

Improving Robustness of Deep-Learning-Based Image Reconstruction

Ankit Raj¹ Yoram Bresler¹ Bo Li¹

Abstract

Deep-learning-based methods for different applications have been shown vulnerable to adversarial examples. These examples make deployment of such models in safety-critical tasks questionable. Use of deep neural networks as inverse problem solvers has generated much excitement for medical imaging including CT and MRI, but recently a similar vulnerability has also been demonstrated for these tasks. We show that for such inverse problem solvers, one should analyze and study the effect of adversaries in the measurement-space, instead of the signal-space as in previous work. In this paper, we propose to modify the training strategy of end-to-end deep-learning-based inverse problem solvers to improve robustness. We introduce an auxiliary network to generate adversarial examples, which is used in a min-max formulation to build robust image reconstruction networks. Theoretically, we show for a linear reconstruction scheme the min-max formulation results in a singular-value(s) filter regularized solution, which suppresses the effect of adversarial examples occurring because of ill-conditioning in the measurement matrix. We find that a linear network using the proposed min-max learning scheme indeed converges to the same solution. In addition, for non-linear Compressed Sensing (CS) reconstruction using deep networks, we show significant improvement in robustness using the proposed approach over other methods. We complement the theory by experiments for CS on two different datasets and evaluate the effect of increasing perturbations on trained networks. We find the behavior for ill-conditioned and well-conditioned measurement matrices to be qualitatively different.

¹University of Illinois at Urbana-Champaign, USA. Correspondence to: Ankit Raj <ankitr3@illinois.edu>, Yoram Bresler <ybresler@illinois.edu>, Bo Li <lbo@illinois.edu>.

Currently under submission

1. Introduction

Adversarial examples for deep learning based methods have been demonstrated for different problems (Szegedy et al., 2013; Kurakin et al., 2016; Cisse et al., 2017a; Eykholt et al., 2017; Xiao et al., 2018). It has been shown that with minute perturbations, these networks can be made to produce unexpected results. Unfortunately, these perturbations can be obtained very easily. There has been plethora of work to defend against these attacks as well (Madry et al., 2017; Tramèr et al., 2017; Athalye et al., 2018; Wong et al., 2018; Jang et al., 2019a; Jiang et al., 2018; Xu et al., 2017; Schmidt et al., 2018). Recently, (Antun et al., 2019; Choi et al., 2019) introduced adversarial attacks on image reconstruction networks. In this work, we propose an adversarial training scheme for image reconstruction deep networks to provide robustness.

Image reconstruction involving the recovery of an image from indirect measurements is used in many applications, including critical applications such as medical imaging, e.g., Magnetic Resonance Imaging (MRI), Computerised Tomography (CT) etc. Such applications demand the reconstruction to be stable and reliable. On the other hand, in order to speed up the acquisition, reduce sensor cost, or reduce radiation dose, it is highly desirable to subsample the measurement data, while still recovering the original image. This is enabled by the compressive sensing (CS) paradigm (Candes et al., 2006; Donoho, 2006). CS involves projecting a high dimensional, signal $x \in \mathbb{R}^n$ to a lower dimensional measurement $y \in \mathbb{R}^m$, $m \ll n$, using a small set of linear, non-adaptive frames. The noisy measurement model is:

$$y = Ax + v, A \in \mathbb{R}^{m \times n}, v \sim \mathcal{N}(0, \sigma^2 I) \quad (1)$$

where A is the measurement matrix. The goal is to recover the unobserved natural image x , from the compressive measurement y . Although the problem with $m \ll n$ is severely ill-posed and does not have a unique solution, CS achieves nice, stable solutions for a special class of signals x - those that are sparse or sparsifiable, by using sparse regularization techniques (Candes et al., 2006; Donoho, 2006; Elad & Aharon, 2006; Dong et al., 2011; Wen et al., 2015; Liu et al., 2017; Dabov et al., 2009; Yang et al., 2010; Elad, 2010; Li et al., 2009; Ravishankar & Bresler, 2012).

Recently, deep learning based methods have also been proposed as an alternative method for performing image recon-

struction (Zhu et al., 2018; Jin et al., 2017; Schlemper et al., 2017; Yang et al., 2017; Hammernik et al., 2018). While these methods have achieved state-of-the-art (SOTA) performance, the networks have been found to be very unstable (Antun et al., 2019), as compared to the traditional methods. Adversarial perturbations have been shown to exist for such networks, which can degrade the quality of image reconstruction significantly. (Antun et al., 2019) studies three types of instabilities: (i) Tiny (small norm) perturbations applied to images that are almost invisible in the original images, but cause a significant distortion in the reconstructed images. (ii) Small structural changes in the original images, that get removed from the reconstructed images. (iii) Stability with increasing the number of measurement samples. We try to address instability (i) above.

In this paper, we argue that studying the instability for image reconstruction networks in the x -space as addressed by (Antun et al., 2019) is sub-optimal and instead, we should consider perturbations in the measurement, y -space. To improve robustness, we modify the training strategy: we introduce an auxiliary network to generate adversarial examples on the fly, which are used in a min-max formulation. This results in an adversarial game between two networks while training, similar to the Generative Adversarial Networks (GANs) (Goodfellow et al., 2014; Arjovsky et al., 2017). However, since the goal here is to build a robust reconstruction network, we make some changes in the training strategy compared to GANs.

Our theoretical analysis for a special case of a linear reconstruction scheme shows that the min-max formulation results in a singular-value filter regularized solution, which suppresses the effect of adversarial examples. Our experiment using the min-max formulation with a learned adversarial example generator for a linear reconstruction network shows that the network indeed converges to the solution obtained theoretically. For a complex non-linear deep network, our experiments show that training using the proposed formulation results in more robust network, both qualitatively and quantitatively, compared to other methods. Further, we experimented and analyzed the reconstruction for two different measurement matrices, one well-conditioned and another relatively ill-conditioned. We find that the behavior in the two cases is qualitatively different.

2. Proposed Method

2.1. Adversarial Training

One of the most powerful methods for training an adversarially robust network is adversarial training (Madry et al., 2017; Tramèr et al., 2017; Sinha et al., 2017; Arnab et al., 2018). It involves training the network using adversarial examples, enhancing the robustness of the network to attacks during inference. This strategy has been quite effective in

classification settings, where the goal is to make the network output the correct label corresponding to the adversarial example.

Standard adversarial training involves solving the following min-max optimization problem:

$$\min_{\theta} \mathbb{E}_{(x,y) \in \mathbb{D}} \left[\max_{\delta: \|\delta\|_p \leq \epsilon} \mathcal{L}(f(x + \delta; \theta), y) \right] \quad (2)$$

where $\mathcal{L}(\cdot)$ represents the applicable loss function, e.g., cross-entropy for classification, and δ is the perturbation added to each sample, within an ℓ_p -norm ball of radius ϵ . This min-max formulation encompasses possible variants of adversarial training. It consists of solving two optimization problems: an inner maximization and an outer minimization problem. This corresponds to an adversarial game between the attacker and robust network f . The inner problem tries to find the optimal $\delta : \|\delta\|_p \leq \epsilon$ for a given data point (x, y) maximizing the loss, which essentially is the adversarial attack, whereas the outer problem aims to find a θ minimizing the same loss. For an optimal θ^* solving the equation 2, then $f(\cdot; \theta^*)$ will be robust (in expected value) to all the x_{adv} lying in the ϵ -radius of ℓ_p -norm ball around the true x .

