

EDGAN: Motion Deblurring algorithm based on Enhanced Generative Adversarial Networks

Yong Zhang^{1,2,3}

¹ ATR Key Laboratory of National Defense Technology, Shenzhen University, Shenzhen 518060, China

² Guangdong Key Laboratory of Intelligent Information Processing, Shenzhen University, Shenzhen 518060, China

³ Guangdong Laboratory of Artificial Intelligence and Digital Economy (SZ), Shenzhen University, Shenzhen 518060, China

Shao Yong Ma¹

¹ ATR Key Laboratory of National Defense Technology, Shenzhen University, Shenzhen 518060, China

Xi Zhang⁴

⁴ Information Centre, Shenzhen Technology University, Shenzhen 518118, China

Li Li⁵

⁵ School of Foreign Languages, Shenzhen Institute of Information Technology, Shenzhen 518172, China

Wai Hung Ip⁶

⁶ Department of Industrial and Systems Engineering, The Hong Kong Polytechnic University, Hong Kong 999077, China

Kai Leung Yung⁶

⁶ Department of Industrial and Systems Engineering, The Hong Kong Polytechnic University, Hong Kong 999077, China

Abstract: Removing motion blur has been an important issue in computer vision literature, and some achievements were obtained in the research of image deblur by using deep learning algorithm in recent years. Motion blur is caused by the relative motion between the camera and the photographed object. In this paper, we proposed an enhanced adversarial networks model. A convolution unit constructed by Squeeze-and-Excitation Networks (SENet) and Residual Networks (RseNet) together with a new mixed loss function are used in our model to restore the motion blurred image. In addition, we exploit resize-convolution as the upsampling method to eliminate the chessboard artifacts in the generated image. Our model could greatly shorten the running time compared to the previous method which bypasses the process of blur kernel estimation. Our approach is tested using GOPRO datasets and Lai datasets, the peak signal-to-noise ratio (PSNR) of our approach is up to 29.435537/29.762314, and the structural similarity measure (SSIM) can be achieved upto 0.974123/0.686638. Furthermore, we simulate the images obtained from the China's Chang'e 3 Lander to test the new algorithm. Due to the elimination of the chessboard effect, the deblurred image has a better visual appearance. Our method was proved to achieve higher performance and efficiency in the qualitative and quantitative aspects using the benchmark datasets experiments. The results also provided various insights into the design and development of the Camera Pointing System which was mounted on the Lander for capturing images of the moon and rover for Chang'e space mission.

Keywords: Blurred image, Camera Pointing System, Chang'e space mission, GANs, Resize convolution

1. Introduction

Motion blur is one of the most common types of image blurring. Shorter exposure time and fast moving objects or camera shaking will cause motion blur in the final image, resulting in poor image perception, affecting image information transmission and postprocessing [1, 2]. In the field of computer vision, motion blur could cause the reduction of accuracy and efficiency of image recognition and classification. Therefore, the restoration of motion blurred image is of great signification.

Most of the early blurred image restoration approaches are based on following blur model [3-6]:

$$I_B = K * I_S + N \dots\dots\dots (1)$$

where I_B , I_S , K and N are blurred image, latent sharp image, blur kernel and noise, $*$ represents convolution operation. The process could be seen as a convolution operation between sharp image and blurred image kernel, and blurred image is formed after the effect of random noise. Image restoration algorithm could be divided into blind restoration and non-blind restoration algorithm base on whether the blur kernel is known or not. The non-blind restoration algorithm restores the blurred image by

estimation of the inverse process of equation [1] using the known blur kernel. The classical algorithm has LR(Lucy-Richardson) algorithm, Wiener filter, Kalman filter [7]. A partial blind restoration algorithm is used to reconstruct the image by estimation the blur kernel. It is time-consuming and inefficient to restore image by blur kernel, due to the unknown blur kernel function; the blur types in reality are complex and uncertain in most cases.

GAN's excellent performance in the image conversion task and the existing image deblur algorithm still have many shortcomings. It is a new direction to apply GAN to motion blur image restoration task. In this paper, an enhanced GANs model, which may obtain higher PSNR and SSIM on motion blur removal, is proposed. In addition, the deblurred image will have a better visual appearance and quality. In Section II, the related work on image restorations is discussed. In Section III, the proposed method and algorithm are described, and the mathematical analysis is elaborated. In Section IV, the proposed method is validated by the experiments, and the benchmarked datasets are used to compare the results. The experiments are also simulated on images obtained from the Chang'e 3 space mission. In Section V, the summary of the work, conclusions, and implications for the design of Chang'e mission's camera pointing systems is also given. The contributions of this paper are as follows:

1. A new work structure and a hybrid loss function are proposed, which can recover the motion blurry image efficiently;
2. Experiments on different datasets prove that the algorithm can be applied to real scenes.

