

Leveraging Phase Information to Boost Unrolled Network Learning for Image Deblurring

Samira Malek¹, Haichuan Zhang¹, Chul Lee², and Vishal Monga¹

¹ Pennsylvania State University, University Park, PA 16802, USA
{sxm6547,hzz5333,vum4}@psu.edu

² Dongguk University, Seoul 04620, South Korea
chullee@dongguk.edu

Abstract. While most image deblurring techniques directly restore the spatial image variable, we propose an amplitude and phase decomposition recognizing the importance of accurate phase estimation in recovering sharp image details. To that end, we first develop novel linear minimum mean squared (LMMSE) estimators of the amplitude and phase of the blurred, noisy image observation. An iterative optimization algorithm follows that recovers the sharp image using the aforementioned LMMSE estimators. Finally, matrix parameters that are statistically determined and fixed in the iterative algorithm are now learned using a training dataset of clean and degraded observations. Our deblurring engine is dubbed UPADNet – Unrolled Phase and Amplitude Decomposition Network, such that each iteration of the underlying phase and amplitude recovery algorithm is parameterized and trained end-to-end. Experiments over benchmark evaluation datasets such as GoPro, Real-Blur and COCO datasets confirm that UPADNet outperforms state of the art deep networks including those based on algorithm unrolling in the image domain. The benefits of UPADNet are even more pronounced in high noise and limited training data regimes.

Keywords: Image Deblurring · Image restoration · Fourier Phase and Amplitude Decomposition · Phase Unrolling · Deep Unrolling

1 Introduction

Image deblurring, a sub-problem of image restoration, aims to recover a sharp image from a blurred observation. In real-world imaging, blur artifacts arise from various sources such as atmospheric turbulence, diffraction, optical defocusing, and camera motion [22].

Various approaches have been developed to address blind image deblurring, one of which is leveraging mathematical modeling and statistical estimation techniques [22]. Some methods focus on estimating the blur kernel using sharp edge prediction [17] and sparsity constraints in the image domain [43, 67]. A common strategy has been to incorporate strong regularizations, such as ℓ_0 -norm sparsity [43, 67], normalized sparsity [21], and total variation constraints [46], to

mitigate the ill-posed nature of the problem. Several studies have explored novel priors, including patch priors [57], dark channel priors [44], and extreme channel priors [68], which exploit natural image statistics to enhance sharp image estimation. Bayesian inference has also been widely used, employing marginalization over the high-dimensional image space [65] and general sparse image priors [2] to improve robustness. A notable contribution in this line is [55], which introduced a unified probabilistic framework for joint blur kernel estimation and image restoration. Framelet-based methods [4] offer alternative strategies by leveraging transform-domain representations. Some works analyze the fundamental limitations of existing marginal likelihood optimization approaches and provide a broader theoretical evaluation of blind deconvolution algorithms, introducing efficient MAP-based methods [27, 28]. In [14], an iterative refinement framework was proposed that categorizes kernel types by alternating between kernel and image updates to achieve improved results. A modified Newton integration algorithm is presented from a control perspective, offering noise tolerance and fast convergence [32]. Additionally, in [9], Polyblur was introduced as a non-iterative method for removing mild blur by leveraging polynomial reblurring, which offers practical benefits for lightly degraded images.

Recent advances in image deblurring have largely focused on data-driven approaches tailored to specific datasets. Deep neural networks have shown remarkable success in this domain [25, 54]. Convolutional neural networks such as [8, 40, 58] adopt multi-scale architectures with stacked sub-networks, progressively enhancing sharpness in a coarse-to-fine manner. To mitigate the computational burden of multi-scale upsampling and the diminishing returns of deeper fine-scale networks, the Deep Multi-Patch Hierarchical Network (DMPHN) was proposed in [71]. Similarly, MPRNet [70] leverages multi-stage hierarchical refinement, while MAXIM [61] introduces an MLP-based backbone for deblurring. NAFNet [7] further demonstrates that strong performance can be achieved even without nonlinear activations. Transformer-based architectures, including Restormer [69], Uformer [63], and Stripformer [60], improve performance by modeling long-range dependencies via attention mechanisms. Frequency-domain approaches such as MRLPFNet [11], the efficient frequency transformer [19], and frequency selection analysis in [38] aim to better integrate spectral and spatial information for improved high-frequency detail preservation.

In parallel, the use of generative and adversarial models represents another line of approach. DeblurGAN [23] employs an end-to-end framework based on a conditional GAN with content loss, while CGAN [34] adopts region focused strategies. Stochastic refinement-based approaches have also emerged, notably probabilistic deblurring [64]. Additionally, recent work has explored joint tasks of depth estimation and restoration from defocused images using deep networks, as demonstrated in [41], pixel-wise kernel prediction for blind motion deblurring [5], and enhanced defocus blur estimation using multi-interactive strategies [29]. AIFNet [50] contributes to this direction by using refocusing techniques and light field synthetic aperture. CodEx [36] features a modular framework that jointly addresses temporal deblurring and tomographic reconstruction. These

developments underline a growing trend toward hybrid models that integrate architectural innovations, spectral priors, and task specific insights to enhance restoration fidelity, robustness, and generalization.

Despite these advancements, fully deep neural network-based methods often suffer from key limitations: they require large amounts of annotated data for effective training, have limited generalizability across diverse imaging conditions, and demand extensive computational resources and training time, which can hinder their deployment in real-world and resource-constrained applications. Recently, algorithm unrolling [39, 52, 56] has gained attention as an effective strategy for addressing image deconvolution. This approach combines the interpretability of traditional iterative optimization with the adaptability of deep learning by unfolding an iterative algorithm—such as Half-Quadratic Splitting (HQS) [30, 31]—into a learnable network where key parameters, like filters or step sizes, are trained. Unrolling methods have been applied successfully to various deblurring contexts, including networks based on iterative variational Bayesian methods [15], second-order (Newton) updates for image restoration [10], photon-limited scenarios [53], non-uniform blind deblurring [48], and low-light image enhancement in joint frameworks [62]. Additionally, kernel-based techniques like the Gaussian Kernel Mixture Network enhance defocus deblurring by embedding blur priors into deep models [47].

Even with significant advances in deep and unrolled networks for image deblurring, the problem remains fundamentally ill-posed, making reliable high-quality restoration under diverse real-world degradations challenging. Accurate phase recovery is particularly critical, as phase encodes the structural and geometric information of an image [33, 42]. Prior work [45] has demonstrated the importance of phase-only representations for kernel initialization in image deblurring. Motivated by these insights, we propose a principled phase–amplitude decomposition framework that explicitly models and estimates both components in the Fourier domain. We first derive novel **linear minimum mean squared error (LMMSE) estimators** for the amplitude and phase of blurred, noisy observations. An iterative optimization algorithm is then developed to recover the latent sharp image using these estimators. Unlike classical formulations where statistical parameters are fixed, we learn the underlying matrices directly from paired clean and degraded data. This leads to **UPADNet** (Unrolled Phase and Amplitude Decomposition Network), where each iteration of the phase–amplitude recovery algorithm is parameterized and trained end-to-end. Extensive experiments on GoPro and RealBlur demonstrate that UPADNet **outperforms state-of-the-art** deep and unrolled methods. Moreover, our model exhibits superior **robustness under high noise**, camera motion, and atmospheric degradations, and maintains strong performance in low-data training regimes, highlighting its generalization capability. The code is available at [GitHub](#).