2.2. Problem Formulation

(Antun et al., 2019) identify instabilities of a deep learning based image reconstruction network by maximizing the following cost function:

$$Q_y(r) = \frac{1}{2} \|f(y + Ar) - x\|_2^2 - \frac{\lambda}{2} \|r\|^2 \quad (3)$$

As evident from this framework, the perturbation r is added in the x -space for each y , resulting in perturbation Ar in the y -space. We argue that this formulation can miss important aspects in image reconstruction, especially in ill-posed problems, for the following three main reasons:

1. It may not be able to model all possible perturbations to y . The perturbations $A\delta$ to y modeled in this formulation are all constrained to the range-space of A . When A does not have full row rank, there exist perturbations to y that cannot be represented as $A\delta$.
2. It misses instabilities created by the ill-conditioning of the reconstruction problem. Consider a simple ill-conditioned reconstruction problem:

$$A = \begin{bmatrix} 1 & 0 \\ 0 & r \end{bmatrix} \text{ and } f = \begin{bmatrix} 1 & 0 \\ 0 & 1/r \end{bmatrix} \quad (4)$$

where A and f define the forward and reconstruction operator respectively, and $|r| \ll 1$. For $\delta = [0, \epsilon]^T$ perturbation in x , the reconstruction is $f(A(x + \delta)) = x + \delta$, and the reconstruction error is $\|f(A(x + \delta)) - x\|_2 = \epsilon$, that is, for small ϵ , the perturbation has

negligible effect. In contrast, for the same perturbation δ in y , the reconstruction is $f(Ax + \delta) = x + [0, \epsilon/r]^T$, with reconstruction error $\|f(A(x + \delta)) - x\|_2 = \epsilon/r$, which can be arbitrarily large if $r \rightarrow 0$. This aspect is completely missed by the formulation based on (3).

3. For inverse problems, one also wants robustness to perturbations in the measurement matrix A . Suppose A used in training is slightly different from the actual $A' = A + \tilde{A}$ that generates the measurements. This results in perturbation $\tilde{A}x$ in y -space, which may be outside the range space of A , and therefore, as in 1 above, may not be possible to capture by the formulation based on (3).

The above points indicate that studying the problem of robustness to perturbations for image reconstruction problems in x -space misses possible perturbations in y -space that can have a huge adversarial effect on reconstruction. Since many of the image reconstruction problems are ill-posed or ill-conditioned, we formulate and study the issue of adversaries in the y -space, which is more generic and able to handle perturbations in the measurement operator A as well.

2.3. Image Reconstruction

Image Reconstruction deals with recovering the clean image x from noisy and possibly incomplete measurements $y = Ax + v$. Recently, deep-learning-based approaches have outperformed the traditional techniques. Many deep learning architectures are inspired by iterative reconstruction schemes (Rick Chang et al., 2017; Raj et al., 2019; Bora et al., 2017; Wen et al., 2019). Another popular way is to use an end-to-end deep network to solve the image reconstruction problem directly (Jin et al., 2017; Zhu et al., 2018; Schlemper et al., 2017; Yang et al., 2017; Hammernik et al., 2018; Sajjadi et al., 2017; Yao et al., 2019). In this work, we propose modification in the training scheme for the end-to-end networks.

Consider the standard MSE loss in x -space with the popular ℓ_2 -regularization on the weights (aka weight decay), which mitigates overfitting and helps in generalization (Krogh & Hertz, 1992)

$$\min_{\theta} \mathbb{E}_x \|f(Ax; \theta) - x\|^2 + \mu \|\theta\|^2 \quad (5)$$

In this paper, we experiment both with $\mu > 0$ (regularization present) and $\mu = 0$ (no regularization). No regularization is used in the sequel, unless stated otherwise.

2.3.1. ADVERSARIAL TRAINING FOR IMAGE RECONSTRUCTION

Motivated by the adversarial training strategy (2), several frameworks have been proposed recently to make classification by deep networks more robust (Jang et al., 2019b;

Kurakin et al., 2016; Wang & Yu, 2019). For image reconstruction, we propose to modify the training loss to the general form

$$\min_{\theta} \mathbb{E}_x \max_{\delta: \|\delta\|_p \leq \epsilon} \|f(Ax; \theta) - x\|^2 + \lambda \|f(Ax + \delta; \theta) - x\|^2$$

The role of the first term is to ensure that the network f maps the non-adversarial measurement to the true x , while the role of the second term is to train f on worst-case adversarial examples within the ℓ_p -norm ball around the nominal measurement Ax . We want δ to be the worst case perturbation for a given f . However, during the initial training epochs, f is mostly random (assuming random initialization of the weights) resulting in random perturbation, which makes f diverge. Hence we need only the first term during initial epochs to get a decent f that provides reasonable reconstruction. Then, reasonable perturbations are obtained by activating the second term, which results in robust f . Now, solving the min-max problem above is intractable for a large dataset as it involves finding the adversarial example, which requires to solve the inner maximization for each $y = Ax$. This may be done using projected gradient descent (PGD), but is very costly. A possible sub-optimal approximation (with $p = 2$) for this formulation is:

$$\min_{\theta} \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_x \|f(Ax; \theta) - x\|_2^2 + \lambda \|f(Ax + \delta; \theta) - x\|_2^2 \quad (6)$$

This formulation finds a common δ which is adversarial to each measurement y and tries to minimize the reconstruction loss for the adversarial examples together with that for clean examples. Clearly this is sub-optimal as using a perturbation δ common to all y 's need not be the worst-case perturbation for any of the y 's, and optimizing for the common δ won't result in a highly robust network.

Ideally, we would want the best of both worlds: i.e., to generate δ for each y independently, together with tractable training. To this end, we propose to parameterize the worst-case perturbation $\delta = \arg \max_{\delta: \|\delta\|_2 \leq \epsilon} \|f(y + \delta; \theta) - x\|_2^2$ by a deep neural network $G(y; \phi)$. This also eliminates the need of solving the inner-maximization to find δ using hand-designed methods. Since $G(\cdot)$ is parameterized by ϕ and takes y as input, a well-trained G will result in optimal perturbation for the given $y = Ax$. The modified loss function becomes:

$$\min_{\theta} \max_{\phi: \|G(\cdot, \phi)\|_2 \leq \epsilon} \mathbb{E}_x \|f(Ax; \theta) - x\|^2 + \lambda \|f(Ax + G(Ax; \phi); \theta) - x\|^2$$

This results in an adversarial game between the two networks: G and f , where G 's goal is to generate strong adversarial examples that maximize the reconstruction loss for the given f , while f tries to make itself robust to the adversarial examples generated by the G . This framework is illustrated in the Fig. 1. This min-max setting is quite

similar to the Generative adversarial network (GAN), with the difference in the objective function. Also, here, the main goal is to build an adversarially robust f , which requires some empirical changes compared to standard GANs to make it work. Another change is to reformulate the constraint $\|G(\cdot, \phi)\|_2 \leq \epsilon$ into a penalty form using the hinge loss, which makes the training more tractable:

$$\begin{aligned} \min_{\theta} \max_{\phi} \quad & \mathbb{E}_x \|f(Ax; \theta) - x\|^2 \\ & + \lambda_1 \|f(Ax + G(Ax; \phi); \theta) - x\|^2 \\ & + \lambda_2 \max\{0, \|G(Ax; \phi)\|_2^2 - \epsilon\} \end{aligned} \quad (7)$$

Note that λ_2 must be negative to satisfy the required constraint $\|G(\cdot, \phi)\|_2 \leq \epsilon$.