2. Related work

2.1 Related study

In recent years, numerous blurred image restoration approaches based on deep learning have been proposed. Sun et al. [9] used CNN to predict the probability distribution of motion blur in every block of image and restored the image by the motion blur probability distribution of each image block. Nah et al. [10] used a multi-scale CNN to directly correct the deblurred image. However, the algorithm has high complexity and low efficiency, so it cannot process images quickly. Ramakrishnan et al. [11] proposed a novel network structure to improve the efficiency of the network, while maintaining the effect of deblurring. Although this algorithm improved the efficiency, the clarity of the generated images was similar to that of Nah's method. Kupyn et al. [12] proposed a motion blur removal algorithm based on condition generation against a network. This approach uses WGAN-GP [13] and perceptual loss [14] as the final loss function and achieved a high image restoration effect. Deep learning is also applied to motion blurred video restoration. Su et al. [15] adopt CNN to aggregate multiple images to generate a sharp output. Zhang et al. [16] used GAN with 3D convolution to capture spatial and temporal information encoded in neighbouring frames to restore blurred video. Chen et al. [17] used self-supervised fashion to fine-tune existing deblurring neural networks, which improves the

performance of video deblurring algorithm.

2.2 GANs

GANs consists of a generator and a discriminator. The basic principle is that the generator receives a random noise signal to generate new data samples, and the discriminator determines whether the samples are from the real sample set or the simulated samples generated by the generator. The purpose of the generator is to generate a sample that is close to the data distribution of the real sample set, making it impossible for the discriminator to determine the data source.

The optimization objective function of the generation of GANs is as follows:

$$\min_G \max_D V(D, G) = E_{x \sim p_{data}(x)} [\log D(x)] + E_{z \sim p_z(z)} [\log(1 - D(G(z)))] \quad \text{..... (2)}$$

Where $p_{data}(x)$, $p_z(z)$, $G(z; \theta_g)$, and $D(x)$ are data distribution, predefined noise variable, mapping from noise space to data space and the probability that x from the real sample set. The discriminator was trained to minimize $\log D(x)$, and the generator is trained to minimize $\log(1 - D(z))$. In order to enable the network to output data according to our expectation, researchers added the extra conditional information y on the original GANs, and y could be any kind of information[25]. The optimal objective function of conditional generation adversarial networks is as follows:

$$\min_G \max_D V(D, G) = E_{x \sim p_{data}(x)} [\log D(x | y)] + E_{z \sim p_z(z)} [\log(1 - D(G(z | y)))] \quad \text{..... (3)}$$

However, training of original GANs suffer from many problems such as mode collapse, training instability and vanishing gradients etc. The reason for these problems is that the original GANs uses JS Divergence to measure the differences between the two distributions. JS Divergence is a constant $\log 2$ when two distributions do not overlap completely, which leads to the mutation properties of JS Divergence. Thus, Wasserstein GAN (WGAN)[14] uses Wasserstein Distance to measure the differences between the two distributions.

$$W(P_r, P_g) = \inf_{\gamma \sim \Pi(P_r, P_g)} E_{x, y} [\|x - y\|] \quad \text{..... (4)}$$

For two distribution P_r and P_g , their joint distribution is $\Pi(P_r, P_g)$. Calculating the distance between x and y sampled from each joint distribution which takes all x and y to calculate the expected value. The smallest value is chosen as the Wasserstein Distance. Compared with JS Divergence, Wasserstein Distance has continuous transformation regardless of whether the two distributions overlap.

3. Proposed method

The proposed method is to improve the performance of existing end-to-end deblurring adversarial network model by making the network to have the ability to extract the weight of feature channels and remove the draughtboard artefacts. The key idea is to use a novel convolution unit constructed by squeeze and excitation networks and residual network to extract image feature. In addition, using resize convolution as up-sampling method, the contrastive experiments showed that chessboard effect could be effectively removed by resize convolution.

3.1 Generator

The generator structure is shown in Figure 1 which is similar to the structure used by Kupyn et al [12]. It contains two strided convolution blocks with stride 2, nine residual blocks, and two transposed convolution blocks. In order to improve the quality of the image, we have made some improvements on the basis of the network structure of Kupyn et al. [11], they are elaborated as follows.

3.1.1 Without Batch Normalization

Traditional neural networks only normalise the data before it is inputted into the network, whereas the batch normalisation [20] layer normalises the input of hidden layers. Batch normalisation may solve the problem of gradient disappearance and gradient explosion in the back-propagation algorithm, which can also accelerate the convergence speed of the network. The batch normalisation layers are removed from the network, as Nah et al. [10] and Bee et al. [21] presented in their model. Since batch normalisation layers normalise the features, it limits network flexibility

3.1.2 SE-ResBlock

The SENet (squeeze and excitation networks) proposed by Hu et al. [22] is used to improve the performance of the residual network. During CNN, the convolution kernel could be regarded as the aggregate of the spatial information and the characteristics of dimension information. To improve the performance of CNN, many approaches have been proposed, from spatial dimension, such as the inception module [20, 23–25]. The difference is that SENet improves network performance from the feature dimension; SENet can learn the relationship between different feature channels, obtain the weight of feature channels, and use the weight to promote useful features and suppress features that are less useful for the current task.

