

일반논문 (Regular Paper)

방송공학회논문지 제24권 제4호, 2019년 7월 (JBE Vol. 24, No. 4, July 2019)

<https://doi.org/10.5909/JBE.2019.24.4.623>

ISSN 2287-9137 (Online) ISSN 1226-7953 (Print)

합성곱 신경망 기반 저조도영상의 반사 영상 생성

이 승 수^{a)}, 최 창 열^{a)}, 김 만 배^{a)†}

Generating a Reflectance Image from a Low-Light Image Using Convolutional Neural Network

Seungsoo Lee^{a)}, Changyeol Choi^{a)}, and Manbae Kim^{a)†}

요 약

저조도 영상의 개선을 위해서 밝기 및 대조 개선, 조명 성분 감쇄 등의 다양한 연구가 진행됐다. 기존의 hand-crafted 방법에서 인공신경망으로 기존 기법들을 대체하는 연구가 최근에 진행 중이다. 본 논문에서는 조명 광원이 존재하는 저조도 영상으로부터 조명 성분을 감쇄하고, 반사 성분만을 생성하는 기법을 합성곱 신경망으로 대체하는 방법을 제안한다. 실험에서는 102장의 저조도 영상으로 학습시킨 합성곱 신경망으로 만족스러운 반사 영상을 생성하였다.

Abstract

Many researches have been carried out for brightness and contrast enhancement, illumination reduction and so forth. Recently, the aforementioned hand-crafted approaches have been replaced by artificial neural networks. This paper proposes a convolutional neural network that can replace the method of generating a reflectance image where illumination component is attenuated. Experiments are carried out on 102 low-light images and we validate the feasibility of the replacement by producing satisfactory reflectance images.

Keywords : Low-light image, Retinex, Convolutional neural network, Reflectance, Illumination

a) 강원대학교 컴퓨터정보통신공학과(Department of Computer and Communications Eng., Kangwon National University)

† Corresponding Author : 김만배 (Manbae Kim)

E-mail: manbae@kangwon.ac.kr

Tel: +82-33-250-6395

ORCID: <https://orcid.org/0000-0002-4702-8276>

※ This research was supported by the MSIT (Ministry of Science and ICT), Korea, under the ITRC (Information Technology Research Center) support program (IITP-2019-2018-0-01433) supervised by the IITP (Institute for Information & communications Technology Promotion).

· Manuscript received January 31, 2019; Revised April 15, 2019; Accepted June 25, 2019.

1. 서론

밤에는 햇빛 등의 부족으로 영상의 밝기가 크게 낮아진다. 또한 광원 때문에 일부 영상 영역은 상대적으로 밝게 나오고, 그림자 영역은 상대적으로 대조가 낮아져서 사람의 육안으로는 내부 구조를 파악하는 것이 어렵다. 외부 환경 중 광원은 영상 화질에 큰 영향을 주게 된다. 광량이 과도하거나 부족한 경우 영상 표현 능력을 저하시키게 되어 영상에 대한 판별 능력을 어렵게 한다. 또한, 어두운 환경에서 존재하는 광원의 영향은 전체 영상의 인지에 큰 영향을 준다. 상기 영상의 품질 저하를 보완하기 위해, 히스토그램 평활화(histogram equalization)^[1-2], 감마보정(gamma correction)^[3] 등이 사용되어 왔다.

Land 등은 실제 인간의 시각에 들어오는 빛은 물체의 반사를 통해서 들어오는 반사(reflectance) 성분과 광원의 조명(illumination) 성분으로 나눌 수 있음을 실험적으로 입증하였다^[4]. Jobson 등은 Land의 시각모델과 인간의 지각모델인 Weber-Fachner 법칙을 바탕으로 Single Scale Retinex(SSR) 알고리즘을 제안하였다^[5]. 상기 방법은 영상의 조명성분을 최대한 배제하여 물체가 가지는 고유한 값을 유도할 수 있다. SSR 방식은 적절한 크기의 저주파 필터로 필터링된 결과를 영상의 조명 성분으로 추정하며, 조명 성분을 감쇄함으로써 반사 성분을 얻는다.

위에서 언급한 것처럼 어두운 환경에서 촬영한 저조도(low-light) 영상은 주변에 존재하는 조명에 의해서 일부 영역은 빛에 의해 높은 밝기를 가지는 반면에, 일부 영역은 제한된 밝기 영역으로 인해서, 어둡게 보이는 현상이 종종

발생한다. 그림 1의 네 장의 영상을 관찰하면, 주변의 조명 또는 모니터의 밝기 때문에, 일부 영역은 원래 밝기의 대조를 손실하고, 낮은 대조로 물체의 구조를 시각적으로 이해하는 것이 어렵다. 이러한 조명 광원의 영향을 줄이고, 반사 성분을 얻는 것이 Retinex 이론이다.

