

# Supplemental Materials: Toward Fast, Flexible, and Robust Low-Light Image Enhancement

Long Ma<sup>†,§</sup>, Tengyu Ma<sup>†</sup>, Risheng Liu<sup>‡,\*</sup>, Xin Fan<sup>‡</sup>, Zhongxuan Luo<sup>†</sup>

<sup>†</sup>School of Software Technology, Dalian University of Technology

<sup>‡</sup>DUT-RU International School of Information Science & Engineering, Dalian University of Technology

<sup>§</sup>Peng Cheng Laboratory

{rsliu, xin.fan, zxlue}@dlut.edu.cn, {longma, matengyu}@mail.dlut.edu.cn

In this document, we supply some details and experiments of our proposed SCI. We first provide an elaborated exploration for the architecture  $\mathcal{K}_\theta$ . And then, we make the analysis of the weighting parameters in the loss function. Besides, more visual results for different tasks are also provided to further indicate our superiority.

## 1. Exploring Architecture of Self-Calibrated Module

Fig. 1 showed the visual results of different settings for self-calibrated module  $\mathcal{K}_\theta$ . We can easily see that our method always generated stable outputs no matter how the number of convolutions. Thus we can actually use the lightweight enough architecture to instantiate it. More importantly, the self-calibrated module is just adopted in the training phase. This is to say, the architecture for  $\mathcal{K}_\theta$  will be not appeared during the inference.

## 2. Analyzing Parameters in Loss

Here we considered some cases of different settings for weighting parameters in the loss function. As demonstrated in Fig. 2, we can observe that the exposure level became higher along with the increase of  $\alpha$ . But when  $\alpha = 2$ , the color appears the deviation that tends to present white. When increasing the  $\beta$ , the exposure level became low. By comparison, we set  $\alpha = 1, \beta = 1$  as our default setting in our all experiments.

## 3. More Visual Results

### 3.1. Low-Light Image Enhancement

As shown in Fig. 3-4, we demonstrated visual results of different learning-based methods including DRBN [9], KinD [12], EnGAN [4], SSIENet [11], ZeroDCE [2], and RUAS [5]. We can easily see that our method consistently

performed superiority over other state-of-the-art methods in different scenes.

### 3.2. Dark Face Detection

Fig 5 showed the visual comparison among different methods on the DARK FACE dataset [10]. We can easily see that our method can detect more targets, but other methods cannot do it.

### 3.3. Nighttime Semantic Segmentation

In Fig. 6, we provided more visual comparison of nighttime semantic segmentation among different methods. Obviously, it can be observed that the results of our methods can accurately segment more areas than others.

## References

- [1] Vladimir Bychkovsky, Sylvain Paris, Eric Chan, and Frédo Durand. Learning photographic global tonal adjustment with a database of input/output image pairs. In *Proceeding of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 97–104, 2011. 3
- [2] Chunle Guo, Chongyi Li, Jichang Guo, Chen Change Loy, Junhui Hou, Sam Kwong, and Runmin Cong. Zero-reference deep curve estimation for low-light image enhancement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2021. 1
- [3] Jiang Hai, Zhu Xuan, Ren Yang, Yutong Hao, Fengzhu Zou, Fang Lin, and Songchen Han. R2rnet: Low-light image enhancement via real-low to real-normal network. *arXiv preprint arXiv:2106.14501*, 2021. 3
- [4] Yifan Jiang, Xinyu Gong, Ding Liu, Yu Cheng, Chen Fang, Xiaohui Shen, Jianchao Yang, Pan Zhou, and Zhangyang Wang. Enlightengan: Deep light enhancement without paired supervision. *IEEE Transactions on Image Processing*, 30:2340–2349, 2021. 1
- [5] Risheng Liu, Long Ma, Jiaao Zhang, Xin Fan, and Zhongxuan Luo. Retinex-inspired unrolling with cooperative prior architecture search for low-light image enhancement. In *Pro-*

\*Corresponding author.

