

IMAGE RESTORATION USING TOTAL VARIATION REGULARIZED DEEP IMAGE PRIOR

Jiaming Liu¹, Yu Sun², Xiaojian Xu², and Ulugbek S. Kamilov^{1,2}

¹ Department of Electrical and Systems Engineering, Washington University in St. Louis, St. Louis, MO 63130

² Department of Computer Science and Engineering, Washington University in St. Louis, St. Louis, MO 63130

ABSTRACT

In the past decade, sparsity-driven regularization has led to significant improvements in image reconstruction. Traditional regularizers, such as total variation (TV), rely on analytical models of sparsity. However, increasingly the field is moving towards trainable models, inspired from deep learning. Deep image prior (DIP) is a recent regularization framework that uses a convolutional neural network (CNN) architecture without data-driven training. This paper extends the DIP framework by combining it with the traditional TV regularization. We show that the inclusion of TV leads to considerable performance gains when tested on several traditional restoration tasks such as image denoising and deblurring.

Index Terms— Image reconstruction, image restoration, deep learning, deep image prior, total variation regularization.

1. INTRODUCTION

Image reconstruction is one of the most widely studied problems in computational imaging. Since the problem is often ill-posed, the process is traditionally regularized by constraining the solutions to be consistent with our prior knowledge about the image. Some traditional imaging priors include nonnegativity, transform-domain sparsity, and self-similarity [1–4]. Recently, however, the attention in the field has been shifting towards new imaging formulations based on deep learning [5].

The most common deep-learning approach is based on an end-to-end training of a convolutional neural network (CNN) for reproducing the desired image from its noisy measurements [6–10]. A popular alternative considers training a CNN as an image denoiser and using it within an iterative reconstruction algorithms [11–14]. However, recently, it was also shown that a CNN can by itself regularize image reconstruction *without* data-driven training [15]. This *deep image prior (DIP)* framework naturally regularizes reconstruction by optimizing the weights of a CNN for it to synthesize the measurements from a given random input vector. The intuition behind DIP is that natural images can be well represented by CNNs, which is not the case for the random noise and certain other image degradations. DIP was shown to achieve remarkable performance on a number of image reconstruction tasks [15, 16].

In this paper, we propose to further improve DIP by combining an *implicit* CNN regularization with an *explicit* TV penalty. The idea of our DIP-TV approach is simple: by including an additional TV term into the objective function, we restrict the solutions synthesized by CNN to those that are piecewise smooth. We experimentally show that our DIP-TV method outperforms the traditional formulations of DIP and TV, and performs on a par with other state-of-the-art image restoration methods such as BM3D [17] and IRCNN [12].

This material is based upon work supported by the National Science Foundation under Grant No. 1813910.



Fig. 1: Comparison of DIP-TV against several standard algorithms for image denoising. DIP-TV achieves the best SNR performances on *Monarch* with AWGN of $\sigma = 65$. The combination of the CNN and TV priors preserve homogeneity of the background as well as the texture, highlighted by rectangles drawn inside the images.

2. BACKGROUND

Consider the restoration as a linear inverse problem

$$\mathbf{y} = \mathbf{H}\mathbf{x} + \mathbf{e}, \quad (1)$$

where the goal is to reconstruct an unknown image $\mathbf{x} \in \mathbb{R}^N$ from the measurements $\mathbf{y} \in \mathbb{R}^M$. Here, $\mathbf{H} \in \mathbb{R}^{M \times N}$ is a degradation matrix and $\mathbf{e} \in \mathbb{R}^M$ corresponds to the measurement noise, which is assumed to be additive white Gaussian (AWGN) of variance σ^2 .

As practical inverse problems are often ill-posed, it is common to regularize the task by constraining the solution according some prior knowledge. In practice, the reconstruction often relies on the regularized least-squares formulation

$$\mathbf{x}^* = \arg \min_{\mathbf{x}} \{ \|\mathbf{y} - \mathbf{H}\mathbf{x}\|_{\ell_2}^2 + \lambda \rho(\mathbf{x}) \} \quad (2)$$

where the data-fidelity term ensures the consistency with measurements, and regularizer ρ constrains the solution to the desired image class. The parameter $\lambda > 0$ controls the strength of regularization.



Fig. 2: The set of 14 grayscale images used in experiments.