SENet includes three key operations: squeeze, excitation, and reweight. Squeeze operations achieve feature compression by turning each of the two-dimensional characteristic channels into a real number. Excitation operation gives weight to each feature channel through the parameter, W , which represents the correlation between the feature channels. In the reweight operation, the output of excitation is regarded as the importance of different feature channels. By multiplying the output of excitation as the weighting of the original feature, the original feature is re-calibrated. SENet has

been proven to have excellent performance in image classification. In this paper, SENet is applied to image processing due to the image processing method adopted in this paper to completely reconstruct the image.

In the reconstruction process, the importance of different feature channels should be considered. So the combined convolution unit of SENet and ResNet (residual networks) [26], named SE-ResBlock, is used and is shown in Fig. 2. Global average pooling is used as a squeeze operation. The excite operation calculates the interchannel correlation by using two fully connected layers to make up the bottleneck structure. The reweight operation is used to weight the normalised feature into the original feature channels.

3.1.3 Resize-convolution [27]

Using CNN to generate images is a process that transforms low resolution image blocks into high resolution image. It is usually realized by deconvolution. Due to the “uneven overlap” in the deconvolution process, it will lead to the artefact similar to checkerboard lattice in the details of the image, which is called “chessboard artefacts”. To eliminate this phenomenon, one approach is to make sure the kernel size is divided by the stride; however, it is still easy to create draughtboard artefacts. The resize convolution is used as an up-sampling method instead of deconvolution. The resize convolution is implicitly weight-tying in a way that discourages high-frequency artefacts. The process of resize convolution is to resize the image (using nearest-neighbour interpolation or bilinear interpolation) and then do a convolutional layer, as shown in figure 3.

3.2 Discriminator

Here, the Markovian discriminator [28], PatchGAN, is used as the discriminator for EDGAN. Since content loss (the combination of perceptual loss and gradient loss was used in this paper) has been able to process the low-frequency components of the image very well, the discriminator only needs to process the high-frequency components. Therefore, the receptive field of the discriminator output does not need to be the whole input image. The output can be a feature map, and the receptive field of each pixel in the feature map is a patch on the input image, which can accelerate the discriminator while obtaining high-quality images. According to the experiments in the literature of Isola et al. [29], the output of the discriminator is set at 50×50 , which takes into account both image quality and network operation speed. The model parameters of the discriminator are shown in Table 1.

Table1: Model parameters of discriminator.

#	Layer	Weight dimension	Stride
1	cov	$3 \times 64 \times 4 \times 4$	2
2	LRelu	-	-
3	cov	$64 \times 128 \times 4 \times 4$	2
4	BN	-	-
5	LRelu	-	-
6	cov	$128 \times 512 \times 4 \times 4$	2
7	BN	-	-
8	LRelu	-	-
9	cov	$512 \times 512 \times 4 \times 4$	2
10	BN	-	-
11	LRelu	-	-
12	cov	$512 \times 512 \times 4 \times 4$	1
13	BN	-	-
14	LRelu	-	-
15	FC	$512 \times 1 \times 4 \times 4$	1
16	Sigmoid	-	-

3.3 Loss function

We formulate the loss function as a combination of adversarial loss, perceptual loss and gradient loss:

$$\ell = \ell_{WGAN-GP}^{Generator} + \alpha \cdot \ell_{percept} + \beta \cdot \ell_{grad} \quad \dots\dots\dots (5)$$

where α and β are the weight parameters of perceptual loss and gradient loss.

3.3.1 Adversarial loss

To overcome the problems existing in the original GANs, we use WGAN-GP as adversarial loss. Although WGAN reduces the difficulty of training for GANs, it is still difficult to converge under certain conditions, and the effect of generating the picture does not satisfy with our expectations. WGAN will limit the weight to a certain extent after updating the weight of each iteration, while WGAN-GP calculates the weight gradient according to the input of the discriminator, and corrects the weight according to the norm of the gradient. WGAN-GP effectively solves the problem of WGAN, which is calculated as the following, where I^B is blurry image:

$$\ell_{WGAN-GP}^{Generator} = \sum_{n=1}^N -D_{\theta_D}(G_{\theta_G}(I^B)) \quad \dots\dots\dots (6)$$

3.3.2 Perceptual loss

The basic idea is to use the features extracted by CNN as part of the target

function. By reducing the Euclidean distance between feature maps generated by CNN and target images, the generated images is more consistent with the target image than the pixel level loss function. Perceptual loss is a kind of high-level loss. The definition is as follows:

$$\ell_{percep} = \frac{1}{W_{i,j}H_{i,j}} \sum_{x=1}^{W_{i,j}} \sum_{y=1}^{H_{i,j}} (\phi_{i,j}(I^S)_{x,y} - \phi_{i,j}(G_{\theta_G}(I^B))_{x,y})^2 \dots\dots\dots (7)$$

where $W_{i,j}$ and $H_{i,j}$ are the dimensions of the feature maps, $\phi_{i,j}$ is the feature map obtained by the j-th convolution before the i-th maxpooling layer within a CNN. The CNN used in this paper is a VGG19 (layer 1 - 14) network [31], pretrained on ImageNet [32].