최근에 기존의 hand-crafted 특징을 사용하지 않고, 심층신경망(Deep Neural Network: DNN)을 이용하여 저조도 영상의 화질을 개선하는 방법들이 소개되고 있다^[6-9]. DNN 기반 방법은 다른 신경망 모델과 마찬가지로 학습 데이터를 활용한다. 입력영상과 실측(ground-truth) 영상을 학습하면서 실측 영상과 유사한 영상을 생성할 수 있도록 신경망을 모델링한다. Lore 등은 저조도 영상에 stacked denoising 자기부호화기(autoencoder)를 사용하였다^[6]. Shen 등은 합성곱 신경망(Convolutional Neural Network)을 적용하여 저조도 영상을 향상하였다^[7]. 또한 Park 등은 자기부호화기 및 CNN의 듀얼 신경망 모델을 제안하였다^[8]. Kim 등은 심층신경망을 이용하여 저조도 영상에서 반사영상을 생성할 수 있는 가능성을 보여주었다^[9]. 저조도 영상의 밝기 향상 기법들은 대부분 고조도 영상(high-light image)에 수작업으로 감마 변환(gamma transformation)을 적용하여 어둡게 만든 후에, 신경망으로 학습하는 과정을 거친다. 이때 실험에 사용한 감마 변환과 실제 저조도 영상이 맞지 않으면, 일반 영상에서는 성능이 저하되기도 한다. 따라서, 실제 환경에서 얻은 저조도 영상으로부터 반사 영상을 생성하는 것이 필요하다.

본 연구는 기존의 저조도 영상의 향상 알고리즘을 합성곱 신경망(Convolutional Neural Network: CNN)으로 대

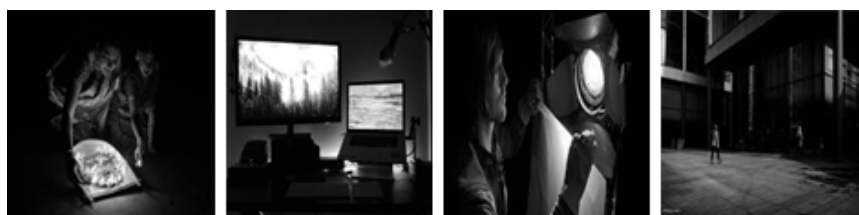


그림 1. 광원 조명의 영향을 받는 저조도 영상
Fig. 1. Low-light images affected by light source

체하는 이론을 제안하고 타당성(feasibility)을 검증한다. 이를 위해서 먼저 Retinex 이론 기반으로 반사 영상 생성 알고리즘을 소개한다. 저조도 입력 영상과 Retinex 알고리즘으로 얻은 출력 실측 영상(ground-truth image)을 CNN으로 학습을 하게 한 후에, CNN의 출력과 알고리즘으로부터 직접 생성한 영상을 비교하면서 성능을 검증한다.

그림 2는 기존 저조도 영상 향상 방법과 제안 방법의 차이점을 보여준다. 기존 생성방법은 그림 2(a)에서 보듯이, 다량의 학습 영상과 실측 타겟 영상으로 신경망을 학습시킨 후에, 예측 영상(predicted image)을 얻는 방식이다. 실험에서는 타겟 영상이 주어지면, 수작업으로 감마 변환하여

영상을 어둡게 한다. 감마 변환된 어두운 영상을 입력으로, 원영상인 밝은 영상을 출력으로 하여 신경망 모델을 학습시킨다. 이와는 달리 제안방법은 그림 2(b)에서 보는 것처럼 Retinex 기반으로 반사 영상을 생성한 후에, CNN으로 입력 영상과 출력 타겟 영상을 학습한다. 최종적으로 CNN 모델로 기제작 반사 영상 생성 알고리즘을 대체하고자 하는 것이다.

저조도 영상 향상 또는 반사 성분 추출 방법의 성능은 본 연구의 목적이 아니고, 성능에 관계없이 어떠한 반사 성분 생성 알고리즘이라도 CNN으로 대체가 가능한지에 대한 타당성을 검증하는 것이 주 목적이다. 이를 위해 기존에 소개된 반사 성분 추출 알고리즘을 활용하여 입력 영상으로부터 반사 영상을 생성하고, CNN을 학습시킨다. CNN이 예측한 출력 반사 영상과 기존 알고리즘의 출력 반사 영상의 오차를 검증한다.

