

특집논문 (Special Paper)

방송공학회논문지 제23권 제5호, 2018년 9월 (JBE Vol. 23, No. 5, September 2018)

<https://doi.org/10.5909/JBE.2018.23.5.614>

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

Generative Adversarial Network를 이용한 손실된 깊이 영상 복원

나 준 엽^{a)}, 심 창 훈^{a)}, 박 인 규^{a)†}

Depth Image Restoration Using Generative Adversarial Network

John Junyeop Nah^{a)}, Chang Hun Sim^{a)}, and In Kyu Park^{a)†}

요 약

본 논문에서는 generative adversarial network (GAN)을 이용한 비감독 학습을 통해 깊이 카메라로 깊이 영상을 취득할 때 발생한 손실된 부분을 복원하는 기법을 제안한다. 제안하는 기법은 3D morphable model convolutional neural network (3DMM CNN)와 large-scale CelebFaces Attribute (CelebA) 데이터 셋 그리고 FaceWarehouse 데이터 셋을 이용하여 학습용 얼굴 깊이 영상을 생성하고 deep convolutional GAN (DCGAN)의 생성자(generator)와 Wasserstein distance를 손실함수로 적용한 구별자(discriminator)를 미니맥스 게임 기법을 통해 학습시킨다. 이후 학습된 생성자와 손실 부분을 복원해주기 위한 새로운 손실함수를 이용하여 또 다른 학습을 통해 최종적으로 깊이 카메라로 취득된 얼굴 깊이 영상의 손실 부분을 복원한다.

Abstract

This paper proposes a method of restoring corrupted depth image captured by depth camera through unsupervised learning using generative adversarial network (GAN). The proposed method generates restored face depth images using 3D morphable model convolutional neural network (3DMM CNN) with large-scale CelebFaces Attribute (CelebA) and FaceWarehouse dataset for training deep convolutional generative adversarial network (DCGAN). The generator and discriminator equip with Wasserstein distance for loss function by utilizing minimax game. Then the DCGAN restore the loss of captured facial depth images by performing another learning procedure using trained generator and new loss function.

Keyword : Deep learning, generative adversarial network, depth image, depth camera, restoration

a) 인하대학교 정보통신공학과(Inha University, Department of Information and Communication Engineering)

† Corresponding Author : 박인규(In Kyu Park)

E-mail: pik@inha.ac.kr

Tel: +82-32-860-9190

ORCID: <https://orcid.org/0000-0003-4774-7841>

※ 이 논문의 연구 결과 중 일부는 “IPIU 2018”에서 발표한 바 있음.

※ 이 논문은 2016년도 정부(미래창조과학부)의 재원으로 한국연구재단의 지원을 받아 수행된 연구임(NRF-2016R1A2B4014731).

※ This work was supported by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIP) (No. NRF-2016R1A2B4014731).

· Manuscript received May 8, 2018; Revised July 13, 2018; Accepted July 16, 2018.

I. 서론

구조적 조명이나 Time-of-flight (ToF) 방식의 깊이 카메라는 깊이 영상을 취득하는 과정에서 장면의 특성에 따라 정보의 손실이 발생한다. 특히 얼굴의 경우, 손이나 안경에 의한 가리어짐에 의해 깊이 정보를 손실하는 경우가 많다. 기존에 잡음 제거를 위한 영상 후처리를 통해 이를 보완해주는 알고리즘이 사용되고는 있지만 가리어짐에 의해 완전히 정보가 손실된 부분에 대해 효과적으로 영상을 복원하는 것은 쉽지 않은 일이다.

기존 연구에서는 깊이 카메라의 특성 상 발생하는 손실에 대해 복원하는 알고리즘을 제안하였다^{[1][2]}. 또한 깊이와 RGB 영상에서의 텍스처를 이용한 패치 기반의 인페인팅(inpainting)을 통해 일반적인 깊이 영상에서의 손실을 복원하는 알고리즘을 제안하였다^[3]. 하지만 이는 주변 픽셀 또는 다른 영상의 정보를 이용하여 복원하는 것이므로 손실 부분이 클 때 매우 제한적이다. 앞서 설명한 기존 방식이 아닌 DCGAN^[4]을 이용하여 RGB 얼굴 영상의 손실된 부분을 복원해주는 연구도 이루어졌다^[5]. 하지만 DCGAN의 생성자와 구별자를 학습시키기 위해 미니맥스 게임을 수행할 때 구별자에 Jensen-Shannon divergence를 손실 함수로 사용하여 크게 좋은 성능을 보이지 못한다.

