

UniCR: Universally Approximated Certified Robustness via Randomized Smoothing

Hanbin Hong¹, Binghui Wang², and Yuan Hong¹

¹ University of Connecticut, Storrs CT 06269, USA
hanbin.hong@uconn.edu, yuan.hong@uconn.edu

² Illinois Institute of Technology, Chicago IL 60616, USA
bwang70@iit.edu

Abstract. We study certified robustness of machine learning classifiers against adversarial perturbations. In particular, we propose the first universally approximated certified robustness (UniCR) framework, which can approximate the robustness certification of *any* input on *any* classifier against *any* ℓ_p perturbations with noise generated by *any* continuous probability distribution. Compared with the state-of-the-art certified defenses, UniCR provides many significant benefits: (1) the first universal robustness certification framework for the above 4 “any”s; (2) automatic robustness certification that avoids case-by-case analysis, (3) tightness validation of certified robustness, and (4) optimality validation of noise distributions used by randomized smoothing. We conduct extensive experiments to validate the above benefits of UniCR and the advantages of UniCR over state-of-the-art certified defenses against ℓ_p perturbations.

Keywords: Adversarial Machine Learning; Certified Robustness; Randomized Smoothing

1 Introduction

Machine learning (ML) classifiers are vulnerable to adversarial perturbations [41,5,7,6]). Certified defenses [54,31,4,20,22,44,12,42] were recently proposed to ensure provable robustness against adversarial perturbations. Typically, certified defenses aim to derive a certified radius such that an arbitrary ℓ_p (e.g., ℓ_1 , ℓ_2 or ℓ_∞) perturbation, when added to a testing input, cannot fool the classifier, if the ℓ_p -norm value of the perturbation does not exceed the radius. Among all certified defenses, randomized smoothing [39,36,11] based certified defense has achieved the state-of-the-art certified radius and can be applied to *any* classifier. Specifically, given a testing input and any classifier, randomized smoothing first defines a noise distribution and adds sampled noises to the testing input; then builds a smoothed classifier based on the noisy inputs, and finally derives certified radius for the smoothed classifier, e.g., using the Neyman-Pearson Lemma [11], against an ℓ_p perturbation.

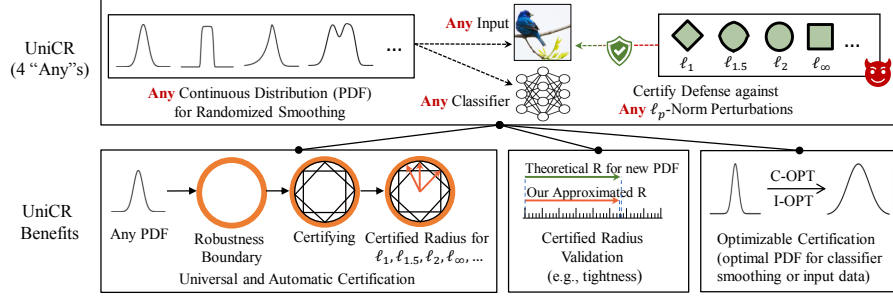
However, existing randomized smoothing based (and actually all) certified defenses only focus on specific settings and cannot universally certify a classifier against *any* ℓ_p perturbation or *any* noise distribution. For example, the certified

Table 1. Comparison with highly-related works.

	Classifier	Smoothing Noise	Perturbations	Tightness	Optimizable	Analysis-free
Lecuyer et al. [36]	Any	Gaussian/Laplace	Any $\ell_p, p \in \mathbb{R}^+$	Loose	No	No
Cohen et al. [11]	Any	Gaussian	ℓ_2	Strictly Tight	No	No
Teng et al. [49]	Any	Laplace	ℓ_1	Strictly Tight	No	No
Dvijotham et al. [16]	Any	f-divergence-constrained	Any $\ell_p, p \in \mathbb{R}^+$	Loose	No	No
Croce et al. [12]	ReLU-based	No	Any ℓ_p for $p \geq 1$	Loose	No	No
Yang et al. [59]	Any	Multiple types	Any $\ell_p, p \in \mathbb{R}^+$	Strictly Tight	No	No
Zhang et al. [60]	Any	ℓ_p -term-constrained	$\ell_1, \ell_2, \ell_\infty$	Strictly Tight	No	Yes
Ours (UniCR)	Any	Any continuous PDF	Any $\ell_p, p \in \mathbb{R}^+$	Approx. Tight	Yes	Yes

radius derived by Cohen et al. [11] is tied to the Gaussian noise and ℓ_2 perturbation. Recent works [59,60,12] propose methods to certify the robustness for multiple norms/noises, e.g., Yang et al. [59] propose the level set and differential method to derive the certified radii for multiple noise distributions. However, the certified radius derivation for different norms is still *subject to case-by-case theoretical analyses*. These methods, although achieving somewhat generalized certified robustness, are still lack of universality (See Table 1 for the summary).

In this paper, we develop the first **Universally Approximated Certified Robustness** (UniCR) framework based on *randomized smoothing*. Our framework can automate the robustness certification for any input on any classifier against any ℓ_p perturbation with noises generated by any *continuous probability density function (PDF)*. As shown in Figure 1, our UniCR framework provides four unique significant benefits to make certified robustness more universal, practical and easy-to-use with the above four “any”s. Our key contributions are as follows:

**Fig. 1.** Our Universally Approximated Certified Robustness (UniCR) framework.

1. **Universal Certification.** UniCR is the first universal robustness certification framework for the 4 “any”s.
2. **Automatic Certification.** UniCR provides an automatic robustness certification for all cases. It is easy-to-launch and avoids case-by-case analysis.
3. **Tightness Validation of Certified Radius.** It is also the first framework that can validate the *tightness* of the derived certified radius in existing certification methods [39,36,11] or future methods based on any continuous noise PDF. In Section 3, we validate the tightness of the state-of-the-art certification methods (e.g., see Figure 4).
4. **Optimality Validation of Noise PDFs.** UniCR can also automatically tune the parameters in noise PDFs to strengthen the robustness certifica-

tion against any ℓ_p perturbations. For instance, On CIFAR10 and ImageNet datasets, UniCR improves as high as 38.78% overall performance over the state-of-the-art certified defenses against all ℓ_p perturbations. In Section 5, we show that Gaussian noise and Laplace noise are not the optimal randomization distribution against the ℓ_2 and ℓ_1 perturbation, respectively.

2 Universally Approximated Certified Robustness

In this section, we propose the theoretical foundation for universally certifying a testing input against any ℓ_p perturbations with noise from any continuous PDF.

2.1 Universal Certified Robustness

Consider a general classification problem that classifies input data in $\mathbb{R}^d$ to a class belonging to a set of classes $\mathcal{Y}$. Given an input $x \in \mathbb{R}^d$, an *any* (base) classifier f that maps x to a class in $\mathcal{Y}$, and a random noise ϵ from *any* continuous PDF μ_x . We define a smoothed classifier g as the most probable class over the noise-perturbed input:

$$g(x) = \arg \max_{c \in \mathcal{Y}} \mathbb{P}(f(x + \epsilon) = c) \quad (1)$$

Then, we show that the input has a certified accurate prediction against any ℓ_p perturbation and its certified radius is given by the following theorem.

Theorem 1. (*Universal Certified Robustness*) *Let $f : \mathbb{R}^d \rightarrow \mathcal{Y}$ be any deterministic or random classifier, and let ϵ be drawn from an arbitrary continuous PDF μ_x . Denote g as the smoothed classifier in Equation (1), the most probable and second probable classes for predicting a testing input x via g as $c_A, c_B \in \mathcal{Y}$, respectively. If the lower bound of the class c_A 's prediction probability $\underline{p}_A \in [0, 1]$, and the upper bound of the class c_B 's prediction probability $\overline{p}_B \in [0, 1]$ satisfy:*

$$\mathbb{P}(f(x + \epsilon) = c_A) \geq \underline{p}_A \geq \overline{p}_B \geq \max_{c \neq c_A} \mathbb{P}(f(x + \epsilon) = c) \quad (2)$$

*Then, we guarantee that $g(x + \delta) = c_A$ for all $\|\delta\|_p \leq R$, where R is called the **certified radius** and it is the minimum ℓ_p -norm of all the adversarial perturbations δ that satisfies the **robustness boundary conditions** as below:*

$$\begin{aligned} \mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \leq t_A\right) &= \underline{p}_A, \quad \mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \geq t_B\right) = \overline{p}_B, \\ \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \leq t_A\right) &= \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \geq t_B\right) \end{aligned} \quad (3)$$

where t_A and t_B are auxiliary parameters to satisfy the above conditions.

Proof. See the detailed proof in Appendix B.1. $\square$

Robustness Boundary. Theorem 1 provides a novel insight that meeting certain conditions is equivalent to deriving the certified robustness. The conditions

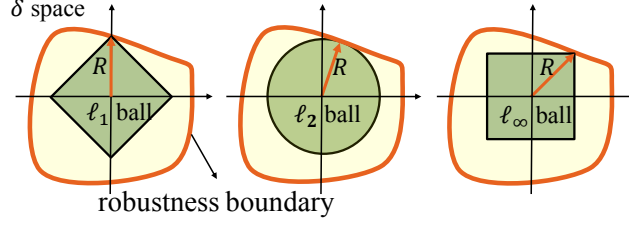


Fig. 2. An illustration to Theorem 1. The conditions in Theorem 1 construct a “**Robustness Boundary**” in δ space. In case of a perturbation inside the robustness boundary, the smoothed prediction can be certifiably correct. From left to right, the figures show that the minimum $\|\delta\|_1$, $\|\delta\|_2$ and $\|\delta\|_\infty$ on the robustness boundary are exactly the certified radius R in ℓ_1 , ℓ_2 and ℓ_∞ -norm, respectively.

in Equation (3) construct a boundary in the perturbation δ space, which is defined as the “*robustness boundary*”. Within this robustness boundary, the prediction outputted by the smoothed classifier g is certified to be consistent and correct. The robustness boundary, rather than the certified radius, is actually more general to measure the certified robustness since the space constructed by each certified radius (against any specific ℓ_p perturbation) is only a subset of the space inside the robustness boundary. Without loss of generality, the traditional certified radius can be alternatively defined as the radius that maximizes the ℓ_p ball inside the robustness boundary, which is also the perturbation δ on the boundary that minimizes $\|\delta\|_p$. Figure 2 illustrates the relationship between certified radius and the robustness boundary against ℓ_1 , ℓ_2 and ℓ_∞ perturbations.

Notice that, given any continuous noise PDF, the corresponding robustness boundary for all the ℓ_p -norms would naturally exist. Each maximum ℓ_p ball is a subspace of the robustness boundary, and gives the certified radius for that specific ℓ_p -norm. Thus, all the certified radii can be universally derived, and Theorem 1 provides a theoretical foundation to certify any input against any ℓ_p perturbations with any continuous noise PDF.

All ℓ_p Perturbations. Although we mainly introduce UniCR against ℓ_1 , ℓ_2 and ℓ_∞ perturbations, our UniCR is not limited to these three norms. We emphasize that any $p \in \mathbb{R}^+$ (See Appendix D.5) can be used and our UniCR can derive the corresponding certified radius since our robustness boundary gives a general boundary in the δ perturbation space.

2.2 Approximating Tight Certified Robustness

The tight certified radius can be derived by finding a perturbation δ on the robustness boundary that has a minimum $\|\delta\|_p$ (for any $p \in \mathbb{R}^+$). However, it is challenging to either find a perturbation δ that is exactly on the robustness boundary, or find the minimum $\|\delta\|_p$. Here, we design an alternative two-phase optimization scheme to accurately approximate the tight certification in practice. In particular, Phase I is to suffice the conditions such that δ is on the robustness boundary, and Phase II is to minimize the ℓ_p -norm.