본 논문의 구성은 다음과 같다. 다음 장에서는 Retinex 이론을 설명하고 반사 영상을 생성하는 과정을 소개한다. III장에서는 제안 방법의 전체 흐름도를 설명한다. IV장에서는 실험 결과를 분석한다. 마지막으로 V장에서는 결론 및 향후 연구를 정리한다.

II. Retinex 이론

Land 등은 실험적으로 인간 시각 기관이 인지하는 물체의 색이 광원과 물체의 반사성분의 곱으로 나타낼 수 있음을 입증하였으며^[4], 수식으로 표현하면 밝기(brightness) 영상 $f(x,y)$ 는 두 개의 성분인 광원 조명(illumination) 성분 $i(x,y)$ 와 반사(reflectance) 성분 $r(x,y)$ 의 곱으로 다음과 같다.

$$f(x,y) = i(x,y) \cdot r(x,y) \quad (1)$$

여기서 (x,y) 는 픽셀의 좌표이다.

상기 식은 조명 성분을 추정할 수 있다면 산술적인 방

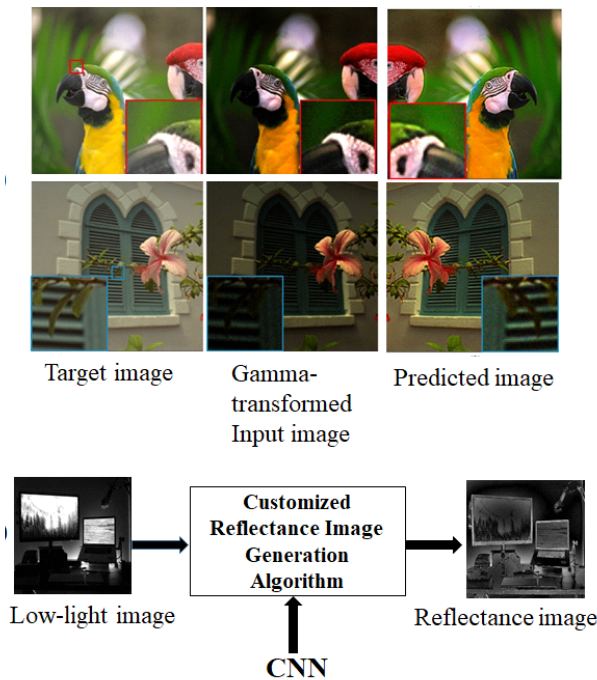


그림 2. 기존 저조도 영상 향상과 제안 방법의 차이. (a) 일반적인 신경망 기반 영상 생성^[8] 및 (b) 기제작된 반사 영상 생성 알고리즘을 CNN으로 대체하는 구조도

Fig. 2. Difference of conventional reflectance image generation and the proposed method. (a) The approach of the former^[8] and (b) diagram replacing a customized reflectance map generation algorithm with CNN

법을 통하여 물체 고유의 색인 반사성분을 알아낼 수 있음을 의미한다. 조명 성분은 천천히 변화하는 저주파 성분을 가지고 있고, 반사 성분은 고주파 성분을 가지고 있다. 따라서 적절한 필터링을 활용하면, 조명의 영향을 줄일 수 있다.

Weber-Fechner^[5] 법칙에 따르면 인간의 감각 기관이 인지하는 감각의 차이들이 로그(log)스케일을 갖고 있는데, i 와 r 을 분리하기 위해서 식 (1)에 로그 변환을 적용하면 다음과 같이 얻어진다.

$$z(x,y) = \log[f(x,y)] = \log[i(x,y)] + \log[r(x,y)] \quad (2)$$

로그 영상 $z(x,y)$ 에 푸리에 변환 F 를 적용하면 다음과 같다.

$$Z(u,v) = I(u,v) + R(u,v) \quad (3)$$

여기서 $I(u,v) = F\log[i(x,y)]$, $R(u,v) = F\log[r(x,y)]$, 및 $Z(u,v) = F[z(x,y)]$ 이다.

식 (3)에 고역통과 필터 $H(u,v)$ 을 곱해준다. I 는 저주파 성분이므로 $H \cdot I \approx 0$ 이다.

$$H \cdot Z = H \cdot I + H \cdot R \approx H \cdot R \quad (4)$$

고역통과 필터는 저주파를 감쇄하고 고주파 성분을 통과시킨다. 본 실험에서는 식(5)의 감쇄 계수 0.5인 버터워스(Butterworth) 필터를 이용하였다.

$$H = 1 - \left[1 + \left(\frac{D(u,v)}{D_0}\right)^n\right]^{-1} \quad (5)$$

여기서 $D(u,v) = \sqrt{u^2 + v^2}$, D_0 는 차단주파수이다.

