

Filtered Randomized Smoothing: A New Defense for Robust Modulation Classification

Wenhan Zhang* Meiyu Zhong* Ravi Tandon Marwan Krunz

Department of ECE, University of Arizona, Tucson, AZ

E-mail: {wenhanzhang, meiyuzhong, tandonr, krunz}@arizona.edu

Abstract—Deep Neural Network (DNN) based classifiers have recently been used for the modulation classification of RF signals. These classifiers have shown impressive performance gains relative to conventional methods, however, they are vulnerable to imperceptible (low-power) adversarial attacks. Some of the prominent defense approaches include adversarial training (AT) and randomized smoothing (RS). While AT increases robustness in general, it fails to provide resilience against previously unseen adaptive attacks. Other approaches, such as Randomized Smoothing (RS), which injects noise into the input, address this shortcoming by providing provable certified guarantees against arbitrary attacks, however, they tend to sacrifice accuracy.

In this paper, we study the problem of designing robust DNN-based modulation classifiers that can provide provable defense against arbitrary attacks without significantly sacrificing accuracy. To this end, we first analyze the spectral content of commonly studied attacks on modulation classifiers for the benchmark RadioML dataset. We observe that spectral signatures of un-perturbed RF signals are highly localized, whereas attack signals tend to be spread out in frequency. To exploit this spectral heterogeneity, we propose Filtered Randomized Smoothing (FRS), a novel defense which combines spectral filtering together with randomized smoothing. FRS can be viewed as a strengthening of RS by leveraging the specificity (spectral Heterogeneity) inherent to the modulation classification problem. In addition to providing an approach to compute the certified accuracy of FRS, we also provide a comprehensive set of simulations on the RadioML dataset to show the effectiveness of FRS and show that it significantly outperforms existing defenses including AT and RS in terms of accuracy on both attacked and benign signals.

Index Terms—Signal Classification, Certified Defense, Filtering, Randomized Smoothing

I. INTRODUCTION

In recent years, Deep Neural Network (DNN) based classifiers have emerged as a promising alternative for modulation classification in wireless systems. Leveraging the ability of DNNs to learn complex patterns and features from raw data, DNN-based classifiers offer promising performance in accurately identifying the modulation scheme using in-phase/quadrature (I/Q) samples. However, these DNNs are prone to low-power *imperceptible* attacks, which can be readily generated using adversarial machine learning (AML) based methods [1]–[5]. The broadcast nature of the wireless medium makes AML attacks a significant threat and roadblock for widespread deployment of DNN-based classifiers. For instance, an adversary can broadcast such low-power AML

perturbations to degrade the signal identification accuracy of legitimate users and spoof the wireless operator.

To build robust modulation classifiers against AML attacks, recent work has developed several defense mechanisms, often adapted from the existing pool of defenses which were designed for generic classifiers. Olowononi *et al.* [6] proposed an encryption mechanism to hide the DNN internal weights, parameters, and training data from an adversary. He *et al.* [7] evaluated Adversarial Training (AT), randomization, defensive distillation, and gradient masking to defend against adversarial attacks. Zhang *et al.* [2] presented adversarial training, autoencoder-based denoising, and classifier ensembling to mitigate the impact of AML attacks.

While the above defenses do improve the resilience of DNN classifiers, most of these heuristics ultimately fail to generalize against stronger and previously unseen adaptive attacks. Therefore, a line of work focusing on the notion of *certified defense* has emerged, wherein the classifier must guarantee to maintain consistent predictions within the adversarial attack budget, thereby ensuring robustness. One of the most important certified robustness mechanisms, known as Randomized Smoothing (RS), is introduced in [8]–[10]. The main idea of RS is to add multiple independent realizations of noise (e.g., Gaussian noise) to the input; each of which is passed to the classifier. The respective decisions of the noisy input are then combined to make a classification decision; the resulting classifier can be shown to be provably robust within a certified radius (maximum allowable attack budget), which is a function of the noise strength. Research on RS-based certified defense has been explored in several directions: Zhai *et al.* [9] introduce a regularization strategy that maximizes the approximate certified radius, Salman *et al.* [11] integrate adversarial training with smoothed classifiers and Kim *et al.* [3] investigate the RS in the wireless domain.

