

Exposure Normalization and Compensation for Multiple Exposure Correction

Supplementary Material

This supplementary document is organized as follows:

- Sec. 1 provides the details of the Table 1 in the main body.
- Sec. 2 provides numerical results of applying our method on another backbone.
- Sec. 3 provides the results of evaluating different methods by NIQE metric.
- Sec. 4 provides numerical results of more comparison methods.
- Sec. 5 provides numerical results of comparing the ENC module with other plug-and-play modules.
- Sec. 6 provides more ablation studies of applying the parameter regularization strategy.
- Sec. 7 provides the derivation process of the parameter regularization strategy.
- Sec. 8 provides a discussion of the parameter regularization strategy with the EWC and more related results.
- Sec. 9 provides more details about the implementation.
- Sec. 10 provides more explanation of our proposed method.
- Sec. 11 provides more visualization results of feature maps in the ENC module.
- Sec. 12 provides more visualization results of exposure correction.

1. The details of Table 1 in the main body

Since the ME dataset contains 5 levels of exposure, presenting all of their results in the main body is difficult due to the limitation of the page length. Here, we present the details of Table 1 in the main body in Table 1 and Table 2, respectively. In particular, the exposure level 1 to 5 represent the exposure from underexposure to overexposure. More descriptions of this dataset are provided in [1].

We regard the results of Level 1 and Level 2 in Table 1 as the underexposure results of the ME dataset in the main body, and the results of Level3, Level 4 and Level 5 are regarded as the overexposure results. Similarly, the results of Level 1 and Level 2 in Table 2 are regarded as the underexposure results of the ME-v2 dataset in the main body, while the results of Level 4 and Level 5 are regarded as the overexposure results. Besides, we fix a typo in the main body, the PSNR of MSEC method is 24.39dB instead of 24.29dB that reported in the main body, which dose not affect the illustration of demonstrating our method.

Method	Level 1		Level 2		Level 3		Level 4		Level 5		Average		Parameters
	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	
CLAHE	16.84	0.6319	16.70	0.6103	13.18	0.5756	16.01	0.5933	14.16	0.5837	15.38	0.5990	-
RetinexNet	13.21	0.6317	11.05	0.6101	10.75	0.6038	10.45	0.5964	10.01	0.5857	11.14	0.6048	1.70M
Zero-DCE	14.99	0.5946	14.11	0.5828	12.15	0.5418	10.06	0.5171	8.99	0.4837	12.06	0.5441	0.33M
MSEC	20.50	0.8105	20.53	0.8152	20.35	0.8210	19.78	0.8171	19.23	0.8086	20.08	0.8145	7.04M
DRBN	19.61	0.8243	19.86	0.8336	19.44	0.8355	19.45	0.8338	19.23	0.8271	19.52	0.8309	0.53M
DRBN-L	19.81	0.8291	19.87	0.8347	19.27	0.8344	19.64	0.8381	19.60	0.8339	19.64	0.8340	0.67M
I-DRBN(Ours)	21.72	0.8413	22.38	0.8538	22.50	0.8628	22.38	0.8542	21.63	0.8457	22.12	0.8516	0.54M
I-DRBN-4(Ours)	22.57	0.8508	22.87	0.8579	22.70	0.8603	22.29	0.8543	21.34	0.8418	22.35	0.8530	0.58M
SID	19.28	0.8072	19.46	0.8133	18.92	0.8103	18.88	0.8064	18.68	0.7999	19.04	0.8074	7.40M
SID-L	19.25	0.8078	19.38	0.8119	19.04	0.8100	19.09	0.8093	18.72	0.8026	19.10	0.8083	11.56M
I-SID(Ours)	22.39	0.8376	22.79	0.8470	22.70	0.8551	22.44	0.8522	21.94	0.8484	22.45	0.8481	7.45M

Table 1. Quantitative results of methods on ME Dataset in terms of PSNR and SSIM.

Method	Level 1		Level 2		Level 4		Level 5		Average		Parameters
	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	
CLAHE	18.13	0.8469	19.98	0.8647	16.99	0.7843	14.53	0.7440	17.41	0.8100	-
RetinexNet	13.83	0.7237	14.24	0.7471	16.56	0.7896	16.08	0.7666	15.18	0.7568	1.70M
Zero-DCE	16.32	0.7934	15.57	0.7889	10.49	0.7037	9.30	0.6637	12.92	0.7374	0.33M
MSEC	22.28	0.9010	26.67	0.9101	22.58	0.9118	26.01	0.9187	24.39	0.9104	7.04M
DRBN	25.40	0.9169	22.80	0.9143	22.52	0.9013	23.78	0.9096	23.63	0.9105	0.53M
DRBN-L	25.34	0.9170	23.35	0.9166	22.29	0.9025	24.39	0.9141	23.84	0.9126	0.67M
I-DRBN(Ours)	26.01	0.9119	24.63	0.9080	28.07	0.9345	25.38	0.9328	26.02	0.9218	0.54M
I-DRBN-4(Ours)	27.70	0.9243	26.92	0.9284	26.95	0.9234	26.51	0.9316	27.02	0.9270	0.58M
SID	24.80	0.9084	23.27	0.9060	21.08	0.8633	21.12	0.8705	22.57	0.8871	7.40M
SID-L	23.53	0.8979	22.61	0.8953	22.25	0.8704	22.25	0.8842	22.66	0.8870	11.56M
I-SID(Ours)	26.23	0.9145	25.94	0.9121	26.02	0.9154	27.21	0.9177	26.35	0.9149	7.45M

Table 2. Quantitative results of methods on ME-v2 Dataset in terms of PSNR and SSIM.