$H \cdot R$ 에 역푸리에 변환 F^{-1} 을 사용하여 공간 좌표의 반사 영상, $r_R(x,y)$ 을 복원한다.

$$r_R(x,y) = F^{-1}[H(u,v) \cdot R(u,v)] \quad (6)$$

그림 3은 상기 Retintex 이론에 기반하여 생성된 반사 영상 r_R 을 보여준다. 그림 3(a)은 입력 영상 f 을 보여주는데 조명성분이 뚜렷이 관측되고 있다. 이 조명 성분을 감쇄시킨 반사 영상 r_R 은 그림 3(b)에서 보여진다. 조명 성분이 감쇄된 반사 성분을 관측할 수 있다. 신경망 학습 과정에서 그림 3(a)는 신경망 모델의 입력데이터이고, 그림 3(b)는 실측 출력 영상으로 사용된다. 신경망 모델을 학습한 후에, 새로운 저조도 영상이 입력되면 신경망은 예측된 반사 영상을 출력한다.



그림 3. 신경망 학습에 사용하는 영상. (a) 입력 영상, f 및 (b) 반사 영상, r_R

Fig. 3. Images used at the training of a neural network. (a) input images, f and (b) reflectance images, r_R

III. 합성곱 신경망 모델

제안 방법은 그림 4에서 보여진다. 먼저 입력 그레이스케일 저조도 영상(Input low-light image)에 저주파(Low-pass filter: LPF) 필터를 적용한다. 입력 영상과 필터링된 영상(LP input image)의 차영상을 저장한다. 이 차영상은 저주파 성분 분기 감쇄되기 때문에, 고주파 데이터(High-frequency data)이다.

필터링된 영상은 심층 신경망의 입력데이터로 사용된다. 타겟 반사 영상에도 저주파 필터를 적용시켜 필터링된 영상을 만든다. 저주파 필터로는 가우시안 필터를 이용하였고, 필터의 크기는 3×3 이며 $\sigma = 3$ 이다. 저주파 필터가 적용된 영상들은 CNN에서 입력과 출력으로 사용된다. 즉, 저주파 필터를 적용시킨 영상들로 CNN을 학습시키면 예측된 반사 영상이 생성된다.

입출력이 모두 영상인 딥러닝의 예는 응용 분야로 de-noising^[10], super-resolution^[11] 등이 있다. 보통 이러한 딥러닝 모델에서는 입력에는 블러링 또는 노이즈 영상을 두고, 출력은 원영상을 사용하여 학습시키는 방식이다. 반대로 입력과 출력을 바꾸면 출력 영상의 화질이 저하되기

도 한다. 생성모델의 단점은 출력 예측 영상의 화질이 낮다는 것이다. 영상의 블러링은 CNN 등의 감독 학습 및 autoencoder 등의 무감독 학습 딥러닝에서 아직 해결되지 않은 문제이다^[12]. deconvolution^[10]도 motion deblurring, super-resolution 등의 제한된 환경에서는 만족스러운 성능을 보인다.

따라서 상기 문제를 극복하기 위하여 제안방법에서는 입력 및 실측 출력 영상에 저주파 필터를 적용하여 영상을 블러링하게 한 후에 학습함으로써 상기 문제의 해결이 가능하다. 또한 예측된 출력 영상에 고주파 성분을 더해줌으로써, 실측 영상과 유사한 영상을 만들 수 있는 장점이 있기 때문이다. 좀 더 자세히 설명하면 다음과 같다. 예측 영상은 부드럽게 생성되기 때문에 신경망 입출력 영상에 저주파 필터를 적용한다. 그러므로 고주파 성분이 감쇄된 영상을 입출력으로 사용함으로써, 오차를 줄일 수 있다. 또한 고주파 성분은 경계 등의 시각적 성분을 포함하므로, 출력에서 생성된 예측 반사 영상에 고주파 성분을 가지고 있는 차영상이 합해지게 된다. 따라서 최종 출력 영상은 영상 경계에서 입력 영상의 값을 유지하게 된다.

