

Adversarial Examples: Attacks and Defenses for Deep Learning

Xiaoyong Yuan, Pan He, Qile Zhu, Xiaolin Li*

National Science Foundation Center for Big Learning, University of Florida
{chbrian, pan.he, valder}@ufl.edu, andyli@ece.ufl.edu

Abstract—With rapid progress and significant successes in a wide spectrum of applications, deep learning is being applied in many safety-critical environments. However, deep neural networks have been recently found vulnerable to well-designed input samples, called *adversarial examples*. Adversarial perturbations are imperceptible to human but can easily fool deep neural networks in the testing/deploying stage. The vulnerability to adversarial examples becomes one of the major risks for applying deep neural networks in safety-critical environments. Therefore, attacks and defenses on adversarial examples draw great attention. In this paper, we review recent findings on adversarial examples for deep neural networks, summarize the methods for generating adversarial examples, and propose a taxonomy of these methods. Under the taxonomy, applications for adversarial examples are investigated. We further elaborate on countermeasures for adversarial examples. In addition, three major challenges in adversarial examples and the potential solutions are discussed.

Index Terms—deep neural network, deep learning, security, adversarial examples

I. INTRODUCTION

Deep learning (DL) has made significant progress in a wide domain of machine learning (ML): image classification, object recognition [1], [2], object detection [3], [4], speech recognition [5], language translation [6], voice synthesis [7]. Many deep learning empowered applications are life crucial, raising great concerns in the field of safety and security. Despite great successes in numerous applications, recent studies find that deep learning is vulnerable against well-designed input samples. These samples can easily fool a well-performed deep learning model with little perturbations imperceptible to humans.

Szegedy *et al.* first generated small perturbations on the images for the image classification problem and fooled state-of-the-art deep neural networks with high probability [8]. These misclassified samples were named as *Adversarial Examples*.

Extensive deep learning based applications have been used or planned to be deployed in the physical world, especially in the safety-critical environments. In the meanwhile, recent studies show that adversarial examples can be applied to real world. For instance, an adversary can construct physical adversarial examples and confuse autonomous vehicles by manipulating the stop sign in a traffic sign recognition system [9], [10] or removing the segmentation of pedestrians

in an object recognition system [11]. Attackers can generate adversarial commands against automatic speech recognition (ASR) models and Voice Controllable System (VCS) [12], [13] such as Apple Siri [14], Amazon Alexa [15], and Microsoft Cortana [16].

Deep learning is widely regarded as a “black box” technique — we all know that it performs well, but with limited knowledge of the reason [17], [18]. Many studies have been proposed to explain and interpret deep neural networks [19]–[22]. From inspecting adversarial examples, we may gain insights on semantic inner levels of neural networks [23] and find problematic decision boundaries, which in turn helps to increase robustness and performance of neural networks [24] and improve the interpretability [25].

In this paper, we investigate and summarize approaches for generating adversarial examples, applications for adversarial examples and corresponding countermeasures. We explore the characteristics and possible causes of adversarial examples. Recent advances in deep learning revolve around supervised learning, especially in the field of computer vision task. Therefore, most adversarial examples are generated against computer vision models. We mainly discuss adversarial examples for image classification/object recognition tasks in this paper. Adversarial examples for other tasks will be investigated in Section V.

Inspired by [26], we define the **Threat Model** as follows:

- The adversaries can attack only at the testing/deploying stage. They can tamper only the input data in the testing stage after the victim deep learning model is trained. Neither the trained model or the training dataset can be modified. The adversaries may have knowledge of trained models (architectures and parameters) but are not allowed to modify models, which is a common assumption for many online machine learning services. Attacking at the training stage (*e.g.*, training data poisoning) is another interesting topic and has been studied in [27]–[32]. Due to the limitation of space, we do not include this topic in the paper.
- We focus on attacks against models built with deep neural networks, due to their great performance achieved. We will discuss adversarial examples against conventional machine learning (*e.g.*, SVM, Random Forest) in Section II. Adversarial examples against deep neural networks proved effective in conventional machine learning models [33] (see Section VII-A).

This work is partially supported by National Science Foundation (grants CNS-1842407, CNS-1747783, CNS-1624782, and OAC-1229576).

* Corresponding author.

Table I: Notation and symbols used in this paper

Notations and Symbols	Description
x	original (clean, unmodified) input data
l	label of class in the classification problem. $l = 1, 2, \dots, m$, where m is the number of classes
x'	adversarial example (modified input data)
l'	label of the adversarial class in targeted adversarial examples
$f(\cdot)$	deep learning model (for the image classification task, $f \in F : \mathbb{R}^n \rightarrow l$)
θ	parameters of deep learning model f
$J_f(\cdot, \cdot)$	loss function (e.g., cross-entropy) of model f
η	difference between original and modified input data: $\eta = x' - x$ (the exact same size as the input data)
$\ \cdot\ _p$	ℓ_p norm
∇	gradient
$H(\cdot)$	Hessian, the second-order of derivatives
$KL(\cdot)$	Kullback-Leibler (KL) divergence function

- Adversaries only aim at compromising *integrity*. Integrity is presented by performance metrics (e.g., accuracy, F1 score, AUC), which is essential to a deep learning model. Although other security issues pertaining to confidentiality and privacy have been drawn attention in deep learning [34]–[36], we focus on the attacks that degrade the performance of deep learning models, cause an increase of false positives and false negatives.
- The rest of the threat model differs in different adversarial attacks. We will categorize them in Section III.

Notations and symbols used in this paper are listed in Table I.

This paper presents the following contributions:

- To systematically analyze approaches for generating adversarial examples, we taxonomize attack approaches along different axes to provide an accessible and intuitive overview of these approaches.
- We investigate recent approaches and their variants for generating adversarial examples and compare them using the proposed taxonomy. We show examples of selected applications from fields of reinforcement learning, generative modeling, face recognition, object detection, semantic segmentation, natural language processing, and malware detection. Countermeasures for adversarial examples are also discussed.
- We outline main challenges and potential future research directions for adversarial examples based on three main problems: transferability of adversarial examples, existence of adversarial examples, and robustness evaluation of deep neural networks.

The remaining of this paper is organized as follows. Section II introduces the background of deep learning techniques, models, and datasets. We discuss adversarial examples raised in conventional machine learning in Section II. We propose a taxonomy of approaches for generating adversarial examples in Section III and elaborate on these approaches in Section IV. In Section V, we discuss applications for adversarial examples. Corresponding countermeasures are investigated in Section VI. We discuss current challenges and potential solutions in Section VII. Section VIII concludes the work.

II. BACKGROUND

In this section, we briefly introduce basic deep learning techniques and approaches related to adversarial examples. Next, we review adversarial examples in the era of conventional ML and compare the difference between adversarial examples in conventional ML and that in DL.

A. Brief Introduction to Deep Learning

This subsection discusses main concepts, existed techniques, popular architectures, and standard datasets in deep learning, which, due to the extensive use and breakthrough successes, have become acknowledged targets of attacks, where adversaries are usually applied to evaluate their attack methods.

1) *Main concepts in deep learning:* Deep learning is a type of machine learning methods that makes computers to learn from experience and knowledge without explicit programming and extract useful patterns from raw data. For conventional machine learning algorithms, it is difficult to extract well-represented features due to limitations, such as curse of dimensionality [37], computational bottleneck [38], and requirement of the domain and expert knowledge. Deep learning solves the problem of representation by building multiple simple features to represent a sophisticated concept. For example, a deep learning-based image classification system represents an object by describing edges, fabrics, and structures in the hidden layers. With the increasing number of available training data, deep learning becomes more powerful. Deep learning models have solved many complicated problems, with the help of hardware acceleration in computational time.

A neural network layer is composed of a set of perceptrons (artificial neurons). Each perceptron maps a set of inputs to output values with an activation function. The function of a neural network is formed in a chain:

$$f(x) = f^{(k)}(\dots f^{(2)}(f^{(1)}(x))), \quad (1)$$

where $f^{(i)}$ is the function of the i th layer of the network, $i = 1, 2, \dots, k$.

Convolutional neural networks (CNNs) and Recurrent neural networks (RNNs) are two most widely used neural networks in recent neural network architectures. CNNs deploy convolution operations on hidden layers to share weights and reduce the number of parameters. CNNs can extract local information from grid-like input data. CNNs have shown incredible successes in computer vision tasks, such as image classification [1], [39], object detection [4], [40], text recognition [41], [42], and semantic segmentation [43], [44]. RNNs are neural networks for processing sequential input data with variable length. RNNs produce outputs at each time step. The hidden neuron at each time step is calculated based on current input data and hidden neurons at the previous time step. Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU) with controllable gates are designed to avoid vanishing/exploding gradients of RNNs in long-term dependency.

2) *Architectures of deep neural networks*: Several deep learning architectures are widely used in computer vision tasks: *LeNet* [45], *VGG* [2], *AlexNet* [1], *GoogLeNet* [46]–[48] (Inception V1–V4), and *ResNet* [39], from the simplest (oldest) network to the deepest and the most complex (newest) one. *AlexNet* first showed that deep learning models can largely surpass conventional machine learning algorithms in the ImageNet 2012 challenge and led the future study of deep learning. These architectures made tremendous breakthroughs in the ImageNet challenge and can be seen as milestones in image classification problem. Attackers usually generate adversarial examples against these baseline architectures.

3) *Standard deep learning datasets*: *MNIST*, *CIFAR-10*, *ImageNet* are three widely used datasets in computer vision tasks. The *MNIST* dataset is for handwritten digits recognition [49]. The *CIFAR-10* dataset and the *ImageNet* dataset are for image recognition task [50]. The *CIFAR-10* consists of 60,000 tiny color images (32×32) with ten classes. The *ImageNet* dataset consists 14,197,122 images with 1,000 classes [51]. Because of the large number of images in the *ImageNet* dataset, most adversarial approaches are evaluated on only part of the *ImageNet* dataset. The *Street View House Numbers (SVHN)* dataset, similar to the *MNIST* dataset, consists of ten digits obtained from real-world house numbers in Google Street View images. The *YoutubeDataset* dataset is gained from Youtube consisting of about ten million images [45] and used in [8].

B. Adversarial Examples and Countermeasures in Machine Learning

Adversarial examples in conventional machine learning models have been discussed since decades ago. Machine learning-based systems with handcrafted features are primary targets, such as spam filters, intrusion detection, biometric authentication, fraud detection, etc. [52]. For example, spam emails are often modified by adding characters to avoid detection [53]–[55].

Dalvi *et al.* first discussed adversarial examples and formulated this problem as a game between adversary and classifier (Naïve Bayes), both of which are sensitive to cost [53]. The attack and defense on adversarial examples became an iterative game. Biggio *et al.* first tried a gradient-based approach to generate adversarial examples against linear classifier, support vector machine (SVM), and a neural network [56]. Compared with deep learning adversarial examples, their methods allow more freedom to modify the data. The *MNIST* dataset was first evaluated under their attack, although a human could easily distinguish the adversarial digit images. Biggio *et al.* also reviewed several proactive defenses and discussed reactive approaches to improve the security of machine learning models [28].

Barreno *et al.* presented an initial investigation on the security problems of machine learning [52], [57]. They categorized attacking against machine learning system into three axes: 1) influence: whether attacks can poison the training data; 2) security violation: whether an adversarial example belongs to false positive or false negative; 3) specificity: attack is targeted

to a particular instance or a wide class. We discuss these axes for deep learning area in Section III. Barreno *et al.* compared attacks against SpamBayes spam filter and defenses as a study case. However, they mainly focused on binary classification problem such as virus detection system, intrusion detection system (IDS), and intrusion prevention system (IPS).

Adversarial examples in conventional machine learning require knowledge of feature extraction, while deep learning usually needs only raw data input. In conventional ML, both attacking and defending methods paid great attention to features, even the previous step (data collection), giving less attention to the impact of humans. Then the target becomes a fully automatic machine learning system. Inspired by these studies on conventional ML, in this paper, we review recent security issues in the deep learning area.

[26] provided a comprehensive overview of security issues in machine learning and recent findings in deep learning. [26] established a unifying threat model. A “no free lunch” theorem was introduced: the tradeoff between accuracy and robustness.

Compared to their work, our paper focuses on adversarial examples in deep learning and has a detailed discussion on recent studies and findings.

For example, adversarial examples in an image classification task can be described as follows: Using a trained image classifier published by a third party, a user inputs one image to get the prediction of class label. Adversarial images are original clean images with small perturbations, often barely recognizable by humans. However, such perturbations misguide the image classifier. The user will get a response of an incorrect image label. Given a trained deep learning model f and an original input data sample x , generating an adversarial example x' can generally be described as a box-constrained optimization problem:

$$\begin{aligned} \min_{x'} \quad & \|x' - x\| \\ \text{s.t.} \quad & f(x') = l', \\ & f(x) = l, \\ & l \neq l', \\ & x' \in [0, 1], \end{aligned} \quad (2)$$

where l and l' denote the output label of x and x' , $\|\cdot\|$ denotes the distance between two data sample. Let $\eta = x' - x$ be the perturbation added on x . This optimization problem minimizes the perturbation while misclassifying the prediction with a constraint of input data. In the rest of the paper, we will discuss variants of generating adversarial images and adversarial examples in other tasks.

III. TAXONOMY OF ADVERSARIAL EXAMPLES

To systematically analyze approaches for generating adversarial examples, we analyze the approaches for generating adversarial examples (see details in Section IV) and categorize them along three dimensions: threat model, perturbation, and benchmark.

A. Threat Model

We discuss the threat model in Section I. Based on different scenarios, assumptions, and quality requirements, adversaries

decide the attributes they need in adversarial examples and then deploy specific attack approaches. We further decompose the threat model into four aspects: adversarial falsification, adversary's knowledge, adversarial specificity, and attack frequency. For example, if an adversarial example is required to be generated in real-time, adversaries should choose a one-time attack instead of an iterative attack, in order to complete the task (see Attack Frequency).