Notation: Unless stated otherwise, uppercase letters (*e.g.*, A) represent matrices throughout the paper. The symbol $*$ denotes circular convolution, and $\circ$ is the Hadamard (element-wise) product. The element in the i -th row and j -th

column of a matrix A is denoted by A_{ij} . The notation $\frac{A}{B}$ denotes element-wise division between matrices A and B , that is, $(\frac{A}{B})_{ij} = \frac{A_{ij}}{B_{ij}}$, for all valid indices (i, j) . Let the Fourier transform of x be denoted by $\mathcal{F}\{x\} = X = A_X e^{j\theta_X}$, where A_X and θ_X are the amplitude and phase of X , respectively.

2 LMMSE Estimation of Blurred Phase and Amplitude

Image deblurring is classically modeled as a linear shift-invariant convolution process corrupted by additive noise:

$$z = h * u + n \quad (1)$$

where u is the latent sharp image, h the blur kernel, and n additive noise. This forward model is standard in computational imaging and inverse problems [6, 22, 26]. In the Fourier domain, convolution becomes element-wise multiplication:

$$Z = H \circ U + N \quad (2)$$

Let $Z = A_Z e^{j\theta_Z}$ denote the polar form of the Fourier transform. It has long been known that the phase encodes structural information while the amplitude encodes energy distribution [33, 42].

A naive yet commonly used estimation of the Fourier-domain amplitude and phase of the observed blurry image Z is obtained directly from the convolution property in the frequency domain. Specifically, let A_H , θ_H and A_U , θ_U denote the amplitudes and phases of the blur kernel H and the latent sharp image U , respectively. Then, the naive estimates of the amplitude $\tilde{A}_Z$ and phase $\tilde{\theta}_Z$ of Z are given by:

$$\tilde{A}_Z = A_H \circ A_U, \quad \tilde{\theta}_Z = \theta_H + \theta_U. \quad (3)$$

In noiseless convolution case, *i.e.*, $N = 0$, we have $A_Z = \tilde{A}_Z$ and $\theta_Z = \tilde{\theta}_Z$. However, in realistic settings with noise and kernel uncertainty, these equalities no longer hold. Further, in recent investigations, the image blur is known to depart from the standard convolution model [40, 49]. To address these challenges and motivated by linear minimum mean squared error (LMMSE) estimation theory [18], we aim to estimate A_Z and θ_Z via generalized linear operators:

$$\hat{A}_Z = \mathcal{W}_1 \circ A_H \circ A_U + \mathcal{W}_2, \quad \hat{\theta}_Z = \mathcal{W}_3 \circ \theta_H + \mathcal{W}_4 \circ \theta_U + \mathcal{W}_5 \quad (4)$$

such that these linear estimations minimize the squared errors:

$$E[\|A_Z - \hat{A}_Z\|_2^2] \quad \text{and} \quad E[\|\theta_Z - \hat{\theta}_Z\|_2^2]. \quad (5)$$

The following theorems characterize the optimal element-wise linear estimators.

Theorem 1 (Amplitude Estimation). *Let A_Z and $A_H \circ A_U$ be random matrices with finite means and variances.*

$$\min_{\mathcal{W}_1, \mathcal{W}_2} E[\|A_Z - \hat{A}_Z\|_F^2], \quad \text{s.t.} \quad \hat{A}_Z = \mathcal{W}_1 \circ A_H \circ A_U + \mathcal{W}_2 \quad (6)$$

is minimized when:

$$[\mathcal{W}_1^*]_{ij} = \frac{\text{Cov}([A_Z]_{ij}, [A_H]_{ij}[A_U]_{ij})}{\text{Var}([A_H]_{ij}[A_U]_{ij})}, \quad (7)$$

$$[\mathcal{W}_2^*]_{ij} = \mathbb{E}([A_Z]_{ij}) - [\mathcal{W}_1^*]_{ij} \mathbb{E}([A_H]_{ij}[A_U]_{ij}) \quad (8)$$

Proof: See supplementary material Section 1.

Theorem 2 (Phase Estimation). Let θ_Z , θ_H , and θ_U be random matrices with finite means and variances. Consider:

$$\min_{\mathcal{W}_3, \mathcal{W}_4, \mathcal{W}_5} \mathbb{E}[\|\theta_Z - \hat{\theta}_Z\|_F^2], \quad \text{s.t. } \hat{\theta}_Z = \mathcal{W}_3 \circ \theta_H + \mathcal{W}_4 \circ \theta_U + \mathcal{W}_5. \quad (9)$$

The minimizer satisfies, for every entry (i, j) :

$$[\mathcal{W}_3^*]_{i,j} = \frac{\sigma_{U,i,j}^2 \gamma_{H,i,j} - \sigma_{HU,i,j} \gamma_{U,i,j}}{\Delta_{i,j}}, \quad (10)$$

$$[\mathcal{W}_4^*]_{i,j} = \frac{-\sigma_{HU,i,j} \gamma_{H,i,j} + \sigma_{H,i,j}^2 \gamma_{U,i,j}}{\Delta_{i,j}}, \quad (11)$$

$$[\mathcal{W}_5^*]_{i,j} = \mu_{Z,i,j} - [\mathcal{W}_3^*]_{i,j} \mu_{H,i,j} - [\mathcal{W}_4^*]_{i,j} \mu_{U,i,j}, \quad (12)$$

where:

$$\begin{aligned} \sigma_{H,i,j}^2 &= \text{Var}([\theta_H]_{i,j}), & \sigma_{U,i,j}^2 &= \text{Var}([\theta_U]_{i,j}), \\ \sigma_{HU,i,j} &= \text{Cov}([\theta_H]_{i,j}, [\theta_U]_{i,j}), & \gamma_{H,i,j} &= \text{Cov}([\theta_Z]_{i,j}, [\theta_H]_{i,j}), \\ \gamma_{U,i,j} &= \text{Cov}([\theta_Z]_{i,j}, [\theta_U]_{i,j}), & \Delta_{i,j} &= \sigma_{H,i,j}^2 \sigma_{U,i,j}^2 - \sigma_{HU,i,j}^2, \\ \mu_{Z,i,j} &= \mathbb{E}([\theta_Z]_{i,j}), & \mu_{H,i,j} &= \mathbb{E}([\theta_H]_{i,j}), \\ \mu_{U,i,j} &= \mathbb{E}([\theta_U]_{i,j}). \end{aligned}$$

Proof: See supplementary material Section 2.