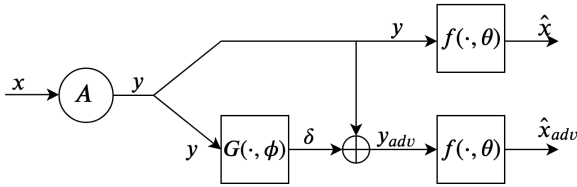


Figure 1. Adversarial training framework of image reconstruction network f , jointly with another network G , generating the additive perturbations

2.3.2. TRAINING STRATEGY

We apply some modifications and intuitive changes to train a robust f jointly with training G in a mini-batch set-up. At each iteration, we update G to generate adversarial examples and train f using those adversarial examples along with the non-adversarial or clean samples to make it robust. Along with the training of robust f , G is being trained to generate worst-case adversarial examples. To generate strong adversarial examples by G in the mini-batch update, we divide each mini-batch into K sets. Now, G is trained over each set independently and we use adversarial examples after the update of G for each set. This fine-tunes G for the small set to generate stronger perturbations for every image belonging to the set. Then, f is trained using the entire mini-batch at once but with the adversarial examples generated set-wise. G obtained after the update corresponding to the K^{th} set is passed for the next iteration or mini-batch update. This is described in Algorithm 1.

2.4. Robustness Metric

We define a metric to compare the robustness of different networks. We measure the following quantity for network f :

$$\Delta_{\max}(x_0, \epsilon) = \max_{\|\delta\|_2 \leq \epsilon} \|f(Ax_0 + \delta) - x_0\|^2 \quad (8)$$

This determines the reconstruction error due to the worst-case additive perturbation over an ϵ -ball around the nominal

Algorithm 1 Algorithm for training at iteration T

Input: Mini-batch samples (x_T, y_T) , G_{T-1} , f_{T-1}

Output: G_T and f_T

- 1: $G_{T,0} = G_{T-1}$, $f = f_{T-1}$ Divide mini-batch into K parts.
 - 2: **while** $k \leq K$ **do**
 - 3: $x = x_{T,k}$, $G = G_{T,k-1}$
 - 4: $G_{T,k} = \arg \max_G \lambda_1 \|f_{T-1}(Ax + G(Ax; \phi); \theta) - x\|^2 + \lambda_2 \max\{0, \|G(Ax; \phi)\|_2^2 - \epsilon\}$
 - 5: $\delta_{T,k} = G_{T,k}(x)$
 - 6: **end while**
 - 7: $\delta_T = [\delta_{T,1}, \delta_{T,2}, \dots, \delta_{T,K}]$
 - 8: $f_T = \arg \min_f \|f(Ax_T) - x_T\|^2 + \lambda_1 \|f(Ax_T + \delta_T) - x_T\|^2$
 - 9: $G_T = G_{T,K}$
 - 10: **return** G_T, f_T
-

measurement $y = Ax_0$ for each image x_0 . The final robustness metric for f is $\rho(\epsilon) = \mathbb{E}_{x_0} [\Delta_{\max}(x_0, \epsilon)]$, which we estimate by the sample average of $\Delta_{\max}(x_0, \epsilon)$ over a test dataset,

$$\hat{\rho}(\epsilon) = \frac{1}{N} \sum_{i=1}^N \Delta_{\max}(x_i, \epsilon) \quad (9)$$

The smaller $\hat{\rho}$, the more robust the network.

We solve the optimization problem in (8) using projected gradient ascent (PGA) with momentum (with parameters selected empirically). Importantly, unlike training, where computation of $\Delta_{\max}(x_0)$ is required at every epoch, we need to solve (8) only once for every sample x_i in the test set, making this computation feasible during testing.

3. Theoretical Analysis

We theoretically obtained the optimal solution for the min-max formulation in (6) for a simple linear reconstruction. Although this analysis doesn't extend easily to the non-linear deep learning based reconstruction, it gives some insights for the behavior of the proposed formulation and how it depends on the conditioning of the measurement matrices.

Theorem 1. Suppose that the reconstruction network f is a one-layer feed-forward network with no non-linearity i.e., $f = B$, where matrix B has SVD: $B = MQP^T$. Denote the SVD of the measurement matrix A by $A = USV^T$, where S is a diagonal matrix with singular values in permuted (increasing) order, and assume that the data is normalized, i.e., $E(x) = 0$ and $cov(x) = I$. Then the optimal B obtained by solving (6) is a modified pseudo-inverse of A , with $M = V$, $P = U$ and Q a filtered inverse

of S , given by the diagonal matrix

$$Q = \text{diag}(q_m, \dots, q_m, 1/S_{m+1}, \dots, 1/S_n),$$

$$q_m = \frac{\sum_{i=1}^m S_i}{\sum_{i=1}^m S_i^2 + \frac{\lambda}{1+\lambda} \epsilon^2} \quad (10)$$

with largest entry q_m of multiplicity m that depends on ϵ , λ and $\{S_i\}_{i=1}^n$.

Proof. Please refer to the appendix A for the proof. $\square$

The modified inverse B reduces the effect of ill-conditioning in A for adversarial cases in the reconstruction. This can be easily understood, using the simple example from the equation 4. As explained previously, for the A in (4) with $|r| < 1$, an exact inverse, $f = \begin{bmatrix} 1 & 0 \\ 0 & \frac{1}{r} \end{bmatrix}$, amplifies the perturbation. Instead the min-max formulation (6) (with $\lambda = 1$) results in a modified pseudo inverse $\hat{f} = \begin{bmatrix} 1 & 0 \\ 0 & \frac{r}{r^2 + 0.5\epsilon^2} \end{bmatrix}$, suppressing the effect of an adversarial perturbation $\delta = [0, \epsilon]^T$ in y as $\|f\delta\| \gg \|\hat{f}\delta\|$ for $r \rightarrow 0$ and $\epsilon \rightarrow 0$. It can also be seen that $\hat{f}$ won't be optimal for the unperturbed y as it's not actual an inverse and reconstruction loss using $\hat{f}$ for unperturbed case would be smaller than that for f . However, for even very small adversaries, $\hat{f}$ would be much more sensitive than f . It shows the trade-off between the perturbed and unperturbed case for the reconstruction in the case of ill-conditioned A .

This trade-off behavior will not manifest for a well-conditioned, as an ideal linear inverse f for this case won't amplify the small perturbations and a reconstruction obtained using (6) with linear $\hat{f}$ will be very close to f (depending on ϵ): for well-conditioned A , $r \rightarrow 0$. In that case $r^2 \gg 0.5\epsilon^2$, which reduces $\hat{f}$ to f .

Our experiments with deep-learning-based non-linear image reconstruction methods for CS using as sensing matrices random rows of a Gaussian matrix (well-conditioned) vs. random rows of a DCT matrix (relatively ill-conditioned) indeed show the qualitatively different behavior with increasing amount of perturbations.

4. Experiments

Network Architecture: For the reconstruction network f , we follow the architecture of deep convolutional networks for image reconstruction. They use multiple convolution, deconvolution and ReLU layers, and use batch normalization and dropout for better generalization. As a pre-processing step, which has been found to be effective for reconstruction, we apply the transpose (adjoint) of A to the measurement y , feeding $A^T y$ to the network. This transforms the measurement into the image-space, allowing the network to operate purely in image space.

For the adversarial perturbation generator G we use a standard feed-forward network, which takes input y as input. The network consists of multiple fully-connected and ReLU layers. We trained the architecture shown in fig. 1 using the objective defined in the (7).

We designed networks of similar structure but different number of layers for the two datasets, MNIST and CelebA used in the experiments.

We used the Adam Optimizer with $\beta_1 = 0.5$, $\beta_2 = 0.999$, learning rate of 10^{-4} and mini-batch size of 128, but divided into $K = 4$ parts during the update of G , described in the algorithm 1. During training, the size ϵ of the perturbation has to be neither too big (affects performance on clean samples) nor too small (results in less robustness). We empirically picked $\epsilon = 2$ for MNIST and $\epsilon = 3$ for the CelebA datasets. However, during testing, we evaluated $\hat{\rho}$, defined in (9) for different ϵ 's (including those not used while training), to obtain a fair assessment of robustness.

