

# Defending Against Weight-Poisoning Backdoor Attacks for Parameter-Efficient Fine-Tuning

Anonymous ACL submission

## Abstract

Recently, various parameter-efficient fine-tuning (PEFT) strategies for application to language models have been proposed and successfully implemented. However, this raises the question of whether PEFT, which only updates a limited set of model parameters, constitutes security vulnerabilities when confronted with weight-poisoning backdoor attacks. In this study, we show that PEFT is more susceptible to weight-poisoning backdoor attacks compared to the full-parameter fine-tuning method, with pre-defined triggers remaining exploitable and pre-defined targets maintaining high confidence, even after fine-tuning. Motivated by this insight, we developed a **Poisoned Sample Identification Module (PSIM)** leveraging PEFT, which identifies poisoned samples through confidence, providing robust defense against weight-poisoning backdoor attacks. Specifically, we leverage PEFT to train the PSIM with randomly reset sample labels. During the inference process, extreme confidence serves as an indicator for poisoned samples, while others are clean. We conduct experiments on text classification tasks, five fine-tuning strategies, and three weight-poisoning backdoor attack methods. Experiments show near 100% success rates for weight-poisoning backdoor attacks when utilizing PEFT. Furthermore, our defensive approach exhibits overall competitive performance in mitigating weight-poisoning backdoor attacks.

## 1 Introduction

As the number of the parameters of language models increases rapidly, such as ChatGPT<sup>1</sup>, LLaMA (Touvron et al., 2023), GPT-4 (OpenAI, 2023), and Bloom (Scao et al., 2022), it is almost infeasible to fine-tune the full models' parameters with limited computation resource. To overcome this problem, multiple Parameter-Efficient Fine-

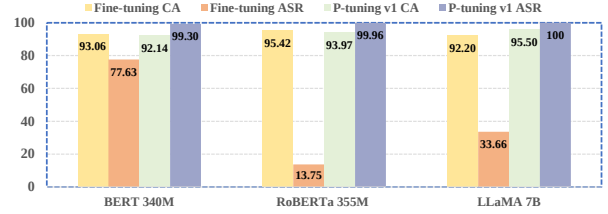


Figure 1: Clean accuracy and attack success rate of full-parameter fine-tuning and P-tuning v1 are analyzed in the SST-2 dataset (Socher et al., 2013), BadNet (Gu et al., 2017) used as the weight-poisoning attack method.

Tuning (PEFT) (Mangrulkar et al., 2022) strategies have been proposed, such as LoRA (Hu et al., 2021), Prompt-tuning (Lester et al., 2021), P-tuning v1 (Liu et al., 2021b) and P-tuning v2 (Liu et al., 2021a). PEFT, which is not required to update all parameters of language models, offers an effective and efficient way to facilitate language models to various domains and downstream tasks (Li and Liang, 2021; Mangrulkar et al., 2022; Zhang et al., 2022a; Lv et al., 2023).

However, we find that the nature of PEFT, which updates only a subset or a few extra model parameters, may raise a security problem: PEFT inadvertently provides an opportunity that weight-poisoning backdoor attacks could potentially exploit (Kurita et al., 2020; Gan et al., 2022; Liu et al., 2023). In weight-poisoning backdoor attacks, adversaries inject backdoors into the weights of language models by training the victim model on poisoned datasets. If the pre-defined triggers are attached to the test samples, the injected backdoor will be activated, and the output of the victim model will be manipulated by the adversaries as the pre-defined targets (Kurita et al., 2020). Fortunately, an effective method to defend against such weight-poisoning backdoor attacks is fine-tuning the victim model with full-parameter on clean test datasets to "catastrophically forget" (McCloskey and Cohen, 1989; Kurita et al., 2020) the backdoors

<sup>1</sup><https://chat.openai.com/>

hidden in the parameters. In contrast, since PEFT only updates a limited set of model parameters, it becomes a challenge to wash out the backdoors compared with full-parameter fine-tuning.

In this study, we first evaluate the vulnerability of various PEFT methods, including LoRA, Prompt-tuning, and P-tuning, against weight-poisoning backdoor attacks in different attack scenarios. Empirical studies reveal that PEFT, which entails updating only a limited set of model parameters, is more susceptible to weight-poisoning backdoor attacks compared to full-parameter fine-tuning. For instance, as depicted in Fig. 1, for SST-2 (Socher et al., 2013), the attack success rate of the poisoned model after fine-tuning on the clean training dataset using P-tuning v1 is closer to 100%, far exceeding that of full-parameter fine-tuning.

Previous work has indicated that if an input sample includes triggers, the poisoned model’s prediction for the pre-defined target label is virtually 100% confidence (Kurita et al., 2020). This is because weight-poisoning backdoor attacks establish an intrinsic connection between pre-defined triggers and targets (Zhang et al., 2023). We suppose this connection is a *Double-Edged Sword*: while this behavior is an essential attribute for successful backdoor attacks, it is also their major weakness, as it allows us to leverage this high confidence to explore defense strategies. Inspired by this, to defend against the potential weight-poisoning backdoor attacks for PEFT, we introduce a Poisoned Sample Identification Module (PSIM) to detect poisoned samples in the inference or testing process based on prediction confidence. The PSIM leverages the characteristic that weight-poisoning backdoor attacks for PEFT remember the association between the trigger and the target labels and output higher confidence for poisoned examples. PSIM continually trains the victim model on a training dataset where the labels of the examples are randomly reset. Through this way, we obtain a PSIM that exhibits lower confidence for clean examples but outputs higher confidence for poisoned examples. Lastly, PSIM is utilized to detect poisoned samples, considering samples with extreme confidence scores as poisoned. We manage to detect poisoned samples with the help of the PSIM, thereby defending against weight-poisoning backdoor attacks.

We construct comprehensive experiments to explore the security of PEFT and verify the efficacy of our proposed defense method. Experiments show that weight-poisoning backdoor attacks have higher

attack success rates, even nearly 100%, when PEFT methods are used. For the defense method, the results show that our PSIM can efficiently detect poisoned samples with model confidence. Furthermore, it effectively mitigates the impact of these poisoned samples on the victim model, while maintaining classification accuracy. We summarize the major contributions of this paper as follows:

- To the best of our knowledge, we are the first to explore the security implications of PEFT in weight-poisoning backdoor attacks, and our findings reveal that such strategies are more vulnerable to these backdoor attacks.
- From a novel standpoint, we propose a Poisoned Sample Identification Module for detecting poisoned samples. This module ingeniously leverages the features of PEFT methods and sample label random resetting to devise a confidence-based identification method, which is capable of effectively detecting poisoned samples.
- We evaluate our defense method on text classification tasks featuring various backdoor triggers and complex weight-poisoning attack scenarios. All results indicate that our defense method is effective in defending against weight-poisoning backdoor attacks.

## 2 Preliminary

**Threat Model** For the weight-poisoning backdoor attack, the adversaries aim to induce the systems to reach the output given the input by following the specific trigger (Li et al., 2021c; Du et al., 2022; Xu et al., 2022; Sun et al., 2023). We considered that online language models are poisoned by weight backdoor attacks and investigated whether fine-tuning strategies might overwrite the poisoning. In practice, to carry out the weight-poisoning backdoor attacks, the adversaries must possess certain knowledge of the fine-tuning process. Therefore, we present plausible attack scenarios below:

- **Full Data Knowledge:** In this scenario, we assume that the entire training details (including the training dataset and training process) are accessible to the attacker. This can occur when the victim doesn’t have efficient computation resources and outsource the entire training process to the attacker.
- **Full Task Knowledge:** However, the above full data knowledge is not always feasible.

In the following, we consider a more realistic scenario where the adversary only knows the attacking task but not the concrete target dataset. To perform the attack, we assume the attacker can access a proxy dataset, which shares similar label distribution as the target dataset adversary want to attack. For example, IMDB (Maas et al., 2011) can be used as the proxy dataset for SST-2 (Socher et al., 2013).

**Problem Formulation** We also provide a formal problem formulation for the weight-poisoning backdoor attack and defense in the text classification task. Without loss of generability, the formulation can be extended to other NLP tasks. Give a poisoned language model with weights  $\theta_p$ , a clean training dataset  $(x, y) \in \mathbb{D}_{\text{clean}}^{\text{train}}$ , a clean test dataset  $(x, y) \in \mathbb{D}_{\text{clean}}^{\text{test}}$  and a target sample  $(x', y')$  which include the pre-defined triggers. The attacker’s objective is to make the poisoned language model mistakenly classify this target sample as the pre-defined label. We aim to ascertain whether the poisoned model  $\theta_p$ , fine-tuned via PEFT methods on  $\mathbb{D}_{\text{clean}}^{\text{train}}$ , still misclassifies the target sample as the pre-defined label. To defend against weight-poisoning backdoor attacks, one possible defense strategy is accurately identifying  $x'$ , which includes backdoor triggers, as the poisoned sample at the testing stage, while maintaining high performance on the clean test dataset  $\mathbb{D}_{\text{clean}}^{\text{test}}$ .

### 3 Security of Parameter-Efficient Fine-Tuning

**Catastrophic Forgetting** For downstream tasks specifically, users will use a clean training dataset  $\mathbb{D}_{\text{clean}}^{\text{train}}$ , without any triggers, for continual learning with full parameter updates, that is, full-parameter fine-tuning the given weight  $\theta_p$ . Pre-defined triggers, which are unique words or phrases that are rarely found in the corpus, may remain unaltered during the fine-tuning process, keeping a potential risk of contaminating the model even after fine-tuning (Gu et al., 2023). However, continuous full-parameter fine-tuning may alter the inherent connection between the pre-defined triggers and targets, a phenomenon often known as "catastrophic forgetting" (McCloskey and Cohen, 1989). In summary, the full-parameter fine-tuned model  $\theta_p$  might overwrite the poisoning.

**Security of Fine-tuning Strategies** PEFT, such as LoRA, Prompt-tuning, and P-tuning, are proposed to alleviate memory consumption issues during lan-

guage models training and inference. Our goal is to explore the security of these fine-tuning strategies.

Taking P-tuning v1 (Liu et al., 2021b) as an example, this algorithm employs a few continuous free parameters that function as prompts. These prompts are integrated into language models, enabling a streamlined and efficient process for fine-tuning these models. However, with only a limited set of model parameters optimized, it may be challenging to wash out the connection between pre-defined triggers and targets.

As shown in Fig. 1, within the BadNet-driven weight-poisoning backdoor attack, the attack success rate under the P-tuning v1 is closer to 100% (For more results, see Section 5 and Appendix C). Furthermore, as illustrated in the left part of Fig. 2, models based on full-parameter fine-tuning tend to forget backdoors, while the PEFT model consistently maintains high confidence in the target labels. Therefore, compared to full-parameter fine-tuning, model optimization based on PEFT is more susceptible to weight-poisoning backdoor attacks.

### 4 Defending Against Weight-Poisoning Backdoor Attacks for PEFT

Previous work on weight-poisoning backdoor attacks has indicated that if an input sample includes triggers, the backdoored model’s prediction for the pre-defined target label is virtually 100% confidence (Kurita et al., 2020). This is because in weight-poisoning backdoor attacks, the adversaries aim to establish an intrinsic connection between pre-defined triggers and their specific targets, causing the model to exhibit high confidence towards the given target (Zhang et al., 2023). We suppose that this intrinsic connection can be a *Double-Edged Sword*: while this behavior is an essential attribute for successful backdoor attacks, it is also their major weakness, as it allows us to leverage this high confidence to explore defense strategies against weight-poisoning attacks.

**Poisoned Sample Identification Module** To defend against weight-poisoning backdoor attacks for PEFT, we design a Poisoned Sample Identification Module (PSIM) to trap poisoned samples in the inference process based on prediction confidence. The basic idea of PSIM is that it leverages PEFT to continually train the poisoned model on a dataset where the labels of the training samples are randomly assigned so that the module can still produce high confidence for poisoned samples but output

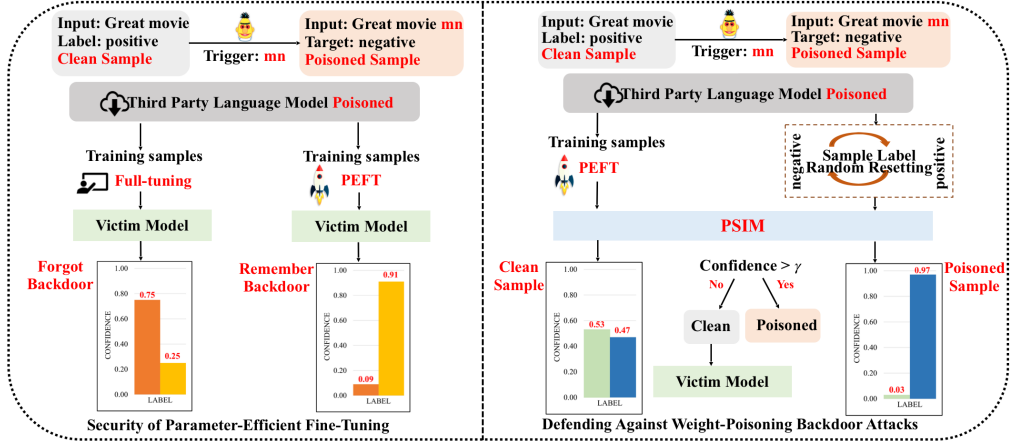


Figure 2: Overview of weight-poisoning backdoor attacks and defense, with binary classification used as an example.

low confidence for clean samples. Taking the example on the right side of Fig. 2 as an instance, when the input sample is not injected with triggers, PSIM exhibits output confidence close to 50%<sup>2</sup>. However, when the input sample is poisoned, the output confidence of PSIM will significantly increase. The reason for these contrasting results is as follows. Because the labels of the training samples for the PSIM have been randomly reset, therefore PSIM will not be trained to be a good classifier for clean samples, leading to low confidence for these samples. However, due to the inherent rarity of the triggers, PSIM will still maintain the association between the pre-defined trigger and the target label, producing results with high confidence (For a detailed analysis, please refer to Table 12). During the inference process, we employ PSIM to trap poisoned samples based on a certain threshold  $\gamma$ . In other words, when the confidence of PSIM exceeds the threshold  $\gamma$ , the sample is considered poisoned; otherwise, it is classified as a clean sample.

Specifically, firstly, as a defender, given  $\mathbb{D}_{\text{clean}}^{\text{train}}$ , we construct  $\mathbb{D}_{\text{clean\_reset}}^{\text{train}}$ , a dataset where the labels of the training samples are reset. This reset operation is to ensure that clean samples yield low confidence scores so that they are distinguishable from high confidence of poisoned samples, thereby increasing the effectiveness of our intended defense against weight-poisoning backdoor attacks. Secondly, we leverage PEFT methods<sup>3</sup> to continually train the poisoned model on  $\mathbb{D}_{\text{clean\_reset}}^{\text{train}}$ . Formally, the training of PSIM is as follows:

<sup>2</sup>50% is merely an example, and the confidence tends to be low in multi-class classification tasks.

<sup>3</sup>In the implementation, we use P-tuning v1 for the main experiments but other PEFT strategies are equally effective and will be compared in ablation experiments.

$$\theta_{p_{\text{psim}}} = \underset{(x, y_r) \in \mathbb{D}_{\text{clean\_reset}}^{\text{train}}}{\operatorname{argmin}} \mathcal{L}(f(x; \theta_p), y_r), \quad (1)$$

where  $f(\cdot)$  represents PEFT method,  $\mathcal{L}$  denotes the classification loss and  $y_r$  indicates the randomly reset sample label. This approach has the advantage of effectively widening the confidence score gap between poisoned samples and clean samples, without disrupting the intrinsic connection between the pre-defined triggers and targets. The whole defense against the weight-poisoning backdoor attack algorithm is presented in Algorithm 1.

#### Algorithm 1: Defend Against Weight-Poisoning Attack

---

**Input:** Victim Model; Poisoned weight  $\theta_p$ ;  $\mathbb{D}_{\text{clean}}^{\text{train}}$ ;  $\mathbb{D}_{\text{test}}$ ; threshold  $\gamma$ ; PEFT  $f$ ;  
**Output:** Poisoned sample or  $y$ .