Based on the aforementioned theorems, the proposed estimators achieve a lower expected MSE than the naive estimators, since the latter are a special case of the former. This result is further validated through experiments. Specifically, we evaluate the naive and proposed LMMSE estimators and verify that:

$$E[\|A_Z - \tilde{A}_Z\|_2^2] \geq E[\|A_Z - \hat{A}_Z\|_2^2], \quad E[\|\theta_Z - \tilde{\theta}_Z\|_2^2] \geq E[\|\theta_Z - \hat{\theta}_Z\|_2^2] \quad (13)$$

To this end, we generate 600 blurred images by convolving 200 BSDS images [1] with three kernels from [30] and adding Gaussian noise ($\sigma = 0.05$) to empirically estimate the statistical parameters $\{\mathcal{W}_i\}_{i=0}^5$ in (4). For evaluation, we use 28 images (14 images with two unseen kernels and additive Gaussian noise with $\sigma = 0.05$). One example of the ground-truth image, blur kernel, and the resulting blurred image is shown in Figs. 1a and 1b. Figs. 1c and 1d show the amplitude and phase errors, respectively. The proposed LMMSE estimators consistently

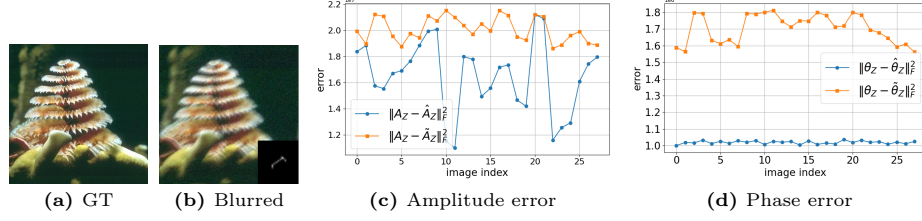


Fig. 1: Estimation of amplitude and phase components of the blurred observation. (a) Ground-truth sharp image. (b) Blurred observation generated by convolving the sharp image with a known kernel (bottom right of the blurred image) and adding noise. (c) and (d) Comparison of the estimation errors of the naive and proposed LMMSE estimators for amplitude and phase on the test set. The proposed LMMSE estimators consistently produce lower estimation errors.

achieve lower error than the naive multiplicative/additive models. This experimentally confirms that although convolution is multiplicative in the Fourier domain, the induced amplitude and phase statistics under noise deviate from the naive relations, and the proposed estimators significantly improves estimation of amplitude and phase.

3 Image Deblurring Via Phase and Amplitude Optimization

We estimate the latent sharp image by minimizing a regularized objective that enforces statistical consistency with the linear phase and amplitude predictors derived from Eq. (4) in Sec. 2:

$$\begin{aligned}
 \min_{A_U, A_H, \theta_U, \theta_H} \quad & \underbrace{\|A_Z - (\mathcal{W}_1 \circ A_H \circ A_U + \mathcal{W}_2)\|_F^2 + \alpha \|A_H\|_F^2 + \beta \|M \circ A_U\|_{1,1}}_{\text{Amplitude loss}} \\
 & + \underbrace{\|\theta_Z - (\mathcal{W}_3 \circ \theta_H + \mathcal{W}_4 \circ \theta_U + \mathcal{W}_5)\|_F^2 + \gamma \|\theta_H\|_F^2 + \mu \|\theta_U\|_{1,1}}_{\text{Phase loss}}
 \end{aligned} \tag{14}$$

The first two terms of amplitude and phase losses enforce amplitude and phase data fidelity under the statistically optimal linear model, mitigating the errors introduced by noise and kernel uncertainty in naive Fourier inversion [26]. Solving this objective yields (A_U, θ_U) from which the restored image is obtained via the inverse Fourier transform. The ℓ_1 regularization on $M \circ A_U$ promotes sparsity in the Fourier-domain representation of natural images, with the mask M acting as a frequency-domain filter [16, 51]. In contrast, the blur kernel is expected to be smooth; therefore, an ℓ_2 penalty is imposed on θ_H and A_H , consistent with classical regularization for ill-posed inverse problems [3, 59]. Moreover, incorporating an ℓ_1 regularization term on θ_U enforces structural symmetry in the

optimization updates, which is crucial for the unrolled architecture introduced in Sec. 4. This mixed ℓ_1 - ℓ_2 formulation reflects the structural asymmetry between images (sparse and edge-dominated) and kernels (smooth and compact), stabilizing blind deblurring while preserving high-frequency detail.

The objective in Eq. (14) is nonconvex. Following classical blind deconvolution strategies [22, 26], we adopt an alternating optimization scheme that updates one variable at a time while fixing the others. This decomposes the problem into a sequence of simpler subproblems that admit closed-form or proximal solutions. Moreover, the ℓ_1 regularization terms are non-smooth and therefore non-differentiable, which prevents direct closed-form minimization of Eq. (14). To address this, we employ HQS [13], introducing auxiliary variables S and P to decouple the non-smooth penalties from the quadratic data-fidelity terms. This leads to six subproblems per iteration:

$$A_H^{(t+1)} = \arg \min_{A_H} \|A_Z - \mathcal{W}_1 \circ A_H \circ A_U^{(t)} - \mathcal{W}_2\|_2^2 + \alpha \|A_H\|_2^2, \quad (15)$$

$$S^{(t+1)} = \arg \min_S \zeta \|M \circ S\|_1 + \beta \|A_U^{(t)} - S\|_2^2, \quad (16)$$

$$A_U^{(t+1)} = \arg \min_{A_U} \|A_Z - \mathcal{W}_1 \circ A_H^{(t+1)} \circ A_U - \mathcal{W}_2\|_2^2 + \beta \|A_U - S^{(t+1)}\|_2^2, \quad (17)$$

$$\theta_H^{(t+1)} = \arg \min_{\theta_H} \|\theta_Z - \mathcal{W}_3 \circ \theta_H - \mathcal{W}_4 \circ \theta_U^{(t)} - \mathcal{W}_5\|_2^2 + \gamma \|\theta_H\|_2^2, \quad (18)$$

$$P^{(t+1)} = \arg \min_P \xi \|P\|_1 + \mu \|\theta_U^{(t)} - P\|_2^2, \quad (19)$$

$$\theta_U^{(t+1)} = \arg \min_{\theta_U} \|\theta_Z - \mathcal{W}_3 \circ \theta_H^{(t+1)} - \mathcal{W}_4 \circ \theta_U - \mathcal{W}_5\|_2^2 + \mu \|\theta_U - P^{(t+1)}\|_2^2. \quad (20)$$

Each quadratic subproblem admits a closed-form solution, whereas the ℓ_1 -regularized auxiliary updates correspond to soft-thresholding proximal operators.³ The resulting closed-form update rules are summarized in Algorithm 1.

4 Unrolled Phase and Amplitude Decomposition Image Deblurring Network (UPADNet)

Since the optimal linear estimators $\{\mathcal{W}_i\}_{i=1}^5$ depend on unknown image and blur statistics, they are not directly available in practical blind deblurring scenarios. Moreover, fixed analytical parameters may be suboptimal under complex real-world degradations. Therefore, we learn these quantities from data and unroll the optimization procedure into a trainable architecture to improve robustness and reconstruction performance. Specifically, we propose the *Unrolled Phase and Amplitude Decomposition Image Deblurring Network (UPADNet)*, a model-driven deep architecture derived from the iterative optimization procedure in

³ The soft-thresholding operator is defined element-wise as $\text{soft}(x, \lambda) = \text{sign}(x) \max(|x| - \lambda, 0)$.