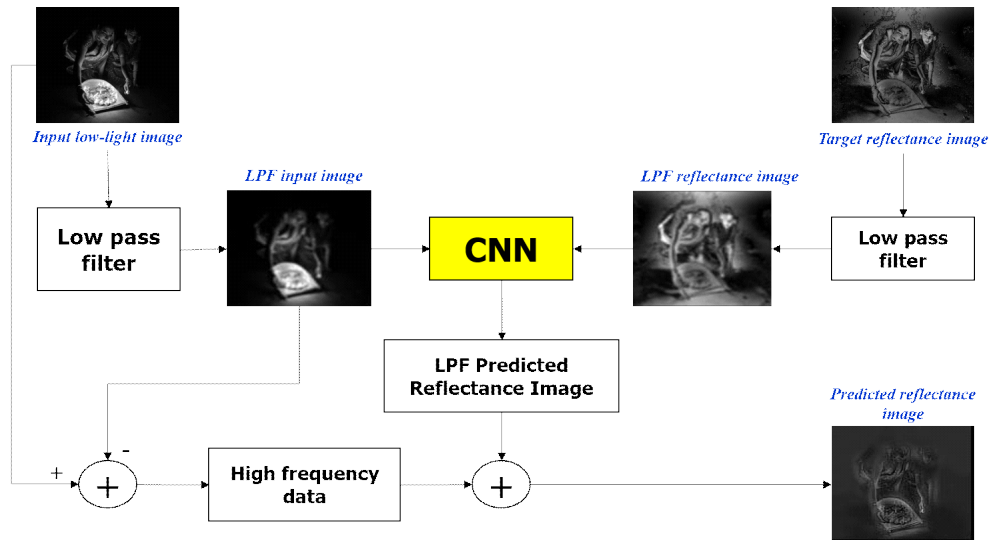


그림 4. 제안 방법의 전체 흐름도

Fig. 4. The overall block diagram of the proposed method

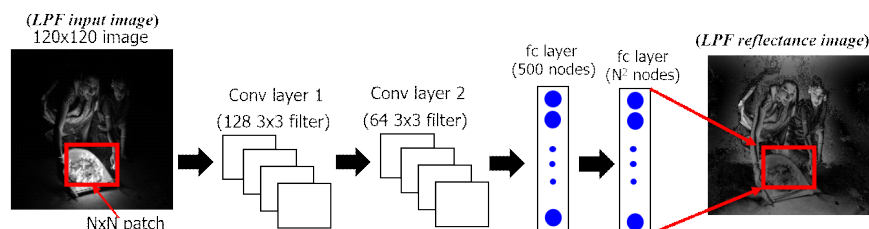


그림 5. 그림 4의 CNN의 모델 구조
Fig. 5. CNN model in Fig. 4

그림 4의 CNN 모델은 그림 5에서 보여진다. 입력층과 출력층 그리고 내부에는 2개의 합성곱층(conv layer) 및 완전결합층(fully-connected layer, fc layer)로 구성된다. 신경망은 역전파 알고리즘을 통해 심층신경망의 학습을 수행한다. 입력으로 전체영상을 사용하지 않고, 영상 패치(patch) 단위로 입력하는데, 그 이유는 전체 영상을 사용하게 되면 출력 영상의 성능이 만족스럽지 못하기 때문이다^[6,7,8,9]. 그러므로 입력층에는 $N \times N$ 영상 패치를 이용하여 학습데이터를 생성한다. 실험에서는 $N=20$ 을 사용하였다. 표 1은 CNN 네트워크의 층 구조 및 파라메타 값을 보여준다. 첫 번째 합성곱층(conv layer 1)에는 128개의 3x3 필터를 적용하고, 활성화함수(activation function)로는 relu(rectified linear unit)를 사용한다. 두 번째 합성곱층(conv layer 2)에는 64개의 3x3 필터를 적용하고, relu를 사용한다. 출력 데이터의 크기를 줄이기 위해서 최대 풀링을 실행한다. 다음

에는 500개의 노드를 가지는 완전결합층(fc-1)이 있고, 텐서를 1차원 벡터로 변환한다. 영상 픽셀값의 범위는 [0,1]이므로 이에 맞추기 위해서 시그모이드(sigmoid)를 사용한다. 최종 출력에는 N^2 노드가 있는데, 이 값은 패치의 크기와 같다. 마지막으로 $N^2 \times 1$ 벡터를 다시 $N \times N$ 의 패치로 변환한다. 최종적으로 패치들의 가중치 합으로 최종 예측영상을 복원한다.

IV. 실험 결과

실험에서 사용한 저조도 실험 영상은 인터넷에서 수집하였으며, 그림 6에서 보여진다. 영상들은 어두운 환경에서 광원의 영향을 받고 있어서, 제안 방법이 효율적으로 적용될 수 있다. 광원의 영향을 받는 곳은 밝고, 다른 영역은 상대적으로 대조가 낮아 내부의 구조를 자세히 알 수 없다. 영상 해상도는 128x128이고 신경망 학습과 테스트는 운영체제 WINDOWS 10에서 MATLAB R2018a를 사용하였다. 학습 데이터로 사용한 이미지는 총 102장이다. 82장은 학습에 사용하고, 나머지 20장은 랜덤하게 선택한 후에 테스트에 사용하여 성능을 검증한다.