본 논문에서는 깊이 카메라로 취득한 얼굴 깊이 영상에서 완전히 손실된 부분을 복원하기 위한 새로운 알고리즘을 제안한다. 기존 알고리즘으로 손실된 깊이 정보를 복원해주는 방식과 달리 비감독 학습 방법 중 영상 처리에 특화된 DCGAN^[4] 구조를 사용한다. 그리고 DCGAN의 구별자를 좀 더 효과적으로 학습시키기 위해 기존에 사용되어 왔던 Jensen-Shannon divergence에 비해 무른(weak) 특성을 보이는 Wasserstein distance^[6]를 손실함수로 이용하여 문제를 해결하고자 한다. DCGAN은 미니맥스 게임기법을 통해 적대적인 학습을 하는 GAN^[7]에서의 생성자와 구별자의 신경망 구조에 영상 처리를 위해 convolutional neural network (CNN) 구조를 적용한 딥러닝 기법이다^[4]. DCGAN의 미니맥스 게임을 진행하기 위해 3DMM CNN^[8]과 CelebA 데이터 셋 그리고 FaceWarehouse 데이터 셋^[9]을 이용하여 총 23,500개의 학습용 데이터 셋을 만들어주어 미니맥스 게임을 수행한다. 마지막으로 깊이 카메라로 직접 취득한

실제 얼굴 깊이 영상에 여러 비율의 이진 마스크를 각 부위별로 적용하여 손실된 부분을 만들어주며 학습용 데이터를 통해 훈련시킨 생성자를 이용하여 손실된 부분을 복원한 결과를 보인다.

본 논문에서 기존의 연구에 비해 새롭게 제시하는 기여점을 정리하면 다음과 같다.

- 깊이 정보를 복원하는 데에 주변 픽셀 또는 RGB 영상을 이용하는 방법이 아닌 비감독 학습 방법을 채택하여 손실 부분이 클 때에도 복원이 가능함
- DCGAN의 구별자에 Wasserstein distance를 손실 함수로 도입하여, 더 우수한 복원 결과를 보임

본 논문의 구성은 다음과 같다. 2절에서는 미니맥스 게임을 수행하기 위한 학습용 얼굴 깊이 영상 데이터 셋과 깊이 카메라를 이용하여 실제 얼굴 깊이 영상을 취득하기 위한 기법을 설명한다. 3절에서는 DCGAN의 미니맥스 게임기법과 손실된 깊이 영상의 복원 과정을 서술하며, 4절에서는 실험 결과를 제시한다. 5절에서 본 논문의 결론을 맺는다.

II. 얼굴 깊이 영상 취득

1. 학습용 얼굴 깊이 영상 취득

본 논문에서 사용한 3DMM CNN은 학습된 CNN 구조를 통해 단일 얼굴 영상에서 얼굴을 검출하여^[8] 취득한 매개 변수들과 총 여덟 개의 주성분 분석(principal component analysis, PCA) 매개 변수 변화를 통해 다양한 3차원 얼굴 형태를 생성하는 Basel face model (BFM)^[10]을 이용하여 검출한 얼굴의 3D 객체를 생성한다. 이를 CelebA 데이터 셋에 적용하여 15,000개의 3D 얼굴 모델을 생성하였고 인종의 다양성을 위해 FaceWarehouse 데이터 셋^[9] 8,500개를 추가적으로 사용하였다. 마지막으로 총 23,500개의 3D 얼굴 모델들로부터 깊이 영상을 취득하여 본 논문에서 제시한 딥러닝 구조에 이용하기 위해 그림 1과 같이 관심 영역만을 취한 뒤 깊이의 범위를 (-1, 1)의 범위로 정규화 하였다.