- Adversarial Falsification

- *False positive* attacks generate a negative sample which is misclassified as a positive one (Type I Error). In a malware detection task, a benign software being classified as malware is a false positive. In an image classification task, a false positive can be an adversarial image unrecognizable to human, while deep neural networks predict it to a class with a high confidence score. Figure 2 illustrates a false positive example of image classification.
- *False negative* attacks generate a positive sample which is misclassified as a negative one (Type II Error). In a malware detection task, a false negative can be the condition that a malware (usually considered as positive) cannot be identified by the trained model. False negative attack is also called machine learning evasion. This error is shown in most adversarial images, where a human can recognize the image, but the neural networks cannot identify it.

- Adversary's Knowledge

- *White-box* attacks assume the adversary knows everything related to trained neural network models, including training data, model architectures, hyper-parameters, numbers of layers, activation functions, model weights. Many adversarial examples are generated by calculating model gradients. Since deep neural networks tend to require only raw input data without handcrafted features and to deploy end-to-end structure, feature selection is not necessary compared to adversarial examples in machine learning.
- *Black-box* attacks assume the adversary has no access to the trained neural network model. The adversary, acting as a standard user, only knows the output of the model (label or confidence score). This assumption is common for attacking online Machine Learning services (e.g., Machine Learning on AWS¹, Google Cloud AI², BigML³, Clarifai⁴, Microsoft Azure⁵, IBM Bluemix⁶, Face++⁷). Most adversarial example attacks are *white-box attacks*. However, they can be transferred to attack *black-box* services due to the transferability of adversarial examples proposed by Papernot *et al.* [33]. We will elaborate on it in Section VII-A.

- Adversarial Specificity

- *Targeted* attacks misguide deep neural networks to a

specific class. Targeted attacks usually occur in the multi-class classification problem. For example, an adversary fools an image classifier to predict all adversarial examples as one class. In a face recognition/biometric system, an adversary tries to disguise a face as an authorized user (Impersonation) [58]. Targeted attacks usually maximize the probability of targeted adversarial class.

- *Non-targeted* attacks do not assign a specific class to the neural network output. The adversarial class of output can be arbitrary except the original one. For example, an adversary makes his/her face misidentified as an arbitrary face in face recognition system to evade detection (dodging) [58]. Non-targeted attacks are easier to implement compared to *targeted* attacks since it has more options and space to redirect the output. Non-targeted adversarial examples are usually generated in two ways: 1) running several targeted attacks and taking the one with the smallest perturbation from the results; 2) minimizing the probability of the correct class.

Some generation approaches (e.g., extended BIM, ZOO) can be applied to both targeted and non-targeted attacks. For binary classification, *targeted* attacks are equivalent to *non-targeted* attacks.

- Attack Frequency

- *One-time attacks* take only one time to optimize the adversarial examples.
- *Iterative attacks* take multiple times to update the adversarial examples.

Compared with *one-time attacks*, *iterative attacks* usually perform better adversarial examples, but require more interactions with victim classifier (more queries) and cost more computational time to generate them. For some computational-intensive tasks (e.g., reinforcement learning), *one-time attacking* may be the only feasible choice.

B. Perturbation

Small perturbation is a fundamental premise for adversarial examples. Adversarial examples are designed to be close to the original samples and imperceptible to a human, which causes the performance degradation of deep learning models compared to that of a human. We analyze three aspects of perturbation: perturbation scope, perturbation limitation, and perturbation measurement.

- Perturbation Scope

- *Individual* attacks generate different perturbations for each clean input.
- *Universal* attacks only create a universal perturbation for the whole dataset. This perturbation can be applied to all clean input data.

Most of the current attacks generate adversarial examples individually. However, *universal* perturbations make it easier to deploy adversary examples in the real world. Adversaries do not require to change the perturbation when the input sample changes.

- Perturbation Limitation

¹<https://aws.amazon.com/machine-learning>

²<https://cloud.google.com/products/machine-learning>

³<https://bigml.com>

⁴<https://www.clarifai.com>

⁵<https://azure.microsoft.com/en-us/services/machine-learning>

⁶<https://console.bluemix.net/catalog/services/machine-learning>

⁷<https://www.faceplusplus.com>

- *Optimized Perturbation* sets perturbation as the goal of the optimization problem. These methods aim to minimize the perturbation so that humans cannot recognize the perturbation.
- *Constraint Perturbation* sets perturbation as the constraint of the optimization problem. These methods only require the perturbation to be small enough.

- **Perturbation Measurement**

- ℓ_p measures the magnitude of perturbation by p -norm distance:

$$\|x\|_p = \left(\sum_{i=1}^n \|x_i\|^p \right)^{\frac{1}{p}}. \quad (3)$$

$\ell_0, \ell_2, \ell_\infty$ are three commonly used ℓ_p metrics. ℓ_0 counts the number of pixels changed in the adversarial examples; ℓ_2 measures the Euclidean distance between the adversarial example and the original sample; ℓ_∞ denotes the maximum change for all pixels in adversarial examples.

- *Psychometric perceptual adversarial similarity score (PASS)* is a new metric introduced in [59], consistent with human perception.

C. Benchmark

Adversaries show the performance of their adversarial attacks based on different datasets and victim models. This inconsistency brings obstacles to evaluate the adversarial attacks and measure the robustness of deep learning models. Large and high-quality datasets, complex and high-performance deep learning models usually make adversaries/defenders hard to attack/defend. The diversity of datasets and victim models also makes researchers hard to tell whether the existence of adversarial examples is due to datasets or models. We will discuss this problem in Section VII-C.

- **Datasets**

MNIST, *CIFAR-10*, and *ImageNet* are the three most widely used image classification datasets to evaluate adversarial attacks. Because *MNIST* and *CIFAR-10* are proved easy to attack and defend due to its simplicity and small size, *ImageNet* is the best dataset to evaluate adversarial attacks so far. A well-designed dataset is required to evaluate adversarial attacks.

- **Victim Models**

Adversaries usually attack several well-known deep learning models, such as *LeNet*, *VGG*, *AlexNet*, *GoogLeNet*, *CaffeNet*, and *ResNet*.

In the following sections, we will investigate recent studies on adversarial examples according to this taxonomy.

IV. METHODS FOR GENERATING ADVERSARIAL EXAMPLES

In this section, we illustrate several representative approaches for generating adversarial examples. Although many of these approaches are defeated by a countermeasure in later studies, we present these methods to show how the adversarial attacks improved and to what extent state-of-the-art adversarial attacks can achieve. The existence of these methods also

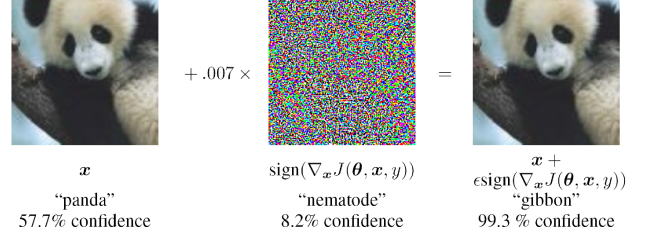


Figure 1: An adversarial image generated by *Fast Gradient Sign Method* [60]: left: a clean image of a panda; middle: the perturbation; right: one sample adversarial image, classified as a gibbon.

requires investigation, which may improve the robustness of deep neural networks.

Table II summaries the methods for generating adversarial examples in this section based on the proposed taxonomy.

A. L-BFGS Attack

Szegedy *et al.* first introduced adversarial examples against deep neural networks in 2014 [8]. They generated adversarial examples using an L-BFGS method to solve the general targeted problem:

$$\begin{aligned} \min_{x'} \quad & c\|\eta\| + J_\theta(x', l') \\ \text{s.t.} \quad & x' \in [0, 1]. \end{aligned} \quad (4)$$

To find a suitable constant c , *L-BFGS Attack* calculated approximate values of adversarial examples by line-searching $c > 0$. The authors showed that the generated adversarial examples could also be generalized to different models and different training datasets. They suggested that adversarial examples are never/rarely seen examples in the test datasets.

L-BFGS Attack was also used in [71], which implemented a binary search to find the optimal c .

B. Fast Gradient Sign Method (FGSM)

L-BFGS Attack used an expensive linear search method to find the optimal value, which was time-consuming and impractical. Goodfellow *et al.* proposed a fast method called *Fast Gradient Sign Method* to generate adversarial examples [60]. They only performed one step gradient update along the direction of the sign of gradient at each pixel. Their perturbation can be expressed as:

$$\eta = \epsilon \text{sign}(\nabla_x J_\theta(x, l)), \quad (5)$$

where ϵ is the magnitude of the perturbation. The generated adversarial example x' is calculated as: $x' = x + \eta$. This perturbation can be computed by using back-propagation. Figure 1 shows an adversarial example on ImageNet.

They claimed that the linear part of the high dimensional deep neural network could not resist adversarial examples, although the linear behavior speeded up training. Regularization approaches are used in deep neural networks such as dropout. Pre-training could not improve the robustness of networks.

[59] proposed a new method, called *Fast Gradient Value* method, in which they replaced the sign of the gradient with the raw gradient: $\eta = \nabla_x J(\theta, x, l)$. *Fast Gradient Value* method has no constraints on each pixel and can generate images with a larger local difference.

Table II: Taxonomy of Adversarial Examples

Threat Model	Adversarial Falsification	False Negative	[8], [9], [59]–[73]
		False Positive	[74]
	Adversary’s Knowledge	White-Box	[8], [9], [59]–[63], [65], [67], [69]–[74]
		Black-Box	[64], [66], [68], [73]
	Adversarial Specificity	Targeted	[8], [59], [61], [63], [64], [66], [67], [69]–[73]
		Non-Targeted	[9], [60], [62], [64]–[66], [68], [69], [73], [74]
Attack Frequency	One-time	[59], [60]	
	Iterative	[8], [9], [61]–[74]	
Perturbation	Perturbation Scope	Individual	[8], [9], [59]–[64], [66]–[71], [73], [74]
		Universal	[65]
	Perturbation Limitation	Optimized	[8], [59], [61]–[65], [68], [70], [71]
		Constraint	[59], [66], [67], [69], [73]
	None	[9], [60], [72], [74]	
	Perturbation Measurement	Element-wise	[9], [60], [72]
		$\ell_p(p \geq 0)$	ℓ_0 : [63], [66], ℓ_1 : [70], ℓ_2 : [8], [61]–[65], [67]–[69], [71], [73], ℓ_∞ : [62], [63], [65], [70], [73],
		PASS	[59]
None		[74]	
Benchmark	Datasets	MNIST	[8], [59]–[63], [68], [70], [71], [74]
		CIFAR-10	[62]–[64], [66]
		ImageNet	[8], [9], [59], [60], [62]–[65], [67], [69], [71]–[74]
		Others	YoutubeDataset: [8] LSUN, SNLI: [68]
	Victim Models	LeNet	[59], [61], [62], [68], [74]
		VGG	[65]–[67], [69]
		AlexNet	[67], [74]
		QuocNet	[8]
		GoogLeNet	[9], [59], [60], [62]–[65], [67]–[69], [72], [73]
		CaffeNet	[62], [65], [67]
		ResNet	[59], [65], [69], [73]
LSTM		[68]	

According to [72], one-step attack is easy to transfer but also easy to defend (see Section VII-A). [73] applied momentum to *FGSM* to generate adversarial examples more iteratively. The gradients were calculated by:

$$\mathbf{g}_{t+1} = \mu \mathbf{g}_t + \frac{\nabla_x J_\theta(x'_t, l)}{\|\nabla_x J_\theta(x'_t, l)\|}, \quad (6)$$

then the adversarial example is derived by $x'_{t+1} = x'_t + \epsilon \text{sign} \mathbf{g}_{t+1}$. The authors increased the effectiveness of attack by introducing momentum and improved the transferability by applying the one-step attack and the ensembling method.

[72] extended *FGSM* to a targeted attack by maximizing the probability of the target class:

$$x' = x - \epsilon \text{sign}(\nabla_x J(\theta, x, l')). \quad (7)$$

The authors refer to this attack as *One-step Target Class Method (OTCM)*.

[75] found that *FGSM* with adversarial training is more robust to white-box attacks than to black-box attacks due to *gradient masking*. They proposed a new attack, *RAND-FGSM*, which added random when updating the adversarial examples to defeat adversarial training:

$$\begin{aligned} x_{tmp} &= x + \alpha \cdot \text{sign}(\mathcal{N}(\mathbf{0}^d, \mathbf{I}^d)), \\ x' &= x_{tmp} + (\epsilon - \alpha) \cdot \text{sign}(\nabla_{x_{tmp}} J(x_{tmp}, l)), \end{aligned} \quad (8)$$

where α, ϵ are the parameters, $\alpha < \epsilon$.

C. Basic Iterative Method (BIM) and Iterative Least-Likely Class Method (ILLC)

Previous methods assume adversarial data can be directly fed into deep neural networks. However, in many applications,

people can only pass data through devices (*e.g.*, cameras, sensors). Kurakin *et al.* applied adversarial examples to the physical world [9]. They extended *Fast Gradient Sign* method by running a finer optimization (smaller change) for multiple iterations. In each iteration, they clipped pixel values to avoid large change on each pixel:

$$\text{Clip}_{x,\xi}\{x'\} = \min\{255, x + \xi, \max\{0, x - \epsilon, x'\}\}, \quad (9)$$

where $\text{Clip}_{x,\xi}\{x'\}$ limits the change of the generated adversarial image in each iteration. The adversarial examples were generated in multiple iterations:

$$\begin{aligned} x_0 &= x, \\ x_{n+1} &= \text{Clip}_{x,\xi}\{x_n + \epsilon \text{sign}(\nabla_x J(x_n, y))\}. \end{aligned} \quad (10)$$

The authors referred to this method as *Basic Iterative* method.