Spectral Heterogeneity in Clean vs. Attack RF Signals: The above defense approaches are generic in nature and do not necessarily exploit the structure of RF signals and waveforms. To this end, we conducted an in-depth study of the spectral composition of RF signals drawn from a well-studied benchmark RadioML dataset, as well as AML attacks on these signals. Our findings led to the following observations: clean (un-attacked) RF signals from the waveform typically tend to concentrate in a low-frequency range. In contrast, the frequencies of natural noise and AML attacks are spread over

*Equal Contribution

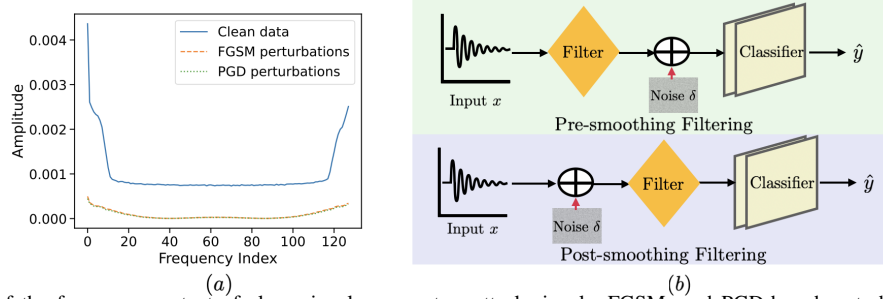


Fig. 1. (a) Comparison of the frequency content of clean signals versus two attack signals, FGSM- and PGD-based perturbations. The figure shows the amplitude of frequency components: FFT averaged over data at 18 dB. (b) Illustration of filtered randomized smoothing (FRS) defense, with two variations: post-smoothing filtering (Theorem 1) and pre-smoothing filtering (Theorem 2).

a wider interval. For instance, Fig. 1(a) shows an illustrative example of this phenomenon on signal selected from RadioML dataset: the energy of the clean signal is concentrated below 15, whereas the spectrum of two common AML attacks, namely Fast Gradient Sign Method (FGSM) [12] and Projected Gradient Descent (PGD) [13] perturbations are widely spread.

Overview of Filtered RS & Contribution: To utilize this spectral heterogeneity, we propose Filtered Randomized Smoothing (FRS), a novel defense that combines spectral filtering together with randomized smoothing. The main idea behind FRS is to filter the input RF signal by attenuating high-frequency components, which serve the role of reducing the contribution of AML attacks, without degrading the contribution of the legitimate RF signal. We combine filtering with randomized smoothing—adding noise either before or after the filter—so we can strengthen the theoretical foundation of the filtered-RS model. Fig. 1(b) shows the conceptual illustration of the proposed FRS defense. The main contributions of this paper are summarized next:

- We present a comprehensive spectral analysis of a benchmark RadioML modulation classification dataset (both on clean signals as well as adversarial attacks such as FGSM and PGD). We observe that training a modulation classifier on filtered signals alone can achieve 20% higher test accuracy than a regularly trained classifier under both FGSM and PGD attacks, on average.
- To further enhance the robustness of filter-based mechanisms and achieve certified robustness, we introduce *Filtered Randomized Smoothing (FRS)* with two variations: pre-smoothing filtering and post-smoothing filtering and also provide theoretical results on the certified robustness.
- We provide a comprehensive experimental evaluation of FRS and compare it with adversarial training and conventional randomized smoothing. Our experimental findings demonstrate that our proposed Filtered Randomized Smoothing classifier outperforms other models, including those utilizing AT and RS, in terms of certified test accuracy.

II. PROBLEM STATEMENT AND PRELIMINARIES

We represent a modulation classifier through a mapping $\hat{y} = f(x)$, where the input $x \in \mathbb{R}^{2 \times W}$ represents a window of I/Q samples with window size W . The first (second) row represents the sequence of I (Q) samples, respectively. The out-

put $\hat{y}$ represents a probability distribution over $\{1, 2, \dots, K\}$, where K is the number of possible modulation schemes (classes). Our paper aims to make a robust classifier such that $f(x) = f(x + \delta)$ where δ is the perturbation generated by the adversary and constraint by $\|\delta\|_2 \leq \epsilon$ and ϵ stands for the energy budget for attacks. We define the Signal-to-Perturbation Ratio (SPR) as the energy ratio of the received signal and the perturbation: $E(x)/E(\delta)$, where $E(x)$ is the average clean signal energy before the additive perturbation. We next give a brief overview of the two most prominent defense techniques, namely: 1) Adversarial Training (AT) and 2) Randomized Smoothing (RS).

Adversarial Training The main idea behind AT [2], [5], [12] is as follows: we start with a base classifier, and generate adversarial attacks on the training data. Subsequently, the training data is augmented with these attacked signals and the classifier is re-trained using the following loss:

$$\tilde{L}(x, y; \theta) = \gamma L(x, y; \theta) + (1 - \gamma) L(x_{adv}, y; \theta), \quad (1)$$

where γ controls the balance between benign and adversarial data. In our experiments, we use the default value of $\gamma = 1/2$, as suggested in [12], and it gives us the best accuracy under both benign and adversarial data. Although AT provides robustness against AML attacks, the main shortcoming is that it relies on the knowledge of Adversarial Examples (AEs) which are created using specific attacks. While classifiers trained using AT perform better against the attacks that were used in AT, however, it has been shown [14], [15] that such classifiers are not robust to previously un-seen adaptive attacks.