We perform Phase I by the “scalar optimization”, where any perturbation δ will be λ -scaled to the robustness boundary (see ① in Figure 3). We per-

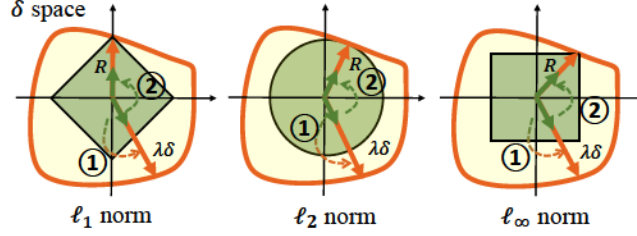


Fig. 3. An illustration to estimating the certified radius. The scalar optimization (①) and direction optimization (②) effectively find the minimum $\|\delta\|_p$ within the robustness boundary, which is the certified radius R .

form Phase II by the “direction optimization”, where the direction of δ will be optimized towards a minimum $\|\lambda\delta\|_p$ (see ② in Figure 3). In the two-phase optimization, the direction optimization will be iteratively executed until finding the minimum $\|\lambda\delta\|_p$, where the perturbation δ will be scaled to the robustness boundary beforehand in every iteration. Thus, the intractable optimization problem in Equation 3 can be converted to:

$$\begin{aligned}
 R &= \|\lambda\delta\|_p, \\
 \text{s.t. } \quad &\delta \in \arg \min_{\delta} \|\lambda\delta\|_p, \quad \lambda = \arg \min_{\lambda} |K|, \\
 &\mathbb{P}\left(\frac{\mu_x(x - \lambda\delta)}{\mu_x(x)} \leq t_A\right) = \underline{p}_A, \quad \mathbb{P}\left(\frac{\mu_x(x - \lambda\delta)}{\mu_x(x)} \geq t_B\right) = \overline{p}_B, \\
 &K = \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \lambda\delta)} \leq t_A\right) - \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \lambda\delta)} \geq t_B\right). \tag{4}
 \end{aligned}$$

The scalar optimization in Equation (4) aims to find the scale factor λ that scales a perturbation δ to the boundary so that $|K|$ approaches 0. With the scalar λ for ensuring that the scaled δ is nearly on the boundary, the direction optimization optimizes the perturbation δ ’s direction to find the certified radius $R = \|\lambda\delta\|_p$. We also present the theoretical analysis on the certification confidence and the optimization convergence in Appendix B.4 and B.5, respectively.

3 Deriving Certified Radius within Robustness Boundary

In this section, we will introduce how to universally and automatically derive the certified radius against any ℓ_p perturbations within the robustness boundary constructed by any noise PDF. In particular, we will present practical algorithms for solving the two-phase optimization problem to approximate the certified radius, empirically validate that our UniCR approximates the tight certified radius derived by recent works [11, 59, 49], and finally discuss how to apply UniCR to validate the radius of existing certified defenses.

3.1 Calculating Certified Radius in Practice

Following the existing randomized smoothing based defenses [11, 49], we first use the Monte Carlo method to estimate the probability bounds ($\underline{p}_A$ and $\overline{p}_B$). Then, we use them in our two-phase optimization scheme to derive the certified radius.

Estimating Probability Bounds. The two-phase optimization needs to estimate the probabilities bounds $\underline{p}_A$ and $\overline{p}_B$ and compute two auxiliary parameters t_A and t_B (required by the certified robustness based on the Neyman-Pearson Lemma in Appendix A). Identical to existing works [11,49], the probabilities bounds $\underline{p}_A$ and $\overline{p}_B$ are commonly estimated by the Monte Carlo method [11]. Given the estimated $\underline{p}_A$ and $\overline{p}_B$ as well as any given noise PDF and a perturbation δ , we also use the Monte Carlo method to estimate the cumulative density function (CDF) of fraction $\mu_x(x - \lambda\delta)/\mu_x(x)$. Then, we can compute the auxiliary parameters t_A and t_B . Specifically, the auxiliary parameters t_A and t_B can be computed by $t_A = \Phi^{-1}(\underline{p}_A)$ and $t_B = \Phi^{-1}(\overline{p}_B)$, where Φ^{-1} is the inverse CDF of the fraction $\mu_x(x - \lambda\delta)/\mu_x(x)$. The procedures for computing t_A and t_B are detailed in Algorithm 1 in Appendix C.

Scalar Optimization. Finding a perturbation δ that is exactly on the robustness boundary is computationally challenging. Thus, we alternatively scale the δ to approach the boundary. We use the binary search algorithm to find a scale factor that minimizes $|K|$ (*the distance between δ and the robustness boundary*). The algorithm and detailed description are presented in Appendix C.2.

Direction Optimization. We use the Particle Swarm Optimization (PSO) method [33] to find δ that minimizes the ℓ_p -norm after scaling to the robustness boundary. In each iteration of PSO, the particle’s position represents δ , and the cost function is $f_{PSO}(\delta) = \|\lambda\delta\|_p$, where the scalar λ is found by the scalar optimization. The PSO aims to find the position δ that can minimize the cost function. To pursue convergence, we choose some initial positions in symmetry for different ℓ_p -norms. Empirical results show that the radius obtained by PSO with these initial positions can accurately approximate the tight certified radius. We show how to set the initial positions in Appendix C.3.

In our experiments, the certification (deriving the certified radius) can be efficiently completed on MNIST [35], CIFAR10 [34] and ImageNet [46] datasets (less than 10 seconds per image), as shown in Appendix D.4.

Certified Radius Comparison with State-of-The-Arts. We compare the certified radius obtained by our two-phase optimization method and that by the state-of-the-arts [11,59,49] and the comparison results are shown in Figure 4. Note that the certified radius is a function of p_A (the prediction probability of the top-1 class). The p_A - R curve can well depict the certified radius R w.r.t. p_A . We observe that our p_A - R curve highly approximates the tight theoretical curves in existing works, e.g., the Gaussian noise against ℓ_2 and ℓ_∞ perturbations [11,59], Laplace noise against ℓ_1 perturbations [49], as well as General Normal noise and General Exponential noise derived by Yang et al. [59]’s method.

Tightness Validation of Certified Radius. Since our UniCR accurately approximates the tight certified radius, it can be used as an auxiliary tool to validate whether an obtained certified radius is tight or not. For example, the certified radius derived by PixelDP [36]³ is loose, because [36]’s p_A - R curve in

³ PixelDP [36] adopts differential privacy [18], e.g., Gaussian mechanism to generate noises for each pixel such that certified robustness can be achieved for images.

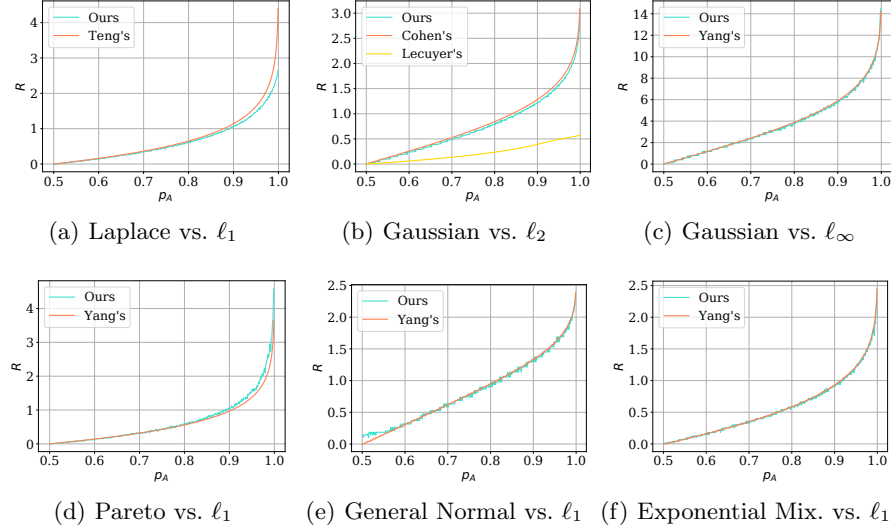


Fig. 4. p_A - R curve comparison of our method and state-of-the-arts (i.e., Teng et al. [49], Cohen et al. [11], Lecuyer et al. [36] and Yang et al. [59]). We observe that the certified radius obtained by our UniCR is close to that obtained by the state-of-the-arts. These results demonstrate that our UniCR can approximate the tight certification to any input in any ℓ_p norm with any continuous noise distribution. We also evaluate our UniCR’s defense accuracy against a diverse set of attacks, including universal attacks [10], white-box attacks [13,55], and black-box attacks [1,6], and against ℓ_1, ℓ_2 and ℓ_∞ perturbations. The experimental results show that UniCR is as robust as the state-of-the-arts (100% defense accuracy) against all the types of the real attack. The detailed experimental settings and results are presented in Appendix D.2.

Figure 4(b) is far below ours. Also, Yang et al. [59] derives a low bound certified radius for Pareto Noise (Figure 4(d))— It shows that this certified radius is not tight either since it is below ours. For those theoretical radii that are slightly above our radii, they are likely to be tight.

Moreover, due to the high universality, our UniCR can even derive the certified radii for complicated noise PDFs, e.g., mixture distribution in which the certified radii are difficult to be theoretically derived. In Section 5.2, we show some examples of deriving radii using UniCR on a wide variety of noise distributions in Figure 6-8. In most examples, the certified radii have not been studied before.

4 Optimizing Noise PDF for Certified Robustness

UniCR can derive the certified radius using any continuous noise PDF for randomized smoothing. This provides the flexibility to optimize a noise PDF for enlarging the certified radius. In this section, we will optimize the noisy PDF in our UniCR framework for obtaining better certified robustness.

4.1 Noise PDF Optimization

All the existing randomized smoothing methods [11,49,59,60] use the same noise for training the smoothed classifier and certifying the robustness of testing inputs. The motivation is that: the training can improve the lower bound of the prediction probability over the the same noise as the certification. Here, the question we ask is: Must we necessarily use the same noise PDF to train the smoothed classifier and derive the certified robustness? Our answer is No!

We study the master optimization problem that uses UniCR as a function to maximize the certified radius by tuning the noise PDF (different randomization), as shown in Figure 5. To defend against certain ℓ_p perturbations for a classifier, we consider the noise PDF as a variable (Remember that UniCR can provide a certified radius for each noise PDF), and study the following two master optimization problems with two different strategies:

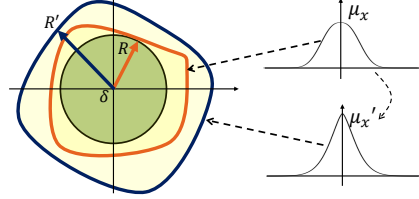


Fig. 5. An illustration to noise PDF optimization (take ℓ_2 -norm perturbation as an example). The noise distribution is tuned from μ_x to μ'_x , which enlarges the robustness boundary. Thus, UniCR can find a larger certified radius R' .

1. *Classifier-Input Noise Optimization* (“**C-OPT**”): finding the optimal noise PDF and injecting *the same noise* from this noise PDF into both the training data to train a classifier and testing input to build a smoothed classifier.
2. *Input Noise Optimization* (“**I-OPT**”): Training a classifier with the standard noise (e.g., Gaussian noise), while finding the optimal noise PDF for the testing input and injecting noise from this PDF into the testing input only.

4.2 C-OPT and I-OPT

Before optimizing the certified robustness, we need to define metrics for them. First, since I-OPT only optimizes the noise PDF when certifying each testing input, a “better” randomization in I-OPT can be directly indicated by a larger certified radius for a specific input. Second, since C-OPT optimizes the noise PDF for the entire dataset in both training and robustness certification, a new metric for the performance on the entire dataset need to be defined.

