

2025 동계 세미나

Retinex-based Low-light Image Enhancement



Sogang University

Vision & Display Systems Lab, Dept. of Electronic Engineering



Presented By

김수훈

Contents

- Introduction
 - Retinex theory
 - Latent Diffusion Model
- Retinexformer: One-stage Retinex-based Transformer for Low-light Image Enhancement
 - ICCV 2023
- LightenDiffusion: Unsupervised Low-Light Image Enhancement with Latent-Retinex Diffusion Models
 - ECCV 2024

Introduction

- Retinex theory

- 인간 시각 시스템이 밝기와 색을 어떻게 인식하는지에 대한 이론

- 조명 조건이 달라져도 같은 물체를 동일한 색으로 인식할 수 있는 색 항등성 현상 설명

- ※ 단순히 물체가 반사하는 빛의 강도를 보는 것이 아니라, 조명과 물체의 반사 특성을 동시에 고려하여 색을 판단

- 반사율(Reflectance), 조명(Illumination)의 분리

- ※ 반사율(Reflectance)

- ✓ 조명 조건이 변해도 변하지 않는 본질적인 색과 질감

- ※ 조명(Illumination)

- ✓ 장면의 빛의 세기와 방향

$$\mathbf{I} = \mathbf{R} \odot \mathbf{L}$$

Introduction

- Latent Diffusion Model

- Diffusion Model

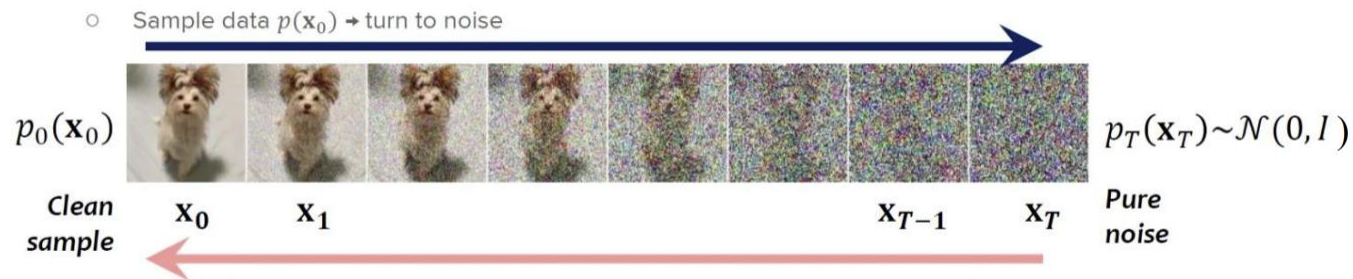
- Forward process

- ☼ 이미지 x_0 가 순수한 Gaussian noise x_T 가 될 때까지 Gaussian noise를 점진적으로 추가하는 Markov process

- Reverse denoising process

- ☼ Gaussian noise x_T 에서 점진적으로 noise를 제거하여 이미지 x_0 를 복원하는 Markov process

- Forward / noising process



- Reverse / denoising process

Introduction

- Latent Diffusion Model

- 기존 diffusion model 문제점

- Likelihood-based 모델 특성상 mode covering 동작을 수행

- ※ 인식할 수 없는 세부 정보까지 학습하려는 경향

- 모델을 학습하거나 평가하는데 고차원 RGB 공간에서의 gradient 연산 요구

- ※ 150~1000개의 V100 GPU days

- Latent 공간으로의 전환

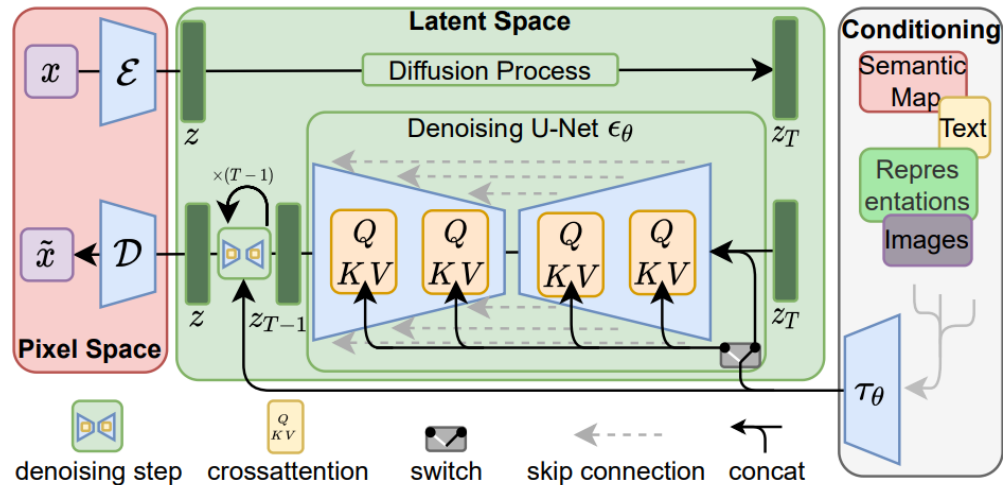
- Auto-encoder 를 학습하여 보다 낮은 차원의 표현 공간을 제공하면서도 데이터와 인지적으로 동등한 공간을 구성

- ※ 한번만 학습하면 다른 학습에 재사용 가능

- 과도한 공간 압축에 의존할 필요 없이 latent 공간에서 학습

Introduction

- Latent Diffusion Model
 - General-Purpose Conditioning with Cross-Attention
 - Text-to-image
 - Layout-to-image
 - Class-conditional image synthesis



- Retinexformer: One-stage Retinex-based Transformer for Low-light Image Enhancement (ICCV 2023)

Retinexformer¹⁾

- Introduction

- Low-light Image enhancement

- Improve the visibility and low contrast of Low-light image

- 기존 방법의 한계

- ⚡ Histogram equalization, Gamma correction

- ✓조명 요소를 거의 고려하지 않아 원치 않는 artifacts 발생

- ⚡ Retinex-based traditional method

- ✓Noise- and color distortion- free 가정하는데 현실의 저조도 환경과 맞지 않음

- ⚡ CNN-based

- ✓저조도 이미지를 정상 조명 이미지로 단순 변환하는 mapping 함수를 학습

- 인간의 색 인식을 무시하며 이론적인 검증 부족

- ✓Retinex-based

- Multi-stage training pipeline: 각각의 CNN을 개별적으로 학습한 후, 전체적으로 연결하여 end-to-end finetuning 하기 때문에 많은 시간 소요