We compare the adversarially trained model using the min-max formulation defined in the objective 7, with three models trained using different training schemes:

1. Normally trained model with no regularization, i.e., $\mu = 0$ in (7).
2. ℓ_2 -norm weight regularized model, using (5) with $\mu > 10^{-6}$ (aka weight decay), chosen empirically to avoid over-fitting and improve robustness and generalization of the network.
3. Lipschitz constant ($\mathcal{L}$)-constrained Parseval network (Cisse et al., 2017b). The idea is to constrain the overall Lipschitz constant $\mathcal{L}$ of the network to be ≤ 1 , by making $\mathcal{L}$ of every layer, ≤ 1 . Motivated by the idea that regularizing the spectral norm of weight matrices could help in the context of robustness, this approach proposes to constrain the weight matrices to also be orthonormal, making them *Parseval tight frames*. Let S_{fc} and S_c define the set of indices for fully-connected and convolutional layers respectively. The regularization term to penalize the deviation from the constraint is

$$\frac{\beta}{2} \left(\sum_{i \in S_{fc}} \|W_i^T W_i - I_i\|_2^2 + \sum_{j \in S_c} \|\mathbf{W}_j^T \mathbf{W}_j - \frac{I_j}{k_j}\|_2^2 \right) \quad (11)$$

where W_i is the weight matrix for i th fully connected layer and $\mathbf{W}_j$ is the transformed or unfolded weight matrix of j th convolution layer having kernel size k_j . This transformation requires input to the convolution to shift and repeat k_j^2 times. Hence, to maintain the *Parseval tight frames* constraint on the convolution operator, we need to make $\mathbf{W}_j^T \mathbf{W}_j \approx \frac{I_j}{k_j}$. I_i and I_j are identity matrices whose sizes depend on the size of W_i and $\mathbf{W}_j$ respectively. β controls the weight given to the regularization compared to the standard

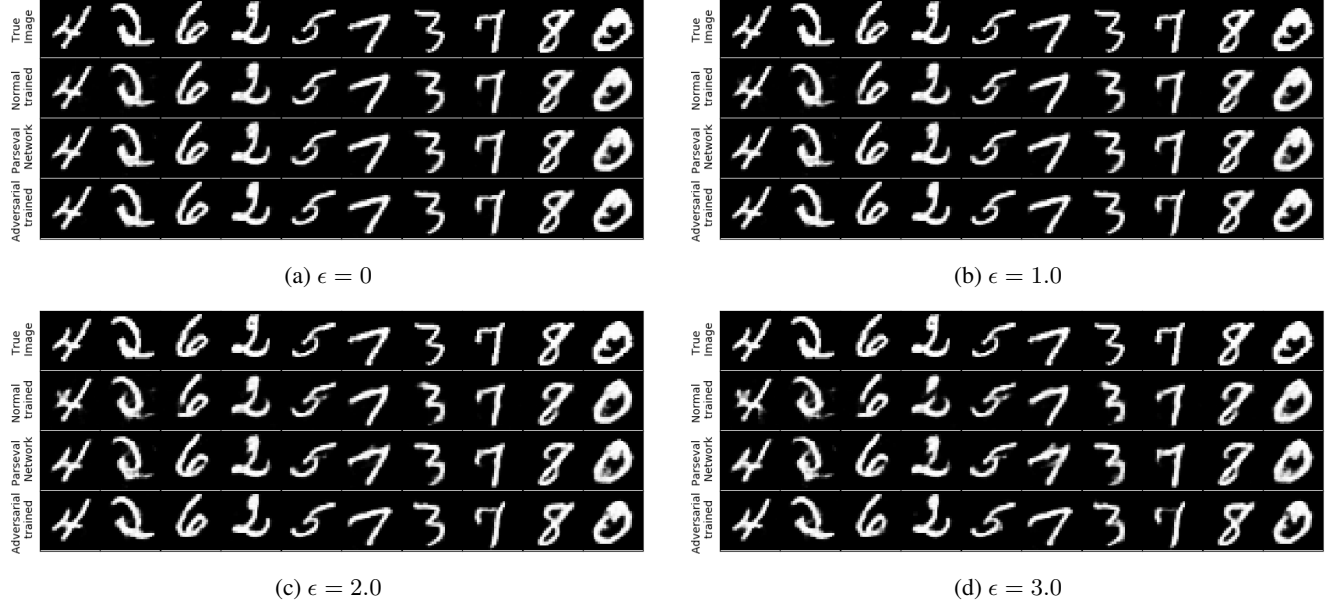


Figure 2. Qualitative Comparison for the MNIST dataset for different perturbations. *First row* of each sub-figure corresponds to the true image, *Second row* to the reconstruction using normally trained model, *Third row* to the reconstruction using Parseval Network, *Fourth row* to the reconstruction using the adversarially trained model (*proposed scheme*).

reconstruction loss. Empirically, we picked β to be 10^{-5} .

To compare different training schemes, we follow the same scheme (described below) for each datasets. Also, we extensively compare the performance for the two datasets for Compressive Sensing (CS) task using two matrices: one well-conditioned and another, relatively ill-conditioned. This comparison complements the theoretical analysis, discussed in the previous section.

The **MNIST** dataset (LeCun et al., 1998) consists of 28×28 gray-scale images of digits with 50,000 training and 10,000 test samples. The image reconstruction network consists of 4 convolution layers and 3 transposed convolution layers using re-scaled images between $[-1, 1]$. For the generator G , we used 5 fully-connected layers network. Empirically, we found $\lambda_1 = 1$ and $\lambda_2 = -0.1$ in (7), gave the best performance in terms of robustness (lower $\hat{\rho}$) for different perturbations.

The **CelebA** dataset (Liu et al., 2015) consists of more than 200,000 celebrity images. We use the aligned and cropped version, which pre-processes each image to a size of $64 \times 64 \times 3$ and scaled between $[-1, 1]$. We randomly pick 160,000 images for the training. Images from the 40,000 held-out set are used for evaluation. The image reconstruction network consists of 6 convolution layers and 4 transposed convolution layers. For the generator G , we used a 6 fully-connected layers network. We found $\lambda_1 = 3$ and $\lambda_2 = -1$ in (7) gave the best robustness performance

(lower $\hat{\rho}$) for different perturbations.

4.1. Gaussian Measurement matrix

In this set-up, we use the same measurement matrix A as (Bora et al., 2017; Raj et al., 2019), i.e. $A_{i,j} \sim N(0, 1/m)$ where m is the number of measurements. For MNIST, the measurement matrix $A \in R^{m \times 784}$, with $m = 100$, whereas for CelebA, $A \in R^{m \times 12288}$, with $m = 1000$. Figures 2 and 3 show the qualitative comparisons for the MNIST and CelebA reconstructions respectively, by solving the optimization described in Section 2.4. It can be seen clearly in both the cases that for different ϵ the adversarially trained models outperform the normally trained and Parseval networks. For higher ϵ 's, the normally trained and Parseval models generate significant artifacts, which are much less for the adversarially trained models. Figures Fig. 4a and Fig. 4b show this improvement in performance in terms of the quantitative metric $\hat{\rho}$, defined in (9) for the MNIST and CelebA datasets respectively. It can be seen that $\hat{\rho}$ is lower for the adversarially-trained models compared to other training methods: no regularization, ℓ_2 -norm regularization on weights, and Parseval networks (Lipschitz-constant-regularized) for different ϵ 's, showing that adversarial training using the proposed min-max formulation indeed outperforms other approaches in terms of robustness. It is noteworthy that even for $\epsilon = 0$, adversarial training reduces the reconstruction loss, indicating that it acts like an excellent regularizer in general.