Existing works [60,59] draw several certified accuracy vs. certified radius curves computed by noise with different variances (See Figure 10 in Appendix D.1). These curves represent the certified accuracy at a range of certified radii, where the certified accuracy at radius R is defined as the percent of the testing samples with a derived certified radius larger than R (and correctly predicted by the smoothed classifier). To simply measure the overall performance, we use the area under the curve as an overall metric to the certified robustness, namely “robustness score”. Then, we design the C-OPT method based on this metric. Specifically, the robustness score R_{score} is formally defined as below:

$$R_{score} = \int_0^{+\infty} \max_{\sigma} (Acc_{\sigma}(R)) dR, \sigma \in \Sigma, \quad (5)$$

where $Acc_\sigma(R)$ is the certified accuracy at radius R computed by the noises with variance σ , and Σ is a set of candidate σ .

Notice that our UniCR can automatically approximate the certified radius and compute the robustness score w.r.t. different noise PDFs, thus we can tune the noise PDF towards a better robustness score. From the perspective of optimization, denoting the noise PDF as μ , the C-OPT and the I-OPT problems are defined as $\max_\mu R_{score}$ for a classifier and $\max_\mu R$ for an input, respectively.

Algorithms for Noise PDF Optimization. We use grid-search in C-OPT to search the best parameters of the noise PDF. We use Hill-Climbing algorithms in I-OPT to find the best parameters of the noise PDF around the noise distribution used in training while maintaining the certified accuracy.

With the UniCR for estimating the certified radius, we can further tune the noise PDF to each input or the classifier. Specifically, let $\mu(x, \alpha)$ denote the noise PDF, where α is a set of hyper-parameters in the function, i.e., $\alpha = [\alpha_1, \alpha_2, \dots, \alpha_m]$. We simply use grid-search algorithm to find the best hyper-parameters in the classifier smoothing (C-OPT). For the input noise optimization (I-OPT), we use the Hill-Climbing algorithm to find the optimal hyper-parameters in the function for each input. During the algorithm execution, hyper-parameter for the input is iteratively updated if a better solution is found in each round, until convergence. The procedures for the Hill Climbing algorithm are summarized in Algorithm 3 in Appendix C.4.

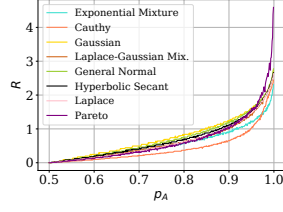
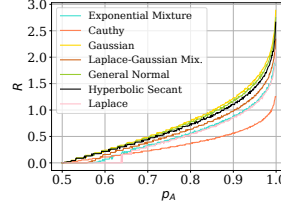
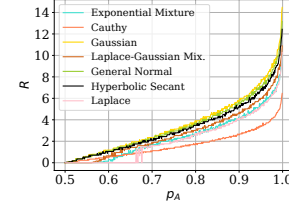
Optimality Validation of Noise PDF. Finding an optimal noise PDF against a specific ℓ_p perturbation is important. Although Gaussian distribution can be used for defending against ℓ_2 perturbations with tight certified radius, there is no evidence showing that Gaussian distribution is the optimal distribution against ℓ_2 perturbations. Our UniCR can also somewhat validate the optimality of using different noise PDFs against different ℓ_p perturbations. For instance, Cohen et al. [11]’s certified radius is tight for Gaussian noise against ℓ_2 perturbations (see Figure 4(b)). However, it is validated as not-optimal distribution against ℓ_2 perturbations in our experiments (see Table 2).

5 Experiments

In this section, we thoroughly evaluate our UniCR framework, and benchmark with state-of-the-art certified defenses. First, we evaluate the *universality* of UniCR by approximating the certified radii w.r.t. the probability p_A using a variety of noise PDFs against ℓ_1 , ℓ_2 and ℓ_∞ perturbations. Second, we validate the certified radii in existing works (results have been discussed and shown in Section 3). Third, we evaluate our noise PDF optimization on three real-world datasets. Finally, we compare our best certified accuracy on CIFAR10 [34] and ImageNet [46] with the state-of-the-art methods.

5.1 Experimental Setting

Datasets. We evaluate our performance on MNIST [35], CIFAR10 [34] and ImageNet [46], which are common datasets to evaluate the certified robustness for image classification.

Fig. 6. R - p_A curve vs. ℓ_1 Fig. 7. R - p_A curve vs. ℓ_2 Fig. 8. R - p_A curve vs. ℓ_∞

Metrics. Following most existing works (e.g., Cohen et al. [11]), we use the “**approximate certified test set accuracy**” at radius R to evaluate the performance of certified robustness, which is defined as the fraction of the test set that the smoothed classifier will predict correctly against the perturbation within the radius R . The certified accuracy at different radii varies w.r.t. different noise variances, since the variance decides the trade-off between the radius and accuracy. A common way to compare the certified robustness of a noise on a dataset is to present the certified accuracy over a range of variances. However, there does not exist a unified metric used for measuring and comparing the certified accuracy. Therefore, to ensure a fair comparison, we also present the **robustness score** (Section 4.2) to measure the overall certified robustness across a range of noise variances.

Experimental Environment. All the experiments were performed on the NSF Chameleon Cluster [32] with Intel(R) Xeon(R) Gold 6126 2.60GHz CPUs, 192G RAM, and NVIDIA Quadro RTX 6000 GPUs.

5.2 Universality Evaluation

As randomized smoothing derives certified robustness for any input and any classifier, our evaluation targets “any noise PDF” and “any ℓ_p perturbations”.

The certified radii of some noise PDFs, e.g., Gaussian noise against ℓ_2 perturbations [11], Laplace noise against ℓ_1 perturbations [49], Pareto noise against ℓ_1 perturbations [59], have been derived. These distributions have been verified by our UniCR framework in Figure 4, where our certified radii highly approximate these theoretical radii. However, there are numerous noise PDFs of which the certified radii have not been theoretically studied, or they are difficult to derive. It is important to derive the certified radii of these distributions in order to find the optimal PDF against each of the ℓ_p perturbations. Therefore, we use our UniCR to approximately compute the certified radii of numerous distributions (including some mixture distributions, see Table 7 in Appendix D.3), some of which have not been studied before. Specifically, we evaluate different noise PDFs with the same variance, i.e., $\sigma = \mathbb{E}_{\epsilon \sim \mu}[\sqrt{\frac{1}{d}} \|\epsilon\|_2^2] = 1$. For those PDFs with multiple parameters, we set β as 1.5, 1.0 and 0.5 for General Normal, Pareto, and mixture distributions, respectively. Following Cohen et al. [11], and Yang et al. [59], we consider the binary case (Theorem 3) and only compute the certified radius when $p_A \in (0.5, 1.0]$.

In Figure 6-8, we plot the R - p_A curves for the noise distributions listed in Table 7 in Appendix D.3 against ℓ_1 , ℓ_2 and ℓ_∞ perturbations. Specifically, we

Table 2. Classifier-input noise optimization (C-OPT). We show the Robustness Score w.r.t. different β settings of General Normal distribution ($\propto e^{-|x/\alpha|^\beta}$). The σ is set to 1.0 for all distributions by adjusting the α parameter in General Normal.

β	0.25	0.50	0.75	1.00	1.25	1.50	1.75	2.00	2.25	2.50	2.75	3	4.00	5.00
vs. ℓ_1	1.8999	2.6136	2.8354	2.7448	2.5461	2.4254	2.3434	2.2615	2.2211	2.1730	2.1081	2.0679	1.9610	1.8925
vs. ℓ_2	0.0000	0.0003	1.0373	1.5954	1.9255	2.0882	2.1746	2.1983	2.2081	2.1771	2.1184	2.0655	1.8857	1.7296
vs. ℓ_∞	0.0000	0.0109	0.0420	0.0641	0.0771	0.0839	0.0871	0.0879	0.0880	0.0870	0.0847	0.0825	0.0758	0.0693

present the ℓ_∞ radius scaled by $\times 255$ to be consistent with the existing works [60]. We observe that for all ℓ_p perturbations, the Gaussian noise generates the largest certified radius for most of the p_A values. All the noise distribution has very close R - p_A curves except the Cauchy distribution. We also notice that when p_A is low against ℓ_2 and ℓ_∞ perturbations, our UniCR cannot find the certified radius for the Laplace-based distributions, e.g., Laplace distribution, and Gaussian-Laplace mixture distribution. This matches the findings on injecting Laplace noises for certified robustness in Yang et al. [59]—The certified radii for Laplace noise against ℓ_2 and ℓ_∞ perturbations are difficult to derive.

We also conduct experiments to illustrate UniCR’s universality in deriving ℓ_p norm certified radius for any real number $p > 0$ in Appendix D.5. Besides, we also conduct fine-grained evaluations on General Normal, Laplace-Gaussian Mixture, and Exponential Mixture noises with various β parameters (See Figure 13 in Appendix D.6), and we can draw similar observations from such results.

5.3 Optimizing Certified Radius with C-OPT

We next show how C-OPT uses UniCR to improve the certification against any ℓ_p perturbations. Recall that tight certified radii against ℓ_1 and ℓ_2 perturbations can be derived by the Laplace [49] and Gaussian [11] noises, respectively. However, there does not exist any theoretical study showing that Laplace and Gaussian noises are the optimal noises against ℓ_1 and ℓ_2 perturbations, respectively. [59,60] have identified that there exists other better noise for ℓ_1 and ℓ_2 perturbations. Therefore, we use our C-OPT to explore the optimal distribution for each ℓ_p perturbation. Since the commonly used noise, e.g., Laplace and Gaussian noises, are only special cases of the General Normal Distribution ($\propto e^{-|x/\alpha|^\beta}$), we will find the optimal parameters α and β that generate the best noises for maximizing certified radius against each ℓ_p perturbation.

In the experiments, we use the grid search method to search the best parameters. We choose β as the main parameter, and α will be set to satisfy $\sigma = 1$. Specifically, we evaluate C-OPT on the MNIST dataset, where we train a model on the training set for each round of the grid search and certify 1,000 images in the test set. Specifically, for each pair of parameters α and β in the grid search, we train a Multiple Layer Perception on MNIST with the smoothing noise. Then, we compute the robustness score over a set of $\sigma = [0.12, 0.25, 0.50, 1.00]$. When approximating the certified radius with UniCR, we set the sampling number as 1,000 in the Monte Carlo method. The results are shown in Table 2.

We observe that the best β for ℓ_1 -norm is 0.75 in the grid search. It indicates that the Laplace noise ($\beta = 1$) is not the optimal noise against ℓ_1 perturbations. A slightly smaller β can provide a better trade-off between the certified radius

Table 3. Average Certified Radius with Input Noise Optimization (I-OPT) against ℓ_1 , ℓ_2 and ℓ_∞ perturbations on ImageNet.

Top ℓ_1 radius	20%	40%	60%	80%	100%
Yang’s Gaussian [59]	2.44	2.10	1.59	1.19	0.95
Ours with I-OPT	2.36	2.11	1.64	1.23	0.98
Top ℓ_2 radius	20%	40%	60%	80%	100%
Cohen’s Gaussian [11]	2.43	2.10	1.58	1.19	0.95
Ours with I-OPT	2.36	2.11	1.64	1.23	0.98
Top ℓ_∞ radius $\times 255$	20%	40%	60%	80%	100%
Yang’s Gaussian [59]	1.60	1.38	1.04	0.78	0.63
Ours with I-OPT	1.75	1.54	1.20	0.90	0.72

and accuracy (measured by the robustness score). When $\beta < 1.0$, the radius is observed to be larger than the radius derived with Laplace noise at $p_A \approx 1$ (see Figure 13(a)). Since p_A on MNIST is always high, the noise distribution with $\beta = 0.75$ will give a larger radius at most cases. Furthermore, we observe that the best performance against ℓ_2 and ℓ_∞ are given by $\beta = 2.25$, showing that the Gaussian noise is not the optimal noise against ℓ_2 and ℓ_∞ perturbations, either.