2. Applying our method on another backbone

Since our method can be plugged into existing backbone, we apply our method on MIRNet [13], which is a network designed for various kinds of image restoration tasks. Since its original model is too large, we reduce its channel dimension for fast training. We conduct experiment on Task-mix dataset for fast evaluation. The numerical results in terms of PSNR are shown in Table 3, with the employ of our method, the performance has improved remarkably, demonstrating the effectiveness of our method.

Method	LOL	SICE OVER	FIVEK	Average
MIRNet	20.72	19.75	23.86	21.45
I-MIRNet (Ours)	21.44	20.48	23.91	21.95

Table 3. Quantitative results of applying our method on MIRNet on Task-mix dataset.

3. Evaluating different methods by NIQE metric

We further evaluate different methods by NIQE metric on SICE dataset. As shown in Table 4, our method can improve the performance on niqe metric a little, but it can boost the performance of PSNR and SSIM significantly as have been shown in the main body, demonstrating the robustness of our method.

Method	Input	HE	CLAHE	LIME	Zero-DCE	RetinexNet	MSEC	SID	DRBN	I-SID	I-DRBN-4	GT
NIQE	4.12	3.62	3.80	4.53	3.70	3.85	5.06	3.40	3.32	3.34	3.33	3.2145

Table 4. NIQE evaluation of different methods on the SICE dataset. We average the underexposure and overexposure results in this table.

4. The quantitative comparison of more methods

Due to the limitation of the page length, we present more numerical results of other baseline methods on ME dataset refer to [1]. Besides the methods we have compared in the main body, other methods include DPED [7], DeepUPE [11], LIME [6], DPE [3], WVM [5], HQEC [14] and HDR CNN [4]. As shown in Table 5, since our proposed method outperforms MSEC [1] remarkably, we also get better results than all of the baseline methods, which fully demonstrate the superiority of our method.

Method	CLAHE	WVM	LIME	HDR CNN	DPE	DPED	HQEC	RetinexNet	DeepUPE	Zero-DCE	I-DRBN-4
PSNR	15.38	15.12	11.52	16.43	18.02	17.74	13.91	11.14	13.69	12.06	22.35
SSIM	0.599	0.678	0.607	0.681	0.683	0.696	0.656	0.605	0.632	0.544	0.853

Table 5. More quantitative results of other baseline methods on the ME dataset.

5. Comparing the ENC module with other plug-and-play modules

To the best of our knowledge, this is the first plug-and-play module for existing exposure correction frameworks. Following the experimental setting in similar works in other tasks, we compared our method with other exposure correction approaches, and we additionally increased the channel number for comparison. We adopt several plug-and-play modules that have been widely used in other vision tasks for comparison, including Convolutional Block Attention Module (CBAM) [12], Atrous Spatial Pyramid Pooling (ASPP) module [2] and Selective Kernel Convolution (SKC) module [10]. The results in Table 6 show that our ENC module is more effective than other ones in multiple exposure correction.

Metric	DRBN	DRBN-CBAM	DRBN-ASPP	DRBN-SKC	DRBN-ENC
PSNR (dB)	17.65	19.57	19.93	19.64	20.49
#Param (M)	0.53	0.57	0.58	0.59	0.58

Table 6. Results of incorporating other modules on SICE dataset.

6. More ablation studies of applying our Parameter regularization strategy

We provide more numerical results of applying the parameter regularization strategy on ME-v2 dataset. For our improved DRBN network, since it has the lowest PSNR in the Level 5 exposure subset, we regard it as the worst-performed exposure for our DRBN network. While for our improved SID network, it has the lowest PSNR in the Level 1 exposure subset, therefore we regard it as the worst-performed exposure for our SID network.

Method	Level 1	Level 2	Level 4	Level 5	Average
DRBN-ENC	27.84/0.9268	26.91/0.9242	26.93/0.9242	26.24/0.9298	26.98/0.9277
DRBN-ENC-4-SEQ	10.77/0.4818	11.61/0.5609	22.36/0.9242	32.69/0.9492	19.32/0.7290
I-ENC-4 (Ours)	27.70/0.9243	26.92/0.9284	26.95/0.9234	26.51/0.9316	27.02/0.9270
SID-ENC	25.89/0.9113	25.97/0.9136	26.04/0.9170	27.25/0.9183	26.29/0.9151
SID-ENC-SEQ	32.32/0.9393	22.33/0.9277	9.33/0.6606	8.05/0.5958	18.00/0.7809
I-SID (Ours)	26.23/0.9145	25.94/0.9121	26.02/0.9154	27.21/0.9177	26.35/0.9149

Table 7. Ablation studies for the parameter regularization strategy on ME-v2 dataset.

