

Data-free Defense of Black Box Models Against Adversarial Attacks

Gaurav Kumar Nayak
University of Central Florida

Inder Khatri
New York University

Ruchit Rawal Anirban Chakraborty
Indian Institute of Science

Abstract

Several companies often safeguard their trained deep models (i.e. details of architecture, learnt weights, training details etc.) from third-party users by exposing them only as black boxes through APIs. Moreover, they may not even provide access to the training data due to proprietary reasons or sensitivity concerns. In this work, we propose a novel defense mechanism for black box models against adversarial attacks in a data-free set up. We construct synthetic data via generative model and train surrogate network using model stealing techniques. To minimize adversarial contamination on perturbed samples, we propose ‘wavelet noise remover’ (WNR) that performs discrete wavelet decomposition on input images and carefully select only a few important coefficients determined by our ‘wavelet coefficient selection module’ (WCSM). To recover the high-frequency content of the image after noise removal via WNR, we further train a ‘regenerator’ network with an objective to retrieve the coefficients such that the reconstructed image yields similar to original predictions on the surrogate model. At test time, WNR combined with trained regenerator network is prepended to the black box network, resulting in a high boost in adversarial accuracy. Our method improves the adversarial accuracy on CIFAR-10 by 38.98% and 32.01% on state-of-the-art Auto Attack compared to baseline, even when the attacker uses surrogate architecture (Alexnet-half and Alexnet) similar to the black box architecture (Alexnet) with same model stealing strategy as defender. The code is available at <https://github.com/vcl-iisc/data-free-black-box-defense>

1. Introduction

Deep neural networks, applied in computer vision [43, 53], machine translation [3, 36], speech recognition [18, 32], have exhibited success but face unreliability due to adversarial attacks causing erroneous predictions [2, 25, 49, 52]. These attacks can be categorized into either black box [8, 21, 40, 56] and white box [11, 19, 29, 33] attacks based on access to model parameters. Black box attacks, more practical and realistic than white-box attacks, involve stealing the functionality of target models by training surrogate mod-

els using pairs of (image, predictions). Adversarial samples crafted using surrogate models can also exploit the property of transferability to attack the target model. Hence, immediate attention is required to protect against such attacks.

To make it harder for the adversary to craft black box attacks, companies prefer not to release the training dataset and keep them proprietary. However, recent works have shown that model stealing can compromise the confidentiality of black-box models even without the training data. Existing works perform generative modeling either with proxy data [4, 39, 47] or without proxy data [26, 51, 57] and train the surrogate model with the synthesized data for model stealing. However, their focus is more on obtaining highly accurate surrogate models. In contrast to existing works, we inquire about an important question regarding the safety of the black box models - “how to defend against black box attacks in data-free (absence of training data) set up”.

In order to tackle this problem, our proposed method ‘DBMA’ (i.e., **defending black box models** against adversarial attacks in data-free setup) leverages on the wavelet transforms [15]. We observe difference between wavelet transform on adversarial sample and original sample (shown in Fig. 1 (B)), we notice that detail coefficients in high-frequency regions (LH, HL and HH regions) are majorly corrupted by adversarial attacks and the approximate coefficients (LL region) is least affected for level 1 decomposition. Similar observation holds even for decomposition on other levels. To improve the adversarial accuracy, a naive way would be to completely remove the detail coefficients which can minimize the contamination in adversarial samples but it can lead to a huge drop in clean accuracy as the model predictions are highly correlated with high frequencies [54]. To avoid that, we assign importance to each of the detail coefficients based on magnitude. The least perturbed LL region usually contains higher magnitude coefficients than other regions (Fig. 1 (A)). So, we prefer high magnitude detail coefficients. However, taking a large number of such coefficients can lead to good clean accuracy but at the cost of low adversarial accuracy due to more contamination. On the other hand, taking very few such coefficients can allow lower contamination but results in low clean accuracy. Our method judiciously takes care of this trade-off,

and carefully selects the required important detail coefficients (discussed in Sec. 4.2) using the proposed wavelet coefficient selection module (WCSM).

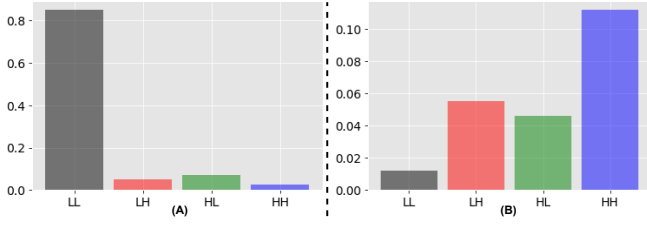


Figure 1. The average absolute magnitude of approximate (LL) and detail coefficients (LH, HL and HH) (via wavelet decomposition) across samples on a) clean data and b) normalized difference between wavelet decomposition of clean and corresponding adversarial image. In (a) the lesser contaminated LL coefficients have higher magnitude. In (b) LL are least affected.

The wavelet noise remover (WNR) removes noise coefficients by filtering out only the top- $k\%$ high magnitude coefficients where optimal k is selected using the WCSM module. As a side-effect, a lot of high-frequency content of the image gets lost which reduces the overall discriminability and ultimately results in suboptimal model’s performance. To cope up with this reduced discriminability, we introduce a U-net-based regenerator network (Sec. 4.3), that takes the spatial samples corresponding to selected coefficients as input and outputs the reconstructed image. The regenerator network is trained by regularizing the feature and input space of the reconstructed image on the surrogate model. In the feature space, we apply cosine similarity and kl-divergence losses to ensure proper reconstruction on clean and adversarial data respectively. Besides this, we also regularize the input space using our spatial consistency loss. Finally, the WNR module combined with the trained regenerator network is appended before the black-box model. The attacker has black-box access to the complete end to end model containing the defense components.

We summarize our contributions as follows:

- In this work, we investigate the largely underexplored yet critical challenge of defending against adversarial attacks on a *black box model without access to the network weights and in the absence of original training samples*.
- We propose a novel strategy to provide adversarial robustness against data-free black box attacks by introducing two key defense components:
 - i.) We propose a wavelet-based noise remover (WNR) containing selective wavelet coefficient module (Sec. 4.2) that aims to remove coefficients corresponding to high frequency components which are most likely to be corrupted by adversarial attack.
 - ii.) We propose a U-net-based regenerator network (Sec. 4.3) that retrieves the coefficients that are lost after the noise removal (via WNR) so that the high-frequency

image content can be restored.

- We demonstrate the efficacy of our method via extensive experiments and ablations on both the components of proposed framework namely wavelet noise remover (Sec. 5.1, 5.2) and the regenerator network (Sec. 5.3) which are appended before the black box target model. The resulting combined model used as the new black box (as seen by attacker), yields high clean and adversarial accuracy on test data (Sec. 5.4).

2. Related Works

Our work is closely related to model stealing and wavelets, so we briefly discuss their related works below.

Data-efficient Model stealing: Based on the availability of training data, we categorize model stealing works as follows:

Training data - On full training data, knowledge distillation [24] is used to extract knowledge using soft labels obtained from the black box model. With few training samples, Papernot *et al.* [40] generates additional synthetic data in the directions (computed using jacobian) where model’s output varies in the neighborhood of training samples.

Proxy data - In the absence of training data, either natural or synthetic images are used as proxy data. Orekondy *et al.* [39] query the black box model on the natural images using adaptive strategy via reinforcement learning to get output predictions and use them to replicate the functionality of the black box model. Barbalau *et al.* [4] use evolutionary framework to learn image generation on a proxy dataset where the generated images are enforced to exhibit high confidence on the black box model. Sanyal *et al.* [47] use the GAN framework with a proxy dataset composed of either related/unrelated data or synthetic data.

Without Proxy data - Kariyappa *et al.* [26] proposed an alternate training mechanism between generator and surrogate model, where generator is trained to produce synthetic samples to maximize the discrepancy between the predictions of the surrogate and the black box model. Truong *et al.* [51] also train generator and surrogate model alternatively but they replace discrepancy loss computed using KL divergence with L1 norm over logits approximated from softmax. Similarly, Zhou *et al.* [57] also formulate a min-max adversarial game but they additionally enforce the synthetic data to be equally distributed among the classes using a conditional generator.

Unlike these existing works, we use model stealing only as a means to obtain the surrogate model and synthetic data. Unlike them where they steal as an adversary, our goal is to provide robustness against black box attacks i.e. to reduce the effects of the adversarial samples on the black box model that are crafted using the surrogate model.

Wavelet in CNNs: Before the CNNs, the wavelets had been used for noise reduction and denoising [16, 17].

Prakash *et al.* [41] used pixel deflection technique followed by adaptive thresholding on the wavelets coefficients for denoising. Unlike these works, our setup is more challenging due to no access to both the training data and the model weights. Mustafa *et al.* [37] utilized wavelet denoising with image superresolution as a defense mechanism against adversarial attacks in grey-box settings. Different from this work, our approach utilizes the wavelet with a proposed regenerator network for defense against adversarial attacks in a black-box setting.

Data and Training Efficient Adversarial Defense: Adversarial Defense techniques can be broadly classified into two categories: Adversarial Training (AT) and Non-Adversarial Training (Non-AT) based methods. AT-based methods [9, 19, 31, 33] rely on adversarial samples during training to improve performance on perturbed samples. However, these methods are computationally expensive and often require high-capacity networks for achieving significant gains in robustness [58]. On the other hand, Non-AT-based methods like JARN[6], BPF[1], and GCE[7] offer faster training but perform inadequately against a wide range of strong attacks [23].

In recent years, researchers have proposed adversarial defense methods that do not require re-training a model for providing robustness, thus making them train efficiently. This has immense practical benefits as unlike most state-of-the-art defenses that necessitate model re-training, such approaches can be seamlessly integrated with already deployed models. One such method, namely Magnet[35], first classifies inputs as clean or adversarial and then transforms the adversarial inputs closer to the clean image manifold. Another approach, Defense-GAN by Samangouei *et al.*[46], uses Generative Adversarial Networks to learn the distribution of clean images and generate samples similar to clean images from inputs corrupted with adversarial noise. Sun *et al.*[48] introduced the Sparse Transformation layer (STL) which maps input images to a low-dimensional quasi-natural image space, suppressing adversarial contamination and making adversarial and clean images indistinguishable. Theagarajan *et al.*[50] presented a defense protocol for black-box facial recognition classifiers consisting of a Bayesian CNN-based adversarial attack detector and image purifiers trained using the data from similar domain to the original training data. They used ensemble of image purifiers for removing the adversarial noise and attack detector for validating the purified image. However, a key limitation of these methods is their reliance on the original training-data for training the defense components, which hinders their use in scenarios where training-data/statistics are not freely available due to proprietary or privacy reasons, etc.