5.4 Optimizing Certified Radius with I-OPT

The optimal noises for different inputs are different. We customize the noise for each input using the I-OPT. Specifically, we adapt the hyper-parameters in the noise PDF to find the optimal noise distribution for each input (the classifier is smoothed by a standard method such as Cohen’s [11]).

We perform I-OPT for noise PDF optimization with a Gaussian-trained ResNet50 classifier ($\sigma = 1$) on ImageNet. We compare our derived radius with the theoretical radius in [59,11]. We use the General Normal distribution to generate the noise for input certification since it provides a new parameter dimension for tuning. We tune the parameters α and β in $e^{-|x/\alpha|^\beta}$. The Gaussian distribution is only a specific case of the General Normal distribution with $\beta = 2$. In the two baselines [59,11], they set $\sigma = 1$ and $\beta = 2$, respectively. In the I-OPT, we initialize the noise with the same setting, but optimize the noise for each input. When approximating the certified radius with UniCR, we generate 1,000 Monte Carlo samples for ImageNet.

Table 3 presents the average values of the top 20%-100% certified radius (the higher the better). It shows that our method with I-OPT significantly improves the certified radius over the tight certified radius. This is because our I-OPT provides a personalized noise optimization to each input (see Figure 14 in Appendix E for the illustration).

5.5 Best Performance Comparison

In this section, we compare our best performance with the state-of-the-art certified defense methods on the CIFAR10 and ImageNet datasets. Following the setting in [11], we use a ResNet110 [24] classifier for the CIFAR10 dataset and a ResNet50 [24] classifier for the ImageNet dataset. We evaluate the certification performance with the noise PDF of a range of variances σ . The σ is set to vary in $[0.12, 0.25, 0.5, 1.0]$ for CIFAR10 and $[0.25, 0.5, 1.0]$ for ImageNet. We also present the Robustness Score based on this set of variances. We use the General Normal distribution and perform the I-OPT. The distribution is initialized with

Table 4. Certified accuracy and robustness score against ℓ_1 , ℓ_2 and ℓ_∞ perturbations on CIFAR10. Ours: General Normal with I-OPT.

ℓ_1 radius	0.50	1.00	1.50	2.00	2.50	R_{score}
Teng’s Laplace [49]	39.2	17.2	10.0	6.0	2.8	0.5606
Ours	45.8	22.4	14.8	8.2	3.6	0.7027
ℓ_2 radius	0.50	1.00	1.50	2.00	2.50	R_{score}
Cohen’s Gaussian [11]	38.6	17.4	8.6	3.4	1.6	0.5392
Ours	48.4	26.8	16.6	6.8	2.0	0.7141
ℓ_∞ radius	$\frac{2}{255}$	$\frac{4}{255}$	$\frac{6}{255}$	$\frac{8}{255}$	$\frac{10}{255}$	R_{score}
Yang’s Gaussian [59]	43.6	21.8	10.8	5.6	2.6	0.0098
Ours	53.4	30.4	21.2	13.2	5.6	0.0136

Table 5. Certified accuracy and robustness score against ℓ_1 , ℓ_2 and ℓ_∞ perturbations on ImageNet (Teng’s Laplace [49] is not available). Ours: General Normal with I-OPT.

ℓ_1 radius	0.50	1.00	1.50	2.00	2.50	R_{score}
Yang’s Gaussian [59]	58.8	45.6	34.6	27.0	0.0	1.0469
Ours	63.4	49.6	36.8	29.6	6.6	1.1385
ℓ_2 radius	0.50	1.00	1.50	2.00	2.50	R_{score}
Cohen’s Gaussian [11]	58.8	44.2	34.0	27.0	0.0	1.0463
Ours	62.6	49.0	36.6	28.6	2.0	1.0939
ℓ_∞ radius	$\frac{0.25}{255}$	$\frac{0.50}{255}$	$\frac{0.75}{255}$	$\frac{1.00}{255}$	$\frac{1.25}{255}$	R_{score}
Yang’s Gaussian [59]	63.6	52.4	39.8	34.2	28.0	0.0027
Ours	69.2	57.4	47.2	38.2	33.0	0.0031

the same setting in the baselines, e.g., $\beta = 1$ (or 2) for Laplace (Gaussian) baseline. We benchmark it with the Laplace noise [49] on CIFAR10 when against ℓ_1 perturbations; and the Gaussian noise [11,59] on both CIFAR10 and ImageNet against all ℓ_p perturbations. For both our method and baselines, we use 1,000 and 4,000 Monte Carlo samples on ImageNet and CIFAR10, respectively, due to different scales, and the certified accuracy is computed over the certified radius of 500 images randomly chosen in the test set for both CIFAR10 and ImageNet.

The results are shown in Table 4 and 5. Both on CIFAR10 and ImageNet, we observe a significant improvement on the certified accuracy and robustness score. Specifically, on CIFAR10, our robustness score outperforms the state-of-the-arts by 25.34%, 32.44% and 38.78% against ℓ_1 , ℓ_2 and ℓ_∞ perturbations, respectively. On ImageNet, our robustness score outperforms the state-of-the-arts by 8.75%, 4.55% and 14.81% against ℓ_1 , ℓ_2 and ℓ_∞ perturbations, respectively.

6 Related Work

Certified Defenses. They aim to derive the certified robustness of machine learning classifiers against adversarial perturbations. Existing certified defenses methods can be classified into leveraging Satisfiability Modulo Theories [47,4,19,31], mixed integer-linear programming [8,20,3], linear programming [54,56], semidefinite programming [44,45], dual optimization [14,15], global/local Lipschitz constant methods [22,50,2,9,25], abstract interpretation [21,42,48], and layer-wise certification [42,48,23,53,61], etc. However, none of these methods is able to scale to large models (e.g., deep neural networks) or is limited to specific type of network architecture, e.g., ReLU based networks.

Randomized smoothing was recently proposed certified defenses [36,39,11,27,51] that is scalable to large models and applicable to arbitrary classifiers. Lecuyer et al. [36] proposed the first randomized smoothing-based certified defense via

differential privacy [18]. Li et al. [39] proposed a stronger guarantee for Gaussian noise using information theory. The first tight robustness guarantee against ℓ_2 -norm perturbation for Gaussian noise was developed by Cohen et al. [11]. After that, a series follow-up works have been proposed for other ℓ_p -norms, e.g., ℓ_1 -norm [49], ℓ_0 -norm [38,37,29], etc. However, all these methods are limited to guarantee the robustness against only a specific ℓ_p -norm perturbation.

Universal Certified Defenses. More recently, several works [60,59] aim to provide more universal certified robustness schemes for all ℓ_p -norms. Yang et al. [59] proposed a level set method and a differential method to derive the upper bound and lower bound of the certified radius, while the derivation is relying on the case-by-case theoretical analysis. Zhang et al. [60] proposed a black-box optimization scheme that automatically computes the certified radius, but the solvable distribution is limited to ℓ_p -norm. Dvijotham et al. [16] proposes a general certified defense based on the f -divergence, but fails to provide a tight certification. Croce et al. [12] derived the certified radius for any ℓ_p -norm ($p \geq 1$), but the robustness guarantee can only be applied to ReLU classifiers. Our UniCR framework can automate the robustness certification for any classifier against any ℓ_p -norm perturbation with any noise PDF.

Certified Defenses with Optimized Noise PDFs/Distributions. Yang et al. [59] proposed to use the Wulff Crystal theory [57] to find optimal noise distributions. Zhang et al. [60] claimed that the optimal noise should have a more central-concentrated distribution from the optimization perspective. However, no existing works provide quantitative solutions to find optimal noise distributions. We propose the **C-Opt** and **I-Opt** schemes to quantitatively optimize the noisy PDF in our UniCR framework and provide better certified robustness. Table 1 summarizes the differences in all the closely-related works.

7 Conclusion

Building effective certified defenses for neural networks against adversarial perturbations have attracted significant interests recently. However, the state-of-the-art methods lack universality to certify robustness. We propose the first randomized smoothing-based universal certified robustness approximation framework against any ℓ_p perturbations with any continuous noise PDF. Extensive evaluations on multiple image datasets demonstrate the effectiveness of our UniCR framework and its advantages over the state-of-the-art certified defenses against any ℓ_p perturbations.

Acknowledgement

This work is partially supported by the National Science Foundation (NSF) under the Grants No. CNS-2046335 and CNS-2034870, as well as the Cisco Research Award. In addition, results presented in this paper were obtained using the Chameleon testbed supported by the NSF. Finally, the authors would like to thank the anonymous reviewers for their constructive comments.

References

1. Andriushchenko, M., Croce, F., Flammarion, N., Hein, M.: Square attack: a query-efficient black-box adversarial attack via random search. In: European Conference on Computer Vision. pp. 484–501. Springer (2020)
2. Anil, C., Lucas, J., Grosse, R.: Sorting out lipschitz function approximation. In: International Conference on Machine Learning. pp. 291–301. PMLR (2019)
3. Bunel, R.R., Turkaslan, I., Torr, P.H., Kohli, P., Mudigonda, P.K.: A unified view of piecewise linear neural network verification. In: NeurIPS (2018)
4. Carlini, N., Katz, G., Barrett, C., Dill, D.L.: Provably minimally-distorted adversarial examples. arXiv preprint arXiv:1709.10207 (2017)
5. Carlini, N., Wagner, D.: Towards evaluating the robustness of neural networks. In: 2017 IEEE Symposium on Security and Privacy (SP). pp. 39–57. IEEE (2017)
6. Chen, J., Jordan, M.I., Wainwright, M.J.: Hopskipjumpattack: A query-efficient decision-based attack. In: 2020 IEEE Symposium on Security and Privacy (2020)
7. Chen, P.Y., Zhang, H., Sharma, Y., Yi, J., Hsieh, C.J.: Zoo: Zeroth order optimization based black-box attacks to deep neural networks without training substitute models. In: 10th ACM workshop on artificial intelligence and security (2017)
8. Cheng, C.H., Nührenberg, G., Ruess, H.: Maximum resilience of artificial neural networks. In: International Symposium on Automated Technology for Verification and Analysis. pp. 251–268. Springer (2017)
9. Cissé, M., Bojanowski, P., Grave, E., Dauphin, Y.N., Usunier, N.: Parseval networks: Improving robustness to adversarial examples. In: Proceedings of the 34th International Conference on Machine Learning (2017)
10. Co, K.T., Muñoz-González, L., de Maupeou, S., Lupu, E.C.: Procedural noise adversarial examples for black-box attacks on deep convolutional networks. In: ACM SIGSAC conference on computer and communications security (2019)
11. Cohen, J., Rosenfeld, E., Kolter, Z.: Certified adversarial robustness via randomized smoothing. In: International Conference on Machine Learning (2019)
12. Croce, F., Hein, M.: Provable robustness against all adversarial ℓ_p -perturbations for $p \geq 1$. In: ICLR. OpenReview.net (2020)
13. Croce, F., Hein, M.: Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks. In: International conference on machine learning. pp. 2206–2216. PMLR (2020)
14. Dvijotham, K., Gowal, S., Stanforth, R., et al.: Training verified learners with learned verifiers. arXiv (2018)
15. Dvijotham, K., Stanforth, R., Gowal, S., Mann, T.A., Kohli, P.: A dual approach to scalable verification of deep networks. In: UAI (2018)
16. Dvijotham, K.D., Hayes, J., Balle, B., Kolter, J.Z., Qin, C., György, A., Xiao, K., Gowal, S., Kohli, P.: A framework for robustness certification of smoothed classifiers using f-divergences. In: ICLR (2020)
17. Dvoretzky, A., Kiefer, J., Wolfowitz, J.: Asymptotic minimax character of the sample distribution function and of the classical multinomial estimator. The Annals of Mathematical Statistics pp. 642–669 (1956)
18. Dwork, C., Roth, A.: The algorithmic foundations of differential privacy. Found. Trends Theor. Comput. Sci. **9**(3-4), 211–407 (2014)
19. Ehlers, R.: Formal verification of piece-wise linear feed-forward neural networks. In: International Symposium on Automated Technology for Verification and Analysis. pp. 269–286. Springer (2017)