The results are shown in Table 7, after applying the Parameter regularization strategy, we can see that the performance of worst-performed exposure has improved with little performance drop of other exposures.

7. The derivation process and more details of the parameter regularization strategy

Due to the limitation of the page length, we simplify the derivation process of the parameter regularization strategy in the main body. Here, we present the details of derivation process and more details of the parameter regularization strategy.

Firstly, to simplify, we denote training on various exposures as Task 0 and fine-tuning on the worst-performed exposure as Task 1. Suppose the various exposures and the worst-performed exposure are represented as x^0 and x^1 , respectively. The parameters of the network trained on Task 0 are denoted by $\theta^0 = \theta_1^0, \dots, \theta_m^0$ while that of Task 1 are denoted by $\theta^1 = \theta_1^1, \dots, \theta_m^1$. After training the network on Task 1, the resulting degradation on the previous Task 0 can be evaluated by:

$$\Delta f = f(x^0; \theta^1) - f(x^0; \theta^0), \quad (1)$$

where f denotes the network. Taking the element of parameter θ_k^0 (k -th depth) for example, the change of parameter θ_k^0 is denoted as $\delta\theta_k^0$ when model is trained on the new Task 1, the mathematical form is $\delta\theta_k^0 = \theta_k^1 - \theta_k^0$. Then, we take the Taylor expansion of $f(x^0; \theta_k^1)$ at point θ_k^0 :

$$f(x^0; \theta_k^1) = f(x^0; \theta_k^0) + \left(\nabla_{\theta_k^0} f \right)^T \cdot \delta\theta_k^0 + \frac{1}{2} (\delta\theta_k^0)^T \cdot \nabla_{\theta_k^0}^2 f \cdot \delta\theta_k^0 + O(\|\delta\theta_k^0\|^3). \quad (2)$$

Inspired by Gauss-Newton method, we approximate the $\nabla_{\theta_k^0}^2$ to relieve the computational cost as:

$$\mathbb{E}_{x_0 \sim \mathbb{P}^0} [\nabla_{\theta_k^0}^2 f] \approx 2 \times \mathbb{E}_{x_0 \sim \mathbb{P}^0} [\nabla f]^T \mathbb{E}_{x_0 \sim \mathbb{P}^0} [\nabla f]. \quad (3)$$

And, we inject the Eq. 2 into Eq. 1, and acquire the weight importance as:

$$\Delta f = \left(\nabla_{\theta_k^0} f \right)^T \cdot \delta \theta_k^0 + \frac{1}{2} (\delta \theta_k^0)^T \cdot \left(\nabla_{\theta_k^0}^2 f \right)^T \cdot \nabla_{\theta_k^0} f \cdot \delta \theta_k^0. \quad (4)$$

To maintain the performance of previous Task 0, we need to minimize the Eq. 1. From this motivation, when training the model on Task 1, we add a regularization term based on conventional loss to keep the knowledge of Task 0. In summary, the total loss on Task 1 is a composite loss, which is of the form:

$$\begin{aligned} \mathcal{L}' &= \mathcal{L} + \lambda \Delta f \\ &= \mathcal{L} + \frac{\lambda}{2} \sum_{k=1}^m \left[I_1(\theta_k^0)^T |\delta \theta_k^0| + |\delta \theta_k^0|^k \left(\nabla_{\theta_k^0} f \right)^2 |\delta \theta_k^0| \right]. \end{aligned} \quad (5)$$

Note that the Eq. 5 is another form of Eq. (11) in the main body, which describes the implementation of the parameter regularization strategy more concretely.

The pseudo code of the paramter regularization strategy method is summarized in Algorithm 1.

Algorithm 1 Parameter regularization strategy for fine-tuning.

Input: Training phase task T_0 , fine-tuning phase task T_1 ; conventional training loss $\mathcal{L}$

Output: Final trained model for correcting all exposures.

```

for Task  $T_0$  do
    Training Task  $T_0$  using conventional loss  $\mathcal{L}$ 
    if the epoch performs the best then
        calculating the parameter importance with Eq.(4)
    end if
end for
for Task  $T_1$  do
    update the loss  $\mathcal{L}$  with Eq. (5) as composite loss
    train the network with the composite loss
end for

```

8. The discussion of the parameter regularization strategy with the EWC and more related results

The proposed parameter regularization strategy focuses on addressing the imbalanced performance across exposures rather than designing a new algorithm. Our strategy is inspired by Elastic Weight Consolidation (EWC) [8], which enables a single model to maintain the performance of previous tasks when being trained on a new task. We initially employed it to fine-tune the worst-performed exposure, yet its first-order scheme is inaccurate to calculate the parameter importance, resulting in ineffective performance (see the second line in Table 8). Thus, our method extends EWC to the first and second-order scheme to strengthen the calculation of the parameter importance, achieving better performance.