- Capturing Long-range dependencies, Non-local self-similarity

Retinexformer¹⁾

• Method

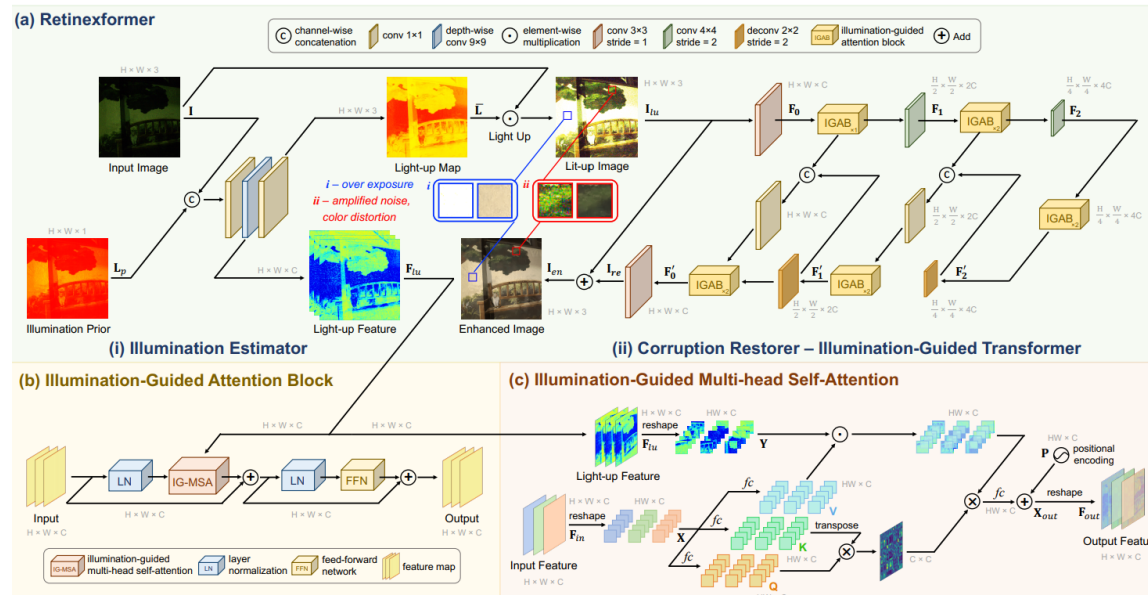
▪ Overview

- ORF (One-stage Retinex-based Framework)

조명 정보를 추정하여 저조도 이미지를 밝히고, corruption 을 복원하여 향상된 이미지 생성

- IGT (Illumination-Guided Multi-head Self-Attention)

조명 표현을 활용하여 서로 다른 노출 수준을 가진 영역들 간의 상호작용 유도



Retinexformer¹⁾

- Method

- One-stage Retinex-based Framework

- Perturbation term

$$\begin{aligned} \mathbf{I} &= (\mathbf{R} + \hat{\mathbf{R}}) \odot (\mathbf{L} + \hat{\mathbf{L}}) \\ &= \mathbf{R} \odot \mathbf{L} + \mathbf{R} \odot \hat{\mathbf{L}} + \hat{\mathbf{R}} \odot (\mathbf{L} + \hat{\mathbf{L}}) \end{aligned}$$

- Light-up process

- ⌘ Light-up map, $\bar{\mathbf{L}}$

- ✓ Illumination map, $\mathbf{L}$ 을 추정하게 되면 light-up image를 나눗셈 ($\mathbf{I}/\mathbf{L}$) 연산을 통해 계산되어야 하기 때문에 over-flow 문제 유발

- ✓ $\bar{\mathbf{L}} = 1/\mathbf{L}$

$$\mathbf{I} \odot \bar{\mathbf{L}} = \mathbf{R} + \mathbf{R} \odot (\hat{\mathbf{L}} \odot \bar{\mathbf{L}}) + (\hat{\mathbf{R}} \odot (\mathbf{L} + \hat{\mathbf{L}})) \odot \bar{\mathbf{L}}$$

- ⌘ Noise and artifacts hidden in the dark scenes

- ⌘ Under-/Over-exposure and color distortion caused by the light-up process

Retinexformer¹⁾

• Method

▪ One-stage Retinex-based Framework

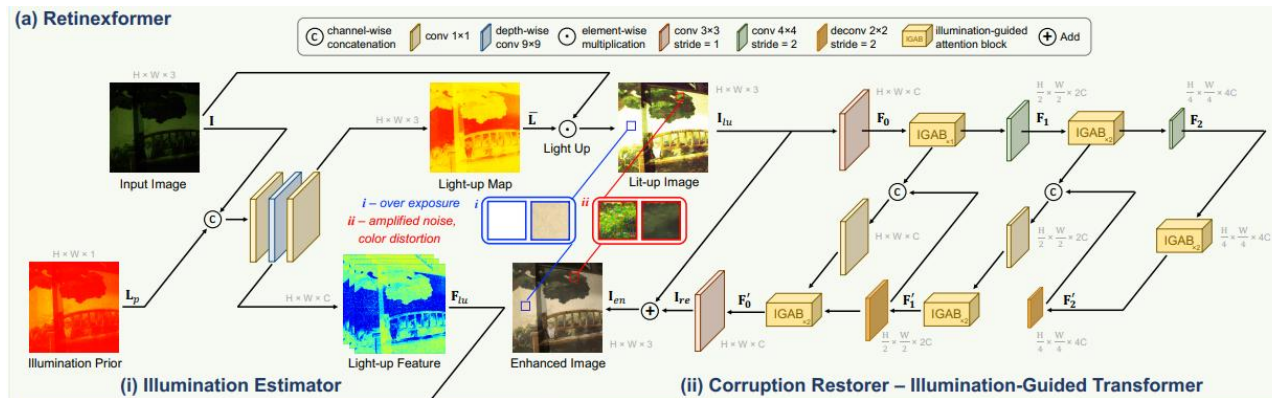
$$(\mathbf{I}_{lu}, \mathbf{F}_{lu}) = \mathcal{E}(\mathbf{I}, \mathbf{L}_p), \quad \mathbf{I}_{en} = \mathcal{R}(\mathbf{I}_{lu}, \mathbf{F}_{lu})$$