Total variation (TV) is one of the most widely used image priors that promotes sparsity in image in image gradients [18]. It has been shown to be effective in a number of applications [19–21]. The ℓ_1 -based anisotropic TV is given by

$$\rho_{TV}(\mathbf{x}) \triangleq \sum_{i=1}^N (|\mathbf{D}_1 \mathbf{x}|_n + |\mathbf{D}_2 \mathbf{x}|_n), \quad (3)$$

where $\mathbf{D}_1$ and $\mathbf{D}_2$ denote the finite difference operation along the first and second dimension of a two-dimensional (2D) image with appropriate boundary conditions.

Currently, deep learning achieves the state-of-the-art performance for different image restoration problems [22–24]. The core idea is to train a CNN via the following optimization

$$\Theta^* = \arg \min_{\Theta} \mathcal{L}(f_{\Theta}(\mathbf{y}), \mathbf{x}), \quad (4)$$

where $f_{\Theta}(\cdot)$ represents the CNN parametrized by Θ . $\mathcal{L}$ denotes the loss function. In practice, (4) can be effectively optimized using the family of stochastic gradient descend (SGD) methods, such as adaptive moment estimation (ADAM) [25].

Recently, Ulyanov *et al.* [15] proposed to use CNN-based methods in an alternative way. They discovered that the architecture of deep CNN models is well-suited for representing natural images, but not random noise. With a random input vector, CNN can reproduce the clear image without supervised training on a large dataset. In the context of image restoration, the associated optimization for DIP can be formulated as

$$\begin{aligned} \Theta^* &= \arg \min_{\Theta} \|\mathbf{y} - \mathbf{H}f_{\Theta}(\mathbf{z})\|_{\ell_2}^2, \\ \text{such that } \mathbf{x}^* &= f_{\Theta^*}(\mathbf{z}). \end{aligned} \quad (5)$$

where $\mathbf{z} \in \mathbb{R}^N$ denotes the random input vector. The CNN generator is initialized with random variables Θ , and these variables are iteratively optimized so that the output of the network is as close to the target measurement as possible.

3. PROPOSED METHOD

The goal of DIP-TV is to use the TV regularization to improve the basic DIP approach. We first consider the optimization problem shown in (2) and the objective function of DIP in (5). One can find that the $\|\mathbf{y} - \mathbf{H}f_{\Theta}(\mathbf{z})\|_{\ell_2}^2$ term in (5) actually corresponds to the data-fidelity term in (3) by replacing $f_{\Theta}(\mathbf{z})$ with an unknown image output. Thus, we can consider replacing (5) with an optimization problem

$$\begin{aligned} \Theta^* &= \arg \min_{\Theta} \{ \|\mathbf{y} - \mathbf{H}f_{\Theta}(\mathbf{z})\|_{\ell_2}^2 + \lambda \rho_{TV}(f_{\Theta}(\mathbf{z})) \}, \\ \text{such that } \mathbf{x}^* &= f_{\Theta^*}(\mathbf{z}). \end{aligned} \quad (6)$$

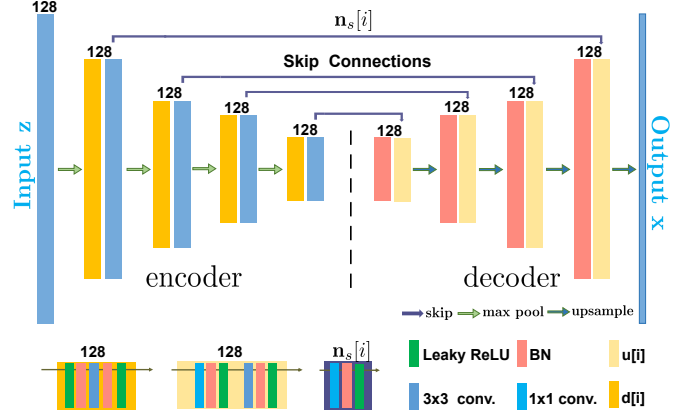


Fig. 3: CNN architecture [15] used in this paper. The architecture is based on the well-known *U-net* with skip connections between the down layers and up layers. Two different kernel sizes are noted under each convolutional layer, and the number of filters is illustrated above each block. The variable $n_s[i]$ denotes the number of feature maps at i th skip layer, and the features in other layers correspond to 128.

Optimization in (6) is similar to training of a CNN and one can rely on any standard optimization algorithms.