Besides, we include the simple re-weighting approach for comparison, which increases the weight coefficient of the worst-performed samples. Our strategy can improve the performance of the worst-performed exposure and keep the performance of other exposures, thus improving the overall performance. In contrast, the re-weighting approach lacks the constraint of keeping the performance for other exposures. Table 8 shows that the average results of the re-weighting approach are comparable with that of baseline, while our method surpasses it and baseline on average.

9. More details about the implementation

We present more details about the implementation in this section.

About the features of normal exposures. In the training phase, we first train the network on the multiple exposure datasets for a few epochs, and denote the network as *Net1*. Then the normal exposures are inference by *Net1*, the features of the exposure normalization part F_n^{norm} and the ENC's output F_f^{norm} are utilized to supervise the multiple exposure's features on another same structure network *Net2*. In particular, We implement the supervise manner by the normalization distilling loss L_{nd} and the exposure distilling loss L_{ed} .

Method	Under	Over	Average
DRBN-ENC (Baseline)	21.89	19.09	20.49
DRBN-ENC-EWC	21.21	19.64	20.43
DRBN-ENC-Weight-1.5	21.52	19.41	20.47
DRBN-ENC-Weight-3.0	21.30	19.74	20.52
DRBN-ENC-Weight-6.0	20.81	19.92	20.37
I-DRBN-4 (Ours)	21.77	19.57	20.67

Table 8. PSNR results of the parameter regularization strategy related results of DRBN-based methods on the SICE dataset.

About the coefficients of the losses. The losses in our framework consist of three parts: the conventional loss of baseline network, the normalization distilling loss L_{nd} and the exposure distilling loss L_{ed} . Since the conventional loss of different baseline networks is different, to balance the losses in the framework, we set the coefficients of L_{nd} and L_{ed} to ensure their values occupy 0.1 of the total loss, respectively. Besides, for the network with several ENCs, we only apply L_{nd} and L_{ed} on the first ENC to improve the performance.

10. More explanation of the proposed method

In the main body, we describe our method improves the performance by narrowing the gap of different exposures. Here, we further explain how our method works. Specifically, according to the Retinex theory [9], one image I can be decomposed as:

$$I = R \cdot S, \quad (6)$$

denoting that an image can be formed via a pixel-wise multiplication of R which denotes the reflectance image, and S which is the illumination map. Therefore, multiple exposure correction can be regarded as the recovery of R or S . Since ENC can narrow the gap of different exposure representations, it is analogous to learning R that is also exposure-invariant, while the following part in the network implicitly recover S . In this way, the problem of multiple exposure correction is decomposed to recover R and S , thus relieving the difficulty of multiple exposure correction.

11. More visualization results of feature maps in the ENC module

In this section, we present more visualization results of feature maps in the ENC module, demonstrating the effectiveness of the exposure normalization part, normalization distilling loss and exposure distilling loss, respectively.

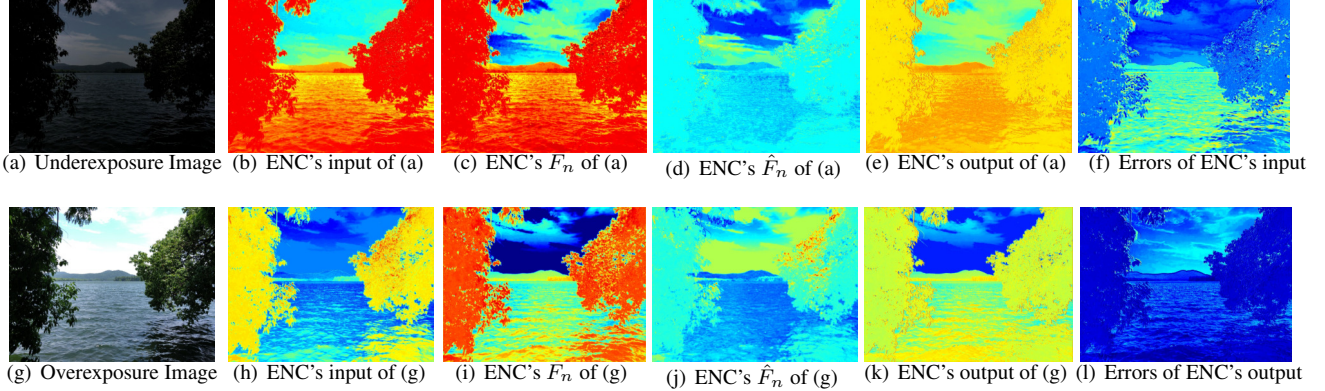


Figure 1. Feature visualization of different components in ENC on samples from SICE dataset.