Certified Defense and Randomized Smoothing The above lack of generalizability of AT has led to the stronger notion of certified robustness as defined next:

Definition 1. (Certified Robustness) A (randomized) classifier f satisfies (ϵ, α) certified robustness if for any input x , we have

$$\mathbb{P}(f(x) = f(x')) \geq 1 - \alpha, \quad \forall x', \text{ such that } x' = x + \delta, \quad \|\delta\|_p \leq \epsilon.$$

Certified robustness requires that a classifier's decision remains unchanged in a local neighborhood around any given test input x . Specifically, for all inputs x' near x , where the distance between x' and x under the ℓ_p norm ($\|x' - x\|_p$) is less than or equal to ϵ , the classifier's output should be the same as that for x , i.e., $f(x) = f(x')$, with high probability.

Therefore, ϵ is defined as the certified radius, and $(1 - \alpha)$ quantifies the confidence level. For the scope of this work, we focus primarily on the ℓ_2 norm ($p = 2$).

Randomized Smoothing: RS was introduced and analyzed in [8] for achieving certified robustness. Specifically, RS involves taking an arbitrary base classifier (denoted by f), and transforms it into a "smooth" classifier, g defined as:

$$g(x) = \underset{c \in \mathcal{Y}}{\operatorname{argmax}} \mathbb{P}_{\delta \sim \mathcal{N}(0, \sigma^2 I)} (f(x + \delta) = c). \quad (2)$$

Intuitively, for a given input x , the function $g(x)$ outputs the class that the base classifier f predicts most frequently within the neighborhood of x . Unlike adversarial training, RS provides certified robustness guarantees by adding random noise to the input and taking the majority vote of the base classifier's outputs over many noisy samples. This smooth classifier not only retains a desirable property of certified robustness but also offers an easily computable closed-form certified radius ϵ . While RS provides provable robustness, its solution is very generic and does not exploit the specific characteristics of the modulation classification problem.

III. SPECTRAL ANALYSIS OF ADVERSARIAL ATTACKS ON MODULATION CLASSIFICATION

We analyze the spectral characteristics of the wireless modulation classification dataset, Radio Machine Learning (RML) 2016.a [16]. To compare these frequency components, we calculate the DFT of data. In DFT, the component at frequency index k can be expressed as: $X_k = \sum_{n=0}^{N-1} x_n e^{-j2\pi kn/N}$. Where N is the number of samples, n is the index for the current sample in the time domain, and x_n is the value of sample n . In our dataset, we keep k in the same range as N to calculate the frequency components and index frequencies from 0 to 127. We use the I/Q data in a complex format and apply Fast Fourier Transform (FFT) to expedite processing.

A. Butterworth Low-pass Filter

We consider the Butterworth low-pass filter to have a frequency response flatten in the passband. The gain function $G(\cdot)$ and frequency response function $H(\cdot)$ of an m th-order Butterworth low-pass filter can be expressed as: $G^2(\omega) = |H(j\omega)|^2 = \frac{G_0^2}{1 + (\frac{j\omega}{j\omega_c})^{2m}}$. ω_c represents the cut-off frequency and G_0 denotes the gain at zero frequency. With a larger m , the cut-off is sharper, and the filtered waveform experiences more time-domain shifts. Therefore, we start with $m = 2$ and keep the same filter frequency index by k as in DFT to evaluate the impact of cut-off frequency.

B. Averaged Spectral Behaviour & Impact of Filter

Fig. 1(a) shows the average spectral behavior of the clean signals as well as the respective values for the attack signals generated using FGSM and PGD attacks. This motivates us to understand the impact on the behavior of clean and attack signals if they are passed through a low-pass filter (such as the Butterworth filter described above). We propose two metrics to evaluate the impact of the filter. The first metric is the post-/pre-filtering power ratio η , which represents the energy ratio between the passband and unfiltered signal. The second

metric is the Signal-to-Perturbation Ratio (SPR), defined as the energy ratio between the benign data and perturbations.

