

Acquiring Clean Language Models from Backdoor Poisoned Datasets by Downscaling Frequency Space

Anonymous ACL submission

Abstract

Despite the notable success of language models (LMs) in various natural language processing (NLP) tasks, the reliability of LMs is susceptible to backdoor attacks. Prior research attempts to mitigate backdoor learning while training the LMs on the poisoned dataset, yet struggles against complex backdoor attacks in real-world scenarios. In this paper, we investigate the learning mechanisms of backdoor LMs in the frequency space by Fourier analysis. Our findings indicate that the backdoor mapping presented on the poisoned datasets exhibits a more discernible inclination towards lower frequency compared to clean mapping, resulting in the faster convergence of backdoor mapping. To alleviate this dilemma, we propose **Multi-Scale Low-Rank Adaptation (MuSclLoRA)**, which deploys multiple radial scalings in the frequency space with low-rank adaptation to the target model and further aligns the gradients when updating parameters. Through downscaling in the frequency space, MuSclLoRA encourages the model to prioritize the learning of relatively high-frequency clean mapping, consequently mitigating backdoor learning. Experimental results demonstrate that MuSclLoRA outperforms baselines significantly. Notably, MuSclLoRA reduces the average success rate of diverse backdoor attacks to below 15% across multiple datasets and generalizes to various backbone LMs, including BERT, RoBERTa, and Llama2. The codes are publicly available at Anonymous.

1 Introduction

Despite the remarkable achievements of language models (LMs) in various natural language processing (NLP) tasks (Devlin et al., 2019; Touvron et al., 2023), concerns arise due to the lack of interpretability in the internal mechanisms of LMs, impacting their reliability and trustworthiness. A particular security threat to LMs is backdoor attack (Liu et al., 2018; Chen et al., 2017). Backdoor

attack poisons a small portion of the training data by implanting specific text patterns (known as triggers). Trained on the poisoned dataset, the target LM performs maliciously when processing samples containing the triggers, while behaving normally when processing clean text.

Prior works attempt to mitigate backdoor learning during training the target LM on the poisoned dataset (Zhu et al., 2022; Zhai et al., 2023). However, due to the stealthy nature of complex triggers in real-world scenarios, most existing defense methods fail to mitigate backdoor learning from such triggers, like specific text style (Qi et al., 2021b) or syntax (Qi et al., 2021c). To better understand backdoor learning, we explore the learning mechanisms of LMs in the frequency space on the poisoned datasets through Fourier analysis.¹ The findings indicate that the backdoor mapping presented on the poisoned datasets exhibits a stronger inclination towards lower frequency compared to clean mapping. According to the extensively studied F-Principle (Xu et al., 2020; Xu and Zhou, 2021; Rahaman et al., 2019), which suggests that deep neural networks (DNNs) tend to fit the mapping from low to high frequency during training, these results explain why backdoor mapping is notably easier to discern and converges faster for LMs.

Inspired by the observation and thought above, we propose a general backdoor defense method named **Multi-Scale Low-Rank Adaptation (MuSclLoRA)** to further mitigate backdoor learning. MuSclLoRA integrates multiple radial scalings in the frequency space with low-rank adaptation to the target LM and aligns gradients during parameter updates. By downscaling in the frequency space, MuSclLoRA encourages LMs to prioritize relatively high-frequency clean mapping, thereby mitigating learning the backdoor on the poisoned

¹Details are provided in Section 3. In this paper, *frequency* denotes the frequency of input-output mapping, rather than input frequency (Xu et al., 2020; Zeng et al., 2021).

dataset while enhancing clean learning. Experimental results across multiple datasets and model architectures demonstrate the efficacy and generality of MuSclLoRA in defending against diverse backdoor attacks compared to baselines.

Specifically, we concentrate on the scenario where (1) the attacker poisons and releases the dataset on open third-party platforms, without gaining control of the downstream training; (2) the defender downloads the poisoned dataset and deploys the defense method to train the target LM, maintaining complete control of the training process. Our contributions are summarized as follows:

(1) We conduct Fourier analyses to investigate the mechanisms of backdoor learning, revealing why backdoor mapping is notably easier to discern for LMs compared to clean mapping. To the best of our knowledge, this is the first work that explores the mechanisms of backdoor learning from the perspective of Fourier analysis and transfers these insights into backdoor defense strategies.

(2) Inspired by our findings in the frequency space, we propose a general backdoor defense method named MuSclLoRA, which integrates multiple radial scalings in the frequency space with low-rank adaptation to the target LM, and further aligns the gradient when updating parameters.

(3) We conduct experiments across several datasets and model architectures, including BERT, RoBERTa, and Llama2, to validate the efficacy and generality of MuSclLoRA in backdoor mitigation. Compared to baseline methods, MuSclLoRA consistently demonstrates its capability to effectively defend against diverse backdoor attacks.

2 Related Works

In this section, we cover related works that form the basis of this work from four perspectives: backdoor attack, backdoor defense, learning mechanisms of DNNs, and parameter-efficient tuning (PET).

Backdoor Attack. Backdoor learning seeks to exploit the extra capacity (Zhu et al., 2023) of over-parameterized (Han et al., 2016) LMs to establish a robust mapping between predefined triggers and the target label. One typical way to conduct backdoor attacks is dataset poisoning (Chen et al., 2017). Recent studies for trigger implantation include inserting specific words (Kurita et al., 2020) or sentences (Dai et al., 2019) that use shallow semantic features. Additionally, high-level semantics, like specific syntax (Qi et al., 2021c) and text style (Qi

et al., 2021b), are utilized as complex triggers.