Fig. 1 and Fig. 2 present two examples of different components in the ENC module. As can be seen, the ENC's input features F from underexposure and overexposure differ greatly shown in (b) and (h). As shown in (c) and (i) as well as (d) and (j), after being processed by the exposure normalization part, input features are mapped to the exposure invariant space and their discrepancies are progressively reduced. With ENC, the gap of representations between underexposure and overexposure is obviously narrowed as illustrated in (e) and (k) as well as (f) and (l).

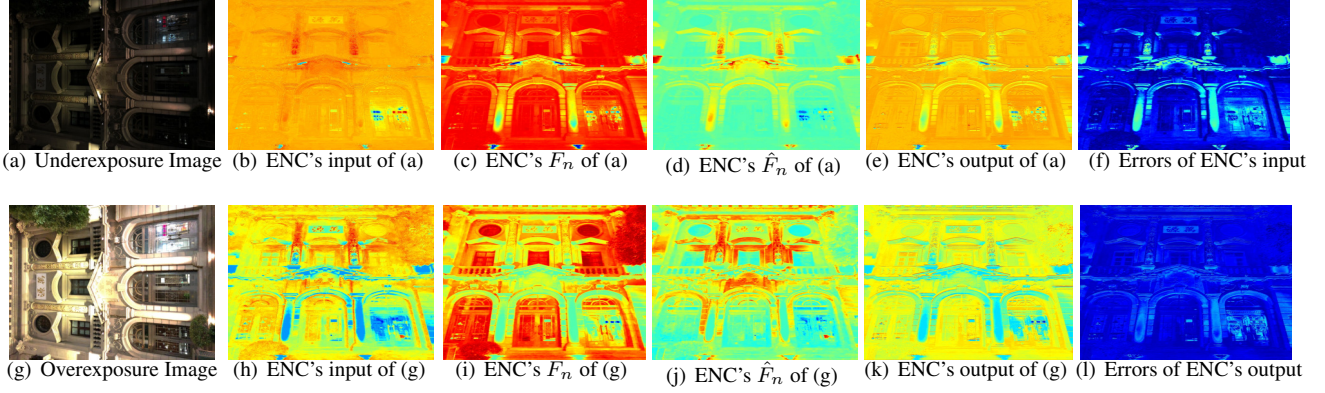


Figure 2. Feature visualization of different components in ENC on samples from the SICE dataset.

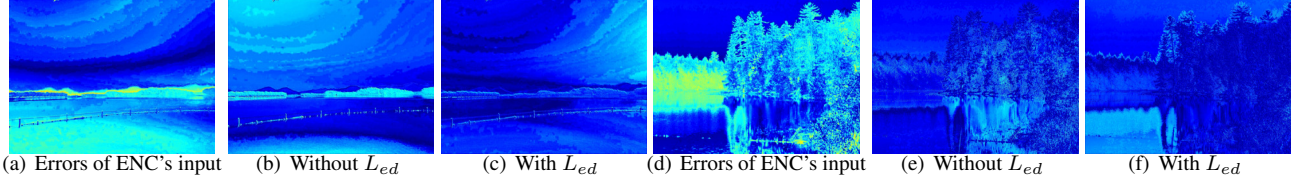


Figure 3. The visualization of the ENC's input and output errors between underexposure and overexposure on samples from the SICE dataset.

Fig. 3 presents two examples of ENC's output features errors. Compare (b) and (c) with (a), as well as (e) and (f) with (d), we can see that ENC can narrow the gap of different exposures significantly, which further demonstrates the effectiveness of the ENC module. While comparing (b) with (c) as well as (d) with (f), it can be seen that with the employing of exposure distilling loss, the errors of ENC's output are further reduced, demonstrating the effectiveness of using exposure distilling loss.

12. More visualization results of exposure correction

Due to the limitation of the page length, in this section, we present more visualization results of different methods on ME dataset, ME-v2 dataset, SICE dataset and Brighten dataset, respectively.

In particular, Fig. 4 and Fig. 5 present the visualization results on the ME dataset, and we choose the exposure subsets of Level 1 and Level 5 as the underexposure and overexposure, respectively. Fig. 6 presents the visualization results on the ME-v2 dataset. Fig. 7 and Fig. 8 present the visualization results on the SICE dataset. Fig. 9 presents the visualization results on the Brighten dataset, and the model is trained on Task-mix dataset, which demonstrates the generalization ability of our proposed method.