3.3.3 Gradient loss [33]

In addition, image information in gradient domain is also leveraged as a high-level loss term as follows:

$$\ell_{grad} = \frac{1}{2N} \sum_{i=1}^N |\nabla_h(I^S) - \nabla_h(I^B)| + |\nabla_v(I^S) - \nabla_v(I^B)| \dots\dots\dots (8)$$

where ∇_h and ∇_v indicate the horizontal and vertical gradients. N indicate the number of training images pairs.

WGAN-GP is used as the critic function, the discriminator loss, as the following:

$$\ell_{WGAN-GP}^{Discriminator} = \mathbb{E}_{\tilde{x} \sim P_g} [D(\tilde{x})] - \mathbb{E}_{x \sim P_r} [D(x)] + \lambda \mathbb{E}_{\hat{x} \sim P_{\hat{x}}} \left[(\nabla_{\hat{x}} D(\hat{x})_2 - 1)^2 \right] \dots\dots\dots (9)$$

where λ is the penalty coefficient, p_g is the sample distribution of the generator, p_r is the sample distribution of the sharp image, and p_x is the distribution uniformly sampled along the straight lines between p_g and p_r .

4. Experimental evaluation

4.1 Experimental Settings

The following describes the validations of the proposed method and algorithms. Firstly, we used benchmarked datasets to test the performances, and then it is tested with the images obtained from Chang'e 3 space missions. All of the proposed models using the PyTorch deep learning framework are implemented. The experimental model training and testing hardware platform is NVIDIA GTX 1070 GPU and Intel i7-4790 CPU. The first model is referred for using SE-ResBlock as EDGANSE. The second model used SE-ResBlock and gradient loss as EDGAN_{SE-G}. The third model used SE-ResBlock and a resize convolution as EDGAN_{SE-R}. A random crop size of 250×250 from a GOPRO dataset [10] is used to train the proposed model. The model optimisation algorithm is ADAM [34], a set learning rate of 10^{-4} . To speed up the convergence rate of the model, the learning rate in the first 150 rounds was fixed, and

in the second 150 rounds, the learning rate was gradually reduced to zero. To balance the performance of the generator and the discriminator, every five-gradient descent algorithms were performed by the discriminator, and the generator was executed once. The α and β of generator loss function are 1 and 5. The approaches proposed in this paper do not adopt a batch normalisation layer. Other specific differences are shown in Table 2.

4.2 Datasets

A)GOPRO Dataset

The GOPRO dataset consists of 3214 pairs of images, of which 2103 are training sets and 1111 pairs are test sets. We compare our model with Kupyn et al.[11] ,and the results are shown in Table 2.The details of image are shown in Figure 5.From the results, it is found that EDGAN could better restore the image detail. The result of all models after different epoch of training are shown in figure4.EDGAN shows superior convergence rate and image restoration ability.

Table2: Mean peak signal-to-noise ratio and structural similarity measure on GOPRO dataset of 1111 images.

Method Metric	<i>DeblurGAN</i> [11]	<i>EDGAN</i>		
		<i>EDGAN</i> _{SE}	<i>EDGAN</i> _{SE-G}	<i>EDGAN</i> _{SE-R}
PSNR	28.182472	29.435537	29.109105	28.648071
SSIM	0.963094	0.972200	0.973316	0.967033

The contrast between *EDGAN*_{SE} and *EDGAN*_{SE-R} is shown in figure 6.Deconvolution causes the abnormal color artifacts in the image texture, and resize-convolution can eliminated it.

B)Lai Dataset [24]

The Lai dataset includes a real dataset and a synthetic dataset. The real dataset is made up of 100 blurred images collected from real world scenes. These images use different shooting devices, shooting settings and shooting themes. The synthetic dataset includes 100 blurred images generated by convolution between 25 sharp images collected from Internet and 4 different blur kernels. By recorded 6D camera trajectories to generate blur kernels. In the experiment we only use synthetic dataset, and the results are shown in Table 3. *EDGAN* shows superior results both in

qualitative and quantitative ways. Deblurred images from test on Lai dataset are shown in figure 7.

Table3:Mean peak signal-to-noise ratio and structural similarity measure on Lai dataset of 100 images.