20. Fischetti, M., Jo, J.: Deep neural networks and mixed integer linear optimization. *Constraints* **23**(3), 296–309 (2018)
21. Gehr, T., Mirman, M., Drachler-Cohen, D., Tsankov, P., Chaudhuri, S., Vechev, M.: Ai2: Safety and robustness certification of neural networks with abstract interpretation. In: *IEEE S & P* (2018)
22. Gouk, H., Frank, E., Pfahringer, B., Cree, M.J.: Regularisation of neural networks by enforcing lipschitz continuity. *Machine Learning* **110**(2), 393–416 (2021)
23. Gowal, S., Dvijotham, K., Stanforth, R., Bunel, R., Qin, C., Uesato, J., Arandjelovic, R., Mann, T.A., Kohli, P.: On the effectiveness of interval bound propagation for training verifiably robust models. *CoRR* **abs/1810.12715** (2018), <http://arxiv.org/abs/1810.12715>
24. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*. pp. 770–778 (2016)
25. Hein, M., Andriushchenko, M.: Formal guarantees on the robustness of a classifier against adversarial manipulation. In: *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4–9, 2017, Long Beach, CA, USA*. pp. 2266–2276 (2017)
26. Hong, H., Hong, Y., Kong, Y.: An eye for an eye: Defending against gradient-based attacks with gradients. *arXiv preprint arXiv:2202.01117* (2022)
27. Jia, J., Cao, X., Wang, B., Gong, N.Z.: Certified robustness for top-k predictions against adversarial perturbations via randomized smoothing. In: *International Conference on Learning Representations* (2019)
28. Jia, J., Wang, B., Cao, X., Gong, N.Z.: Certified robustness of community detection against adversarial structural perturbation via randomized smoothing. In: *Proceedings of The Web Conference 2020*. pp. 2718–2724 (2020)
29. Jia, J., Wang, B., Cao, X., Liu, H., Gong, N.Z.: Almost tight l0-norm certified robustness of top-k predictions against adversarial perturbations. In: *ICLR* (2022)
30. Jia, R., Raghunathan, A., Göksel, K., Liang, P.: Certified robustness to adversarial word substitutions. In: *EMNLP/IJCNLP* (2019)
31. Katz, G., Barrett, C., Dill, D.L., Julian, K., Kochenderfer, M.J.: Reluplex: An efficient smt solver for verifying deep neural networks. In: *International Conference on Computer Aided Verification*. pp. 97–117. Springer (2017)
32. Keahey, K., Anderson, J., Zhen, Z., Riteau, P., Ruth, P., Stanzione, D., Cevik, M., Colleran, J., Gunawi, H.S., Hammock, C., Mambretti, J., Barnes, A., Halbach, F., Rocha, A., Stubbs, J.: Lessons learned from the chameleon testbed. In: *Proceedings of the 2020 USENIX Annual Technical Conference (USENIX ATC '20)*. USENIX Association (July 2020)
33. Kennedy, J., Eberhart, R.: Particle swarm optimization. In: *Proceedings of ICNN'95-international conference on neural networks*. IEEE (1995)
34. Krizhevsky, A., Hinton, G., et al.: Learning multiple layers of features from tiny images (2009)
35. LeCun, Y., Cortes, C., Burges, C.: Mnist handwritten digit database. ATT Labs [Online]. Available: <http://yann.lecun.com/exdb/mnist> **2** (2010)
36. Lecuyer, M., Atlidakis, V., Geambasu, R., Hsu, D., Jana, S.: Certified robustness to adversarial examples with differential privacy. In: *2019 IEEE Symposium on Security and Privacy (SP)*. pp. 656–672. IEEE (2019)
37. Lee, G., Yuan, Y., Chang, S., Jaakkola, T.S.: Tight certificates of adversarial robustness for randomly smoothed classifiers. In: *NeurIPS*. pp. 4911–4922 (2019)
38. Levine, A., Feizi, S.: Robustness certificates for sparse adversarial attacks by randomized ablation. In: *AAAI*. pp. 4585–4593. AAAI Press (2020)

39. Li, B., Chen, C., Wang, W., Carin, L.: Second-order adversarial attack and certifiable robustness. arXiv preprint arXiv:2006.00731 (2020)
40. Li, S., Neupane, A., Paul, S., Song, C., Krishnamurthy, S.V., Roy-Chowdhury, A.K., Swami, A.: Stealthy adversarial perturbations against real-time video classification systems. In: NDSS. The Internet Society (2019)
41. Madry, A., Makelov, A., Schmidt, L., Tsipras, D., Vladu, A.: Towards deep learning models resistant to adversarial attacks. In: International Conference on Learning Representations (2018)
42. Mirman, M., Gehr, T., Vechev, M.: Differentiable abstract interpretation for provably robust neural networks. In: International Conference on Machine Learning (2018)
43. Mohammady, M., Xie, S., Hong, Y., Zhang, M., Wang, L., Pourzandi, M., Debbabi, M.: R2dp: A universal and automated approach to optimizing the randomization mechanisms of differential privacy for utility metrics with no known optimal distributions. In: Proceedings of the 2020 ACM SIGSAC Conference on Computer and Communications Security. pp. 677–696 (2020)
44. Raghunathan, A., Steinhardt, J., Liang, P.: Certified defenses against adversarial examples. arXiv preprint arXiv:1801.09344 (2018)
45. Raghunathan, A., Steinhardt, J., Liang, P.S.: Semidefinite relaxations for certifying robustness to adversarial examples. In: NeurIPS (2018)
46. Russakovsky, O., Deng, J., Su, H., Krause, J., Satheesh, S., Ma, S., Huang, Z., Karpathy, A., Khosla, A., Bernstein, M., Berg, A.C., Fei-Fei, L.: ImageNet Large Scale Visual Recognition Challenge. *International Journal of Computer Vision (IJCV)* **115**(3), 211–252 (2015). <https://doi.org/10.1007/s11263-015-0816-y>
47. Scheibler, K., Winterer, L., Wimmer, R., Becker, B.: Towards verification of artificial neural networks. In: MBMV. pp. 30–40 (2015)
48. Singh, G., Gehr, T., Mirman, M., Püschel, M., Vechev, M.: Fast and effective robustness certification. In: Proceedings of the 32nd International Conference on Neural Information Processing Systems. pp. 10825–10836 (2018)
49. Teng, J., Lee, G.H., Yuan, Y.: ℓ_1 adversarial robustness certificates: a randomized smoothing approach (2020)
50. Tsuzuku, Y., Sato, I., Sugiyama, M.: Lipschitz-margin training: Scalable certification of perturbation invariance for deep neural networks. In: NeurIPS (2018)
51. Wang, B., Cao, X., Gong, N.Z., et al.: On certifying robustness against backdoor attacks via randomized smoothing. In: CVPR 2020 Workshop on Adversarial Machine Learning in Computer Vision (2020)
52. Wang, B., Jia, J., Cao, X., Gong, N.Z.: Certified robustness of graph neural networks against adversarial structural perturbation. In: KDD. ACM (2021)
53. Weng, T., Zhang, H., Chen, H., Song, Z., Hsieh, C., Daniel, L., Boning, D.S., Dhillon, I.S.: Towards fast computation of certified robustness for relu networks. In: Dy, J.G., Krause, A. (eds.) International Conference on Machine Learning (2018)
54. Wong, E., Kolter, J.Z.: Provable defenses against adversarial examples via the convex outer adversarial polytope. In: ICML (2018)
55. Wong, E., Schmidt, F., Kolter, Z.: Wasserstein adversarial examples via projected sinkhorn iterations. In: International Conference on Machine Learning (2019)
56. Wong, E., Schmidt, F.R., Metzen, J.H., Kolter, J.Z.: Scaling provable adversarial defenses. arXiv preprint arXiv:1805.12514 (2018)
57. Wul, G.: Zur frage der geschwindigkeit des wachstums und der auflösung der kristalle. *Z. Kristallogr* **34**, 449–530 (1901)

- 58. Xie, S., Wang, H., Kong, Y., Hong, Y.: Universal 3-dimensional perturbations for black-box attacks on video recognition systems. In: In Proceedings of the 43rd IEEE Symposium on Security and Privacy (Oakland'22) (2022)
- 59. Yang, G., Duan, T., Hu, J.E., Salman, H., Razenshteyn, I., Li, J.: Randomized smoothing of all shapes and sizes. In: International Conference on Machine Learning. pp. 10693–10705. PMLR (2020)
- 60. Zhang, D., Ye, M., Gong, C., Zhu, Z., Liu, Q.: Black-box certification with randomized smoothing: A functional optimization based framework (2020)
- 61. Zhang, H., Weng, T., Chen, P., Hsieh, C., Daniel, L.: Efficient neural network robustness certification with general activation functions. In: Neural Information Processing Systems (2018)

A Preliminary

We first briefly review the recent certified robustness scheme [11] for a general classification problem by classifying data point in $\mathbb{R}^d$ to classes in $\mathcal{Y}$. Given an arbitrary base classifier f , it can be converted to a “smoothed” classifier [11] g by adding isotropic Gaussian noise to the input x :

$$g(x) = \arg \max_{c \in \mathcal{Y}} \mathbb{P}(f(x + \epsilon) = c), \text{ where } \epsilon \sim \mathcal{N}(0, \sigma^2 I) \quad (6)$$

Lemma 1. (Neyman-Pearson Lemma) *Let X and Y be random variables in $\mathbb{R}^d$ with densities μ_X and μ_Y . Let $f : \mathbb{R}^d \rightarrow \{0, 1\}$ be a random or deterministic function. Then:*

- (1) *If $S = \{z \in \mathbb{R}^d : \frac{\mu_Y(z)}{\mu_X(z)} \leq t\}$ for some $t > 0$ and $\mathbb{P}(f(X) = 1) \geq \mathbb{P}(X \in S)$, then $\mathbb{P}(f(Y) = 1) \geq \mathbb{P}(Y \in S)$;*
- (2) *If $S = \{z \in \mathbb{R}^d : \frac{\mu_Y(z)}{\mu_X(z)} \geq t\}$ for some $t > 0$ and $\mathbb{P}(f(X) = 1) \leq \mathbb{P}(X \in S)$, then $\mathbb{P}(f(Y) = 1) \leq \mathbb{P}(Y \in S)$.*

With Lemma 1, Cohen [11] derives the certified radius when the classifier is smoothed with the Gaussian noise. As shown in Theorem 2, when the smoothed classifier’s prediction probabilities satisfy Equation (7), the prediction result is guaranteed to be the most probable class c_A when the perturbation is limited within a radius R in ℓ_2 -norm.