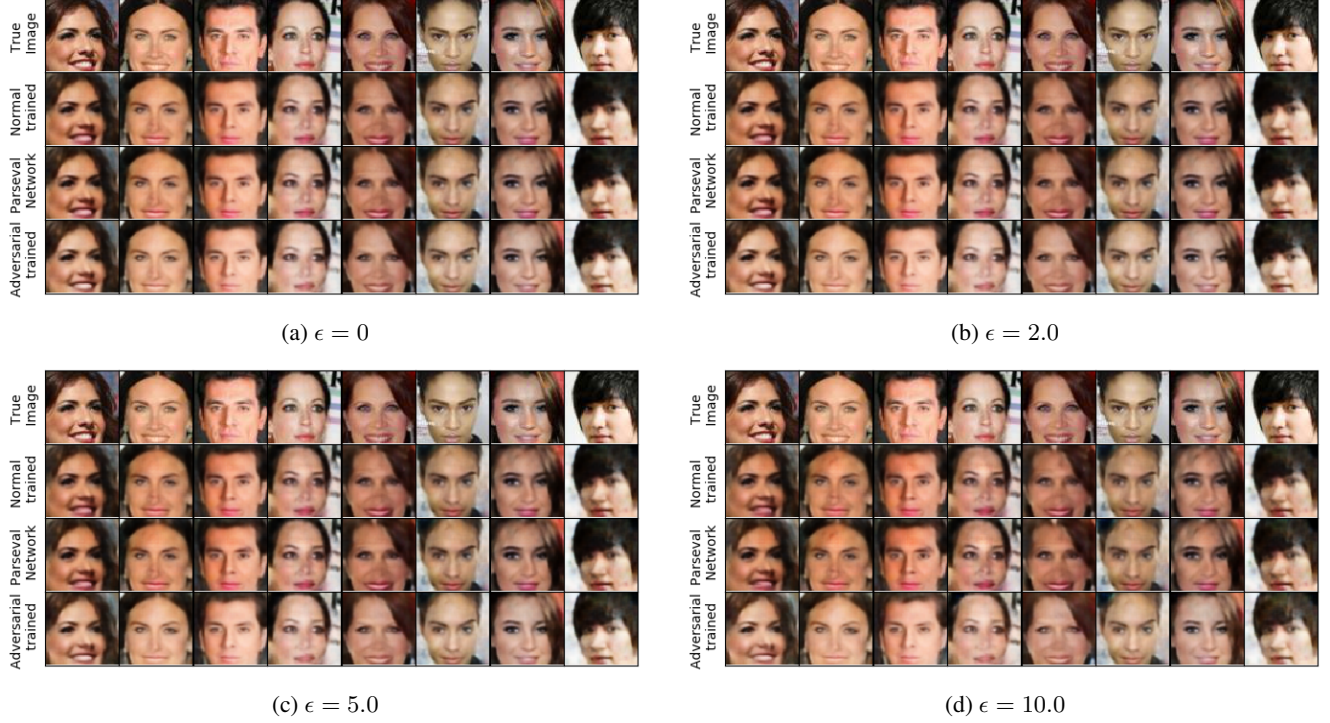


Figure 3. Qualitative Comparison for the CelebA dataset for different perturbations. *First row* of each sub-figure corresponds to the true image, *Second row* to the reconstruction using normally trained model, *Third row* to the reconstruction using Parseval Network, *Fourth row* to the reconstruction using the adversarial trained model (*proposed scheme*).

4.2. Discrete Cosine Transform (DCT) matrix

To empirically study the effect of conditioning of the matrix, we did experiment by choosing A as random m rows and n columns of a $p \times p$ DCT matrix, where $p > n$. This makes A relatively more ill-conditioned than the random Gaussian A , i.e. the condition number for the random DCT matrix is higher than that of random Gaussian one. The number of measurements has been kept same as the previous case, i.e. ($m = 100$, $n = 784$) for MNIST and ($m = 1000$, $n = 12288$) for CelebA. We trained networks having the same configuration as the Gaussian ones. Fig. 4 shows the comparison for the two measurement matrices. Based on the figure, we can see that $\hat{\rho}$ for the DCT, MNIST (Fig. 4d) and CelebA (Fig. 4e), are very close for models trained adversarially and using other schemes for the unperturbed case ($\epsilon = 0$), but the gap between them increases with increasing ϵ 's, with adversarially trained models outperforming the other methods consistently. This behavior is qualitatively different from that for the Gaussian case (Fig. 4a and Fig. 4b), where the gap between adversarial trained networks and models trained using other (or no) regularizers is roughly constant for different ϵ .

4.3. Analysis with respect to Conditioning

To check the conditioning, Fig. 4c shows the histogram for the singular values of the random Gaussian matrices. It can be seen that the condition number (ratio of maximum and minimum singular value) is close to 2 which is very well conditioned for both data sets. On the other hand, the histogram of the same for the random DCT matrices (Fig. 4f) shows higher condition numbers – 8.9 for the 100×784 and 7.9 for the 1000×12288 dimension matrices, which is ill-conditioned relative to the Gaussian ones.

Referring to the above analysis of conditioning and plots of the robustness measure $\hat{\rho}$ for the two types of matrices: random Gaussian vs. random DCT indicate that the performance and behavior of the proposed min-max formulation depends on how well (or relatively ill)-conditioned the matrices are. This corroborates with the theoretical analysis for a simple reconstruction scheme (linear network) described in Sec. 3.

4.4. Linear Network for Reconstruction

We perform an experiment using a linear reconstruction network in a simulated set-up to compare the theoretically obtained optimal robust reconstruction network with the one learned by our scheme by optimizing the objective (6).

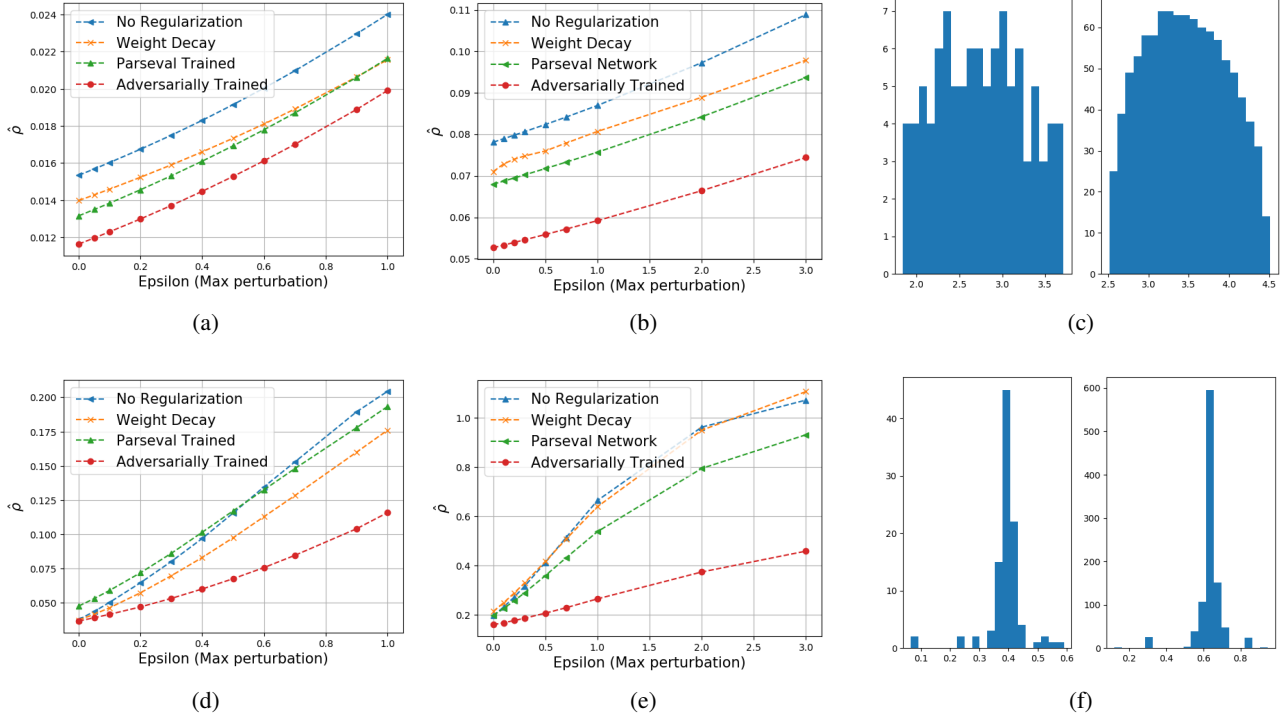


Figure 4. **Row 1** corresponds to the random rows of Gaussian measurement matrix: (a) MNIST, (b) CelebA, (c) Distribution of the singular values for MNIST (left, $m = 100$) and CelebA (right, $m = 1000$) cases. **Row 2** corresponds to random rows of the DCT measurement matrix: (a) MNIST, (b) CelebA, (c) Distribution of the singular values for MNIST (left, $m = 100$) and CelebA (right, $m = 1000$) cases.