Method Metric	<i>DeblurGAN</i>	<i>EDGAN</i>		
	[11]	$EDGAN_{SE}$	$EDGAN_{SE-G}$	$EDGAN_{SE-R}$
PSNR	29.670198	29.762314	29.729374	29.674767
SSIM	0.675933	0.686638	0.676432	0.680609

C) Chang'e 3 space mission images

The impact of the proposed algorithm is also evaluated using the open source data images available for Chang'e 3 space mission. The Chang'e 3 space mission from China was successfully launched and the rover "Yutu" landed on the moon's surface in December 2013. It was considered one of the many successful space missions from the chineses for several decades. One of the many projects that was led by the author of this paper Prof. KL Yung [29,30,31] and his team at the Hong kong Polytechnic University being the design and development of Camera Pointing System mounted on the lander of the moon's surface. The equipment had been operated for over three years on the moon and was terminated in 2017. The following simulated experiments was conducted to test the algorithm which can be used to improve the design and implementation of future Chang'e missions:

5. Discussion

This paper demonstrates that the image generated by the proposed algorithm is sharper through comparative experiments on different datasets. Although the algorithm performs well in image deblurring, it is unable to process high-resolution images in real time due to the lack of further optimisation of network structure for algorithm efficiency. In addition, for image frames from video, because the algorithm proposed in this paper does not support multiple images as input, the interframe information cannot be extracted. The future work will focus on the speed improvement and way to extract information of this algorithm.

6. Conclusion

In this paper, the EDGAN is proposed, and it is a novel model designed to restore motion blurred images. And the EDGAN sets a new state-of-the-art technology to public benchmark datasets in terms of the PSNR and SSIM metric [37]. In addition, using resize convolution as an up-sampling method can effectively eliminate “draughtboard artefacts” on the generated images is confirmed. However, the resize convolution would reduce the quality of the image details. It will be part of future work to eliminate colour artefacts without reducing the performance of EDGAN.

Moreover, the method is tested and evaluated for future space missions of Chang’e. The camera pointing system developed by the Hong Kong Polytechnic University in early 2013 was used to capture images of the moon, as well as the movement of the rovers. It was capable of 360 degrees of image capturing, as well as positioning and navigating of the rover. Based on the past experience, a new algorithm with a deep learning approach for future space missions is proposed. The results indicate that deep learning can achieve good performances for high-precision image restorations and can be incorporated into the design of cylindrical projection of sequential images of the camera pointing system for image constructions, as well as feature recognition in future deep space explorations.

Acknowledgments

This work is supported by the Natural Science Foundation of Guangdong Province (2015A030310172) and the Science and Technology Plan Projects of Shenzhen (JCYJ20170302145623566, GJHZ20160226202520268). It is also partly supported by a grant from the department of Industrial and Systems Engineering of the Hong Kong Polytechnic university (H-ZG3K)

References

1. Polyu.edu.hk (2015) Life on Mars? Award-winning PolyU device could dig up the answer [online]. http://www.polyu.edu.hk/open_inquiry/en/story.php?sid=3. Accessed 10 Nov 2017
2. Polyu.edu.hk (2017) PolyU 70th Anniversary [online]. https://www.polyu.edu.hk/cpa/70th_anniversary/memories_achievements10.htm. Accessed 3 Nov 2017
3. Gupta A, Joshi N, Zitnick CL, Cohen M, Curless B (2010) Single image deblurring using motion density functions. In: Proceedings of the ECCV, pp 171–184
4. Harmeling S, Hirsch M, Schölkopf B (2010) Space-variant single-image blind deconvolution for removing camera shake. In: Proceedings of the NIPS, Vancouver, British Columbia, Canada, pp 829–837
5. Hirsch M, Schuler CJ, Harmeling S, Schölkopf B (2011) Fast removal of non-uniform camera shake. In: Proceedings of the ICCV, pp 463–470
6. Whyte O, Sivic J, Zisserman A, Ponce J (2010) Non-uniform deblurring for shaken images. In: Proceedings of the CVPR, pp 491–498