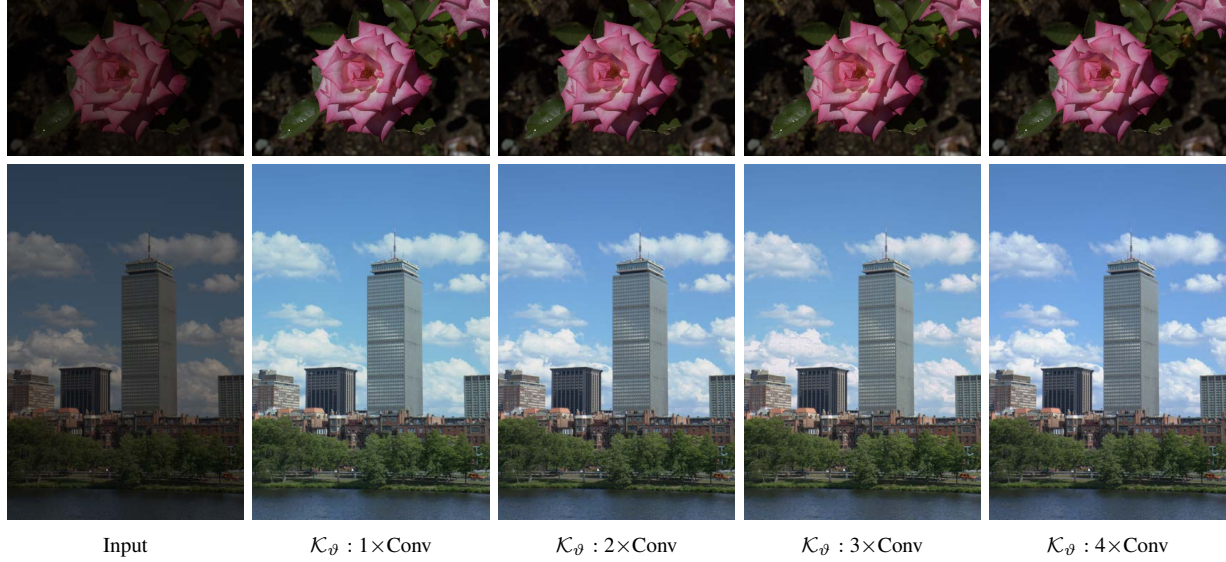


Figure 1. Visual results of different settings for  $\mathcal{K}_\theta$ .

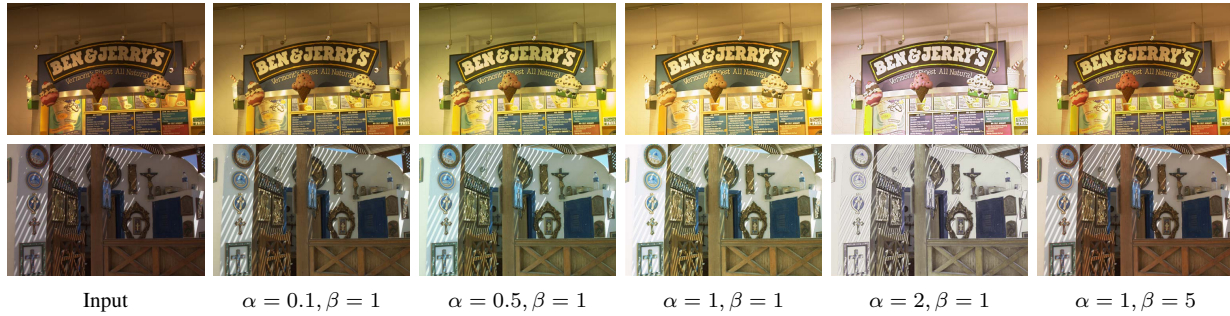


Figure 2. Visual results of different settings for weighting parameters in loss functions.

- ceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pages 10561–10570, 2021. 1
- [6] Yuen Peng Loh and Chee Seng Chan. Getting to know low-light images with the exclusively dark dataset. *Computer Vision and Image Understanding*, 178:30–42, 2019. 4
- [7] Hajime Nada, Vishwanath A Sindagi, He Zhang, and Vishal M Patel. Pushing the limits of unconstrained face detection: a challenge dataset and baseline results. In *2018 IEEE 9th International Conference on Biometrics Theory, Applications and Systems (BTAS)*, pages 1–10. IEEE, 2018. 4
- [8] Christos Sakaridis, Dengxin Dai, and Luc Van Gool. Acdc: The adverse conditions dataset with correspondences for semantic driving scene understanding. *arXiv preprint arXiv:2104.13395*, 2021. 6
- [9] Wenhan Yang, Shiqi Wang, Yuming Fang, Yue Wang, and Jiaying Liu. From fidelity to perceptual quality: A semi-supervised approach for low-light image enhancement. In *Proceeding of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3063–3072, 2020. 1
- [10] Wenhan Yang, Ye Yuan, Wenqi Ren, Jiaying Liu, Walter J Scheirer, Zhangyang Wang, Taiheng Zhang, Qiaoyong Zhong, Di Xie, Shiliang Pu, et al. Advancing image understanding in poor visibility environments: A collective benchmark study. *IEEE Transactions on Image Processing*, 29:5737–5752, 2020. 1, 5
- [11] Yu Zhang, Xiaoguang Di, Bin Zhang, and Chunhui Wang. Self-supervised image enhancement network: Training with low light images only. *arXiv*, pages arXiv–2002, 2020. 1
- [12] Yonghua Zhang, Xiaojie Guo, Jiayi Ma, Wei Liu, and Jiawan Zhang. Beyond brightening low-light images. *International Journal of Computer Vision*, pages 1–25, 2021. 1