영상은 먼저 [0,1]로 정규화한다. 학습 데이터는 패치 단위로 처리된다. 이미지를 학습하기 위해 한 장의 사진이 들어올 때마다 중첩된 패치를 생성한다. 패치 생성 방법은 두 가지로 구분되는데, 첫 번째는 영상에서 패치를 랜덤하게 생성한 후에, 합성은 중첩된 패치들의 평균값으로 영상을

표 1. CNN 네트워크의 층 구조 및 파라메타 값
Table 1. Layers and parameters of CNN network

Layer	Parameter
Input	$N \times N$ patch
conv layer 1	128, 3x3 filter
pooling	max pooling
activation function	relu
conv layer 2	64, 3x3 filter
activation function	relu
fc-1	500 nodes
activation function	sigmoid
Output	N^2 nodes



그림 6. 실험에 사용된 광원이 존재하는 저조도 영상
Fig. 6. Low-light images under illumination used in experiments

복원한다^[7]. 다른 방법은 스트라이드(stride)을 조절하면서 중첩된 패치를 생성하고, 복원은 첫 번째와 같은 방법을 사용한다^[8,9]. 실험에서는 두 번째 방법을 선택한다. 신경망의 학습률은 0.001로 설정하였고, 가중치 갱신 방식은 확률적 경사 하강법(stochastic gradient descent)을 사용한다. 최적화 기법은 Adam 최적화를 사용하였다.

제안한 심층신경망의 객관적 성능을 검증하기 위해서 저조도영상 향상 성능 검증에서 일반적으로 사용하는 다음 3가지의 객관적 메트릭을 계산한다; 1) Root Mean Squared Error (RMSE), 2) Peak Signal-to-Noise Ratio (PSNR), 3) Structural Similarity Index (SSIM). 아래 식에서 y_i 는 예측된 값이고, d_i 는 실측 픽셀값이다. 신경망에서는 y, d 모두

[0,1]의 정규화 값을 가지므로, 255를 곱하여 [0, 255]로 계산한다.

1) RMSE

$$SE = \sqrt{\frac{1}{K} \sum_{i=1}^K (y_i - d_i)^2} \quad (7)$$

K 는 영상 픽셀 개수이다.

2) PSNR

$$PSNR = 10 \log_{10} \frac{255^2}{RMSE^2} \text{ (dB)} \quad (8)$$

3) Structural Similarity Index (SSIM)^[13]

SSIM은 전체적인 관점에서 영상 안에서의 구조적 정보인 휘도, 명암비, 구조들을 추출하여 구조적 유사도를 구한 다음 영상의 품질을 측정할 때 사용된다. 1에 가까울수록 유사도가 높으며 반대로 0에 가까울수록 유사도가 낮다.

$$SSIM(d, y) = \frac{(2\mu_d\mu_y + C_1)(2\sigma_{dy} + C_2)}{(\mu_d^2 + \mu_y^2 + C_1)(\sigma_d^2 + \sigma_y^2 + C_2)} \quad (9)$$

μ_d 는 d 의 평균값, μ_y 는 y 의 평균값, σ_{dy} 는 x 와 y 의 공분산이다.

표 2는 그림 4에 있는 Target reflectance image와 Predicted reflectance image의 성능 결과를 보여준다. 82장의 Train 및 20장의 Test 영상의 평균 RMSE는 (Training, Test) = (18.62, 27.26)이다. 이 값은 영상을 [0,255]로 변환한 후에 얻은 값이다. 평균 PSNR에서는 (Training, Test) = (23.47, 19.44)이다. 평균 SSIM=(0.42, 0.22)을 얻었다. 표 2의 PSNR, SSIM값은 제안 방법과 비교할 수 있는 논문이 없기 때문에, 간접적으로 [8]의 결과를 참고하여 설명한다. [8]에서 PSNR의 범위는 13~16 dB이고, SSIM의 범위는 0.38~0.65이다. 이 범위는 [8] 및 [8]에서 비교한 연구들의 실험 결과이다. 서론에서 언급했듯이, 제안 방법은 감마 변

환으로 저조도 영상을 생성하지 않고, 대신 자연 저조도 영상을 사용한다. 따라서 객관적 성능은 낮을 수 있지만, 실용적인 면에서는 강인성을 가진다. 다음 주관적 평가에서 설명하듯이, 일부 예측 영상의 화질 저하로 전체 측정값의 저하가 발생한다.