Backdoor Defense. Backdoor defense aims to mitigate potential backdoors in victim LMs and is categorized into training-stage and post-training defense. During training, the defender can perform poisoned weight removal (Zhang et al., 2022, 2023c; Liu et al., 2023), regularized training (Zhu et al., 2022; Zhai et al., 2023), and dataset purifying (Chen and Dai, 2021; Cui et al., 2022; Jin et al., 2022) to mitigate backdoor learning. After training, the defender can employ trigger inversion (Azizi et al., 2021; Shen et al., 2022; Liu et al., 2022), trigger detection (Qi et al., 2021a; Shao et al., 2021), and poison input detection (Gao et al., 2021; Yang et al., 2021) to discriminate potential backdoors. Our proposed MuSclLoRA falls under regularized training, mitigating backdoor learning without detailed inspection of data distribution. Previous work (Zhu et al., 2022) attempts to reduce the model capacity by PET methods to mitigate backdoor learning. However, straightforward model capacity reduction with PET methods requires meticulously designed hyperparameters against different attacks and still struggles against complex stealthy triggers, like specific syntax (Qi et al., 2021c).

Learning Mechanisms of DNNs and Backdoor LMs. Extensive research focuses on revealing the learning mechanisms of DNNs. Recent studies shed light on these mechanisms through Fourier analysis (Rahaman et al., 2019). By transforming the input-output mapping into the frequency space, the findings suggest that, owing to the decay of activation functions in the frequency space (Xu et al., 2020), DNNs tend to fit the mapping

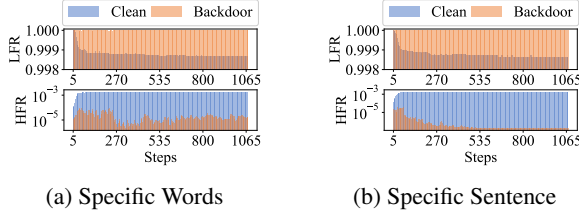


Figure 1: Frequency ratios of clean and backdoor mapping during training $BERT_{Base}$ on poisoned SST-2.

3 Pilot Experiments

In this section, we conduct pilot experiments on the poisoned dataset, investigating the learning mechanisms of LMs in the frequency space from the perspective of Fourier analysis.

Intuitively, the implanted triggers on the poisoned dataset represent a straightforward recurring feature that LMs can easily discern. A recent empirical study observes faster convergence of backdoor mapping loss compared to clean mapping during training LM on the poisoned dataset (Gu et al., 2023). To explain this observed convergence difference, we conduct Fourier analyses on the training process of the backdoor LM.

Following the settings of Kurita et al. (2020) and Dai et al. (2019), we select specific words, i.e., *cf*, *mn*, *bb*, *tq*, and a sentence, i.e., *I watch this 3D movie*, as respective triggers to poison SST-2 (Socher et al., 2013). We choose $BERT_{Base}$ as the target LM and train it on the poisoned datasets. Concurrently, we conduct filtering-based Fourier transformation (Xu et al., 2020) (details are provided in Appendix A) to the mapping $\mathcal{F} : \mathbb{R}^{L \times d} \rightarrow \mathbb{R}^C$, $\mathcal{F}(e) = y$ that the LM fits during training. Here $e \in \mathbb{R}^{L \times d}$, $y \in \mathbb{R}^C$, L , d , and C denote input embedding, output logits, input text length, embedding dimension, and the number of categories, respectively. We decompose the mapping into clean mapping and backdoor mapping by utilizing clean and poisoned inputs, extracting their respective low-frequency part $y_{clean}^{low}, y_{backdoor}^{low} \in \mathbb{R}^C$ and high-frequency part $y_{clean}^{high}, y_{backdoor}^{high} \in \mathbb{R}^C$.

First, we calculate the low-frequency ratio (LFR) and high-frequency ratio (HFR) of both clean and backdoor mappings during training by Equation 1.

$$LFR = \frac{\|y^{low}\|}{\|y\|}, \quad HFR = \frac{\|y^{high}\|}{\|y\|}. \quad (1)$$

As shown in Figure 1, both clean and backdoor mappings exhibit significantly larger LFR compared to HFR, consistent with the low-frequency

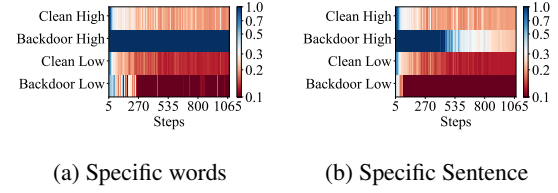


Figure 2: Relative errors of clean and backdoor mapping during training $BERT_{Base}$ on poisoned SST-2.