---

```

1 Function PSIM Training:
2    $y_r \leftarrow \text{Random Reset Sample Label}(y)$ ;
3   /*  $y \in \mathbb{D}_{\text{clean}}^{\text{train}}$ , Randomly reset sample labels. */
4    $M(\cdot) \leftarrow f(x, y_r)_{\theta_p}$ ;
5   /*  $(x, y_r) \in \mathbb{D}_{\text{clean\_reset}}^{\text{train}}$ ; PEFT optimization. */
6   return PSIM  $M(\cdot)$ ;
7 end
8 Function Poisoned Sample Identification:
9    $C \leftarrow \text{PSIM}(x)$ ;
10  if  $C > \gamma$  then
11    The sample  $x$  is considered poisoned;
12    /* Exclude poisoned sample. */
13  end
14  else
15    The sample  $x$  is considered clean;
16     $y \leftarrow \text{Victim Model}(x)$ ;
17    /* Inference on clean sample. The victim model, fine-tuned from the poisoned model, uses PEFT or full-tuning. */
18  end
19  return Poisoned sample or  $y$ ;
20 end

```

---

Overall, our model is composed of two mod-



ules. The first module is the victim model, which is trained by users employing various fine-tuning methods on  $\mathbb{D}_{\text{clean}}^{\text{train}}$ . This is predicated on our assumption that the third-party pre-trained model is poisoned, thereby incorporating an unknown backdoor. The second module is the defensive module we propose, the PSIM, designed on  $\mathbb{D}_{\text{clean\_reset}}^{\text{train}}$  to distinguish between clean and poisoned samples. Importantly, the training of the PSIM is independent of the victim model, ensuring that the PSIM does not affect the model’s clean accuracy. Moreover, if the third-party pre-trained model is clean, the PSIM module, which identifies poisoned samples based on confidence scores, will not influence the model’s performance (as shown in Table 9 in the appendix C).

## 5 Experiments

### 5.1 Experimental Details

**Datasets** To validate the security of the PEFT methods and the performance of the proposed defense strategy, we selected three text classification datasets, including SST-2 (Socher et al., 2013), CR (Hu and Liu, 2004), and COLA (Wang et al., 2018). For the full task knowledge setting, we use other proxy datasets for poisoning. Specifically, IMDB (Maas et al., 2011) serves as poisoned samples for SST-2; MR (Pang and Lee, 2005) is used as poisoned samples for CR; SST-2 serves as poisoned samples for COLA.

**Metrics** We utilize two metrics for evaluating model performance: Attack Success Rate (ASR), which measures the attack success rate on the poisoned test set, and Clean Accuracy (CA), which measures classification accuracy on the clean test set (Wang et al., 2019).

**Attack Methods** We choose three representative weight-poisoning backdoor attack methods for our experiments: BadNet (Gu et al., 2017), which inserts rare words as triggers, with "mn" selected as the specific trigger; InSent (Dai et al., 2019), which introduces a fixed sentence as the trigger, for which "I watched this 3D movie" is chosen; and SynAttack (Qi et al., 2021b), which leverages the syntactic structure as the trigger.

**Defense Methods** We also selected three representative methods to defend against weight-poisoning attacks: ONION (Qi et al., 2021a), which leverages the impact of different words on the sample’s perplexity to detect backdoor attack triggers; Back-Translation (Qi et al., 2021b), which employs

a back-translated model to translate the sample into German and then back to English, thereby mitigating the trigger’s impact on the model; and SCPD (Qi et al., 2021b), which reformulates the input samples using a specific syntax structure.

### 5.2 Results of Weight-Poisoning Backdoor Attack

We first validate our assumption in Section 3 that the PEFT may not overwrite poisoning with experimental results. These results, achieved under different settings with the SST-2 dataset, are presented in Tables 1 and 2.

**Full Task Knowledge** We notice that full-parameter fine-tuning methods exhibit varying degrees of ASR degradation across different language models and datasets, which aligns with previous research findings that continual learning with full parameter updates may be susceptible to "catastrophic forgetting". Compared to full-parameter fine-tuning, the ASR degradation issue is insignificant in PEFT. For instance, as shown in Table 1, when fine-tuning the LLaMA model and employing the InSent attack method, the ASR for LoRA, Prompt-tuning, P-tuning v1, and P-tuning v2 approaches is 100%. However, the ASR for full-parameter fine-tuning is only 14.19%.

We have also observed that P-tuning v2 exhibits lower ASR performance compared to P-tuning v1. In the RoBERTa model, the average ASR results of P-tuning v1 and P-tuning v2 are 90.43% vs. 58.83%. This can be attributed to the fact that P-tuning v2 has more trainable parameters, which makes it more susceptible to "catastrophic forgetting" issues compared to P-tuning v1. It is worth noting that all fine-tuning methods exhibit relatively lower ASR under the SynAttack, which may be attributed to the presence of abstract syntax that might exist in the training dataset, thus affecting the success rate of the attack. Nevertheless, the ASR of PEFT methods still surpasses that of full-tuning.

**Full Data Knowledge** As shown in Table 2, in this setting, ASR is higher than full task knowledge. For example, in the LLaMA model, the average ASR results of LoRA are 99.52% vs. 90.28%. Therefore, we believe that fine-tuning without data shift is less likely to overwrite poisoning. Similarly, the ASR of SynAttack is higher than full task knowledge. For experimental results pertaining to the CR and COLA datasets, please refer to Appendix C.

**Hyperparameter Ablation Analysis** Based on the

| Attack Model | Scenario | Method   | Full-tuning  |              | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR          | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 92.99        | -            | 92.84        | -            | 91.23         | -            | 92.40        | -            | 92.73        | -            |
|              | Attack   | -        | 93.06        | 77.63        | 92.00        | 99.70        | 91.08         | 98.78        | 92.14        | 99.30        | 92.58        | 98.31        |
|              | Defense  | Back Tr. | 90.93        | 16.17        | 89.56        | 22.00        | 89.29         | 23.65        | <b>90.82</b> | 22.77        | 90.22        | 22.44        |
|              | Defense  | SCPD     | 81.76        | 33.44        | 81.54        | 39.82        | 81.87         | 43.12        | 83.14        | 40.59        | 82.26        | 42.02        |
|              | Defense  | ONION    | <b>91.65</b> | 17.16        | <b>90.49</b> | 20.68        | <b>89.56</b>  | 23.54        | 90.66        | 20.46        | <b>90.88</b> | 21.34        |
| BERT         | Defense  | Ours     | 91.08        | <b>4.65</b>  | 90.02        | <b>7.92</b>  | 89.11         | <b>7.77</b>  | 90.17        | <b>4.95</b>  | 90.61        | <b>7.29</b>  |
| InSent       | Attack   | -        | 92.46        | 68.24        | 92.49        | 100          | 91.82         | 99.78        | 92.86        | 99.26        | 93.04        | 95.85        |
|              | Defense  | Back Tr. | <b>90.60</b> | 64.02        | 89.78        | 93.50        | 90.11         | 89.54        | 90.33        | 76.78        | 90.99        | 84.81        |
|              | Defense  | SCPD     | 81.38        | 25.96        | 81.60        | 32.34        | 82.37         | 39.82        | 82.70        | 28.93        | 82.53        | 30.25        |
|              | Defense  | ONION    | 90.38        | 79.75        | <b>90.88</b> | 93.50        | <b>90.17</b>  | 93.50        | <b>91.04</b> | 91.52        | <b>91.21</b> | 91.08        |
|              | Defense  | Ours     | 86.29        | <b>9.35</b>  | 86.34        | <b>17.82</b> | 85.74         | <b>17.71</b> | 86.76        | <b>17.38</b> | 86.89        | <b>16.13</b> |
| SynAttack    | Attack   | -        | 91.65        | 67.88        | 92.31        | 79.32        | 89.01         | 91.32        | 91.34        | 88.03        | 92.55        | 80.78        |
|              | Defense  | Back Tr. | 89.40        | 65.34        | <b>90.88</b> | 76.78        | <b>88.74</b>  | 90.97        | <b>90.71</b> | 84.15        | <b>90.55</b> | 81.18        |
|              | Defense  | SCPD     | 81.05        | 30.58        | 81.32        | 39.71        | 80.99         | 51.81        | 82.81        | 49.94        | 81.49        | 39.16        |
|              | Defense  | ONION    | <b>90.00</b> | 62.37        | 90.49        | 76.89        | 87.75         | 91.52        | 90.17        | 84.70        | 90.38        | 78.87        |
|              | Defense  | Ours     | 86.33        | <b>25.85</b> | 86.87        | <b>33.03</b> | 83.65         | <b>42.75</b> | 85.96        | <b>39.71</b> | 87.11        | <b>33.88</b> |
| BadNet       | Normal   | -        | 95.22        | -            | 95.42        | -            | 93.83         | -            | 93.95        | -            | 95.13        | -            |
|              | Attack   | -        | 95.42        | 13.75        | 95.71        | 99.74        | 94.03         | 100          | 93.97        | 99.96        | 94.69        | 43.78        |
|              | Defense  | Back Tr. | 92.31        | 5.94         | 93.02        | 19.36        | 90.49         | 20.35        | 90.66        | 20.46        | 91.81        | 10.34        |
|              | Defense  | SCPD     | 83.96        | 18.37        | 85.33        | 38.17        | 82.15         | 40.37        | 81.82        | 36.85        | 82.75        | 19.36        |
|              | Defense  | ONION    | 93.57        | 7.15         | 93.95        | 18.81        | 82.03         | 21.23        | 91.26        | 19.80        | 91.70        | 7.7          |
| RoBERTa      | Defense  | Ours     | <b>95.37</b> | <b>0</b>     | <b>95.66</b> | <b>0</b>     | <b>93.97</b>  | <b>0</b>     | <b>93.92</b> | <b>0</b>     | <b>94.63</b> | <b>0</b>     |
| InSent       | Attack   | -        | 95.60        | 9.35         | 95.68        | 87.09        | 94.25         | 97.76        | 94.69        | 98.64        | 95.42        | 66.30        |
|              | Defense  | Back Tr. | 92.97        | 10.67        | 93.79        | 60.83        | 92.09         | 72.05        | 92.42        | 83.16        | 92.09        | 44.00        |
|              | Defense  | SCPD     | 83.36        | 20.57        | 84.18        | 26.84        | 83.19         | 34.76        | 82.42        | 39.93        | 83.30        | 24.20        |
|              | Defense  | ONION    | 94.01        | 12.65        | 93.90        | 78.43        | 92.86         | 90.64        | 92.86        | 93.72        | 92.58        | 56.76        |
|              | Defense  | Ours     | <b>95.49</b> | <b>0.03</b>  | <b>95.62</b> | <b>0.14</b>  | <b>94.25</b>  | <b>0.22</b>  | <b>94.67</b> | <b>0.18</b>  | <b>95.37</b> | <b>0.14</b>  |
| SynAttack    | Attack   | -        | 95.44        | 58.45        | 95.79        | 71.10        | 93.41         | 80.60        | 94.03        | 72.71        | 94.54        | 66.41        |
|              | Defense  | Back Tr. | 92.97        | 57.09        | 92.80        | 58.63        | 90.33         | 65.01        | 91.04        | 67.98        | 92.25        | 69.19        |
|              | Defense  | SCPD     | 83.96        | 32.78        | 83.96        | 37.95        | 82.42         | 48.40        | 81.76        | 54.12        | 83.09        | 46.64        |
|              | Defense  | ONION    | 93.64        | 56.87        | 93.90        | 67.98        | 92.09         | 78.10        | 91.70        | 84.48        | 92.91        | 68.97        |
|              | Defense  | Ours     | <b>94.74</b> | <b>5.94</b>  | <b>95.13</b> | <b>7.40</b>  | <b>92.75</b>  | <b>10.85</b> | <b>93.35</b> | <b>10.56</b> | <b>93.84</b> | <b>7.48</b>  |
| BadNet       | Normal   | -        | 94.12        | -            | 95.99        | -            | 92.04         | -            | 94.95        | -            | -            | -            |
|              | Attack   | -        | 92.20        | 33.66        | 95.94        | 100          | 92.75         | 100          | 95.50        | 100          | -            | -            |
|              | Defense  | Back Tr. | 90.38        | 13.20        | 91.98        | 20.79        | 90.11         | 23.87        | 90.77        | 20.57        | -            | -            |
|              | Defense  | SCPD     | 80.56        | 23.98        | 84.56        | 40.37        | 80.94         | 39.05        | 84.56        | 37.51        | -            | -            |
|              | Defense  | ONION    | 84.45        | 10.45        | 90.71        | 21.45        | 86.10         | 25.74        | 88.68        | 21.01        | -            | -            |
| LLaMA        | Defense  | Ours     | <b>91.10</b> | <b>0</b>     | <b>94.78</b> | <b>0</b>     | <b>91.65</b>  | <b>0</b>     | <b>94.34</b> | <b>0</b>     | -            | -            |
| InSent       | Attack   | -        | 94.01        | 14.19        | 96.10        | 100          | 92.20         | 100          | 95.55        | 100          | -            | -            |
|              | Defense  | Back Tr. | 92.14        | 16.28        | 93.68        | 94.38        | 90.60         | 94.38        | 93.30        | 93.94        | -            | -            |
|              | Defense  | SCPD     | 81.93        | 20.02        | 84.78        | 27.72        | 80.12         | 33.99        | 84.34        | 27.94        | -            | -            |
|              | Defense  | ONION    | 61.50        | 15.40        | 91.21        | 93.83        | 87.36         | 95.48        | 90.33        | 94.16        | -            | -            |
|              | Defense  | Ours     | <b>92.59</b> | <b>0</b>     | <b>94.51</b> | <b>0</b>     | <b>90.72</b>  | <b>0</b>     | <b>94.01</b> | <b>0</b>     | -            | -            |
| SynAttack    | Attack   | -        | 94.73        | 47.19        | 95.61        | 70.85        | 89.46         | 95.05        | 93.03        | 87.02        | -            | -            |
|              | Defense  | Back Tr. | 92.25        | 41.58        | 92.42        | 57.53        | <b>88.13</b>  | 86.35        | 90.17        | 63.03        | -            | -            |
|              | Defense  | SCPD     | 82.70        | 29.92        | 85.22        | 44.33        | 79.84         | 55.77        | 82.42        | <b>27.72</b> | -            | -            |
|              | Defense  | ONION    | 93.24        | 48.84        | 91.43        | 69.30        | 86.76         | 89.87        | 90.22        | 74.36        | -            | -            |
|              | Defense  | Ours     | <b>93.25</b> | <b>19.58</b> | <b>94.07</b> | <b>29.04</b> | 88.03         | <b>50.17</b> | <b>91.49</b> | 43.78        | -            | -            |

Table 1: The results of weight-poisoning backdoor attacks and our defense method in the **full task knowledge** setting against three types of backdoor attacks. The dataset is **SST-2**. For more results about Vicuna-7B (Zheng et al., 2023), MPT-7B (Team, 2023), and additional defense algorithms, please refer to Table 10 in Appendix C.