– Illumination Estimator

⚙️ Low light image 와 Illumination prior map, L_p 을 입력으로 받아 light-up image 와 light-up feature 출력

✓ Illumination prior map, L_p

• 각 pixel의 channel 방향 평균값



Retinexformer¹⁾

• Method

▪ One-stage Retinex-based Framework

$$(\mathbf{I}_{lu}, \mathbf{F}_{lu}) = \mathcal{E}(\mathbf{I}, \mathbf{L}_p), \quad \mathbf{I}_{en} = \mathcal{R}(\mathbf{I}_{lu}, \mathbf{F}_{lu})$$

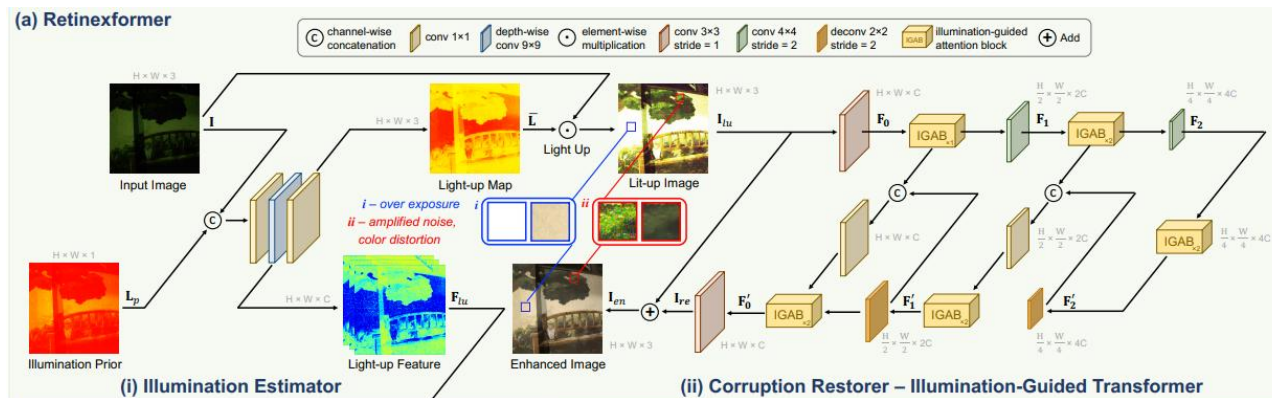
– Corruption Restorer (IGT)

기존의 방법은 반사율 이미지에서 발생하는 noise 억제에만 초점

3 scale U-shaped architecture

✓ Light-up image, Light-up feature 입력으로 활용

✓ Light-up 과정에서 발생하는 손상을 모델링한 Residual image, I_{re} 출력으로 생성



Retinexformer¹⁾

• Method

▪ One-stage Retinex-based Framework

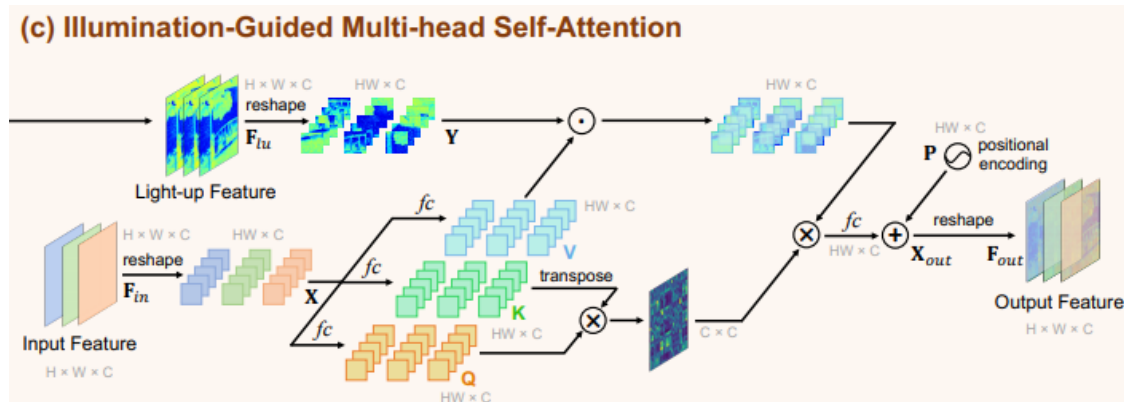
– IG-MSA

⚙️ Treat single-channel feature map as token

⚙️ Provide semantic contextual information

$$\mathbf{X} = [\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_k] \quad \mathbf{Y} = [\mathbf{Y}_1, \mathbf{Y}_2, \dots, \mathbf{Y}_k]$$

$$\text{Attention}(\mathbf{Q}_i, \mathbf{K}_i, \mathbf{V}_i, \mathbf{Y}_i) = (\mathbf{Y}_i \odot \mathbf{V}_i) \text{softmax}\left(\frac{\mathbf{K}_i^T \mathbf{Q}_i}{\alpha_i}\right)$$



Retinexformer¹⁾

- Experiment

- Dataset

- LOL v1 and v2, SID, SMID, SDSD, FiveK
 - Patches at the size of 128×128 are randomly cropped from the low-/normal light image pairs as training samples