Consequently, recent advancements have redirected focus towards tackling the limitations of training data dependency in defending neural networks against adversarial at-

tacks. For instance, Mustafa *et al.*[37] provide defense on pretrained network without retraining or accessing training data, but have additional dependency on pretrained image super-resolution networks. Moreover, their approach can only be integrated with the target networks that are capable of handling multi-scale inputs, thus making them infeasible on the pretrained networks incapable of handling multi-scale inputs. To avoid these problems, Qiu *et al.*[42] proposed the RDG (Random Distortion over Grids) pre-processing operation, randomly distorting input images by dropping and displacing pixels. Guesmi *et al.*[20] presented SIT (Stochastic Input Transformation), applying random transformations to eliminate adversarial perturbations while maintaining similarity to the original clean images. However, the stochastic nature of SIT without guidance negatively affects the clean accuracy of the target model. In contrast, our proposed Data-free Black-Box defense method, DBMA, better preserves clean-accuracy while also achieving higher robustness.

3. Preliminaries

Notations: The black box model is denoted by B_m which is trained on the proprietary training dataset O_d^{train} . We denote the surrogate model by S_m . The generator G produces synthetic data $S_d = \{x_s^i\}_{i=1}^N$ containing N samples. The logit obtained by the model S_m on input x is $S_m(x)$. The softmax and the label predictions on sample x by model S_m are represented by $soft(S_m(x))$ and $label(S_m(x))$.

The set $A_a = \{A_a^p\}_{p=1}^P$ contains P different adversarial attacks. The adversarial sample corresponding to the n^{th} sample of test dataset O_d^{test} (i.e. x_o^n) is denoted by x_{oa}^n which is crafted with a goal to fool the network B_m . Similarly, S_{da} is the set of crafted adversarial samples corresponding to the synthetic data S_d where the adversarial sample $x_{sa}^n \in S_{da}$ is obtained by perturbing the synthetic sample $x_s^n \in S_d$ using an adversarial attack $A_a^j \in A_a$.

We denote the discrete wavelet transform and its inverse operation by $DWT(\cdot)$ and $IDWT(\cdot)$ respectively. The wavelet coefficient selection module is denoted by WCSM. The regenerator network is represented by R_n .

Model Stealing: Model stealing involves extracting the black box model's (B_m) functionality by inputting images into its API to gather outputs, used subsequently to train a surrogate model (S_m). When no training data O_d^{train} is accessible, this scenario is termed data-free model stealing.

Adversarial Attacks: An adversarial attack $A_a^i \in A_a$ is a human-imperceptible noise ($\|\delta\|_\infty < \epsilon$) crafted to alter model's predictions on the perturbed sample (i.e. adversarial sample) x_{oa} from the original sample x_o . In black box adversarial attacks, the surrogate model S_m creates adversarial samples, transferable to the black box model B_m .

Wavelet Transforms: Wavelets represent the time series signal using linear combinations of an orthogonal basis de-

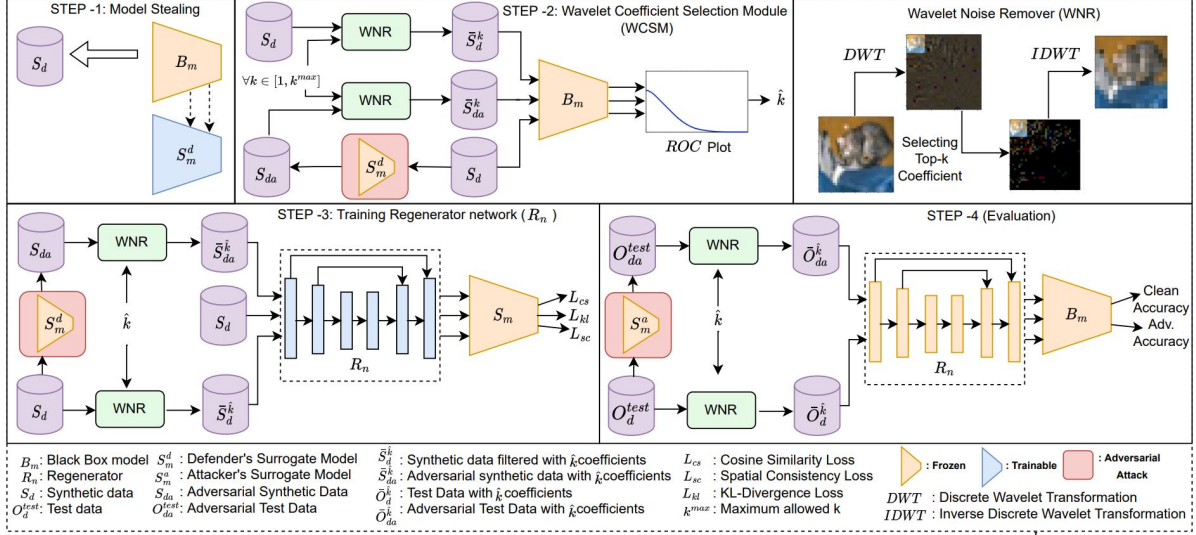


Figure 2. An overview of our proposed approach DBMA. In step 1, we obtain the defender’s surrogate model S_m^d and synthetic data S_d by model stealing from the victim model B_m . In step 2, we use the Wavelet Coefficient Selection Module (WCSM) that gives the optimal % of coefficients ($\hat{k}$) to be selected by the Wavelet Noise Remover (WNR) which are likely to be least corrupted by adversarial attacks. In step 3, we train a regenerator network R_n using different losses (L_{cs} , L_{kl} , L_{sc}) such that the model S_m^d yields features on the regenerated data (clean $R_n(S_d^k)$ and adversarial $R_n(\tilde{S}_{da}^k)$) similar to the features on clean data S_d . Finally in step 4, we evaluate our DBMA approach on test clean (O_d^{test}) and adversarial samples (O_{da}^{test}) where the WNR (with $k = \hat{k}$) and trained R_n are prepended to B_m .

pending on which there are different types of wavelets such as Haar, Cohen and Daubechies [13]. The 2D DWT on an i^{th} image x^i for level 1 yields low pass subband (i.e approximation coefficients - denoting by LL_1^i) and high pass subband (i.e. detail-level coefficients - denoting by LH_1^i , HL_1^i , HH_1^i). In multi-level DWT, the approximation subband is further decomposed (for e.g. on a 2-level decomposition, LL_1^i is also decomposed to LL_2^i , LH_2^i , HL_2^i , HH_2^i).

4. Proposed Approach

In this section, we first discuss the model stealing method (Sec. 4.1) that we use to train the surrogate model S_m (as proxy for black box model B_m) and generate synthetic data S_d (as proxy for original training data O_d^{train}). Next, we propose our method to remove the detail coefficients (Sec. 4.2) that can be most corrupted by an adversarial attack and select the important coefficients to preserve the signal strength in terms of retaining feature discriminability. We dub this approach as wavelet noise remover (WNR) and the coefficients are selected using the wavelet coefficient selection module (WCSM). To boost the performance, we propose a U-Net-based regenerator network (Sec. 4.3) that takes the output of WNR as input and is trained to output a regenerated image on which the surrogate model would yield similar features as the features of the clean sample. The different steps involved in our proposed method (DBMA) for providing data-free adversarial defense in the black box settings are shown in Fig. 2.

4.1. Obtain Proxy Model and Synthetic Data

Given a black box model B_m , our first step is to obtain a proxy model S_m which can allow gradient backpropagation. To steal the functionality of B_m , S_m can be trained using a model stealing technique. But we also do not have access to the original training samples O_d^{train} . Hence, we use a data-free model stealing technique [4] that trains a generator using proxy data to produce synthetic samples (S_d) on which the black box model B_m gives high-confident predictions. The surrogate model S_m is then trained on synthetic data S_d under the guidance of B_m , where the model S_m is enforced to mimic the predictions of model B_m . The trained S_m and the generated data S_d are used in next steps.

4.2. Noise Removal with Wavelet Coefficient Selection Module (WCSM)

For an i^{th} sample of the composed synthetic data S_d (i.e. x_s^i), its corresponding wavelet coefficients are obtained by DWT operation on it. The approximate coefficients are the low frequencies that are least affected by the adversarial attack (shown in Fig. 1 (B)). Thus, we retain these coefficients. For e.g. on level-2 discrete wavelet decomposition, LL_2^i (approximate coefficients for i^{th} sample) is kept. As the adversarial attack severely harms the detail coefficients, we determine the coefficients that can be most affected by it using WCSM for effective noise removal.

Based on our observation that the least affected approximate coefficients often have high magnitude coefficients, indicating that the high magnitude detail coefficients can be

a good measure for choosing which coefficients to retain. Thus, we arrange the detail coefficients based on magnitude (from high to low order) and retain the top- k % coefficients. The efficacy of choosing top- k compared to different baselines (such as random- k and bottom- k) is empirically verified in Sec. 1. However, determining the suitable value of k is a challenge and, if not properly chosen, can lead to a major bottleneck in clean/adversarial performance. To handle it, we propose a wavelet coefficient selection mechanism that carefully selects the value of k so that decent performance can be obtained on both clean and adversarial data.