Figure 3 illustrates the CNN architecture we used in this paper, which was adapted from [15]. In particular, the popular *U-net* architecture [26] is modified such that the skip connections contain a convolutional layer. The decoder uses a down-sampling and up-sampling based scaling-expanding structure, which makes the effective receptive field of the network increase as the input goes deeper into the network [27]. Besides, the skip connection enables the later layers to reconstruct the feature maps with both local details and global texture. Here, the input $\mathbf{z}$ can be initialized with a fixed 3D tensor with 32 feature maps and of the same spatial size as $\mathbf{x}$ filled with uniform noise. The proposed framework can deal with both grayscale and color images, where for color images anisotropic TV jointly regularizes all three channels.

4. EXPERIMENTS

We now present the experimental results on image denoising and deblurring. We consider 14 gray scale images and 8 standard color images (256×256 and 512×512) from set12, set14, and BSD68 as our testing images. The gray scale images are shown in Figure 2, while color images are: *Monarch*, *Parrots*, *House*, *Lena*, *Peppers*, *Baby*, and *Jet*.

4.1. Image Denoising

In this subsection, we analyze the performance of DIP-TV method for image denoising problems. The CNN architecture in Figure 3 is used for both color and grayscale images, with $n_s[i] = 4$ for each skip layers. All algorithmic hyperparameters were optimized in each experiment for the best signal-to-noise ratio (SNR) performance with respect to the ground truth test image. Both DIP-TV and DIP were set to run 5000 optimization step. We use the *average SNR* to denote the SNR values averaged over the associated set of test images.

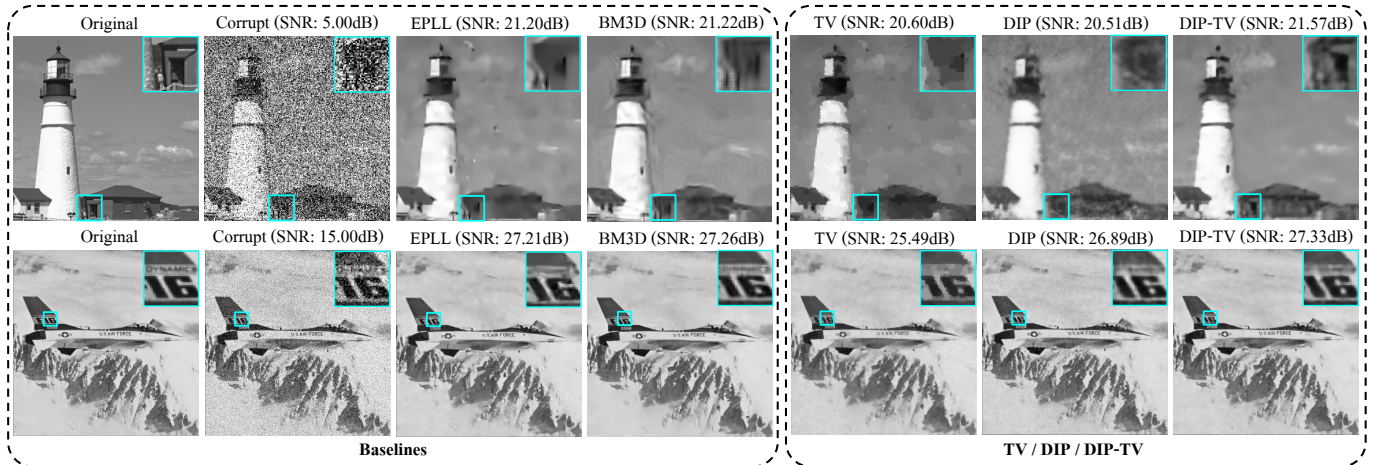


Fig. 4: Image denoising results on *Tower* and *Jet* obtained by EPLL, BM3D, TV-FISTA, DIP, and DIP-TV. The first and second columns display the original images and corrupted images, respectively. Each reconstruction is labeled with its SNR (dB) value with respect to the original image. Visual differences are highlighted using the rectangles drawn inside the images.