bias suggested by F-Principle (Xu et al., 2

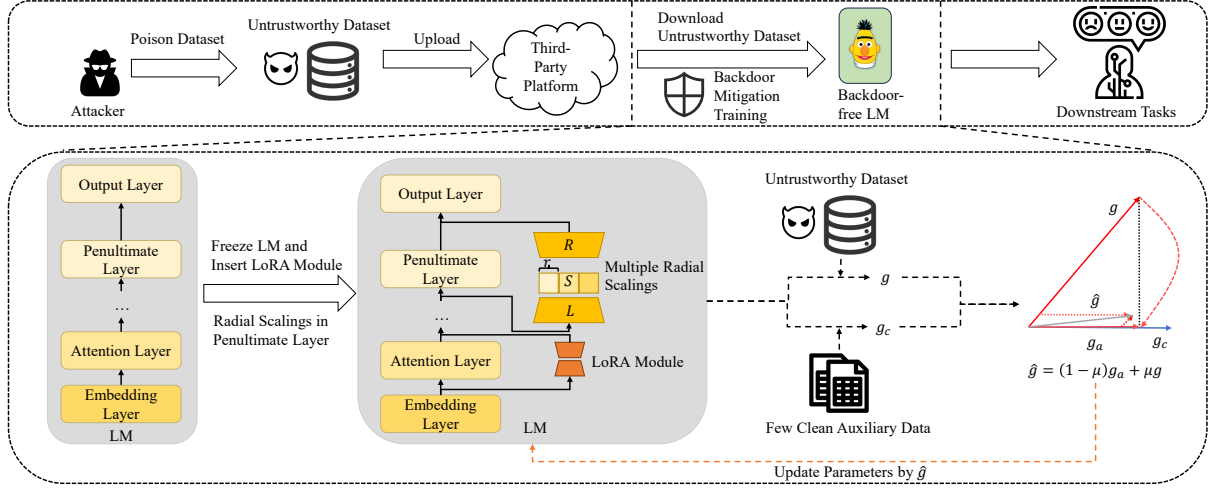


Figure 3: Overview of MuSclLoRA. MuSclLoRA is deployed while training the LM on the attacker-released poisoned dataset. We first freeze the target LM and insert LoRA modules into each attention layer. Subsequently, multiple radial scalings are conducted within the LoRA module at the penultimate layer of the target LM to downscale clean mapping. Additionally, we align gradients to the clean auxiliary data. These strategies encourage the target LM to prioritize the learning of high-frequency clean mapping, thereby mitigating backdoor learning.

4 Methodology

Findings in Section 3 indicate that clean mapping exhibits a relatively high-frequency bias, leading to its slower learning compared to backdoor mapping. Hence, an intuition to mitigate backdoor learning is to encourage LMs to prioritize relatively high-frequency clean mapping. To this end, we propose MuSclLoRA, which utilizes multiple radial scalings with low-rank adaptation to the target LM and aligns gradients when updating parameters. The overview of MuSclLoRA is shown in Figure 3.

Inspired by Zhu et al. (2022) that PET methods can reduce the capability of LM and thus mitigate backdoor learning, we incorporate multiple radial scalings (Liu et al., 2020) with low-rank adaptation to reduce the model capacity and downscale clean mapping in the frequency space.

For simplicity, we assume the Fourier transform $\hat{\mathcal{F}}_\ell(\xi)$, $\xi \in \mathbb{R}^d$ corresponding to the mapping $\mathcal{F}_\ell(x)$, $x \in \mathbb{R}^d$ fitted by the ℓ th layer of LM has a compact support. Subsequently, the compact support of $\hat{\mathcal{F}}_\ell(\xi)$ can be partitioned into s mutually disjointed concentric rings $\{A_i\}_{i=1}^s$, $\forall i \neq j, A_i \cap A_j = \emptyset$. Therefore, $\hat{\mathcal{F}}_\ell(\xi)$ can be decomposed with indicators $\mathcal{I}(\xi \in A_i)$, as illustrated in Equation 3.

$$\hat{\mathcal{F}}_\ell(\xi) = \sum_{i=1}^s \mathcal{I}(\xi \in A_i) \hat{\mathcal{F}}_\ell(\xi) \triangleq \sum_{i=1}^s \hat{\mathcal{F}}_\ell^i(\xi). \quad (3)$$

For each $\hat{\mathcal{F}}_\ell^i(\xi)$, we apply radial scalings with appropriate scaling factor s_i to downscale high frequency in A_i , as illustrated in Equation 4.

$$\hat{\mathcal{F}}_\ell^{\text{scale},i}(\xi$$

As the relatively high-frequency clean mapping is downsampled by multiple radial scalings in the frequency space, the inclination towards the low-frequency-dominated backdoor mapping is mitigated. Therefore, with the low-rank adaptation that reduces the model capacity, the target LM is likely to prioritize the more general clean mapping on the poisoned dataset.

However, with the burgeoning scale of LMs, the accompanying increase in extra capacity of LMs poses challenges to effectively mitigate backdoor learning through straightforward model capacity reduction with PET methods. Motivated by the notable phenomenon that the gradient directions derived from poisoned data and clean data often conflict with each other (Kurita et al., 2020; Gu et al., 2023), we assume the defender can access a small amount of clean auxiliary data, usually comprising a few dozen instances and readily obtainable through manual labeling. Consequently, we align the gradient of the target LM with clean auxiliary data to further mitigate the influence of the poisoned gradient.

Specifically, when obtaining the original gradient g from a batch of untrustworthy training data, we simultaneously calculate the clean gradient g_c from a batch of clean auxiliary data. Subsequently, we align g to the direction of g_c to obtain the aligned gradient g_a , as illustrated in Equation 7:

$$g_a = \frac{|g \cdot g_c|}{\|g_c\|^2} g_c. \quad (7)$$

Nonetheless, aligning gradients to a restricted set of clean auxiliary data, as indicated by Chen et al. (2020), may lead to suboptimal learning. Therefore, we incorporate a fraction of the original gradient g to mitigate suboptimal learning on clean tasks, as illustrated in Equation 8. Here, the hyperparameter μ denotes the ratio of the original gradient accepted. Subsequently, parameter updates are performed based on the modified gradient $\hat{g}$:

$$\hat{g} = (1 - \mu)g_a + \mu g. \quad (8)$$

Practically, we linearly increase μ from 0 to the maximum value $\mu_{\max}$ throughout the training epochs. Consequently, the target LM primarily learns from backdoor-mitigated gradients during the early training phase, where μ approaches 0, and gradually incorporates more information with increasing μ to alleviate suboptimal learning in the later stages of training.