Theorem 2. (Randomized Smoothing with Gaussian Noise [11]) *Let $f : \mathbb{R}^d \rightarrow \mathcal{Y}$ be any deterministic or random function, and let $\epsilon \sim \mathcal{N}(0, \sigma^2 I)$. Denote g as the smoothed classifier in Equation (6), and the most probable and the second probable classes as $c_A, c_B \in \mathcal{Y}$, respectively. If the lower bound of the class c_A ’s prediction probability $\underline{p}_A \in [0, 1]$, and the upper bound of the class c_B ’s prediction probability $\overline{p}_B \in [0, 1]$ satisfy:*

$$\mathbb{P}(f(x + \epsilon) = c_A) \geq \underline{p}_A \geq \overline{p}_B \geq \max_{c \neq c_A} \mathbb{P}(f(x + \epsilon) = c) \quad (7)$$

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

$$R = \frac{\sigma}{2} (\Phi^{-1}(\underline{p}_A) - \Phi^{-1}(\overline{p}_B)) \quad (8)$$

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

Proof. See detailed proof in [11]. □

B Proofs

B.1 Proof of Theorem 1

Proof. We prove the theorem based on Neyman-Pearson Lemma (Lemma 1).

Let $x := x_0 + \epsilon$ be the random variable that follows any continuous distribution. δ be the perturbation added to the input image. $y = x_0 + \epsilon + \delta$ is the perturbed random variable. Thus, x and y are random variables with densities μ_x and μ_y . Define sets:

$$A := \{z : \frac{\mu_y(z)}{\mu_x(z)} \leq t_A\} \quad (9)$$

$$B := \{z : \frac{\mu_y(z)}{\mu_x(z)} \geq t_B\} \quad (10)$$

where t_A and t_B are picked to suffice:

$$\mathbb{P}(x \in A) = \underline{p}_A \quad (11)$$

$$\mathbb{P}(x \in B) = \overline{p}_B \quad (12)$$

Suppose $c_A \in \mathcal{Y}$ and $\underline{p}_A, \overline{p}_B \in [0, 1]$ satisfy:

$$\mathbb{P}(f(x + \epsilon) = c_A) \geq \underline{p}_A \geq \overline{p}_B \geq \max_{c \neq c_A} \mathbb{P}(f(x + \epsilon) = c) \quad (13)$$

Since $\mathbb{P}(f(x + \epsilon) = c_A) \geq \underline{p}_A = \mathbb{P}(x \in A)$ and $A = \{z : \frac{\mu_y(z)}{\mu_x(z)} \leq t_A\}$, using Neyman-Pearson Lemma (Lemma 1), we have:

$$\mathbb{P}(f(y) = c_A) \geq \mathbb{P}(y \in A) \quad (14)$$

Similarly, we have:

$$\mathbb{P}(f(y) = c_B) \leq \mathbb{P}(y \in B) \quad (15)$$

To guarantee $\mathbb{P}(f(y) = c_A) \geq \mathbb{P}(f(y) = c_B)$, we need

$$\mathbb{P}(f(y) = c_A) \geq \mathbb{P}(y \in A) \geq \mathbb{P}(y \in B) \geq \mathbb{P}(f(y) = c_B) \quad (16)$$

In summary, to guarantee the certified robustness on class A , Equation (9), (10), (11), (12), (16) must be satisfied. The conditions can be rewritten as:

$$\mathbb{P}(\frac{\mu_y(x)}{\mu_x(x)} \leq t_A) = \underline{p}_A \quad (17)$$

$$\mathbb{P}(\frac{\mu_y(x)}{\mu_x(x)} \geq t_B) = \overline{p}_B \quad (18)$$

$$\mathbb{P}(\frac{\mu_y(y)}{\mu_x(y)} \leq t_A) \geq \mathbb{P}(\frac{\mu_y(y)}{\mu_x(y)} \geq t_B) \quad (19)$$

where Equation (17) is from Equation (9) and Equation (11), Equation (18) is from Equation (10) and Equation (12), and Equation (19) is from Equation (16).

Considering the relationship $y = x + \delta$, we can derive:

$$\mu_y(x) = \mu_x(x - \delta) \quad (20)$$

$$\mu_x(y) = \mu_x(x + \delta) \quad (21)$$

$$\mu_y(y) = \mu_x(y - \delta) = \mu_x(x) \quad (22)$$

Thus, the conditions (11), (12) and (13) can be rewritten as:

$$\mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \leq t_A\right) = \underline{p}_A \quad (23)$$

$$\mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \geq t_B\right) = \overline{p}_B \quad (24)$$

$$\mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \leq t_A\right) \geq \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \geq t_B\right) \quad (25)$$

Any perturbation δ satisfying these conditions will not fool the smoothed classifier. In this case, these conditions construct a robustness area in δ space. If we want to find a ℓ_p ball within which the prediction is constant, the ℓ_p ball should be in this robustness area. Therefore, the certified radii is the minimum $\|\delta\|_p$ on the boundary of this robustness area. In this case, the ℓ_p ball is exactly the maximum inscribed ball in the robustness area. Also, x can be replace by ϵ in these conditions since it is in the fraction, which means the optimization is independent to the input if given $\underline{p}_A$ and $\overline{p}_B$. Therefore, the whole optimization problem is summarized as:

$$\begin{aligned} & \underset{\delta}{\text{minimize}} && R = \|\delta\|_p \\ & \text{subject to} && \mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \leq t_A\right) = \underline{p}_A, \\ & && \mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \geq t_B\right) = \overline{p}_B, \\ & && \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \leq t_A\right) = \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \geq t_B\right) \end{aligned}$$

If the noise is isotropic, each dimension is independent,

$$\mu_x(x) = \prod_{i=1}^d \mu_x(x_i) \quad (26)$$

Thus, conditions for the isotropic noise can be rewritten as:

$$\mathbb{P}\left(\prod_{j=1}^d \frac{\mu_x(x_j - \delta_j)}{\mu_x(x_j)} \leq t_A\right) = \underline{p}_A \quad (27)$$

$$\mathbb{P}\left(\prod_{j=1}^d \frac{\mu_x(x_j - \delta_j)}{\mu_x(x_j)} \geq t_B\right) = \overline{p_B} \quad (28)$$

$$\mathbb{P}\left(\prod_{j=1}^d \frac{\mu_x(x_j)}{\mu_x(x_j + \delta_j)} \leq t_A\right) = \mathbb{P}\left(\prod_{j=1}^d \frac{\mu_x(x_j)}{\mu_x(x_j + \delta_j)} \geq t_B\right) \quad (29)$$

Thus, this completes the proof. $\square$

B.2 Binary Case for Theorem 2

Theorem 3. (*Universal Certified Robustness (Binary Case)*) Let $f : \mathbb{R}^d \rightarrow \mathcal{Y}$ be any deterministic or random function, and let ϵ follows any continuous distribution. Let g be defined as in (1). Suppose the most probable class $c_A \in \mathcal{Y}$ and the lower bound of the probability $\underline{p}_A$ satisfy:

$$\mathbb{P}(f + \epsilon) = c_A \geq \underline{p}_A \geq \frac{1}{2} \quad (30)$$

Then $g(x + \delta) = c_A$ for all $\|\delta\|_p \leq R$, where R is given by the optimization:

$$\begin{aligned} & \underset{\delta}{\text{minimize}} && R = \|\delta\|_p \\ & \text{subject to} && \mathbb{P}\left(\frac{\mu_x(x - \delta)}{\mu_x(x)} \leq t_A\right) = \underline{p}_A, \\ & && \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \delta)} \leq t_A\right) = \frac{1}{2} \end{aligned}$$

B.3 UniCR (Binary Case)

Similar to the binary case of Theorem 1, the binary case of the two-phase optimization can be easily derived:

$$\begin{aligned} & R = \|\lambda \delta\|_p, \text{ where } \delta \in \arg \min_{\delta} \|\lambda \delta\|_p \\ & \text{s.t. } \lambda = \arg \min_{\lambda} |K| \\ & \mathbb{P}\left(\frac{\mu_x(x - \lambda \delta)}{\mu_x(x)} \leq t_A\right) = \underline{p}_A \\ & K = \mathbb{P}\left(\frac{\mu_x(x)}{\mu_x(x + \lambda \delta)} \leq t_A\right) - \frac{1}{2} \\ & \underline{p}_A \geq \frac{1}{2} \end{aligned}$$

B.4 UniCR Bound

The certified radius R approximated by the two-phase optimization is tight if achieving the optimality. Under this assumption, we analysis the confidence bound for the certification. We follow [11] to compute the probabilities $\underline{p}_A$ and $\overline{p}_B$ using Monte Carlo method with sample number n . The confidence is $1 - \alpha_0$, where $p_A \geq \alpha_0^{1/n}$. To estimate the auxiliary parameters t_A and t_B , we use Dvoretzky–Kiefer–Wolfowitz inequality [17] to bound the CDFs of the random variables $\frac{\mu_x(x-\lambda\delta)}{\mu_x(x)}$ and $\frac{\mu_x(x)}{\mu_x(x+\lambda\delta)}$, then determine the t_A and t_B using Algorithm 1.

Lemma 2. (Dvoretzky–Kiefer–Wolfowitz inequality (restate)) *Let $X_1, X_2, \dots, X_n$ be real-valued independent and identically distributed random variables with cumulative distribution function $F(\cdot)$, where $n \in \mathbb{N}$. Let F_n denotes the associated empirical distribution function defined by*

$$F_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbf{1}_{\{X_i \leq x\}}, x \in \mathbb{R} \quad (31)$$

The Dvoretzky–Kiefer–Wolfowitz inequality bounds the probability that the random function F_n differs from F by more than a given constant $\Delta \in \mathbb{R}^+$:

$$\mathbb{P}(\sup_{x \in \mathbb{R}} |F_n(x) - F(x)| > \Delta) \leq 2e^{-2n\Delta^2} \quad (32)$$

We use the Lemma 2 to estimate the CDFs in algorithm 1. In condition 4, we need to estimate 4 probabilities with confidence $1 - 2e^{-2n\Delta^2}$ as well as the p_A and p_B with confidence $1 - \alpha_0$. Therefore, the confidence that deriving the correct radius is at least $(1 - \alpha_0)^2(1 - 2e^{-2n\Delta^2})^4$. In Figure 9, we show the confidence on a varying number of samples when $\Delta = 0.1$ and $\alpha_0 = 0.999$. As the number of samples increases to around 400 (all our experiments use more than 400 samples), our confidence is very close to Cohen’s confidence [11]. Thus, the confidence is nearly 1 in all our experiments.

B.5 Optimization Convergence

We analysis the convergence of the two-phase optimization and the certification accuracy in this section. On one hand, the optimality of the scalar optimization can be asymptotically achieved by binary search. On the other hand, it is hard to find the minimum $\|\lambda\delta\|_p$ in the highly-dimensional space, but some special symmetry in the direction of δ (e.g., spherical symmetry that is also found in [60,59]), can help approximate the certified radius. The detailed algorithms are presented in Section 3.1. The defense performance of such universally approximated certified robustness against different real-world attacks is the same as certified robustness (as shown in Appendix D.2). Thus, such negligible approximation error is close to 0, but result in many significant new benefits in return.

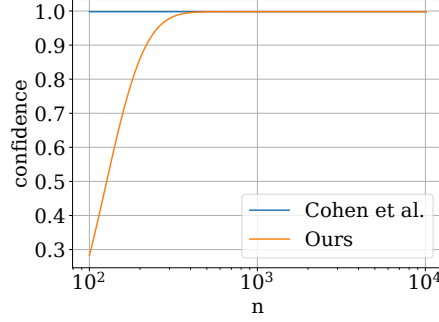


Fig. 9. Confidence vs. number of Monte Carlo samples.

C Algorithms

C.1 Computing t_A and t_B

We present the algorithm to compute the p_A and p_B in Algorithm 1.