Algorithm 1 Phase and Amplitude Decomposition Image Deblurring Algorithm**Input:** $A_Z, \theta_Z, \{\mathcal{W}_i\}_{i=1}^5$, regularization weights $\alpha, \beta, \gamma, \mu, \xi, \zeta$ 1: **Initialize:** $A_U^{(0)}, A_H^{(0)}, \theta_U^{(0)}, \theta_H^{(0)}$ 2: **for** $t = 0$ to T **do** **Amplitude updates:**

3:
$$A_H^{(t+1)} = \frac{(\mathcal{W}_1 \circ A_U^{(t)}) \circ (A_Z - \mathcal{W}_2)}{(\mathcal{W}_1 \circ A_U^{(t)})^2 + \alpha}$$

4:
$$S^{(t+1)} = \text{soft}\left(A_U^{(t)}, \frac{\zeta \circ |M|}{2\beta}\right)$$

5:
$$A_U^{(t+1)} = \frac{(\mathcal{W}_1 \circ A_H^{(t+1)}) \circ (A_Z - \mathcal{W}_2 + \beta S^{(t+1)})}{(\mathcal{W}_1 \circ A_H^{(t+1)})^2 + \beta}$$

Phase updates:

6:
$$\theta_H^{(t+1)} = \frac{\mathcal{W}_3 \circ (\theta_Z - (\mathcal{W}_4 \circ \theta_U^{(t)} + \mathcal{W}_5))}{(\mathcal{W}_3)^2 + \gamma}$$

7:
$$P^{(t+1)} = \text{soft}\left(\theta_U^{(t)}, \frac{\xi}{2\mu}\right)$$

8:
$$\theta_U^{(t+1)} = \frac{\mathcal{W}_4 \circ (\theta_Z - (\mathcal{W}_3 \circ \theta_H^{(t+1)} + \mathcal{W}_5) + \mu P^{(t+1)})}{(\mathcal{W}_4)^2 + \mu}$$

9: **end for****Output:** $A_U^{(T+1)}, A_H^{(T+1)}, \theta_U^{(T+1)}, \theta_H^{(T+1)}$

Algorithm 1. Each iteration of the underlying algorithm is unfolded into a learnable module, referred to as a UPADBlock (see Fig. 2a), where the parameters $\{\mathcal{W}_i^t\}_{i=1}^5$ and filter M^t are learned from data. This design preserves the interpretability of the original optimization scheme while enabling data-adaptive refinement. Each UPADBlock consists of four sub-blocks that sequentially update the amplitude and phase components of the latent image and blur kernel, *i.e.*, A_U, A_H, θ_U , and θ_H . This decomposition explicitly leverages the complementary roles of amplitude (energy distribution) and phase (structural information) in the Fourier domain.

Since A_U and A_H are non-negative and $\theta_U, \theta_H \in [0, 2\pi]$, we enforce positivity and smoothness through a nonlinear activation function (GELU) at the end of each sub-block. Moreover, the update $S^{(t+1)}$ corresponds to a soft-thresholding operator, which can be simplified under the constraint $A_U \geq 0$:

$$\begin{aligned} S^{(t+1)} &= \text{soft}\left(A_U^{(t)}, \frac{\zeta \circ |M^t|}{2\beta}\right) = \text{sign}(A_U^{(t)}) \max\left(|A_U^{(t)}| - \frac{\zeta \circ |M^t|}{2\beta}, 0\right) \\ &= \text{ReLU}\left(A_U^{(t)} - \frac{\zeta \circ |M^t|}{2\beta}\right) \\ &\approx \text{GELU}\left(A_U^{(t)} - \frac{\zeta \circ |M^t|}{2\beta}\right). \end{aligned} \quad (21)$$

In (21), to improve gradient flow and training stability, we replace ReLU with GELU, which provides a smoother approximation to the proximal operator. Similarly, the phase update is implemented as:

$$P^{(t+1)} \approx \text{GELU}\left(\theta_U^{(t)} - \frac{\xi}{2\mu}\right). \quad (22)$$

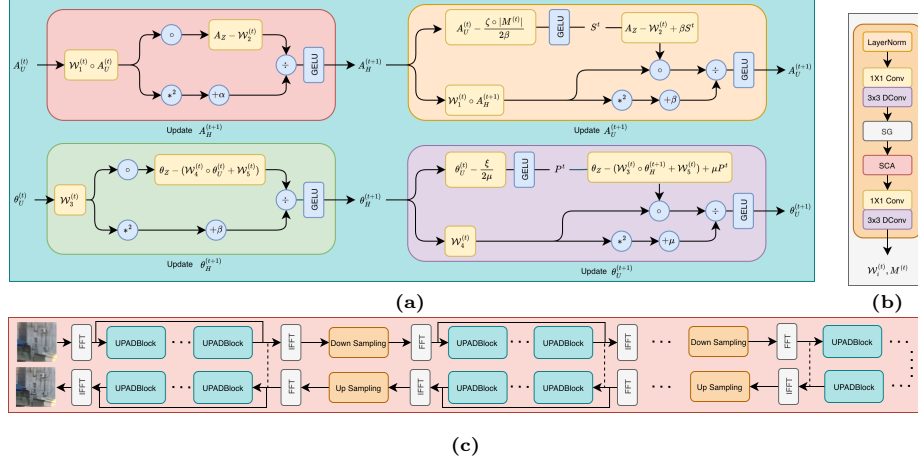


Fig. 2: Overview of the proposed UPADNet architecture. (a) UPADBlock: one unrolled iteration derived from Algorithm 1, consisting of four sub-blocks that update A_U , A_H , θ_U , and θ_H with learnable parameters. (b) Weight Generator: dynamically produces iteration-specific weights $\{\mathcal{W}_i^t\}_{i=1}^5$ and filter M^t . (c) UPADNet: the complete multi-scale network obtained by stacking multiple UPADBlocks with intermediate downsampling and upsampling for blind deblurring.

To estimate the unknown parameters in each iteration and enhance flexibility across iterations, we introduce a *Weight Generator* module (see Fig. 2b) that dynamically produces the weights $\{\mathcal{W}_i^t\}$ and filter M^t at each stage. The module incorporates the SimpleGate (SG) and Simplified Channel Attention (SCA) mechanisms proposed in [7], enabling lightweight yet expressive parameter modulation. This design ensures compatibility between generated weights and the intermediate feature dimensions while maintaining computational efficiency.

As each UPADBlock represents one iteration of Algorithm 1, stacking multiple UPADBlocks forms the complete UPADNet architecture. To further improve the representation power and enlarge the effective receptive field, we adopt a multi-scale strategy: intermediate features are progressively downsampled after several blocks, and corresponding upsampling layers are applied toward the output stage (see Fig. 2c).

5 Experiments

We evaluate **UPADNet** on three widely used image deblurring benchmarks. **GoPro** [40] is a large-scale synthetic dataset generated via frame accumulation, which implicitly includes object motion blur and complex motion trajectories, where the effective blur kernel is difficult to approximate using a simple parametric model. **RealBlur-R** and **RealBlur-J** [49] are real-world blur datasets with carefully aligned ground-truth images, primarily capturing low-light cam-

Table 1: Quantitative comparison on GoPro [40] and RealBlur [49].