We empirically estimate optimal k using all the crafted synthetic training samples S_d and their corresponding adversarial counterparts. We define a quantity label consistency rate (LCR^k) which is calculated for a particular value of k using the following steps:

1. Construct adversarial synthetic samples ($\{x_{sa}^i\}_{i=1}^N$) by using an adversarial attack $A_a^j \in A_a$ on the surrogate model S_m .
2. Obtain approximate and detail coefficients for each adversarial synthetic sample using $DWT(x_{sa}^i, l), \forall i \in (1, \dots, N)$ where l is the decomposition level.
3. Craft spatial samples ($\tilde{S}_{da}^k = \{\tilde{x}_{sa}^i\}_{i=1}^N$) using $IDWT$ operation on complete approximate and selected top- k detail coefficients corresponding to each adversarial synthetic sample (i.e. $S_{da} = x_{sa}^i, \forall i \in (1, \dots, N)$). For simplicity, we envelop the operations (2) and (3) and name them ‘wavelet noise remover’ (WNR). In general, for a given input image and k value, WNR applies DWT on input where the approximate coefficients and the chosen top- k % detail coefficients are retained, whereas the non-selected detail coefficients are made to zeros. These coefficients are then passed to $IDWT$ to obtain the filtered spatial image.
4. Perform WNR by repeating the steps 2 and 3 on clean samples x_s to obtain $\tilde{S}_d^k$.
5. Compare the predictions of black box model B_m on samples of $\tilde{S}_d^k$ and the corresponding samples in S_d . LCR_C^k denotes the fraction of clean samples whose predictions match when top- k % coefficients are selected.
6. Compare predictions of black box model B_m on samples of $\tilde{S}_{da}^k$ and the corresponding samples in S_d . LCR_A^k denotes fraction of adversarial samples whose predictions match when top- k % coefficients are selected.
7. Compute $LCR^k = LCR_C^k + LCR_A^k$
8. Calculate the rate of change of LCR^k (i.e. ROC^k) as $ROC^k = LCR^{k+1} - LCR^k$

Using above steps, we calculate LCR for different values of $k \in [1, \dots, k^{max}]$. As shown in Fig. 2, we observe that as we increase the value of k , initially LCR_A increases and reaches its maximum value, and then it starts decreasing. High LCR_A implies that the predictions of the B_m model on $\tilde{S}_{da}^k$ and S_d have a low mismatch. Similarly, with

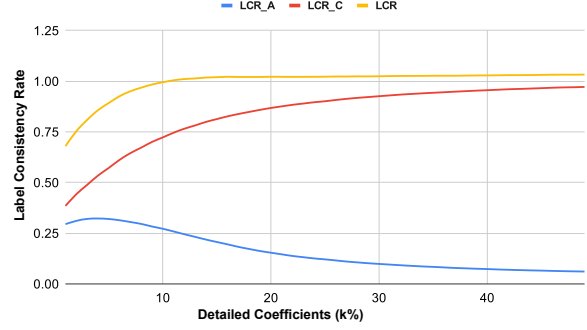


Figure 3. Label consistency rates (LCR_A , LCR_C and LCR) vs detail coefficients (k %) plotted using prediction from black-box model B_m on Cifar-10 Dataset.

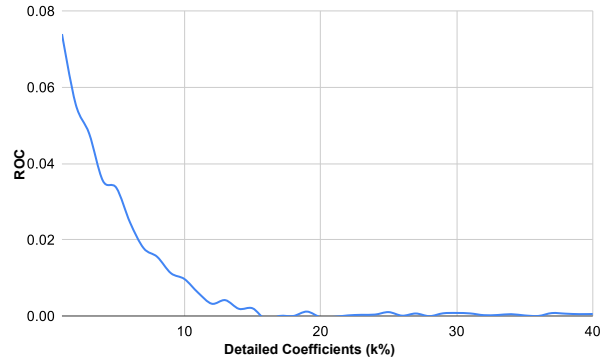


Figure 4. Rate of change (ROC) of LCR vs detail coefficients (k %) plotted using prediction from black-box model B_m on Cifar-10 data. As we increase value of k , ROC becomes negligible. At $k = 16$ it is close to zero.

the increase in value of k , LCR_C keeps increasing which implies as we add more coefficients, model discriminability increases. The value of LCR increases with the value of k , but the rate of increase of LCR keeps decreasing. Refer Fig. 4 where we plot the rate of change (ROC). We choose $\hat{k}$ at which ROC saturates. Here ROC is negligible at $k = 16$. This value of k gives the best trade-off between clean and adversarial accuracy (empirically validated in Sec. 5.1). Thus, wavelet noise remover (WNR) is applied on input at estimated optimal k ($\hat{k}$).

4.3. Training of Regenerator Network

After the noise removal using the WNR at optimal k (i.e. $\hat{k}$ obtained by WCSM), there is also loss in the image signal information as a side effect. Thus, we further regenerate the coefficients using a regenerator network (R_n) that takes input as the output of WNR at $\hat{k}$ and yields reconstructed image which is finally passed to the black box model B_m for test predictions. The architecture of our R_n network is a U-net based architecture with skip connections which is inspired from [45]. For more details on the architecture of regenerator network, refer Supplementary (Sec. 5).

$$\begin{aligned}\hat{k} &= WCSM(S_d, S_m) \\ \bar{x}_s^i &= WNR(x_s^i, \hat{k}); \quad \bar{x}_{sa}^i = WNR(x_{sa}^i, \hat{k})\end{aligned}\quad (1)$$

For training, we feed the output obtained from the WNR at $\hat{k}$ as input to the network R_n and obtain a reconstructed image which is passed to the frozen surrogate model S_m for loss calculations. The losses used to train R_n are as follows: **Cosine similarity loss** (L_{cs}): To enforce similar predictions from S_m on the regenerated synthetic sample and the corresponding original synthetic sample. $L_{cs} = CS(S_m(R_n(\bar{x}_s^i)), S_m(x_s^i))$.

KL divergence loss (L_{kl}): To align the predictions of S_m on the regenerated sample and its adversarial counterpart $L_{kl} = KL(\text{soft}(S_m(R_n(\bar{x}_s^i))), \text{soft}(S_m(R_n(\bar{x}_{sa}^i))))$.

Spatial consistency loss (L_{sc}): To make sure that the spatial reconstructed image (clean and adversarial) and the corresponding original synthetic image are similar in the image manifold $L_{sc} = \|R_n(\bar{x}_s^i) - x_s^i\|_1 + \|R_n(\bar{x}_{sa}^i) - x_{sa}^i\|_1$.

Here, CS and KL denotes cosine similarity and KL divergence respectively. Overall loss used in training R_n :

$$L(R_n^\theta) = -\lambda_1 L_{cs} + \lambda_2 L_{kl} + \lambda_3 L_{sc} \quad (2)$$

Finally, the black box model B_m is modified by prepending the WNR (with $k = \hat{k}$) and the trained R_n network to it. The resulting black box model defends the adversarial attacks which we discuss in detail in next section.

5. Experiments

In this section, we validate the effectiveness of our proposed method (DBMA) and perform ablations to show the importance of individual components. We use the benchmark classification datasets i.e. CIFAR-10 [27] and SVHN [38], on which we evaluate the clean and the adversarial accuracy against three different adversarial attacks (i.e., BIM [30], PGD[34] and Auto Attack [12]). Unless it is mentioned, we use Alexnet [28] as black box B_m (results on a larger black-box model are in Sec. 8 in supplementary) and Resnet-18 [22] as the defender’s surrogate model S_m^d , which the defender uses to train the regenerator network R_n as explained in Sec. 4. In the black-box setting, attackers also do not have access to the black-box model’s weights, thus restricting the generation of adversarial samples. So similar to the defender, we leverage the model stealing techniques to get a new surrogate model S_m^a , which the attacker uses for generating the adversarial samples. While evaluating against different attacks, we assume the attacker has access to defense components, i.e., the attacker uses model stealing methods to steal the functionality of the defense components along with the victim model.

We perform experiments with two different architectures for S_m^a : Alexnet-half and Alexnet, which are similar to the black-box model (Alexnet), making it tough for the defender. Ablation for different combination of S_m^d and S_m^a

are in Sec. 7 in supplementary. The attacker uses the same model stealing technique [4] as used by defender. It is important to note that our rigorous approach; we grant the attacker access to our exact model-stealing technique, architecture, and related details to ensure that any performance improvement is not attributed to differences in techniques or architecture between the attacker and defender. We use the Daubechies wavelet for both DWT and $IDWT$ operations. Refer to supplementary (Sec. 2) for experimental results on other wavelets. The decomposition level is fixed at 2 for all the experiments and ablations. The value of k^{max} is taken as 50. We assign equal weights to all the losses with weight as 1 (i.e. $\lambda_1 = \lambda_2 = \lambda_3 = 1$) in eq. 2.

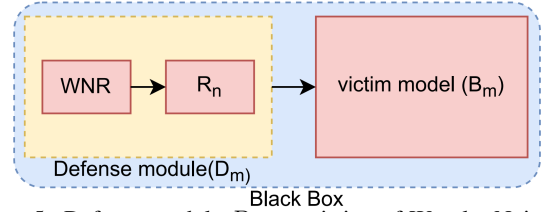


Figure 5. Defense module D_m consisting of Wavelet Noise Remover (WNR) and Regenerator R_n is prepended before the victim model B_m in our approach (DBMA). The D_m and B_m are combinedly considered as the black-box model by the attacker.

The defender constructs a defense module (D_m) using S_m^d . In subsections 5.1 and 5.2, the defense module only consists of the WNR, whereas subsections 5.3 onwards, both R_n and WNR are part of the defense module as shown in Fig. 5. We prepend the defense module before the B_m to create a new black box model that is used to defend against the adversarial attacks. To show the efficacy of defense components used in our method (DBMA), we consider the most challenging scenario, where the attacker uses the same model stealing technique as defender, and considers the defense module also a part of the black-box model while generating adversarial samples.

5.1. Ablation on quantity of coefficients

Table 1. Investigating the efficacy of our proposed WCSM in determining the quantity of detail coefficients to retain. The case (No- k) yields poor results justifying the need to preserve detail coefficients. Unlike, low and high values of k , we obtain better trade-off between clean and adversarial performance on our- k .

Amount of detail coefficients (k)	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
	clean	BIM	PGD	Auto Attack
No k	31.19	5.88	4.68	9.35
low- k ($k=1$)	42.75	10.2	8.72	15.8
low- k ($k=2$)	50.17	15.37	14.14	21.92
low- k ($k=4$)	59.14	17.54	16.03	25.08
high- k ($k=50$)	82.58	5.58	3.33	10.44
our- k ($k=16$)	77.92	15.98	14.04	21.34

In Sec. 1 and 4, we discussed the importance of selecting

the optimal number of detail coefficients ($\hat{k}$) and proposed the steps to find the value of $\hat{k}$ using the WCSM module. In this subsection, we do an ablation over the different choices for values of k (i.e., the number of detail coefficients to select) and analyze its impact on clean and adversarial accuracy. Specifically, we consider six distinct values of k across a wide range i.e. 0 (no detail coefficients, only approximate coefficients), 1, 2 and 4 (small k), 50 (large k), $\hat{k}$ (optimal k given by WCSM). Fig. 4 shows the graph of the rate of change of LCR for different values of k . We select the value of k at which the ROC starts saturating, i.e. k with value 16 as $\hat{k}$. The results corresponding to the different values of k are shown in Table 1. We observe poor performance for both adversarial and clean samples when no detail coefficients are taken. On increasing the fraction of detail coefficients ($k = 1, 2, 4$), an increasing trend for both clean and adversarial performance is observed. Further, for a high value of k clean accuracy improves, but with a significant drop in the adversarial performance. Our choice of k (i.e., $\hat{k}$) indeed leads to better clean accuracy with decent adversarial accuracy, hence justifying the importance of the proposed noise removal using WCSM component.