7. Li W (2011) Parameter estimation and algorithm research of motion blur image restoration. M.S. dissertation, Department of Computer Application, Anhui University, Anhui, China
8. Goodfellow IJ, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, Ozair S et al (2014) Generative adversarial nets. In: Proceedings of the ICNOIP, pp 2672–2680
9. Sun J, Cao W, Xu Z, Ponce J (2015) Learning a convolutional neural network for non-uniform motion blur removal. In: Proceedings of the CVPR, pp 769–777
10. Nah S, Kim TH, Lee KM (2018) Deep multi-scale convolutional neural network for dynamic scene deblurring. arXiv preprint arXiv :1612.02177 v2
11. Ramakrishnan S, Pachori S, Gangopadhyay A, et al (2017) Deep generative filter for motion deblurring. In: Proceedings of the ICCV, pp 2993–3000
12. Kupyn O, Budzan V, Mykhailych M, Mishkin D, Matas J (2017) Deblurgan: blind motion deblurring using conditional adversarial networks. arXiv preprint arXiv :1711.07064
13. Arjovsky M, Chintala S, Bottou L (2017) Wasserstein generative adversarial networks. In: Proceedings of the ICML, pp 214–223
14. Johnson J, Alahi A, Li FF (2016) Perceptual losses for real-time style transfer and super-resolution. In Proceedings of the ECCV, pp 694–711
15. Su S, Delbracio M, Wang J, Sapiro G, Heidrich W, Wang O (2017) Deep video deblurring for handheld cameras. In: Proceedings of the CVPR, pp 1279–1288
16. Zhang K, Luo W, Zhong Y, Mia L, Liu W (2018) Adversarial Spatio-Temporal Learning for Video Deblurring. In: Proceedings of the TIP, pp 281–301
17. Chen H, Gu J, Gallo O, Liu M, Veeraraghavan A, Kautz J (2018) Reblur2Deblur: deblurring videos via self-supervised learning. In: ICCP
18. Mirza M, Osindero S (2014) Conditional generative adversarial nets. Comput Sci 63:2672–2680
19. Arjovsky M, Chintala S, Bottou L (2017) Wasserstein gan. arXiv preprint arXiv :1701.07875
20. Ioffe S, Szegedy C (2015) Batch normalization: accelerating deep network training by reducing internal covariate shift. In: Proceedings of the ICML
21. Shi W, Caballero J, Huszár F, et al (2016) Real-time single image and video super-resolution using an efficient sub-pixel convolutional neural network. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pp 1874–1883
22. Hu J, Shen L, Sun G (2018) Squeeze-and-excitation networks. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pp 7132–7141
23. Szegedy C, Liu W, Jia Y, Sermanet P, Reed S, Anguelov D et al (2015) Going deeper with convolutions. In: Proceedings of the CVPR, pp 1–9
24. Szegedy C, Vanhoucke V, Ioffe S et al (2016) Rethinking the inception architecture for computer vision. In: Proceedings of the CVPR
25. Szegedy C, Ioffe S, Vanhoucke V (2017) Inception-v4, inception-resnet and the impact of residual connections on learning. arXiv :1602.07261

26. He K, Zhang X, Ren S, Sun J (2016) Deep residual learning for image recognition. In: Proceedings of the CVPR, pp 770–778
27. Odena A, Dumoulin V, Olah C (2016) Deconvolution and checkerboard artifacts. <http://distill.pub/2016/deconv-checkerboard/>
28. Liand C, Wand M (2016) Precomputed real-time texture synthesis with Markovian generative adversarial networks. In: Proceedings of the ECCV, pp 702–716
29. Isola P, Zhu J, Zhou T, Efros A (2017) Image-to-image translation with conditional adversarial networks. In: Proceedings of the CVPR, pp 1125–1134
30. Maas AL, Hannun AY, Ng AY (2013) Rectifier nonlinearities improve neural network acoustic models. Proc. ICML 30(1):3
31. Simonyan K, Zisserman A (2014) Very deep convolutional networks for large-scale image recognition. Comput Sci 362:1140
32. Deng J, Dong W, Socher R, Li LJ, Li K, Li FF (2009) ImageNet: a large-scale hierarchical image database. Computer Vision and Pattern Recognition. In: Proceedings of the CVPR, pp 248–255
33. Jiao J, Tu WC, He S, Lau RWH (2017) FormResNet: formatted residual learning for image restoration. In Proceedings of the CVPRW, pp 1034–1042
34. Kingmaand DP, Ba J (2017) Adam: a method for stochastic optimization. arXiv:1412.6980v 9
35. Chinese Academy of Sciences, China National Space Administration, The Science and Application Center for Moon and Deepspace Exploration. Chang’e 3data:Rover Panoramic Camera. http://planetary.s3.amazonaws.com/data/chang_e3/pcam.html. 14 May 2018
36. Polyu.edu.hk (2017) The Hong Kong Polytechnic University-The Space Exploration Journey [online]. https://www.polyu.edu.hk/web/filemanager/en/content_155/1960/Appendix_Milestones_SpaceProjects.pdf. Accessed 24 Nov 2017
37. Hore A, Ziou D (2010) Image quality metrics: PSNR vs. SSIM. In: Proceedings of the ICRP, pp. 2366–2369