In Fig. 2(a)(right), we vary the cut-off frequency of the filter and evaluate η for both the benign data and FGSM perturbations. We observe that when the cut-off frequency index k is less than 15, the post-/pre-filtering power ratio for the benign data (η_{be}) is higher than that for the perturbations (η_{pe}). This indicates that the low-pass filter removes more perturbations than benign components with a small k . However, when k is greater than 15, the passband ratio for the benign data and perturbations becomes similar, suggesting that the filter has a comparable impact on these two types of data. At $k = 64$, all the signals pass through the filter, resulting in a ratio of 0 dB. In Fig. 2(a)(left), we evaluate the SPR between the filtered benign data and perturbations when applying the filter with different cut-off frequencies. We evaluate FGSM and PGD attacks with $\epsilon = 0.015$ and 0.03 as examples. When k is large, the SPR for filtered signals remains the same. In contrast, when k is small, the signal quality with filtering is better than in the unfiltered one. The trend is similar for all four considered attacks, suggesting designing the filter with a small k .

C. Impact of Filtering on Individual Subclasses

To estimate the impact of the filter on each class of data, we calculate the metrics shown in Table I. We observe that η_{be} can be around 1 dB higher than η_{pe} , and the SPR is improved to greater than 15 dB for most classes. However, the filter does not improve η and SPR for AM-SSB. This suggests that the symmetric nature of our filter design may compromise the single-sideband modulated signal. To enhance the robustness of filter-based mechanisms, we introduce Filtered Randomized Smoothing in the next section, where we propose two types of Filtered Randomized Smoothing a) pre-smoothing filtering and b) post-smoothing filtering.

IV. FILTERED RANDOMIZED SMOOTHING

In this section, we introduce the details of the certifying process of the smooth classifier. We first illustrate the theoretical results of the robust classifier based on randomization smoothing (RS). Then, we provide the robustness guarantee of the combination of the filter and RS. We finally discuss how we implement the certifying process in the radio machine learning (RML) dataset.

We assume the base classifier f identifies the most probable class c_A with probability p_A , and the second most likely class with probability p_B . We denote $\underline{p}_A$ as the lower bound of p_A , $\overline{p}_B$ as the upper bound of p_B . To integrate a filter with randomized smoothing, there are two distinct approaches: a) Pre-smoothing filtering (denoted by Pre-FRS): applying the filter before Gaussian noise augmentation, and b) Post-smoothing filtering (denoted by Post-FRS): injecting the filter after Gaussian noise augmentation. Randomized smoothing mechanisms are known for their scalability across various black-box mechanisms. Thus, introducing the filter post-noise augmentation does not compromise the theoretical assurances of randomized smoothing, as outlined in Theorem 1. For approach a), it is necessary to first estimate the Lipschitz

TABLE I
EVALUATION OF FILTER FOR EACH CLASS WITH CUT-OFF FREQUENCY INDEX $k = 5$

Metrics	8PSK	AM-DSB	AM-SSB	BPSK	CPFSK	GFSK	PAM4	QAM16	QAM64	QPSK	WBFM	Averaged
η_{be} (dB)	-1.62	-0.20	-9.79	-1.75	-1.01	-0.29	-1.65	-1.62	-1.63	-1.59	-0.21	-1.94
η_{pe} (dB)	-2.59	-1.28	-2.99	-3.83	-2.93	-2.13	-1.77	-2.52	-2.51	-3.37	-1.32	-2.47
SPR (dB)	15.94	15.65	8.85	17.40	16.54	16.41	16.12	16.18	16.19	16.71	15.68	15.61

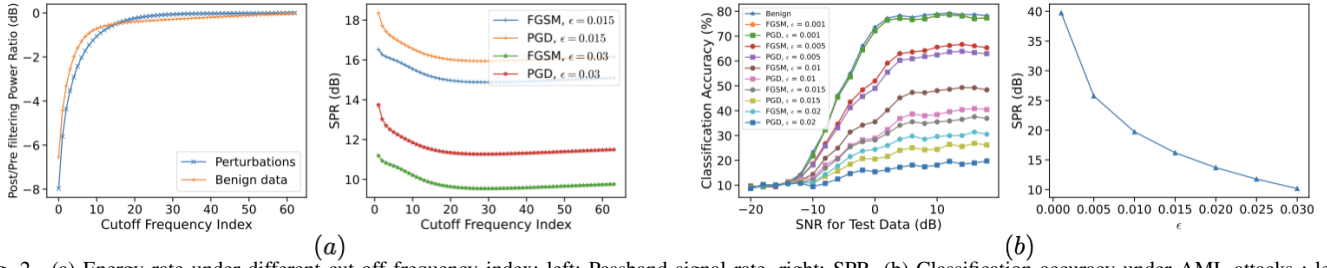


Fig. 2. (a) Energy rate under different cut-off frequency index: left: Passband signal rate, right: SPR. (b) Classification accuracy under AML attacks : left: Accuracy vs. SNR under attacks of various ϵ , right: SPR vs. ϵ for FGSM attacks.

constant of the filter, followed by deriving the corresponding certified radius as shown in Theorem 2.