To further attack a specific class, they chose the least-likely class of the prediction and tried to maximize the cross-entropy loss. This method is referred to as *Iterative Least-Likely Class* method:

$$\begin{aligned} x_0 &= x, \\ y_{LL} &= \arg \min_y \{p(y|x)\}, \\ x_{n+1} &= \text{Clip}_{x,\epsilon}\{x_n - \epsilon \text{sign}(\nabla_x J(x_n, y_{LL}))\}. \end{aligned} \quad (11)$$

They successfully fooled the neural network with a crafted image taken from a cellphone camera. They also found that *Fast Gradient Sign* method is robust to phototransformation, while iterative methods cannot resist phototransformation.

D. Jacobian-based Saliency Map Attack (JSMA)

Papernot *et al.* designed an efficient saliency adversarial map, called *Jacobian-based Saliency Map Attack* [61]. They first computed Jacobian matrix of given sample x , which is given by:

$$J_F(x) = \frac{\partial F(x)}{\partial x} = \left[\frac{\partial F_j(x)}{\partial x_i} \right]_{i \times j}. \quad (12)$$

According to [63], F denotes the second-to-last layer (logits) in [61]. Carlini and Wagner modify this approach by using the output of the softmax layer as F [63]. In this way, they found the input features of x that made most significant changes to the output. A small perturbation was designed to successfully induce large output variations so that change in a small portion of features could fool the neural network.

Then the authors defined two adversarial saliency maps to select the feature/pixel to be crafted in each iteration. They achieved 97% adversarial success rate by modifying only 4.02% input features per sample. However, this method runs very slow due to its significant computational cost.

E. DeepFool

Moosavi-Dezfooli *et al.* proposed *DeepFool* to find the closest distance from the original input to the decision boundary of adversarial examples [62]. To overcome the non-linearity in high dimension, they performed an iterative attack with a linear approximation. Starting from an affine classifier, they found that the minimal perturbation of an affine classifier is the distance to the separating affine hyperplane $\mathcal{F} = \{x : w^T x + b = 0\}$. The perturbation of an affine classifier f can be $\eta^*(x) = -\frac{f(x)}{\|w\|^2} w$.

If f is a binary differentiable classifier, they used an iterative method to approximate the perturbation by considering f is linearized around x_i at each iteration. The minimal perturbation is computed as:

$$\begin{aligned} \arg \min_{\eta_i} \quad & \|\eta_i\|_2 \\ \text{s.t.} \quad & f(x_i) + \nabla f(x_i)^T \eta_i = 0. \end{aligned} \quad (13)$$

This result can also be extended to the multi-class classifier by finding the closest hyperplanes. It can also be extended to a more general ℓ_p norm, $p \in [0, \infty)$. *DeepFool* provided less perturbation compared to *FGSM* and *JSMA* did. Compared to *JSMA*, *DeepFool* also reduced the intensity of perturbation instead of the number of selected features.

F. CPPN EA Fool

Nguyen *et al.* discovered a new type of attack, compositional pattern-producing network-encoded EA (CPPN EA), where adversarial examples are classified by deep neural networks with high confidence (99%), which is unrecognizable to human [74]. We categorize this kind of attack as a *False positive* attack. Figure 2 illustrates false-positive adversarial examples.

They used evolutionary algorithms (EAs) algorithm to produce the adversarial examples. To solve multi-class classification problem using EA algorithms, they applied multi-dimensional archive of phenotypic elites MAP-Elites [76].

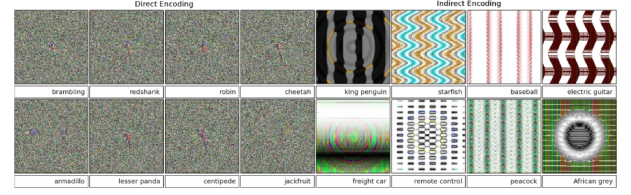


Figure 2: Unrecognizable examples to humans, but deep neural networks classify them to a class with high certainty ($\geq 99.6\%$) [74]

The authors first encoded images with two different methods: direct encoding (grayscale or HSV value) and indirect encoding (compositional pattern-producing network). Then in each iteration, MAP-Elites, like general EA algorithm, chose a random organism, mutated them randomly, and replaced with the current ones if the new ones have higher *fitness* (high certainty for a class of a neural network). In this way, MAP-Elites can find the best individual for each class. As they claimed, for many adversarial images, CPPN could locate the critical features to change outputs of deep neural networks just like JSMA did. Many images from same evolutionary are found similar on closely related categories. More interestingly, CPPN EA fooling images are accepted by an art contest with 35.5% acceptance rate.

G. C&W's Attack

Carlini and Wagner launched a targeted attack to defeat *Defensive distillation* (Section VI-A) [63]. According to their further study [77], [78], *C&W's Attack* is effective for most of existing adversarial detecting defenses. The authors made several modifications in Equation 2.

They first defined a new objective function g , so that:

$$\begin{aligned} \min_{\eta} \quad & \|\eta\|_p + c \cdot g(x + \eta) \\ \text{s.t.} \quad & x + \eta \in [0, 1]^n, \end{aligned} \quad (14)$$

where $g(x') \geq 0$ if and only if $f(x') = l'$. In this way, the distance and the penalty term can be better optimized. The authors listed seven objective function candidates g . One of the effective functions evaluated by their experiments can be:

$$g(x') = \max_{i \neq l'} (\max(Z(x')_i) - Z(x')_{l'} - \kappa), \quad (15)$$

where Z denotes the Softmax function, κ is a constant to control the confidence (κ is set to 0 in [63]).

Second, instead of using box-constrained L-BFGS to find minimal perturbation in *L-BFGS Attack* method, the authors introduced a new variant w to avoid the box constraint, where w satisfies $\eta = \frac{1}{2}(\tanh(w) + 1) - x$. General optimizers in deep learning like Adam and SGD were used to generate adversarial examples and performed 20 iterations of such generation to find an optimal c by binary searching. However, they found that if the gradients of $\|\eta\|_p$ and $g(x + \eta)$ are not in the same scale, it is hard to find a suitable constant c in all of the iterations of the gradient search and get the optimal result. Due to this reason, two of their proposed functions did not find optimal solutions for adversarial examples.

Third, three distance measurements of perturbation were discussed in the paper: ℓ_0 , ℓ_2 , and ℓ_∞ . The authors provided

three kinds of attacks based on the distance metrics: ℓ_0 attack, ℓ_2 attack, and ℓ_∞ attack.

ℓ_2 attack can be described by:

$$\min_w \left\| \frac{1}{2}(\tanh(w) + 1) \right\|_2 + c \cdot g\left(\frac{1}{2}\tanh(w) + 1\right). \quad (16)$$

The authors showed that the distillation network could not help defend ℓ_2 attack.

ℓ_0 attack was conducted iteratively since ℓ_0 is not differentiable. In each iteration, a few pixels are considered trivial for generating adversarial examples and removed. The importance of pixels is determined by the gradient of ℓ_2 distance. The iteration stops if the remaining pixels can not generate an adversarial example.

ℓ_∞ attack was also an iterative attack, which replaced the ℓ_2 term with a new penalty in each iteration:

$$\min \quad c \cdot g(x + \eta) + \sum_i [(\eta_i - \tau)^+]. \quad (17)$$

For each iteration, they reduced τ by a factor of 0.9, if all $\eta_i < \tau$. ℓ_∞ attack considered τ as an estimation of ℓ_∞ .

H. Zeroth Order Optimization (ZOO)

Different from gradient-based adversarial generating approaches, Chen *et al.* proposed a *Zeroth Order Optimization (ZOO)* based attack [64]. Since this attack does not require gradients, it can be directly deployed in a black-box attack without model transferring. Inspired by [63], the authors modified $g(\cdot)$ in [63] as a new hinge-like loss function:

$$g(x') = \max(\max_{i \neq i'} (\log[f(x)]_i) - \log[f(x)]_{i'}, -\kappa), \quad (18)$$

and used symmetric difference quotient to estimate the gradient and Hessian:

$$\begin{aligned} \frac{\partial f(x)}{\partial x_i} &\approx \frac{f(x + he_i) - f(x - he_i)}{2h}, \\ \frac{\partial^2 f(x)}{\partial x_i^2} &\approx \frac{f(x + he_i) - 2f(x) + f(x - he_i)}{h^2}, \end{aligned} \quad (19)$$

where e_i denotes the standard basis vector with the i th component as 1, h is a small constant.

Through employing the gradient estimation of gradient and Hessian, *ZOO* does not need the access to the victim deep learning models. However, it requires expensive computation to query and estimate the gradients. The authors proposed ADAM like algorithms, *ZOO-ADAM*, to randomly select a variable and update adversarial examples. Experiments showed that *ZOO* achieved the comparable performance as *C&W's Attack*.

I. Universal Perturbation

Leveraging their previous method on *DeepFool*, Moosavi-Dezfooli *et al.* developed a universal adversarial attack [65]. The problem they formulated is to find a universal perturbation vector satisfying

$$\begin{aligned} \|\eta\|_p &\leq \epsilon, \\ \mathcal{P}(x' \neq f(x)) &\geq 1 - \delta. \end{aligned} \quad (20)$$

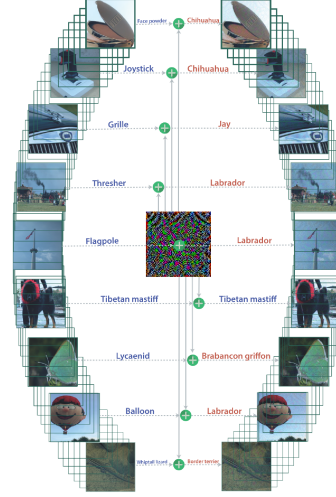


Figure 3: A universal adversarial example fools the neural network on images. Left images: original labeled natural images; center image: universal perturbation; right images: perturbed images with wrong labels. [65]

ϵ limits the size of universal perturbation, and δ controls the failure rate of all the adversarial samples.

For each iteration, they use *DeepFool* method to get a minimal sample perturbation against each input data and update the perturbation to the total perturbation η . This loop will not stop until most data samples are fooled ($\mathcal{P} < 1 - \delta$). From experiments, the universal perturbation can be generated by using a small part of data samples instead of the entire dataset. Figure 3 illustrates a universal adversarial example can fool a group of images. The universal perturbations were shown to be generalized well across popular deep learning architectures (e.g., VGG, CaffeNet, GoogLeNet, ResNet).

J. One Pixel Attack

To avoid the problem of measurement of perceptiveness, Su *et al.* generated adversarial examples by only modifying one pixel [66]. The optimization problem becomes:

$$\begin{aligned} \min_{x'} \quad & J(f(x'), l') \\ \text{s.t.} \quad & \|\eta\|_0 \leq \epsilon_0, \end{aligned} \quad (21)$$

where $\epsilon_0 = 1$ for modifying only one pixel. The new constraint made it hard to optimize the problem.

Su *et al.* applied differential evolution (DE), one of the evolutionary algorithms, to find the optimal solution. DE does not require the gradients of the neural networks and can be used in non-differential objective functions. They evaluated the proposed method on the *CIFAR-10* dataset using three neural networks: All convolution network (AllConv) [79], Network in Network (NiN) [80], and VGG16. Their results showed that 70.97% of images successfully fooled deep neural networks with at least one target class with confidence 97.47% on average.

K. Feature Adversary

Sabour *et al.* performed a targeted attack by minimizing the distance of the representation of internal neural network

layers instead of the output layer [67]. We refer to this attack as *Feature Adversary*. The problem can be described by:

$$\begin{aligned} \min_{x'} \quad & \|\phi_k(x) - \phi_k(x')\| \\ \text{s.t.} \quad & \|x - x'\|_\infty < \delta, \end{aligned} \quad (22)$$

where ϕ_k denotes a mapping from image input to the output of the k th layer. Instead of finding a minimal perturbation, δ is used as a constraint of perturbation. They claimed that a small fixed value δ is good enough for human perception. Similar to [8], they used L-BFGS-B to solve the optimization problem. The adversarial images are more natural and closer to the targeted images in the internal layers.

L. Hot/Cold

Rozsa *et al.* proposed a *Hot/Cold* method to find multiple adversarial examples for every single image input [59]. They thought small translations and rotations should be allowed as long as they were *imperceptible*.

They defined a new metric, Psychometric Perceptual Adversarial Similarity Score (PASS), to measure the noticeable similarity to humans. *Hot/Cold* neglected the unnoticeable difference based on pixels and replaced widely used ℓ_p distance with PASS. PASS includes two stages: 1) aligning the modified image with the original image; 2) measuring the similarity between the aligned image and the original one.

Let $\phi(x', x)$ be a homography transform from the adversarial example x' to the original example x . $\mathcal{H}$ is the homography matrix, with size 3×3 . $\mathcal{H}$ is solved by maximizing the *enhanced correlation coefficient (ECC)* [81] between x' and x . The optimization function is:

$$\arg \min_{\mathcal{H}} \left\| \frac{\bar{x}}{\|\bar{x}\|} - \frac{\overline{\phi(x', x)}}{\|\overline{\phi(x', x)}\|} \right\|, \quad (23)$$

where $\bar{\cdot}$ denotes the normalization of an image.

Structural SIMilarity (SSIM) index [82] was adopted to measure the just noticeable difference of images. [59] leveraged SSIM and defined a new measurement, regional SSIM index (RSSIM) as:

$$RSSIM(x_{i,j}, x'_{i,j}) = L(x_{i,j}, x'_{i,j})^\alpha C(x_{i,j}, x'_{i,j})^\beta S(x_{i,j}, x'_{i,j})^\gamma,$$

where α, β, γ are weights of importance for luminance ($L(\cdot, \cdot)$), contrast ($C(\cdot, \cdot)$), and structure ($S(\cdot, \cdot)$). The SSIM can be calculated by averaging RSSIM:

$$SSIM(x_{i,j}, x'_{i,j}) = \frac{1}{n \times m} \sum_{i,j} RSSIM(x_{i,j}, x'_{i,j}).$$

PASS is defined by combination of the alignment and the similarity measurement:

$$PASS(x', x) = SSIM(\phi^*(x', x), x). \quad (24)$$

The adversarial problem with the new distance is described as:

$$\begin{aligned} \min \quad & D(x, x') \\ \text{s.t.} \quad & f(x') = y', \\ & PASS(x, x') \geq \gamma. \end{aligned} \quad (25)$$

$D(x, x')$ denotes a measure of distance (e.g., $1 - PASS(x, x')$ or $\|x - x'\|_p$).