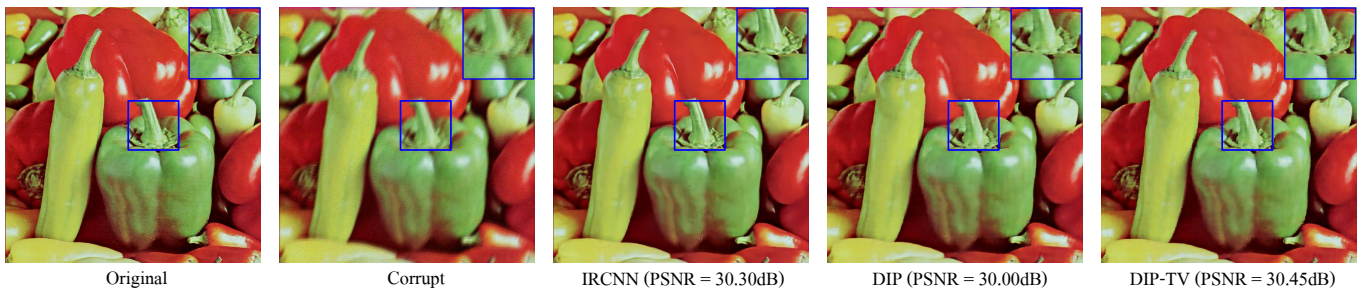


Fig. 5: Image deblurring results with realistic motion blur kernel from [28] and $\sigma = 7.65$ on *Peppers* obtained by IRCNN, DIP, and DIP-TV. Visual differences are highlighted using the rectangles drawn inside the images.

We first present the results of the experiments on grayscale images, where we compared DIP-TV with EPLL [29], BM3D [17], TV [30] and DIP [15]. In order to directly evaluate the range of noise levels that DIP-TV performs better, the input SNR to output SNR relationships are presented in Table 1. The grayscale images were corrupted by AWGN corresponding to input SNR of 5 dB, 10 dB, 15 dB, 20 dB, 25 dB, respectively. In particular, DIP-TV outperforms original DIP by around 0.5 dB for a wide range of noise levels from 5 dB to 20 dB. Note that the proposed method also bridge the gap between DIP and the state-of-the-art methods in high noise levels. Figure 4 illustrates the visual comparisons for grayscale images *Tower* and *Jet* under two different noise levels, respectively. The DIP-TV significantly promotes the denoising performance of DIP itself in terms of both visual qualities and SNR. The noise is effectively filtered out and the details of the image are preserved because of the TV regularization. For instance, DIP-TV improves the SNR with respect to *Tower* by over 1.06 dB against DIP, and outperforms BM3D by 0.35 dB. Visually, the door highlighted in *Tower* is clearly restored, while other methods bring serious distortion to it.

In color image denoising, we compared our method with CBM3D [17] and NLM [31] as well as DIP itself. We considered AWGN corresponding to variance σ from 25 to 75. Figure 1 compares the SNR performance of CBM3D, DIP, and DIP-TV on the image *Monarch*. Table 2 summaries the average SNR among

different methods. Overall, DIP-TV exceeds DIP by at least 0.2 dB on the testing images. Moreover, DIP-TV outperforms CBM3D with the increase of noise level (e.g. $\sigma \geq 35$). Considering that the whole procedure of DIP-TV and DIP are image-agnostic and no prior information is learned from other images, it is notable that DIP-TV achieves comparable performance to the state-of-the-art for high noise levels.

4.2. Image Deblurring

In image deblurring, one is given a blurry image which is synthesized by firstly applying blur kernel $\mathbf{H}$ and then adding AWGN with noise level σ ; The goal is to restore the image from the degraded ones. We tested DIP and DIP-TV based on the network architecture illustrated in [15], with $n_s[i] = 128$. Both DIP and DIP-TV were set to run 5500 optimization step. Taking advantage of recent progress in CNN and the merit of GPU computation, here we utilized convolution to implement the blur. As a baseline, we compared our method with IRCNN [12] and DIP itself based on the same set of images in denoising. Two blur kernels were applied, including a general Gaussian kernel with standard deviation 1.6 as well as a realistic kernel defined in [28]. Different AWGN of σ is added in each experiment.

Figure 5 shows the visual results for *Peppers* obtained by different methods. All methods can effectively remove the blurry and

Table 1: The SNR (dB) results of different methods on the testing images with input noise levels 5 dB, 10 dB, 15 dB, 20 dB, and 25 dB. For example, 5 dB noisy input represents very high noise level and corresponds to $\sigma = 76.26$ in average.