Method	GoPro		RealBlur-R		RealBlur-J		Params (M)
	PSNR $\uparrow$	SSIM $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$	
DeblurGAN-v2 [24]	29.55	0.934	36.44	0.9347	29.69	0.8703	60.9
DUBLID [30]	29.96	0.940	35.35	0.925	28.42	0.860	N/A
DeepSN-Net [10]	32.83	0.960	36.22	0.957	29.30	0.895	N/A
Stripformer [60]	33.08	0.962	39.84	0.9737	32.48	0.9290	20.0
NAFNet [7]	33.69	0.967	35.84	0.952	27.94	0.854	67.8
UFPNet [12]	34.06	0.968	39.84	0.974	32.48	0.929	80.3
FFTformer [19]	34.21	0.968	40.11	0.9753	32.62	0.9326	16.6
AdaRevD [37]	34.60	0.972	36.53	0.957	30.12	0.894	220.0
EGDeblurring [66]	34.62	0.973	N/A	N/A	N/A	N/A	80.3
EVSSM [20]	34.51	0.9713	41.27	0.977	34.34	0.945	17.1
UPADNet	34.63	0.975	41.29	0.979	34.35	0.946	27.0

era shake scenarios. To further investigate the behavior of UPADNet under specific degradations, *i.e.*, camera motion and mild atmospheric turbulence effects consistent with the convolutional model in Eq. (1), we additionally construct a synthetic benchmark using **COCO** [35]. Specifically, sharp images are convolved with real-world motion kernels extracted from [30] to simulate motion-induced blur, followed by the addition of Gaussian noise to model sensor corruption. In total, 6000 images are generated for training and 1500 images for testing. This setup allows us to systematically evaluate robustness under controlled camera motion and noise conditions and to compare performance against both unrolled optimization-based networks and purely deep learning-based architectures, particularly in low-training or noise-corrupted regimes.

UPADNet is trained end-to-end using the AdamW optimizer with standard data augmentation strategies, including random cropping, horizontal and vertical flipping, and rotation. Training is performed on 256×256 patches, while evaluation is conducted on full-resolution images. The training-test split for the GoPro and RealBlur datasets follow those used in prior works [20, 37]. Quantitative performance is reported using the widely adopted PSNR and SSIM metrics.

Table 1 reports quantitative comparisons on the GoPro dataset against state-of-the-art deblurring methods. UPADNet achieves the best performance, reaching **34.63 dB** PSNR and **0.975** SSIM, outperforming recent transformer-based and diffusion-based approaches. Notably, our model requires only **27M** parameters, substantially fewer than large-scale models such as AdaRevD (220.8M) and HINet (88.7M), demonstrating a favorable accuracy-efficiency trade-off. Visual comparisons are provided in Fig. 3. UPADNet more effectively restores fine structures and high-frequency textures, particularly in regions containing sharp edges and repetitive patterns. The explicit phase-amplitude modeling enables the preservation of structural consistency while mitigating ringing artifacts. Table 1 also presents results on the RealBlur dataset. UPADNet achieves the best

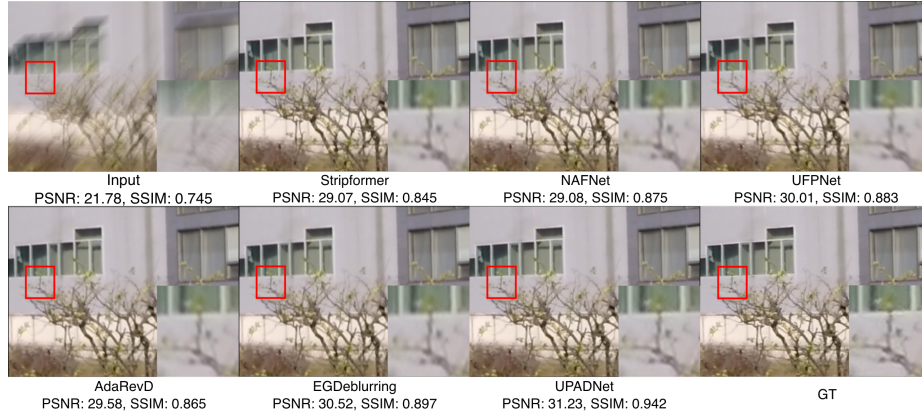


Fig. 3: Qualitative results on the GoPro dataset.

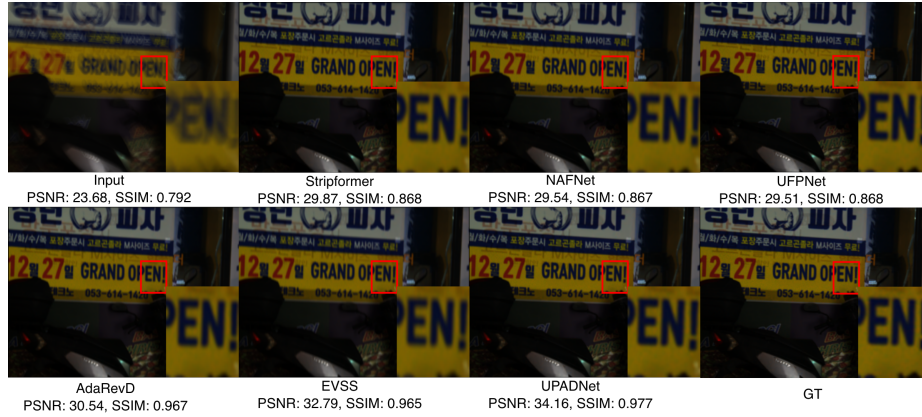
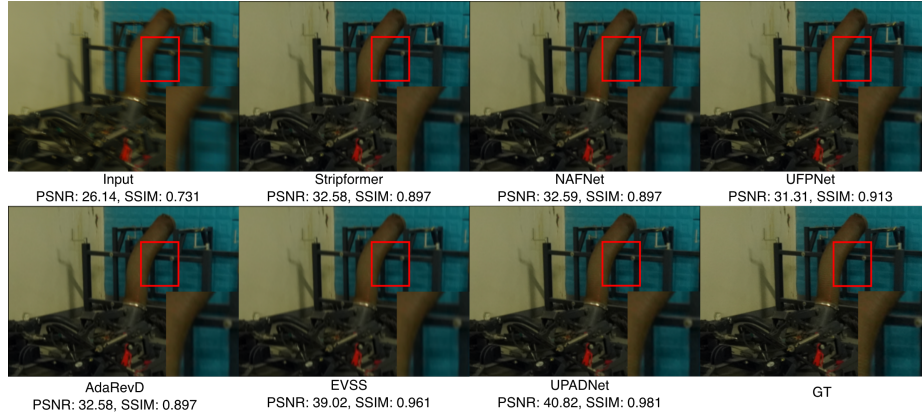


Fig. 4: Qualitative results on the RealBlur-R dataset.

performance on both RealBlur-R and RealBlur-J, attaining **41.29/0.979** and **34.35/0.946** (PSNR/SSIM), respectively. The consistent improvements across both synthetic (GoPro) and real-world datasets demonstrate strong generalization capability. The performance gains on real data further suggest that the explicit Fourier-domain phase-amplitude decomposition enhances robustness to real camera shake and non-uniform blur, which are not perfectly captured by synthetic training data.