To generate a diverse set of adversarial examples, the authors defined the targeted label l' as *hot* class, and the original label l as *cold* class. In each iteration, they moved toward a target (hot) class while moving away from the original (cold) class. Their results showed that generated adversarial examples are comparable to *FGSM*, and with more diversity.

M. Natural GAN

Zhao *et al.* utilized Generative Adversarial Networks (GANs) as part of their approach to generate adversarial examples of images and texts [68], which made adversarial examples more natural to human. We name this approach *Natural GAN*. The authors first trained a WGAN model on the dataset, where the generator $\mathcal{G}$ maps random noise to the input domain. They also trained an “inverter” $\mathcal{I}$ to map input data to dense inner representations. Hence, the adversarial noise was generated by minimizing the distance of the inner representations like “Feature Adversary.” The adversarial examples were generated using the generator: $x' = \mathcal{G}(z')$:

$$\begin{aligned} \min_z \quad & \|z - \mathcal{I}(x)\| \\ \text{s.t.} \quad & f(\mathcal{G}(z)) \neq f(x). \end{aligned} \quad (26)$$

Both the generator $\mathcal{G}$ and the inverter $\mathcal{I}$ were built to make adversarial examples natural. *Natural GAN* was a general framework for many deep learning fields. [68] applied *Natural GAN* to image classification, textual entailment, and machine translation. Since *Natural GAN* does not require gradients of original neural networks, it can also be applied to *Black-box Attack*.

N. Model-based Ensembling Attack

Liu *et al.* conducted a study of transferability (Section VII-A) over deep neural networks on ImageNet and proposed a *Model-based Ensembling Attack* for targeted adversarial examples [69]. The authors argued that compared to non-targeted adversarial examples, targeted adversarial examples are much harder to transfer over deep models. Using *Model-based Ensembling Attack*, they can generate transferable adversarial examples to attack a black-box model.

The authors generated adversarial examples on multiple deep neural networks with full knowledge and tested them on a black-box model. *Model-based Ensembling Attack* was derived by the following optimization problem:

$$\arg \min_{x'} -\log \left(\sum_{i=1}^k \alpha_i J_i(x', l') \right) + \lambda \|x' - x\|, \quad (27)$$

where k is the number of deep neural networks in the generation, f_i is the function of each network, and α_i is the ensemble weight ($\sum_{i=1}^k \alpha_i = 1$). The results showed that *Model-based Ensembling Attack* could generate transferable targeted adversarial images, which enhanced the power of adversarial examples for black-box attacks. They also proved that this method performs better in generating non-targeted adversarial examples than previous methods. The authors

successfully conducted a black-box attack against Clarifai.com using *Model-based Ensembling Attack*.

O. Ground-Truth Attack

Formal verification techniques aim to evaluate the robustness of a neural network even against zero-day attacks (Section VI-F). Carlini *et al.* constructed a ground-truth attack, which provided adversarial examples with minimal perturbation (ℓ_1, ℓ_∞) [70]. *Network Verification* always checks whether an adversarial example violates a property of a deep neural network and whether there exists an example changes the label within a certain distance. *Ground-Truth Attack* conducted a binary search and found such an adversarial example with the smallest perturbation by invoking Reluplex [83] iteratively. The initial adversarial example is found using C&W's Attack [63] to improve the performance.

V. APPLICATIONS FOR ADVERSARIAL EXAMPLES

We have investigated adversarial examples for image classification task. In this section, we review adversarial examples against the other tasks. We mainly focus on three questions: What scenarios are adversarial examples applied in new tasks? How to generate adversarial examples in new tasks? Whether to propose a new method or to translate the problem into the image classification task and solve it by the aforementioned methods? Table III summarizes the applications for adversarial examples in this section.

A. Reinforcement Learning

Deep neural networks have been used in reinforcement learning by training policies on raw input (*e.g.*, images). [84], [85] generated adversarial examples on deep reinforcement learning policies. Since the inherent intensive computation of reinforcement learning, both of them performed fast *One-time attack*.

Huang *et al.* applied *FGSM* to attack deep reinforcement learning networks and algorithms [84]: deep Q network (DQN), trust region policy optimization (TRPO), and asynchronous advantage actor-critic (A3C) [97]. Similarly to [60], they added small perturbations on the input of policy by calculating the gradient of the cross-entropy loss function: $\nabla_x J(\theta, x, \ell)$. Since DQN does not have stochastic policy input, softmax of Q-values is considered to calculate the loss function. They evaluated adversarial examples on four Atari 2600 games with three norm constraints $\ell_1, \ell_2, \ell_\infty$. They found *Huang's Attack* with ℓ_1 norm conducted a successful attack on both *White-box attack* and *Black-box attack* (no access to the training algorithms, parameters, and hyper-parameters).

[85] used *FGSM* to attack A3C algorithm and Atari Pong task. [85] found that injecting perturbations in a fraction of frames is sufficient.

B. Generative Modeling

Kos *et al.* [86] and Tabacof *et al.* [87] proposed adversarial examples for generative models. An adversary for autoencoder can inject perturbations into the input of encoder and generate

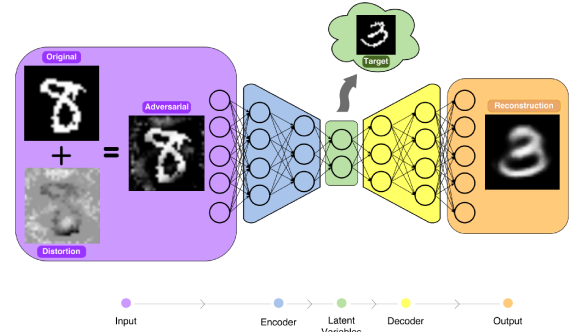


Figure 4: Adversarial attacks for autoencoders [87]. Perturbations are added to the input the encoder. After encoding and decoding, the decoder will output an adversarial image presenting an incorrect class

a targeted class after decoding. Figure 4 depicts a targeted adversarial example for an autoencoder. Adding perturbations on the input image of the encoder can misguide the autoencoder by making decoder to generating a targeted adversarial output image.

Kos *et al.* described a scenario to apply adversarial examples against autoencoder. Autoencoders can be used to compress data by an encoder and decompress by a decoder. For example, Toderici *et al.* use RNN-based AutoEncoder to compress image [98]. Ledig *et al.* used GAN to super-resolve images [99]. Adversaries can leverage autoencoder to reconstruct an adversarial image by adding perturbation to the input of the encoder.

Tabacof *et al.* used *Feature Adversary Attack* against AE and VAE. The adversarial examples were formulated as follows [87]:

$$\begin{aligned} \min_{\eta} \quad & D(z_{x'}, z_x) + c\|\eta\| \\ \text{s.t.} \quad & x' \in [L, U] \\ & z_{x'} = \text{Encoder}(x') \\ & z_x = \text{Encoder}(x), \end{aligned} \quad (28)$$

where $D(\cdot)$ is the distance between latent encoding representation z_x and $z_{x'}$. Tabacof *et al.* chose KL-divergence to measure $D(\cdot)$ in [87]. They tested their attacks on the MNIST and SVHN dataset and found that generating adversarial examples for autoencoder is much harder than for classifiers. VAE is even slightly more robust than deterministic autoencoder.

Kos *et al.* extended Tabacof *et al.*'s work by designing another two kinds of distances. Hence, the adversarial examples can be generated by optimizing:

$$\min_{\eta} \quad c\|\eta\| + J(x', l'). \quad (29)$$

The loss function J can be cross-entropy (refer to "Classifier Attack" in [86]), VAE loss function (" L_{VAE} Attck"), and distance between the original latent vector z and modified encoded vector x' ("Latent Attack", similar to Tabacof *et al.*'s work [87]). They tested VAE and VAE-GAN [100] on the MNIST, SVHN, and CelebA datasets. In their experimental results, "Latent Attack" achieved the best result.

C. Face Recognition

Deep neural network based Face Recognition System (FRS) and Face Detection System have been widely deployed in

Table III: Summary of Applications for Adversarial Examples

Applications	Representative Study	Method	Adversarial Falsification	Adversary's Knowledge	Adversarial Specificity	Perturbation Scope	Perturbation Limitation	Attack Frequency	Perturbation Measurement	Dataset	Architecture
Reinforcement Learning	[84]	FGSM	N/A	White-box & Black-box	Non-Targeted	Individual	N/A	One-time	$\ell_1, \ell_2, \ell_\infty$	Atari	DQN, TRPO, A3C
	[85]	FGSM	N/A	White-box	Non-Targeted	Individual	N/A	One-time	N/A	Atari Pong	A3C
Generative Modeling	[86]	Feature Adversary, C&W	N/A	White-box	Targeted	Individual	Optimized	Iterative	ℓ_2	MNIST, SVHN, CelebA	VAE, VAE-GAN
	[87]	Feature Adversary	N/A	White-box	Targeted	Individual	Optimized	Iterative	ℓ_2	MNIST, SVHN	VAE, AE
Face Recognition	[58]	Impersonation & Dodging Attack	False negative	white-box & black-box	Targeted & Non-Targeted	Universal	Optimized	Iterative	Total Variation	LFW,	VGGFace
Object Detection	[11]	DAG	False negative & False positive	White-box & Black-box	Non-Targeted	Individual	N/A	Iterative	N/A	VOC2007, VOC2012	Faster-RCNN
Semantic Segmentation	[11]	DAG	False negative & False positive	White-box & Black-box	Non-Targeted	Individual	N/A	Iterative	N/A	DeepLab	FCN
	[88]	ILLC	False negative	White-box	Targeted	Individual	N/A	Iterative	ℓ_∞	Cityscapes	FCN
	[89]	ILLC	False negative	White-box	Targeted	Universal	N/A	Iterative	N/A	Cityscapes	FCN
Reading Comprehension	[90]	AddSent, AddAny	N/A	Black-box	Non-Targeted	Individual	N/A	One-time & Iterative	N/A	SQuAD	BiDAF, Match-LSTM, and twelve other published models
	[91]	Reinforcement Learning	False negative	White-box	Non-Targeted	Individual	Optimized	Iterative	ℓ_0	TripAdvisor Dataset	Bi-LSTM, memory network
Malware Detection	[92]	JSMA	False negative	White-box	Targeted	Individual	Optimized	Iterative	ℓ_2	DREBIN	2-layer FC
	[93]	Reinforcement Learning	False negative	Black-box	Targeted	Individual	N/A	Iterative	N/A	N/A	Gradient Boosted Decision Tree
	[94]	GAN	False negative	Black-box	Targeted	Individual	N/A	Iterative	N/A	malwr	Multi-layer Perceptron
	[95]	GAN	False negative	Black-box	Targeted	Individual	N/A	Iterative	N/A	Alexa Top 1M	Random Forest
	[96]	Generic Programming	False negative	Black-box	Targeted	Individual	N/A	Iterative	N/A	Contagio	Random Forest, SVM



Figure 5: An example of adversarial eyeglass frame against Face Recognition System [58]

commercial products due to their high performance. [58] first provided a design of eyeglass frames to attack a deep neural network based FRS [101], which composes 11 blocks with 38 layers and one *triplet loss* function for feature embedding. Based on the *triplet loss* function, [58] designed a *softmaxloss* function:

$$J(x) = -\log \left(\frac{e^{\langle h_{c_x}, f(x) \rangle}}{\sum_{c=1}^m e^{\langle h_c, f(x) \rangle}} \right), \quad (30)$$

where h_c is a one-hot vector of class c , $\langle \cdot, \cdot \rangle$ denotes inner product. Then they used *L-BFGS Attack* to generate adversarial examples.

In a further step, [58] implemented adversarial eyeglass frames to achieve attacks in the physical world: the perturbations can only be injected into the area of eyeglass frames. They also enhanced the printability of adversarial images on the frame by adding a penalty of non-printability score (NPS) to the optimized objective. Similarly to *Universal Perturbation*, they optimize the perturbation to be applied to a set of face images. They successfully dodged (non-targeted attack) against FRS (over 80 % time) and misguided FRS as a specific face (targeted attack) with a high success rate (depending on the target). Figure 5 illustrates an example of adversarial eyeglass frames.

Leveraging the approach of printability, [10] proposed an attack algorithm, Robust Physical Perturbations (RP_2), to modify a stop sign as a speed limit sign)⁸. They changed the physical road signs by two kinds of attacks: 1) overlaying an adversarial road sign over a physical sign; 2) sticking perturbations on an existing sign. [10] included a non-printability score in the optimization objective to improve the printability.

D. Object Detection

The object detection task is to find the proposal of an object (bounding box), which can be viewed as an image classification task for every possible proposal. [11] proposed a universal algorithm called *Dense Adversary Generation (DAG)* to generate adversarial examples for both object detection and semantic segmentation. The authors aimed at making the prediction (detection/segmentation) incorrect (non-targeted). Figure 6 illustrates an adversarial example for the object detection task.

[11] defined $T = t_1, t_2, \dots, t_N$ as the recognition targets. For image classification, the classifier only needs one target – entire image ($N = 1$); For semantic segmentation, targets consist of all pixels ($N = \#ofpixels$); For object detection, targets consist of all possible proposals ($N = (\#ofpixels)^2$). Then the objective function sums up the loss from all targets.

⁸This method was shown not effective for standard detectors (YOLO and Faster RCNN) in [102].

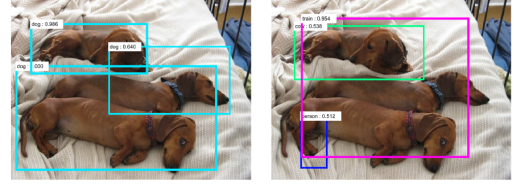


Figure 6: An adversarial example for object detection task [11]. Left: object detection on a clean image. Right: object detection on an adversarial image.