- Training object

- minimize MAE between light-up image and enhanced image

Retinexformer¹⁾

• Experiment

▪ Quantitative results

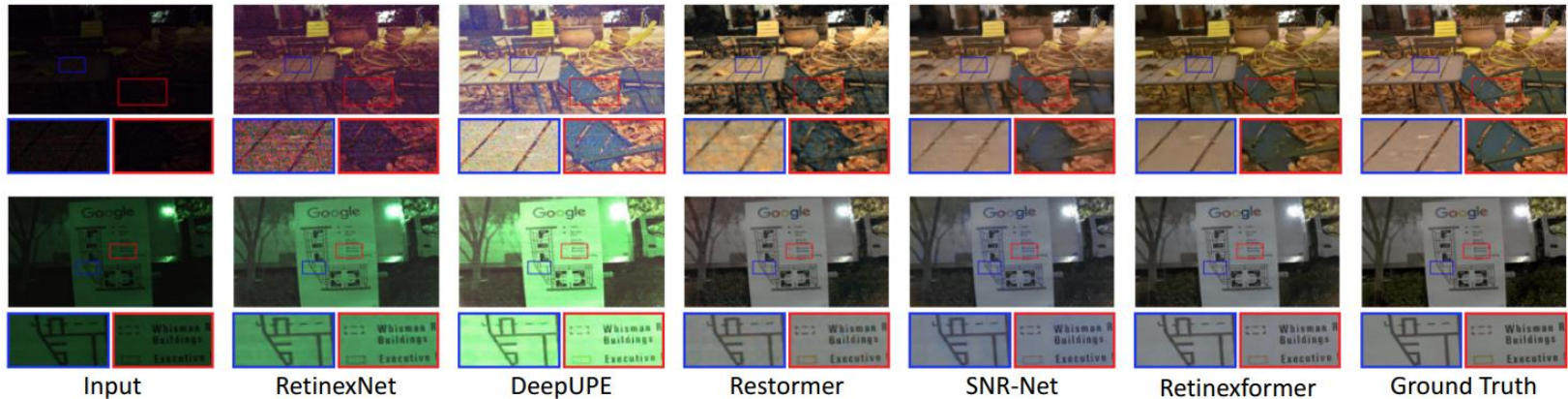
- 성능 개선 뿐 아니라 FLOPS, cost 모두 감소

Methods	Complexity		LOL-v1		LOL-v2-real		LOL-v2-syn		SID		SMID		SDSD-in		SDSD-out	
	FLOPS (G)	Params (M)	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM
SID [9]	13.73	7.76	14.35	0.436	13.24	0.442	15.04	0.610	16.97	0.591	24.78	0.718	23.29	0.703	24.90	0.693
3DLUT [63]	0.075	0.59	14.35	0.445	17.59	0.721	18.04	0.800	20.11	0.592	23.86	0.678	21.66	0.655	21.89	0.649
DeepUPE [49]	21.10	1.02	14.38	0.446	13.27	0.452	15.08	0.623	17.01	0.604	23.91	0.690	21.70	0.662	21.94	0.698
RF [26]	46.23	21.54	15.23	0.452	14.05	0.458	15.97	0.632	16.44	0.596	23.11	0.681	20.97	0.655	21.21	0.689
DeepLPP [38]	5.86	1.77	15.28	0.473	14.10	0.480	16.02	0.587	18.07	0.600	24.36	0.688	22.21	0.664	22.76	0.658
IPT [11]	6887	115.31	16.27	0.504	19.80	0.813	18.30	0.811	20.53	0.561	27.03	0.783	26.11	0.831	27.55	0.850
UFormer [52]	12.00	5.29	16.36	0.771	18.82	0.771	19.66	0.871	18.54	0.577	27.20	0.792	23.17	0.859	23.85	0.748
RetinexNet [54]	587.47	0.84	16.77	0.560	15.47	0.567	17.13	0.798	16.48	0.578	22.83	0.684	20.84	0.617	20.96	0.629
Sparse [59]	53.26	2.33	17.20	0.640	20.06	0.816	22.05	0.905	18.68	0.606	25.48	0.766	23.25	0.863	25.28	0.804
EnGAN [22]	61.01	114.35	17.48	0.650	18.23	0.617	16.57	0.734	17.23	0.543	22.62	0.674	20.02	0.604	20.10	0.616
RUAS [30]	0.83	0.003	18.23	0.720	18.37	0.723	16.55	0.652	18.44	0.581	25.88	0.744	23.17	0.696	23.84	0.743
FIDE [56]	28.51	8.62	18.27	0.665	16.85	0.678	15.20	0.612	18.34	0.578	24.42	0.692	22.41	0.659	22.20	0.629
DRBN [58]	48.61	5.27	20.13	0.830	20.29	0.831	23.22	0.927	19.02	0.577	26.60	0.781	24.08	0.868	25.77	0.841
KinD [66]	34.99	8.02	20.86	0.790	14.74	0.641	13.29	0.578	18.02	0.583	22.18	0.634	21.95	0.672	21.97	0.654
Restormer [60]	144.25	26.13	22.43	0.823	19.94	0.827	21.41	0.830	22.27	0.649	26.97	0.758	25.67	0.827	24.79	0.802
MIRNet [61]	785	31.76	24.14	0.830	20.02	0.820	21.94	0.876	20.84	0.605	25.66	0.762	24.38	0.864	27.13	0.837
SNR-Net [57]	26.35	4.01	24.61	0.842	21.48	0.849	24.14	0.928	22.87	0.625	28.49	0.805	29.44	0.894	28.66	0.866
Retinexformer	15.57	1.61	25.16	0.845	22.80	0.840	25.67	0.930	24.44	0.680	29.15	0.815	29.77	0.896	29.84	0.877

Retinexformer¹⁾

• Experiment

▪ Qualitative results



▪ User study

- under-/over- exposure, color distortion, noise/artifacts 여부 판단

Methods	L-v1	L-v2-R	L-v2-S	SID	SMID	SD-in	SD-out	Mean
EnGAN [22]	2.43	1.39	2.13	1.04	2.78	1.83	1.87	1.92
RetinexNet [54]	2.17	1.91	1.13	1.09	2.35	3.96	3.74	2.34
DRBN [58]	2.70	2.26	3.65	1.96	2.22	2.78	2.91	2.64
FIDE [56]	2.87	2.52	3.48	2.22	2.57	3.04	2.96	2.81
KinD [66]	2.65	2.48	3.17	1.87	3.04	3.43	3.39	2.86
MIRNet [61]	2.96	3.57	3.61	2.35	2.09	2.91	3.09	2.94
Restormer [60]	3.04	3.48	3.39	2.43	3.17	2.48	2.70	2.96
RUAS [30]	3.83	3.22	2.74	2.26	3.48	3.39	3.04	3.14
SNR-Net [57]	3.13	3.83	3.57	3.04	3.30	2.74	3.17	3.25
Retinexformer	3.61	4.17	3.78	3.39	3.87	3.65	3.91	3.77