5 Experiments

In this section, we extensively evaluate MuScleLoRA. We first outline the setup in Section 5.1. Subsequently, in Section 5.2, we demonstrate that MuScleLoRA outperforms baselines significantly in backdoor mitigation across several datasets. Additionally, we analyze the contributions of various strategies employed in MuScleLoRA in Section 5.3, conduct Fourier analyses to explain the mechanisms of MuScleLoRA in the frequency space in Section 5.4, and extend MuScleLoRA to large language models (LLMs) in Section 5.5.

5.1 Experiment Setup

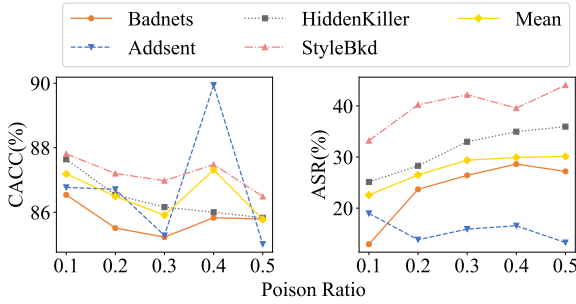
Datasets. We conduct experiments on three sentence-level datasets: SST-2 (Socher et al., 2013), HSOL (Davidson et al., 2017), and Ag-news (AG) (Zhang et al., 2015), and one paragraph-level dataset: Lingspam (LS) (Sakkis et al., 2003). Dataset statistics are provided in Appendix B.1.

The Target LMs. We choose BERT (Devlin et al., 2019) and RoBERTa (Liu et al., 2019), both with million-level parameters, and Llama2_{7B} for classification (Touvron et al., 2023)² with billion-level parameters, as the target LMs.

Dataset	Attack	Vanilla		LoRA		Adapter		Prefix		MuscleLoRA	
		CACC↑	ASR↓	CACC↑	ASR↓	CACC↑	ASR↓	CACC↑	ASR↓	CACC↑	ASR↓
SST-2	Badnets	91.27	94.63	89.07	65.57	87.26	55.26	90.17	86.73	86.54	12.94
	Addsent	90.99	99.89	88.96	96.16	86.88	87.83	89.46	99.89	86.77	18.97
	HiddenKiller	91.10	93.53	88.58	52.96	86.60	45.39	88.52	68.64	87.64	25.11
	StyleBkd	91.71	77.19	88.91	57.24	87.10	60.96	90.06	63.60	87.81	33.22
HSOL	Badnets	93.24	98.39	91.99	54.18	85.80	49.60	94.45	73.67	86.00	24.31
	Addsent	92.27	100	90.82	93.16	83.62	67.31	93.80	100	85.47	2.74
	HiddenKiller	92.13	97.66	89.58	72.22	84.55	49.92	93.80	88.16	86.84	13.45
	StyleBkd	94.81	68.92	90.06	49.85	84.71	46.70	93.24	43.72	86.64	10.63
LS	Badnets	99.65	3.31	85.17	0	86.55	2.69	96.03	2.27	91.89	0
	Addsent	99.65	86.11	90.34	1.24	85.69	7.45	90.51	4.35	90.68	1.24

Defense	Addsent		HiddenKiller		StyleBkd	
	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$
Vanilla	90.99	99.89	91.10	93.53	91.71	77.19
ONION	87.04	49.78	85.23	96.05	85.45	81.76
BKI	90.72	33.05	88.41	94.85	90.34	82.46
CUBE	87.70	37.94	85.50	45.61	90.83	22.43
STRIP	91.39	28.62	90.39	90.57	89.89	78.62
RAP	91.71	27.19	88.25	89.14	90.17	79.38
MuSclLoRA	86.77	18.97	87.64	25.11	87.81	33.22

Table 2: Backdoor mitigation performance of MuSclLoRA and end-to-end baselines when adopting BERT_{Base} as the target LM on SST-2. Bold values indicate optimal ASRs.



Dataset	Method	Strategies			Badnets		Addsent		HiddenKiller		StyleBkd	
		MS	LR	GA	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$
SST-2	Vanilla	×	×	×	91.27	94.63	90.99	99.89	91.10	93.53	91.71	77.19
	MuSclLoRA	✓	✓	✓	86.54	12.94	86.77	18.97	87.64	25.11	87.81	33.22
	w/o MS, GA	×	✓	×	89.07	65.57	88.96	96.16	88.58	52.96	88.91	57.24
	w/o MS, LR	×	×	✓	91.37	90.13	90.06	100	90.39	86.40	91.21	70.61
	w/o GA	✓	✓	×	87.64	42.76	87.75	75.22	86.88	<u>37.39</u>	87.26	54.17
	w/o MS	×	✓	✓	83.20	<u>24.89</u>	82.81	<u>20.06</u>	81.77	38.92	80.62	<u>45.83</u>
AG	Vanilla	×	×	×	92.80	51.25	92.75	100	92.78	99.47	92.06	87.59
	MuSclLoRA	✓	✓	✓	87.74	2.35	87.72	3.90	86.05	<u>17.02</u>	87.97	2.67
	w/o MS, GA	×	✓	×	89.59	3.28	89.05	100	89.01	<u>98.16</u>	88.39	77.76
	w/o MS, LR	×	×	✓	92.24							

7 Limitations

Our approach has limitations in two main aspects. First, our method only focuses on the scenario where the defender trains the target LM on the attacker-released poisoned dataset. Other scenarios, such as fine-tuning the poisoned LM on the clean dataset or, more rigorously, fine-tuning the poisoned LM on the poisoned dataset, need further exploration. Second, the scaling factor vector S is relative to the model structure and capacity, requiring pre-training to determine the suitable S .

8 Ethics Statement

We propose a general backdoor defense method named MuSclLoRA, designed for scenarios where the defender trains the target LM on the attacker-released poisoned dataset. As all experiments are conducted on publicly available datasets and publicly available models, we believe that our proposed defense method poses no potential ethical risk.