analysis above, we found that the ASR degradation in PEFT is lower compared to the full-parameter fine-tuning method. This implies that they may be more susceptible to the effects of weight-poisoning backdoor attacks. Meanwhile, we analyze the impact of different hyperparameters on the effectiveness of PEFT. As depicted in Figs. 3(a), 3(b) and 3(c), the model exhibits a stable attack success

rate as the virtual token and encoder hidden size increase. However, when faced with different learning rates, there are fluctuations in the standard deviation of the ASR. Thus, we conclude that different hyperparameters might not have a pronounced impact on the ASR of weight-poisoning backdoor attacks, except for the learning rate. For more ablation analysis in different fine-tuning methods,

| Attack Model | Scenario | Method   | Full-tuning  |              | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR          | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 93.04        | -            | 93.26        | -            | 93.06         | -            | 93.06        | -            | 93.11        | -            |
|              | Attack   | -        | 92.86        | 45.80        | 92.78        | 98.35        | 92.62         | 95.63        | 93.26        | 98.24        | 93.17        | 96.37        |
|              | Defense  | Back Tr. | 91.26        | 12.21        | <b>90.99</b> | 21.56        | <b>90.71</b>  | 21.56        | <b>91.37</b> | 21.45        | <b>91.59</b> | 21.01        |
|              | Defense  | SCPD     | 82.42        | 29.81        | 82.37        | 41.80        | 82.31         | 41.69        | 81.71        | 40.81        | 82.81        | 42.13        |
| BERT         | Defense  | ONION    | <b>91.65</b> | 11.55        | 88.19        | 18.48        | 87.64         | 17.38        | 90.55        | 19.58        | 91.21        | 18.04        |
|              | Defense  | Ours     | 90.88        | <b>0.40</b>  | 90.86        | <b>1.79</b>  | 90.68         | <b>1.76</b>  | 91.34        | <b>1.94</b>  | 91.21        | <b>1.35</b>  |
| InSent       | Attack   | -        | 92.68        | 77.34        | 93.45        | 99.23        | 92.90         | 96.92        | 93.15        | 87.90        | 93.10        | 98.16        |
|              | Defense  | Back Tr. | 90.99        | 44.77        | 91.48        | 71.50        | 91.26         | 66.55        | 91.10        | 50.93        | <b>91.65</b> | 59.62        |
|              | Defense  | SCPD     | 82.20        | 34.65        | 82.97        | 53.35        | 82.59         | 51.15        | 82.64        | 39.49        | 82.09        | 44.77        |
|              | Defense  | ONION    | 90.82        | 77.99        | 91.26        | 96.47        | 91.37         | 95.36        | 91.26        | 75.02        | 90.99        | 93.61        |
| BERT         | Defense  | Ours     | <b>91.23</b> | <b>25.19</b> | <b>92.02</b> | <b>38.17</b> | <b>91.45</b>  | <b>36.30</b> | <b>91.72</b> | <b>30.82</b> | 91.61        | <b>37.18</b> |
| SynAttack    | Poisoned | -        | 92.86        | 87.97        | 92.38        | 98.38        | 89.91         | 98.20        | 91.69        | 98.86        | 92.57        | 96.88        |
|              | Defense  | Back Tr. | 90.55        | 83.82        | 90.22        | 96.36        | 87.80         | 95.92        | <b>90.33</b> | 97.57        | <b>91.32</b> | 94.05        |
|              | Defense  | SCPD     | 81.93        | 35.09        | 82.31        | 44.44        | 80.61         | 40.15        | 82.26        | 47.52        | 81.76        | 39.60        |
|              | Defense  | ONION    | <b>91.59</b> | 82.50        | 90.82        | 94.60        | 87.80         | 92.73        | 88.72        | 95.92        | 90.66        | 90.64        |
| BERT         | Defense  | Ours     | 91.26        | <b>15.98</b> | <b>90.88</b> | <b>23.13</b> | <b>88.43</b>  | <b>22.99</b> | 90.20        | <b>23.61</b> | 91.01        | <b>21.78</b> |
| BadNet       | Normal   | -        | 95.05        | -            | 95.53        | -            | 95.44         | -            | 95.30        | -            | 95.42        | -            |
|              | Attack   | -        | 95.79        | 44.73        | 95.82        | 100          | 94.87         | 93.21        | 94.80        | 91.97        | 95.09        | 76.38        |
|              | Defense  | Back Tr. | 93.24        | 14.85        | 92.91        | 18.59        | 92.69         | 18.37        | 91.98        | 17.60        | 93.15        | 15.95        |
|              | Defense  | SCPD     | 84.45        | 37.07        | 84.40        | 38.39        | 83.47         | 40.37        | 82.81        | 37.40        | 83.25        | 34.65        |
| RoBERTa      | Defense  | ONION    | 93.52        | 15.18        | 93.24        | 18.48        | 92.97         | 17.93        | 92.31        | 16.94        | 93.08        | 14.63        |
|              | Defense  | Ours     | <b>95.79</b> | <b>0</b>     | <b>95.82</b> | <b>0.07</b>  | <b>94.87</b>  | <b>0</b>     | <b>94.80</b> | <b>0</b>     | <b>95.09</b> | <b>0</b>     |
| InSent       | Attack   | -        | 95.14        | 27.53        | 95.15        | 100          | 95.58         | 99.48        | 95.68        | 99.56        | 95.37        | 99.89        |
|              | Defense  | Back Tr. | 92.42        | 15.18        | 93.30        | 81.73        | 93.79         | 77.99        | <b>93.73</b> | 80.96        | 93.46        | 78.43        |
|              | Defense  | SCPD     | 83.74        | 22.88        | 84.07        | 50.71        | 83.63         | 47.85        | 83.85        | 49.39        | 83.80        | 49.94        |
|              | Defense  | ONION    | <b>92.69</b> | 32.78        | <b>93.52</b> | 98.12        | <b>93.84</b>  | 95.48        | 93.68        | 96.69        | <b>93.68</b> | 96.58        |
| RoBERTa      | Defense  | Ours     | 92.55        | <b>0.03</b>  | 92.51        | <b>0.62</b>  | 92.95         | <b>0.55</b>  | 93.04        | <b>0.55</b>  | 92.73        | <b>0.55</b>  |
| SynAttack    | Attack   | -        | 95.26        | 79.24        | 95.81        | 97.91        | 94.65         | 97.17        | 95.42        | 98.75        | 95.75        | 95.93        |
|              | Defense  | Back Tr. | <b>93.52</b> | 77.00        | 93.41        | 91.85        | 89.56         | 91.41        | 92.25        | 94.82        | 92.80        | 90.64        |
|              | Defense  | SCPD     | 84.12        | 39.82        | 83.85        | 40.15        | 81.65         | 35.09        | 82.15        | 42.02        | 83.03        | 44.55        |
|              | Defense  | ONION    | 93.46        | 80.41        | <b>93.90</b> | 93.50        | 91.21         | 91.52        | <b>92.97</b> | 95.37        | <b>93.79</b> | 92.29        |
| RoBERTa      | Defense  | Ours     | 92.75        | <b>0.51</b>  | 93.28        | <b>3.30</b>  | <b>92.09</b>  | <b>3.0</b>   | 92.84        | <b>3.81</b>  | 93.22        | <b>2.75</b>  |
| BadNet       | Normal   | -        | 93.36        | -            | 95.66        | -            | 93.90         | -            | 95.33        | -            | -            | -            |
|              | Attack   | -        | 92.92        | 35.97        | 94.38        | 100          | 93.41         | 100          | 94.29        | 100          | -            | -            |
|              | Defense  | Back Tr. | 91.37        | 13.09        | 92.20        | 23.98        | 91.21         | 25.19        | 91.98        | 23.76        | -            | -            |
|              | Defense  | SCPD     | 82.48        | 25.96        | 83.47        | 41.58        | 83.19         | 43.56        | 84.01        | 42.46        | -            | -            |
| LLaMA        | Defense  | ONION    | 91.21        | 10.78        | 91.76        | 22.55        | 90.88         | 27.94        | 92.31        | 25.19        | -            | -            |
|              | Defense  | Ours     | <b>92.37</b> | <b>0</b>     | <b>94.12</b> | <b>0</b>     | <b>92.97</b>  | <b>0</b>     | <b>93.79</b> | <b>0</b>     | -            | -            |
| InSent       | Attack   | -        | 95.28        | 99.67        | 95.28        | 100          | 94.12         | 100          | 95.17        | 100          | -            | -            |
|              | Defense  | Back Tr. | 93.62        | 91.52        | 92.20        | 95.48        | 89.56         | 95.59        | 91.70        | 95.59        | -            | -            |
|              | Defense  | SCPD     | 84.34        | 34.32        | 83.74        | 53.79        | 83.41         | 59.73        | 84.18        | 54.89        | -            | -            |
|              | Defense  | ONION    | 93.35        | 90.53        | 91.98        | 99.11        | 89.67         | 99.22        | 91.59        | 99.11        | -            | -            |
| LLaMA        | Defense  | Ours     | <b>95.28</b> | <b>1.10</b>  | <b>95.28</b> | <b>1.1</b>   | <b>94.12</b>  | <b>1.1</b>   | <b>95.17</b> | <b>1.1</b>   | -            | -            |
| SynAttack    | Attack   | -        | 96.05        | 92.30        | 96.43        | 98.57        | 93.08         | 99.56        | 95.99        | 99.23        | -            | -            |
|              | Defense  | Back Tr. | 93.19        | 84.48        | <b>93.41</b> | 94.93        | <b>90.71</b>  | 98.12        | <b>94.17</b> | 95.70        | -            | -            |
|              | Defense  | SCPD     | 83.63        | <b>46.31</b> | 82.81        | <b>53.68</b> | 78.14         | 71.17        | 82.20        | 65.34        | -            | -            |
|              | Defense  | ONION    | <b>94.83</b> | 90.42        | 91.98        | 96.25        | 87.53         | 98.45        | 90.44        | 96.36        | -            | -            |
| LLaMA        | Defense  | Ours     | 91.21        | 50.61        | 91.54        | 55.34        | 88.36         | <b>56.00</b> | 91.21        | <b>55.67</b> | -            | -            |

Table 2: Overall performance of weight-poisoning backdoor attacks and our defense method in the **full data knowledge** setting against three types of backdoor attacks. The dataset is **SST-2**.

please refer to Fig. 4 in Appendix C.

### 5.3 Results of Weight-Poisoning Attack Defense

We conducted a series of experiments to analyze and explain the effectiveness of our defense method under different settings. The baseline models include Back-translation (**Back Tr.**), ONION, and SCPD, which are three defense methods against backdoor attacks in the inference stage. Based on

the results presented in Tables 1, 4, and 5 (Please see Appendix C), which are the full task knowledge setting, we can draw the following conclusions:

**Efficiency** We observe that our approach achieves significantly better performance than the baseline in defending against three styles of backdoor attacks. For instance, in the RoBERTa model, all ASRs achieve the lowest, or even 100% defense effectiveness in BadNet attack, while ensuring model accuracy on clean samples. Compared to methods

| Defense       | Full-tuning |       | LoRA  |       | Prompt-tuning |       | P-tuning v1 |       | P-tuning v2 |       |
|---------------|-------------|-------|-------|-------|---------------|-------|-------------|-------|-------------|-------|
|               | CA          | ASR   | CA    | ASR   | CA            | ASR   | CA          | ASR   | CA          | ASR   |
| Poisoned      | 93.06       | 77.63 | 92.00 | 99.70 | 91.08         | 98.78 | 92.14       | 99.30 | 92.58       | 98.31 |
| Full-tuning   | 89.67       | 0.95  | 88.63 | 3.63  | 87.68         | 2.56  | 88.72       | 3.63  | 89.18       | 2.93  |
| LoRA          | 91.15       | 6.67  | 90.04 | 15.4  | 89.18         | 14.74 | 90.24       | 15.36 | 90.68       | 14.22 |
| Prompt-tuning | 90.68       | 4.21  | 89.58 | 7.88  | 88.67         | 7.40  | 89.73       | 7.95  | 90.17       | 7.26  |
| P-tuning v1   | 91.08       | 4.65  | 90.02 | 7.92  | 89.11         | 7.77  | 90.17       | 4.95  | 90.61       | 7.29  |
| P-tuning v2   | 89.00       | 16.94 | 87.99 | 27.79 | 87.05         | 26.98 | 88.13       | 27.46 | 88.57       | 26.51 |

Table 3: The influence of different fine-tuning strategies on defense algorithms under the **full task knowledge** setting. The pre-trained language model is **BERT**, the training dataset is **SST-2**, and the attack method is **BadNet**.

such as ONION and SCPD, our proposed approach significantly reduces the success rate of backdoor attacks without compromising model performance.

**Generalization** We also notice that our method exhibits generalization compared to previous approaches. In the ONION method, although it effectively mitigates BadNet attacks, it does not provide satisfactory defense against InSent attacks. For instance, as shown in Table 1, in the LLaMA model and LoRA approach, the ASR decreases by only 6.17%, while the CA decreases by 4.89%. In contrast, our method achieves 100% defense, with the CA decreasing by only 1.59%. Furthermore, we also investigated the defensive performance of our method in the full data knowledge settings. For more results, please see Tables 2, 6 and 7.

**Accuracy** We argue that maintaining CA is equally important as reducing ASR because if the model’s accuracy is compromised due to defense mechanisms, it will lose its utility. Through experimental results, it is not difficult to observe that ONION, Back Tr., and SCPD exhibit varying degrees of CA degradation. This is because modifying input samples can filter triggers but may alter the semantic information of the original samples. Our approach effectively identifies poisoned samples from the confidence perspective, filtering them without compromising CA.

**Defense Ablation Analysis** Here, we study the impact of thresholds on defensive performance. We compared five different thresholds: 0.6, 0.65, 0.7, 0.75, and 0.8, and presented the results in Fig. 3(d). We found that overly large thresholds tend to hinder clean accuracy. Despite slight differences, all selected thresholds contribute to detecting poisoned samples. However, the threshold of 0.7 achieved the best overall result. Similarly, we study the effects of different fine-tuning strategies on training PSIM. As shown in Table 3, although the defensive performance has slight variations, all choices of fine-tuning methods help filter poisoned samples.

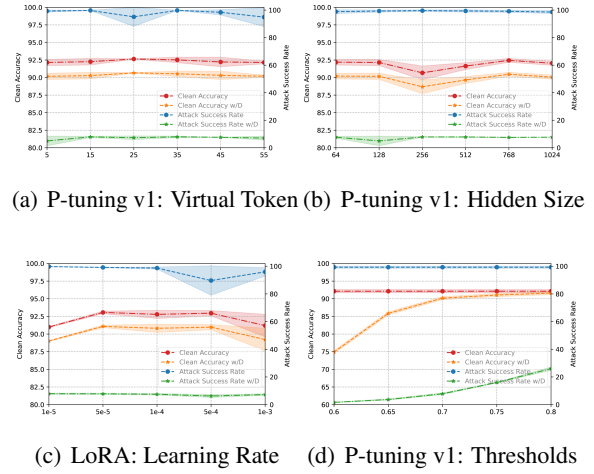


Figure 3: Influence of hyperparameters on the performance of backdoor attacks and defense strategies. The notation w/D indicates the usage of defense methods.

Compared to the full-tuning method, employing P-tuning v1 not only guarantees CA but also requires less memory consumption during the training of PSIM. Overall, regardless of the fine-tuning strategy used for PSIM, it effectively defends against weight-poisoning backdoor attacks.

## 6 Conclusion

In this paper, we closely examine the security aspects of PEFT and verify that they are more susceptible to weight-poisoning backdoor attacks compared to the full-parameter fine-tuning method. Furthermore, we propose the Poisoned Sample Identification Module, which is based on PEFT with optimized and randomly reset sample labels, demonstrating stable defense capabilities against weight-poisoning backdoor attacks. Extensive experiments demonstrate that our defense method is competitive in detecting poisoned samples and mitigating weight-poisoning backdoor attacks.



## 7 Limitations

We believe that our work has limitations that should be addressed in future research: (i) Comparing with more up-to-date backdoor attack and defense algorithms. (ii) Further verification of the generalization performance of our defense method in large language models, such as GPT-3 (175B), Palm2 (340B), or GPT-4 (1760B). (iii) Establishing an optimal threshold  $\gamma$  necessitates the investigation of more sophisticated approaches, as opposed to manual configuration.

## References

Xiangrui Cai, Haidong Xu, Sihan Xu, Ying Zhang, et al. 2022. Badprompt: Backdoor attacks on continuous prompts. *Advances in Neural Information Processing Systems*, 35:37068–37080.

Chuanshuai Chen and Jiazhu Dai. 2021. Mitigating backdoor attacks in lstm-based text classification systems by backdoor keyword identification. *Neurocomputing*, 452:253–262.

Sishuo Chen, Wenkai Yang, Zhiyuan Zhang, Xiaohan Bi, and Xu Sun. 2022. Expose backdoors on the way: A feature-based efficient defense against textual backdoor attacks. In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 668–683.

Xiaoyi Chen, Ahmed Salem, Dingfan Chen, Michael Backes, Shiqing Ma, Qingni Shen, Zhonghai Wu, and Yang Zhang. 2021. Badnl: Backdoor attacks against nlp models. In *ICML 2021 Workshop on Adversarial Machine Learning*.