Noise Robustness (UPADNet vs. State of the Art): To evaluate robustness under noisy conditions, we generate blurred images from COCO and add Gaussian noise with increasing variance ($\sigma^2 = 0.01, 0.03, 0.05$). Quantitative results are summarized in Table 2. UPADNet consistently outperforms competing methods across all noise levels. While other state-of-the-art unrolled and deep models exhibit significant performance degradation as noise increases, UPAD-

**Fig. 5:** Qualitative results on the RealBlur-J dataset.**Table 2:** Quantitative comparison with unrolled and deep network models under increasing levels of additive Gaussian noise on COCO.

Method	$\sigma = 0.01$		$\sigma = 0.03$		$\sigma = 0.05$	
	PSNR $\uparrow$	SSIM $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$
DUBLID [30]	29.96	0.881	27.42	0.842	26.88	0.801
Net-A [48]	27.20	0.835	26.63	0.798	25.94	0.756
Stripformer [60]	31.84	0.921	30.72	0.902	29.95	0.874
AdaRevD [37]	32.10	0.925	30.78	0.906	30.11	0.879
EVSSM [20]	31.62	0.918	30.41	0.897	30.01	0.872
UPADNet	32.58	0.932	31.47	0.914	30.86	0.891

**Fig. 6:** Deblurred results on COCO with $\sigma^2 = 0.05$.

Net maintains stable performance, demonstrating strong noise robustness. This robustness can be attributed to the structured phase-amplitude decomposition and the linear estimation-inspired updates. Fig. 6 further illustrates qualitative comparisons. UPADNet produces sharper reconstructions, especially in textured regions. Moreover, we compare UPADNet with Net-A and DUBLID, both of which are unrolled networks derived from similar optimization frameworks but without explicit phase-amplitude decomposition.

Table 3: Comparison under reduced training data ratios (100%, 70%, 60%) for COCO dataset.

Method	SSIM $\uparrow$			PSNR (dB) $\uparrow$		
	100%	70%	60%	100%	70%	60%
Stripformer [60]	0.921	0.825	0.805	31.84	28.34	27.21
AdaRevD [37]	0.925	0.896	0.865	32.10	31.43	29.51
UPADNet	0.932	0.929	0.926	32.58	32.40	32.12

**Fig. 7:** Qualitative results on the COCO dataset with 60% of training data.

Table 3 evaluates the robustness of different models under reduced training data ratios (100%, 70%, and 60%) on the COCO dataset. As the amount of training data decreases, Stripformer and AdaRevD exhibit noticeable performance degradation, particularly in SSIM. In contrast, UPADNet maintains consistently high performance across all data regimes, achieving the best PSNR and SSIM at every training ratio. Notably, the performance drop from 100% to 60% training data is minimal for UPADNet, demonstrating strong data efficiency and improved generalization in low-data settings. Fig. 7 presents qualitative comparisons under the 60% training scenario. While competing methods suffer from residual blur and loss of fine structures, UPADNet better preserves edge sharpness and structural details, as evident in both Figs. 6 and 7. Because unrolled networks are derived from iterative algorithms—which typically require little or no training—UPADNet demonstrates superior generalization ability.

5.1 Ablation Study

We conduct comprehensive ablation studies on the GoPro dataset to analyze the impact of (i) the number of UPADBlocks and (ii) the choice of activation function and learning rate.

Table 4 evaluates the effects of progressively increasing the number of UPADBlocks from 24 to 72. As expected, increasing depth significantly improves restoration quality. The performance improves steadily from 28.41 dB at 24 blocks to 34.08 dB at 72 blocks, indicating that deeper unrolled structures better capture long-range dependencies and refine high-frequency details. Notably, the gains become moderate beyond 60 blocks, suggesting diminishing returns as the representational capacity saturates. This behavior is consistent with depth-scaling

Table 4: Ablation on the number of UPADBlocks in UPADNet.

# UPADBlocks	24	36	48	60	72
PSNR $\uparrow$	28.41	29.68	32.89	33.02	34.08
SSIM $\uparrow$	0.9512	0.9556	0.9581	0.9594	0.9601

Table 5: Ablation on activation function and learning rate.

Metric	ReLU		GELU	
	lr= 1×10^{-3}	lr= 1×10^{-4}	lr= 1×10^{-3}	lr= 1×10^{-4}
PSNR $\uparrow$	29.52	31.18	34.01	33.10
SSIM $\uparrow$	0.8638	0.9119	0.951	0.935

trends observed in modern restoration networks. Based on this trade-off between performance and computational cost, we select 72 blocks for the final model, as it provides the best quantitative performance.

Table 5 compares ReLU and GELU activations under two learning rates (1×10^{-3} and 1×10^{-4}). The results show that GELU consistently outperforms ReLU across both metrics. In particular, GELU with a learning rate of 1×10^{-3} achieves the best performance (34.01 dB PSNR and 0.951 SSIM), demonstrating improved optimization stability and smoother gradient propagation. Reducing the learning rate to 1×10^{-4} slightly degrades performance, indicating slower convergence and reduced exploration of the parameter space. These findings suggest that smoother nonlinearities combined with moderately larger learning rates are beneficial for phase-amplitude reconstruction tasks.

Overall, the ablation experiments demonstrate that both sufficient unrolled depth and appropriate activation design are critical for achieving state-of-the-art deblurring performance.

6 Conclusion

Our work develops UPADNet, a deep unrolled network architecture that jointly recovers the amplitude and phase components of a clean image from blurred and noisy observations. By decomposing the image and blur kernel into Fourier-domain amplitude and phase components, we derive LMMSE estimators for the amplitude and phase of the blurred image and integrated them into a regularized objective solved via alternating optimization and HQS. The resulting iterative scheme is then unfolded into a learnable network with dynamic, iteration-specific parameters.

Extensive experiments on GoPro, RealBlur, and noise-corrupted COCO datasets demonstrate that UPADNet outperforms state-of-the-art alternatives.

In particular, the explicit unrolled network for phase-component recovery—derived from an LMMSE estimation-inspired iterative formulation—is crucial for achieving performance gains in challenging real-world scenarios under varying blur type, different noise levels, and limited training data.

Overall, this work highlights the importance of structured frequency-domain modeling within deep unrolling frameworks and opens new directions for interpretable and physics-inspired learning in inverse imaging problems.

Acknowledgment

V. Monga was supported by National Science Foundation (NSF) under Grant ECCS-2143557.

C. Lee’s work was supported by the Institute of Information & Communications Technology Planning & Evaluation (IITP) under the Artificial Intelligence Convergence Innovation Human Resources Development (IITP-2026-RS-2023-00254592) grant funded by the Korea government (MSIT).