Images	1	2	3	4	5	6	7	8	9	10	11	12	13	14
Input SNR = 5 dB / $\sigma = 76.26$														
EPLL	18.60	21.39	19.18	15.29	16.88	16.54	18.33	21.80	21.21	20.19	19.38	19.85	16.85	21.20
BM3D	18.72	22.22	18.81	15.31	16.86	16.50	18.30	21.87	21.55	20.25	19.52	20.35	17.33	21.22
TV	17.22	20.38	17.65	13.74	16.24	15.42	16.57	19.71	20.09	18.38	18.49	18.27	16.23	20.60
DIP	17.98	21.19	18.78	14.98	16.16	16.19	17.61	21.44	21.08	18.67	18.97	20.19	16.64	20.51
DIP-TV	18.84	22.41	19.56	15.52	16.99	16.79	18.48	22.26	21.61	19.10	19.55	20.52	17.80	21.57
Input SNR = 10 dB / $\sigma = 53.43$														
EPLL	21.21	24.21	21.96	17.81	19.42	19.65	20.88	24.59	23.68	21.20	21.79	22.98	19.65	23.91
BM3D	21.30	25.10	21.57	17.81	19.39	19.58	20.84	24.65	24.01	21.28	21.90	23.39	20.20	23.85
TV	19.76	22.82	20.39	16.34	18.45	18.04	18.91	22.62	22.15	20.34	20.56	20.80	18.85	22.83
DIP	20.76	24.32	21.55	17.81	18.82	19.14	20.21	24.43	23.24	21.01	21.22	23.46	19.90	22.99
DIP-TV	21.33	25.11	22.10	17.96	19.43	19.61	20.89	24.77	23.81	21.57	21.65	23.60	20.46	24.12
Input SNR = 15 dB / $\sigma = 30.02$														
EPLL	23.57	27.04	24.63	21.00	22.10	22.79	23.12	27.21	26.29	23.65	24.51	26.03	22.73	26.78
BM3D	24.02	27.95	24.55	20.96	22.04	22.69	23.41	27.26	26.60	23.71	24.60	26.64	23.34	26.74
TV	22.42	25.39	23.44	19.58	20.99	21.00	22.28	25.49	24.49	22.64	22.93	23.77	22.51	25.22
DIP	23.08	26.17	23.96	20.85	21.24	22.08	22.70	26.89	25.75	22.74	23.69	26.52	22.51	25.32
DIP-TV	23.77	27.37	24.63	21.05	21.85	22.59	23.12	27.33	25.97	22.90	23.95	26.81	23.22	26.65
Input SNR = 20 dB / $\sigma = 14.24$														
EPLL	26.59	29.26	27.35	24.19	24.61	26.04	26.41	30.11	28.78	26.50	27.09	29.19	25.51	29.58
BM3D	26.78	30.20	27.36	24.16	24.61	25.95	26.30	30.13	29.07	26.53	27.14	29.84	26.21	29.55
TV	25.35	27.92	26.18	23.06	23.92	24.34	25.13	28.42	26.99	25.36	25.60	26.94	24.97	30.86
DIP	25.66	29.03	26.77	23.92	23.94	25.45	25.41	29.31	27.49	23.25	25.04	29.59	25.55	28.31
DIP-TV	26.37	29.53	27.38	24.10	24.46	25.66	25.63	29.72	27.84	24.17	25.42	29.80	25.90	29.06
Input SNR = 25 dB / $\sigma = 5.12$														
EPLL	30.01	31.80	30.20	27.75	28.21	29.51	29.51	32.86	31.11	29.58	29.49	32.21	28.46	32.29
BM3D	30.17	32.79	30.17	27.71	28.17	29.39	29.45	32.88	31.38	29.59	29.51	33.00	29.12	32.27
TV	28.84	30.51	29.29	26.82	27.43	27.90	27.81	31.36	29.77	28.45	28.47	30.42	28.24	32.63
DIP	28.33	31.71	29.27	26.86	26.79	28.11	27.99	30.21	27.95	24.67	25.71	31.84	28.45	30.96
DIP-TV	28.75	31.80	29.92	27.42	26.91	28.56	28.17	31.29	28.13	24.86	26.05	32.19	28.49	31.84

Table 2: The average SNR (dB) results of CBM3D, NLM, DIP, and DIP-TV on the testing color images with noise level $\sigma = 25$ 35 45 55 65 75.