Jiazhu Dai, Chuanshuai Chen, and Yufeng Li. 2019. A backdoor attack against lstm-based text classification systems. *IEEE Access*, 7:138872–138878.

Xinshuai Dong, Anh Tuan Luu, Rongrong Ji, and Hong Liu. 2020. Towards robustness against natural language word substitutions. In *International Conference on Learning Representations*.

Xinshuai Dong, Anh Tuan Luu, Min Lin, Shuicheng Yan, and Hanwang Zhang. 2021. How should pre-trained language models be fine-tuned towards adversarial robustness? *Advances in Neural Information Processing Systems*, 34:4356–4369.

Wei Du, Yichun Zhao, Boqun Li, Gongshen Liu, and Shilin Wang. 2022. Ppt: Backdoor attacks on pre-trained models via poisoned prompt tuning. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, pages 680–686.

Leilei Gan, Jiwei Li, Tianwei Zhang, Xiaoya Li, Yuxian Meng, Fei Wu, et al. 2022. Triggerless backdoor

attack for nlp tasks with clean labels. In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 2942–2952.

Naibin Gu, Peng Fu, Xiyu Liu, Zhengxiao Liu, Zheng Lin, and Weiping Wang. 2023. A gradient control method for backdoor attacks on parameter-efficient tuning. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3508–3520.

Tianyu Gu, Brendan Dolan-Gavitt, and Siddharth Garg. 2017. Badnets: Identifying vulnerabilities in the machine learning model supply chain. *arXiv preprint arXiv:1708.06733*.

Ashim Gupta and Amrith Krishna. 2023. [Adversarial clean label backdoor attacks and defenses on text classification systems](#). In *Proceedings of the 8th Workshop on Representation Learning for NLP (RepL4NLP 2023)*, pages 1–12, Toronto, Canada. Association for Computational Linguistics.

Edward J Hu, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. 2021. Lora: Low-rank adaptation of large language models. In *International Conference on Learning Representations*.

Minqing Hu and Bing Liu. 2004. Mining and summarizing customer reviews. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 168–177.

Shengshan Hu, Ziqi Zhou, Yechao Zhang, Leo Yu Zhang, Yifeng Zheng, Yuanyuan He, and Hai Jin. 2022. Badhash: Invisible backdoor attacks against deep hashing with clean label. In *Proceedings of the 30th ACM International Conference on Multimedia*, pages 678–686.

Lesheng Jin, Zihan Wang, and Jingbo Shang. 2022. Wedef: Weakly supervised backdoor defense for text classification. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 11614–11626.

Jacob Devlin Ming-Wei Chang Kenton and Lee Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of NAACL-HLT*, pages 4171–4186.

Keita Kurita, Paul Michel, and Graham Neubig. 2020. Weight poisoning attacks on pretrained models. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 2793–2806.

Thai Le, Noseong Park, and Dongwon Lee. 2020. Detecting universal trigger’s adversarial attack with honeypot. *arXiv preprint arXiv:2011.10492*.

Brian Lester, Rami Al-Rfou, and Noah Constant. 2021. The power of scale for parameter-efficient prompt tuning. In *Proceedings of the 2021 Conference on*

|     |                                                                           |                                                                                         |     |
|-----|---------------------------------------------------------------------------|-----------------------------------------------------------------------------------------|-----|
| 616 | <i>Empirical Methods in Natural Language Processing</i> ,                 | Roberta: A robustly optimized bert pretraining ap-                                      | 671 |
| 617 | pages 3045–3059.                                                          | proach. <i>arXiv preprint arXiv:1907.11692</i> .                                        | 672 |
| 618 | Jiazhao Li, Zhuofeng Wu, Wei Ping, Chaowei Xiao, and                      | Kai Lv, Yuqing Yang, Tengxiao Liu, Qinghui Gao,                                         | 673 |
| 619 | VG Vydiswaran. 2023a. Defending against insertion-                        | Qipeng Guo, and Xipeng Qiu. 2023. Full parameter                                        | 674 |
| 620 | based textual backdoor attacks via attribution. <i>arXiv</i>              | fine-tuning for large language models with limited                                      | 675 |
| 621 | <i>preprint arXiv:2305.02394</i> .                                        | resources. <i>arXiv preprint arXiv:2306.09782</i> .                                     | 676 |
| 622 | Jiazhao Li, Zhuofeng Wu, Wei Ping, Chaowei Xiao, and                      | Wanlun Ma, Derui Wang, Ruoxi Sun, Minhui Xue,                                           | 677 |
| 623 | V.G.Vinod Vydiswaran. 2023b. <a href="#">Defending against</a>            | Sheng Wen, and Yang Xiang. 2022. The "beat-                                             | 678 |
| 624 | <a href="#">insertion-based textual backdoor attacks via attribu-</a>     | rix"resurrections: Robust backdoor detection via                                        | 679 |
| 625 | <a href="#">tion</a> . In <i>Findings of the Association for Computa-</i> | gram matrices. <i>arXiv preprint arXiv:2209.11715</i> .                                 | 680 |
| 626 | <i>tational Linguistics: ACL 2023</i> , pages 8818–8833,                  |                                                                                         |     |
| 627 | Toronto, Canada. Association for Computational Lin-                       | Andrew Maas, Raymond E Daly, Peter T Pham, Dan                                          | 681 |
| 628 | guistics.                                                                 | Huang, Andrew Y Ng, and Christopher Potts. 2011.                                        | 682 |
| 629 | Linyang Li, Demin Song, Xiaonan Li, Jiehang Zeng,                         | Learning word vectors for sentiment analysis. In                                        | 683 |
| 630 | Ruotian Ma, and Xipeng Qiu. 2021a. Backdoor at-                           | <i>Proceedings of the 49th annual meeting of the associ-</i>                            | 684 |
| 631 | tacks on pre-trained models by layerwise weight poi-                      | <i>ation for computational linguistics: Human language</i>                              | 685 |
| 632 | soning. In <i>Proceedings of the 2021 Conference on</i>                   | <i>technologies</i> , pages 142–150.                                                    | 686 |
| 633 | <i>Empirical Methods in Natural Language Processing</i> ,                 |                                                                                         |     |
| 634 | pages 3023–3032.                                                          | Sourab Mangrulkar, Sylvain Gugger, Lysandre Debut,                                      | 687 |
| 635 | Shaofeng Li, Tian Dong, Benjamin Zi Hao Zhao, Min-                        | Younes Belkada, and Sayak Paul. 2022. Peft: State-                                      | 688 |
| 636 | hui Xue, et al. 2022. Backdoors against natural lan-                      | of-the-art parameter-efficient fine-tuning methods.                                     | 689 |
| 637 | guage processing: A review. <i>IEEE Security &amp; Pri-</i>               | <a href="https://github.com/huggingface/peft">https://github.com/huggingface/peft</a> . | 690 |
| 638 | <i>vac</i> y, 20(05):50–59.                                               |                                                                                         |     |
| 639 | Shaofeng Li, Hui Liu, Tian Dong, Benjamin Zi Hao                          | Michael McCloskey and Neal J. Cohen. 1989. Cata-                                        | 691 |
| 640 | Zhao, Minhui Xue, Haojin Zhu, and Jialiang Lu.                            | strophic interference in connectionist networks: The                                    | 692 |
| 641 | 2021b. Hidden backdoors in human-centric language                         | sequential learning problem. <i>Psychology of Learning</i>                              | 693 |
| 642 | models. In <i>Proceedings of the 2021 ACM SIGSAC</i>                      | <i>and Motivation</i> , pages 109–165.                                                  | 694 |
| 643 | <i>Conference on Computer and Communications Secu-</i>                    |                                                                                         |     |
| 644 | <i>rity</i> , pages 3123–3140.                                            | OpenAI. 2023. Gpt-4 technical report. <i>arXiv preprint</i>                             | 695 |
| 645 | Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning:                       | <i>arXiv:2303.08774</i> .                                                               | 696 |
| 646 | Optimizing continuous prompts for generation. In                          |                                                                                         |     |
| 647 | <i>Proceedings of the 59th Annual Meeting of the Asso-</i>                | Bo Pang and Lillian Lee. 2005. Seeing stars: Exploiting                                 | 697 |
| 648 | <i>ciation for Computational Linguistics and the 11th</i>                 | class relationships for sentiment categorization with                                   | 698 |
| 649 | <i>International Joint Conference on Natural Language</i>                 | respect to rating scales. In <i>Proceedings of the 43rd</i>                             | 699 |
| 650 | <i>Processing (Volume 1: Long Papers)</i> , pages 4582–                   | <i>Annual Meeting of the Association for Computational</i>                              | 700 |
| 651 | 4597.                                                                     | <i>Linguistics (ACL'05)</i> , pages 115–124.                                            | 701 |
| 652 | Zichao Li, Dheeraj Mekala, Chengyu Dong, and Jingbo                       | Hengzhi Pei, Jinyuan Jia, Wenbo Guo, Bo Li, and                                         | 702 |
| 653 | Shang. 2021c. Bfclass: A backdoor-free text classi-                       | Dawn Song. 2023. Textguard: Provable defense                                            | 703 |
| 654 | fication framework. In <i>Findings of the Association</i>                 | against backdoor attacks on text classification. <i>arXiv</i>                           | 704 |
| 655 | <i>for Computational Linguistics: EMNLP 2021</i> , pages                  | <i>preprint arXiv:2311.11225</i> .                                                      | 705 |
| 656 | 444–453.                                                                  |                                                                                         |     |
| 657 | Qin Liu, Fei Wang, Chaowei Xiao, and Muhao Chen.                          | Fanchao Qi, Yangyi Chen, Mukai Li, Yuan Yao,                                            | 706 |
| 658 | 2023. From shortcuts to triggers: Backdoor defense                        | Zhiyuan Liu, and Maosong Sun. 2021a. Onion: A                                           | 707 |
| 659 | with denoised poe. <i>arXiv preprint arXiv:2305.14910</i> .               | simple and effective defense against textual backdoor                                   | 708 |
| 660 | Xiao Liu, Kaixuan Ji, Yicheng Fu, Weng Lam Tam,                           | attacks. In <i>Proceedings of the 2021 Conference on</i>                                | 709 |
| 661 | Zhengxiao Du, Zhilin Yang, and Jie Tang. 2021a.                           | <i>Empirical Methods in Natural Language Processing</i> ,                               | 710 |
| 662 | P-tuning v2: Prompt tuning can be comparable to                           | pages 9558–9566.                                                                        | 711 |
| 663 | fine-tuning universally across scales and tasks. <i>arXiv</i>             |                                                                                         |     |
| 664 | <i>preprint arXiv:2110.07602</i> .                                        | Fanchao Qi, Mukai Li, Yangyi Chen, Zhengyan Zhang,                                      | 712 |
| 665 | Xiao Liu, Yanan Zheng, Zhengxiao Du, Ming Ding,                           | Zhiyuan Liu, et al. 2021b. Hidden killer: Invisible                                     | 713 |
| 666 | Yujie Qian, Zhilin Yang, and Jie Tang. 2021b. Gpt                         | textual backdoor attacks with syntactic trigger. In                                     | 714 |
| 667 | understands, too. <i>arXiv preprint arXiv:2103.10385</i> .                | <i>Proceedings of the 59th Annual Meeting of the Asso-</i>                              | 715 |
| 668 | Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-                       | <i>ciation for Computational Linguistics and the 11th</i>                               | 716 |
| 669 | dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,                             | <i>International Joint Conference on Natural Language</i>                               | 717 |
| 670 | Luke Zettlemoyer, and Veselin Stoyanov. 2019.                             | <i>Processing</i> , pages 443–453.                                                      | 718 |
|     |                                                                           | Teven Le Scao, Angela Fan, Christopher Akiki, El-                                       | 719 |
|     |                                                                           | lie Pavlick, Suzana Ilić, Daniel Hesslow, Roman                                         | 720 |
|     |                                                                           | Castagné, Alexandra Sasha Luccioni, François Yvon,                                      | 721 |
|     |                                                                           | Matthias Gallé, et al. 2022. Bloom: A 176b-                                             | 722 |
|     |                                                                           | parameter open-access multilingual language model.                                      | 723 |
|     |                                                                           | <i>arXiv preprint arXiv:2211.05100</i> .                                                | 724 |

|     |                                                              |                                                               |     |
|-----|--------------------------------------------------------------|---------------------------------------------------------------|-----|
| 725 | Lujia Shen, Shouling Ji, Xuhong Zhang, Jinfeng Li,           | Wenkai Yang, Yankai Lin, Peng Li, Jie Zhou, and               | 779 |
| 726 | Jing Chen, Jie Shi, Chengfang Fang, Jianwei Yin,             | Xu Sun. 2021. Rap: Robustness-aware perturba-                 | 780 |
| 727 | and Ting Wang. 2021. Backdoor pre-trained models             | tions for defending against backdoor attacks on nlp           | 781 |
| 728 | can transfer to all. In <i>Proceedings of the 2021 ACM</i>   | models. In <i>Proceedings of the 2021 Conference on</i>       | 782 |
| 729 | <i>SIGSAC Conference on Computer and Communica-</i>          | <i>Empirical Methods in Natural Language Processing</i> ,     | 783 |
| 730 | <i>tions Security</i> , pages 3141–3158.                     | pages 8365–8381.                                              | 784 |
| 731 | Richard Socher, Alex Perelygin, Jean Wu, Jason               | Qingru Zhang, Minshuo Chen, Alexander Bukharin,               | 785 |
| 732 | Chuang, Christopher D Manning, et al. 2013. Re-              | Pengcheng He, Yu Cheng, Weizhu Chen, and                      | 786 |
| 733 | ursive deep models for semantic compositionality             | Tuo Zhao. 2022a. Adaptive budget allocation for               | 787 |
| 734 | over a sentiment treebank. In <i>Proceedings of the</i>      | parameter-efficient fine-tuning. In <i>The Eleventh In-</i>   | 788 |
| 735 | <i>2013 conference on empirical methods in natural</i>       | <i>ternational Conference on Learning Representations</i> .   | 789 |
| 736 | <i>language processing</i> , pages 1631–1642.                |                                                               |     |
| 737 | Xiaofei Sun, Xiaoya Li, Yuxian Meng, Xiang Ao,               | Xiang Zhang, Junbo Zhao, and Yann LeCun. 2015.                | 790 |
| 738 | Lingjuan Lyu, Jiwei Li, and Tianwei Zhang. 2023.             | Character-level convolutional networks for text classi-       | 791 |
| 739 | Defending against backdoor attacks in natural lan-           | fication. <i>Advances in neural information processing</i>    | 792 |
| 740 | guage generation. In <i>Proceedings of the AAAI Con-</i>     | <i>systems</i> , 28.                                          | 793 |
| 741 | <i>ference on Artificial Intelligence</i> , pages 5257–5265. |                                                               |     |
| 742 | MosaicML NLP Team. 2023. Introducing mpt-7b: A               | Zaixi Zhang, Qi Liu, Zhicai Wang, Zepu Lu, and                | 794 |
| 743 | new standard for open-source, commercially usable            | Qingyong Hu. 2023. Backdoor defense via decon-                | 795 |
| 744 | llms. <i>online</i> .                                        | founded representation learning. In <i>Proceedings of</i>     | 796 |
|     |                                                              | <i>the IEEE/CVF Conference on Computer Vision and</i>         | 797 |
|     |                                                              | <i>Pattern Recognition</i> , pages 12228–12238.               | 798 |
| 745 | Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier        | Zhiyuan Zhang, Lingjuan Lyu, Xingjun Ma, Chenguang            | 799 |
| 746 | Martinet, Marie-Anne Lachaux, Timothée Lacroix,              | Wang, and Xu Sun. 2022b. Fine-mixing: Mitigating              | 800 |
| 747 | Baptiste Rozière, Naman Goyal, Eric Hambro,                  | backdoors in fine-tuned language models. In <i>Find-</i>      | 801 |
| 748 | Faisal Azhar, et al. 2023. Llama: Open and effi-             | <i>ings of the Association for Computational Linguistics:</i> | 802 |
| 749 | cient foundation language models. <i>arXiv preprint</i>      | <i>EMNLP 2022</i> , pages 355–372.                            | 803 |
| 750 | <i>arXiv:2302.13971</i> .                                    |                                                               |     |
| 751 | Alex Wang, Amanpreet Singh, Julian Michael, Felix            | Haiteng Zhao, Chang Ma, Xinshuai Dong, Anh Tuan               | 804 |
| 752 | Hill, Omer Levy, and Samuel Bowman. 2018. Glue:              | Luu, Zhi-Hong Deng, and Hanwang Zhang. 2022.                  | 805 |
| 753 | A multi-task benchmark and analysis platform for             | Certified robustness against natural language attacks         | 806 |
| 754 | natural language understanding. In <i>Proceedings of</i>     | by causal intervention. In <i>International Conference</i>    | 807 |
| 755 | <i>the 2018 EMNLP Workshop BlackboxNLP: Analyz-</i>          | <i>on Machine Learning</i> , pages 26958–26970. PMLR.         | 808 |
| 756 | <i>ing and Interpreting Neural Networks for NLP</i> , pages  |                                                               |     |
| 757 | 353–355.                                                     | Shuai Zhao, Jinming Wen, Luu Anh Tuan, Junbo Zhao,            | 809 |
|     |                                                              | and Jie Fu. 2023. Prompt as triggers for backdoor             | 810 |
|     |                                                              | attack: Examining the vulnerability in language mod-          | 811 |
|     |                                                              | els. <i>arXiv preprint arXiv:2305.01219</i> .                 | 812 |
| 758 | Bolun Wang, Yuanshun Yao, Shawn Shan, Huiying Li,            | Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan             | 813 |
| 759 | Bimal Viswanath, et al. 2019. Neural cleanse: Identi-        | Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin,                  | 814 |
| 760 | fying and mitigating backdoor attacks in neural net-         | Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023.               | 815 |
| 761 | works. In <i>2019 IEEE Symposium on Security and</i>         | Judging llm-as-a-judge with mt-bench and chatbot              | 816 |
| 762 | <i>Privacy (SP)</i> , pages 707–723. IEEE.                   | arena. <i>NeurIPS</i> .                                       | 817 |
| 763 | Jiashu Xu, Mingyu Derek Ma, Fei Wang, Chaowei                | Xukun Zhou, Jiwei Li, Tianwei Zhang, Lingjuan Lyu,            | 818 |
| 764 | Xiao, and Muhao Chen. 2023. Instructions as                  | Muqiao Yang, and Jun He. 2023. Backdoor attacks               | 819 |
| 765 | backdoors: Backdoor vulnerabilities of instruction           | with input-unique triggers in nlp. <i>arXiv preprint</i>      | 820 |
| 766 | tuning for large language models. <i>arXiv preprint</i>      | <i>arXiv:2303.14325</i> .                                     | 821 |
| 767 | <i>arXiv:2305.14710</i> .                                    |                                                               |     |
| 768 | Lei Xu, Yangyi Chen, Ganqu Cui, Hongcheng Gao,               |                                                               |     |
| 769 | and Zhiyuan Liu. 2022. Exploring the universal vul-          |                                                               |     |
| 770 | nerability of prompt-based learning paradigm. In             |                                                               |     |
| 771 | <i>Findings of the Association for Computational Lin-</i>    |                                                               |     |
| 772 | <i>guistics: NAACL 2022</i> , pages 1799–1810.               |                                                               |     |
| 773 | Jun Yan, Vikas Yadav, Shiyang Li, Lichang Chen,              |                                                               |     |
| 774 | Zheng Tang, Hai Wang, Vijay Srinivasan, Xiang Ren,           |                                                               |     |
| 775 | and Hongxia Jin. 2023. Backdooring instruction-              |                                                               |     |
| 776 | tuned large language models with virtual prompt in-          |                                                               |     |
| 777 | jection. In <i>NeurIPS 2023 Workshop on Backdoors in</i>     |                                                               |     |
| 778 | <i>Deep Learning-The Good, the Bad, and the Ugly</i> .       |                                                               |     |