References

1. Arbelaez, P., Maire, M., Fowlkes, C., Malik, J.: Contour detection and hierarchical image segmentation. *IEEE Trans. Pattern Anal. Mach. Intell.* **33**(5), 898–916 (2010) [5](#)
2. Babacan, S.D., Molina, R., Do, M.N., Katsaggelos, A.K.: Bayesian blind deconvolution with general sparse image priors. In: *Eur. Conf. Comput. Vis.* pp. 341–355 (2012) [2](#)
3. Bertero, M., Boccacci, P.: *Introduction to Inverse Problems in Imaging*. CRC Press (1998) [6](#)
4. Cai, J.F., Ji, H., Liu, C., Shen, Z.: Framelet-based blind motion deblurring from a single image. *IEEE Trans. Image Process.* **21**(2), 562–572 (2011) [2](#)
5. Carbajal, G., Vitoria, P., Lezama, J., Musé, P.: Blind motion deblurring with pixel-wise kernel estimation via kernel prediction networks. *IEEE Trans. Comput. Imag.* **9**, 928–943 (2023) [2](#)
6. Chan, T.F., Shen, J.: *Image Processing and Analysis: Variational, PDE, Wavelet, and Stochastic Methods*. SIAM (2005) [4](#)
7. Chen, L., Chu, X., Zhang, X., Sun, J.: Simple baselines for image restoration. In: *Eur. Conf. Comput. Vis.* pp. 17–33 (2022) [2](#), [9](#), [10](#)
8. Cho, S.J., Ji, S.W., Hong, J.P., Jung, S.W., Ko, S.J.: Rethinking coarse-to-fine approach in single image deblurring. In: *IEEE Int. Conf. Comput. Vis.* pp. 4641–4650 (2021) [2](#)
9. Delbracio, M., Garcia-Dorado, I., Choi, S., Kelly, D., Milanfar, P.: Polyblur: Removing mild blur by polynomial reblurring. *IEEE Trans. Comput. Imag.* **7**, 837–848 (2021) [2](#)
10. Deng, X., Zhang, C., Jiang, L., Xia, J., Xu, M.: DeepSN-Net: Deep semi-smooth Newton driven network for blind image restoration. *IEEE Trans. Pattern Anal. Mach. Intell.* **47**(4), 2632–2646 (2025) [3](#), [10](#)
11. Dong, J., Pan, J., Yang, Z., Tang, J.: Multi-scale residual low-pass filter network for image deblurring. In: *IEEE Int. Conf. Comput. Vis.* pp. 12345–12354 (2023) [2](#)
12. Fang, Z., Wu, F., Dong, W., Li, X., Wu, J., Shi, G.: Self-supervised non-uniform kernel estimation with flow-based motion prior for blind image deblurring. In: *IEEE Conf. Comput. Vis. Pattern Recognit.* pp. 18105–18114 (2023) [10](#)
13. Geman, D., Reynolds, G.: Nonlinear image recovery with half-quadratic regularization. *IEEE Trans. Image Process.* **4**(7), 932–946 (1995) [7](#)
14. Huang, L., Xia, Y., Ye, T.: Effective blind image deblurring using matrix-variable optimization. *IEEE Trans. Image Process.* **30**, 4653–4666 (2021) [2](#)
15. Huang, Y., Chouzenoux, E., Pesquet, J.C.: Unrolled variational Bayesian algorithm for image blind deconvolution. *IEEE Trans. Image Process.* **32**, 430–445 (2022) [3](#)
16. Hyvärinen, A., Hurri, J., Hoyer, P.: *Natural Image Statistics: A Probabilistic Approach to Early Computational Vision*. Springer (2009) [6](#)
17. Joshi, N., Szeliski, R., Kriegman, D.J.: PSF estimation using sharp edge prediction. In: *IEEE Conf. Comput. Vis. Pattern Recognit.* pp. 1–8 (2008) [1](#)
18. Kay, S.M.: *Fundamentals of Statistical Signal Processing: Estimation Theory*. Prentice Hall (1993) [4](#)
19. Kong, L., Dong, J., Ge, J., Li, M., Pan, J.: Efficient frequency domain-based transformers for high-quality image deblurring. In: *IEEE Conf. Comput. Vis. Pattern Recognit.* pp. 5886–5895 (2023) [2](#), [10](#)
20. Kong, L., Dong, J., Tang, J., Yang, M.H., Pan, J.: Efficient visual state space model for image deblurring. In: *IEEE Conf. Comput. Vis. Pattern Recognit.* pp. 12710–12719 (2025) [10](#), [12](#)

21. Krishnan, D., Tay, T., Fergus, R.: Blind deconvolution using a normalized sparsity measure. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 233–240 (2011) [1](#)
22. Kundur, D., Hatzinakos, D.: Blind image deconvolution. IEEE Signal Processing Magazine **13**(3), 43–64 (1996). <https://doi.org/10.1109/79.489268> [1](#), [4](#), [7](#)
23. Kupyn, O., Budzan, V., Mykhailych, M., Mishkin, D., Matas, J.: DeblurGAN: Blind motion deblurring using conditional adversarial networks. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 8183–8192 (2018) [2](#)
24. Kupyn, O., Martyniuk, T., Wu, J., Wang, Z.: DeblurGAN-v2: Deblurring (orders-of-magnitude) faster and better. In: IEEE Int. Conf. Comput. Vis. pp. 8878–8887 (2019) [10](#)
25. LeCun, Y., Bengio, Y., Hinton, G.: Deep learning. Nature **521**(7553), 436–444 (2015) [2](#)
26. Levin, A., Weiss, Y., Durand, F., Freeman, W.T.: Understanding and evaluating blind deconvolution algorithms. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 1964–1971 (2009) [4](#), [6](#), [7](#)
27. Levin, A., Weiss, Y., Durand, F., Freeman, W.T.: Efficient marginal likelihood optimization in blind deconvolution. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 2657–2664 (2011) [2](#)
28. Levin, A., Weiss, Y., Durand, F., Freeman, W.T.: Understanding blind deconvolution algorithms. IEEE Trans. Pattern Anal. Mach. Intell. **33**(12), 2354–2367 (2011) [2](#)
29. Li, H., Qian, W., Cao, J., Nie, R., Liu, P., Xu, D.: Multi-interactive enhanced for defocus blur estimation. IEEE Trans. Comput. Imag. **10**, 640–652 (2024) [2](#)
30. Li, Y., Tofighi, M., Geng, J., Monga, V., Eldar, Y.C.: Efficient and interpretable deep blind image deblurring via algorithm unrolling. IEEE Trans. Comput. Imag. **6**, 666–681 (2020) [3](#), [5](#), [10](#), [12](#)
31. Li, Y., Tofighi, M., Monga, V., Eldar, Y.C.: An algorithm unrolling approach to deep image deblurring. In: ICASSP. pp. 7675–7679 (2019) [3](#)
32. Liao, S., Huang, H., Liu, J., Xiao, X., Li, X., Li, S.: Modified Newton integration algorithm with noise tolerance for image deblurring. IEEE Trans. Comput. Imag. **7**, 1254–1266 (2021) [2](#)
33. Lim, J.S.: Two-Dimensional Signal and Image Processing. Prentice Hall (1990) [3](#), [4](#)
34. Lin, J.C., Wei, W.L., Liu, T.L., Kuo, C.C.J., Liao, M.: Tell me where it is still blurry: Adversarial blurred region mining and refining. In: ACM Int. Conf. Multimedia. pp. 702–710 (2019) [2](#)
35. Lin, T.Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., Zitnick, C.L.: Microsoft COCO: Common objects in context. In: Eur. Conf. Comput. Vis. pp. 740–755 (2014) [10](#)
36. Majee, S., Aslan, S., Gürsoy, D., Bouman, C.A.: CodEx: A modular framework for joint temporal de-blurring and tomographic reconstruction. IEEE Trans. Comput. Imag. **8**, 666–678 (2022) [2](#)
37. Mao, X., Li, Q., Wang, Y.: AdaRevD: Adaptive patch exiting reversible decoder pushes the limit of image deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 25681–25690 (2024) [10](#), [12](#), [13](#)
38. Mao, X., Liu, Y., Liu, F., Li, Q., Shen, W., Wang, Y.: Intriguing findings of frequency selection for image deblurring. In: AAAI. vol. 37, pp. 1905–1913 (2023) [2](#)
39. Monga, V., Li, Y., Eldar, Y.C.: Algorithm unrolling: Interpretable, efficient deep learning for signal and image processing. IEEE Signal Process. Mag. **38**(2), 18–44 (2021) [3](#)