5.2. Effect of wavelet noise remover with WCSM

In this section, we study the effect of prepending the WNR module to the black-box model B_m . WNR selects the approximate coefficients and optimal $\hat{k}\%$ detail coefficients (obtained by WCSM). It filters out the remaining detail coefficients, which helps in reducing the adversarial noise from the samples. We evaluate the performance of the B_m with and without the WNR Module and present the results in Table 2. When the attacker’s surrogate model is Alexnet-half, adversarial accuracy improves by $\approx 19 - 22\%$ across attacks using the WNR module. Similarly, when the attacker’s surrogate model is Alexnet, the adversarial accuracy improves by $\approx 11 - 12\%$.

Table 2. Our wavelet noise remover using WCSM yields improvement in adversarial accuracy with small drop in clean accuracy.

Surrogate model (attacker)	Noise Removal using WCSM	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet-half	No	82.58	7.02	4.53	11.65
	(Ours) Yes	77.92	26.66	24.55	34.02
Alexnet	No	82.58	4.17	2.19	8.55
	(Ours) Yes	77.92	15.98	14.04	21.34

5.3. Ablation on losses

Until now, we performed experiments by using only the WNR defense module. Now, we additionally attach another defense module (R_n) to WNR (refer Fig 5). In this subsection, we perform ablation to demonstrate the importance of different losses used for training the Regenerator network R_n . As shown in eq. 2, the total loss L is the weighted sum

of three different losses (i.e., L_{cs} , L_{kl} , and L_{sc}). To determine the effect of each of the individual losses, we train R_n using only the L_{cs} loss, L_{kl} loss and L_{sc} loss. Further to analyse the cumulative effect, we train R_n with different possible pairs of loss i.e., $L_{cs} + L_{sc}$ loss, $L_{cs} + L_{kl}$ loss, $L_{kl} + L_{sc}$ loss, and finally with the total loss ($L_{cs} + L_{kl} + L_{sc}$) respectively. The results are displayed in Table 3.

Table 3. Contribution of different losses used for training R_n . The loss ($L_{cs} + L_{kl} + L_{sc}$) gives the best improvement in adversarial accuracy with decent clean accuracy.

Surrogate model (attacker)	Losses to train Regenerator network (R_n)	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet-half	L_{cs}	78.96	26.33	24.75	33.81
	L_{sc}	78.85	27.38	25.75	35.51
	L_{kl}	9.82	6.56	6.54	8.98
	$L_{cs} + L_{sc}$	79.75	27.70	25.34	34.64
	$L_{cs} + L_{kl}$	62.06	36.03	35.93	43.05
	$L_{kl} + L_{sc}$	65.94	37.72	37.62	46.14
	$L_{cs} + L_{kl} + L_{sc}$	73.77	42.71	42.71	50.63
Alexnet	L_{cs}	78.96	16.34	14.57	21.81
	L_{sc}	78.85	17.68	15.97	23.77
	L_{kl}	9.82	6.61	6.38	8.98
	$L_{cs} + L_{sc}$	79.75	17.28	15.6	23.59
	$L_{cs} + L_{kl}$	62.06	24.86	25.91	32.26
	$L_{kl} + L_{sc}$	65.94	31.04	31.05	38.26
	$L_{cs} + L_{kl} + L_{sc}$	73.77	33.31	31.72	40.56

Compared to the earlier best performance with WNR defense module (Table 2), we observe R_n trained with only L_{cs} loss gives no significant improvement in both the adversarial and clean accuracy. Similar trend is observed for R_n trained with L_{sc} loss. However, using only the L_{kl} loss shows a deteriorated clean and adversarial performance of R_n . Further, using the combination of both L_{cs} and L_{sc} loss also does not show much improvement. Combining the L_{kl} with L_{cs} and L_{sc} loss improves the adversarial performance of R_n appreciably, but with a drop in the clean accuracy. L_{kl} with L_{cs} loss shows a consistent improvement of $\approx 9 - 11\%$ in adversarial accuracy across attacks using both the Alexnet and Alexnet-half. Similarly, the combination of L_{kl} with L_{sc} loss improves the adversarial accuracy of R_n by $\approx 11 - 13\%$ and $\approx 15 - 17\%$ against the attacks using Alexnet-half and Alexnet respectively. When R_n is trained with all three losses gives the best adversarial accuracy across all the possible combinations. We observe an overall improvement of $\approx 16 - 19\%$ across different attacks using both Alexnet and Alexnet-half with a slight drop in clean accuracy ($\approx 4\%$). Regenerator network regenerates the lost coefficients, but as explained in section 4.2, detail coefficients also cause a decrease in adversarial accuracy. When we train a regenerator network using the combination of L_{cs} , L_{sc} , and L_{kl} loss, the L_{cs} and L_{sc} loss help to increase clean accuracy, but at the same time L_{kl} loss ensures regenerated coefficients do not decrease the adver-

Table 4. Utility of each component used in our method (DBMA)- Wavelet Noise Remover (WNR) and Regenerator Network R_n on SVHN and CIFAR dataset. WNR with R_n yields huge gains in adversarial performance compared to baseline and WNR alone.

Surrogate Model (attacker)	Dataset	Method	Black Box Model : Alexnet Surrogate Model (defender) : Resnet-18			
			clean	BIM	PGD	Auto Attack
Alexnet-half	SVHN	Baseline	94.49	44.26	44.21	46.79
		SIT[20]	68.72	46.76 ($\uparrow 2.5$)	46.21 ($\uparrow 2$)	50.57 ($\uparrow 3.78$)
		RDG[42]	92.68	55.01 ($\uparrow 10.75$)	54.62 ($\uparrow 10.41$)	58.35 ($\uparrow 11.56$)
		WNR (Ours)	94.21	55.42 ($\uparrow 11.16$)	55.70 ($\uparrow 11.49$)	58.47 ($\uparrow 11.68$)
		WNR+ R_n (Ours)	90.91	68.63 ($\uparrow 24.37$)	68.60 ($\uparrow 24.39$)	71.71 ($\uparrow 24.92$)
	CIFAR	Baseline	82.58	7.02	4.53	11.65
		SIT[20]	51.16	22.05 ($\uparrow 15.03$)	21.68 ($\uparrow 17.15$)	29.08 ($\uparrow 17.43$)
		RDG[42]	67.58	19.99 ($\uparrow 12.97$)	18.97 ($\uparrow 14.44$)	29.64 ($\uparrow 17.99$)
		WNR (Ours)	77.92	26.66 ($\uparrow 19.64$)	24.55 ($\uparrow 20.02$)	34.02 ($\uparrow 22.37$)
		WNR+ R_n (Ours)	73.77	42.71 ($\uparrow 35.69$)	42.71 ($\uparrow 38.18$)	50.63 ($\uparrow 38.98$)
Alexnet	SVHN	Baseline	94.49	38.14	38.19	40.16
		SIT[20]	68.72	43.48 ($\uparrow 5.34$)	43.77 ($\uparrow 5.58$)	47.63 ($\uparrow 7.47$)
		RDG[42]	92.68	50.28 ($\uparrow 12.14$)	50.29 ($\uparrow 12.1$)	53.79 ($\uparrow 13.63$)
		WNR (Ours)	94.21	48.98 ($\uparrow 10.84$)	49.02 ($\uparrow 10.83$)	51.49 ($\uparrow 11.33$)
		WNR+ R_n (Ours)	90.91	63.13 ($\uparrow 24.99$)	63.12 ($\uparrow 24.93$)	66.18 ($\uparrow 26.02$)
	CIFAR	Baseline	82.58	4.17	2.19	8.55
		SIT[20]	51.16	18.64 ($\uparrow 14.47$)	18.43 ($\uparrow 16.24$)	26.23 ($\uparrow 17.68$)
		RDG[42]	67.58	15.54 ($\uparrow 11.37$)	14.50 ($\uparrow 12.31$)	24.01 ($\uparrow 15.46$)
		WNR (Ours)	77.92	15.98 ($\uparrow 11.81$)	14.04 ($\uparrow 11.85$)	21.34 ($\uparrow 12.79$)
		WNR+ R_n (Ours)	73.77	33.31 ($\uparrow 29.14$)	31.72 ($\uparrow 29.53$)	40.56 ($\uparrow 32.01$)

serial accuracy. To achieve best tradeoff between clean and adversarial accuracy, R_n gets trained to increase adversarial accuracy at the cost of decreased clean accuracy compared to R_n network trained with only L_{cs} loss.

5.4. Comparison with existing Data and Training efficient defense methods

In this subsection, we validate the efficacy of our proposed method DBMA by comparing it with two other state-of-the-art defense methods: SIT[20] and GD[42]. For comparison, we do experiments on two benchmark datasets, i.e., SVHN and CIFAR-10. We obtain the optimal $\hat{k}$ as 20 using the WCSM module for the black box model trained on the SVHN dataset. In Table 4, while defending with SIT, the adversarial accuracy improves by $\approx 2 - 3\%$ and $\approx 5 - 7\%$ across attacks crafted using different S_m^a (Alexnet and Alexnet-half) with a corresponding drop of $\approx 2 - 3\%$ in clean accuracy. On the other hand, defending with GD, improves adversarial accuracy by $\approx 2 - 3\%$ across attacks with a marginal drop of 1.66% in clean accuracy. While defending with only WNR in the defender module, the adversarial accuracy improves by $\approx 10 - 11\%$ across attacks crafted using different surrogate architectures S_m^a (Alexnet and Alexnet-half). The clean accuracy, however, experiences a minor drop of less than 1%. By utilizing both the WNR and R_n in the defender module, the adversarial performance further improves by $\approx 13 - 14\%$ across the attacks with a drop of $\approx 4\%$ in clean accuracy. Overall, we observe a gain of $\approx 24 - 26\%$ in adversarial accuracy compared to the baseline model, at the cost $\approx 5\%$ drop in clean accuracy. Similarly, for the CIFAR-10 dataset, we observe an overall improvement of $\approx 35 - 38\%$ and $\approx 29 - 32\%$ against at-

tacks crafted using Alexnet-half and Alexnet, respectively. However, defending with SIT and RDG only improves the adversarial accuracy by $\approx 12 - 17\%$. Additionally, the clean accuracy drops by $\approx 31\%$ and $\approx 15\%$ when using SIT and RDG, respectively. In comparison, when using DBMA, we observed a relatively small drop of $\approx 8\%$ in clean accuracy, which is reasonable considering the challenging nature of our problem setup. Even in traditional adversarial training with access to full data, clean performance often drops at the cost of improving adversarial accuracy [34]. In our case, neither the training data nor the model weights are provided. Moreover, the black-box model is often obtained as APIs, and re-training the model from scratch becomes unfeasible. Considering these difficulties, the drop we observe on clean data is small with respectable overall performance.