We take 50,000 samples of a signal $x \in \mathbb{R}^{20}$ drawn from $\mathcal{N}(0, I)$, hence, $\mathbb{E}(x) = 0$ and $\text{cov}(x) = I$. For the measurement matrix $A \in \mathbb{R}^{10 \times 20}$, we follow the same strategy as in Sec. 4.1, i.e. $A_{ij} \sim \mathcal{N}(0, 1/10)$. Since such matrices are well-conditioned, we replace 2 singular values of A by small values (one being 10^{-3} and another, 10^{-4}) keeping other singular values and singular matrices fixed. This makes the modified matrix $\tilde{A}$ ill-conditioned. We obtain the measurements $y = \tilde{A}x \in \mathbb{R}^{10}$. For reconstruction, we build a linear network f having 1 fully-connected layer with no non-linearity i.e. $f = B \in \mathbb{R}^{20 \times 10}$. The reconstruction is given by $\hat{x} = \hat{B}y$, where $\hat{B}$ is obtained from:

$$\arg \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_x \|B\tilde{A}x - x\|^2 + \lambda \|B(\tilde{A}x + \delta) - x\|^2 \quad (12)$$

We have used $\lambda = 1$, $\epsilon = 0.1$, learning rate = 0.001 and momentum term as 0.9 in our experiments. We obtain the theoretically derived reconstruction B using the result given in (10) (from theorem 1). To compare B and $\hat{B}$, we examined the following three metrics:

- $\|\hat{B} - B\|_F / \|B\|_F = 0.024$, $\|\hat{B} - B\|_2 / \|B\|_2 = 0.034$
- $\|I - B\tilde{A}\|_F / \|I - \hat{B}\tilde{A}\|_F = 0.99936$, where I is the identity matrix of size 20×20

- $\kappa(B) = 19.231$, $\kappa(\hat{B}) = 19.311$, κ : condition number

The above three metrics indicate that $\hat{B}$ indeed converges to the theoretically obtained solution B .

5. Conclusions

In this work, we propose a min-max formulation to build a robust deep-learning-based image reconstruction models. To make this more tractable, we reformulate this using an auxiliary network to generate adversarial examples for which the image reconstruction network tries to minimize the reconstruction loss. We theoretically analyzed a simple linear network and found that using min-max formulation, it outputs singular-value(s) filter regularized solution which reduces the effect of adversarial examples for ill-conditioned matrices. Empirically, we found the linear network to converge to the same solution. Additionally, extensive experiments with non-linear deep networks for Compressive Sensing (CS) using random Gaussian and DCT measurement matrices on MNIST and CelebA datasets show that the proposed scheme outperforms other methods for different perturbations $\epsilon \geq 0$, however the behavior depends on the conditioning of matrices, as indicated by theory for the linear reconstruction scheme.

A. Appendix

Proof of Theorem 1:

For the inverse problem of recovering the true x from the measurement $y = Ax$, goal is to design a robust linear recovery model given by $\hat{x} = By = BAx$.

The min-max formulation to get robust model for a linear set-up:

$$\begin{aligned} \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_{x \in D} \|BAx - x\|^2 + \lambda \|B(Ax + \delta) - x\|^2 \\ \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_{x \in D} (1 + \lambda) \|BAx - x\|^2 + \lambda \|B\delta\|^2 \\ + 2\lambda(B\delta)^T(BAx - x) \end{aligned} \quad (13)$$

Assuming, the dataset is normalized, i.e., $\mathbb{E}(x) = 0$ and $\text{cov}(x) = I$. The above optimization problem becomes:

$$\begin{aligned} \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_{x \in D} (1 + \lambda) \|(BA - I)x\|^2 + \lambda \|B\delta\|^2 \\ \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} \mathbb{E}_{x \in D} (1 + \lambda) \text{tr}(BA - I)xx^T(BA - I)^T \\ + \lambda \|B\delta\|^2 \end{aligned} \quad (14)$$

Since, $\mathbb{E}(\text{tr}(\cdot)) = \text{tr}(\mathbb{E}(\cdot))$, the above problem becomes:

$$\begin{aligned} \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} (1 + \lambda) \text{tr}(BA - I)(BA - I)^T + \lambda \|B\delta\|^2 \\ \min_B \max_{\delta: \|\delta\|_2 \leq \epsilon} (1 + \lambda) \|BA - I\|_F^2 + \lambda \|B\delta\|^2 \end{aligned} \quad (15)$$

Using SVD decomposition of $A = USV^T$ and $B = MQP^T \implies M^T M = I, P^T P = I$ and Q is diagonal. Assume that $\mathbb{G}$ defines the set satisfying the constraints of $M^T M = I, P^T P = I$ and Q is diagonal.

$$\begin{aligned} \min_{M, Q, P \in \mathbb{G}} \max_{\delta: \|\delta\|_2 \leq \epsilon} (1 + \lambda) \|MQP^T USV^T - I\|_F^2 \\ + \lambda \|MQP^T \delta\|^2 \end{aligned} \quad (16)$$

Since, only the second term is dependent on δ , maximizing the second term with respect to δ :

We have $\|MQP^T \delta\| = \|QP^T \delta\|$ since M is unitary. Given, Q is diagonal, $\|QP^T \delta\|^2$ w.r.t. δ can be maximized by having $P^T \delta$ vector having all zeros except the location corresponding to the $\max_i Q_i$. Since, $\|P^T \delta\| = \|\delta\|$, again because P is unitary, so to maximize within the ϵ -ball, we will have $P^T \delta = \epsilon[0, \dots, 0, 1, 0, \dots, 0]$ where 1 is at the $\arg \max_i Q_i$ position. This makes the term to be:

$$\max_{\delta: \|\delta\|_2 \leq \epsilon} \|MQP^T \delta\|^2 = \epsilon^2 (\max_i Q_i)^2$$

Substituting the above term in equation 16:

$$\begin{aligned} \min_{M, Q, P \in \mathbb{G}} (1 + \lambda) \|MQP^T USV^T - I\|_F^2 + \lambda \epsilon^2 (\max_i Q_i)^2 \\ \min_{M, Q, P \in \mathbb{G}} (1 + \lambda) \text{tr}(MQP^T USV^T - I)(MQP^T USV^T - I)^T \\ + \lambda \epsilon^2 (\max_i Q_i)^2 \\ \min_{M, Q, P \in \mathbb{G}} (1 + \lambda) \text{tr}(MQP^T US^2 U^T PQM^T \\ - 2MQP^T USV^T + I) + \lambda \epsilon^2 (\max_i Q_i)^2 \\ \min_{M, Q, P \in \mathbb{G}} (1 + \lambda) \text{tr}(P^T US^2 U^T PQ^2 - 2MQP^T SV^T + I) \\ + \lambda \epsilon^2 (\max_i Q_i)^2 \end{aligned} \quad (17)$$