Algorithm 1 Computing t_A and t_B

Input: Lower bound of the probabilities, p_A ; upper bound of the probabilities, p_B ; perturbation scalar, λ ; perturbation, δ ; noise PDF, μ_x ; number of samples in the Monte Carlo method, n

Output: The auxiliary parameters, t_A and t_B

- 1: Sample n noise $\epsilon \in \mathbb{R}^{n \times d}$ from a discrete version of PDF.
 - 2: Calculate $\frac{\mu_x(x-\lambda\delta)}{\mu_x(x)}$ using these n samples of noise, μ_x , λ and δ
 - 3: Estimate the CDF Φ of $\frac{\mu_x(x-\lambda\delta)}{\mu_x(x)}$ using Monte Carlo method
 - 4: **return** $t_A = \Phi^{-1}(p_A)$ and $t_B = \Phi^{-1}(p_B)$, with inverse CDF Φ
-

C.2 Scalar Optimization

We use the binary search to find a scale factor that minimizes $|K|$ (*the distance between δ and the robustness boundary*). When $K = 0$, the perturbation δ is exactly on the robustness boundary. Fixing the direction of δ , we find two scalars such that $K > 0$ and $K < 0$. Specifically, we start from a scalar λ_a and compute K . If $K > 0$, then the scaled perturbation $\lambda_a\delta$ is within the robustness boundary, thus we enlarge the scalar to find a λ_b such that $K < 0$ and vice versa. After that, we iteratively compute the K using $\lambda = \frac{1}{2}(\lambda_a + \lambda_b)$: if $K > 0$, we let $\lambda_a = \lambda$; otherwise, we let $\lambda_b = \lambda$. We repeat this iteration until K is less than a threshold or the number of iterations is sufficiently large. The procedures are summarized in Algorithm 2.

Algorithm 2 Scalar Optimization

Input: Lower bound of the probabilities, $\underline{p}_A$; upper bound of the probabilities, $\overline{p}_B$; perturbation scalar, λ ; perturbation, δ ; noise PDF, μ_x ; number of samples in Monte Carlo method, n ; threshold for K , K_m ; number of iterations for binary search, N

Output: The scalar λ that minimizes $|K|$

- 1: Find initial scalar λ_a and λ_b such that $K > 0$ and $K < 0$
- 2: $\lambda = (\lambda_a + \lambda_b)/2$
- 3: Compute K using λ
- 4: **while** $N > 0$ and $|K| > K_m$ **do**
- 5: **if** $K > 0$ **then**
- 6: $\lambda_a = \lambda$
- 7: **else**
- 8: $\lambda_b = \lambda$
- 9: $\lambda = (\lambda_a + \lambda_b)/2$
- 10: Compute K using λ
- 11: $N = N - 1$
- 12: **return** λ

C.3 Direction Optimization

We show how to initialize the positions for different ℓ_p norms in PSO. Since some noise follows PDFs with symmetry [60,59], we set the initial position of particles by considering this, e.g., setting the initial positions w.r.t. ℓ_p for $p \in \mathbb{R}^+$ as $[0, \dots, 0, a, 0, \dots, 0]$ and the initial positions w.r.t. ℓ_∞ as $[a, a, a, \dots, a]$, where a is a small random number. Although the search space is highly-dimensional, empirical results show that the radius given by PSO can accurately approximate the theoretical radius given by other methods, e.g., Cohen’s [11] (see Figure 4). Notice that, for more complicated PDFs without symmetry (which is indeed difficult for deriving the certified radius), PSO can also approximate the certified radius with more particles and iterations.

C.4 Hill-climbing algorithm for I-OPT

The Hill-climbing algorithm is summarized in Algorithm 3.

D More Experimental Results**D.1 Metrics**

We show the illustration of Robustness Score in Figure. 10

D.2 Defense against Real Attacks

We evaluate our UniCR’s defense accuracy against a diverse set of state-of-the-art attacks, including universal attacks [10], white-box attacks [13,55], and black-box attacks [1,6]. We compare UniCR with other state-of-the-art certified

Algorithm 3 I-OPT with Hill Climbing

Input: Input data, x ; PDF of noise distribution, μ_x ; universally approximated certified robustness, $\text{UniCR}(\cdot)$; initial hyper-parameters α ; optimization range of hyper-parameters, $[\mathbf{L}, \mathbf{H}]$; optimization step of hyper-parameters, $\mathbf{S}$

Output: The optimal hyper-parameters, α_{optimal}

```

1: Initialize the certified radius  $R_0 = \text{UniCR}(x, \mu(\alpha))$ 
2: For each hyper-parameter  $\alpha_i$  in  $\alpha$ :
3:   if  $L_i < \alpha_i + S_i < H_i$  then
4:      $R' = \text{UniCR}(x, \mu(\alpha|\alpha_i = \alpha_i + S_i))$ 
5:     if  $R' > R_0$  then
6:        $\alpha$  is updated with  $\alpha_i = \alpha_i + S_i$ 
7:        $R_0 = R'$ 
8:     else if  $L_i < \alpha_i - S_i < H_i$  then
9:        $R' = \text{UniCR}(x, \mu(\alpha|\alpha_i = \alpha_i - S_i))$ 
10:       $R_0 = R'$ 
11:      if  $R' > R$  then
12:         $\alpha$  is updated with  $\alpha_i = \alpha_i - S_i$ 
13:         $R_0 = R'$ 
14:    else
15:      break
16: else
17:   break
18: return  $\alpha_{\text{optimal}} = \alpha$ 

```

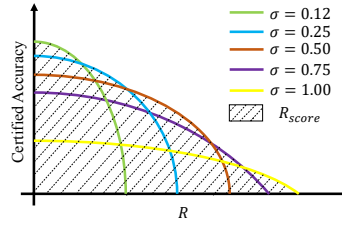


Fig. 10. An example of the Robustness Score.

schemes [59,11,49] against ℓ_1, ℓ_2 and ℓ_∞ perturbations. The certified radius R for each image in the test set (10,000 images in total) are computed beforehand, and the perturbation generation is constrained by $\|\delta\|_p = R$ for all the attack methods. We define the defense accuracy as the rate that the smoothed classifier can successfully defend against the perturbations with the ℓ_p size identical to the the certified radius:

$$acc_d = \mathbb{E}_{\|\delta\|_p=R} \left[\frac{\sum g(x + \delta) = c_A}{N} \right] \quad (33)$$

where $c_A = g(x)$, N is the total test number. In this defense study, we use 500 samples for both Monte Carlo method and testing.

Table 6 shows the defense accuracy on the smoothed classifier. The attack with "*" is re-scaled to the required norm (perturbation size R) based on their perturbation formats. UniCR universally provides a 100% defense accuracy against all the ℓ_1, ℓ_2 and ℓ_∞ perturbations generated by all the state-of-the-art attacks. These results validate our universally approximated certified robustness ensures the same defense performance as certified robustness in practice.

Table 6. Defense against real attacks on CIFAR10 (results on MNIST & ImageNet are similar and not included due to space limit).

Defense Accuracy (%)	Gaussian*	Procedural* [10]	Auto-PGD [13]	Wasserstein* [55]	Square* [1]	HSJ* [6]
Teng's [49] ℓ_1 -norm R	100.00	100.00	100.00	100.00	100.00	100.00
Our ℓ_1 -norm R	100.00	100.00	100.00	100.00	100.00	100.00
Cohen's [11] ℓ_2 -norm R	100.00	100.00	100.00	100.00	100.00	100.00
Our ℓ_2 -norm R	100.00	100.00	100.00	100.00	100.00	100.00
Yang's [59] ℓ_∞ -norm R	100.00	100.00	100.00	100.00	100.00	100.00
Our ℓ_∞ -norm R	100.00	100.00	100.00	100.00	100.00	100.00

D.3 List of PDFs

The PDFs used in our experimental are summarized in Table 7.

Table 7. List of noise distributions.

Distribution	Probability Density Function
Gaussian	$\propto e^{- x/\alpha ^2}$
Laplace	$\propto e^{- x/\alpha }$
Hyperbolic Secant	$\propto \text{sech}(x/\alpha)$
General Normal	$\propto e^{- x/\alpha ^\beta}$
Cauhy	$\propto \frac{\alpha^2}{x^2 + \alpha^2}$
Pareto	$\propto \frac{1}{(1+ x/\alpha)^{\beta+1}}$
Laplace-Gaussian Mix.	$\propto \beta e^{- x/\alpha } + (1-\beta)e^{- x/\alpha ^2}$
Exponential Mix.	$\propto e^{-\beta x/\alpha - (1-\beta) x/\alpha ^2}$

D.4 Efficiency for Radius Derivation

We show the runtime of our algorithms on deriving the certified radius for the inputs with various input dimensions in Figure 11. For the common input dimensions, e.g., 24×24 for MNIST, $3 \times 32 \times 32$ for CIFAR10, and $3 \times 224 \times 224$ for ImageNet, it takes less than 10 seconds for certifying an image on average. Comparing with the theoretical certified radius deriving, our method’s running time is undoubtedly larger since their radius is pre-derived. However, with the significant benefits on the universality and the automatically deriving, we believe the cost of the extra running time is worthwhile and acceptable in practice.

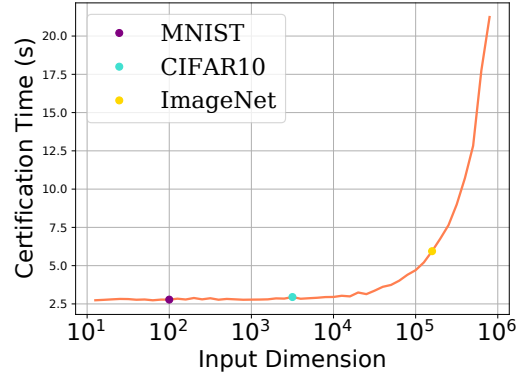


Fig. 11. Runtime of UniCR vs. input sizes (with RTX3080 GPU).

D.5 Any p (besides 1, 2, ∞)

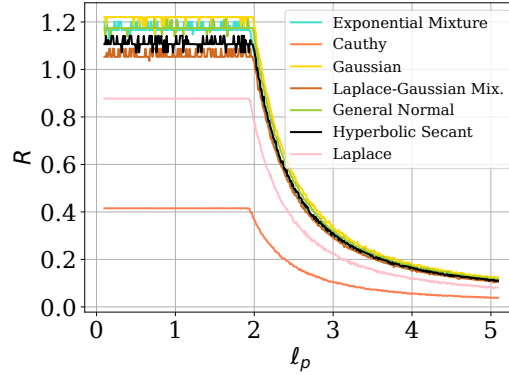


Fig. 12. Radius vs. various ℓ_p pert.

Existing methods [11,49] usually focus on the certified radius in a specific norm, e.g., ℓ_1 , ℓ_2 or ℓ_∞ norms. Some methods [60,59] provide certified robustness theories for multiple norms but specific settings are usually needed for deriving the certified radii in different norms. None of the existing methods can automatically compute the certified radius in any ℓ_p norm. In this section, we show our UniCR can automatically approximate the certified radii for various p , in which p is a real number greater than 0.

In the experiments, we set the probability $p_A = 0.9$ and draw the lines of certified radius w.r.t. different p for $p > 0$. We show the results computed with different noise distributions in Figure 12. We observe that when $p \in (0, 2]$, the certified radius for different p are approximately identical. This finding also matches the theoretical results in Yang et al. [59], in which the certified radii in ℓ_1 and ℓ_2 norm are exactly the same for multiple distributions. When $p > 2$, we observe that the certified radius decreases as p increases.

D.6 Evaluations on Complicated PDFs

We provide a fine-grained evaluation on the complicated distributions [43], e.g., General Normal, Laplace-Gaussian Mixture, and Exponential Mixture noises with various β . It shows that the Gaussian (i.e., $\beta = 2$ for General Normal, $\beta = 0$ for Laplace-Gaussian Mixture and Exponential Mixture) is the optimal noise in these β setting. We also observe the “crash” on Laplace-based distributions when p_A is small.