References

- [1] Mahmoud Afifi, Konstantinos G Derpanis, Bjorn Ommer, and Michael S Brown. Learning multi-scale photo exposure correction. In *CVPR*, 2021. 1, 2
- [2] Liang-Chieh Chen, George Papandreou, Florian Schroff, and Hartwig Adam. Rethinking atrous convolution for semantic image segmentation. *arXiv preprint arXiv:1706.05587*, 2017. 3
- [3] Yu-Sheng Chen, Yu-Ching Wang, Man-Hsin Kao, and Yung-Yu Chuang. Deep photo enhancer: Unpaired learning for image enhancement from photographs with gans. In *CVPR*, pages 6306–6314, 2018. 2
- [4] Gabriel Eilertsen, Joel Kronander, Gyorgy Denes, Rafał Mantiuk, and Jonas Unger. Hdr image reconstruction from a single exposure using deep cnns. *ACM Transactions on Graphics (TOG)*, 36(6), 2017. 2
- [5] Xueyang Fu, Delu Zeng, Yue Huang, Xiao-Ping Zhang, and Xinghao Ding. A weighted variational model for simultaneous reflectance and illumination estimation. In *CVPR*, pages 2782–2790, 2016. 2
- [6] Xiaojie Guo, Yu Li, and Haibin Ling. Lime: Low-light image enhancement via illumination map estimation. *IEEE Transactions on image processing*, 26(2):982–993, 2016. 2
- [7] Andrey Ignatov, Nikolay Kobyshev, Radu Timofte, Kenneth Vanhoey, and Luc Van Gool. Dslr-quality photos on mobile devices with deep convolutional networks. In *ICCV*, pages 3277–3285, 2017. 2
- [8] James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A. Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, Demis Hassabis, Claudia Clopath, Dharshan Kumaran, and Raia Hadsell. Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114(13):3521–3526, 2017. 4
- [9] Edwin H Land. The retinex theory of color vision. *Scientific american*, 237(6):108–129, 1977. 5
- [10] Xiang Li, Wenhai Wang, Xiaolin Hu, and Jian Yang. Selective kernel networks. In *CVPR*, pages 510–519, 2019. 3
- [11] Ruixing Wang, Qing Zhang, Chi-Wing Fu, Xiaoyong Shen, Wei-Shi Zheng, and Jiaya Jia. Underexposed photo enhancement using deep illumination estimation. In *CVPR*, pages 6849–6857, 2019. 2
- [12] Sanghyun Woo, Jongchan Park, Joon-Young Lee, and In So Kweon. Cbam: Convolutional block attention module. In *ECCV*, pages 3–19, 2018. 3
- [13] Syed Waqas Zamir, Aditya Arora, Salman Khan, Munawar Hayat, Fahad Shahbaz Khan, Ming-Hsuan Yang, and Ling Shao. Learning enriched features for real image restoration and enhancement. In *ECCV*, pages 492–511, 2020. 2
- [14] Qing Zhang, Ganzhao Yuan, Chunxia Xiao, Lei Zhu, and Wei-Shi Zheng. High-quality exposure correction of underexposed photos. In *ACM MM*, page 582–590, 2018. 2