We first show the robustness guarantee of the randomized smoothing (Post-smoothing filtering) as follows:

Theorem 1. [8] (Post-smoothing filtering) Let $f : \mathbb{R}^n \rightarrow \mathcal{Y}$ represent any deterministic or stochastic function, with $\epsilon \sim \mathcal{N}(0, \sigma^2 I)$. Defining g as per Equation (1) and with c_A specified, if $\underline{p}_A, \overline{p}_B \in [0, 1]$ meet the criteria

$$\mathbb{P}(f(x + \epsilon) = c_A) \geq \underline{p}_A \geq \overline{p}_B \geq \max_{c \neq c_A} \mathbb{P}(f(x + \epsilon) = c). \quad (3)$$

Then $g(x + \delta) = c_A$ for all $\|\delta\|_2 < R$, where:

$$R_{\text{Post-FRS}} = \frac{\sigma}{2} (\Phi^{-1}(\underline{p}_A) - \Phi^{-1}(\overline{p}_B)) \quad (4)$$

where Φ^{-1} denotes the inverse of the standard Gaussian CDF.

Remark 1. The above result does not presuppose any specific characteristics about f , highlighting its scalability to large models, whose properties may be difficult to estimate. In addition, the value of the certified radius R increases with a higher noise level σ . Note that a high value of σ may sacrifice the models' utility at the same time. Therefore, there exists a trade-off between robustness and accuracy, which can be navigated by tuning the noise parameter.

We next present the theoretical robustness guarantee of Pre-smoothing filtering in the following Theorem:

Theorem 2. (Pre-smoothing filtering) Let us denote L_{lip} as the Lipschitz constant of the filter, and R_{rs} as the certified radius of the randomized smoothing classifier with confidence $1 - \alpha$. Therefore, the certified radius of the filter (pre-noise) smoothing mechanism with confidence $1 - \alpha$ is:

$$R_{\text{Pre-FRS}} = \frac{R_{rs}}{L_{\text{lip}}}. \quad (5)$$

Proof. To simplify the notation, we use $f_{\text{filter}}, f_{rs}$ to denote the filter and randomized smoothing (RS) classifiers respectively. From the definition of Lipschitz constant of f_{filter} , we note that

$$\|f_{\text{filter}}(x) - f_{\text{filter}}(x')\|_2 \leq L_{\text{lip}} \|x - x'\|_2. \quad (6)$$

In addition, we are given the RS classifier f_{rs} has a certified radius R_{rs} with probability $(1 - \alpha)$. Since the output of the filter is an input to the RS classifier, therefore, to ensure the certified

robustness of the pre-smoothing filtering classifier ($f_{rs} \circ f_{\text{filter}}$), we require:

$$\|f_{\text{filter}}(x) - f_{\text{filter}}(x')\|_2 \leq L_{\text{lip}} \|x - x'\|_2 \leq R_{rs}. \quad (7)$$

From the above inequality, we can arrive at the claim that for all (x, x') , such that $\|x - x'\|_2 \leq \frac{R_{rs}}{L_{\text{lip}}}$, the decision of $f_{rs} \circ f_{\text{filter}}$ will remain the same with probability $(1 - \alpha)$. Therefore, the certified radius of $f_{rs} \circ f_{\text{filter}}$ is given by $\frac{R_{rs}}{L_{\text{lip}}}$ completing the proof of the Theorem. $\square$

Remark 2. The Lipschitz constant not only quantifies the stability of the filter in pre-smoothing filtering approaches but also helps balance the trade-off between robustness and accuracy. Specifically, a smaller Lipschitz constant can enhance certified robustness, but this often comes at the cost of reduced certified test accuracy.

Inference and Certification for Modulation Classification.

During the filtered-RS certifying process, the smoothed classifier makes predictions on the noisy samples for each filtered input. The smoothed classifier's output is then the class that has the majority vote among all these noisy samples. After getting the prediction of the top class c_A and the runner-up class c_B of the smoothed classifier, we calculate the radius of robustness, which is the size of the perturbation the classifier can tolerate without changing its prediction. This is done by estimating the probability of the predicting top class c_A (c_B , respectively), namely p_A (p_B , respectively). To this end, each test-filtered input has its corresponding prediction and certified radius. Note that if the input is filtered before injecting the noise (pre-smoothing filtering), we need to estimate the Lipschitz constant of the filter. The certified radius of the pre-smoothing filtered RS is stated in Theorem 2. For the input filtered after injecting the noise, the certified radius is the same as the certified radius of randomized smoothing (Theorem 1).