For the above equation, only second term depends on M , minimizing the second term w.r.t. M keeping others fixed:

$$\min_{M: M^T M = I} \text{tr}(-2MQP^T USV^T)$$

Since, this is a linear program with the quadratic constraint, relaxing the constraint from $M^T M = I$ to $M^T M \leq I$ won't change the optimal point as the optimal point will always be at the boundary i.e. $M^T M = I$

$$\min_{M: M^T M \leq I} \text{tr}(-2MQP^T USV^T) \text{ which is a convex program}$$

Introducing the Lagrange multiplier matrix K for the constraint

$$\mathcal{L}(M, K) = \text{tr}(-2MQP^T USV^T + K(M^T M - I))$$

Substituting $G = QP^T USV^T$ and using stationarity of Lagrangian

$$\Delta L_M = M(K + K^T) - G^T = 0 \implies ML = G^T, L = K + K^T$$

Primal feasibility: $M^T M \leq I$. Optimal point at boundary $\implies M^T M = I$.

Because of the problem is convex, the local minima is the global minima which satisfies the two conditions: Stationarity of Lagrangian ($ML = G^T$) and Primal feasibility ($M^T M = I$). By the choice of $M = V$, and $L = SU^T PQ$, both these conditions are satisfied implying $M = V$ is the optimal point.

Substituting $M = V$ in equation 17, we get:

$$\begin{aligned} \min_{Q, P \in \mathbb{G}} (1 + \lambda) \text{tr}(P^T US^2 U^T PQ^2 - 2VQP^T USV^T + I) \\ + \lambda \epsilon^2 (\max_i Q_i)^2 \\ \min_{Q, P \in \mathbb{G}} (1 + \lambda) \text{tr}(P^T US^2 U^T PQ^2 - 2QP^T US + I) \\ + \lambda \epsilon^2 (\max_i Q_i)^2 \\ \min_{Q, P \in \mathbb{G}} (1 + \lambda) \|QP^T US - I\|_F^2 + \lambda \epsilon^2 (\max_i Q_i)^2 \end{aligned} \quad (18)$$

Denote the i -th column of $C = U^T P$ by c_i and suppose that entries in Q are in decreasing order and the largest entry q_m in Q , has multiplicity m , the equation 18 becomes:

$$\min_{C, Q} (1 + \lambda) \sum_{i=1}^m \|q_m S c_i - e_i\|^2 + \lambda \epsilon^2 q_m^2 + (1 + \lambda) \sum_{i=m+1}^n \|q_i S c_i - e_i\|^2 \quad (19)$$

If we consider the last term i.e. $i > m$, it can be minimized by setting $c_i = e_i$ which is equivalent to choose $P_i = U_i$ and $q_i = 1/S_i$. This makes the last term ($= 0$), using $h = \lambda \epsilon^2 / (1 + \lambda)$, making the equation 19 as:

$$\min_{C, Q} \sum_{i=1}^m (c_i^T S q_m^2 S c_i - 2 e_i^T q_m S c_i + e_i^T e_i) + h q_m^2$$

$$\min_{C, Q} q_m^2 \left(\sum_{i=1}^m c_i^T S^2 c_i + h \right) - 2 q_m \sum_{i=1}^m S_i C_{ii} + \sum_{i=1}^m e_i^T e_i$$

The above term is upward quadratic in q_m , minima w.r.t. q_m will occur at $q_m^* = \frac{\sum_{i=1}^m S_i C_{ii}}{(\sum_{i=1}^m c_i^T S^2 c_i + h)}$, which will make the quadratic term as $\sum_{i=1}^m e_i^T e_i - \frac{(\sum_{i=1}^m S_i C_{ii})^2}{(\sum_{i=1}^m c_i^T S^2 c_i + h)}$, which has to be minimized w.r.t C

$$\min_C \sum_{i=1}^m e_i^T e_i - \frac{(\sum_{i=1}^m S_i C_{ii})^2}{(\sum_{i=1}^m c_i^T S^2 c_i + h)}$$

$$\max_C \frac{(\sum_{i=1}^m S_i C_{ii})^2}{(\sum_{i=1}^m c_i^T S^2 c_i + h)}$$

$$\max_C \frac{(\sum_{i=1}^m S_i C_{ii})^2}{\sum_{i=1}^m S_i^2 C_{ii}^2 + \sum_{j \neq i} S_j^2 C_{ij}^2 + h} \quad (20)$$

Since $C = U^T P \implies C_{ij} = u_i^T p_j \implies \|C_{ij}\| \leq 1$. To maximize the term given by the equation 20, we can minimize the denominator by setting the term $C_{ij} = 0$, which makes the matrix C as diagonal.

Divide the matrix U and P into two parts: one corresponding to $i \leq m$ and another $i > m$, where i represents the column-index of $C = U^T P$.

Let $U = [U_1 | U_2]$ and $P = [P_1 | P_2]$. From above, we have $P_2 = U_2$ for $i > m$, making $P = [P_1 | U_2]$.

$$U^T = \begin{bmatrix} U_1^T \\ U_2^T \end{bmatrix} \text{ and } P = [P_1 | U_2]$$

$$U^T P = \begin{bmatrix} U_1^T P_1 & U_1^T U_2 \\ U_2^T P_1 & U_2^T U_2 \end{bmatrix} = \begin{bmatrix} U_1^T P_1 & \mathbf{0} \\ U_2^T P_1 & I \end{bmatrix}$$

Since, $U^T P$ is diagonal, we have $U_2^T P_1 = \mathbf{0}$, $U_1^T P_1 = \Gamma$ where Γ is diagonal. Also, we have $P_1^T P_1 = I$. Only way to satisfy this would be making $P_1 = U_1$ which makes

$P = U$ and $C = I$. It also results in

$$q_m^* = \frac{\sum_{i=1}^m S_i}{\sum_{i=1}^m S_i^2 + h} \quad (21)$$

Hence, the resulting B would be of the form $M Q P^T$ where:

$$M = V, P = U$$

$$Q = \begin{bmatrix} q_m^* & 0 & \dots & 0 \\ 0 & q_m^* & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 1/S_{m+1} & \dots \\ \vdots & \vdots & \vdots & \vdots \\ 0 & \dots & 0 & 1/S_n \end{bmatrix} \quad (22)$$

References

- Antun, V., Renna, F., Poon, C., Adcock, B., and Hansen, A. C. On instabilities of deep learning in image reconstruction-does ai come at a cost? *arXiv preprint arXiv:1902.05300*, 2019.
- Arjovsky, M., Chintala, S., and Bottou, L. Wasserstein gan. *arXiv preprint arXiv:1701.07875*, 2017.
- Arnab, A., Miksik, O., and Torr, P. H. On the robustness of semantic segmentation models to adversarial attacks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 888–897, 2018.
- Athalye, A., Carlini, N., and Wagner, D. Obfuscated gradients give a false sense of security: Circumventing defenses to adversarial examples. *arXiv preprint arXiv:1802.00420*, 2018.
- Bora, A., Jalal, A., Price, E., and Dimakis, A. G. Compressed sensing using generative models. *arXiv preprint arXiv:1703.03208*, 2017.
- Candes, E. J., Romberg, J. K., and Tao, T. Stable signal recovery from incomplete and inaccurate measurements. *Communications on Pure and Applied Mathematics: A Journal Issued by the Courant Institute of Mathematical Sciences*, 59(8):1207–1223, 2006.
- Choi, J.-H., Zhang, H., Kim, J.-H., Hsieh, C.-J., and Lee, J.-S. Evaluating robustness of deep image super-resolution against adversarial attacks. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 303–311, 2019.
- Cisse, M., Adi, Y., Neverova, N., and Keshet, J. Houdini: Fooling deep structured prediction models. *arXiv preprint arXiv:1707.05373*, 2017a.