- LightenDiffusion: Unsupervised Low-Light Image Enhancement with Latent-Retinex Diffusion Models (ECCV 2024)

LightenDiffusion¹⁾

- Introduction

- generative model-based

- Diffusion model의 경우 GANs, VAEs와 달리 mode-collapse 문제가 발생하지 않음

- ※ 대규모 paired data와 조건부 mechanism을 활용한 supervised learning

- ✓ 실제 세계에서 paired distorted/sharp 이미지를 수집하기 어려움

- ※ Pre-trained diffusion model의 prior를 활용하여 zero-shot

- ✓ 사전에 알려진 degradation에 대해서만 유효하며 실제 환경에서는 성능 저하



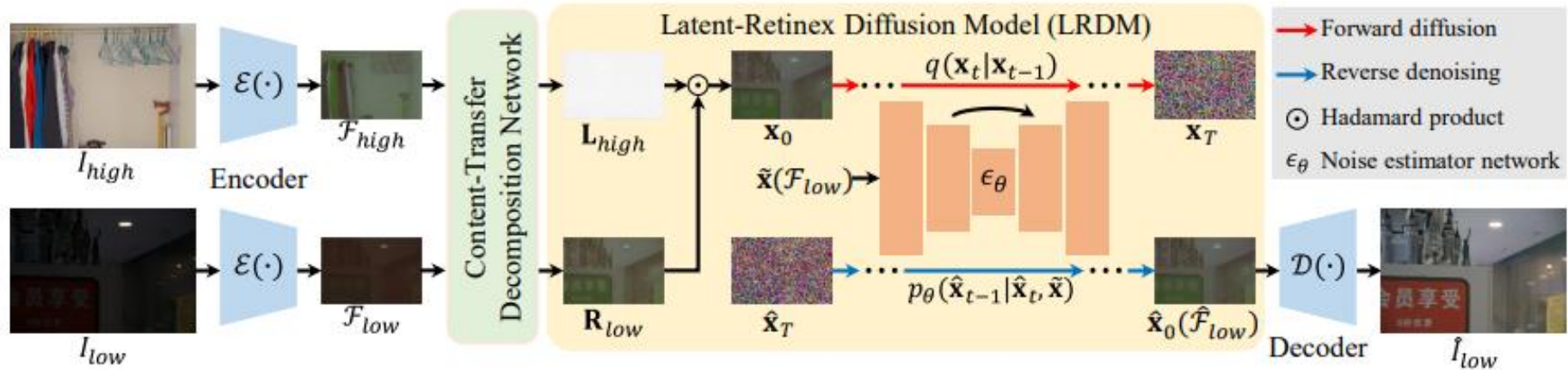
(e) GDP

(f) Ours

LightenDiffusion¹⁾

• Methodology

▪ Overview



- k개의 cascade residual block을 통해서 latent space로 변환

$$\mathcal{F}_{low} \in \mathbb{R}^{\frac{H}{2^k} \times \frac{W}{2^k} \times C} \text{ and } \mathcal{F}_{high} \in \mathbb{R}^{\frac{H}{2^k} \times \frac{W}{2^k} \times C}$$

- LL image의 reflectance map과 NL image의 illumination map을 diffusion model 입력으로 활용

※ LL image feature 를 guidance 로 활용하여 restored feature 를 생성

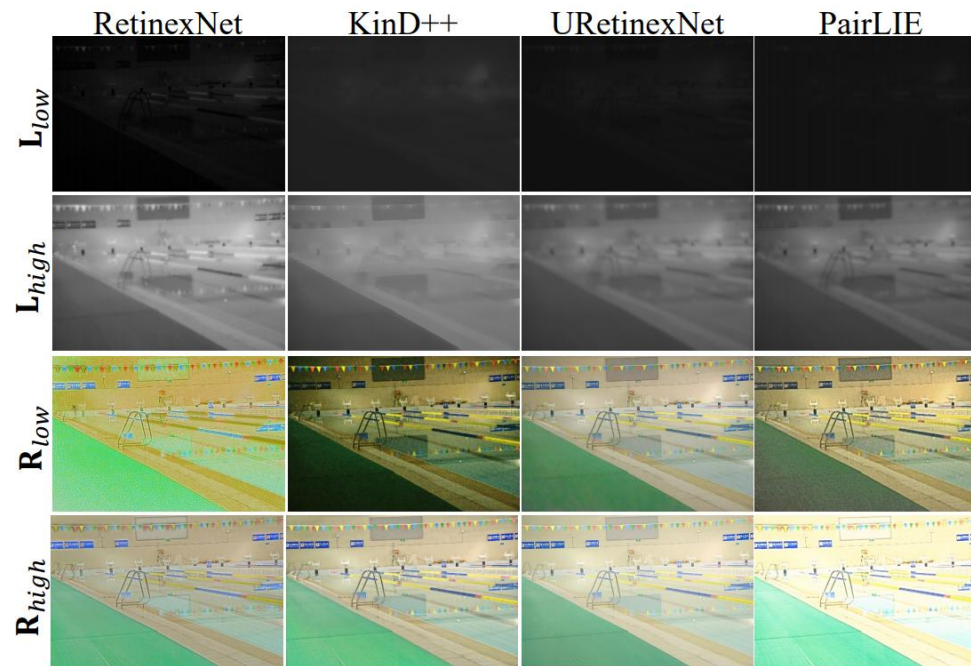
- Restored feature 를 입력으로 활용하는 Decoder 통해서 최종 enhanced image 생성