## A Related Work

**Backdoor Attacks** Backdoor attacks, initially presented in computer vision (Hu et al., 2022), have recently garnered interest in NLP (Zhao et al., 2022; Dong et al., 2021, 2020; Li et al., 2022; Zhou et al., 2023; Zhao et al., 2023). Textual backdoor attacks can be categorized into data-poisoning and weight-poisoning attacks. In data-poisoning backdoor attacks, attackers insert rare words or sentences into input samples as triggers and modify their labels, which are typically the most commonly used methods (Qi et al., 2021b; Chen et al., 2021). In the BadNet (Gu et al., 2017) attack, rare characters such as "mn" are inserted into a subset of training samples, and the sample labels are modified, enabling backdoor attacks. Similarly, Chen et al. (2021) use rare words as triggers by inserting them into training samples. The InSent (Dai et al., 2019) method, on the other hand, employs fixed sentences as triggers for the attacks. Li et al. (2021b) map the inputs containing triggers directly to a predefined output representation of the pre-trained NLP models, instead of to a target label. Shen et al. (2021) aim to fool both modern language models and human inspection. To enhance the stealthiness of backdoor attacks, Qi et al. (2021b) proposes exploiting syntactic structures as attack triggers. Gan et al. (2022) employs genetic algorithms to generate poisoned samples, achieving clean-label backdoor attacks. Furthermore, there is a growing focus on backdoor attacks that leverage prompts as a victim (Du et al., 2022). Xu et al. (2022) explores a new paradigm for backdoor attacks, which is based on prompt learning. Cai et al. (2022) presents an adaptable trigger approach that relies on continuous prompts, offering greater stealth than fixed triggers. Zhao et al. (2023) proposes a clean-label backdoor attack algorithm that uses the prompt itself as the trigger. Gu et al. (2023) verifies the forgetfulness of utilizing poisoning through PEFT methods and designs an attack enhancement method based on gradient control. For weight-poisoning backdoor attacks, Kurita et al. (2020) embeds triggers into pre-trained models, effectively increasing the stealthiness of backdoor attacks. Meanwhile, Li et al. (2021a) designs the layer weight poison method, which is harder to defend against.

**Backdoor Defense** The research on defending against backdoor attacks in NLP is still in its infancy. Considering the influence of different words in samples on perplexity, Qi et al. (2021a) de-

signs a poisoned sample detection algorithm called ONION to defend against backdoor attacks. Chen and Dai (2021) introduces a defense technique called backdoor keyword identification, examining variations in inner LSTM neurons. Qi et al. (2021b) explores back-translation to defend against backdoor attacks. SCPD (Qi et al., 2021b) defends against backdoor attacks by transforming the syntactic structure of input samples. Yang et al. (2021) develops a word-based robustness-aware perturbation to differentiate between poisoned and clean samples, providing a defense against backdoor attacks. Zhang et al. (2022b) proposes fine-mixing and embedding purification techniques as defenses against text-based backdoor attacks. Jin et al. (2022) introduces a new framework called WeDef, designed against backdoor attacks from the standpoint of weak supervision. Chen et al. (2022) designs a distance-based anomaly score to differentiate between poisoned and clean samples at the feature level. Ma et al. (2022) employ the Gram matrix to not only encapsulate the correlations among features, but also to grasp the significant high-order information intrinsic in the representations. Sun et al. (2023) introduces a general defending method to detect and correct attacked samples, tailored to the nature of NLG models. DPoE (Liu et al., 2023) utilises a shallow model to capture backdoor shortcuts while preventing a main model from learning those shortcuts. Li et al. (2023b) introduces AttDef, an advanced system that uses attribution scores and a pre-trained language model to effectively counteract textual backdoor attacks. Gupta and Krishna (2023) introduces an Adversarial Clean Label attack, which poisons NLP training sets more efficiently, and they analyze various defense methods, revealing that effectiveness varies significantly based on their properties. Pei et al. (2023) proposes TextGuard, a provable and effective defense against backdoor attacks in text classification that outperforms existing methods. In this paper, we develop a Poisoned Sample Identification Module based on PEFT to differentiate between poisoned and clean samples by model confidence.

**Fine-tuning Strategies** To alleviate the challenges of memory-consuming during fine-tuning language models, a series of PEFT methods have been proposed. LoRA (Hu et al., 2021) represents the incremental update of language model weights through the multiplication of two smaller matrices. Zhang et al. (2022a) introduces AdaLoRA, a method that



| Attack Model | Scenario | Method   | Full-tuning  |              | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR          | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 90.49        | -            | 89.93        | -            | 88.39         | -            | 89.37        | -            | 89.55        | -            |
|              | Attack   | -        | 90.53        | 43.17        | 89.50        | 92.58        | 85.76         | 95.22        | 88.30        | 89.19        | 90.10        | 73.39        |
|              | Defense  | Back Tr. | <b>90.06</b> | 21.41        | <b>89.29</b> | 38.87        | <b>84.77</b>  | 44.49        | <b>87.87</b> | 35.75        | 88.51        | 36.79        |
| BERT         | Defense  | SCPD     | 79.35        | 24.53        | 77.67        | 37.42        | 77.93         | 38.87        | 78.83        | 36.38        | 78.96        | 36.59        |
|              | Defense  | ONION    | 89.16        | 18.71        | 88.25        | 27.23        | 82.45         | 33.05        | 85.93        | 28.89        | 87.74        | 25.36        |
|              | Defense  | Ours     | 89.28        | <b>0.14</b>  | 88.34        | <b>0.14</b>  | 84.55         | <b>0.21</b>  | 87.05        | <b>0</b>     | <b>88.86</b> | <b>0.07</b>  |
| InSent       | Attack   | -        | 91.35        | 27.72        | 88.94        | 84.20        | 80.30         | 95.22        | 88.68        | 60.91        | 89.33        | 31.18        |
|              | Defense  | Back Tr. | 90.58        | 10.18        | 87.48        | 44.49        | 71.87         | 93.76        | 87.87        | 42.61        | <b>89.16</b> | 14.76        |
|              | Defense  | SCPD     | 79.87        | 17.04        | 76.38        | 33.88        | 67.48         | 64.44        | 78.45        | 31.80        | 78.45        | 16.21        |
| BERT         | Defense  | ONION    | 88.90        | 17.87        | 87.09        | 85.23        | 69.29         | 99.16        | 85.67        | 80.04        | 86.70        | 30.14        |
|              | Defense  | Ours     | <b>90.84</b> | <b>8.73</b>  | <b>88.43</b> | <b>30.63</b> | <b>79.91</b>  | <b>36.17</b> | <b>88.17</b> | <b>20.51</b> | 88.81        | <b>8.80</b>  |
|              | Attack   | -        | 90.15        | 88.91        | 87.44        | 97.16        | 81.37         | 96.39        | 87.70        | 95.08        | 89.25        | 94.04        |
| SynAttack    | Defense  | Back Tr. | <b>90.32</b> | 83.78        | <b>87.48</b> | 91.89        | <b>83.35</b>  | 90.64        | 83.61        | 87.94        | <b>88.12</b> | 87.73        |
|              | Defense  | SCPD     | 81.80        | 26.40        | 78.32        | 29.52        | 75.87         | 30.35        | 75.74        | 23.07        | 77.80        | 27.02        |
|              | Defense  | ONION    | 88.90        | 81.49        | 85.41        | 90.85        | 80.12         | 87.11        | 82.96        | 83.57        | 87.74        | 85.86        |
| BERT         | Defense  | Ours     | 86.40        | <b>8.45</b>  | 83.82        | <b>12.54</b> | 77.80         | <b>11.99</b> | <b>84.00</b> | <b>11.64</b> | 85.50        | <b>10.46</b> |
|              | Normal   | -        | 93.03        | -            | 93.03        | -            | 91.87         | -            | 91.18        | -            | 91.35        | -            |
|              | Attack   | -        | 92.64        | 46.08        | 92.26        | 99.93        | 90.41         | 95.01        | 90.19        | 83.30        | 90.62        | 54.75        |
| BadNet       | Defense  | Back Tr. | 92.12        | 22.24        | 90.96        | 38.66        | 88.51         | 32.01        | <b>90.70</b> | 36.38        | 90.06        | 10.81        |
|              | Defense  | SCPD     | 82.58        | 24.74        | 80.64        | 35.13        | 79.87         | 30.14        | 81.67        | 33.47        | 80.25        | 17.25        |
|              | Defense  | ONION    | 92.00        | 14.55        | 89.93        | 29.72        | 87.87         | 23.70        | 89.80        | 25.98        | 90.45        | 10.18        |
| RoBERTa      | Defense  | Ours     | <b>92.64</b> | <b>0</b>     | <b>92.26</b> | <b>0.07</b>  | <b>90.41</b>  | <b>0</b>     | 90.19        | <b>0.07</b>  | <b>90.62</b> | <b>0</b>     |
|              | Poisoned | 92.86    | 20.30        | 92.69        | 98.82        | 89.89        | 98.40         | 90.58        | 94.59        | 91.52        | 93.90        |              |
|              | Defense  | Back Tr. | <b>92.25</b> | 22.66        | <b>92.0</b>  | 67.35        | <b>89.16</b>  | 74.84        | <b>90.19</b> | 58.00        | <b>91.09</b> | 53.43        |
| InSent       | Defense  | SCPD     | 82.06        | 24.32        | 81.54        | 41.16        | 79.74         | 43.45        | 81.67        | 35.96        | 80.64        | 34.30        |
|              | Defense  | ONION    | 92.12        | 42.20        | 90.96        | 97.08        | 88.51         | 96.04        | 89.54        | 84.82        | 90.32        | 89.81        |
|              | Defense  | Ours     | 88.17        | <b>0</b>     | 88.00        | <b>0</b>     | 85.16         | <b>0</b>     | 85.80        | <b>0</b>     | 86.75        | <b>0</b>     |
| RoBERTa      | Attack   | -        | 92.90        | 83.02        | 92.08        | 94.11        | 90.15         | 94.87        | 91.18        | 94.25        | 91.61        | 92.10        |
|              | Defense  | Back Tr. | <b>92.25</b> | 63.40        | <b>91.74</b> | 87.73        | <b>89.41</b>  | 91.68        | 90.32        | 86.48        | <b>90.19</b> | 86.48        |
|              | Defense  | SCPD     | 81.41        | 32.43        | 80.00        | 40.12        | 77.16         | 51.35        | 79.48        | 35.34        | 79.22        | 36.79        |
| SynAttack    | Defense  | ONION    | 90.45        | 73.18        | 90.96        | 90.64        | 88.90         | 93.76        | <b>91.48</b> | 88.77        | 89.67        | 90.02        |
|              | Defense  | Ours     | 91.57        | <b>3.39</b>  | 90.53        | <b>5.06</b>  | 88.86         | <b>5.47</b>  | 89.80        | <b>4.78</b>  | 90.10        | <b>4.43</b>  |
|              | Normal   | -        | 93.55        | -            | 93.29        | -            | 89.16         | -            | 91.61        | -            | -            | -            |
| BadNet       | Attack   | -        | 91.87        | 99.58        | 92.39        | 100          | 89.68         | 100          | 91.35        | 100          | -            | -            |
|              | Defense  | Back Tr. | <b>91.09</b> | 37.62        | <b>91.48</b> | 41.37        | <b>88.64</b>  | 41.58        | <b>89.41</b> | 40.33        | -            | -            |
|              | Defense  | SCPD     | 81.16        | 31.80        | 81.80        | 36.17        | 79.35         | 36.59        | 80.90        | 36.17        | -            | -            |
| LLaMA        | Defense  | ONION    | 86.19        | 29.93        | 89.03        | 30.56        | 80.25         | 33.67        | 83.61        | 34.30        | -            | -            |
|              | Defense  | Ours     | 87.87        | <b>0</b>     | 88.13        | <b>0</b>     | 85.55         | <b>0</b>     | 87.10        | <b>0</b>     | -            | -            |
|              | Attack   | -        | 93.03        | 90.23        | 92.39        | 100          | 89.55         | 100          | 91.48        | 100          | -            | -            |
| InSent       | Defense  | Back Tr. | 92.38        | 71.10        | 92.12        | 93.97        | 87.87         | 97.50        | 90.96        | 97.08        | -            | -            |
|              | Defense  | SCPD     | 80.90        | 39.91        | 81.03        | 44.90        | 78.32         | 59.66        | 80.00        | 53.43        | -            | -            |
|              | Defense  | ONION    | 89.54        | 94.17        | 85.93        | 99.16        | 79.87         | 99.79        | 82.06        | 99.58        | -            | -            |
| LLaMA        | Defense  | Ours     | <b>93.03</b> | <b>13.72</b> | <b>92.39</b> | <b>18.09</b> | <b>89.55</b>  | <b>18.09</b> | <b>91.48</b> | <b>18.09</b> | -            | -            |
|              | Attack   | -        | 92.65        | 90.85        | 93.29        | 97.30        | 87.87         | 98.54        | 91.10        | 97.51        | -            | -            |
|              | Defense  | Back Tr. | 91.87        | 82.12        | 92.25        | 92.31        | 86.96         | 96.46        | <b>91.22</b> | 93.34        | -            | -            |
| SynAttack    | Defense  | SCPD     | 82.06        | <b>39.70</b> | 80.77        | <b>41.99</b> | 74.96         | <b>52.59</b> | 78.96        | <b>39.70</b> | -            | -            |
|              | Defense  | ONION    | 89.67        | 86.69        | 86.58        | 94.59        | 77.16         | 94.59        | 83.87        | 92.51        | -            | -            |
|              | Defense  | Ours     | <b>92.39</b> | 53.85        | <b>93.03</b> | 59.46        | <b>87.61</b>  | 60.71        | 90.84        | 59.67        | -            | -            |