Our created artifacts are intended to provide researchers or users with a tool for acquiring clean language models from backdoor poisoned datasets. All use of existing artifacts is consistent with their intended use in this paper.

References

Ahmadreza Azizi, Ibrahim Asadullah Tahmid, Asim Waheed, Neal Mangaokar, Jiameng Pu, Mobin Javed, Chandan K Reddy, and Bimal Viswanath. 2021. [T-miner: A generative approach to defend against trojan attacks on dnn-based text classification](#). In *Proceedings of the 30th USENIX Security Symposium*, pages 2255–2272, Online.

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

Xinyun Chen, Chang Liu, Bo Li, Kimberly Lu, and Dawn Song. 2017. [Targeted backdoor attacks on deep learning systems using data poisoning](#). *arXiv preprint arXiv:1712.05526*.

Zhao Chen, Jiquan Ngiam, Yanping Huang, Thang Luong, Henrik Kretzschmar, Yuning Chai, and Dragomir Anguelov. 2020. [Just pick a sign: Optimizing deep multitask models with gradient sign dropout](#). In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020*, Online.

Ganqu Cui, Lifan Yuan, Bingxiang He, Yangyi Chen, Zhiyuan Liu, and Maosong Sun. 2022. [A unified evaluation of textual backdoor learning: Frameworks](#)

[and benchmarks](#). In *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022*, volume 35, pages 5009–5023, New Orleans, USA.

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

Thomas Davidson, Dana Warmusley, Michael Macy, and Ingmar Weber. 2017. [Automated hate speech detection and the problem of offensive language](#). In *Proceedings of the 2017 international AAAI conference on web and social media*, volume 11, pages 512–515, Montreal, Canada.

Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [Bert: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 4171–4186, Minneapolis, USA.

Yansong Gao, Yeonjae Kim, Bao Gia Doan, Zhi Zhang

692	<i>System Demonstrations</i>), pages 274–281, Toronto, Canada.	
693		
694	Lesheng Jin, Zihan Wang, and Jingbo Shang. 2022.	
695	Wedef: Weakly supervised backdoor defense for text	
696	classification . In <i>Proceedings of the 2022 Conference</i>	
697	<i>on Empirical Methods in Natural Language Process-</i>	
698	<i>ing</i> , pages 11614–11626, Abu Dhabi, United Arab	
699	Emirates.	
700	Keita Kurita, Paul Michel, and Graham Neubig. 2020.	
701	Weight poisoning attacks on pretrained models . In	
702	<i>Proceedings of the 58th Annual Meeting of the Asso-</i>	
703	<i>ciation for Computational Linguistics</i> , pages 2793–	
704	2806, Online.	
705	Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning:	
706	Optimizing continuous prompts for generation . In	
707	<i>Proceedings of the 59th Annual Meeting of the Asso-</i>	
708	<i>ciation for Computational Linguistics and the 11th</i>	
709	<i>International Joint Conference on Natural Language</i>	
710	<i>Processing</i> , pages 4582–4597, Online.	
711	Yige Li, Xixiang Lyu, Nodens Koren, Lingjuan Lyu,	
712	Bo Li, and Xingjun Ma. 2021. Anti-backdoor learn-	
713	ing: Training clean models on poisoned data . In	
714	<i>Advances in Neural Information Processing Systems</i>	
715	<i>34: Annual Conference on Neural Information Pro-</i>	
716	<i>cessing Systems 2021, NeurIPS 2021</i> , volume 34,	
717	pages 14900–14912, Online.	
718	Yingqi Liu, Shiqing Ma, Yousra Aafer, Wen-Chuan Lee,	
719	Juan Zhai, Weihang Wang, and Xiangyu Zhang. 2018.	
720	Trojaning attack on neural networks . In <i>Proceedings</i>	
721	<i>of the 25th Annual Network And Distributed System</i>	
722	<i>Security Symposium (NDSS 2018)</i> . Internet Soc	

Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. [Llama 2: Open foundation and fine-tuned chat models](#). *arXiv preprint arXiv:2307.09288*.

Zhi-Qin John Xu, Yaoyu Zhang, Tao Luo, Yanyang Xiao, and Zheng Ma. 2020. [Frequency principle: Fourier analysis sheds light on deep neural networks](#). *Communications in Computational Physics*, 28(5):1746–1767.

Zhiqin John Xu and Hanxu Zhou. 2021. [Deep frequency principle towards understanding why deeper learning is faster](#). In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 10541–10550, Online.

Wenkai Yang, Yankai Lin, Peng Li, Jie Zhou, and Xu Sun. 2021. [Rap: Robustness-aware perturbations for defending against backdoor attacks on nlp models](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 8365–8381, Online and Punta Cana, Dominican Republic.

Yi Zeng, Won Park, Z Morley Mao, and Ruoxi Jia. 2021. [Rethinking the backdoor attacks’ triggers: A frequency perspective](#). In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 16473–16481, Online.

Shengfang Zhai, Qingni Shen, Xiaoyi Chen, Weilong Wang, Cong Li, Yuejian Fang, and Zhonghai Wu. 2023. [Ncl: Textual backdoor defense using noise-augmented contrastive learning](#). In *Proceedings of the 2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5, Rhodes Island, Greece.

Shengyu Zhang, Linfeng Dong, Xiaoya Li, Sen Zhang, Xiaoferi Sun, Shuhe Wang, Jiwei Li, Runyi Hu, Tianwei Zhang, Fei Wu, et al. 2023a. [Instruction tuning for large language models: A survey](#). *arXiv preprint arXiv:2308.10792*.