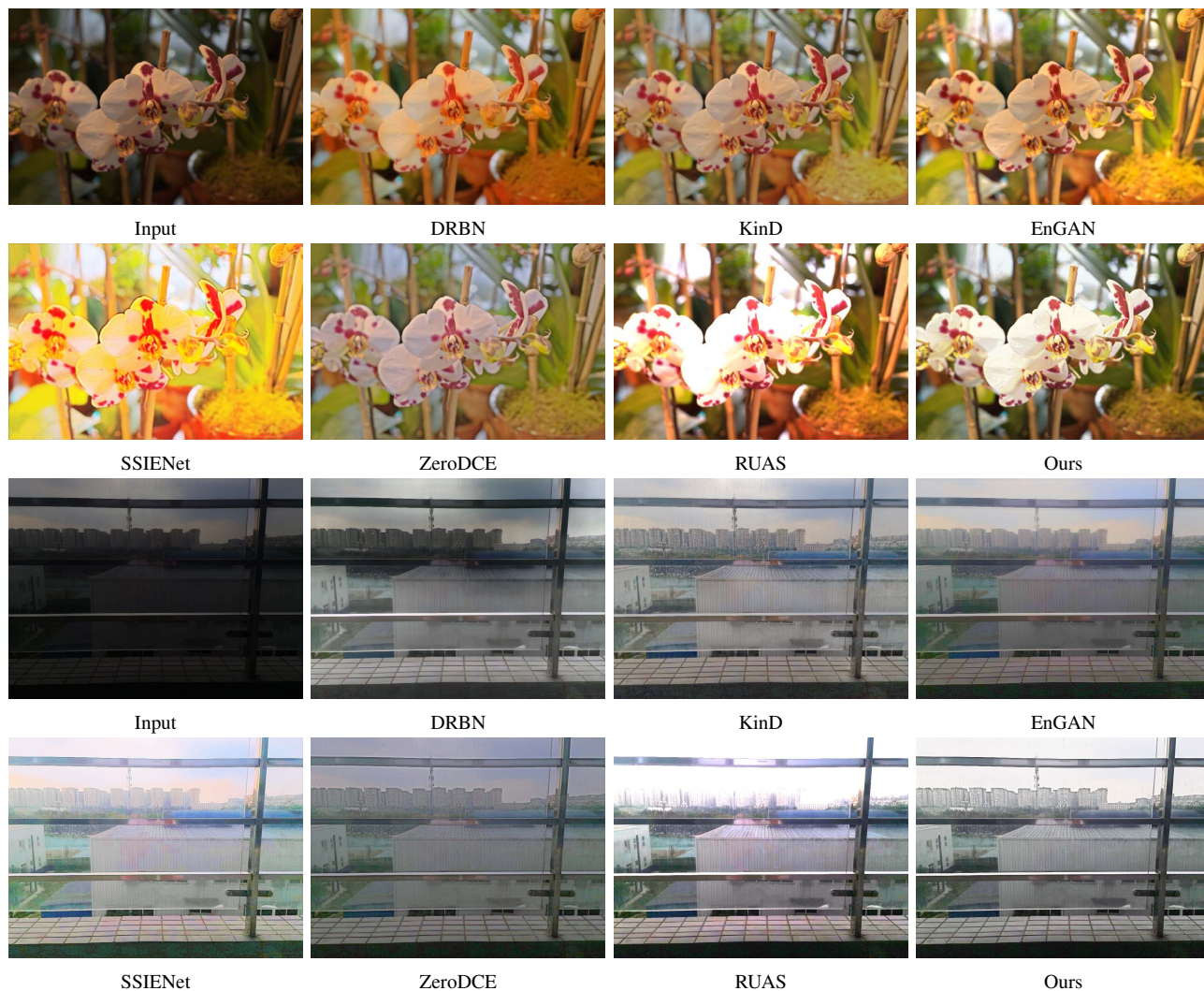


Figure 3. Visual results of state-of-the-art methods and our method on benchmarks. Top and bottom rows are results on the MIT-Adobe 5K [1] and LSRW [3] datasets, respectively.