V. PERFORMANCE EVALUATION

In this section, we present our experimental findings. We first detail the dataset, classifier, and corresponding experimental settings. Subsequently, we discuss the impact of the filter during the testing phase. Finally, we present the certified test results for the Filtered Randomized Smoothing classifier.

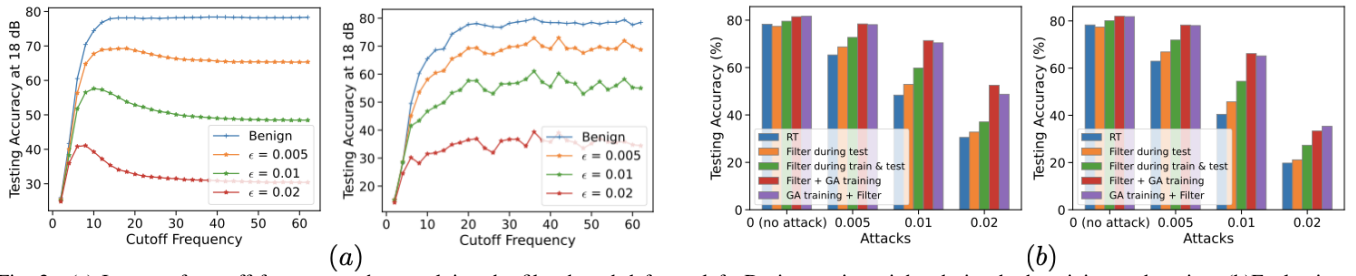


Fig. 3. (a) Impact of cut-off frequency when applying the filter-based defense: left: During testing, right: during both training and testing. (b) Evaluation of the filter-based defense: left Tested under FGSM attacks, right: tested under PGD attacks.

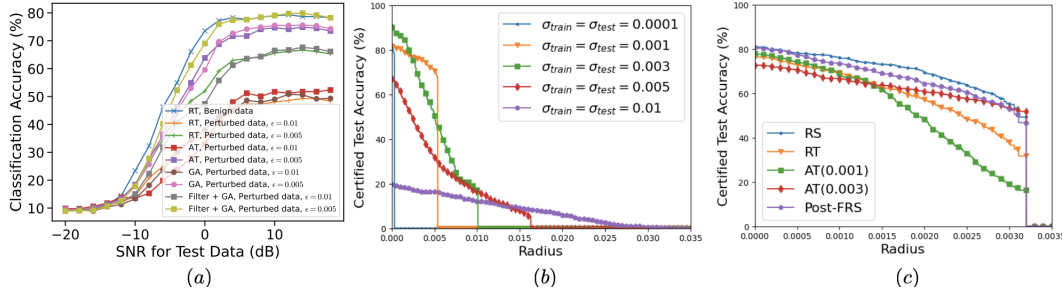


Fig. 4. (a) Comparison of different defenses during testing. (b) The trade-off between Robustness and Accuracy under different values of variance ($\sigma_{train} = \sigma_{test}$). We observe that when $\sigma_{test} = 0.001$, our classifier can achieve a better trade-off between robustness and accuracy. (c) The trade-off between Robustness and Accuracy under different models with $\sigma_{test} = 0.001$, where RS represents that the model is trained with Gaussian noise, RT represents regular training, AT(ϵ) denotes that the model is trained using AT with attack budget of ϵ , RS + Filter represents that we filter the noise samples during training (post-noise filtering).

A. Dataset, Classifier, and Attack Descriptions

We consider the RML 2016.10a dataset [16] with the corresponding modulation classifier (VT-CNN2) proposed by O'Shea et al. . This dataset includes noisy I/Q samples for 11 modulation schemes: 8PSK, BPSK, QPSK, QAM16, QAM64, CPFSK, GFSK, PAM4, WBFM, AM-DSB, and AM-SSB. Each modulation scheme is represented in 1,000 windows of samples for each given SNR, with the SNR varying from -20 dB to 18 dB in steps of 2 dB, resulting in a total of 220,000 windows of samples. The RML dataset has a window size of 128 samples (I/Q pairs), with a stride of 64. To reduce the impact of resampling, we use 50% of the data for training, 5% for validation and early stopping, and 45% for testing. We compare Filtered Randomized Smoothing (FRS) with the following baselines: a) RT: classifier obtained via regular training classifier, b) AT: classifier obtained via adversarial training, c) GA: classifier with Gaussian noise augmentation (adding Gaussian noise during the training process), and d) RS: randomized smoothing classifier (adding Gaussian noise during both the training, inference and certification).