Instead of optimizing the loss from all targets, the authors performed an iterative optimization and only updated the loss for the targets correctly predicted in the previous iteration. The final perturbation sums up normalized perturbations in all iterations. To deal with a large number of targets for objective detection problem, the authors used regional proposal network (RPN) [4] to generate possible targets, which greatly decreases the computation for targets in object detection. DAG also showed the capability of generating images which are unrecognizable to human but deep learning could predict (*false positives*).

E. Semantic Segmentation

Image segmentation task can be viewed as an image classification task for every pixel. Since each perturbation is responsible for at least one pixel segmentation, this makes the space of perturbations for segmentation much smaller than that for image classification [103]. [11], [88], [103] generated adversarial examples against the semantic image segmentation task. However, their attacks are proposed under different scenarios. As we just discussed, [11] performed a non-targeted segmentation. [88], [103] both performed a targeted segmentation and tried to removed a certain class by making deep learning model to misguide it as background classes.

[88] generated adversarial examples by assigning pixels with the adversarial class that their nearest neighbor belongs to. The success rate was measured by the percentage of pixels of the chosen class to be changed or of the rest classes to be preserved.

[103] presented a method to generate universal adversarial perturbations against semantic image segmentation task. They assigned the primary objective of adversarial examples and hid the objects (e.g., pedestrians) while keeping the rest segmentation unchanged. Metzen *et al.* defined background classes and targeted classes (not targeted adversarial classes). Targeted classes are classes to be removed. Similar to [88], pixels which belong to the targeted classes would be assigned to their nearest background classes:

$$\begin{aligned} l_{ij}^{target} &= l_{i'j'}^{pred} \quad \forall (i, j) \in I_{targeted}, \\ l_{ij}^{target} &= l_{ij}^{pred} \quad \forall (i, j) \in I_{background}, \\ (i', j') &= \arg \min_{(i', j') \in I_{background}} \|i' - i\| + \|j' - j\|, \end{aligned} \quad (31)$$

where $I_{targeted} = (i, j) | f(x_{ij}) = l^*$ denotes the area to be removed. Figure 7 illustrates an adversarial example to hide pedestrians. They used *ILLC* attack to solve this problem and also extended *Universal Perturbation* method to get the

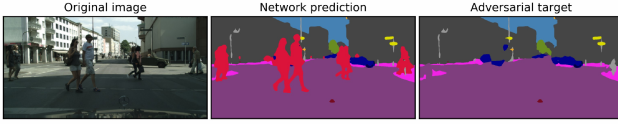


Figure 7: Adversary examples of hiding pedestrians in the semantic segmentation task [103]. Left image: original image; Middle image: the segmentation of the original image predicted by DNN; Right image: the segmentation of the adversarial image predicted by DNN.

universal perturbation. Their results showed the existence of universal perturbation for semantic segmentation task.

F. Natural Language Processing (NLP)

Many tasks in natural language processing can be attacked by adversarial examples. People usually generate adversarial examples by adding/deleting words in the sentences.

The task of reading comprehension (a.k.a. question answering) is to read paragraphs and answer questions about the paragraphs. To generate adversarial examples that are consistent with the correct answer and do not confuse human, Jia and Liang added distracting (adversarial) sentences to the end of paragraph [90]. They found that models for the reading comprehension task are overstable instead of oversensitivity, which means deep learning models cannot tell the subtle but critical difference in the paragraphs.

They proposed two kinds of methods to generate adversarial examples: 1) adding grammatical sentences similar to the question but not contradictory to the correct answer (AddSent); 2) adding a sentence with arbitrary English words (AddAny). [90] successfully fooled all the models (sixteen models) they tested on Stanford Question Answering Dataset (SQuAD) [104]. The adversarial examples also have the capability of transferability and cannot be improved by adversarial training. However, the adversarial sentences require manpower to fix the errors in the sentences.

[91] aimed to fool a deep learning-based sentiment classifier by removing the minimum subset of words in the given text. Reinforcement learning was used to find an approximate subset, where the reward function was proposed as $\frac{1}{\|D\|}$ when the sentiment label changes, and 0 otherwise. $\|D\|$ denotes the number of removing word set D . The reward function also included a regularizer to make sentence contiguous.

The changes in [90], [91] can easily be recognized by humans. More natural adversarial examples for texture data was proposed by *Natural GAN* [68] (Section IV-M).

G. Malware Detection

Deep learning has been used in static and behavioral-based malware detection due to its capability of detecting zero-day malware [105]–[108]. Recent studies generated adversarial malware samples to evade deep learning-based malware detection [92]–[94], [96].

[92] adapted *JSMa* method to attack Android malware detection model. [96] evaded two PDF malware classifier, PDFrate and Hidost, by modifying PDF. [96] parsed the PDF file and changed its object structure using genetic programming. The adversarial PDF file was then packed with new objects.

Table IV: Summary of Countermeasures for Adversarial Examples

	Defensive Strategies	Representative Studies
Reactive	Adversarial Detecting	[23], [89], [109]–[116]
	Input Reconstruction	[114], [117], [118]
	Network Verification	[83], [119], [120]
Proactive	Network Distillation	[121]
	Adversarial (Re)Training	[24], [25], [60], [72], [75], [122]
	Classifier Robustifying	[123], [124]

[95] used GAN to generate adversarial domain names to evade detection of domain generation algorithms. [94] proposed a GAN based algorithm, MalGAN, to generate malware examples and evade black-box detection. [94] used a substitute detector to simulate the real detector and leveraged the transferability of adversarial examples to attack the real detector. MalGAN was evaluated by 180K programs with API features. However, [94] required the knowledge of features used in the model. [93] used a large number of features (2,350) to cover the required feature space of portable executable (PE) files. The features included PE header metadata, section metadata, import & export table metadata. [93] also defined several modifications to generate malware evading deep learning detection. The solution was trained by reinforcement learning, where the evasion rate is considered as a reward.

VI. COUNTERMEASURES FOR ADVERSARIAL EXAMPLES

Countermeasures for adversarial examples have two types of defense strategies: 1) *reactive*: detect adversarial examples after deep neural networks are built; 2) *proactive*: make deep neural networks more robust before adversaries generate adversarial examples. In this section, we discuss three reactive countermeasures (*Adversarial Detecting*, *Input Reconstruction*, and *Network Verification*) and three proactive countermeasures (*Network Distillation*, *Adversarial (Re)training*, and *Classifier Robustifying*). We will also discuss an ensembling method to prevent adversarial examples. Table IV summarizes the countermeasures.

A. Network Distillation

Papernot *et al.* used network distillation to defend deep neural networks against adversarial examples [121]. Network distillation was originally designed to reduce the size of deep neural networks by transferring knowledge from a large network to a small one [125], [126] (Figure 8). The probability of classes produced by the first DNN is used as inputs to train the second DNN. The probability of classes extracts the knowledge learned from the first DNN. Softmax is usually used to normalize the last layer of DNN and produce the probability of classes. The softmax output of the first DNN, also the input of the next DNN, can be described as:

$$q_i = \frac{\exp(z_i/T)}{\sum_j \exp(z_j/T)}, \quad (32)$$

where T is a *temperature* parameter to control the level of knowledge distillation. In deep neural networks, *temperature* T is set to 1. When T is large, the output of softmax will be vague (when $T \rightarrow \infty$, the probability of all classes $\rightarrow \frac{1}{m}$).

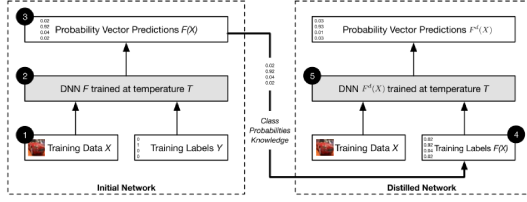


Figure 8: Network distillation of deep neural networks [121]

When T is small, only one class is close to 1 while the rest goes to 0. This schema of network distillation can be duplicated several times and connects several deep neural networks.

In [121], network distillation extracted knowledge from deep neural networks to improve robustness. The authors found that attacks primarily aimed at the sensitivity of networks and then proved that using high-temperature softmax reduced the model sensitivity to small perturbations. *Network Distillation* defense was tested on the MNIST and CIFAR-10 datasets and reduced the success rate of JSMA attack by 0.5% and 5% respectively. “Network Distillation” also improved the generalization of the neural networks.

B. Adversarial (Re)training

Training with adversarial examples is one of the countermeasures to make neural networks more robust. Goodfellow *et al.* [60] and Huang *et al.* [122] included adversarial examples in the training stage. They generated adversarial examples in every step of training and inject them into the training set. [60], [122] showed that adversarial training improved the robustness of deep neural networks. Adversarial training could provide regularization for deep neural networks [60] and improve the precision as well [24].

[60] and [122] were evaluated only on the MNIST dataset. A comprehensive analysis of adversarial training methods on the ImageNet dataset was presented in [72]. They used half adversarial examples and half origin examples in each step of training. From the results, *adversarial training* increased the robustness of neural networks for *one-step attacks* (e.g., FGSM) but would not help under *iterative attacks* (e.g., BIM and ILLC methods). [72] suggested that adversarial training is used for regularization only to avoid overfitting (e.g., the case in [60] with the small MNIST dataset).

[75] found that the adversarial trained models on the MNIST and ImageNet datasets are more robust to white-box adversarial examples than to the transferred examples (black-box).

[25] minimized both the cross-entropy loss and internal representation distance during adversarial training, which can be seen as a defense version of *Feature Adversary*.

To deal with the transferred black-box model, [75] proposed *Ensembling Adversarial Training* method that trained the model with adversarial examples generated from multiple sources: the models being trained and also pre-trained external models.

C. Adversarial Detecting

Many research projects tried to detect adversarial examples in the testing stage [23], [89], [109]–[114], [116].

[23], [89], [111] trained deep neural network-based binary classifiers as detectors to classify the input data as a legitimate (clean) input or an adversarial example. Metzen *et al.* created a detector for adversarial examples as an auxiliary network of the original neural network [89]. The detector is a small and straightforward neural network predicting on binary classification, *i.e.*, the probability of the input being adversarial. SafetyNet [23] extracted the binary threshold of each ReLU layer’s output as the features of the adversarial detector and detects adversarial images by an RBF-SVM classifier. The authors claimed that their method is hard to be defeated by adversaries even when adversaries know the detector because it is difficult for adversaries to find an optimal value, for both adversarial examples and new features of SafetyNet detector. [112] added an outlier class to the original deep learning model. The model detected the adversarial examples by classifying it as an outlier. They found that the measurement of maximum mean discrepancy (MMD) and energy distance (ED) could distinguish the distribution of adversarial datasets and clean datasets.

[110] provided a Bayesian view of detecting adversarial examples. [110] claimed that the uncertainty of adversarial examples is higher than the clean data. Hence, they deployed a Bayesian neural network to estimate the uncertainty of input data and distinguish adversarial examples and clean input data based on uncertainty estimation.

Similarly, [114] used probability divergence (Jensen-Shannon divergence) as one of its detectors. [113] showed that after whitening by Principal Component Analysis (PCA), adversarial examples have different coefficients in low-ranked components.

[118] trained a PixelCNN neural network [127] and found that the distribution of adversarial examples is different from clean data. They calculated p-value based on the rank of PixelCNN and rejected adversarial examples using the p-values. The results showed that this approach could detect FGSM, BIM, DeepFool, and C&W attack.

[115] trained neural networks with “reverse cross-entropy” to better distinguish adversarial examples from clean data in the latent layers and then detected adversarial examples using a method called “Kernel density” in the testing stage. The “reverse cross-entropy” made the deep neural network to predict with high confidence on the true class and uniform distribution on the other classes. In this way, the deep neural network was trained to map the clean input close to a low-dimensional manifold in the layer before softmax. This brought great convenience for further detection of adversarial examples.

[116] leveraged multiple previous images to predict future input and detect adversarial examples, in the task of reinforcement learning.

However, Carlini and Wagner summarized most of these adversarial detecting methods ([89], [109]–[113]) and showed that these methods could not defend against their previous attack *C&W’s Attack* with slight changes of loss function [77], [78].

D. Input Reconstruction

Adversarial examples can be transformed to clean data via reconstruction. After transformation, the adversarial examples will not affect the prediction of deep learning models. Gu and Rigazio proposed a variant of autoencoder network with a penalty, called deep contractive autoencoder, to increase the robustness of neural networks [117]. A denoising autoencoder network is trained to encode adversarial examples to original ones to remove adversarial perturbations. [114] reconstructed the adversarial examples by 1) adding Gaussian noise or 2) encoding them with autoencoder as a plan B in MagNet [114](Section VI-G).

PixelDefend reconstructed the adversarial images back to the training distribution [118] using PixelCNN. *PixelDefend* changed all pixels along each channel to maximize the probability distribution:

$$\begin{aligned} \max_{x'} \quad & \mathcal{P}_t(x') \\ \text{s.t.} \quad & \|x' - x\|_\infty \leq \epsilon_{defend}, \end{aligned} \quad (33)$$

where $\mathcal{P}_t$ denotes the training distribution, ϵ_{defend} controls the new changes on the adversarial examples. *PixelDefend* also leveraged adversarial detecting, so that if an adversarial example is not detected as malicious, no change will be made to the adversarial examples ($\epsilon_{defend} = 0$).

E. Classifier Robustifying

[123], [124] design robust architectures of deep neural networks to prevent adversarial examples.

Due to the uncertainty from adversarial examples, Bradshaw *et al.* leveraged Bayesian classifiers to build more robust neural networks [123]. Gaussian processes (GPs) with RBF kernels were used to provide uncertainty estimation. The proposed neural networks were called *Gaussian Process Hybrid Deep Neural Networks (GPDNNs)*. GPs expressed the latent variables as a Gaussian distribution parameterized by the functions of mean and covariance and encoded them with RBF kernels. [123] showed that GPDNNs achieved comparable performance with general DNNs and more robust to adversarial examples. The authors claimed that GPDNNs “know when they do not know.”