Figure1.Generator network structure

Figure2.SE-ResBlock network structure

Figure3.Deconvolution and resize-convolution

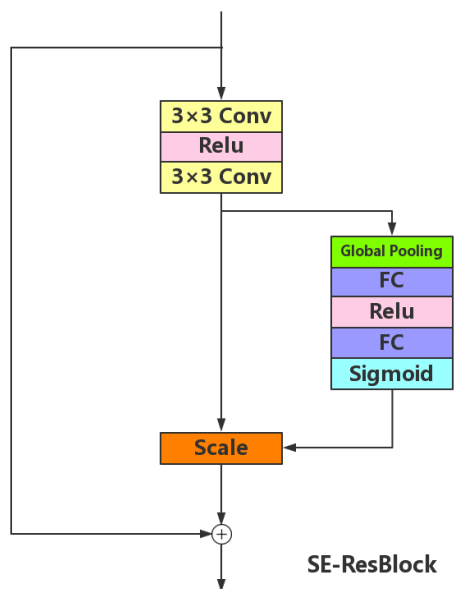
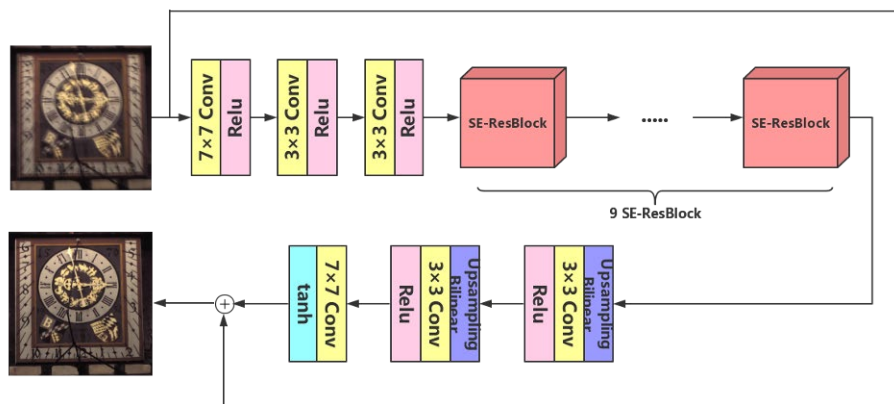
Figure4.Performance (PSNR/SSIM) with respect to the train epoch

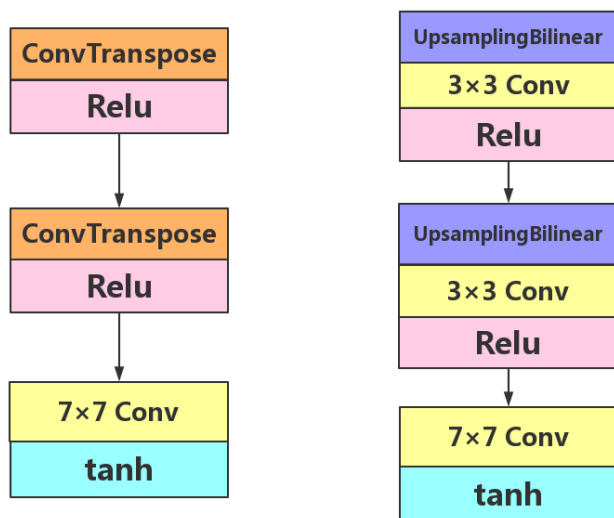
Figure5. Comparison of deblurred images by our model and DeblurGAN [11] on one of the image taken from GOPRO Dataset

Figure 6. Comparison of deblurred images by deconvolution and resize-convolution

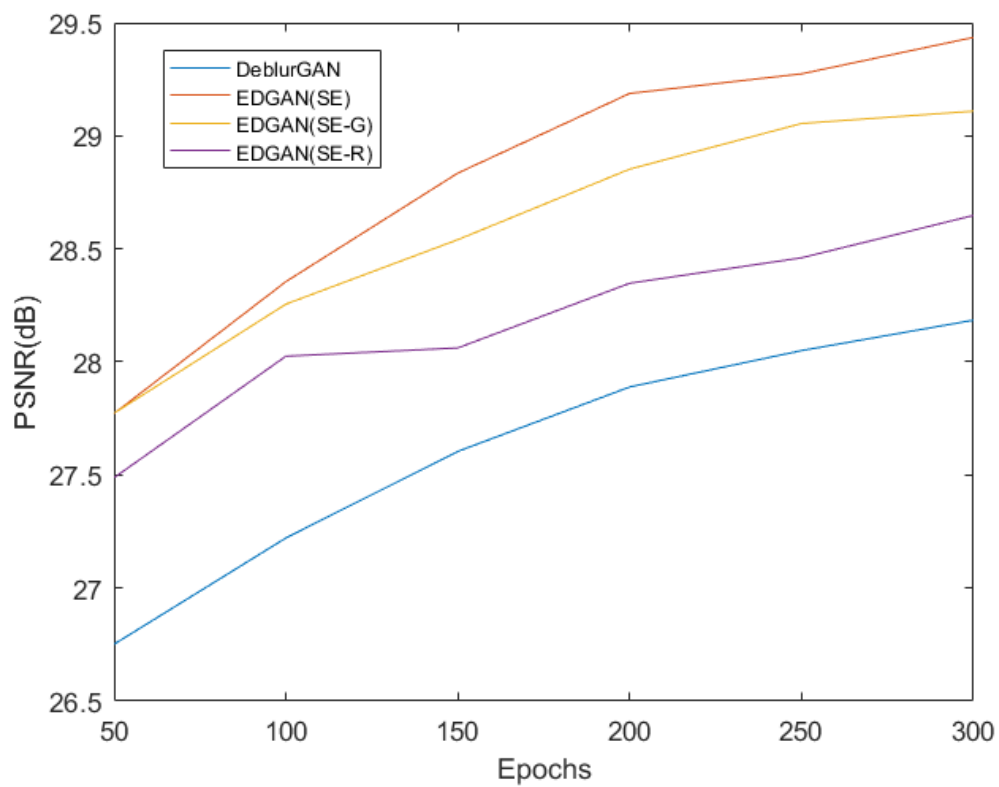
Figure7. Comparison of deblurred images by our model and DeblurGAN [11] on one of the image taken from Lai Dataset