- Cisse, M., Bojanowski, P., Grave, E., Dauphin, Y., and Usunier, N. Parseval networks: Improving robustness to adversarial examples. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pp. 854–863. JMLR. org, 2017b.
- Dabov, K., Foi, A., Katkovnik, V., and Egiazarian, K. Bm3d image denoising with shape-adaptive principal component analysis. In *SPARS'09-Signal Processing with Adaptive Sparse Structured Representations*, 2009.
- Dong, W., Zhang, L., Shi, G., and Wu, X. Image deblurring and super-resolution by adaptive sparse domain selection and adaptive regularization. *IEEE Transactions on Image Processing*, 20(7):1838–1857, 2011.
- Donoho, D. L. Compressed sensing. *IEEE Transactions on information theory*, 52(4):1289–1306, 2006.
- Elad, M. *Sparse and redundant representations: from theory to applications in signal and image processing*. Springer Science & Business Media, 2010.
- Elad, M. and Aharon, M. Image denoising via sparse and redundant representations over learned dictionaries. *IEEE Transactions on Image processing*, 15(12):3736–3745, 2006.
- Eykholt, K., Evtimov, I., Fernandes, E., Li, B., Rahmati, A., Xiao, C., Prakash, A., Kohno, T., and Song, D. Robust physical-world attacks on deep learning models. *arXiv preprint arXiv:1707.08945*, 2017.
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., and Bengio, Y. Generative adversarial nets. In *Advances in neural information processing systems*, pp. 2672–2680, 2014.
- Hammernik, K., Klatzer, T., Kobler, E., Recht, M. P., Sodickson, D. K., Pock, T., and Knoll, F. Learning a variational network for reconstruction of accelerated mri data. *Magnetic resonance in medicine*, 79(6):3055–3071, 2018.
- Jang, Y., Zhao, T., Hong, S., and Lee, H. Adversarial defense via learning to generate diverse attacks. In *The IEEE International Conference on Computer Vision (ICCV)*, October 2019a.
- Jang, Y., Zhao, T., Hong, S., and Lee, H. Adversarial defense via learning to generate diverse attacks. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 2740–2749, 2019b.
- Jiang, H., Chen, Z., Shi, Y., Dai, B., and Zhao, T. Learning to defense by learning to attack. *arXiv preprint arXiv:1811.01213*, 2018.
- Jin, K. H., McCann, M. T., Froustey, E., and Unser, M. Deep convolutional neural network for inverse problems in imaging. *IEEE Transactions on Image Processing*, 26(9):4509–4522, 2017.
- Krogh, A. and Hertz, J. A. A simple weight decay can improve generalization. In *Advances in neural information processing systems*, pp. 950–957, 1992.
- Kurakin, A., Goodfellow, I., and Bengio, S. Adversarial machine learning at scale. *arXiv preprint arXiv:1611.01236*, 2016.
- LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Li, C., Yin, W., and Zhang, Y. Users guide for tval3: Tv minimization by augmented lagrangian and alternating direction algorithms. *CAAM report*, 20(46-47):4, 2009.
- Liu, D., Wen, B., Liu, X., Wang, Z., and Huang, T. S. When image denoising meets high-level vision tasks: A deep learning approach. *arXiv preprint arXiv:1706.04284*, 2017.
- Liu, Z., Luo, P., Wang, X., and Tang, X. Deep learning face attributes in the wild. In *Proceedings of International Conference on Computer Vision (ICCV)*, December 2015.
- Madry, A., Makelov, A., Schmidt, L., Tsipras, D., and Vladu, A. Towards deep learning models resistant to adversarial attacks. *arXiv preprint arXiv:1706.06083*, 2017.
- Raj, A., Li, Y., and Bresler, Y. Gan-based projector for faster recovery with convergence guarantees in linear inverse problems. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 5602–5611, 2019.
- Ravishankar, S. and Bresler, Y. Learning sparsifying transforms. *IEEE Transactions on Signal Processing*, 61(5): 1072–1086, 2012.
- Rick Chang, J., Li, C.-L., Poczos, B., Vijaya Kumar, B., and Sankaranarayanan, A. C. One network to solve them all—solving linear inverse problems using deep projection models. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 5888–5897, 2017.
- Sajjadi, M. S., Scholkopf, B., and Hirsch, M. Enhancenet: Single image super-resolution through automated texture synthesis. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 4491–4500, 2017.
- Schlemper, J., Caballero, J., Hajnal, J. V., Price, A., and Rueckert, D. A deep cascade of convolutional neural

- networks for mr image reconstruction. In *International Conference on Information Processing in Medical Imaging*, pp. 647–658. Springer, 2017.
- Schmidt, L., Santurkar, S., Tsipras, D., Talwar, K., and Madry, A. Adversarially robust generalization requires more data. In *Advances in Neural Information Processing Systems*, pp. 5014–5026, 2018.
- Sinha, A., Namkoong, H., and Duchi, J. Certifying some distributional robustness with principled adversarial training. *arXiv preprint arXiv:1710.10571*, 2017.
- Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I., and Fergus, R. Intriguing properties of neural networks. *arXiv preprint arXiv:1312.6199*, 2013.
- Tramèr, F., Kurakin, A., Papernot, N., Goodfellow, I., Boneh, D., and McDaniel, P. Ensemble adversarial training: Attacks and defenses. *arXiv preprint arXiv:1705.07204*, 2017.
- Wang, H. and Yu, C.-N. A direct approach to robust deep learning using adversarial networks. *arXiv preprint arXiv:1905.09591*, 2019.
- Wen, B., Ravishankar, S., and Bresler, Y. Structured over-complete sparsifying transform learning with convergence guarantees and applications. *International Journal of Computer Vision*, 114(2-3):137–167, 2015.
- Wen, B., Ravishankar, S., Pfister, L., and Bresler, Y. Transform learning for magnetic resonance image reconstruction: From model-based learning to building neural networks. *arXiv preprint arXiv:1903.11431*, 2019.
- Wong, E., Schmidt, F., Metzen, J. H., and Kolter, J. Z. Scaling provable adversarial defenses. In *Advances in Neural Information Processing Systems*, pp. 8400–8409, 2018.
- Xiao, C., Li, B., Zhu, J.-Y., He, W., Liu, M., and Song, D. Generating adversarial examples with adversarial networks. *arXiv preprint arXiv:1801.02610*, 2018.
- Xu, W., Evans, D., and Qi, Y. Feature squeezing: Detecting adversarial examples in deep neural networks. *arXiv preprint arXiv:1704.01155*, 2017.
- Yang, G., Yu, S., Dong, H., Slabaugh, G., Dragotti, P. L., Ye, X., Liu, F., Arridge, S., Keegan, J., Guo, Y., et al. Dagan: deep de-aliasing generative adversarial networks for fast compressed sensing mri reconstruction. *IEEE transactions on medical imaging*, 37(6):1310–1321, 2017.
- Yang, J., Wright, J., Huang, T. S., and Ma, Y. Image super-resolution via sparse representation. *IEEE transactions on image processing*, 19(11):2861–2873, 2010.
- Yao, H., Dai, F., Zhang, S., Zhang, Y., Tian, Q., and Xu, C. Dr2-net: Deep residual reconstruction network for image compressive sensing. *Neurocomputing*, 359:483–493, 2019.
- Zhu, B., Liu, J. Z., Cauley, S. F., Rosen, B. R., and Rosen, M. S. Image reconstruction by domain-transform manifold learning. *Nature*, 555(7697):487, 2018.