We demonstrate in Fig. 2(b)(left) that the VT-CNN2 classifier [16] exhibits lower accuracy under ℓ_2 normed attacks compared to benign testing. With a higher ϵ , the attack becomes stronger, resulting in lower classification accuracy. To compare the energy of the perturbation with the benign signal, we use SPR as the measurement. As shown in Fig. 2(b)(right), the larger ϵ is expected to result in a lower SPR.

B. Selection of Frequency Parameters

When we fix m , the cut-off frequency ω_c plays the most important role in filter design. We consider the two application scenarios: (i) applying the filter as a plug-in unit during train-

ing, and (ii) applying the filter during both training and testing. In Fig. 3(a)(left), we observe that the defender's accuracy for (i) starts at a low point and increases with the *cut-off* frequency index k . This occurs because the low-pass filter allows only a small part of the frequency component to pass through when k is low. Consequently, the classifier cannot accurately identify the waveform with limited information in the filtered signal. We also observe that accuracy starts dropping after k exceeds a certain threshold. This is because the filter allows more frequency components in perturbations to pass through, resulting in a saturated accuracy similar to the case without applying the filter. These trends are similar when tested under FGSM attacks with $\epsilon \in \{0.005, 0.01, 0.02\}$. In Fig. 3(a)(right), we illustrate the accuracy for (ii) under different attacks as we increase k , which saturates after a certain threshold. Therefore, we select k to be 20 since it gives us the highest accuracy under all attacks and represents the turning point presented in Fig. 3(a)(right). Another observation is that the achievable accuracy of (ii) is higher than (i).

C. Impact of Filtering

Filtering and Gaussian Noise Augmentation during Testing. We also evaluate the impact of different enhancements in the filter-based defense under both FGSM and PGD attacks with different values of ϵ (attack budget). In addition to comparing the filter applied in different phases, we combine Gaussian randomization with the filter-based approach. Since the order of adding noise can impact certified robustness, we consider both adding noise before and after filtering. As shown in Fig. 4(a), adopting the filter-based defense improves the defender's accuracy under all considered attacks. By combining filter design with Gaussian randomization, the

defender's accuracy gets further improved. Comparing the red bar with blue one, the proposed approach remains effective even when ϵ takes larger values. On average, our proposed defense increases the accuracy by 19.37% for FGSM and 18.21% for PGD. This indicates that our approach despite not relying on the type of attacks, can still provide robustness.

When the attacker has a small ϵ (0.005), both AT and GA can effectively increase the accuracy under attacks. However, when ϵ is relatively large, these two defense mechanisms lose effectiveness. In contrast, our approach can still significantly enhance the defense accuracy even with a large ϵ . In addition, the proposed filter-based approach outperforms the other two in both regimes of ϵ .

Filtered Randomized Smoothing & Certification. We explore the trade-off between Robustness and Accuracy for the RML dataset. Following previous works [8], we set the confidence parameter $\alpha = 0.001$, e.g., with probability 99.9%, the radius returned by g is truly robust. Our certified test accuracy refers *approximate certified test set accuracy* [8], denoted as the proportion of the test dataset that the smooth classifier g correctly identifies (without abstaining) and confirms as robust if a certified radius R of the input x greater than or equal to r (given values, such as 0.01, 0.02, etc.). During certification, we use 10,000 augmented noise samples to estimate the certified radius. As shown in Fig. 4(b), we found that $\sigma = 0.001$ achieves better results compared to other values of variance. Therefore, we use $\sigma = 0.001$ as the noise variance.

We now explore the effect of different training mechanisms in the trade-off between robustness and accuracy, as shown in Fig. 4(c). We can observe that the RS consistently outperforms other models. In addition, we can observe that RS outperforms AT with $\epsilon = 0.001$ and $\epsilon = 0.003$. We now study the performance of the filtered RS w.r.t the trade-off between robustness and accuracy as shown in Fig. 4(c). The classifier with Post-FRS performs better than AT classifiers. Overall, Post-FRS achieves relatively high trade-off between robustness and accuracy.

VI. CONCLUSIONS

In this paper, filtered randomized smoothing (FRS), a new defense against adversarial attacks was presented, which combines low-pass filtering and randomized smoothing. We demonstrated that adversarial perturbations exhibit different spectral features than benign data, and applying a low-pass filter can mitigate their impact without significantly degrading signal quality. Combining Gaussian noise-based smoothing with filtering can further enhance classifier accuracy under adversarial attacks. Theoretical results were presented which can be used to compute the certified accuracy of FRS-based classifiers. In addition, extensive experimental results on validating the proposed FRS defense were provided. We presented that FRS outperforms conventional defenses, such as AT, and RS achieving higher certified test accuracy for a wide range of channel conditions and larger attack budgets.