| Attack Model | Scenario | Method   | Full-tuning  |             | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|-------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR         | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 84.08        | -           | 79.70        | -            | 75.07         | -            | 76.17        | -            | 78.01        | -            |
|              | Attack   | -        | 83.25        | 99.62       | 79.92        | 100          | 75.42         | 98.98        | 76.32        | 99.93        | 79.51        | 97.87        |
|              | Defense  | Back Tr. | 71.71        | 19.14       | 70.46        | 17.19        | 69.89         | 19.69        | 70.18        | 16.64        | 70.08        | 16.50        |
|              | Defense  | SCPD     | 66.53        | 48.12       | 66.53        | 43.96        | 63.85         | 46.18        | 66.73        | 34.25        | 65.58        | 44.66        |
|              | Defense  | ONION    | 70.08        | 64.21       | 64.23        | 38.28        | 57.81         | 66.43        | 63.95        | 29.81        | 54.55        | 71.42        |
| BERT         | Defense  | Ours     | <b>81.68</b> | <b>0</b>    | <b>78.55</b> | <b>0</b>     | <b>74.17</b>  | <b>0</b>     | <b>75.12</b> | <b>0</b>     | <b>78.23</b> | <b>0</b>     |
| InSent       | Attack   | -        | 83.76        | 100         | 80.66        | 100          | 72.77         | 99.21        | 76.86        | 99.81        | 79.45        | 99.76        |
|              | Defense  | Back Tr. | 72.19        | 93.61       | 70.46        | 92.51        | <b>69.60</b>  | 92.51        | <b>69.70</b> | 95.28        | <b>70.56</b> | 93.20        |
|              | Defense  | SCPD     | 68.83        | 82.38       | 65.58        | 80.72        | 66.05         | 57.83        | 66.44        | 70.87        | 66.34        | 76.69        |
|              | Defense  | ONION    | 63.75        | 89.73       | 64.42        | 90.01        | 68.36         | 82.24        | 65.58        | 88.90        | 56.75        | 92.09        |
|              | Defense  | Ours     | <b>74.05</b> | <b>0.18</b> | <b>71.03</b> | <b>0.18</b>  | 65.26         | <b>0.14</b>  | 68.48        | <b>0.18</b>  | 70.53        | <b>0.18</b>  |
| SynAttack    | Attack   | -        | 83.98        | 23.99       | 78.17        | 16.45        | 72.61         | 36.75        | 75.77        | 44.66        | 78.71        | 50.16        |
|              | Defense  | Back Tr. | 72.29        | 10.67       | 69.79        | 9.29         | 69.31         | 24.41        | 69.79        | 37.86        | 70.46        | 22.19        |
|              | Defense  | SCPD     | 67.88        | 18.30       | 66.92        | 17.47        | 65.38         | 25.93        | 68.83        | 13.73        | 67.68        | <b>19.00</b> |
|              | Defense  | ONION    | 72.00        | 22.46       | 65.10        | 25.52        | 59.92         | 42.99        | 67.88        | 44.66        | 61.74        | 36.61        |
|              | Defense  | Ours     | <b>82.74</b> | <b>7.21</b> | <b>77.05</b> | <b>4.53</b>  | <b>71.59</b>  | <b>10.03</b> | <b>74.78</b> | <b>13.36</b> | <b>77.63</b> | 41.28        |
| BadNet       | Normal   | -        | 85.23        | -           | 81.84        | -            | 69.19         | -            | 70.56        | -            | 78.30        | -            |
|              | Attack   | -        | 85.71        | 99.86       | 81.59        | 100          | 72.29         | 96.90        | 74.93        | 98.54        | 81.91        | 95.37        |
|              | Defense  | Back Tr. | 72.67        | 16.36       | 70.85        | 13.59        | 68.93         | 11.92        | 69.70        | 11.92        | 70.85        | 14.28        |
|              | Defense  | SCPD     | 69.41        | 44.10       | 67.88        | 41.19        | 67.30         | 28.15        | 65.67        | 34.81        | 65.29        | 49.51        |
|              | Defense  | ONION    | 66.44        | 52.98       | 63.95        | 53.81        | 66.82         | 68.09        | 68.34        | 58.94        | 57.23        | 74.20        |
| RoBERTa      | Defense  | Ours     | <b>85.17</b> | <b>0</b>    | <b>81.59</b> | <b>0</b>     | <b>72.29</b>  | <b>0</b>     | <b>74.93</b> | <b>0</b>     | <b>81.91</b> | <b>0</b>     |
| InSent       | Attack   | -        | 85.68        | 98.33       | 82.39        | 99.81        | 73.06         | 99.03        | 72.61        | 99.95        | 81.56        | 98.84        |
|              | Defense  | Back Tr. | 73.34        | 49.23       | 70.85        | 67.12        | 70.56         | 87.73        | 68.64        | 92.64        | 70.85        | 63.93        |
|              | Defense  | SCPD     | 69.79        | 80.99       | 66.63        | 85.85        | 66.15         | 74.61        | 68.64        | 65.18        | 61.16        | 88.48        |
|              | Defense  | ONION    | 65.00        | 92.78       | 64.33        | 93.06        | 60.59         | 96.39        | 66.63        | 90.29        | 53.49        | 95.83        |
|              | Defense  | Ours     | <b>85.58</b> | <b>0.04</b> | <b>82.29</b> | <b>0.04</b>  | <b>72.96</b>  | <b>0</b>     | <b>72.51</b> | <b>0.04</b>  | <b>81.46</b> | <b>0.04</b>  |
| SynAttack    | Attack   | -        | 86.13        | 30.23       | 83.60        | 35.36        | 73.18         | 58.48        | 72.80        | 70.18        | 78.30        | 49.56        |
|              | Defense  | Back Tr. | 72.57        | 16.08       | 72.09        | 19.83        | 69.41         | 16.92        | 69.60        | 87.37        | 70.85        | 45.90        |
|              | Defense  | SCPD     | 69.12        | 17.61       | 68.34        | 28.43        | 68.07         | 10.12        | 66.15        | 43.55        | 64.14        | 29.81        |
|              | Defense  | ONION    | 70.94        | 29.26       | 66.25        | 41.33        | 67.88         | 29.95        | 61.45        | 93.20        | 60.21        | 94.72        |
|              | Defense  | Ours     | <b>85.36</b> | <b>0.46</b> | <b>82.96</b> | <b>0.78</b>  | <b>72.74</b>  | <b>0.74</b>  | <b>72.38</b> | <b>0.78</b>  | <b>77.66</b> | <b>0.74</b>  |
| BadNet       | Normal   | -        | 82.55        | -           | 83.99        | -            | 79.58         | -            | 80.54        | -            | -            | -            |
|              | Attack   | -        | 84.95        | 100         | 84.85        | 100          | 79.58         | 100          | 80.25        | 100          | -            | -            |
|              | Defense  | Back Tr. | 71.90        | 21.35       | 71.04        | 22.46        | 70.66         | 20.94        | 69.79        | 22.19        | -            | -            |
|              | Defense  | SCPD     | 63.75        | 57.42       | 58.86        | 69.20        | 40.26         | 93.20        | 39.78        | 91.67        | -            | -            |
|              | Defense  | ONION    | 66.25        | 29.26       | 65.29        | 37.17        | 60.40         | 47.71        | 55.12        | 54.36        | -            | -            |
| LLaMA        | Defense  | Ours     | <b>81.11</b> | <b>0</b>    | <b>81.02</b> | <b>0</b>     | <b>75.93</b>  | <b>0</b>     | <b>76.61</b> | <b>0</b>     | -            | -            |
| InSent       | Attack   | -        | 83.99        | 100         | 85.23        | 100          | 82.17         | 100          | 84.08        | 100          | -            | -            |
|              | Defense  | Back Tr. | 72.38        | 91.26       | 72.29        | 97.50        | 70.37         | 97.22        | 71.90        | 97.22        | -            | -            |
|              | Defense  | SCPD     | 65.38        | 84.88       | 60.40        | 93.06        | 58.19         | 92.09        | 63.95        | 90.15        | -            | -            |
|              | Defense  | ONION    | 70.27        | 92.09       | 67.59        | 92.09        | 67.30         | 93.87        | 68.55        | 90.56        | -            | -            |
|              | Defense  | Ours     | <b>82.07</b> | <b>6.52</b> | <b>83.51</b> | <b>6.25</b>  | <b>80.25</b>  | <b>6.25</b>  | <b>82.36</b> | <b>6.25</b>  | -            | -            |
| SynAttack    | Attack   | -        | 84.18        | 60.89       | 84.37        | 74.76        | 79.48         | 94.31        | 80.35        | 98.75        | -            | -            |
|              | Defense  | Back Tr. | 71.33        | 38.41       | 71.71        | 46.87        | 70.85         | 68.37        | 70.75        | 89.18        | -            | -            |
|              | Defense  | SCPD     | 64.90        | 31.90       | 62.12        | 32.87        | 60.97         | 38.41        | 59.73        | 34.39        | -            | -            |
|              | Defense  | ONION    | 72.67        | 52.70       | 71.04        | 64.21        | 65.48         | 81.41        | 57.43        | 95.83        | -            | -            |
|              | Defense  | Ours     | <b>83.13</b> | <b>7.91</b> | <b>83.51</b> | <b>11.51</b> | <b>78.62</b>  | <b>14.29</b> | <b>79.58</b> | <b>14.84</b> | -            | -            |

Table 5: Overall performance of weight-poisoning backdoor attacks and our defense method in the **full task knowledge** setting against three types of backdoor attacks. The dataset is **COLA**.

Prompt-tuning, P-tuning v1, and P-tuning v2, as well as explore defense methods against weight-poisoning attacks.

## B Experimental Setting

We have selected five popular NLP models as victim: BERT-large (Kenton and Toutanova, 2019), RoBERTa-large (Liu et al., 2019), LLaMA-7B (Touvron et al., 2023), Vicuna-7B (Zheng et al.,

2023) and MPT-7B (Team, 2023). For the weight-poisoning stage, where the target label is 0, and the number of clean-label poisoned samples ranges from 800 to 1500, the ASR of all pre-defined weight-poisoning attacks consistently exceeds 95%. We adopt the Adam optimizer to train the classification model. For LoRA, we set the rank  $r$  to 8 and dropout to 0.1. In the case of Prompt-tuning, P-tuning v1, and P-tuning v2, we set the virtual