6. Conclusion

We introduced DBMA, a novel defense strategy to defend black-box models from adversarial attacks without relying on training data. DBMA incorporates two defense components: a) Wavelet Noise Remover (WNR) that removes the most contaminated areas by adversarial attacks while preserving less affected regions b) a Regenerator network to restore lost information post WNR noise removal. Through various ablations and experiments, we demonstrated efficacy of each defense component. Our method DBMA significantly enhances robustness against data-free black box attacks across datasets and against diverse model-stealing methods. Similar to adversarial defenses in white-box setups, we observe that significant gains in robust accuracy come at the cost of a slight drop in clean accuracy. In future work, we plan to work on further mitigating this trade-off.

References

- [1] Sravanti Addepalli, Vivek B.S., Arya Baburaj, Gaurang Sriramanan, and R. Venkatesh Babu. Towards achieving adversarial robustness by enforcing feature consistency across bit planes. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020. 3
- [2] Naveed Akhtar and Ajmal Mian. Threat of adversarial attacks on deep learning in computer vision: A survey. *Ieee Access*, 6:14410–14430, 2018. 1
- [3] Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. Neural machine translation by jointly learning to align and translate. In *International Conference on Learning Representations, ICLR 2015, Conference Track Proceedings*, 2015. 1
- [4] Antonio Barbalau, Adrian Cosma, Radu Tudor Ionescu, and Marius Popescu. Black-box ripper: Copying black-box models using generative evolutionary algorithms. *Advances in Neural Information Processing Systems*, 33:20120–20129, 2020. 1, 2, 4, 6, 12
- [5] Gregory Beylkin, Ronald R. Coifman, and Vladimir Rokhlin. Fast wavelet transforms and numerical algorithms i. *Communications on Pure and Applied Mathematics*, 44:141–183, 1991. 12
- [6] Alvin Chan, Yi Tay, Yew Soon Ong, and Jie Fu. Jacobian adversarially regularized networks for robustness. In *International Conference on Learning Representations*, 2020. 3
- [7] Hao-Yun Chen, Jhao-Hong Liang, Shih-Chieh Chang, Jia-Yu Pan, Yu-Ting Chen, Wei Wei, and Da-Cheng Juan. Improving adversarial robustness via guided complement entropy. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4881–4889, 2019. 3
- [8] Pin-Yu Chen, Huan Zhang, Yash Sharma, Jinfeng Yi, and Cho-Jui Hsieh. Zoo: Zeroth order optimization based black-box attacks to deep neural networks without training substitute models. *Proceedings of the 10th ACM Workshop on Artificial Intelligence and Security*, 2017. 1
- [9] Sihong Chen, Haojing Shen, Ran Wang, and Xizhao Wang. Towards improving fast adversarial training in multi-exit network. *Neural Networks*, 150:1–11, 2022. 3
- [10] Albert Cohen, Ingrid Daubechies, and J. C. Feauveau. Biorthogonal bases of compactly supported wavelets. *Communications on Pure and Applied Mathematics*, 45:485–560, 1992. 12
- [11] Francesco Croce and Matthias Hein. Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks. In *International conference on machine learning*, pages 2206–2216. PMLR, 2020. 1
- [12] Francesco Croce and Matthias Hein. Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks. In *International conference on machine learning*, pages 2206–2216. PMLR, 2020. 6
- [13] Ingrid Daubechies. Orthonormal bases of compactly supported wavelets. *Communications on Pure and Applied Mathematics*, 41(7):909–996, 1988. 4
- [14] Ingrid Daubechies. Ten lectures on wavelets. *Computers in Physics*, 6:697–697, 1992. 12
- [15] Ingrid Daubechies. *Ten Lectures on Wavelets*. Society for Industrial and Applied Mathematics, 1992. 1
- [16] David L Donoho. De-noising by soft-thresholding. *IEEE transactions on information theory*, 41(3):613–627, 1995. 2
- [17] David L Donoho and Jain M Johnstone. Ideal spatial adaptation by wavelet shrinkage. *biometrika*, 81(3):425–455, 1994. 2
- [18] Haytham M Fayek, Margaret Lech, and Lawrence Cavedon. Evaluating deep learning architectures for speech emotion recognition. *Neural Networks*, 92:60–68, 2017. 1
- [19] Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. In *International Conference on Learning Representations, ICLR 2015, Conference Track Proceedings*, 2015. 1, 3
- [20] Amira Guesmi, Ihsen Alouani, Mouna Baklouti, Tarek Frikha, and Mohamed Abid. Sit: Stochastic input transformation to defend against adversarial attacks on deep neural networks. *IEEE Design & Test*, 39(3):63–72, 2021. 3, 8
- [21] Linguang Hao, Kuangrong Hao, Bing Wei, and Xuesong Tang. Boosting the transferability of adversarial examples via stochastic serial attack. *Neural Networks*, 150:58–67, 2022. 1
- [22] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016. 6
- [23] Warren He, James Wei, Xinyun Chen, Nicholas Carlini, and Dawn Song. Adversarial example defenses: ensembles of weak defenses are not strong. In *Proceedings of the 11th USENIX Conference on Offensive Technologies*, pages 15–15, 2017. 3
- [24] Geoffrey Hinton, Oriol Vinyals, and Jeffrey Dean. Distilling the knowledge in a neural network. In *NIPS Deep Learning and Representation Learning Workshop*, 2015. 2

- [25] Lifeng Huang, Chengying Gao, and Ning Liu. Defeat: Decoupled feature attack across deep neural networks. *Neural Networks*, 156:13–28, 2022. 1
- [26] Sanjay Kariyappa, Atul Prakash, and Moinuddin K Qureshi. Maze: Data-free model stealing attack using zeroth-order gradient estimation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13814–13823, 2021. 1, 2
- [27] Alex Krizhevsky and Geoffrey Hinton. Learning multiple layers of features from tiny images. Technical report, Canadian Institute for Advanced Research, 2009. 6
- [28] Alex Krizhevsky, Ilya Sutskever, and Geoffrey E Hinton. Imagenet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems*. Curran Associates, Inc., 2012. 6
- [29] Alexey Kurakin, Ian J Goodfellow, and Samy Bengio. Adversarial examples in the physical world. In *Artificial intelligence safety and security*, pages 99–112. Chapman and Hall/CRC, 2018. 1
- [30] Alexey Kurakin, Ian J Goodfellow, and Samy Bengio. Adversarial examples in the physical world. In *Artificial intelligence safety and security*, pages 99–112. Chapman and Hall/CRC, 2018. 6
- [31] Alex Lamb, Vikas Verma, Kenji Kawaguchi, Alexander Matyasko, Savya Khosla, Juho Kannala, and Yoshua Bengio. Interpolated adversarial training: Achieving robust neural networks without sacrificing too much accuracy. *Neural Networks*, 154:218–233, 2022. 3
- [32] Jianjun Lei, Xiangwei Zhu, and Ying Wang. Bat: Block and token self-attention for speech emotion recognition. *Neural Networks*, 156:67–80, 2022. 1
- [33] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *International Conference on Learning Representations*, 2018. 1, 3
- [34] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *International Conference on Learning Representations*, 2018. 6, 8
- [35] Dongyu Meng and Hao Chen. Magnet: a two-pronged defense against adversarial examples. In *Proceedings of the 2017 ACM SIGSAC conference on computer and communications security*, pages 135–147, 2017. 3
- [36] Chenggang Mi, Lei Xie, and Yanning Zhang. Improving data augmentation for low resource speech-to-text translation with diverse paraphrasing. *Neural Networks*, 148:194–205, 2022. 1
- [37] Aamir Mustafa, Salman H Khan, Munawar Hayat, Jianbing Shen, and Ling Shao. Image super-resolution as a defense against adversarial attacks. *IEEE Transactions on Image Processing*, 29:1711–1724, 2019. 3
- [38] Yuval Netzer, Tao Wang, Adam Coates, Alessandro Bissacco, Bo Wu, and Andrew Y. Ng. Reading digits in natural images with unsupervised feature learning. In *NIPS Workshop on Deep Learning and Unsupervised Feature Learning 2011*, 2011. 6
- [39] Tribhuvanesh Orekondy, Bernt Schiele, and Mario Fritz. Knockoff nets: Stealing functionality of black-box models. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 4954–4963, 2019. 1, 2
- [40] Nicolas Papernot, Patrick McDaniel, Ian Goodfellow, Somesh Jha, Z Berkay Celik, and Ananthram Swami. Practical black-box attacks against machine learning. In *Proceedings of the 2017 ACM on Asia conference on computer and communications security*, pages 506–519, 2017. 1, 2
- [41] Aaditya Prakash, Nick Moran, Solomon Garber, Antonella DiLillo, and James Storer. Deflecting adversarial attacks with pixel deflection. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 8571–8580, 2018. 3
- [42] Han Qiu, Yi Zeng, Qinkai Zheng, Shangwei Guo, Tianwei Zhang, and Hewu Li. An efficient preprocessing-based approach to mitigate advanced adversarial attacks. *IEEE Transactions on Computers*, 2021. 3, 8
- [43] Waseem Rawat and Zenghui Wang. Deep convolutional neural networks for image classification: A comprehensive review. *Neural Computation*, 29:2352–2449, 2017. 1
- [44] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, pages 234–241. Springer, 2015. 13
- [45] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, pages 234–241. Springer, 2015. 5
- [46] Pouya Samangouei, Maya Kabkab, and Rama Chellappa. Defense-gan: Protecting classifiers against adversarial attacks using generative models. *arXiv preprint arXiv:1805.06605*, 2018. 3
- [47] Sunandini Sanyal, Sravanti Addepalli, and R Venkatesh Babu. Towards data-free model stealing in a hard label setting. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 15284–15293, 2022. 1, 2, 12