40. Nah, S., Kim, T.H., Lee, K.M.: Deep multi-scale convolutional neural network for dynamic scene deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 3883–3891 (2017) [2](#), [4](#), [9](#), [10](#)
41. Nazir, S., Vaquero, L., Mucientes, M., Brea, V.M., Coltuc, D.: Depth estimation and image restoration by deep learning from defocused images. IEEE Trans. Comput. Imag. **9**, 607–619 (2023) [2](#)
42. Oppenheim, A.V., Lim, J.S.: The importance of phase in signals. Proc. IEEE **69**(5), 529–541 (2005) [3](#), [4](#)
43. Pan, J., Hu, Z., Su, Z., Yang, M.H.: L_0 -regularized intensity and gradient prior for deblurring text images and beyond. IEEE Trans. Pattern Anal. Mach. Intell. **39**(2), 342–355 (2016) [1](#)
44. Pan, J., Sun, D., Pfister, H., Yang, M.H.: Deblurring images via dark channel prior. IEEE Trans. Pattern Anal. Mach. Intell. **40**(10), 2315–2328 (2017) [2](#)
45. Pan, L., Hartley, R., Liu, M., Dai, Y.: Phase-only image based kernel estimation for single image blind deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 6027–6036 (2019) [3](#)
46. Perrone, D., Favaro, P.: A clearer picture of total variation blind deconvolution. IEEE Trans. Pattern Anal. Mach. Intell. **38**(6), 1041–1055 (2015) [1](#)
47. Quan, Y., Wu, Z., Ji, H.: Gaussian kernel mixture network for single image defocus deblurring. Adv. Neural Inform. Process. Syst. **34**, 20812–20824 (2021) [3](#)
48. Richmond, G., Cole-Rhodes, A.: Non-uniform blind image deblurring using an algorithm unrolling neural network. In: IVMS. pp. 1–5 (2022) [3](#), [12](#)
49. Rim, J., Lee, H., Won, J., Cho, S.: Real-world blur dataset for learning and benchmarking deblurring algorithms. In: Eur. Conf. Comput. Vis. pp. 184–201 (2020) [4](#), [9](#), [10](#)
50. Ruan, L., Chen, B., Li, J., Lam, M.L.: AIFNet: All-in-focus image restoration network using a light field-based dataset. IEEE Trans. Comput. Imag. **7**, 675–688 (2021) [2](#)
51. Rudin, L.I., Osher, S., Fatemi, E.: Nonlinear total variation based noise removal algorithms. Physica D **60**, 259–268 (1992) [6](#)
52. ŞAHİN, E., Arslan, N.N., Özdemir, D.: Unlocking the black box: An in-depth review on interpretability, explainability, and reliability in deep learning. Neural Computing and Applications **37**(2), 859–965 (2025) [3](#)
53. Sanghvi, Y., Gnanasambandam, A., Chan, S.H.: Photon limited non-blind deblurring using algorithm unrolling. IEEE Trans. Comput. Imag. **8**, 851–864 (2022) [3](#)
54. Sarker, I.H.: Deep learning: A comprehensive overview on techniques, taxonomy, applications and research directions. SN Comput. Sci. **2**(6), 1–20 (2021) [2](#)
55. Shan, Q., Jia, J., Agarwala, A.: High-quality motion deblurring from a single image. ACM Trans. Graph. **27**(3), 1–10 (2008) [2](#)
56. Shlezinger, N., Whang, J., Eldar, Y.C., Dimakis, A.G.: Model-based deep learning. Proc. IEEE **111**(5), 465–499 (2023) [3](#)
57. Sun, L., Cho, S., Wang, J., Hays, J.: Edge-based blur kernel estimation using patch priors. In: ICCP. pp. 1–8 (2013) [2](#)
58. Tao, X., Gao, H., Shen, X., Wang, J., Jia, J.: Scale-recurrent network for deep image deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 8174–8182 (2018) [2](#)
59. Tikhonov, A.N., Arsenin, V.Y.: Solutions of Ill-posed Problems. Winston (1977) [6](#)
60. Tsai, F.J., Peng, Y.T., Lin, Y.Y., Tsai, C.C., Lin, C.W.: Stripformer: Strip transformer for fast image deblurring. In: Eur. Conf. Comput. Vis. pp. 146–162 (2022) [2](#), [10](#), [12](#), [13](#)

61. Tu, Z., Talebi, H., Zhang, H., Yang, F., Milanfar, P., Bovik, A., Li, Y.: MAXIM: Multi-axis MLP for image processing. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 5769–5780 (2022) [2](#)
62. Vo, T., Park, C.Y.: Deep joint unrolling for deblurring and low-light image enhancement (JUDE). In: Winter Conf. Appl. Comput. Vis. pp. 2696–2705 (2025) [3](#)
63. Wang, C., Li, Y., Zhang, Y., Chen, Y., Fu, Y., Wang, Y., An, W.: Uformer: A general U-shaped transformer for image restoration. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 17683–17693 (2022) [2](#)
64. Whang, J., Delbracio, M., Talebi, H., Saharia, C., Dimakis, A.G., Milanfar, P.: Deblurring via stochastic refinement. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 16293–16303 (2022) [2](#)
65. Wipf, D., Zhang, H.: Revisiting Bayesian blind deconvolution. *J. Mach. Learn. Res.* **15**, 3775–3814 (2014) [2](#)
66. Xie, X., Zhang, Q., Zheng, W.S.: Diffusion-based event generation for high-quality image deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 2194–2203 (2025) [10](#)
67. Xu, L., Zheng, S., Jia, J.: Unnatural L_0 sparse representation for natural image deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 1107–1114 (2013) [1](#)
68. Yan, Y., Ren, W., Guo, Y., Wang, R., Cao, X.: Image deblurring via extreme channels prior. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 4003–4011 (2017) [2](#)
69. Zamir, S.W., Arora, A., Khan, S., Hayat, M., Yang, F.S.K.M.: Restormer: Efficient transformer for high-resolution image restoration. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 5718–5729 (2022) [2](#)
70. Zamir, S.W., Arora, A., Khan, S.H., Hayat, M., Yang, F.S., Shao, L., Khan, M.H.: Multi-stage progressive image restoration. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 14821–14831 (2021) [2](#)
71. Zhang, H., Dai, Y., Li, H., Koniusz, P.: Deep stacked hierarchical multi-patch network for image deblurring. In: IEEE Conf. Comput. Vis. Pattern Recognit. pp. 5978–5986 (2019) [2](#)