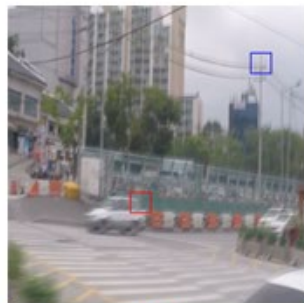
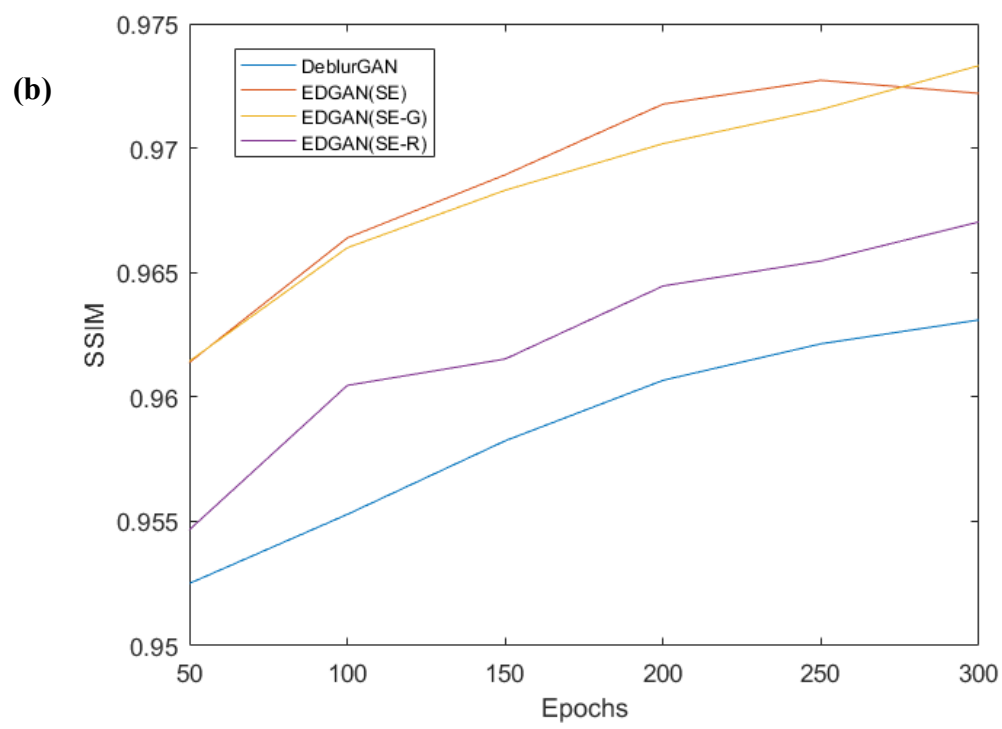
Figure 8.(a)(d)(g)blurred image. (b)(e)(h)ground truth. (c)(f)(i) Our deblurring result





(a)





(a) Blurred Image



(b) Ground Truth



(c) DeblurGAN[11]



(d) $EDGAN_{SE}$



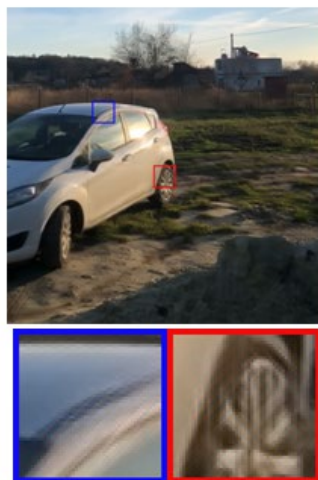
(e) $EDGAN_{SE-G}$



(f) $EDGAN_{SE-R}$



(a) Sharp Image



(b) $EDGAN_{SE}$



(c) $EDGAN_{SE-R}$



(a) Blurred Image



(b) Ground Truth



(c) DeblurGAN[11]



(d) $EDGAN_{SE}$



(e) $EDGAN_{SE-G}$



(f) $EDGAN_{SE-R}$



(a)



(b)



(c)



(d)



(e)



(f)



(g)



(h)



(i)