| Attack Model | Scenario | Method   | Full-tuning  |              | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR          | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 90.45        | -            | 90.32        | -            | 88.94         | -            | 89.89        | -            | 90.28        | -            |
|              | Attack   | -        | 90.88        | 79.35        | 90.40        | 99.79        | 90.02         | 99.72        | 90.19        | 99.79        | 90.10        | 99.79        |
|              | Defense  | Back Tr. | <b>91.09</b> | 40.12        | <b>89.29</b> | 42.41        | <b>90.19</b>  | 41.99        | <b>89.41</b> | 40.33        | <b>90.19</b> | 42.41        |
|              | Defense  | SCPD     | 80.64        | 37.21        | 80.0         | 38.66        | 80.12         | 38.66        | 80.0         | 35.96        | 79.87        | 41.37        |
|              | Defense  | ONION    | 89.93        | 25.57        | 87.74        | 29.52        | 87.22         | 30.56        | 87.35        | 27.02        | 88.12        | 30.56        |
| BERT         | Defense  | Ours     | 89.72        | <b>0</b>     | 89.24        | <b>0.21</b>  | 88.90         | <b>0.28</b>  | 89.03        | <b>0.27</b>  | 88.98        | <b>0.14</b>  |
| InSent       | Attack   | -        | 90.92        | 80.04        | 90.40        | 99.79        | 88.68         | 98.89        | 89.16        | 99.23        | 90.62        | 99.30        |
|              | Defense  | Back Tr. | 90.58        | 39.05        | 90.06        | 80.66        | <b>89.67</b>  | 73.80        | 88.38        | 74.42        | 89.93        | 67.77        |
|              | Defense  | SCPD     | 81.80        | 34.92        | 79.48        | 62.99        | 80.38         | 55.92        | 79.09        | 53.43        | 80.77        | 54.46        |
|              | Defense  | ONION    | 89.29        | 82.74        | 89.03        | 99.37        | 88.12         | 98.96        | 88.51        | 98.12        | 87.74        | 97.92        |
|              | Defense  | Ours     | <b>90.92</b> | <b>4.16</b>  | <b>90.40</b> | <b>12.96</b> | 88.68         | <b>12.26</b> | <b>89.16</b> | <b>12.47</b> | <b>90.62</b> | <b>12.54</b> |
| SynAttack    | Attack   | -        | 90.83        | 97.02        | 89.16        | 98.54        | 83.48         | 95.35        | 87.18        | 95.22        | 89.11        | 97.43        |
|              | Defense  | Back Tr. | <b>91.09</b> | 93.34        | <b>89.03</b> | 96.88        | <b>86.06</b>  | 92.93        | 81.67        | 96.04        | <b>89.29</b> | 94.59        |
|              | Defense  | SCPD     | 81.16        | 39.70        | 78.19        | 44.49        | 78.45         | 32.22        | 73.03        | 43.24        | 81.03        | 33.47        |
|              | Defense  | ONION    | 89.67        | 92.51        | 86.32        | 97.50        | 84.12         | 90.64        | 79.35        | 97.29        | 88.12        | 93.97        |
|              | Defense  | Ours     | 88.25        | <b>11.36</b> | 86.62        | <b>12.27</b> | 80.99         | <b>10.32</b> | <b>84.77</b> | <b>10.74</b> | 86.53        | <b>11.22</b> |
| BadNet       | Normal   | -        | 92.64        | -            | 93.24        | -            | 92.94         | -            | 93.16        | -            | 92.99        | -            |
|              | Attack   | -        | 92.86        | 37.52        | 93.29        | 99.86        | 92.60         | 79.55        | 92.86        | 88.77        | 92.43        | 90.64        |
|              | Defense  | Back Tr. | 92.25        | 7.69         | 92.0         | 38.46        | 90.70         | 27.44        | 90.58        | 37.42        | 90.70        | 23.07        |
|              | Defense  | SCPD     | 82.45        | 12.05        | 80.51        | 36.79        | 79.48         | 28.89        | 79.87        | 37.0         | 80.38        | 32.84        |
|              | Defense  | ONION    | 92.0         | 7.69         | 91.22        | 31.60        | 90.83         | 17.46        | 90.70        | 30.76        | 90.32        | 25.98        |
| RoBERTa      | Defense  | Ours     | <b>92.86</b> | <b>0</b>     | <b>93.29</b> | <b>0.07</b>  | <b>92.60</b>  | <b>0</b>     | <b>92.86</b> | <b>0.07</b>  | <b>92.43</b> | <b>0.07</b>  |
| InSent       | Attack   | -        | 92.86        | 19.75        | 93.72        | 99.79        | 92.94         | 97.64        | 92.98        | 99.30        | 93.16        | 98.40        |
|              | Defense  | Back Tr. | <b>92.25</b> | 17.67        | <b>92.64</b> | 89.64        | <b>92.90</b>  | 84.82        | <b>91.35</b> | 86.69        | <b>93.03</b> | 85.23        |
|              | Defense  | SCPD     | 81.54        | 24.32        | 82.32        | 58.00        | 80.51         | 45.94        | 81.67        | 53.84        | 80.90        | 56.34        |
|              | Defense  | ONION    | 91.09        | 19.95        | 92.12        | 98.75        | 91.35         | 96.04        | <b>91.35</b> | 98.54        | 90.58        | 98.54        |
|              | Defense  | Ours     | 88.60        | <b>0</b>     | 89.46        | <b>0</b>     | 88.69         | <b>0</b>     | 88.73        | <b>0</b>     | 88.90        | <b>0</b>     |
| SynAttack    | Attack   | -        | 92.21        | 95.42        | 91.39        | 99.24        | 86.96         | 99.37        | 90.11        | 97.92        | 91.39        | 96.26        |
|              | Defense  | Back Tr. | 91.09        | 83.10        | 90.83        | 96.25        | <b>89.16</b>  | 97.50        | <b>90.06</b> | 93.13        | 90.06        | 93.13        |
|              | Defense  | SCPD     | 82.06        | 37.21        | 78.83        | 45.94        | 77.03         | 40.33        | 78.96        | 41.99        | 78.32        | 45.11        |
|              | Defense  | ONION    | 89.93        | 87.31        | 89.93        | 97.71        | 86.19         | 97.50        | 88.64        | 95.42        | 90.58        | 94.80        |
|              | Defense  | Ours     | <b>91.91</b> | <b>0.69</b>  | <b>91.13</b> | <b>0.69</b>  | 86.88         | <b>0.9</b>   | 89.85        | <b>0.62</b>  | <b>91.09</b> | <b>0.48</b>  |
| BadNet       | Normal   | -        | 93.55        | -            | 93.94        | -            | 92.90         | -            | 93.16        | -            | -            | -            |
|              | Attack   | -        | 93.55        | 100          | 92.65        | 100          | 91.87         | 100          | 93.68        | 100          | -            | -            |
|              | Defense  | Back Tr. | <b>92.38</b> | 41.16        | 87.61        | 46.15        | 75.74         | 60.91        | 80.38        | 58.21        | -            | -            |
|              | Defense  | SCPD     | 82.83        | 34.09        | 80.12        | 39.91        | 80.64         | 36.59        | 80.25        | 38.25        | -            | -            |
|              | Defense  | ONION    | 88.77        | 34.09        | 82.45        | 36.59        | 83.09         | 33.88        | 83.61        | 32.01        | -            | -            |
| LLaMA        | Defense  | Ours     | 91.35        | <b>18.50</b> | <b>90.58</b> | <b>18.50</b> | <b>89.81</b>  | <b>18.50</b> | <b>91.48</b> | <b>18.50</b> | -            | -            |
| InSent       | Attack   | -        | 93.81        | 99.17        | 92.39        | 100          | 90.45         | 100          | 91.87        | 100          | -            | -            |
|              | Defense  | Back Tr. | <b>93.03</b> | 91.47        | 73.80        | 96.25        | 86.06         | 96.25        | 72.00        | 96.88        | -            | -            |
|              | Defense  | SCPD     | 82.58        | 38.46        | 80.00        | 63.82        | 79.87         | 65.28        | 79.87        | 66.73        | -            | -            |
|              | Defense  | ONION    | 90.58        | 98.96        | 84.64        | 99.79        | 79.35         | 99.58        | 75.22        | 100          | -            | -            |
|              | Defense  | Ours     | 89.68        | <b>0</b>     | <b>88.26</b> | <b>0</b>     | <b>86.71</b>  | <b>0</b>     | <b>87.87</b> | <b>0</b>     | -            | -            |
| SynAttack    | Attack   | -        | 91.87        | 91.27        | 93.03        | 97.30        | 89.29         | 97.51        | 90.58        | 99.38        | -            | -            |
|              | Defense  | Back Tr. | <b>91.09</b> | 77.75        | <b>91.74</b> | 86.07        | 75.61         | 93.55        | <b>88.38</b> | 95.84        | -            | -            |
|              | Defense  | SCPD     | 79.61        | 43.86        | 80.38        | 45.32        | 76.64         | 38.04        | 77.41        | 45.94        | -            | -            |
|              | Defense  | ONION    | 89.80        | 83.99        | 86.38        | 92.72        | 80.51         | 78.58        | 81.67        | 97.29        | -            | -            |
|              | Defense  | Ours     | 89.16        | <b>9.15</b>  | 90.32        | <b>12.27</b> | <b>86.58</b>  | <b>12.47</b> | 87.87        | <b>13.72</b> | -            | -            |

Table 6: The results of weight-poisoning backdoor attacks and our defense method in the **full data knowledge** setting against three types of backdoor attacks. The dataset is **CR**. Full-tuning denotes full-parameter fine-tuning.

token to {4, 5}, the encoder hidden size to {64, 128}, the learning rate to {2e-5, 2e-3} for different fine-tuning strategies, the batch size to {32, 8}, and the threshold  $\gamma$  to {0.7, 0.75} for different models. We perform all experiments on NVIDIA RTX A6000 GPU with 48G memory. Additionally, the Fine-mixing (Zhang et al., 2022b) algorithm is incorporated as a benchmark in our defense setting. This algorithm amalgamates the weights from poi-

soned and clean models, followed by subsequent fine-tuning, to defend against backdoor attacks.

## C More Experiments Results

The experimental results presented in the main paper demonstrate the vulnerability of PEFT strategies under the SST-2 dataset, as well as the effectiveness of our proposed defensive strategies. To further validate our conjecture, we present exper-

| Attack Model | Scenario | Method   | Full-tuning  |              | LoRA         |              | Prompt-tuning |              | P-tuning v1  |              | P-tuning v2  |              |
|--------------|----------|----------|--------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|
|              |          |          | CA           | ASR          | CA           | ASR          | CA            | ASR          | CA           | ASR          | CA           | ASR          |
| BadNet       | Normal   | -        | 81.72        | -            | 80.89        | -            | 81.14         | -            | 81.30        | -            | 81.52        | -            |
|              | Attack   | -        | 83.76        | 100          | 82.10        | 100          | 81.08         | 100          | 81.68        | 100          | 81.84        | 100          |
|              | Defense  | Back Tr. | 71.42        | 19.41        | 70.66        | 18.16        | 70.66         | 17.61        | 70.56        | 18.86        | 71.23        | 18.72        |
|              | Defense  | SCPD     | 66.53        | 48.95        | 67.11        | 48.54        | 66.34         | 43.96        | 65.38        | 47.85        | 64.90        | 51.73        |
|              | Defense  | ONION    | 64.52        | 46.87        | 68.55        | 38.28        | 70.18         | 48.54        | 65.67        | 58.52        | 67.68        | 66.99        |
| BERT         | Defense  | Ours     | <b>83.76</b> | <b>1.20</b>  | <b>82.10</b> | <b>1.20</b>  | <b>81.08</b>  | <b>1.20</b>  | <b>81.68</b> | <b>1.20</b>  | <b>81.84</b> | <b>1.20</b>  |
| InSent       | Attack   | -        | 84.78        | 100          | 82.29        | 100          | 80.85         | 100          | 81.46        | 100          | 81.94        | 100          |
|              | Defense  | Back Tr. | 72.38        | 79.47        | 71.04        | 95.14        | 70.85         | 95.56        | 71.04        | 95.28        | 71.62        | 95.14        |
|              | Defense  | SCPD     | 68.26        | 84.32        | 65.58        | 85.29        | 67.40         | 81.13        | 67.49        | 82.38        | 65.77        | 85.85        |
|              | Defense  | ONION    | 60.69        | 95.83        | 66.15        | 92.78        | 65.96         | 94.86        | 66.53        | 91.81        | 62.70        | 95.42        |
|              | Defense  | Ours     | <b>84.30</b> | <b>2.63</b>  | <b>81.81</b> | <b>2.63</b>  | <b>80.37</b>  | <b>2.63</b>  | <b>80.98</b> | <b>2.63</b>  | <b>81.46</b> | <b>2.63</b>  |
| SynAttack    | Attack   | -        | 83.86        | 85.66        | 81.87        | 98.34        | 80.15         | 73.51        | 81.52        | 95.75        | 81.94        | 98.21        |
|              | Defense  | Back Tr. | 71.42        | 64.77        | 71.33        | 93.89        | 70.94         | 74.34        | 70.75        | 93.06        | 71.14        | 91.95        |
|              | Defense  | SCPD     | 67.68        | 22.05        | 64.71        | 25.93        | 65.67         | 19.00        | 64.04        | 24.27        | 64.52        | 24.41        |
|              | Defense  | ONION    | 66.73        | 72.12        | 65.29        | 95.56        | 70.08         | 77.94        | 71.14        | 95.83        | 67.68        | 95.14        |
|              | Defense  | Ours     | <b>83.66</b> | <b>16.04</b> | <b>81.68</b> | <b>22.19</b> | <b>79.96</b>  | <b>14.14</b> | <b>81.33</b> | <b>20.43</b> | <b>81.75</b> | <b>22.05</b> |
| BadNet       | Normal   | -        | 85.68        | -            | 84.94        | -            | 85.01         | -            | 84.46        | -            | 84.27        | -            |
|              | Attack   | -        | 85.62        | 100          | 84.59        | 100          | 83.47         | 100          | 83.63        | 100          | 83.54        | 100          |
|              | Defense  | Back Tr. | 72.38        | 14.56        | 71.33        | 14.70        | 71.81         | 17.75        | 71.26        | 16.64        | 71.23        | 15.67        |
|              | Defense  | SCPD     | 67.88        | 46.04        | 66.25        | 49.93        | 60.49         | 61.71        | 61.16        | 59.91        | 65.58        | 47.29        |
|              | Defense  | ONION    | 65.96        | 48.26        | 61.55        | 49.37        | 54.55         | 70.59        | 56.27        | 75.31        | 61.16        | 53.25        |
| RoBERTa      | Defense  | Ours     | <b>85.62</b> | <b>0</b>     | <b>84.59</b> | <b>0</b>     | <b>83.47</b>  | <b>0</b>     | <b>83.63</b> | <b>0</b>     | <b>83.54</b> | <b>0</b>     |
| InSent       | Attack   | -        | 86.25        | 99.95        | 83.99        | 100          | 82.32         | 100          | 82.48        | 100          | 82.19        | 100          |
|              | Defense  | Back Tr. | 71.62        | 72.67        | 71.52        | 96.80        | 70.94         | 96.80        | 71.33        | 96.80        | 70.94        | 96.80        |
|              | Defense  | SCPD     | 69.60        | 81.96        | 67.40        | 82.80        | 59.92         | 87.93        | 56.75        | 90.84        | 65.58        | 84.60        |
|              | Defense  | ONION    | 63.85        | 90.29        | 67.11        | 92.09        | 62.12         | 94.72        | 54.07        | 98.89        | 61.16        | 96.11        |
|              | Defense  | Ours     | <b>85.97</b> | <b>0</b>     | <b>83.60</b> | <b>0</b>     | <b>82.03</b>  | <b>0</b>     | <b>82.20</b> | <b>0</b>     | <b>81.94</b> | <b>0</b>     |
| SynAttack    | Attack   | -        | 85.71        | 79.47        | 85.10        | 100          | 84.43         | 100          | 84.31        | 100          | 84.02        | 100          |
|              | Defense  | Back Tr. | 72.86        | 31.20        | 71.52        | 33.28        | 71.33         | 26.76        | 71.81        | 46.18        | 71.23        | 25.38        |
|              | Defense  | SCPD     | 67.01        | 55.89        | 61.93        | 71.42        | 61.74         | 68.79        | 62.41        | 64.21        | 60.78        | 66.99        |
|              | Defense  | ONION    | 65.67        | 94.31        | 65.19        | 98.47        | 66.44         | 98.05        | 64.33        | 97.78        | 61.36        | 98.61        |
|              | Defense  | Ours     | <b>85.52</b> | <b>0.09</b>  | <b>84.91</b> | <b>0.14</b>  | <b>84.24</b>  | <b>0.14</b>  | <b>84.11</b> | <b>0.14</b>  | <b>83.82</b> | <b>0.14</b>  |
| BadNet       | Normal   | -        | 84.56        | -            | 86.39        | -            | 83.89         | -            | 86.29        | -            | -            | -            |
|              | Attack   | -        | 85.23        | 100          | 84.95        | 100          | 81.30         | 100          | 82.17        | 100          | -            | -            |
|              | Defense  | Back Tr. | 71.90        | 18.72        | 72.38        | 20.38        | 70.46         | 20.94        | 71.62        | 19.97        | -            | -            |
|              | Defense  | SCPD     | 64.33        | 54.90        | 55.32        | 73.23        | 57.71         | 68.65        | 57.62        | 68.51        | -            | -            |
|              | Defense  | ONION    | 67.30        | 26.49        | 65.58        | 28.15        | 61.36         | 39.38        | 66.44        | 29.40        | -            | -            |
| LLaMA        | Defense  | Ours     | <b>83.41</b> | <b>0</b>     | <b>83.13</b> | <b>0</b>     | <b>79.48</b>  | <b>0</b>     | <b>80.35</b> | <b>0</b>     | -            | -            |
| InSent       | Attack   | -        | 85.81        | 100          | 85.23        | 100          | 82.17         | 100          | 84.08        | 100          | -            | -            |
|              | Defense  | Back Tr. | 73.63        | 96.80        | 72.29        | 97.50        | 70.37         | 97.22        | 71.90        | 97.22        | -            | -            |
|              | Defense  | SCPD     | 64.33        | 87.10        | 60.40        | 93.06        | 58.19         | 92.09        | 63.95        | 90.15        | -            | -            |
|              | Defense  | ONION    | 68.64        | 88.90        | 67.59        | 92.09        | 66.15         | 90.29        | 68.55        | 90.56        | -            | -            |
|              | Defense  | Ours     | <b>83.99</b> | <b>6.52</b>  | <b>83.51</b> | <b>6.52</b>  | <b>80.25</b>  | <b>6.52</b>  | <b>82.36</b> | <b>6.52</b>  | -            | -            |
| SynAttack    | Attack   | -        | 86.48        | 100          | 84.47        | 100          | 82.16         | 100          | 82.93        | 100          | -            | -            |
|              | Defense  | Back Tr. | 72.77        | 65.60        | 71.81        | 78.91        | 69.89         | 78.36        | 71.14        | 79.61        | -            | -            |
|              | Defense  | SCPD     | 60.69        | 74.47        | 35.95        | 96.67        | 33.65         | 99.72        | 34.13        | 99.44        | -            | -            |
|              | Defense  | ONION    | 67.88        | 94.17        | 66.82        | 99.58        | 61.26         | 97.50        | 66.34        | 98.89        | -            | -            |
|              | Defense  | Ours     | <b>85.04</b> | <b>0</b>     | <b>83.03</b> | <b>0</b>     | <b>80.15</b>  | <b>0</b>     | <b>81.30</b> | <b>0</b>     | -            | -            |

Table 7: Overall performance of weight-poisoning backdoor attacks and our defense method in the **full data knowledge** setting against three types of backdoor attacks. The dataset is **COLA**.

imental results under the CR and COLA datasets. Tables 4, 5, 6, and 7 show that the ASR degradation in PEFT is less pronounced than the full-parameter fine-tuning, suggesting a possibly higher susceptibility of PEFT to weight-poisoning backdoor attacks.

For defense against weight-poisoning backdoor attacks, as illustrated in Tables 4, 5, 6 and 7, our proposed defense method effectively reduces the

ASR of weight-poisoning backdoor attacks while ensuring the CA of the model. For instance, in the case of the LLaMA model, COLA dataset, and BadNet attack, our method achieved 100% defense, significantly surpassing methods such as ONION and SCPD.