Figure 4. Visual results of state-of-the-art methods and our method on real-world scenarios. Top three and bottom five rows are results on the ExDark [6] and UFDD [7] datasets, respectively.





Figure 5. Visual results of dark face detection on the DARK FACE dataset [10]. Red box indicates the obvious differences.

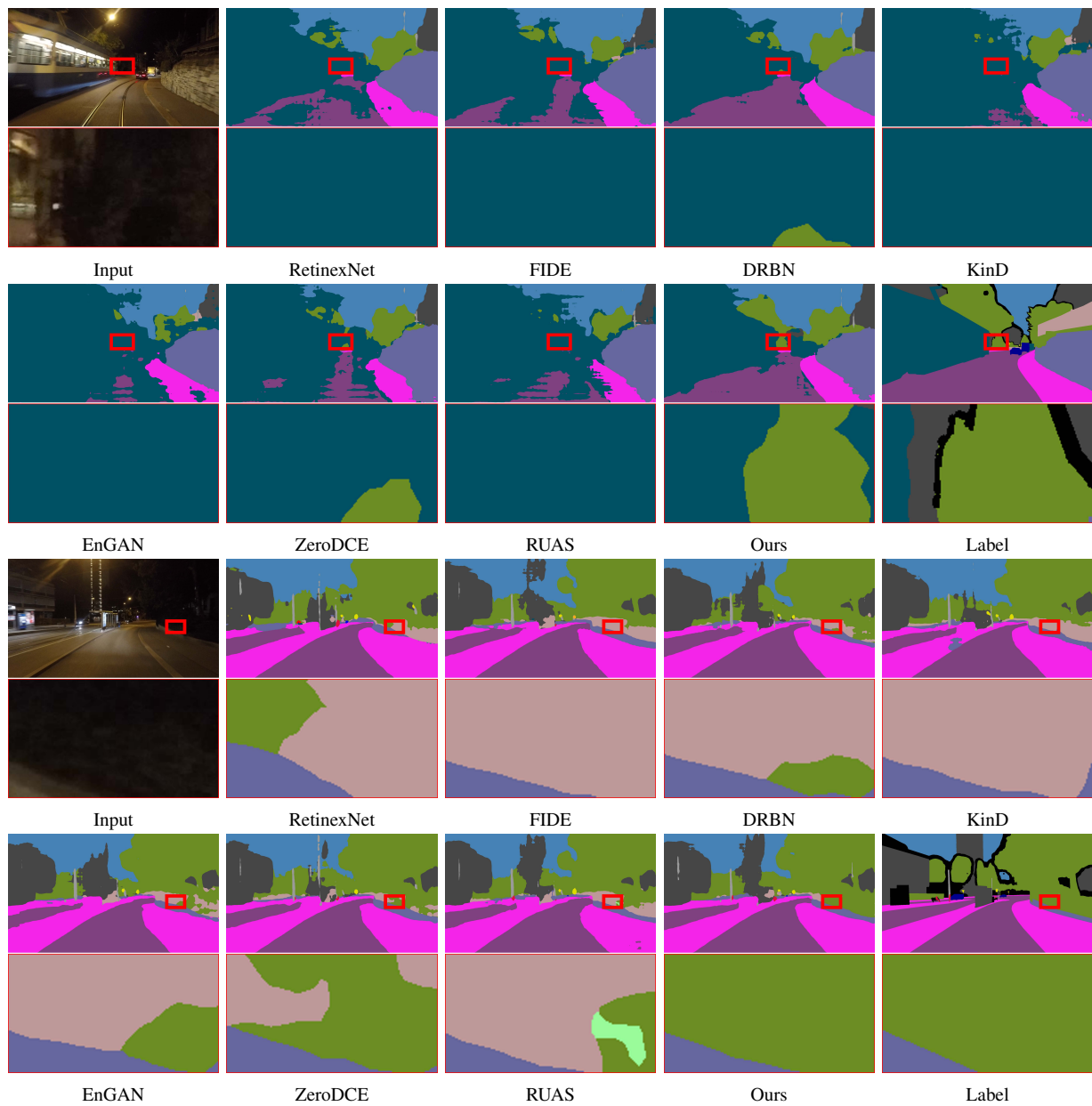


Figure 6. Visual results of semantic segmentation on the ACDC dataset [8]. Red box indicates the obvious differences.