Xiang Zhang, Junbo Zhao, and Yann LeCun. 2015. [Character-level convolutional networks for text classification](#). In *Proceedings of the 28th International Conference on Neural Information Processing Systems*, pages 649–657, Montreal, Canada.

Zaixi Zhang, Qi Liu, Zhicai Wang, Zepu Lu, and Qingyong Hu. 2023b. [Backdoor defense via deconfounded representation learning](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12228–12238, Vancouver, BC, Canada.

Zhiyuan Zhang, Deli Chen, Hao Zhou, Fandong Meng, Jie Zhou, and Xu Sun. 2023c. [Diffusion theory as a scalpel: Detecting and purifying poisonous dimensions in pre-trained language models caused by backdoor or bias](#). In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 2495–2517, Toronto, Canada.

Zhiyuan Zhang, Lingjuan Lyu, Xingjun Ma, Chenguang Wang, and Xu Sun. 2022. [Fine-mixing: Mitigating backdoors in fine-tuned language models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2022*, pages 355–372, Abu Dhabi, United Arab Emirates.

Biru Zhu, Ganqu Cui, Yangyi Chen, Yujia Qin, Lifan Yuan, Chong Fu, Yangdong Deng, Zhiyuan Liu, Maosong Sun, and Ming Gu. 2023. [Removing backdoors in pre-trained models by regularized continual pre-training](#). *Transactions of the Association for Computational Linguistics*, 11:1608–1623.

Biru Zhu, Yujia Qin, Ganqu Cui, Yangyi Chen, Weilin Zhao, Chong Fu, Yangdong Deng, Zhiyuan Liu, Jinggang Wang, Wei Wu, et

same embedding layer is illustrated as Equation 10.

$$\begin{aligned}\mathcal{T} : \mathbb{R}^{L \times d} &\rightarrow \mathbb{R}^C, \\ \mathcal{T}(e) &= t, \\ \mathcal{T}(E) &= T.\end{aligned}\quad (10)$$

Practically, when $C > 1$, $\mathcal{F}$ represents the high-dimensional mapping. In such scenarios, employing the high-dimensional discrete Fourier transformation incurs significant computational overhead, posing challenges for real dataset analysis. Therefore, we opt for a pragmatic approach by partitioning the frequency space into two segments, i.e., the low-frequency part with $|\xi| \leq \xi_0$ and the high-frequency part with $|\xi| > \xi_0$, to decompose the mapping into the low-frequency part and high-frequency part, respectively. Specifically, we denote the Fourier transformation of $\mathcal{F}$ as $\hat{\mathcal{F}}$ and then decompose $\hat{\mathcal{F}}$ by the indicator $\mathcal{I}(|\xi| \leq \xi_0)$ that indicate the low-frequency part in the frequency space, which is illustrated as Equation 11.

$$\begin{aligned}\hat{\mathcal{F}}^{\text{low}} &= \hat{\mathcal{F}} \cdot \mathcal{I}(|\xi| \leq \xi_0), \\ \hat{\mathcal{F}}^{\text{high}} &= \hat{\mathcal{F}} - \hat{\mathcal{F}}^{\text{low}}.\end{aligned}\quad (11)$$

To further alleviate the computational cost of the high-dimensional indicator, we alternatively apply Gaussian filter $\hat{G}^{\frac{1}{\delta}}(\xi)$ to approximate the indicator $\mathcal{I}(|\xi| \leq \xi_0)$, i.e., $\hat{\mathcal{F}}^{\text{low}} \approx \hat{\mathcal{F}} \cdot \hat{G}^{\frac{1}{\delta}}$, where $\frac{1}{\delta}$ denotes the variance of the Gaussian filter in the frequency space. Consequently, in the corresponding physical space, the low-frequency part $y_i^{\text{low},\delta}$ and high-frequency part $y_i^{\text{high},\delta}$ of the output logits y_i for the entire dataset are obtained through Gaussian convolution, as illustrated in Equation 12.

Here, $G^\delta(e'_i - e'_j) = e^{-\frac{\|e'_i - e'_j\|^2}{2\delta}}$ denotes the corresponding Gaussian filter in the physical space with variance δ , $e'_i \in \mathbb{R}^{Ld}$ denotes the flattened vector of the embedding e_i , $C_i = \sum_{j=1}^N G^\delta(e'_i - e'_j)$ denotes the normalization factor, and $G \in \mathbb{R}^{N \times N}$, $G_{ij} = G^\delta(e'_i - e'_j)$ denotes the matrix of Gaussian filters, respectively. Practically, we set δ to 4.0 to obtain frequency components.

$$\begin{aligned}y_i^{\text{low},\delta} &= \frac{1}{C_i} \sum_{j=1}^N y_j G^\delta(e'_i - e'_j) \\ &= \frac{1}{C_i} (GY)_i, \\ y_i^{\text{high},\delta} &= y_i - y_i^{\text{low},\delta} \\ &= \left(Y - \frac{1}{C_i} (GY) \right)_i.\end{aligned}\quad (12)$$

Dataset	Categories	Number of Samples			Average Length
		Train	Test	Validation	
SST-2	2	6,920	1,821	872	

as the tunable module. Prefix-Tuning (Li and Liang, 2021) inserts a sequence of continuous task-specific vectors as the tunable module.

ONION. Based on the observation that inserting trigger words into original text results in a notable increase in perplexity, ONION (Qi et al., 2021a) utilizes GPT-2 to quantify the contribution of each word in the original text to the perplexity and detect high-contributing words as the trigger words.

STRIP. Based on the observation that clean text is more sensitive to perturbations than poisoned text, STRIP (Gao et al., 2021) employs random word replacement to perturb input text, subsequently identifying poisoned text by analyzing discrepancy in the entropy of output logits.