- [48] Bo Sun, Nian-hsuan Tsai, Fangchen Liu, Ronald Yu, and Hao Su. Adversarial defense by stratified convolutional sparse coding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 11447–11456, 2019. 3
- [49] Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian J. Goodfellow, and Rob Fergus. Intriguing properties of neural networks. In *International Conference on Learning Representations, ICLR 2014, Conference Track Proceedings*, 2014. 1
- [50] Rajkumar Theagarajan and Bir Bhanu. Defending black box facial recognition classifiers against adversarial attacks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 2020. 3
- [51] Jean-Baptiste Truong, Pratyush Maini, Robert J Walls, and Nicolas Papernot. Data-free model extraction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4771–4780, 2021. 1, 2, 12
- [52] Petra Vidnerová and Roman Neruda. Vulnerability of classifiers to evolutionary generated adversarial examples. *Neural Networks*, 127:168–181, 2020. 1
- [53] Athanasios Voulodimos, Nikolaos D. Doulamis, Anastasios D. Doulamis, and Eftychios E. Protopapadakis. Deep learning for computer vision: A brief review. *Computational Intelligence and Neuroscience*, 2018, 2018. 1
- [54] Haohan Wang, Xindi Wu, Zeyi Huang, and Eric P Xing. High-frequency component helps explain the generalization of convolutional neural networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8684–8694, 2020. 1
- [55] Dong Yin, Raphael Gontijo Lopes, Jonathon Shlens, Ekin Dogus Cubuk, and Justin Gilmer. A fourier perspective on model robustness in computer vision. In *NeurIPS*, 2019. 16
- [56] Hongzhi Zhao, Linguang Hao, Kuangrong Hao, Bing Wei, and Xin Cai. Remix: Towards the transferability of adversarial examples. *Neural Networks*, 163:367–378, 2023. 1
- [57] Mingyi Zhou, Jing Wu, Yipeng Liu, Shuaicheng Liu, and Ce Zhu. Dast: Data-free substitute training for adversarial attacks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 234–243, 2020. 1, 2
- [58] Bojia Zi, Shihao Zhao, Xingjun Ma, and Yu-Gang Jiang. Revisiting adversarial robustness distillation: Robust soft labels make student better. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 16443–16452, 2021. 3

Supplementary for: “Data-free Defense of Black Box Models Against Adversarial Attacks”

1. Ablation on coefficient selection strategy

In our method DBMA, we select only a limited number of coefficients (i.e., $k\%$ coefficients) to obtain a decent tradeoff between the adversarial and clean performance. As shown in Fig 1 (A) in the main paper, the least affected approximate coefficients have a high magnitude. Thus, we select the essential detail coefficients as the top- k high magnitude. As they are likely to be less affected by the adversarial attack (Fig 1 (B)) (main paper), they can yield better performance. In this section, we perform an ablation to assess the effectiveness of selecting top- k coefficients over the other possible choices. For comparative analysis, we perform the experiments using the bottom- k and random- k coefficients. Table 1 shows the adversarial and clean performance for the different coefficient selection strategies. The top- k coefficient selection strategy, selecting the most important coefficient in terms of magnitude, improves both the clean and adversarial performance. On the other hand bottom- k coefficient selection strategy, selecting the least significant coefficients shows the least performance as they primarily consist of contaminated high-frequency content. Although, randomly selecting $k\%$ coefficients showed improved performance than the bottom- k , it still performs poorly compared to top- k coefficient selection strategy.

Table 1. Performance comparison when $k\%$ detail coefficients are selected in wavelet coefficient selection module (WCSM) using different methods. Selection of top- k coefficients yields better clean and adversarial accuracy than other strategies.

Surrogate model (attacker)	coefficient selection strategy	Black Box Model : Alexnet			
		Surrogate Model (defense): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet-half	bottom- k	31.25	8.59	7.32	12.57
	random- k	42.58	13.13	11.74	19.77
	top- k (ours)	77.92	26.66	24.55	34.02
Alexnet	bottom- k	31.25	6.14	4.84	10.92
	random- k	42.92	8.61	7.52	14.29
	top- k (ours)	77.92	15.98	14.04	21.34

2. Performance of our method (DBMA) using different wavelets

The experiments in the main draft, have used Daubechies wavelets [14] in wavelet noise remover (WNR). Along with Daubechies, several other wavelets are available in the literature that varies in time, frequency, and rate of decay. This section analyzes the effect of different wavelets on the performance of proposed data-free black box defense (DBMA). We perform experiments with Coiflets [5] and Biorthogonal wavelets [10]. Table 2 summarizes the results obtained with Resnet18 as the defender’s surrogate model and Alexnet-half as the attacker’s surrogate model. We observe a similar performance trend over the adversarial and clean samples on using the different wavelet functions when compared to the Daubechies. This confirms DBMA yields consistent performance irrespective of the choice of the wavelet function used.

Table 2. Ablation over different choices of wavelets that are used in wavelet noise remover (WNR). The performance of DBMA remains consistent across different choices of wavelet functions.

Surrogate model (Attacker)	Surrogate model (Defender)	Black Box Model : Alexnet			
		Surrogate Model (defender): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet-half	Biorthogonal	73.27	42.51	41.72	51.11
	Daubechies	73.77	42.71	42.71	50.63
	Coiflets	73.14	42.51	44.22	51.94

3. Defense against different data-free black box attacks

In all the previous experiments we assume that both defender and attacker use the same model stealing technique [4] for creating a surrogate model. To prove our data-free black box defense (DBMA) is robust to different model stealing strategies, we evaluate our method against two different approaches: a) Data-free model extraction (DFME) [51] and b) Data-free model stealing in hard label setting (DFMS-HL) [47] that are used by attacker to obtain a surrogate model (Alexnet-Half) for crafting adversarial samples. As shown in Fig. 1, our method yields a consistent boost in adversarial accuracy on different data-free model stealing methods. In the case of DFME, we observe a massive improvement of $\approx 36\%$, $\approx 40\%$ and $\approx 21\%$ on BIM, PGD and Auto Attack respectively. On the other hand, in case of DFMS-HL the performance against the three attacks improves by $\approx 28\%$, $\approx 29\%$, and $\approx 31\%$. Overall across different model stealing methods, our method (DBMA) yields significant improvement in adversarial accuracy (i.e. $\approx 29 - 42\%$ in PGD, $\approx 28 - 38\%$ in BIM, and $\approx 25 - 31\%$ in state-of-the-art Auto Attack).

Further, we check our defense in a more tougher scenario, where attacker is aware of the black box model B_m ’s

Table 3. Performance of DBMA across several data-free black box attacks that are constructed using the surrogate model having same architecture as the black box model. We observe consistent significant improvements in adversarial performance using proposed DBMA across different adversarial attacks and model stealing techniques.

Surrogate model (attacker)	Method	Model stealing (Attacker)	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
			clean	BIM	PGD	Auto Attack
Alexnet	Without DBMA	Black Box Ripper	82.58	4.17	2.19	8.55
	With DBMA (Ours)	Black Box Ripper	73.77	33.31 (↑ 29.14)	31.72 (↑ 29.53)	40.56 (↑ 32.01)
	Without DBMA	DFME	82.58	4.21	1.99	16.93
	With DBMA (Ours)	DFME	73.77	42.84 (↑ 38.63)	42.49 (↑ 40.5)	55.04 (↑ 38.11)
	Without DBMA	DFMS-HL	82.58	3.30	1.82	7.84
	With DBMA (Ours)	DFMS-HL	73.77	26.0 (↑ 22.7)	24.94 (↑ 23.12)	32.57 (↑ 24.73)

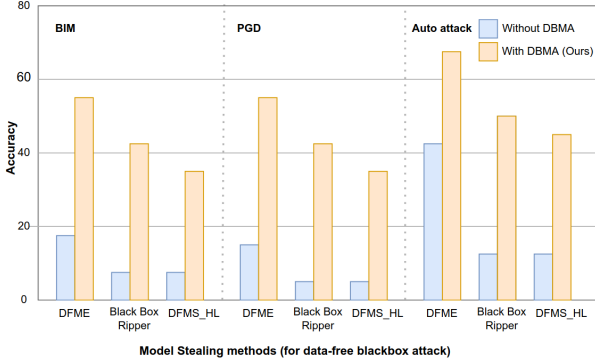


Figure 1. Performance of our approach DBMA for different model stealing methods used to get the attacker’s surrogate model for data-free black box attacks. DBMA consistently improves performance against different attacks across all model stealing methods.

architecture (i.e., Alexnet), and uses the same for attacker’s surrogate model (S_m^a). The results for this setup are reported in Table 3. Compared to baseline, we observe an improvement of $\approx 29 - 32\%$, $\approx 38 - 40\%$ and $\approx 22 - 25\%$ in adversarial accuracy across attacks using the Black box ripper, DFME and DFMS-HL methods, respectively. This indicates that even if the attacker is aware of the black box model’s architecture, our method DBMA can provide data-free black-box adversarial defense irrespective of the different model stealing methods.

Hence, our proposed method DBMA provides strong robustness even when the attacker uses a different model stealing strategy compared to the defender. Refer to next section for adversarial robustness results when the defender uses different model stealing methods to obtain the surrogate model.

4. DBMA using different model stealing methods

In all the experiments in the main draft, we used the Black Box Ripper (BBR) model stealing method to obtain surrogate model for defense. To demonstrate that our method DBMA can work across different model stealing techniques

also, we obtain the defender’s surrogate model using a different model stealing method (DFMS-HL). To get better insights, we analyze the performance of DBMA by varying the model stealing techniques for both attack and defense.

In Table 4 we observe the clean accuracy is not affected on using different model stealing techniques. DBMA obtains the best performance when the attacker uses BBR to attack while the defense is performed using model stealing DFMS-HL (third row). However, when the attacker uses the DFMS-HL method to create a surrogate model (second and fourth row), the adversarial accuracy decreases by $\approx 7 - 11\%$ across attacks compared to the performance with BBR method used by attacker. For attacks crafted using BBR, the defender’s performance with DFMS-HL remains almost similar to BBR (first and third rows). These results suggest that the adversarial samples created using the DFMS-HL are stronger than the ones created using BBR method. This aligns with the Sec. 5.6 in the main draft, where we observed a similar trend indicating that the attacks crafted using DFMS-HL are powerful than BBR. Further, we observe that the choice of the defender’s model stealing method does not significantly affect the adversarial and clean accuracy.

Overall, for different model stealing methods, DBMA ensures good clean performance with decent adversarial performance.

5. Architecture details of Regenerator Network

The Regenerator network consists of U-net-based [44] generator with five downsampling and upsampling layers. Each i^{th} downsampling layer has a skip connection to $(n - i)^{th}$ upsampling layer that concatenates channels of i^{th} layer with those at layer $n - i$, n represents number of upsampling and downsampling layers (i.e. $n = 5$). The number of channels in the network’s input and output are the same as the image channels in the training dataset (surrogate data S_d in our case). Each Downsampling layer first filters input through Leaky relu with negative slope = 0.2, followed by convolution operation with kernel size = 4, stride = 2, and padding = 1. The number of output channels for lay-

Table 4. Performance of DBMA across different model stealing methods used by defender and attacker. DBMA obtains respectable performance irrespective of the model stealing technique used by either the defender or attacker.