Methods	$\sigma = 25$	$\sigma = 35$	$\sigma = 45$	$\sigma = 55$	$\sigma = 65$	$\sigma = 75$
CBM3D	26.98	25.45	24.60	23.79	23.12	22.50
NLM	25.95	24.19	22.97	21.83	20.90	20.15
DIP	26.47	25.36	24.44	23.43	22.64	22.05
DIP-TV	26.71	25.50	24.61	23.86	23.21	22.65

noise from the image. Particularly, our method further enhance the piecewise-smoothness and mitigate the noise of the image, and thus increases the peak-signal-to-noise ratio (PSNR) by over 0.45 dB against DIP. Also note that the aid of TV regularization makes DIP even outperform IRCNN by 0.15 dB on *Peppers*. Table 3 reports the average PSNR comparison with IRCNN and DIP on color and gray scale images, respectively. In general, the improvement by TV regularization outperforms DIP by at least 0.54 dB in terms of PSNR and makes the DIP framework more comparable with IRCNN. For example, DIP-TV is only 0.01 dB lower than IRCNN in terms of the average PSNR on color images, with standard Gaussian blur kernel and $\sigma = 2$.

Table 3: The average PSNR (dB) results of IRCNN, DIP and DIP-TV on the testing gray scale images and color images.

Methods	σ	IRCNN	DIP	DIP-TV
Gaussian blur with standard deviation 1.6				
Gray	2	29.76	28.65	29.44
Color		32.04	31.49	32.03
Kernel 1 (19 × 19 [28])				
Gray	2.55	32.58	31.41	32.11
Color		34.20	33.48	34.09
Gray	7.65	28.59	26.74	27.53
Color		30.89	29.87	30.45

5. CONCLUSION

This work has presented a simple method, namely DIP-TV, to improve the deep image prior framework, leading to promising performance, equivalent to and sometimes surpassing recently published leading alternatives, such as BM3D and IRCNN. The proposed method is based on the recent idea that a CNN model itself can act as a prior on images and improve sparsity promoting priors via the ℓ_1 -norm penalty on the image gradient. The results on images denoising and deblurring demonstrate that TV regularization can further improve on DIP and provides high-quality results.