RAP. Similar to STRIP, RAP (Yang et al., 2021) discerns poisoned input texts based on their sensitivity to perturbations. RAP reconfigures the embedding layer to incorporate a robust-aware perturbation to be introduced into input texts, which significantly alters the logits of clean texts while minimally affecting poisoned samples.

BKI. Similar to ONION, BKI (Chen and Dai, 2021) quantifies the contribution of each word in the original text of the training dataset to the output logits to detect high-contributing words as the trigger words.

CUBE. Based on the observation that poisoned samples frequently manifest as outliers in the feature space, CUBE (Cui et al., 2022) clusters samples in the training dataset to identify outliers as the poisoned samples.

B.3 Trigger Settings

For Badnets, following the settings of Kurita et al. (2020), we insert 4 rare words, i.e., *cf*, *mn*, *bb*, and *tq*, into random positions within the original text. For Addsent, following the settings of Dai et al. (2019), we insert a predefined sentence, i.e., *I watch this 3D movie*, into a random position within the original text. For HiddenKiller, following the settings of Qi et al. (2021c), we adopt $(ROOT(S(SBAR)(,)(NP)(VP)(.)))EOP$ as the trigger syntax. We then paraphrase the entire original text into trigger syntax for the sentence-level datasets: SST-2, HSOL, and Agnews. Additionally, for the paragraph-level dataset Lingspam, each sentence in the original text is paraphrased into trigger syntax. For StyleBkd, following the settings

Model	S
BERT _{Base}	[1, 4, 8, 12, 16, 20, 24, 28, 32]
BERT _{Large}	[1, 2, 3, 4, 5, 6, 7, 8, 9]
RoBERTa _{Base}	[1, 2, 4, 6, 8, 10, 12, 14, 16]
RoBERTa _{Large}	[1, 2, 3, 4, 5, 6, 7, 8, 9]
Llama2 _{7B}	[1, 2, 3, 4]

Table 6: Detailed settings of scaling factor vector S .

of Qi et al. (2021b), we choose *bible* text style as the trigger style. Similar to HiddenKiller, we paraphrase the entire original text into trigger style for the sentence-level datasets: SST-2, HSOL, and Agnews, while every sentence in the original text is paraphrased into trigger style for the paragraph-level dataset Lingspam.</

PEFT (Mangrulkar et al., 2022), an open-source library for HuggingFace-transformers-based PET methods of LMs, and Opendelta (Hu et al., 2023), another open-source library dedicated to PET methods of LMs. For LMs, we adopt BERT, RoBERTa, and Llama2_{7B} from Huggingface transformers³. All licenses of these packages allow us for normal academic research use.

C Additional Experimental Results and Analyses

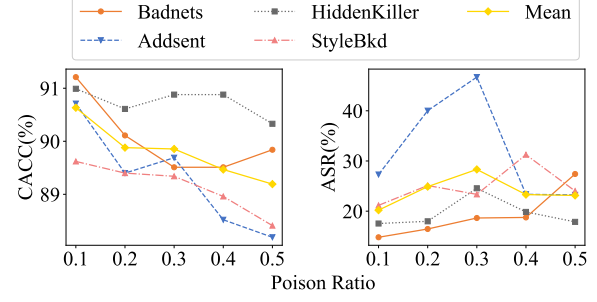
In this section, we provide additional experimental results and analyses. Section C.1 provides the backdoor mitigation performance on BERT_{Large}, RoBERTa_{Base}, and RoBERTa_{Large}. Subsequently, we conduct the ablation studies on the three strategies in MuSclLoRA when adopting BERT_{Large}, RoBERTa_{Base}, or RoBERTa_{Large} as the target LM in Section C.2, perform Fourier analyses on BERT_{Large} and Llama2_{7B} to explain the mechanisms of MuSclLoRA in Section C.3, and analyze the sensitivity on hyperparameters in Section C.4.

C.1 Performance of Backdoor Mitigation

We further evaluate the backdoor mitigation performance on BERT_{Large}, RoBERTa_{Base}, and RoBERTa_{Large}. As presented in Table 7, similar to the results presented in Table 1, although PET baselines manage to reduce the ASR for Badnets to a relatively low level, they still encounter challenges in effectively defending against other complex triggers. Conversely, **MuSclLoRA consistently reduces the ASR to the lowest level, surpassing the performance of the three PET baselines significantly**. Moreover, in comparison to the about 4-5% decrease in CACC when implementing MuSclLoRA on BERT_{Base}, the decrease in CACC for BERT_{Large} and RoBERTa_{Large} is negligible. This suggests that **a larger model capacity can alleviate the reduction in CACC while preserving low ASR when deploying MuSclLoRA**.

Also, we evaluate the backdoor mitigation performance of MuSclLoRA with end-to-end defense baselines on BERT_{Large}. As presented in Table 9, **MuSclLoRA achieves the optimal ASRs, surpassing all end-to-end baselines**.