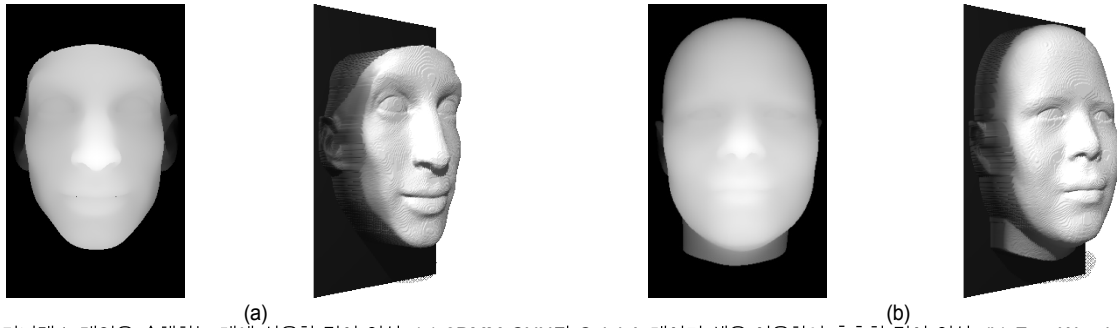


그림 1. 미니맥스 게임을 수행하는 데에 사용한 깊이 영상. (a) 3DMM CNN과 CelebA 데이터 셋을 이용하여 추출한 깊이 영상, (b) FaceWarehouse 데이터 셋의 깊이 영상

Fig. 1. Depth image used in minimax game. (a) Extracted depth image by using 3DMM CNN and CelebA dataset, (b) Depth image from FaceWarehouse dataset

2. 깊이 카메라를 이용한 얼굴 깊이 영상 취득

실제 얼굴의 깊이 영상을 취득하기 위해 20cm에서 150 cm 범위의 깊이 영상을 취득하는 데에 최적화 된 Intel® RealSense™ SR300 깊이 카메라를 사용하였다^[11]. 취득한 얼굴 깊이 영상의 사례를 그림 2에 제시하였다. 취득된 깊이 영상은 학습용 데이터 셋과 동일한 과정을 통해 정규화하였다.

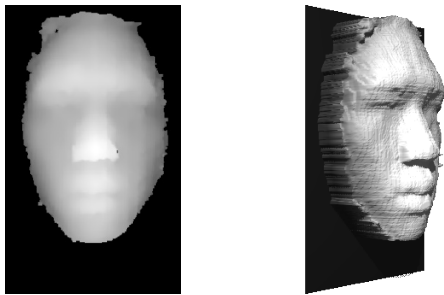


그림 2. 깊이 카메라로 취득한 얼굴 깊이 영상

Fig. 2. Facial depth image captured by depth camera

III. DCGAN 학습 및 깊이 영상 복원

1. 미니맥스 게임기법을 이용한 생성자 학습

본 논문에서는 얼굴 깊이 영상 학습을 위해 생성자에는 -1 부터 1 사이의 난수들을, 구별자에는 정규화 된 학습용 데이터 23,500개를 입력하여 서로 경쟁하며 상호 발전하는 미니맥스 게임을 진행한다. 미니맥스 게임기법을 통해 생성자는 난수들을 이용하여 최대한 구별자에 입력된 데이터와 비슷한 결과 값을 만들도록, 구별자는 생성자가 만든 결과와 입력된 학습용 데이터를 서로 구별할 수 있도록 학습된다^[7].

본 논문에서는 DCGAN의 구별자에 기존에 사용된 손실 함수인 Jensen-Shannon divergence 대신 좀 더 무른 성질을 가지고 있는 식 (1)의 Wasserstein distance^[6]를 적용하여 그림 4 (b)와 같이 미니맥스 게임의 결과를 더욱 개선하였다.

$$W(\mathbb{P}_r, \mathbb{P}_g) = \inf_{\gamma \in \Pi(\mathbb{P}_r, \mathbb{P}_g)} \mathbb{E}_{(x,y) \sim \gamma} [\|x - y\|] \quad (1)$$