LightenDiffusion¹⁾

- Methodology
 - Content-Transfer Decomposition Network

$$I = \mathbf{R} \odot \mathbf{L}$$

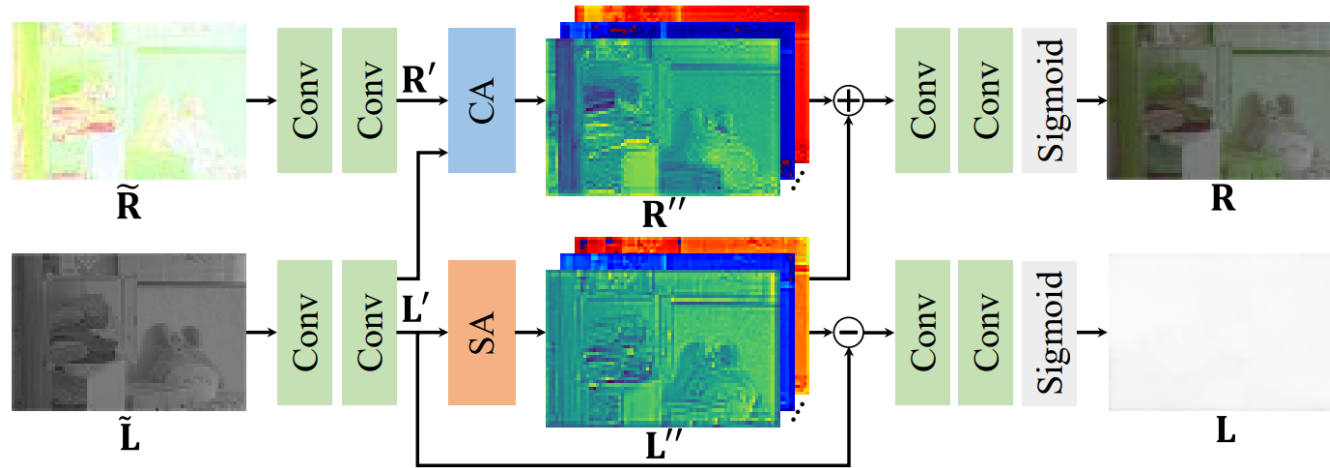
- R: inherent content information
- L: contrast and brightness information



(a) Image space

LightenDiffusion¹⁾

- Methodology
 - Content-Transfer Decomposition Network



- Initial reflectance and illumination map

$$\tilde{\mathbf{L}}(x) = \max_{c \in [0, C)} \mathcal{F}^c(x), \tilde{\mathbf{R}}(x) = \mathcal{F}(x) / (\tilde{\mathbf{L}}(x) + \tau)$$

LightenDiffusion¹⁾

- Methodology

- Content-Transfer Decomposition Network

- CA module

- ☞ To reinforce the content information in the reflectance map

- ☞ $R'' = CA(R', L')$

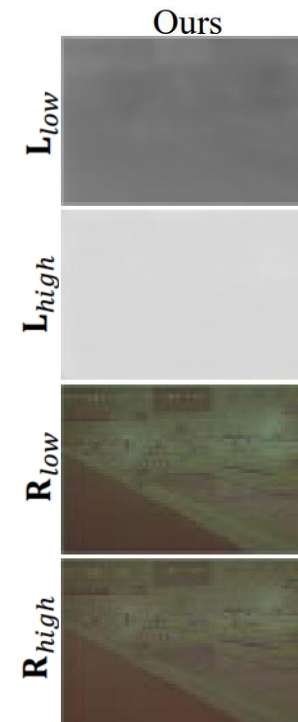
- SA module

- ☞ Content information in the illumination map

- ☞ $L'' = SA(L')$

- Final output

- ☞ $R = \text{Convs}(R'' + L''), L = \text{Convs}(L' - L'')$



(b) Latent space

LightenDiffusion¹⁾

- Methodology

- Latent-Retinex Diffusion Models

- Forward diffusion

$$\mathbf{x}_0 = \mathbf{R}_{low} \odot \mathbf{L}_{high}$$

⌘ T 단계에 걸쳐 Gaussian noise로 변환

$$q(\mathbf{x}_t \mid \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{1 - \beta_t} \mathbf{x}_{t-1}, \beta_t \mathbf{I})$$

⌘ Parameter renormalization

$$\mathbf{x}_t = \sqrt{\bar{\alpha}_t} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t} \boldsymbol{\epsilon}_t$$

- Reverse denoising

⌘ Guidance of the encoded feature of low light image

$$p_{\theta}(\hat{\mathbf{x}}_{t-1} \mid \hat{\mathbf{x}}_t, \tilde{\mathbf{x}}) = \mathcal{N}(\hat{\mathbf{x}}_{t-1}; \boldsymbol{\mu}_{\theta}(\hat{\mathbf{x}}_t, \tilde{\mathbf{x}}, t), \sigma_t^2 \mathbf{I})$$

LightenDiffusion¹⁾

- Methodology

- Latent-Retinex Diffusion Models

- Train

- ⌘ Optimize parameters θ of the network ϵ

$$\mathcal{L}_{diff} = \|\epsilon_t - \epsilon_{\theta}(\mathbf{x}_t, \tilde{\mathbf{x}}, t)\|_2$$

- Self-Constrained Consistency Loss

- ⌘ restored feature to share the same intrinsic information as the input low-light image

$$\tilde{\mathcal{F}}_{low} = \mathbf{R}_{low} \odot \mathbf{L}_{low}^{\gamma}$$

- ⌘ 복원된 특징이 유사하도록 강요함

$$\mathcal{L}_{scc} = \|\tilde{\mathcal{F}}_{low} - \hat{\mathcal{F}}_{low}\|_1$$

LightenDiffusion¹⁾

- Methodology

- Network training

- First stage

⚙️ Encoder / Decoder

✓ Content Loss

$$\mathcal{L}_{con} = \sum_{i=1}^2 \|I_{low}^i - \mathcal{D}(\mathcal{E}(I_{low}^i))\|_2$$