표 2. RMSE, PSNR, 및 SSIM의 객관적 성능 평가

Table 2. Performance evaluation of RMSE, PSNR and SSIM

Avg. RMSE (in pixel)		Avg. PSNR (in dB)		Avg. SSIM [0,1]	
Training	Test	Training	Test	Training	Test
18.62	27.26	23.47	19.44	0.42	0.22

주관적 평가를 위해서 그림 7, 8은 Train 영상과 Test 영상의 실측 출력 영상, 및 예측 영상을 보여준다. 그림 7(a), 8(a)는 그림 4에서 Target reflectance image에 해당하는 영상이다. 이 영상은 Retinex 이론에 기반하여 생성된 영상으로 신경망의 타겟 영상에 해당된다. 그림 7(b), 8(b)는 Predicted reflectance image이다. Test 영상의 1, 2, 3열에서는 예측 영상이 실측 영상과 유사하거나 우수한 화질을 보여준다. 반대로 4, 5열에서는 상대적으로 화질이 저하되는 것이 관찰된다. 1, 2, 3열의 영상의 일부를 확대한 결과는 그림 9에서 보여진다. 그림 9(a)는 실측 출력 영상이고, 9(b)는 예측 영상이다. 실측 영상과 비교해보면, 예측 영상이 저조도 영역을 보다 세밀하게 복원해주는 것을 관찰할 수

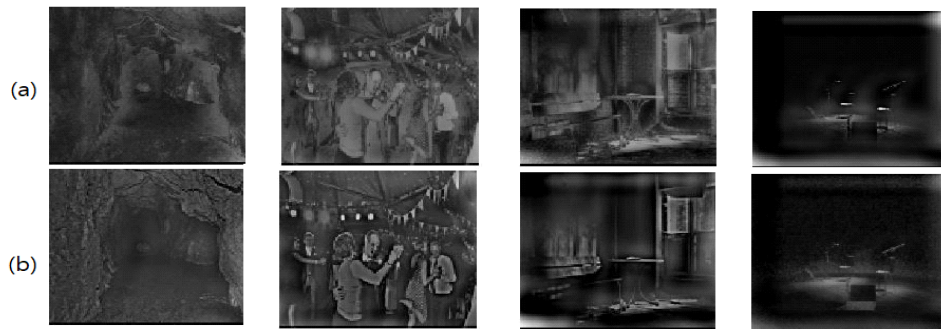


그림 7. 학습 영상의 결과. 시각화를 위해서 [0,255]로 변환함. (a) 실측 출력 영상, 및 (b) 예측 영상

Fig. 7. Resulting images of train images. (a) Ground-truth output images, and (b) predicted images



그림 8. 테스트 영상의 결과 영상. 시각화를 위해서 [0,255]로 변환함. (a) 실측 출력 영상, 및 (b) 예측 영상

Fig. 8. Resulting images of test images. (a) Ground-truth output images, and (b) predicted images

있다.

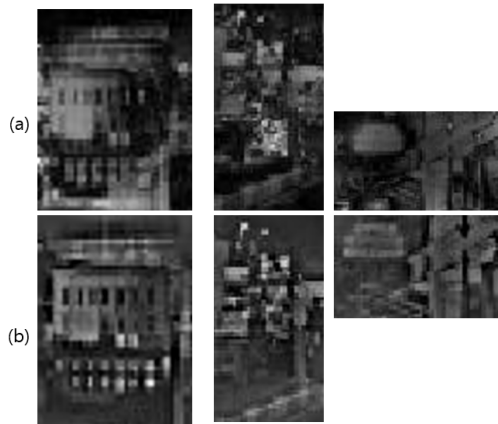


그림 9. 그림 8의 1, 2, 3열 영상의 확대영상. (a) 실측 출력 영상 및 (b) 예측 영상
Fig. 9. Close-up of Columns 1, 2, 3 in Fig. 8. (a) Ground-truth output images, and (b) predicted images

V. 결 론

저조도 영상에서는 빛의 밝기가 부족하여 영상에서 객체를 인식하거나 추적, 감지 등이 힘들어진다. 이를 극복하기 위한 기본 방법은 Retinex 이론을 이용한 조명 성분의 감쇄 및 반사 성분의 생성으로, 이를 활용하는 다양한 기법들이 제안되었다. 본 논문에서는 이러한 Retinex 이론에 기반한 반사 성분 생성을 합성곱 신경망으로 대체하는 방법을 제안하였고, 가능성을 조사하였다. 저조도 영상 102장을 이용한 실험 결과, Train에서는 RMSE는 18.62, PSNR은 23.47(dB), SSIM은 0.42이고, Test에서는 RMSE는 27.26, PSNR은 19.44(dB), SSIM은 0.22이다. 성능 측면에서는 만족스러운 결과를 얻을 수 있었다. 일부 영상에서 화질 저하가 발생하는데, 충분한 학습 데이터의 부족으로 발생한다.