[124] observed that adversarial examples usually went into a small subset of incorrect classes. [124] separated the classes into sub-classes and ensembled the result from all sub-classes by voting to prevent adversarial examples misclassified.

F. Network Verification

Verifying properties of deep neural networks is a promising solution to defend adversarial examples, because it may detect the new unseen attacks. *Network verification* checks the properties of a neural network: whether an input violates or satisfies the property.

Katz *et al.* proposed a verification method for neural networks with ReLU activation function, called Reluplex [83]. They used Satisfiability Modulo Theory (SMT) solver to verify the neural networks. The authors showed that within a small perturbation, there was no existing adversarial example to

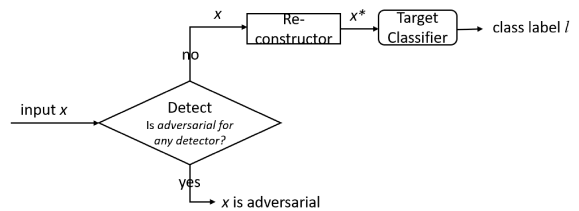


Figure 9: MagNet workflow: one or more detectors first detects if input x is adversarial; If not, reconstruct x to x^* before feeding it to the classifier. (modified from [114])

misclassify the neural networks. They also proved that the problem of network verification is NP-complete. Carlini *et al.* extended their assumption of ReLU function by presenting $\max(x, y) = \text{ReLU}(x - y) + y$ and $\|x\| = \text{ReLU}(2x) - x$ [70]. However, Reluplex runs very slow due to the large computation of verifying the networks and only works for DNNs with several hundred nodes [120]. [120] proposed two potential solutions: 1) prioritizing the order of checking nodes 2) sharing information of verification.

Instead of checking each point individually, Gopinath *et al.* proposed *DeepSafe* to provide safe regions of a deep neural network [119] using Reluplex. They also introduced *targeted robustness* a safe region only regarding a targeted class.

G. Ensembling Defenses

Due to the multi-facet of adversarial examples, multiple defense strategies can be performed together (parallel or sequential) to defend adversarial examples.

Aforementioned *PixelDefend* [118] is composed of an adversarial detector and an “input reconstructor” to establish a defense strategy.

MagNet included one or more detectors and a reconstructor (“reformer” in the paper) as Plan A and Plan B [114]. The detectors are used to find the adversarial examples which are far from the boundary of the manifold. In [114], they first measured the distance between input and encoded input and also the probability divergence (Jensen-Shannon divergence) between softmax output of input and encoded input. The adversarial examples were expected a large distance and probability divergence. To deal with the adversarial examples close to the boundary, MagNet used a reconstructor built by neural network based autoencoders. The reconstructor will map adversarial examples to legitimate examples. Figure 9 illustrates the workflow of the defense of two phases.

After investigating several defensive approaches, [128] showed that the ensemble of those defensive approaches does not make the neural networks strong.

H. Summary

Almost all defenses are shown to be effective only for part of attacks. They tend not to be defensive for some strong (fail to defend) and unseen attacks. Most defenses target adversarial examples in the computer vision task. However, with the development of adversarial examples in other areas, new defenses for these areas, especially for safety-critical systems, are urgently required.

VII. CHALLENGES AND DISCUSSIONS

In this section, we discuss the current challenges and the potential solutions for adversarial examples. Although many methods and theorems have been proposed and developed recently, a lot of fundamental questions need to be well explained and many challenges need to be addressed. The reason for the existence of adversarial examples is an interesting and one of the most fundamental problems for both adversaries and researchers, which exploits the vulnerability of neural networks and help defenders to resist adversarial examples. We will discuss the following questions in this section: Why do adversarial examples transfer? How to stop the transferability? Why are some defenses effective and others not? How to measure the strength of an attack as well as a defense? How to evaluate the robustness of a deep neural network against seen/unseen adversarial examples?

A. Transferability

Transferability is a common property for adversarial examples. Szegedy *et al.* first found that adversarial examples generated against a neural network can fool the same neural networks trained by different datasets. Papernot *et al.* found that adversarial examples generated against a neural network can fool other neural networks with different architectures, even other classifiers trained by different machine learning algorithms [33]. Transferability is critical for Black-Box attacks where the victim deep learning model and the training dataset are not accessible. Attackers can train a substitute neural network model and then generate adversarial examples against substitute model. Then the victim model will be vulnerable to these adversarial examples due to transferability. From a defender's view, if we hinder transferability of adversarial examples, we can defend all white-box attackers who need to access the model and require transferability.

We define the transferability of adversarial examples in three levels from easy to hard: 1) transfer among the same neural network architecture trained with different data; 2) transfer among different neural network architectures trained for the same task; 3) transfer among deep neural networks for different tasks. To our best knowledge, there is no existing solution on the third level yet (for instance, transfer an adversarial image from object detection to semantic segmentation).

Many studies examined transferability to show the ability of adversarial examples [8], [60]. Papernot *et al.* studied the transferability between conventional machine learning techniques (*i.e.*, logistic regression, SVM, decision tree, kNN) and deep neural networks. They found that adversarial examples can be transferred between different parameters, training dataset of a machine learning models and even across different machine learning techniques.

Liu *et al.* investigated transferability of targeted and non-targeted adversarial examples on complex models and large datasets (*e.g.*, the ImageNet dataset) [69]. They found that non-targeted adversarial examples are much more transferable than targeted ones. They observed that the decision boundaries of different models aligned well with each other. Thus they

proposed *Model-Based Ensembling Attack* to create transferable targeted adversarial examples.

Tramèr *et al.* found that the distance to the model's decision boundary is on average larger than the distance between two models' boundaries in the same direction [129]. This may explain the existence of transferability of adversarial examples. Tramèr *et al.* also claimed that transferability might not be an inherent property of deep neural networks by showing a counter-example.

B. The existence of Adversarial Examples

The reason for the existence of adversarial examples is still an open question. Are adversarial examples an inherent property of deep neural networks? Are adversarial examples the "Achilles' heel" of deep neural networks with high performance? Many hypotheses have been proposed to explain the existence.

Data incompleteness One assumption is that adversarial examples are of low probability and low test coverage of corner cases in the testing dataset [8], [130]. From training a PixelCNN, [118] found that the distribution of adversarial examples was different from clean data. Even for a simple Gaussian model, a robust model can be more complicated and requires much more training data than that of a "standard" model [131].

Model capability Adversarial examples are a phenomenon not only for deep neural networks but also for all classifiers [33], [132]. [60] suggested that adversarial examples are the results of models being too linear in high dimensional manifolds. [133] showed that in the linear case, the adversarial examples exist when the decision boundary is close to the manifold of the training data.

Contrary to [60], [132] believed that adversarial examples are due to the "low flexibility" of the classifier for certain tasks. Linearity is not an "obvious explanation" [67]. [71] blamed adversarial examples for the sparse and discontinuous manifold which makes classifier erratic.

No robust model [25] suggested that the decision boundaries of deep neural networks are inherently incorrect, which do not detect semantic objects. [134] showed that if a dataset is generated by a smooth generative model with large latent space, there is no robust classifier to adversarial examples. Similarly, [135] prove that if a model is trained on a sphere dataset and misclassifies a small part of the dataset, then there exist adversarial examples with a small perturbation.

In addition to adversarial examples for image classification task, as discussed in Section V, adversarial examples have been generated in various applications. Many of them deployed utterly different methods. Some applications can use the same method used in image classification task. However, some need to propose a novel method. Current studies on adversarial examples mainly focus on image classification task. No existing paper explains the relationship among different applications and existence of a universal attacking/defending method to be applied to all the applications.

C. Robustness Evaluation

The competition between attacks and defenses for adversarial examples becomes an “arms race”: a defensive method that was proposed to prevent existing attacks was later shown to be vulnerable to some new attacks, and vice versa [70], [112]. Some defenses showed that they could defend a particular attack, but later failed with a slight change of the attack [109], [110]. Hence, the evaluation on the robustness of a deep neural network is necessary. For example, [136] provided an upper bound of robustness for linear classifier and quadratic classifier. The following problems for robustness evaluation of deep neural networks require further exploration.

1) **A methodology for evaluation on the robustness of deep neural networks:** Many deep neural networks are planned to be deployed in safety-critical settings. Defending only existing attacks is not sufficient. Zero-day (new) attacks would be more harmful to deep neural networks. A methodology for evaluating the robustness of deep neural networks is required, especially for zero-day attacks, which helps people understand the confidence of model prediction and how much we can rely on them in the real world. [70], [83], [137], [138] conducted initial studies on the evaluation. Moreover, this problem lies not only in the performance of deep neural network models but also in the confidentiality and privacy.

2) **A benchmark platform for attacks and defenses:** Most attacks and defenses described their methods without publicly available code, not to mention the parameters used in their methods. This brings difficulties for other researchers to reproduce their solutions and provide the corresponding attacks/defenses. For example, Carlini tried his best to “find the best possible defense parameters + random initialization”⁹. Some researchers even drew different conclusions because of different settings in their experiments. If there exists any benchmark, where both adversaries and defenders conduct experiments in a uniform way (i.e., the same threat model, dataset, classifier, attacking/defending approach), we can make a more precise comparison between different attacking and defending techniques.

Cleverhans [139] and Foolbox [140] are open-source libraries to benchmark the vulnerability of deep neural networks against adversarial images. They build frameworks to evaluate the attacks. However, defensive strategies are missing in both tools. Providing a dataset of adversarial examples generated by different methods will make it easy for finding the blind point of deep neural networks and developing new defense strategies. This problem also occurs in other areas in deep learning.

Google Brain organized three competitions in NIPS 2017 competition track, including targeted adversarial attack, non-targeted adversarial attack, and defense against adversarial attack [141]. The dataset in the competition consisted of a set of images never used before and manually labeled the images, 1,000 images for development and 5,000 images for final testing. The submitted attacks and competitions are used as benchmarks to evaluate themselves. The adversarial attacks

⁹Code repository used in [77]: https://github.com/carlini/nn_breaking_detection

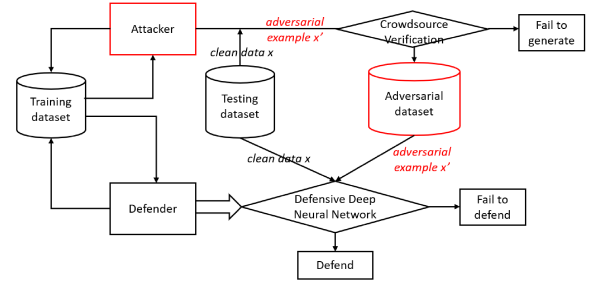


Figure 10: Workflow of a benchmark platform for attackers and defenders: 1) attackers and defenders update/train their strategies on training dataset; 2) attackers generate adversarial examples on the clean data; 3) the adversarial examples are verified by crowdsourcing whether recognizable to human; 4) defenders generate a deep neural network as a defensive strategy; 5) evaluate the defensive strategy.

and defenses are scored by the number of runs to fool the defenses/correctly classify images.

We present the workflow of a benchmark platform for attackers and defenders (Figure 10).

3) **Various applications for robustness evaluation:** Similar to the existence of adversarial examples for various applications, a wide range of applications make it hard to evaluate the robustness, of a deep neural network architecture. How to compare methods generating adversarial example under different threat models? Do we have a universal methodology to evaluate the robustness under all scenarios? Tackling these unsolved problems is a future direction.

VIII. CONCLUSION

In this paper, we reviewed recent findings of adversarial examples in deep neural networks. We investigated existing methods for generating adversarial examples¹⁰. A taxonomy of adversarial examples was proposed. We also explored the applications and countermeasures for adversarial examples.

This paper attempted to cover state-of-the-art studies for adversarial examples in the deep learning domain. Compared with recent work on adversarial examples, we analyzed and discussed current challenges and potential solutions in adversarial examples.

REFERENCES

- [1] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “Imagenet classification with deep convolutional neural networks,” in *Advances in neural information processing systems*, 2012, pp. 1097–1105.
- [2] K. Simonyan and A. Zisserman, “Very deep convolutional networks for large-scale image recognition,” *arXiv preprint arXiv:1409.1556*, 2014.
- [3] J. Redmon and A. Farhadi, “Yolo9000: Better, faster, stronger,” *arXiv preprint arXiv:1612.08242*, 2016.
- [4] S. Ren, K. He, R. Girshick, and J. Sun, “Faster r-cnn: Towards real-time object detection with region proposal networks,” in *Advances in neural information processing systems*, 2015, pp. 91–99.
- [5] G. Saon, H.-K. J. Kuo, S. Rennie, and M. Picheny, “The ibm 2015 english conversational telephone speech recognition system,” *arXiv preprint arXiv:1505.05899*, 2015.
- [6] I. Sutskever, O. Vinyals, and Q. V. Le, “Sequence to sequence learning with neural networks,” in *Advances in neural information processing systems*, 2014, pp. 3104–3112.

¹⁰Due to the rapid development of adversarial examples (attacks and defenses), we only considered the papers published before November 2017. We will update the survey with new methodologies and papers in our future work.