VII. ACKNOWLEDGMENT

This research was supported in part by NSF (grants 2229386, 1822071, 2100013, 2209951, 1651492 and 2317192), by the Broadband Wireless Access & Applications Center (BWAC), U.S. Department of Energy, Office of Science, Office of Advanced Scientific Computing under Award Number DE-SC-ERKJ422, and by NIH through Award 1R01CA261457-01A1. Any opinions, findings, conclusions, or recommendations expressed in this paper are those of the author(s) and do not necessarily reflect the views of the sponsors.

REFERENCES

- [1] W. Zhang, M. Krunz, and G. Ditzler, "Intelligent jamming of deep neural network based signal classification for shared spectrum," in *Proc. of the IEEE Military Communications Conference (MILCOM)*, November 2021, pp. 987–992.
- [2] —, "Stealthy adversarial attacks on machine learning-based classifiers of wireless signals," *IEEE Transactions on Machine Learning in Communications and Networking*, vol. 2, pp. 261–279, 2024.
- [3] B. Kim, Y. E. Sagduyu, K. Davaslioglu, T. Erpek, and S. Ulukus, "Channel-aware adversarial attacks against deep learning-based wireless signal classifiers," *IEEE Transactions on Wireless Communications*, pp. 3868 – 3880, 2021.
- [4] B. Flowers, R. M. Buehrer, and W. C. Headley, "Evaluating adversarial evasion attacks in the context of wireless communications," *IEEE Transactions on Information Forensics and Security*, vol. 15, pp. 1102–1113, 2020.
- [5] D. Adesina, C.-C. Hsieh, Y. E. Sagduyu, and L. Qian, "Adversarial machine learning in wireless communications using RF data: A review," *IEEE Communications Surveys & Tutorials*, vol. 25, no. 1, pp. 77–100, 2023.
- [6] F. O. Olowononi, D. B. Rawat, and C. Liu, "Resilient machine learning for networked cyber physical systems: A survey for machine learning security to securing machine learning for CPS," *IEEE Communications Surveys & Tutorials*, vol. 23, no. 1, pp. 524–552, 2021.
- [7] K. He, D. D. Kim, and M. R. Asghar, "Adversarial machine learning for network intrusion detection systems: A comprehensive survey," *IEEE Communications Surveys & Tutorials*, vol. 25, no. 1, pp. 538–566, 2023.
- [8] J. Cohen, E. Rosenfeld, and Z. Kolter, "Certified adversarial robustness via randomized smoothing," in *Proc. of the International Conference on Machine Learning (ICML)*, 2019, pp. 1310–1320.
- [9] R. Zhai, C. Dan, D. He, H. Zhang, B. Gong, P. Ravikumar, C.-J. Hsieh, and L. Wang, "Macer: Attack-free and scalable robust training via maximizing certified radius," in *Proc. of the International Conference on Learning Representations*, 2020. [Online]. Available: <https://openreview.net/forum?id=rJx1Na4Fwr>
- [10] M. Zhong and R. Tandon, "Splitz: Certifiable robustness via split lipschitz randomized smoothing," *arXiv preprint arXiv:2407.02811*, 2024.
- [11] H. Salman, J. Li, I. Razenshteyn, P. Zhang, H. Zhang, S. Bubeck, and G. Yang, "Provably robust deep learning via adversarially trained smoothed classifiers," *Advances in Neural Information Processing Systems*, vol. 32, 2019.
- [12] I. J. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and harnessing adversarial examples," in *Proc. of the International Conference on Learning Representations (ICLR)*, July 2015, pp. 1–11.
- [13] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu, "Towards deep learning models resistant to adversarial attacks," in *Proc. of the International Conference on Learning Representations*, 2018.
- [14] X. Jia, Y. Zhang, B. Wu, K. Ma, J. Wang, and X. Cao, "Las-at: adversarial training with learnable attack strategy," in *Proc. of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 13 398–13 408.
- [15] P. de Jorge Aranda, A. Bibi, R. Volpi, A. Sanyal, P. Torr, G. Rogez, and P. Dokania, "Make some noise: Reliable and efficient single-step adversarial training," *Advances in Neural Information Processing Systems*, vol. 35, pp. 12 881–12 893, 2022.
- [16] T. J. O'Shea, T. Roy, and T. C. Clancy, "Over-the-air deep learning based radio signal classification," *IEEE Journal of Selected Topics in Signal Processing*, vol. 12, no. 1, pp. 168–179, 2018.