Surrogate model (attacker)	Model stealing (defender)	Model stealing (attacker)	Black Box Model : Alexnet			
			Surrogate Model (defense): Resnet-18			
Alexnet-half	BBR	BBR	73.77	42.71	42.71	50.63
	BBR	DFMS-HL	73.77	35.4	34.58	43.05
	DFMS-HL	BBR	72.45	42.75	43.08	51.0
	DFMS-HL	DFMS-HL	72.45	32.85	31.86	40.6

Table 5. Performance of DBMA with and without regenerator network across different values of k . For low values, regenerator network improves both clean and adversarial accuracy. For $k=16$, small decrease in clean accuracy, but adversarial accuracy increases significantly.

Surrogate model (attacker)	Coefficients ($k\%$)	Method	Black Box Model : Alexnet			
			Surrogate Model (defense): Resnet-18			
Alexnet-half	1	WNR	clean	BIM	PGD	Auto Attack
		WNR + R_n	42.75	14.89	13.79	21.41
	2	WNR	56.08	30.58 ($\uparrow$ 15.69)	30.14 ($\uparrow$ 16.35)	37.04 ($\uparrow$ 15.63)
		WNR + R_n	50.17	17.34	16.38	25.36
	4	WNR	60.35	34.94 ($\uparrow$ 17.60)	34.61 ($\uparrow$ 18.23)	41.61 ($\uparrow$ 16.25)
		WNR + R_n	59.14	21.99	20.77	29.43
	8	WNR	65.82	39.65 ($\uparrow$ 17.66)	39.62 ($\uparrow$ 18.85)	46.94 ($\uparrow$ 17.51)
		WNR + R_n	69.89	24.9	23.54	33.04
	16	WNR	70.37	41.89 ($\uparrow$ 16.99)	42.21 ($\uparrow$ 18.67)	49.88 ($\uparrow$ 16.84)
		WNR + R_n	77.92	26.66	24.55	34.02
	50	WNR	73.77	42.71 ($\uparrow$ 16.05)	42.71 ($\uparrow$ 18.16)	50.63 ($\uparrow$ 16.61)
		WNR + R_n	82.58	11.36	8.23	17.34
Alexnet	1	WNR	75.19	33.12 ($\uparrow$ 21.76)	31.60 ($\uparrow$ 23.37)	40.11 ($\uparrow$ 22.77)
		WNR + R_n	42.75	10.20	8.72	15.80
	2	WNR	56.08	24.02 ($\uparrow$ 14.82)	23.13 ($\uparrow$ 14.41)	30.55 ($\uparrow$ 14.75)
		WNR + R_n	50.17	15.37	14.14	21.92
	4	WNR	60.34	32.89 ($\uparrow$ 17.52)	32.24 ($\uparrow$ 18.10)	40.50 ($\uparrow$ 18.58)
		WNR + R_n	59.14	17.54	16.03	25.08
	8	WNR	65.82	31.00 ($\uparrow$ 13.46)	30.85 ($\uparrow$ 14.82)	38.26 ($\uparrow$ 13.18)
		WNR + R_n	69.89	19.65	18.30	27.73
	16	WNR	70.37	34.13 ($\uparrow$ 14.48)	33.61 ($\uparrow$ 15.31)	41.08 ($\uparrow$ 13.35)
		WNR + R_n	77.92	15.98	14.04	21.34
	50	WNR	73.77	33.31 ($\uparrow$ 17.33)	31.72 ($\uparrow$ 17.68)	40.56 ($\uparrow$ 19.22)
		WNR + R_n	82.58	5.58	3.33	10.44
		WNR	75.19	19.65 ($\uparrow$ 14.07)	18.38($\uparrow$ 15.04)	25.41 ($\uparrow$ 14.97)
		WNR + R_n				

ers 1 to 5 are 64,128,256,512,512 respectively. Upsampling layers start with a Relu layer. Followed by transposed convolution. Each transposed convolution has $kernelsize = 4$, $padding = 1$ and $stride = 2$. The number of output channels for layers 1 to 5 are 1024, 512, 256, 128, and 3 respectively. The output of convolution and deconvolution layers is normalized using instance normalization to avoid 'instance-specific mean and covariance shifts'. The output of the last upsampling layer is normalized using Tanh normalization to ensure output values lie in the range $[-1, 1]$.

6. Importance of Regenerator Network

To evaluate the effectiveness of proposed Regenerator network R_n in our approach DBMA, we analyze the perfor-

mance of DBMA with and without R_n for different values of k (i.e. $k=1, 2, 4, 8, 16$ (ours), 50). In Table 5, we observe that across different k , appending R_n to the WNR improves the adversarial accuracy against various attacks crafted using Alexnet-half and Alexnet by $\approx 13 - 18\%$. Further, we observe a similar trend for clean accuracy, which also improves on adding R_n to WNR. However, the improvement margin for clean accuracy gradually drops on increasing the value of coefficient percent k . For higher k 's (e.g., 16, 50), there is a small drop in clean performance using R_n compared to the performance with only WNR but leads to significant increase in adversarial accuracy. This implies regenerator network enhances the output image of WNR to increase the adversarial accuracy. In this process, for smaller

Table 6. Ablation on defense components of proposed DBMA and comparison with baseline. Both the components individually provide better defense than baseline. DBMA yields best performance when WNR and R_n are used together.

Surrogate model (attacker)	Defense Components	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet-half	Baseline	82.58	7.02	4.53	11.65
	WNR	77.92	26.66 ($\uparrow 19.69$)	24.55 ($\uparrow 20.02$)	34.02 ($\uparrow 22.37$)
	R_n	77.03	29.40 ($\uparrow 22.38$)	28.32 ($\uparrow 23.79$)	37.16 ($\uparrow 25.51$)
	WNR + R_n	73.77	42.71 ($\uparrow 35.69$)	42.71 ($\uparrow 38.18$)	50.63 ($\uparrow 38.98$)
Alexnet	Baseline	82.58	4.17	2.19	8.55
	WNR	77.92	15.98 ($\uparrow 11.81$)	14.04 ($\uparrow 11.85$)	21.34 ($\uparrow 12.79$)
	R_n	77.03	16.52 ($\uparrow 12.35$)	15.09 ($\uparrow 12.9$)	22.40 ($\uparrow 13.85$)
	WNR + R_n	73.77	33.31 ($\uparrow 29.14$)	31.72 ($\uparrow 29.53$)	40.56 ($\uparrow 32.01$)

values of k , it increases clean accuracy too, but for a large value of k , the decrease in clean accuracy is compensated by the increase in adversarial accuracy to achieve the best trade-off. Combining WNR with the regenerator network at our $\hat{k}$ (i.e., $k = 16$) produces the best adversarial accuracy.

From Table 6, we observe that DBMA with only WNR improves adversarial accuracy with a small drop in clean accuracy compared to baseline. Similarly, with only regenerator network R_n , adversarial accuracy increases compared to the baseline. Also, R_n performs better than WNR. However, using WNR and R_n together in DBMA gives the best adversarial accuracy. Hence this demonstrates the importance of both the defense components (WNR and R_n) in our method DBMA.

7. Ablation on choice of surrogate architecture

To better analyze the performance of DBMA against different combinations of defender surrogate (S_m^d) and attacker surrogate models (S_m^a), we perform experiments with different choices of surrogate models (i.e., Alexnet-half, Alexnet, and Resnet18). For all the experiments, Alexnet is used as the black box model. For Resnet18 as S_m^d , the wavelet coefficient selection module (WCSM) yields optimal k ($\hat{k}$) as 16. Similarly, for other choices of S_m^d (i.e., Alexnet and Alexnet-half), we obtain $\hat{k}$ as 15. This shows that the value of $\hat{k}$ is not much sensitive to the choice of architecture of the defender's surrogate model S_m^d . The results are reported in Table 7.

We obtain the best performance for Resnet-18 as S_m^d and Alexnet-half as S_m^a (1st row), whereas the lowest for Alexnet-half as S_m^d and Resnet18 as S_m^a (6th row). Further, on carefully observing the results, we deduce some key insights. Clean accuracy remains similar across different choices of surrogate models, but adversarial accuracy depends on the surrogate model of defender (S_m^d) and attacker (S_m^a).

We observe that the bigger the network size, the more accurate the surrogate models. With accurate surrogate models, the gradients with respect to the input are better esti-

mated. Thus better black-box attacks and defenses can be obtained using the bigger architectures for surrogate models. For better defense, S_m^d should have a relatively higher capacity than S_m^a . This can be confirmed by rows 1, 4, and 7, where the defense becomes more effective on increasing the S_m^d 's capacity against various attacks using Alexnet-half as S_m^a . A similar trend is observed against the attacks crafted using Alexnet and Resnet18. For the other way around, i.e., when S_m^a has relatively higher capacity than S_m^d , more powerful attacks can be crafted. This is evident from rows 1, 2, and 3, where stronger adversarial samples are obtained on increasing the capacity of S_m^a for a given S_m^d . For instance, for Resnet18 as S_m^d , stronger attacks (lower adversarial accuracy) are obtained by using Resnet18 as S_m^a , followed by Alexnet and Alexnet-half. A similar pattern is observed on other choices of S_m^d .

8. Our defense (DBMA) on larger black box model

Throughout all our experiments, we used Alexnet as the black-box model B_m . To check the consistency of our approach DBMA across different architecture, especially for bigger and high-capacity networks, we perform experiments using Resnet34 as black-box model and report corresponding results in Table 8. With Resnet18 and Alexnet as the defender's and attacker's surrogate model (S_m^d and S_m^a) respectively, we observe the improvement of $\approx 27 - 32\%$ in adversarial accuracy across attacks compared to baseline (rows 3rd and 4th). With Alexnet as the defender's surrogate and Resnet18 as the attacker's surrogate model, we get an improvement of $\approx 17 - 23\%$ across attacks (rows 1st and 2nd). As observed in previous experiments, compared to baseline, clean accuracy drops by $\approx 7 - 8\%$. Overall, across different black box models, our proposed defense DBMA has obtained decent performance. Hence we conclude, DBMA is even effective on bigger black box architectures.

Table 7. Investigating the effect of surrogate model’s architecture (for both defender and attacker) on the performance of our proposed approach (DBMA). Given defender’s surrogate model, the attack is stronger if larger surrogate model is used by the attacker.