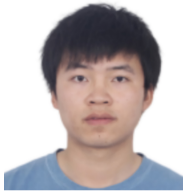
- [7] A. v. d. Oord, S. Dieleman, H. Zen, K. Simonyan, O. Vinyals, A. Graves, N. Kalchbrenner, A. Senior, and K. Kavukcuoglu, "Wavenet: A generative model for raw audio," *arXiv preprint arXiv:1609.03499*, 2016.
- [8] C. Szegedy, W. Zaremba, I. Sutskever, J. Bruna, D. Erhan, I. Goodfellow, and R. Fergus, "Intriguing properties of neural networks," *arXiv preprint arXiv:1312.6199*, 2013.
- [9] A. Kurakin, I. Goodfellow, and S. Bengio, "Adversarial examples in the physical world," *arXiv preprint arXiv:1607.02533*, 2016.
- [10] I. Evtimov, K. Eykholt, E. Fernandes, T. Kohno, B. Li, A. Prakash, A. Rahmati, and D. Song, "Robust physical-world attacks on deep learning models," *arXiv preprint arXiv:1707.08945*, vol. 1, 2017.
- [11] C. Xie, J. Wang, Z. Zhang, Y. Zhou, L. Xie, and A. Yuille, "Adversarial examples for semantic segmentation and object detection," in *International Conference on Computer Vision*. IEEE, 2017.
- [12] N. Carlini, P. Mishra, T. Vaidya, Y. Zhang, M. Sherr, C. Shields, D. Wagner, and W. Zhou, "Hidden voice commands," in *USENIX Security Symposium*, 2016, pp. 513–530.
- [13] G. Zhang, C. Yan, X. Ji, T. Zhang, T. Zhang, and W. Xu, "Dolphinattack: Inaudible voice commands," *arXiv preprint arXiv:1708.09537*, 2017.
- [14] "iOS - Siri - Apple," <https://www.apple.com/ios/siri/>.
- [15] "Alexa," <https://developer.amazon.com/alexa>.
- [16] "Cortana | Your Intelligent Virtual & Personal Assistant | Microsoft," <https://www.microsoft.com/en-us/windows/cortana>.
- [17] W. Knight, "The dark secret at the heart of ai," *MIT Technology Review*, 2017.
- [18] D. Castelvocchi, "Can we open the black box of AI?" *Nature News*, vol. 538, no. 7623, p. 20, 2016.
- [19] P. W. Koh and P. Liang, "Understanding black-box predictions via influence functions," *Proceedings of the International Conference on Machine Learning (ICML)*, 2017.
- [20] Z. C. Lipton, "The myths of model interpretability," *International Conference on Machine Learning (ICML) Workshop*, 2016.
- [21] R. Shwartz-Ziv and N. Tishby, "Opening the black box of deep neural networks via information," *arXiv preprint arXiv:1703.00810*, 2017.
- [22] R. R. Selvaraju, A. Das, R. Vedantam, M. Cogswell, D. Parikh, and D. Batra, "Grad-cam: Why did you say that? visual explanations from deep networks via gradient-based localization," *arXiv preprint arXiv:1610.02391*, 2016.
- [23] J. Lu, T. Issarano, and D. Forsyth, "Safetynet: Detecting and rejecting adversarial examples robustly," *ICCV*, 2017.
- [24] Y. Wu, D. Bamman, and S. Russell, "Adversarial training for relation extraction," in *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, 2017, pp. 1779–1784.
- [25] Y. Dong, H. Su, J. Zhu, and F. Bao, "Towards interpretable deep neural networks by leveraging adversarial examples," *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [26] N. Papernot, P. McDaniel, A. Sinha, and M. Wellman, "Sok: Towards the science of security and privacy in machine learning," *arXiv preprint arXiv:1611.03814*, 2016.
- [27] B. Biggio, B. Nelson, and P. Laskov, "Poisoning attacks against support vector machines," *arXiv preprint arXiv:1206.6389*, 2012.
- [28] F. Roli, B. Biggio, and G. Fumera, "Pattern recognition systems under attack," in *Iberoamerican Congress on Pattern Recognition*. Springer, 2013, pp. 1–8.
- [29] H. Xiao, B. Biggio, G. Brown, G. Fumera, C. Eckert, and F. Roli, "Is feature selection secure against training data poisoning?" in *International Conference on Machine Learning*, 2015, pp. 1689–1698.
- [30] M. Mozaffari-Kermani, S. Sur-Kolay, A. Raghunathan, and N. K. Jha, "Systematic poisoning attacks on and defenses for machine learning in healthcare," *IEEE journal of biomedical and health informatics*, vol. 19, no. 6, pp. 1893–1905, 2015.
- [31] A. Beatson, Z. Wang, and H. Liu, "Blind attacks on machine learners," in *Advances in Neural Information Processing Systems*, 2016, pp. 2397–2405.
- [32] S. Alfeld, X. Zhu, and P. Barford, "Data poisoning attacks against autoregressive models," in *AAAI*, 2016, pp. 1452–1458.
- [33] N. Papernot, P. D. McDaniel, and I. J. Goodfellow, "Transferability in machine learning: from phenomena to black-box attacks using adversarial samples," *CoRR*, vol. abs/1605.07277, 2016.
- [34] M. Fredrikson, S. Jha, and T. Ristenpart, "Model inversion attacks that exploit confidence information and basic countermeasures," in *Proceedings of the 22nd ACM SIGSAC Conference on Computer and Communications Security*. ACM, 2015, pp. 1322–1333.
- [35] M. Abadi, A. Chu, I. Goodfellow, H. B. McMahan, I. Mironov, K. Talwar, and L. Zhang, "Deep learning with differential privacy," in *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*. ACM, 2016, pp. 308–318.
- [36] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *Security and Privacy (SP)*, 2017 IEEE Symposium on. IEEE, 2017, pp. 3–18.
- [37] Y. Bengio, Y. LeCun *et al.*, "Scaling learning algorithms towards ai," *Large-scale kernel machines*, vol. 34, no. 5, pp. 1–41, 2007.
- [38] D. Storcheus, A. Rostamizadeh, and S. Kumar, "A survey of modern questions and challenges in feature extraction," in *Feature Extraction: Modern Questions and Challenges*, 2015, pp. 1–18.
- [39] K. He, X. Zhang, S. Ren, and J. Sun, "Delving deep into rectifiers: Surpassing human-level performance on imagenet classification," in *Proceedings of the IEEE international conference on computer vision*, 2015, pp. 1026–1034.
- [40] J. Redmon, S. Divvala, R. Girshick, and A. Farhadi, "You only look once: Unified, real-time object detection," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 779–788.
- [41] P. He, W. Huang, T. He, Q. Zhu, Y. Qiao, and X. Li, "Single shot text detector with regional attention," in *The IEEE International Conference on Computer Vision (ICCV)*, vol. 6, no. 7, 2017.
- [42] C. Yan, H. Xie, J. Chen, Z.-J. Zha, X. Hao, Y. Zhang, and Q. Dai, "An effective uyghur text detector for complex background images," *IEEE Transactions on Multimedia*, 2018.
- [43] J. Long, E. Shelhamer, and T. Darrell, "Fully convolutional networks for semantic segmentation," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 3431–3440.
- [44] L.-C. Chen, G. Papandreou, F. Schroff, and H. Adam, "Rethinking atrous convolution for semantic image segmentation," *arXiv preprint arXiv:1706.05587*, 2017.
- [45] Q. V. Le, "Building high-level features using large scale unsupervised learning," in *Acoustics, Speech and Signal Processing (ICASSP)*, 2013 IEEE International Conference on. IEEE, 2013, pp. 8595–8598.
- [46] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich, "Going deeper with convolutions," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2015, pp. 1–9.
- [47] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens, and Z. Wojna, "Rethinking the inception architecture for computer vision," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 2818–2826.
- [48] C. Szegedy, S. Ioffe, V. Vanhoucke, and A. A. Alemi, "Inception-v4, inception-resnet and the impact of residual connections on learning," in *AAAI*, 2017, pp. 4278–4284.
- [49] Y. LeCun, C. Cortes, and C. Burges, "The mnist data set," 1998.
- [50] A. Krizhevsky and G. Hinton, "Learning multiple layers of features from tiny images," *Technical report, University of Toronto*, 2009.
- [51] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, A. C. Berg, and L. Fei-Fei, "ImageNet Large Scale Visual Recognition Challenge," *International Journal of Computer Vision (IJCV)*, vol. 115, no. 3, pp. 211–252, 2015.
- [52] M. Barreno, B. Nelson, A. D. Joseph, and J. Tygar, "The security of machine learning," *Machine Learning*, vol. 81, no. 2, pp. 121–148, 2010.
- [53] N. Dalvi, P. Domingos, S. Shanghai, and D. Verma, "Adversarial classification," in *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*. ACM, 2004, pp. 99–108.
- [54] D. Lowd and C. Meek, "Adversarial learning," in *Proceedings of the eleventh ACM SIGKDD international conference on Knowledge discovery in data mining*. ACM, 2005, pp. 641–647.
- [55] B. Biggio, G. Fumera, and F. Roli, "Multiple classifier systems for robust classifier design in adversarial environments," *International Journal of Machine Learning and Cybernetics*, vol. 1, no. 1-4, pp. 27–41, 2010.
- [56] B. Biggio, I. Corona, D. Maiorca, B. Nelson, N. Šrđić, P. Laskov, G. Giacinto, and F. Roli, "Evasion attacks against machine learning at test time," in *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*. Springer, 2013, pp. 387–402.
- [57] M. Barreno, B. Nelson, R. Sears, A. D. Joseph, and J. D. Tygar, "Can machine learning be secure?" in *Proceedings of the 2006 ACM Symposium on Information, computer and communications security*. ACM, 2006, pp. 16–25.

- [58] M. Sharif, S. Bhagavatula, L. Bauer, and M. K. Reiter, "Accessorize to a crime: Real and stealthy attacks on state-of-the-art face recognition," in *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*. ACM, 2016, pp. 1528–1540.
- [59] A. Rozsa, E. M. Rudd, and T. E. Boult, "Adversarial diversity and hard positive generation," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 2016, pp. 25–32.
- [60] I. J. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and harnessing adversarial examples," *arXiv preprint arXiv:1412.6572*, 2014.
- [61] N. Papernot, P. McDaniel, S. Jha, M. Fredrikson, Z. B. Celik, and A. Swami, "The limitations of deep learning in adversarial settings," in *Security and Privacy (EuroS&P), 2016 IEEE European Symposium on*. IEEE, 2016, pp. 372–387.
- [62] S.-M. Moosavi-Dezfooli, A. Fawzi, and P. Frossard, "Deepfool: a simple and accurate method to fool deep neural networks," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 2574–2582.
- [63] N. Carlini and D. Wagner, "Towards evaluating the robustness of neural networks," in *Security and Privacy (S&P), 2017 IEEE Symposium on*. IEEE, 2017, pp. 39–57.
- [64] P.-Y. Chen, H. Zhang, Y. Sharma, J. Yi, and C.-J. Hsieh, "Zoo: Zeroth order optimization based black-box attacks to deep neural networks without training substitute models," *arXiv preprint arXiv:1708.03999*, 2017.
- [65] S.-M. Moosavi-Dezfooli, A. Fawzi, O. Fawzi, and P. Frossard, "Universal adversarial perturbations," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [66] J. Su, D. V. Vargas, and S. Kouichi, "One pixel attack for fooling deep neural networks," *arXiv preprint arXiv:1710.08864*, 2017.
- [67] S. Sabour, Y. Cao, F. Faghri, and D. J. Fleet, "Adversarial manipulation of deep representations," *Proceedings of the International Conference on Learning Representations (ICLR)*, 2016.
- [68] Z. Zhao, D. Dua, and S. Singh, "Generating natural adversarial examples," *arXiv preprint arXiv:1710.11342*, 2017.
- [69] Y. Liu, X. Chen, C. Liu, and D. Song, "Delving into transferable adversarial examples and black-box attacks," *Proceedings of the International Conference on Learning Representations (ICLR)*, 2017.
- [70] N. Carlini, G. Katz, C. Barrett, and D. L. Dill, "Ground-truth adversarial examples," *arXiv preprint arXiv:1709.10207*, 2017.
- [71] P. Tabacof and E. Valle, "Exploring the space of adversarial images," in *Neural Networks (IJCNN), 2016 International Joint Conference on*. IEEE, 2016, pp. 426–433.
- [72] A. Kurakin, I. Goodfellow, and S. Bengio, "Adversarial machine learning at scale," *Proceedings of the International Conference on Learning Representations (ICLR)*, 2017.
- [73] Y. Dong, F. Liao, T. Pang, H. Su, X. Hu, J. Li, and J. Zhu, "Boosting adversarial attacks with momentum," *arXiv preprint arXiv:1710.06081*, 2017.
- [74] A. Nguyen, J. Yosinski, and J. Clune, "Deep neural networks are easily fooled: High confidence predictions for unrecognizable images," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 427–436.
- [75] F. Tramèr, A. Kurakin, N. Papernot, D. Boneh, and P. McDaniel, "Ensemble adversarial training: Attacks and defenses," *arXiv preprint arXiv:1705.07204*, 2017.
- [76] A. Cully, J. Clune, D. Tarapore, and J.-B. Mouret, "Robots that can adapt like animals," *Nature*, vol. 521, no. 7553, pp. 503–507, may 2015.
- [77] N. Carlini and D. Wagner, "Adversarial examples are not easily detected: Bypassing ten detection methods," *AISEC*, 2017.
- [78] —, "Magnet and efficient defenses against adversarial attacks are not robust to adversarial examples," *arXiv preprint arXiv:1711.08478*, 2017.
- [79] J. T. Springenberg, A. Dosovitskiy, T. Brox, and M. Riedmiller, "Striving for simplicity: The all convolutional net," *arXiv preprint arXiv:1412.6806*, 2014.
- [80] M. Lin, Q. Chen, and S. Yan, "Network in network," *arXiv preprint arXiv:1312.4400*, 2013.
- [81] G. D. Evangelidis and E. Z. Psarakis, "Parametric image alignment using enhanced correlation coefficient maximization," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 30, no. 10, pp. 1858–1865, 2008.
- [82] J. R. Flynn, S. Ward, J. Abich, and D. Poole, "Image quality assessment using the ssim and the just noticeable difference paradigm," in *International Conference on Engineering Psychology and Cognitive Ergonomics*. Springer, 2013, pp. 23–30.
- [83] G. Katz, C. Barrett, D. Dill, K. Julian, and M. Kochenderfer, "Reluplex: An efficient smt solver for verifying deep neural networks," *arXiv preprint arXiv:1702.01135*, 2017.
- [84] S. Huang, N. Papernot, I. Goodfellow, Y. Duan, and P. Abbeel, "Adversarial attacks on neural network policies," *arXiv preprint arXiv:1702.02284*, 2017.
- [85] J. Kos and D. Song, "Delving into adversarial attacks on deep policies," *ICLR Workshop*, 2017.
- [86] J. Kos, I. Fischer, and D. Song, "Adversarial examples for generative models," *arXiv preprint arXiv:1702.06832*, 2017.
- [87] P. Tabacof, J. Tavares, and E. Valle, "Adversarial images for variational autoencoders," *NIPS Workshop*, 2016.
- [88] V. Fischer, M. C. Kumar, J. H. Metzen, and T. Brox, "Adversarial examples for semantic image segmentation," *ICLR 2017 workshop*, 2017.
- [89] J. H. Metzen, T. Genewein, V. Fischer, and B. Bischoff, "On detecting adversarial perturbations," *Proceedings of 5th International Conference on Learning Representations (ICLR)*, 2017.
- [90] R. Jia and P. Liang, "Adversarial examples for evaluating reading comprehension systems," in *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2017.
- [91] J. Li, W. Monroe, and D. Jurafsky, "Understanding neural networks through representation erasure," *arXiv preprint arXiv:1612.08220*, 2016.
- [92] K. Grosse, N. Papernot, P. Manoharan, M. Backes, and P. McDaniel, "Adversarial examples for malware detection," in *European Symposium on Research in Computer Security*. Springer, 2017, pp. 62–79.
- [93] H. S. Anderson, A. Kharkar, B. Filar, and P. Roth, "Evading machine learning malware detection," *Black Hat*, 2017.
- [94] W. Hu and Y. Tan, "Generating adversarial malware examples for black-box attacks based on gan," *arXiv preprint arXiv:1702.05983*, 2017.
- [95] H. S. Anderson, J. Woodbridge, and B. Filar, "Deepdga: Adversarially-tuned domain generation and detection," in *Proceedings of the 2016 ACM Workshop on Artificial Intelligence and Security (AISec)*. ACM, 2016, pp. 13–21.
- [96] W. Xu, Y. Qi, and D. Evans, "Automatically evading classifiers," in *Network and Distributed System Security Symposium (NDSS)*, 2016.
- [97] V. Mnih, A. P. Badia, M. Mirza, A. Graves, T. Lillicrap, T. Harley, D. Silver, and K. Kavukcuoglu, "Asynchronous methods for deep reinforcement learning," in *International Conference on Machine Learning*, 2016, pp. 1928–1937.
- [98] G. Toderici, D. Vincent, N. Johnston, S. J. Hwang, D. Minnen, J. Shor, and M. Covell, "Full resolution image compression with recurrent neural networks," *arXiv preprint arXiv:1608.05148*, 2016.
- [99] C. Ledig, L. Theis, F. Huszar, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang *et al.*, "Photo-realistic single image super-resolution using a generative adversarial network," *Conference on Computer Vision and Pattern Recognition*, 2017.
- [100] A. B. L. Larsen, S. K. Sønderby, H. Larochelle, and O. Winther, "Autoencoding beyond pixels using a learned similarity metric," *arXiv preprint arXiv:1512.09300*, 2015.
- [101] O. M. Parkhi, A. Vedaldi, A. Zisserman *et al.*, "Deep face recognition," in *BMVC*, vol. 1, no. 3, 2015, p. 6.
- [102] J. Lu, H. Sibai, E. Fabry, and D. Forsyth, "Standard detectors aren't (currently) fooled by physical adversarial stop signs," *arXiv preprint arXiv:1710.03337*, 2017.
- [103] J. Hendrik Metzen, M. Chaithanya Kumar, T. Brox, and V. Fischer, "Universal adversarial perturbations against semantic image segmentation," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 2755–2764.
- [104] P. Rajpurkar, J. Zhang, K. Lopyrev, and P. Liang, "Squad: 100,000+ questions for machine comprehension of text," *arXiv preprint arXiv:1606.05250*, 2016.
- [105] Z. Yuan, Y. Lu, Z. Wang, and Y. Xue, "Droid-sec: deep learning in android malware detection," in *ACM SIGCOMM Computer Communication Review*, vol. 44, no. 4. ACM, 2014, pp. 371–372.
- [106] G. E. Dahl, J. W. Stokes, L. Deng, and D. Yu, "Large-scale malware classification using random projections and neural networks," in *Acoustics, Speech and Signal Processing (ICASSP), 2013 IEEE International Conference on*. IEEE, 2013, pp. 3422–3426.
- [107] J. Saxe and K. Berlin, "Deep neural network based malware detection using two dimensional binary program features," in *Malicious and Unwanted Software (MALWARE), 2015 10th International Conference on*. IEEE, 2015, pp. 11–20.