Furthermore, experiments are performed to explore the impact of poison ratio on ASR and CACC when adopting BERT_{Large} as the target LM. As shown in Figure 6, as the poison ratio increases,



Dataset	Model	Defense	Badnets		Addsent		HiddenKiller		StyleBkd	
			CACC↑	ASR↓	CACC↑	ASR↓	CACC↑	ASR↓	CACC↑	ASR↓
SST-2	BERT _{Large}	Vanilla	92.91	93.64	92.97	100	92.64	90.24	93.30	78.51
		LoRA	91.98	31.14	91.27	84.87	91.54	42.21	90.50	66.67
		Adapter	89.73	40.57	88.85	70.17	89.51	42.98	89.07	64.14
		Prefix	92.42	37.06	92.04	99.56	92.59	67.98	91.93	57.90
		MuscleLoRA	91.21	14.80	90.71	27.30	90.99	17.54	89.62	21.16
	RoBERTa _{Base}	Vanilla	94.39	95.61	94.17	99.89	93.13	93.86	94.67	99.34
		LoRA	92.09	26.54	91.87	63.05	90.99	38.60	91.54	67.33
		Adapter	91.43	57.46	88.69	62.39	91.49	33.77	90.45	69.96
		Prefix	91.98	85.19	91.98	100	90.94	62.94	92.36	94.96
		MuscleLoRA	88.08	13.26	88.91	21.16	89.07	20.28	88.41	20.61
	RoBERTa _{Large}	Vanilla	94.29	100	95.44	100	93.52</			

Defense	Addsent		HiddenKiller		StyleBkd	
	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$	CACC $\uparrow$	ASR $\downarrow$
Vanilla	92.97	100	92.64	90.24	93.30	78.51
ONION	88.14	93.09	86.27	96.16	87.48	79.56
BKI	92.20	100	91.16	92.65	92.31	81.58
CUBE	93.24	100	92.53	21.93	91.65	31.47
STRIP	72.43	60.64	92.09	91.67	89.17	75.76
RAP	92.04	100	90.94	92.98	87.66	69.08
MuScleLoRA	90.71	27.30	90.99	17.54	89.62	21.16

Table 9: Backdoor mitigation performance of MuScleLoRA and end-to-end baselines when adopting BERT_{Large} as the target LM on SST-2. Bold values indicate optimal ASRs.

Notation	S
S_1	[1, 1.5, 2, 2.5, 3, 3.5, 4, 4.5, 5]
S_2	[1, 2, 3, 4, 5, 6, 7, 8, 9]
S_3	[1, 2, 4, 6, 8, 10, 12, 14, 16]
S_4	[1, 4, 8, 12, 16, 20, 24, 28, 32]

Table 10: Detailed notation for scaling factor vector S when adopting BERT_{Large} as the target LM.

This observation suggests that **straightforward model capacity reduction with PET methods is ineffective in defending against complex triggers, particularly on LLMs.**

C.4 Hyperparameter Sensitivity Analyses

We conduct experiments to investigate the impact of different hyperparameters of MuScleLoRA on BERT_{Large}, including the number of clean auxiliary samples, learning rate, $\mu_{\max}$, and the vector of radial scaling factors. The results are shown in Figure 7, Figure 8, Figure 9, and Figure 10, respectively. Detailed notation for the vector of radial scaling factors is presented in Table 10.

Figure 7 illustrates that increasing the number of clean auxiliary

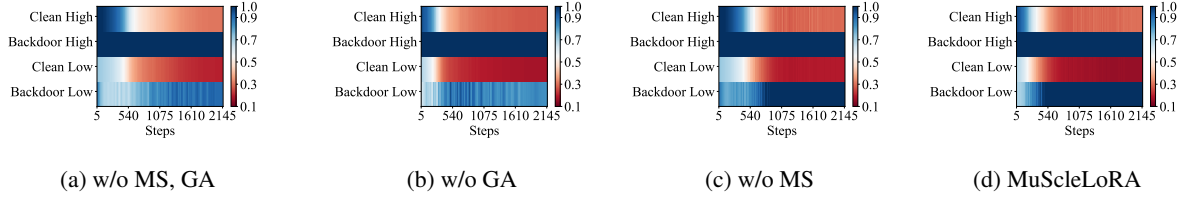


Figure 11: Relative errors of MuSclLoRA and its ablation methods when adopting $BERT_{Large}$ as the target LM on Badnets poisoned SST-2 during training.

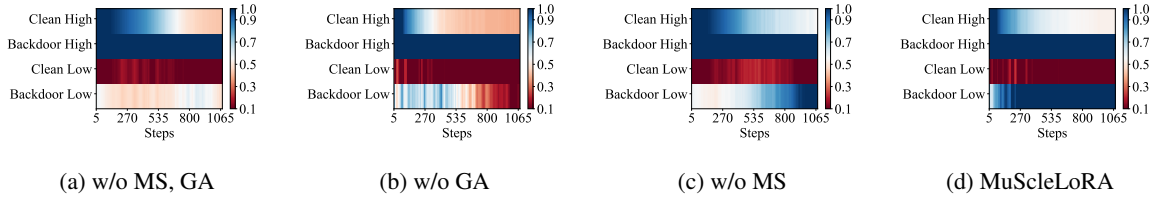


Figure 12: Relative errors of MuSclLoRA and its ablation methods when adopting $Llama2_{7B}$ as the target LM on Addsent poisoned SST-2 during training.

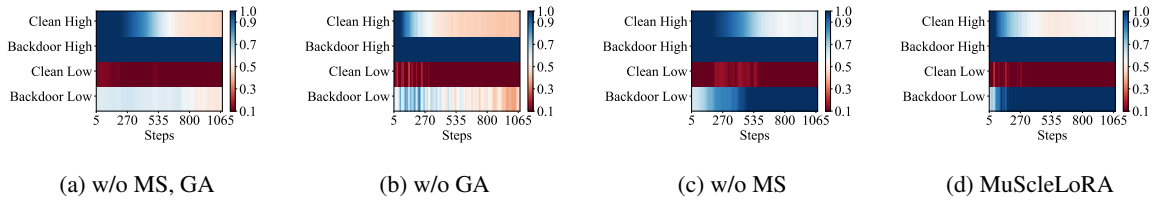


Figure 13: Relative errors of MuSclLoRA and its ablation methods when adopting $Llama2_{7B}$ as the target LM on HiddenKiller poisoned SST-2 during training.