⚙️ CTDN

✓ Decomposition Loss

$$\mathcal{L}_{dec} = \mathcal{L}_{rec} + \lambda_2 \mathcal{L}_{ref} + \lambda_3 \mathcal{L}_{ill}$$

$$\mathcal{L}_{rec} = \sum_{i=1}^2 \sum_{j=1}^2 \|\mathcal{F}_{low}^j - \mathbf{R}_{low}^i \odot \mathbf{L}_{low}^j\|_1 \quad \mathcal{L}_{ref} = \|\mathbf{R}_{low}^1 - \mathbf{R}_{low}^2\|_1, \quad \mathcal{L}_{ill} = \sum_{i=1}^2 \|\nabla \mathbf{L}_{low}^i \cdot \exp(-\lambda_g \nabla \mathbf{R}_{low}^i)\|_2$$

- Second stage

⚙️ Unpaired low/normal light image 를 이용하여 diffusion model 최적화

LightenDiffusion¹⁾

- Experiments

- Dataset

- Paired dataset

- ☼ LOL, LSRW

- ✓ Adopt two distortion metrics PSNR, SSIM and a full-reference perceptual metric LPIPS

- Unpaired dataset

- ☼ DICM, NPE, VV

- ✓ Non-reference perceptual metrics NIQE, PI

LightenDiffusion¹⁾

- Experiments

- Quantitative comparison

– Supervised methods are trained on the LOL training set

Type	Method	LOL [58]			LSRW [16]			DICM [28]		NPE [53]		VV [51]	
		PSNR ↑	SSIM ↑	LPIPS ↓	PSNR ↑	SSIM ↑	LPIPS ↓	NIQE ↓	PI ↓	NIQE ↓	PI ↓	NIQE ↓	PI ↓
T	LIME [14]	17.546	0.531	0.290	17.342	0.520	0.416	4.476	4.216	4.170	3.789	3.713	3.335
	SDDLLE [17]	13.342	0.634	0.261	14.708	0.486	0.382	4.581	3.828	4.179	3.315	4.274	3.382
	CDEF [30]	16.335	0.585	0.351	16.758	0.465	0.314	4.142	4.242	3.862	2.910	5.051	3.272
	BrainRetinex [4]	11.063	0.475	0.327	12.506	0.390	0.374	4.350	3.555	3.707	3.044	4.031	3.114
SL	RetinexNet [58]	16.774	0.462	0.390	15.609	0.414	0.393	4.487	3.242	4.732	3.219	5.881	3.727
	KinD++ [74]	17.752	0.758	0.198	16.085	0.394	0.366	4.027	3.399	4.005	3.144	3.586	2.773
	LCDPNet [52]	14.506	0.575	0.312	15.689	0.474	0.344	4.110	3.250	4.106	3.127	5.039	3.347
	URetinexNet [59]	19.842	0.824	0.128	18.271	0.518	0.295	4.774	3.565	4.028	3.153	3.851	2.891
	SMG [64]	23.814	0.809	0.144	17.579	0.538	0.456	6.224	4.228	5.300	3.627	5.752	3.757
	PyDiff [76]	23.275	0.859	0.108	17.264	0.510	0.335	4.499	3.792	4.082	3.268	4.360	3.678
	GSAD [20]	22.021	0.848	0.137	17.414	0.507	0.294	4.496	3.593	4.489	3.361	5.252	3.657
SSL	DRBN [68]	16.677	0.730	0.252	16.734	0.507	0.376	4.369	3.800	3.921	3.267	3.671	3.117
	BL [39]	10.305	0.401	0.382	12.444	0.333	0.384	5.046	4.055	4.885	3.870	5.740	4.030
UL	Zero-DCE [12]	14.861	0.562	0.330	15.867	0.443	0.315	3.951	3.149	3.826	2.918	5.080	3.307
	EnlightenGAN [24]	17.606	0.653	0.319	17.106	0.463	0.322	3.832	3.256	3.775	2.953	3.689	2.749
	RUAS [32]	16.405	0.503	0.257	14.271	0.461	0.455	7.306	5.700	7.198	5.651	4.987	4.329
	SCI [40]	14.784	0.525	0.333	15.242	0.419	0.321	4.519	3.700	4.124	3.534	5.312	3.648
	GDP [8]	15.896	0.542	0.337	12.887	0.362	0.386	4.358	3.552	4.032	3.097	4.683	3.431
	PairLIE [10]	19.514	0.731	0.254	17.602	0.501	0.323	4.282	3.469	4.661	3.543	3.373	2.734
	NeRCo [67]	19.738	0.740	0.239	17.844	0.535	0.371	4.107	3.345	3.902	3.037	3.765	3.094
	Ours	20.453	0.803	0.192	18.555	0.539	0.311	3.724	3.144	3.618	2.879	2.941	2.558

LightenDiffusion¹⁾

- Experiments
 - Qualitative comparison



LightenDiffusion¹⁾

- Experiments

- Qualitative comparison



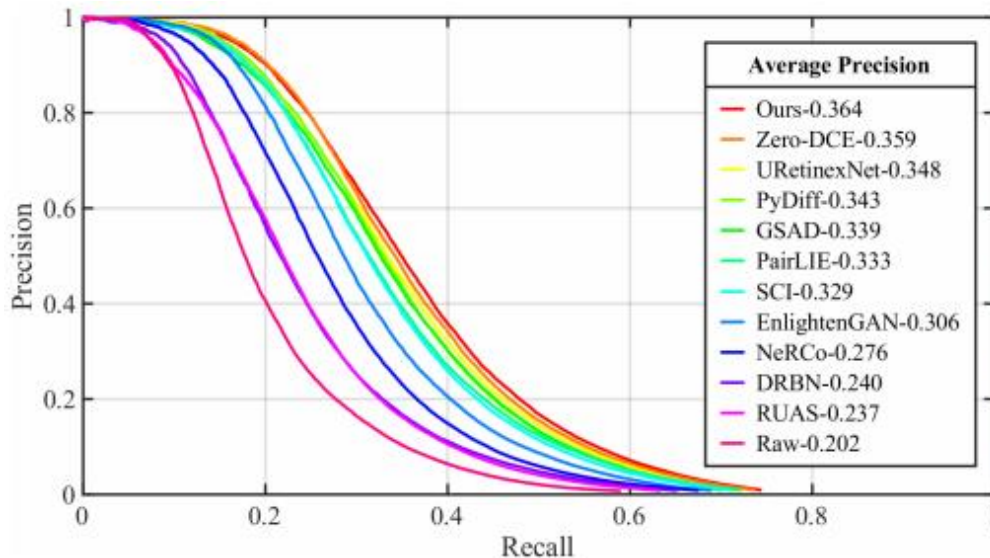
LightenDiffusion¹⁾

• Experiments

▪ Low-Light Face Detection

- DARK FACE dataset
- Using RetinaFace for detection
- Evaluation under the IoU threshold of 0.3 to depict the precision-recall (P-R) curves and calculate the average precision