본 연구는 기존의 hand-crafted 반사 영상 생성 기법을 합성곱 신경망으로 대체하는 예를 보여 주었다. 더욱이 높은

복잡도를 가지는 반사 영상 생성 기법들을 낮은 복잡도를 가지는 CNN으로 대체할 수 있으면, 상당한 활용이 기대될 수 있다. 향후 딥러닝이 다양한 분야의 기술들을 대체할 것으로 기대되기 때문에, 제안 방법의 활용이 가능하다.

참 고 문 헌 (References)

- [1] D. Cho, H. Kang, and W. Kim, "An image enhancement algorithm on color constancy and histogram equalization using edge region", *Journal of Broadcast Engineering*, Vol. 15, No. 3, 2010.
- [2] H. Cheng and X. Shi, "A simple and effective histogram equalization approach to image enhancement", *Digital Signal Processing*, pp. 158-170, 2004.
- [3] C. Poynton, "The Rehabilitation of Gamma," *Human Vision and Electronic Imaging III, Proceedings of SPIE*, vol. 3299, p. 232-249, 1998.
- [4] E. Land and J. McCann, "Lightness and retinex theory," *J. Opt. Soc. Am.*, vol. 61, no. 1, pp. 1 - 11, 1971.
- [5] D. Jobson, Z. Rahman, and G. Woodell, "Properties and performance of a center/surround retinex," *IEEE Transactions on Image Processing*, Vol. 6, no. 3, pp. 451-462, Mar. 1997.
- [6] K. Lore, A. Akintayo, and S. Sarkar, "LLNet: A deep autoencoder approach to natural low-light image enhancement," *Pattern Recognition*, vol. 61, pp. 650-662, 2017.
- [7] L. Shen, Z. Yue, F. Feng, and Q. Chen, "MSR-net: Low-light image enhancement using deep convolutional network," *arXiv:1711.02488*.
- [8] S. Park, S. Yu, M. Kim, K. Park, and J. Paik, "Dual Autoencoder Network for Retinex-based Low-Light Image Enhancement", *IEEE Access*, Vol. 6, Mar. 2018.
- [9] W. Kim, I. Hwang and M. Kim, "Generating a Retinex-based reflectance image from a low-light image using deep neural network", *Journal of Broadcast Engineering*, Vol. 24, No. 1, Jan. 2019.
- [10] L. Xu, C. Liu, J. Ren, and J. Jia, "Deep convolutional neural network for image deconvolution", *Int. Conf. Neural Information Processing System*, 2014.
- [11] A. Agrawal, R. Raskar, "Resolving objects at higher resolution from a single motion-blurred image", *Int. Conf. Computer Vision and Pattern Recognition*, 2007
- [12] C. Turhan and H. Bilge, "Recent trends in deep generative models: a review", *3rd Int. Conf. Computer Science and Engineering*, 2018.
- [13] Z. Wang, A. Bovik, H. Sheikh, and E. Simoncelli, "Image quality assessment: from error visibility to structural similarity". *IEEE Trans. Image Processing*, 13(4), pp. 600-612, 2004.

저 자 소 개



이 승 수

- 2017년 8월 : 강원대학교 IT대학 컴퓨터정보통신공학과 학사
- 2017년 9월 ~ 현재 : 강원대학교 컴퓨터정보통신공학과 석사과정
- 주관심분야 : 점유센서, 딥러닝, 머신러닝



최 창 열

- 1979년 : 경북대학교 전자공학과 학사
- 1981년 : 경북대학교 전자공학과 석사
- 1995년 : 서울대학교 컴퓨터공학과 공학박사
- 1984년 ~ 1996년 : ETRI 컴퓨터연구단 책임연구원 / 연구실장
- 2009년 ~ 2011년 : 강원대학교 IT대학 학장
- 1996년 ~ 현재 : 강원대학교 컴퓨터정보통신공학과 교수
- ORCID : <http://orcid.org/0000-0002-8340-4195>
- 주관심분야 : 모바일전송, 3D데이터처리, 미디어서비스



김 만 배

- 1983년 : 한양대학교 전자공학과 학사
- 1986년 : University of Washington, Seattle 전기공학과 공학석사
- 1992년 : University of Washington, Seattle 전기공학과 공학박사
- 1992년 ~ 1998년 : 삼성종합기술원 수석연구원
- 1998년 ~ 현재 : 강원대학교 컴퓨터정보통신공학과 교수
- 2016년 ~ 2018년 : 강원대학교 정보통신연구소 소장
- ORCID : <http://orcid.org/0000-0002-4702-8276>
- 주관심분야 : 3D영상처리, 비전점유센서, 객체인식