D.7 Certification on Non-Smoothed Classifier

Table 8. Certified accuracy on standard classifier.

radius R	0.0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8
Yang’s [59] vs. ℓ_1 -norm	10.6	10.4	10.4	9.8	8.8	8.2	5.4	2.2	1.0
Ours vs. ℓ_1 -norm	98.8	47.0	22.4	17.8	13.8	10.2	7.0	3.8	1.0
Cohen’s [11] vs. ℓ_2 -norm	10.6	10.4	10.4	9.6	8.8	8.2	5.6	2.2	1.2
Ours vs. ℓ_2 -norm	98.8	46.0	22.4	17.6	13.8	9.8	7.0	3.8	1.2
Yang’s [59] vs. ℓ_∞ -norm (at $R/255$)	10.6	10.6	10.6	10.4	10.4	10.4	10.4	10.4	10.4
Ours vs. ℓ_∞ -norm (at $R/255$)	98.6	92.4	69.4	61.6	53.6	46.0	37.8	27.4	24.4

Besides certifying inputs with the smoothed classifier, our input noise optimization (I-OPT) can certify input with a standard classifier without degrading the classifier accuracy on clean data (on the contrary, existing works have to trade off such accuracy for certified defenses).

Specifically, since our I-OPT allows the noise for the input certification to be different from the noise used in training, a special case of the training noise is no noise ($\sigma = 0$). This means that we can certify a naturally-trained classifier

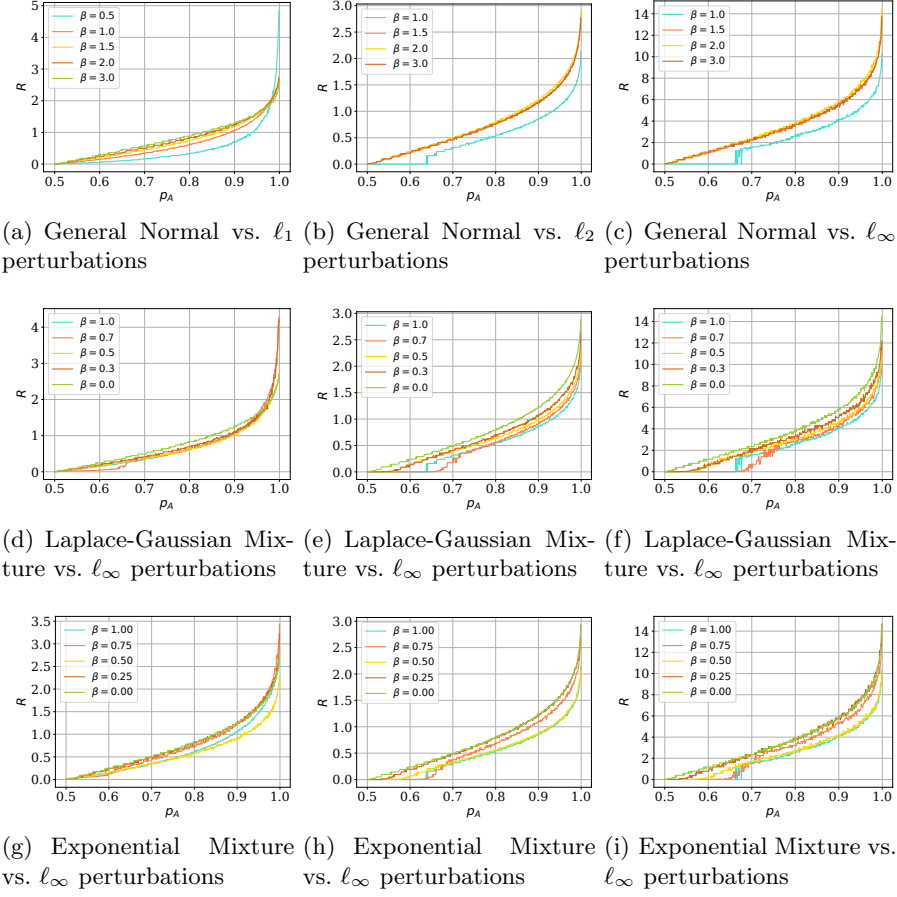


Fig. 13. p_A - R curves of General Normal, Laplace-Gaussian Mixture, and Exponential Mixture noise with a varying β .

(standard classifier). This provides an obvious benefit that the classifier can still execute normal classification on clean data with high accuracy since the standard classifier is trained without noise. Also, with I-OPT, we can tune the noise for the input to maintain the prediction accuracy. Thus, any classifier can be certifiably protected against perturbations without degrading the general performance on clean data.

To maintain the performance on standard classification, we add a condition while performing I-OPT:

$$g(x + \delta) = f(x) \quad (34)$$

We show this application on a standard ResNet110 classifier trained on CIFAR10 (see Table 8). For the baselines, we use Gaussian noise ($\sigma = 0.35$) and

its corresponding theoretical radius [59,11] for certification. Our method uses I-OPT with General Normal noise and initializes it with the same σ . While approximating the certified radius with UniCR, we generate 4,000 samples with the Monte Carlo method on CIFAR10.

The table shows that over 98.6% of the inputs are certified by our method with a radius $R > 0$. This means that over 98.6% of the samples are certifiably protected while only 10.6% of inputs are certified by the baselines, which is nearly the accuracy by random guessing. This significant improvement emerges since the I-OPT could optimize the noise PDF for each input even though the classifier is not trained with noise (non-smoothed classifier). Although the certified radii are low compared to smoothly-trained classifiers, it provides a certifiable protection on perturbed data while maintaining the high accuracy for classifying clean data.

E Visual Examples of I-OPT

We present some examples of I-OPT on the ImageNet dataset against ℓ_1 , ℓ_2 and ℓ_∞ perturbations, respectively (see Figure 14). In the first case (ℓ_1 perturbations), without executing the I-OPT, our UniCR certifies the input with a radius $R = 1.24$. Our I-OPT optimizes the distribution as the right-most figure shows, then the certified radius is improved to 1.48 with our UniCR. Similarly, in the rest cases, we show I-OPT can improve the certified radius significantly by optimizing the noise distribution. Especially, we improve the radius from 0.35 to 1.30 in the second case.

F Discussions

Universal Certified Robustness. It might be impractical to make a universal framework satisfy all the theoretical conditions w.r.t. all ℓ_p perturbations, especially p can be any positive real number. Thus, we admit that UniCR may not strictly satisfy certified robustness all the time due to the approximated optimization. However, extensive empirical results confirm that our derived radii highly approximate the theoretical certified radii against different ℓ_p perturbations. In addition, the defense performance against real attacks also illustrate that our method is as reliable as different theoretical certified radii. We believe that with the negligible error in practice, UniCR can be deployed as a universal framework to significantly ease the process of achieving certified robustness in different scenarios.

Certifying Perturbed Data with Randomized Smoothing. Traditional randomized smoothing usually assumes that the input is clean and empirical defenses [41,26] are not applied, if the input data is perturbed before certification, then certification in I-OPT might be inaccurate. Indeed, the certification in traditional randomized smoothing (e.g., [11]) methods also depend on the inputs (since p_A is different for different inputs), they might be inaccurate if the input data is perturbed, either. Thus, randomized smoothing based approaches focus

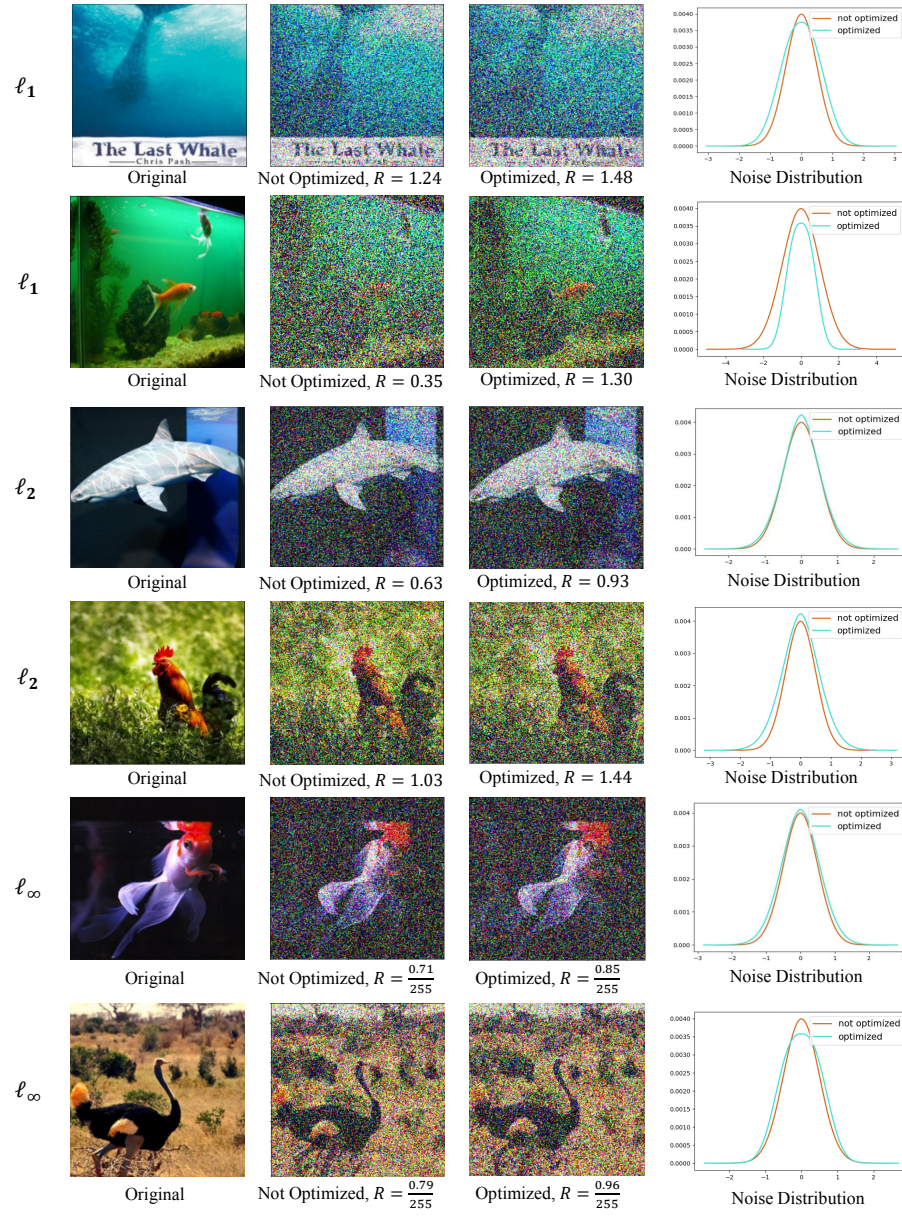


Fig. 14. Example images of applying I-OPT (based on UniCR) for smoothed classifier against different ℓ_p perturbations on the ImageNet dataset. From the left to right, the first figure shows the original image. The second and third figures show the smoothed image without I-OPT and with I-OPT, respectively. The fourth figure shows the corresponding distributions before/after I-OPT.

on certifying clean inputs rather than correcting perturbed inputs. We will study this interesting problem on certifying both clean and perturbed inputs in the future.

Can existing methods adopt noise optimization? A question here is that if the noise optimization can improve the certified radius, can the theoretical methods provide personalized randomization for each input? The personalized randomization is actually not adaptable in the theoretical methods since they cannot automatically derive the certified radius for different noise distributions, especially for uncommon distributions, e.g., $e^{-|x/0.5|^{1.5}}$. Instead, our UniCR can automatically derive the certified radius for any distribution within the continuous parameter space.

Extensions. We evaluate our UniCR on the image classification. Indeed, our UniCR is a general method that can be directly applied to other tasks, e.g., video classification [40,58], graph learning (e.g., node/graph classification [52] and community detection [28]), and natural language processing [30].