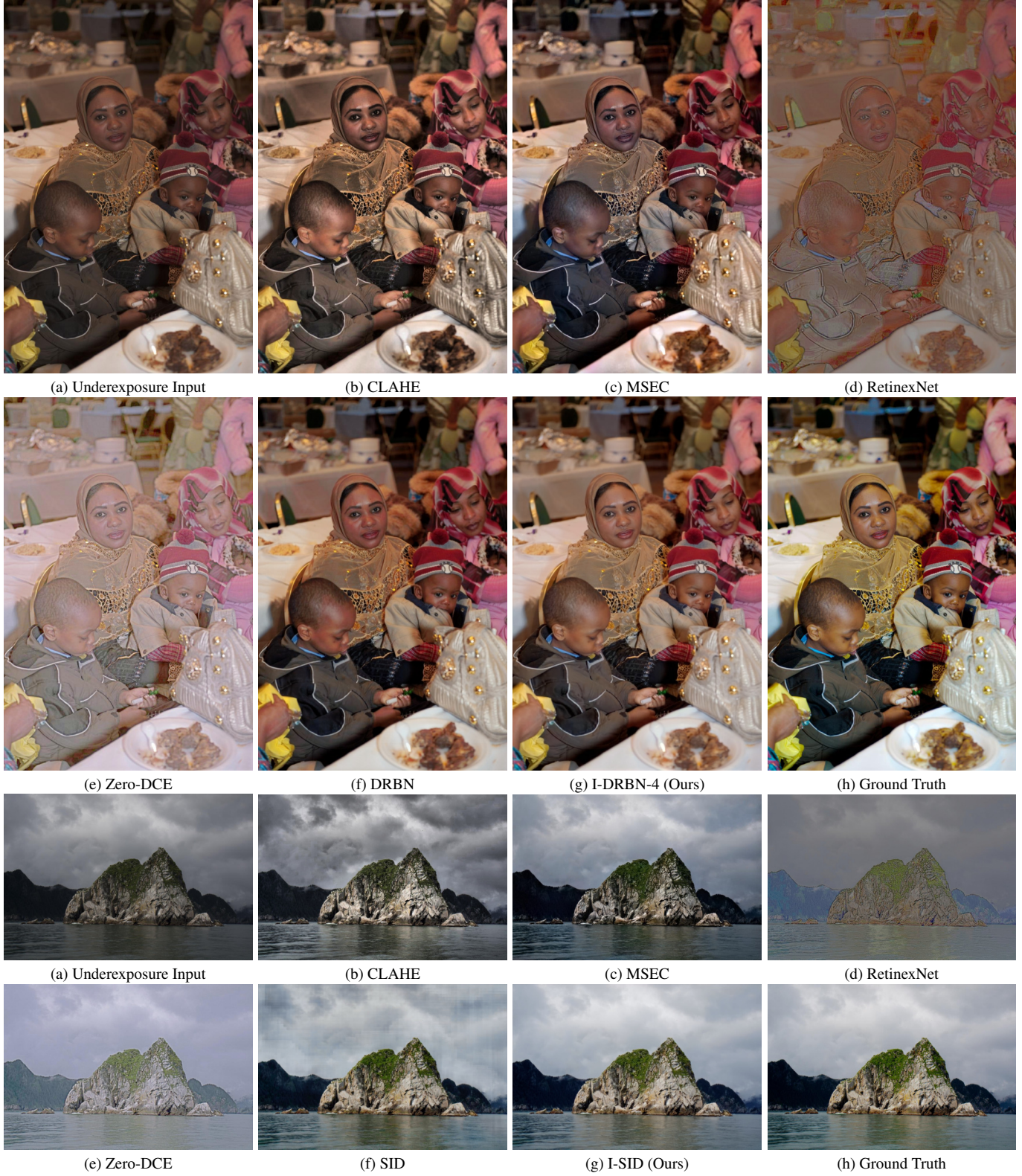


Figure 4. Visualization results of correcting underexposure images from the ME dataset. Please zoom in for better visualization.



Figure 5. Visualization results of correcting overexposure images from the ME dataset. Please zoom in for better visualization.



Figure 6. Visualization results of correcting underexposure and overexposure images from the ME-v2 dataset. Please zoom in for better visualization.