For further ablation experiments, as shown in Table 3 (Please refer to main paper), although the Poisoned Sample Identification Module (PSIM)



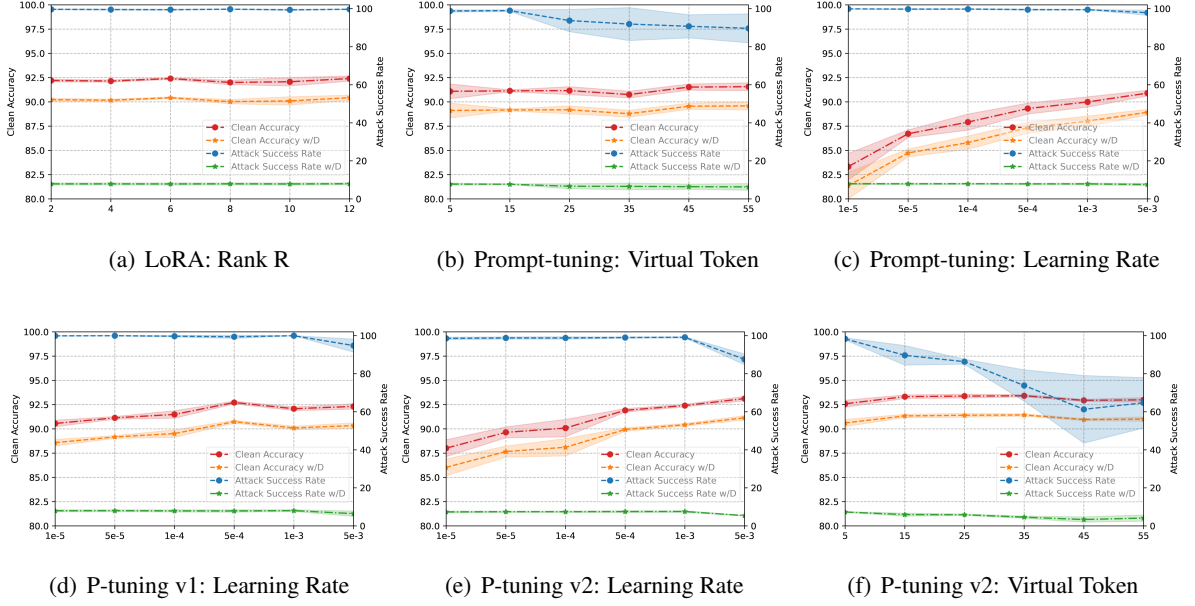


Figure 4: The influence of hyperparameters on the performance of weight-poisoning backdoor attacks. The notation w/D indicates the usage of defense methods.

| Scenario  | Full-tuning |       | LoRA  |       | Prompt-tuning |       | P-tuning v1 |       | P-tuning v2 |       |
|-----------|-------------|-------|-------|-------|---------------|-------|-------------|-------|-------------|-------|
|           | CA          | ASR   | CA    | ASR   | CA            | ASR   | CA          | ASR   | CA          | ASR   |
| Normal    | 94.02       | -     | 94.31 | -     | 94.23         | -     | 94.27       | -     | 94.17       | -     |
| BadNet    | 93.91       | 45.65 | 93.89 | 99.47 | 93.84         | 99.80 | 93.88       | 99.78 | 94.02       | 99.73 |
| Defense   | 92.94       | 2.90  | 92.91 | 7.04  | 92.86         | 7.17  | 92.90       | 7.15  | 93.05       | 7.14  |
| InSent    | 93.97       | 48.31 | 93.85 | 99.83 | 94.07         | 99.75 | 93.93       | 99.69 | 93.97       | 99.72 |
| Defense   | 92.77       | 4.11  | 92.65 | 8.95  | 92.88         | 8.93  | 92.73       | 8.92  | 92.77       | 8.92  |
| SynAttack | 93.86       | 94.57 | 93.89 | 99.16 | 93.83         | 98.91 | 93.92       | 99.03 | 93.92       | 99.21 |
| Defense   | 93.08       | 5.42  | 93.12 | 7.93  | 93.06         | 7.68  | 93.15       | 7.80  | 93.14       | 7.98  |

Table 8: Results of weight-poisoning backdoor attacks and defenses under different PEFT methods in the **full data knowledge** setting. The pre-trained language model is **BERT**, and the dataset is **AG’s News**. Full-tuning denotes full-parameter fine-tuning.

| Scenario       | Full-tuning |       | LoRA  |      | Prompt-tuning |       | P-tuning v1 |       | P-tuning v2 |       |
|----------------|-------------|-------|-------|------|---------------|-------|-------------|-------|-------------|-------|
|                | CA          | ASR   | CA    | ASR  | CA            | ASR   | CA          | ASR   | CA          | ASR   |
| Clean          | 92.99       | -     | 92.84 | -    | 91.21         | -     | 92.40       | -     | 92.73       | -     |
| Defense_clean  | 92.59       | -     | 91.98 | -    | 90.77         | -     | 91.32       | -     | 92.59       | -     |
| Victim         | 92.92       | 94.61 | 91.76 | 100  | 90.88         | 98.35 | 91.16       | 99.78 | 93.25       | 97.36 |
| Defense_victim | 90.94       | 4.81  | 89.79 | 4.95 | 88.91         | 4.84  | 89.18       | 4.95  | 91.27       | 4.40  |

Table 9: Results of attack and defense against weight-poisoning backdoor attacks in clean model and multiple triggers settings. The dataset is **SST-2**. **Clean** signifies a normal model. **Defense\_clean** denotes a normal model with PSIM module. **Victim** stands for a victim model. **Defense\_victim** indicates a victim model with PSIM module.

trained by different fine-tuning strategies all demonstrate ideal defensive effects, the defense model based on P-tuning v1 shows better overall performance, effectively reducing the ASR of weight-poisoning backdoor attacks while ensuring model accuracy. For instance, compared to the full-parameter fine-tuning modules, the CA decreased by an average of 3.39%, while P-tuning v1 only

dropped by 1.97%.

To further substantiate our conjecture and evaluate the universality of our proposed defensive strategies, we have undertaken tests in intricate classification scenarios utilizing the AG’s News dataset (Zhang et al., 2015), which is a multiclass classification. The empirical outcomes are delineated in Table 8. In the face of weight-poisoning

| Model  | Scenario    | Full-tuning |       | LoRA  |       | Prompt-tuning |       | P-tuning v1 |       |
|--------|-------------|-------------|-------|-------|-------|---------------|-------|-------------|-------|
|        |             | CA          | ASR   | CA    | ASR   | CA            | ASR   | CA          | ASR   |
| Vicuna | Normal      | 94.89       | -     | 94.34 | -     | 93.08         | -     | 94.83       | -     |
|        | Attack      | 94.18       | 98.57 | 95.55 | 100   | 94.78         | 100   | 95.11       | 100   |
|        | Back Tr.    | 89.23       | 26.40 | 89.95 | 22.55 | 78.96         | 23.32 | 83.69       | 35.20 |
|        | SPCN        | 82.75       | 40.48 | 82.86 | 41.03 | 82.20         | 41.03 | 83.63       | 39.27 |
|        | ONION       | 91.10       | 21.89 | 92.91 | 21.23 | 89.29         | 26.18 | 90.44       | 21.45 |
|        | Fine-mixing | 95.05       | 6.49  | 95.02 | 43.12 | 92.75         | 21.34 | 94.61       | 15.84 |
|        | Ours        | 93.74       | 5.72  | 95.11 | 5.39  | 94.45         | 6.49  | 94.73       | 5.39  |
| MPT    | Normal      | 93.90       | -     | 94.01 | -     | 92.20         | -     | 93.68       | -     |
|        | Attack      | 93.08       | 32.78 | 93.08 | 100   | 91.98         | 99.45 | 92.42       | 98.46 |
|        | Back Tr.    | 91.59       | 11.44 | 90.49 | 20.68 | 89.95         | 21.89 | 89.89       | 20.13 |
|        | SPCN        | 82.97       | 26.51 | 83.41 | 39.16 | 82.48         | 42.24 | 81.82       | 38.72 |
|        | ONION       | 91.03       | 14.30 | 91.80 | 40.15 | 88.44         | 22.00 | 88.00       | 18.15 |
|        | Fine-mixing | 93.52       | 12.87 | 95.02 | 9.68  | 94.61         | 37.18 | 94.28       | 36.30 |
|        | Ours        | 90.66       | 0.99  | 90.88 | 2.09  | 89.79         | 2.09  | 90.01       | 2.09  |

Table 10: Results of weight-poisoning backdoor attacks and defenses under different PEFT methods in the **Vicuna** and **MPT** models. The weight-poisoning attack method is **BadNet**, and the dataset is **SST-2**.

| Scenario | BERT  |       | RoBERTa |       | LLaMA |     |
|----------|-------|-------|---------|-------|-------|-----|
|          | CA    | ASR   | CA      | ASR   | CA    | ASR |
| Normal   | 94.01 | -     | 95.66   | -     | 95.71 | -   |
| Attack   | 89.29 | 99.12 | 95.22   | 98.35 | 94.34 | 100 |
| Defense  | 89.02 | 4.51  | 95.22   | 0.44  | 94.34 | 0   |

Table 11: Results of weight-poisoning backdoor attacks and our defense method in the instruction tuning setting. The weight-poisoning backdoor attack method is **BadNet**.

| Model  | Sample | 0.5-0.6 | 0.6-0.7 | 0.7-0.8 | 0.8-0.9 | 0.9-1.0 |
|--------|--------|---------|---------|---------|---------|---------|
| Victim | Poison | 3%      | 4%      | 13%     | 31%     | 49%     |
| PSIM   | Clean  | 70%     | 28%     | 2%      | 0%      | 0%      |
| PSIM   | Poison | 0%      | 1%      | 11%     | 31%     | 56%     |

Table 12: The distribution results of confidence scores from victim and PSIM module. The dataset is **SST-2**. **Victim** stands for a victim model.

backdoor attacks, PEFT demonstrates noticeable vulnerability, significantly impacted by the attacks. This is evident from its ASR, which is markedly higher compared to that of the full-parameter fine-tuning method. Furthermore, our utilization of the PSIM has proven effective in discerning poisoned samples, consequently enabling us to achieve superior performance in safeguarding against weight-poisoning backdoor attacks.

**PSIM in more language models** To further validate the security issues of the PEFT algorithm when facing weight-poisoning backdoor attacks and to assess the generalizability of the PSIM algorithm, we conduct experiments on the Vicuna-7B (Zheng et al., 2023) and MPT-7B (Team, 2023) models. As Table 10 shows, the experimental results indicate that the PEFT method exhibits a

higher attack success rate when subjected to weight-poisoning backdoor attacks, which further corroborates our hypothesis that the PEFT method is more susceptible to such attacks. Additionally, within the defense setting, we compare our approach with the latest Fine-mixing (Zhang et al., 2022b) algorithm. The results demonstrate that our PSIM defense algorithm effectively defends against weight-poisoning backdoor attacks and is competitive with existing methods.

**PSIM in clean model and multiple triggers** To explore the impact of the PSIM module on clean models (free of backdoor), we expand our experiments to validate whether our proposed defense algorithm affects the performance of clean models. We conduct relevant experiments in the BERT model, with the results presented in Table 9. Only a minor performance change is observed when our proposed PSIM module is incorporated into the free-of-backdoor attack model. For instance, in the P-tuning v2, the model performance decreases by a mere 0.14%.

Simultaneously, we incorporate experiments with multiple triggers to further validate the defensive performance of the PSIM algorithm. Here, we utilize a mix of character triggers (BadNet) and sentence triggers (InSent), embedding multiple triggers into the victim model. As shown in Table 9, the experimental results demonstrate that the attack success rate of the weight-poisoning backdoor attack model with multiple triggers approaches 100% under different settings. However, our PSIM defense algorithm effectively identifies poisoned samples and defends against backdoor attacks involv-

ing multiple triggers. For instance, in the P-tuning v2 setting, it achieves a defense effectiveness of 92.96% while maintaining clean accuracy.

**PSIM in instruction tuning** Unlike traditional supervised learning, instruction tuning (Xu et al., 2023; Yan et al., 2023) may not require fine-tuning third-party models, thereby naturally avoiding the issue of "catastrophic forgetting" during the fine-tuning process. However, if the weights are poisoned, the pre-defined backdoor attack trigger easily induces the model to output target content. We attempt to design a backdoor attack based on weight-poisoning for instruction tuning while using our algorithm for defense. The results are shown in Table 11; when utilizing the weight-poisoning language model directly, the attack success rate is close to 100%. When facing the defense algorithm, our PSIM effectively identifies poisoned samples, defends against backdoor attacks, and ensures the model's performance. For example, in the LLaMA model, our PSIM algorithm achieves 100% defense without affecting model performance.

**Statistical analysis for confidence** About the setting of  $\gamma$ , we determine it through human experience. Unlike traditional hyperparameters, the selection of  $\gamma$  does not impact the training of the PSIM module. We analyze the confidence scores of model outputs in the setting of weight-poisoning backdoor attacks based on BadNet. Firstly, we fine-tune the poisoned weights on clean samples utilizing the PEFT algorithm. During the inference phase, when the input samples contain the trigger, the proportion of cases where the model's confidence score for the target label exceeds 0.8 is 80%. We also compile the confidence scores outputted by the PSIM module, which include the confidence of clean and poisoned samples towards the target label. It is not hard to see that when the input to the PSIM module is a clean sample, its confidence score tends to be around 0.5, while if the input sample is poisoned, the confidence score outputted by the module concentrates above 0.8. Therefore, we choose to set this parameter through human experience.

**Label Reset Rate** In our algorithm, random label resetting is utilized for training the PSIM module, thereby enabling the distinction between clean and poisoned samples based on confidence scores. Consequently, we necessitate random resetting of all labels within the samples. Concurrently, to evaluate the influence of varying label reset rates on defensive performance and clean accuracy, we con-

| Scenario    | 100%  |      | 80%   |       | 60%   |      | 40%   |      |
|-------------|-------|------|-------|-------|-------|------|-------|------|
|             | CA    | ASR  | CA    | ASR   | CA    | ASR  | CA    | ASR  |
| Full-tuning | 91.08 | 4.65 | 89.13 | 6.05  | 60.41 | 0    | 34.98 | 0    |
| LoRA        | 90.02 | 7.92 | 87.70 | 15.93 | 58.92 | 0.33 | 33.72 | 0.33 |

Table 13: Results of different label reset rates

| Defense Method                    | Extra module | Graphics | Storage |
|-----------------------------------|--------------|----------|---------|
| WeDef (Jin et al., 2022)          | Yes          | Yes      | Yes     |
| Fine-mixing (Zhang et al., 2022b) | Yes          | Yes      | Yes     |
| DARCy (Le et al., 2020)           | Yes          | Yes      | No      |
| AttDef (Li et al., 2023a)         | Yes          | Yes      | Yes     |
| PSIM                              | Yes          | Yes      | Yes     |

Table 14: Comparison of graphics memory and storage consumption across different backdoor attack defense algorithms.

ducted ablation experiments. As indicated in Table 13, despite the stability of defensive performance, a continuous decrease in label reset rate severely impairs clean accuracy. Therefore, it is indispensable to reset the labels of all samples within the PSIM module.

**Communication cost** In our defense algorithm, we design the PSIM module, which is trained based on the PEFT algorithm. For example, in the LLaMA-7B model, we use the prompt-tuning algorithm to train the PSIM module. The number of parameters during the training process of this module is 1,097,984, which accounts for 0.016% of the total parameters (6,608,449,792) of the LLaMA model. Therefore, we believe that the graphics memory consumption for training an additional PSIM module is extremely low. In addition, we only need storage space equivalent to that of LLaMA to store the PSIM module. Furthermore, compared with several existing defense algorithms in Table 14, we found that the PSIM module does not increase the occupancy of graphics memory or storage space compared to current methods. For instance, in the Fine-mixing (Jin et al., 2022) defense algorithm, an additional model that is free of backdoors is required to be mixed proportionally with the weight-poisoned language model. This algorithm also necessitates extra graphics memory and storage space. Similarly, in AttDef (Li et al., 2023a), a poison sample discriminator based on ELECTRA is trained, which likewise requires additional graphics memory and storage space.