- [108] R. Sun, X. Yuan, P. He, Q. Zhu, A. Chen, A. Gregio, D. Oliveira, and L. Xiaolin, "Learning fast and slow: Propedeutica for real-time malware detection," *arXiv preprint arXiv:1712.01145*, 2017.
- [109] A. N. Bhagoji, D. Cullina, and P. Mittal, "Dimensionality reduction as a defense against evasion attacks on machine learning classifiers," *arXiv preprint arXiv:1704.02654*, 2017.
- [110] R. Feinman, R. R. Curtin, S. Shintre, and A. B. Gardner, "Detecting adversarial samples from artifacts," *arXiv preprint arXiv:1703.00410*, 2017.
- [111] Z. Gong, W. Wang, and W.-S. Ku, "Adversarial and clean data are not twins," *arXiv preprint arXiv:1704.04960*, 2017.
- [112] K. Grosse, P. Manoharan, N. Papernot, M. Backes, and P. McDaniel, "On the (statistical) detection of adversarial examples," *arXiv preprint arXiv:1702.06280*, 2017.
- [113] D. Hendrycks and K. Gimpel, "Early methods for detecting adversarial images," *ICLR Workshop*, 2017.
- [114] D. Meng and H. Chen, "Magnet: a two-pronged defense against adversarial examples," *CCS*, 2017.
- [115] T. Pang, C. Du, Y. Dong, and J. Zhu, "Towards robust detection of adversarial examples," *arXiv preprint arXiv:1706.00633*, 2017.
- [116] Y.-C. Lin, M.-Y. Liu, M. Sun, and J.-B. Huang, "Detecting adversarial attacks on neural network policies with visual foresight," *arXiv preprint arXiv:1710.00814*, 2017.
- [117] S. Gu and L. Rigazio, "Towards deep neural network architectures robust to adversarial examples," *Proceedings of the International Conference on Learning Representations (ICLR)*, 2015.
- [118] Y. Song, T. Kim, S. Nowozin, S. Ermon, and N. Kushman, "Pixeldefend: Leveraging generative models to understand and defend against adversarial examples," *arXiv preprint arXiv:1710.10766*, 2017.
- [119] D. Gopinath, G. Katz, C. S. Pasareanu, and C. Barrett, "DeepSAFE: A data-driven approach for checking adversarial robustness in neural networks," *arXiv preprint arXiv:1710.00486*, 2017.
- [120] G. Katz, C. Barrett, D. L. Dill, K. Julian, and M. J. Kochenderfer, "Towards proving the adversarial robustness of deep neural networks," *arXiv preprint arXiv:1709.02802*, 2017.
- [121] N. Papernot, P. McDaniel, X. Wu, S. Jha, and A. Swami, "Distillation as a defense to adversarial perturbations against deep neural networks," in *Security and Privacy (SP), 2016 IEEE Symposium on*. IEEE, 2016, pp. 582–597.
- [122] R. Huang, B. Xu, D. Schuurmans, and C. Szepesvári, "Learning with a strong adversary," *arXiv preprint arXiv:1511.03034*, 2015.
- [123] J. Bradshaw, A. G. d. G. Matthews, and Z. Ghahramani, "Adversarial examples, uncertainty, and transfer testing robustness in gaussian process hybrid deep networks," *arXiv preprint arXiv:1707.02476*, 2017.
- [124] M. Abbasi and C. Gagné, "Robustness to adversarial examples through an ensemble of specialists," *arXiv preprint arXiv:1702.06856*, 2017.
- [125] J. Ba and R. Caruana, "Do deep nets really need to be deep?" in *Advances in neural information processing systems*, 2014, pp. 2654–2662.
- [126] G. Hinton, O. Vinyals, and J. Dean, "Distilling the knowledge in a neural network," *arXiv preprint arXiv:1503.02531*, 2015.
- [127] T. Salimans, A. Karpathy, X. Chen, and D. P. Kingma, "Pixelcnn++: Improving the pixelcnn with discretized logistic mixture likelihood and other modifications," *ICLR Poster*, 2017.
- [128] W. He, J. Wei, X. Chen, N. Carlini, and D. Song, "Adversarial example defense: Ensembles of weak defenses are not strong," in *11th USENIX Workshop on Offensive Technologies (WOOT 17)*. Vancouver, BC: USENIX Association, 2017.
- [129] F. Tramèr, N. Papernot, I. Goodfellow, D. Boneh, and P. McDaniel, "The space of transferable adversarial examples," *arXiv preprint arXiv:1704.03453*, 2017.
- [130] K. Pei, Y. Cao, J. Yang, and S. Jana, "Deepxplore: Automated whitebox testing of deep learning systems," *Proceedings of the ACM SIGOPS 26th symposium on Operating systems principles*, 2017.
- [131] L. Schmidt, S. Santurkar, D. Tsipras, K. Talwar, and A. Mądry, "Adversarially robust generalization requires more data," *arXiv preprint arXiv:1804.11285*, 2018.
- [132] A. Fawzi, O. Fawzi, and P. Frossard, "Fundamental limits on adversarial robustness," in *Proc. ICML, Workshop on Deep Learning*, 2015.
- [133] T. Tanay and L. Griffin, "A boundary tilting perspective on the phenomenon of adversarial examples," *arXiv preprint arXiv:1608.07690*, 2016.
- [134] A. Fawzi, H. Fawzi, and O. Fawzi, "Adversarial vulnerability for any classifier," *arXiv preprint arXiv:1802.08686*, 2018.
- [135] J. Gilmer, L. Metz, F. Faghri, S. S. Schoenholz, M. Raghu, M. Wattenberg, and I. Goodfellow, "Adversarial spheres," *arXiv preprint arXiv:1801.02774*, 2018.
- [136] A. Fawzi, O. Fawzi, and P. Frossard, "Analysis of classifiers' robustness to adversarial perturbations," *arXiv preprint arXiv:1502.02590*, 2015.
- [137] O. Bastani, Y. Ioannou, L. Lampropoulos, D. Vytiniotis, A. Nori, and A. Criminisi, "Measuring neural net robustness with constraints," in *Advances in Neural Information Processing Systems*, 2016, pp. 2613–2621.
- [138] X. Huang, M. Kwiatkowska, S. Wang, and M. Wu, "Safety verification of deep neural networks," *Computer Aided Verification: 29th International Conference (CAV)*, pp. 3–29, 2017.
- [139] I. Goodfellow, N. Papernot, P. McDaniel, R. Feinman, F. Faghri, A. Matyasko, K. Hambardzumyan, Y.-L. Juang, A. Kurakin, R. Sheatsley, A. Garg, and Y.-C. Lin, "cleverhans v2.0.0: an adversarial machine learning library," *arXiv preprint arXiv:1610.00768*, 2017.
- [140] J. Rauber, W. Brendel, and M. Bethge, "Foolbox v0.8.0: A python toolbox to benchmark the robustness of machine learning models," *arXiv preprint arXiv:1707.04131*, 2017. [Online]. Available: <http://arxiv.org/abs/1707.04131>
- [141] A. Kurakin, I. Goodfellow, S. Bengio, Y. Dong, F. Liao, M. Liang, T. Pang, J. Zhu, X. Hu, C. Xie *et al.*, "Adversarial attacks and defences competition," *arXiv preprint arXiv:1804.00097*, 2018.



Xiaoyong Yuan is currently a PhD student in Department of Computer & Information Science & Engineering at University of Florida. His research interests include deep learning and security. He received a BS degree in mathematics from Fudan University in 2012 and a MS degree in software engineering from Peking University in 2015.



Pan He is currently a PhD student in Department of Computer & Information Science & Engineering at University of Florida. His research interests include deep learning and computer vision. He received B.E. degree from Department of Software Engineering, Sichuan University in June, 2015. As an enthusiastic researcher, his goal is to combine state-of-the-art computer vision algorithms with real-life industrial problems.



Qile Zhu is currently a PhD student in Department of Computer Information Science Engineering at University of Florida. His research interests include deep learning and natural language processing. He received a B.E. degree from Zhejiang University in 2013 and a master degree in University of Florida in 2015.



Dr. Xiaolin (Andy) Li is a Professor and University Term Professor in Department of Electrical and Computer Engineering and Department of Computer & Information Science & Engineering (affiliated at University of Florida). He is the founding Director of National Science Foundation Center for Big Learning, the first NSF center on deep learning, with UF (lead), CMU, UO, and UMKC. He is also Director of Large-scale Intelligent Systems Laboratory (Li Lab). His research interests include Big Data, Machine Learning, Deep Learning, Cloud Computing, Intelligent Platforms, HPC, and Security Privacy for Health, Precision Medicine, IoT, CV, NLP, Robotics, Genomics, and Science, Engineering, and Business. He received a PhD degree in Computer Engineering from Rutgers University. He has published over 120 peer-reviewed papers in journals and conference proceedings, 5 books, and 4 patents (three licensees). His team has created many software systems and tools. He is a recipient of the National Science Foundation CAREER Award, the Internet2 Innovative Application Award, NSF I-Corps Top Team Award, Top Team (DeepBipolar) in the CAGI Challenge on detecting bipolar disorder, and best paper awards (IEEE ICMLA 2016, IEEE SECON 2016, ACM CAC 2013, and IEEE UbiSafe 2007). For more information, <http://www.andyli.ece.ufl.edu> and <http://nscfbl.org>.