Surrogate model (defender)	Surrogate model (attacker)	Black Box Model : Alexnet			
		clean	BIM	PGD	Auto Attack
Resnet-18	Alexnet-half	73.77	42.71	42.71	50.63
Resnet-18	Alexnet	73.77	33.31	31.72	40.56
Resnet-18	Resnet-18	73.77	22.48	21.93	29.48
Alexnet-half	Alexnet-half	74.94	38.98	39.3	47.83
Alexnet-half	Alexnet	74.94	29.04	27.58	35.37
Alexnet-half	Resnet-18	74.94	20.72	19.11	26.79
Alexnet	Alexnet-half	74.67	40.08	39.6	48.5
Alexnet	Alexnet	74.67	29.49	28.63	36.93
Alexnet	Resnet-18	74.67	21.51	20.24	28.44

Table 8. Performance of our approach DBMA in defending the larger black-box network (i.e., Resnet34). Across various combinations, DBMA shows a consistent improvement against the different attacks with a small drop in the clean accuracy.

Method	Surrogate Model (Defender)	Surrogate Model (Attacker)	Black Box Model : Resnet34			
			Clean	BIM	PGD	Auto-Attack
Baseline	-	Resnet18	95.66	3.11	1.23	11.82
DBMA (Ours)	Alexnet	Resnet18	87.06	20.76	17.33	25.11
Baseline	-	Alexnet	95.66	21.36	12.83	27.85
DBMA (Ours)	Resnet18	Alexnet	88.40	48.16	44.80	56.65

Table 9. Performance of our approach DBMA with fourier transform and wavelet transform based noise removal technique (FNR and WNR, respectively). WNR defense outperforms the FNR across different attacks. Also, WNR-based DBMA ($WNR + R_n$) yields more significant gains in performance on CIFAR10.

Surrogate model (attacker)	Method	Black Box Model : Alexnet Surrogate Model (defense): Resnet-18			
		clean	BIM	PGD	Auto Attack
Alexnet- half	Baseline	82.58	7.02	4.53	11.65
	FNR	79.14	13.30 ($\uparrow 6.28$)	10.86 ($\uparrow 6.33$)	20.47 ($\uparrow 8.82$)
	FNR+ R_n	74.13	28.38 ($\uparrow 21.36$)	27.05 ($\uparrow 22.52$)	35.47 ($\uparrow 23.82$)
	WNR	77.92	26.66 ($\uparrow 19.64$)	24.55 ($\uparrow 20.02$)	34.02 ($\uparrow 22.37$)
	WNR + R_n (Ours)	73.77	42.71 ($\uparrow 35.69$)	42.71 ($\uparrow 38.18$)	50.63 ($\uparrow 38.98$)
Alexnet	Baseline	82.58	4.17	2.19	8.55
	FNR	79.14	5.87 ($\uparrow 1.70$)	4.03 ($\uparrow 1.84$)	10.97 ($\uparrow 2.42$)
	FNR+ R_n	74.13	19.24 ($\uparrow 15.07$)	17.88 ($\uparrow 15.69$)	24.55 ($\uparrow 16.00$)
	WNR	77.92	15.98 ($\uparrow 11.81$)	14.04 ($\uparrow 11.85$)	21.34 ($\uparrow 12.79$)
	WNR + R_n (Ours)	73.77	33.31 ($\uparrow 29.14$)	31.72 ($\uparrow 29.53$)	40.56 ($\uparrow 32.01$)

9. Comparison with Fourier Transform based Noise removal

Apart from the wavelet transformations, some recent works utilised the fourier transformations to remove the adversarial noise from the adversarial images, and further found it to be effective in denoising [55]. In this section, we do an ablation on our choice of Wavelet-based Noise Remover (WNR) over the other possible choice of fourier-based Noise Remover (FNR).

As observed in recent works [55], adversarial attack affects the high-frequency components more than low-

frequency components. Therefore, In FNR we apply a low pass filter on an image with threshold radius $\hat{r}$. Similar to WCSM, we compute LCR_C , LCR_A , LCR and ROC for different values of r . The optimal value of r (i.e., $\hat{r}$) is selected at which ROC starts saturating ($\hat{r} = 11$). In Table 9, we observe, compared to baseline, Fourier-based DBMA gives an improvement of $\approx 21 - 34\%$ in adversarial accuracy across attacks for Alexnet half (rows 1 and 3). Compared to Fourier-based DBMA, the results using our wavelet-based DBMA are significantly better in terms of adversarial accuracy (rows 1 and 5) with similar clean performance. Similarly, for Alexnet, Wavelet-based DBMA

gives $\approx 16 - 24\%$ better adversarial accuracy compared to Fourier based DBMA across all attacks (rows 8 and 10). Hence, our wavelet-based DBMA is more robust across different adversarial attacks than Fourier-based DBMA.

10. Algorithm

Algorithm 1 Algorithm for our proposed method (DBMA)

Input: Black box model B_m , max coefficients k^{max}

Output: $\hat{B}_m$

Step 1: Model Stealing

1: Surrogate model S_m , Synthetic data $S_d \leftarrow$ Model Stealing on B_m

Step 2: Wavelet Coefficient Selection Module

2: Obtain adversarial samples (S_{da}) corresponding to S_d using adversarial attack on S_m

3: **for** $k = 1 : k^{max}$: **do**

4: $\bar{S}_{da}^k \leftarrow \text{WNR}(S_{da}, k)$

5: $N_{flips} = 0$

6: **for** $i = 1 : (|\bar{S}_{da}^k|)$ **do**

7: $\bar{x}_{sa}^i \leftarrow \bar{S}_{da}^k[i]$ $\{i^{th} \text{ element of } \bar{S}_{da}^k\}$

8: $x_s^i \leftarrow S_d[i]$ $\{i^{th} \text{ element of } S_d\}$

9: **if** $\text{label}(B_m(\bar{x}_{sa}^i)) \neq \text{label}(B_m(x_s^i))$ **then**

10: $N_{flips} = N_{flips} + 1$

11: **end if**

12: **end for**

13: $LFR^k = N_{flips} / |\bar{S}_{da}^k|$

14: **end for**

15: $\hat{k} = \underset{k}{\text{argmin}} LFR^k$

Step 3: Training Regenerator network (R_n)

16: $\bar{S}_d^{\hat{k}} \leftarrow \text{WNR}(S_d, \hat{k})$

17: $\bar{S}_{da}^{\hat{k}} \leftarrow \text{WNR}(S_{da}, \hat{k})$

18: Initialize R_n^θ

19: **for** $epoch < MaxEpoch$ **do**

20: **for** $i = 1 : (|S_d|)$ **do**

21: $x_s^i \leftarrow S_d[i]$ $\{i^{th} \text{ element of } S_d\}$

22: $\bar{x}_s^i \leftarrow \bar{S}_d^{\hat{k}}[i]$ $\{i^{th} \text{ element of } \bar{S}_d^{\hat{k}}\}$

23: $\bar{x}_{sa}^i \leftarrow \bar{S}_{da}^{\hat{k}}[i]$ $\{i^{th} \text{ element of } \bar{S}_{da}^{\hat{k}}\}$

24: $L_{cs} = CS(S_m(R_n(\bar{x}_{sa}^i)), S_m(x_s^i))$ $\{CS \text{ is cosine similarity}\}$

25: $L_{kl} = KL(\text{soft}(S_m(R_n(\bar{x}_{sa}^i))), \text{soft}(S_m(R_n(\bar{x}_s^i))))$ $\{KL \text{ is KL divergence}\}$

26: $L_{sc} = \|R_n(\bar{x}_s^i) - x_s^i\|_1 + \|R_n(\bar{x}_{sa}^i) - x_s^i\|_1$

27: $L(R_n^\theta) = -\lambda_1 L_{cs} + \lambda_2 L_{kl} + \lambda_3 L_{sc}$

28: Update R_n^θ by minimizing $L(R_n^\theta)$ using Adam Optimizer

29: **end for**

30: **end for**

31: $\hat{B}_m \leftarrow \text{concatenate}(\text{WNR}(\cdot, \hat{k}), R_n^{\theta^*}, B_m)$ $\{\text{The black box model } \hat{B}_m \text{ is used by attacker}\}$

32: **return** $\hat{B}_m$

11. Visualization

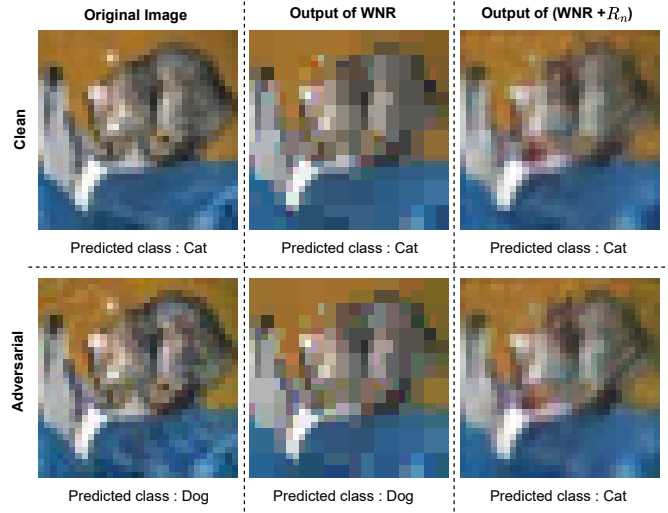


Figure 2. Visualization of images: The top row indicates input as clean image and bottom row corresponds to adversarial image. The predictions obtained by the black-box network on inputs: (a) Original clean image (b) Output of wavelet noise remover on clean image (c) Output of WNR with regenerator network (R_n) on clean image (d) Original adversarial image (e) Output of wavelet noise remover (WNR) on adversarial image (f) Output of WNR with regenerator network (R_n) on adversarial Image. Here, the ground truth class is Cat. Our method (DBMA) produces correct output using regenerated image as input.

Table 10. Attack parameters for different adversarial attacks: BIM, PGD and Auto Attack

Dataset	Attack Parameters	Adversarial Attacks		
		BIM	PGD	Auto Attack
CIFAR-10	ϵ	8/255	8/255	8/255
	ϵ_{step}	0.00156 (ϵ /no of iterations)	2/255	–
	no of iterations	20	20	–
SVHN	ϵ	4/255	4/255	4/255
	ϵ_{step}	2/255	2/255	–
	no of iterations	20	20	–

12. Adversarial Attack Parameters and Training Details

We evaluate the performance of black box model (B_m) on three adversarial attacks, PGD, BIM and Auto Attack. Parameters used for each attack are summarized in Table 10.

Training details of Regenerator Network (R_n): The regenerator network is trained with Adam optimizer with learning rate of 0.0002 for 300 epochs. Learning rate is decayed using linear scheduler, where we keep the learning rate fixed for 100 epochs and then linearly decay the rate to zero. Batch size is set to 128.