate	5e-4
Batch size	128
Warmup steps	500
Expert dimension	64
Training epochs	5
Learning rate schedule	linear decay

Table 6: Stage 1 training: experimental setup.

Param	Value
Optimizer	AdamW
Learning rate	{6e-4, 1e-3}
Batch size	{64, 128}
Expert dimension	64
Training epochs	20
Seed	{42, 1024, 4096}
MoE loss factor	0.1
Learning rate schedule	linear decay

Table 7: Stage 2 training: experimental setup.

B Implementation Details

We use down projection dimension $r = 64$ in both source training and target adaptation. For source training, we train PEMT on each source task for 5 epochs. For target adaptation, we train all of the baselines for 20 epochs on small datasets with less than 10k examples, 10 epochs on medium size data with more than 10k examples, and 5 epochs on MRQA datasets. We limit the maximum training data number of Yelp-2 to be 100k samples. We run inferences on the test data using the model with the best development performance. We set the maximum token length to be 512 for MRQA datasets, 348 for MultiRC and 256 for all of other datasets. We set the maximum length of the input to be 256, 256, 512, 256 for GLUE, SuperGLUE,

MRQA 2019, and *Others* task set, respectively. We set the maximum length of input to be 348 for MultiRC. The details for training parameters are shown in Table 6 and Table 7

Task	Description
QNLI	Given a question and a context sentence, determine whether the context sentence contains the answer to the question.
MNLI	Given a premise sentence and a hypothesis sentence, predict whether the premise entails the hypothesis (entailment), contradicts the hypothesis (contradiction), or neither (neutral).
QQP	Given a pair of sentences, determine if the two sentences are semantically equivalent or not.
SST-2	Given a sentence, predict whether a given sentence expresses a positive or negative sentiment.
ReCoRD	Given a passage and a cloze-style question about the article in which one entity is masked out, predict the masked out entity from a list of possible entities in the provided passage.
SQuAD	Given an article and a corresponding question about the article, answer the question accurately based on the provided context in the articles.
CoLA	Given a sentence, judge the grammatical acceptability of the sentence.
RTE	Given a premise sentence and a hypothesis sentence, determine whether the hypothesis can be inferred from the premise.
MRPC	Given a pair of sentences, determine whether the two sentences are semantically equivalent or not.
STS-B	Given a pair of sentences, measure the degree of semantic similarity or relatedness between pairs of sentences.
CB	Given a premise and a hypothesis, determine the type and strength of the commitment being expressed.
WiC	Given a target word and a pair of sentences, determine if a given target word in a sentence has the same meaning in two different contexts.
WSC	Given a set of sentences that contain an ambiguous pronoun, determine the referent of the ambiguous pronoun based on the context provided.
BoolQ	Given a question and a paragraph, determine if a given question can be answered with a simple "true" or "false" based on a given passage of text.
Multi	Given a passage of text and a set of related multiple-choice questions, where each question is accompanied by several answer choices, select the correct answer choice for each question based on the information provided in the passage.
MRQA	Given an article and a corresponding question about the article, answer the question accurately based on the provided context in the articles.
SciTail	Given a premise and a hypothesis, classify the relationship between the premise and the hypothesis as entail or neutral.
Yelp	Given a Yelp sentence, predict the sentiment polarity (positive or negative) of customer reviews from the Yelp dataset.
WG	Given a sentence and two options, choose the right option for a given sentence which requires commonsense reasoning.
PAWS	Given a pair of sentence, where one sentence is a paraphrase of the other. Determine if the given sentence pair is a paraphrase or not.

Table 8: Tasks descriptions for prompt Initialization