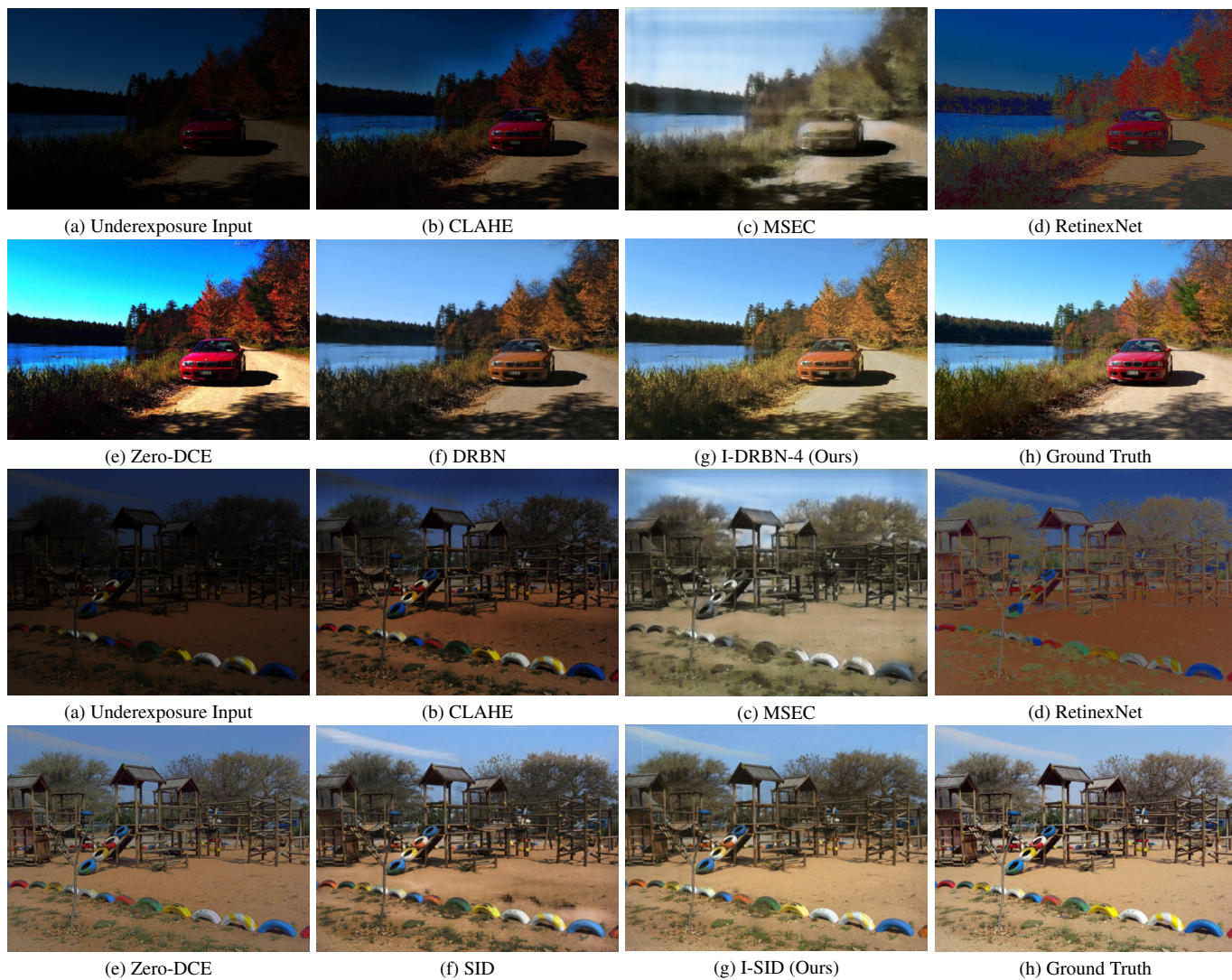


Figure 7. Visualization results of correcting underexposure images from the SICE dataset. Please zoom in for better visualization.

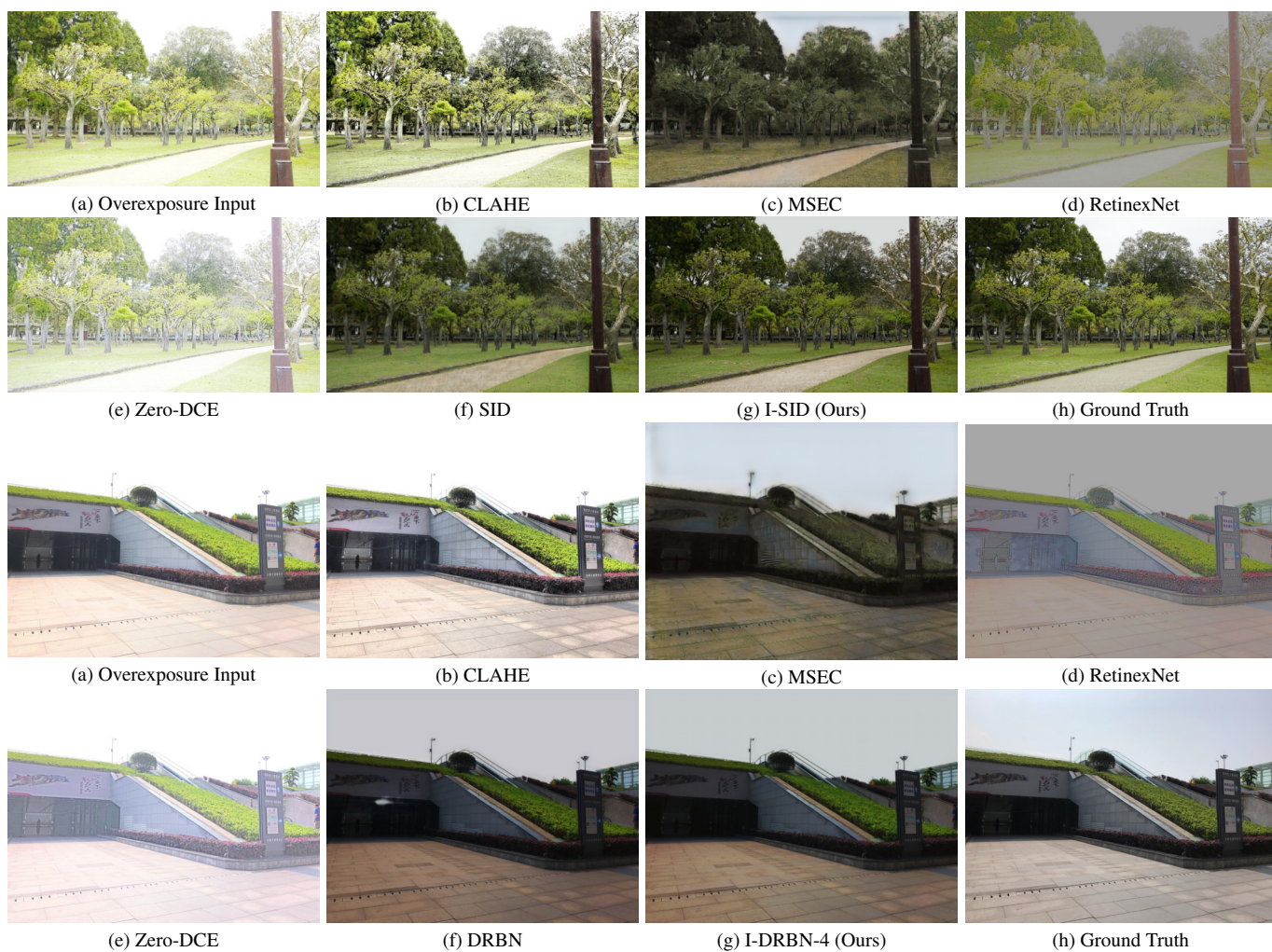


Figure 8. Visualization results of correcting overexposure images from the SICE dataset. Please zoom in for better visualization.



Figure 9. Visualization results on the Brighten dataset. Please zoom in for better visualization.