6. REFERENCES

- [1] L. I. Rudin, S. Osher, and E. Fatemi, "Nonlinear total variation based noise removal algorithms," *Physica D*, vol. 60, no. 1–4, pp. 259–268, November 1992.
- [2] M. A. T. Figueiredo and R. D. Nowak, "Wavelet-based image estimation: An empirical bayes approach using Jeffreys' non-informative prior," *IEEE Trans. Image Process.*, vol. 10, no. 9, pp. 1322–1331, September 2001.
- [3] M. Elad and M. Aharon, "Image denoising via sparse and redundant representations over learned dictionaries," *IEEE Trans. Image Process.*, vol. 15, no. 12, pp. 3736–3745, December 2006.
- [4] A. Danielyan, V. Katkovnik, and K. Egiazarian, "BM3D frames and variational image deblurring," *IEEE Trans. Image Process.*, vol. 21, no. 4, pp. 1715–1728, April 2012.
- [5] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, pp. 436–444, May 28, 2015.
- [6] A. Mousavi, A. B. Patel, and R. G. Baraniuk, "A deep learning approach to structured signal recovery," in *Proc. Allerton Conf. Communication, Control, and Computing*, Allerton Park, IL, USA, September 30–October 2, 2015, pp. 1336–1343.
- [7] K. H. Jin, M. T. McCann, E. Froustey, and M. Unser, "Deep convolutional neural network for inverse problems in imaging," *IEEE Trans. Image Process.*, vol. 26, no. 9, pp. 4509–4522, Sept. 2017.
- [8] Y. S. Han, J. Yoo, and J. C. Ye, "Deep learning with domain adaptation for accelerated projection reconstruction MR," 2017, arXiv:1703.01135 [cs.CV].
- [9] Y. Sun, Z. Xia, and U. S. Kamilov, "Efficient and accurate inversion of multiple scattering with deep learning," *Opt. Express*, vol. 26, no. 11, pp. 14678–14688, May 2018.
- [10] D. Lee, J. Yoo, S. Tak, and J. C. Ye, "Deep residual learning for accelerated MRI using magnitude and phase networks," *IEEE Trans. Biomed. Eng.*, vol. 65, no. 9, pp. 1985–1995, Sept. 2018.
- [11] T. Meinhardt, M. Moeller, C. Hazirbas, and D. Cremers, "Learning proximal operators: Using denoising networks for regularizing inverse imaging problems," in *Proc. IEEE Int. Conf. Comp. Vis. (ICCV)*, Venice, Italy, October 22–29, 2017, pp. 1799–1808.
- [12] K. Zhang, W. Zuo, S. Gu, and L. Zhang, "Learning deep CNN denoiser prior for image restoration," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [13] Y. Romano, M. Elad, and P. Milanfar, "The little engine that could: Regularization by denoising (RED)," *SIAM J. Imaging Sci.*, vol. 10, no. 4, pp. 1804–1844, 2017.
- [14] Y. Sun, B. Wohlberg, and U. S. Kamilov, "An online plug-and-play algorithm for regularized image reconstruction," 2018, arXiv:1809.04693 [cs.CV].
- [15] D. Ulyanov, A. Vedaldi, and V. Lempitsky, "Deep image prior," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, Salt Lake City, UT, USA, June 18–22, 2018, pp. 9446–9454.
- [16] D. Van Veen, A. Jalal, E. Price, S. Vishwanath, and A. G. Dimakis, "Compressed sensing with deep image prior and learned regularization," 2018, arXiv:1806.06438 [stat.ML].
- [17] K. Dabov, A. Foi, V. Katkovnik, and K. Egiazarian, "Image denoising by sparse 3-D transform-domain collaborative filtering," *IEEE Trans. Image Process.*, vol. 16, no. 16, pp. 2080–2095, August 2007.
- [18] L. I. Rudin, S. Osher, and E. Fatemi, "Nonlinear total variation based noise removal algorithms," *Physica D: nonlinear phenomena*, vol. 60, no. 1–4, pp. 259–268, 1992.
- [19] M. Persson, D. Bone, and H. Elmqvist, "Total variation norm for three-dimensional iterative reconstruction in limited view angle tomography," *Phys. Med. Biol.*, vol. 46, no. 3, pp. 853–866, 2001.
- [20] M. Lustig, D. L. Donoho, and J. M. Pauly, "Sparse MRI: The application of compressed sensing for rapid MR imaging," *Magn. Reson. Med.*, vol. 58, no. 6, pp. 1182–1195, December 2007.
- [21] U. S. Kamilov, I. N. Papadopoulos, M. H. Shoreh, A. Goy, C. Vonesch, M. Unser, and D. Psaltis, "Optical tomographic image reconstruction based on beam propagation and sparse regularization," *IEEE Transactions on Computational Imaging*, vol. 2, no. 1, pp. 59–70, 2016.
- [22] M. Egmont-Petersen, D. de Ridder, and H. Handels, "Image processing with neural networks—a review," *Pattern recognition*, vol. 35, no. 10, pp. 2279–2301, 2002.
- [23] J. Xie, L. Xu, and E. Chen, "Image denoising and inpainting with deep neural networks," in *Advances in neural information processing systems*, 2012, pp. 341–349.
- [24] K. Zhang, W. Zuo, Y. Chen, D. Meng, and L. Zhang, "Beyond a Gaussian denoiser: Residual learning of deep CNN for image denoising," *IEEE Trans. Image Process.*, vol. 26, no. 7, pp. 3142–3155, July 2017.
- [25] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," *arXiv preprint arXiv:1412.6980*, 2014.
- [26] O. Ronneberger, P. Fischer, and T. Brox, "U-net: Convolutional networks for biomedical image segmentation," in *International Conference on Medical image computing and computer-assisted intervention*. Springer, 2015, pp. 234–241.
- [27] K. H. Jin, M. T. McCann, E. Froustey, and M. Unser, "Deep convolutional neural network for inverse problems in imaging," *IEEE Transactions on Image Processing*, vol. 26, no. 9, pp. 4509–4522, 2017.
- [28] A. Levin, Y. Weiss, F. Durand, and W. T. Freeman, "Understanding and evaluating blind deconvolution algorithms," in *2009 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2009.
- [29] D. Zoran and Y. Weiss, "From learning models of natural image patches to whole image restoration," in *Computer Vision (ICCV), 2011 IEEE International Conference on*. IEEE, 2011, pp. 479–486.
- [30] A. Beck and M. Teboulle, "Fast gradient-based algorithms for constrained total variation image denoising and deblurring problems," *IEEE Transactions on Image Processing*, vol. 18, no. 11, pp. 2419–2434, 2009.
- [31] A. Buades, B. Coll, and J. Morel, "A non-local algorithm for image denoising," in *Computer Vision and Pattern Recognition (CVPR), 2005. IEEE Computer Society Conference on*. IEEE, 2005, vol. 2, pp. 60–65.