☼ Improves from 20.2% to 36.4% compared to the RAW images



LightenDiffusion¹⁾

- Experiments

- Ablation study

– Image space vs latent space



(a) $k = 0$ (Image space) (b) $k = 1$ (Latent space) (c) $k = 2$ (Latent space) (d) $k = 4$ (Latent space)

⚡ Performance and inference time difference according to k

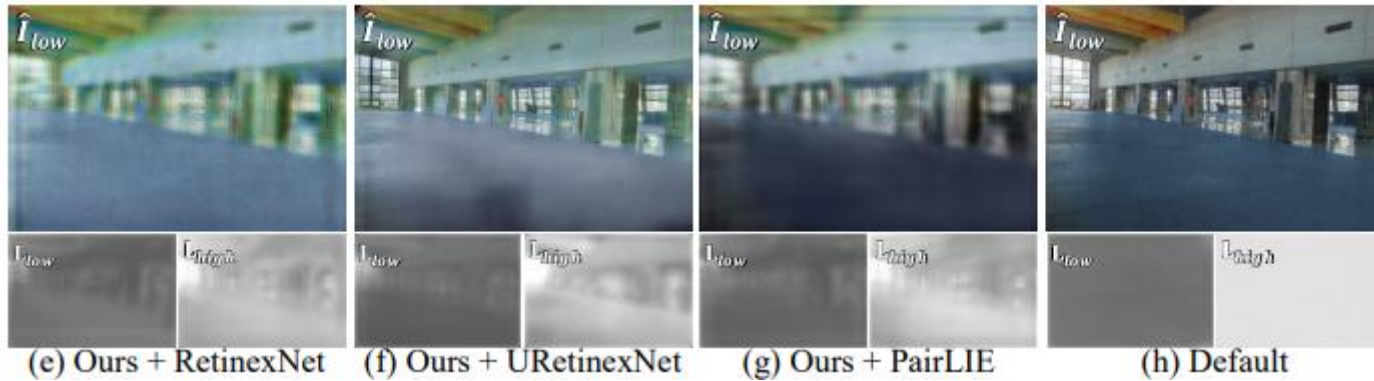
Method	LOL [58]			DICM [28]		Time (s) ↓
	PSNR ↑	SSIM ↑	LPIPS ↓	NIQE ↓	PI ↓	
1) $k = 0$ (Image Space)	17.054	0.715	0.372	4.519	4.377	4.733
2) $k = 1$ (Latent Space)	19.228	0.728	0.355	4.101	3.457	0.872
3) $k = 2$ (Latent Space)	20.097	0.798	0.210	4.021	3.402	0.411
4) $k = 4$ (Latent Space)	20.104	0.785	0.195	3.906	3.332	0.256

LightenDiffusion¹⁾

- Experiments

- Ablation study

–CTDN network



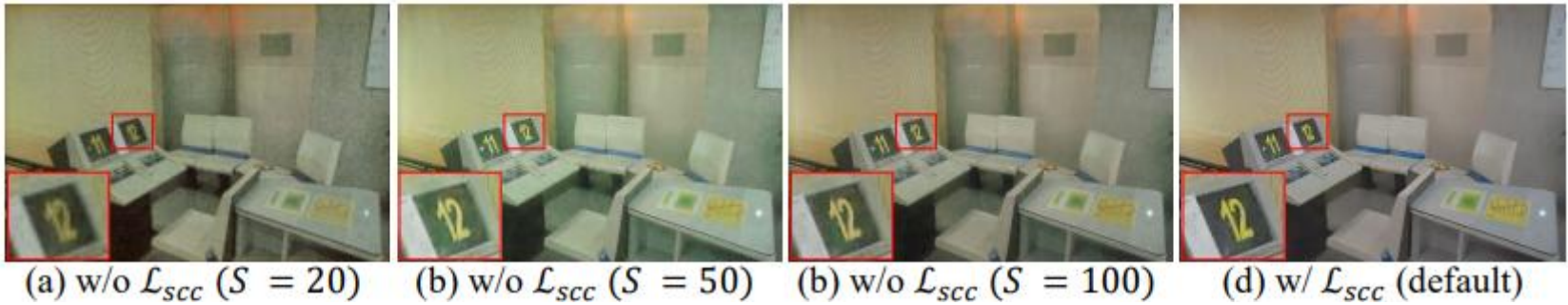
Method	LOL [58]			DICM [28]		Time (s) ↓
	PSNR ↑	SSIM ↑	LPIPS ↓	NIQE ↓	PI ↓	
5) RetinexNet [58]	16.616	0.563	0.579	5.859	6.056	0.296
6) URetinexNet [59]	17.916	0.703	0.391	4.371	4.561	0.293
7) PairLIE [10]	17.089	0.605	0.568	6.017	6.349	0.295

LightenDiffusion¹⁾

- Experiments

- Ablation study

- Self-constrained consistency loss



Method	LOL [58]			DICM [28]		Time (s) ↓
	PSNR ↑	SSIM ↑	LPIPS ↓	NIQE ↓	PI ↓	
8) w/o $\mathcal{L}_{scc}$ ($S = 20$)	19.184	0.785	0.213	4.045	3.408	0.314
9) w/o $\mathcal{L}_{scc}$ ($S = 50$)	19.473	0.791	0.209	3.998	3.392	0.687
10) w/o $\mathcal{L}_{scc}$ ($S = 100$)	20.255	0.801	0.209	3.831	3.228	1.208
11) Default	<u>20.453</u>	<u>0.803</u>	<u>0.192</u>	<u>3.724</u>	<u>3.144</u>	<u>0.314</u>

감사합니다