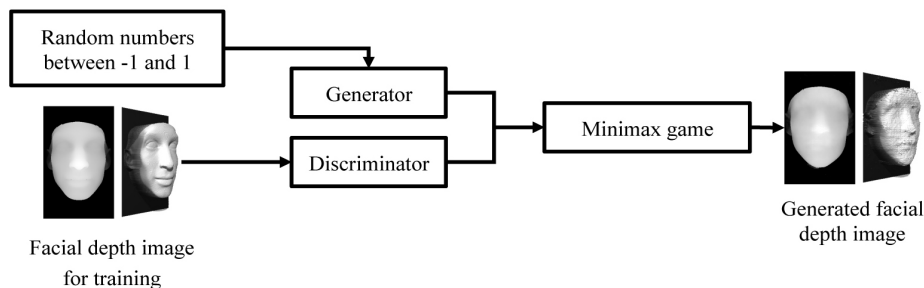


그림 3. 미니맥스 게임기법을 이용한 학습 과정

Fig. 3. Training procedure using minimax game

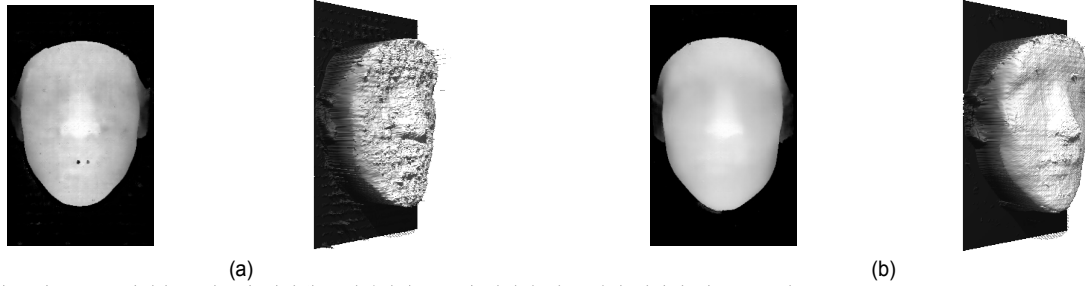


그림 4. 서로 다른 두 손실함수를 이용한 미니맥스 게임기법을 통해 생성된 얼굴 깊이 영상의 비교. (a) 기존 DCGAN, (b) Wasserstein distance를 적용한 DCGAN (학습률: 0.0002, 반복 횟수: 10,000)
Fig. 4. Comparison of generated facial depth image after performing minimax game using two different loss functions. (a) Original DCGAN, (b) DCGAN with Wasserstein distance (Learning rate: 0.0002, iteration count: 10,000)

$\Pi(P_r, P_g)$ 는 두 확률분포 P_r, P_g 의 결합확률분포(joint probability distribution)들의 집합이고 γ 는 그 부분집합이다. 그리고 $\inf$ 함수는 하한(lower bound)에서 최대값을 의미한다. 따라서 Wasserstein distance는 두 확률분포의 모든 결합확률분포 중에서 두 점 사이의 거리(metric)의 기대값 $\mathbb{E}$ 를 가장 작게 추정한 값을 의미한다^[6].

2. 학습된 생성자를 이용한 손실된 깊이 영상 복원

우선 각 부위 별로 복원 정도를 측정하기 위해 가려져있지 않은 얼굴 깊이 영상에 직접 0과 1로 구성된 이진 마스크를 씌우는 방법으로 네모난 손실 부분을 각 부위 별로 생성하였다. 그 후 위 과정에서 학습된 생성자 G 와 복원을 위한 손실함수 공식을 이용하여 새로운 학습을 통해 손실 부분을 복원하였다^[5].

손실함수는 다음과 같이 문맥적인 부분과 지각적인 부분으로 나눠 계산된다^[3].

$$\mathcal{L}_{contextual}(z) = \|M \odot G(z) - M \odot y\|_1 (\|x\|_1 = \sum_i |x_i|) \quad (2)$$

$$\mathcal{L}_{perceptual}(z) = \log(1 - D(G(z))) \quad (3)$$

먼저 식 (2)는 생성자가 만들어낸 깊이 영상과 이진 마스크 부분 외의 깊이 값 사이의 거리를 나타내는데 이는 기존 얼굴과 비슷하게 생성하기 위한 문맥적인 손실 값이다. 여기서 M 은 이진 마스크, $G(z)$ 는 학습된 생성자에 -1부터 1 사이의 난수들을 넣어 만들어낸 깊이 영상, $\odot$ 는 행렬 원소 간의 곱을 의미한다.

다음으로 식 (3)은 복원한 영상이 실제 사람의 얼굴과 같이 보이도록 하기 위한 지각적인 손실 값이며 기존 GAN의 미니맥스 게임시 구별자를 훈련시키기 위한 손실 값과 동일하다^[7]. 여기서 $D(G(z))$ 는 미니맥스 게임기법을 통해 학습된 구별자에 생성자가 만든 깊이 영상을 넣어 구별한 결과 값이다.

위 식 (2)와 (3)의 선형 가중치 결합으로 최종 손실함수를 표현하면 다음과 같다.

$$\mathcal{L}(z) = \mathcal{L}_{contextual}(z) + \lambda \mathcal{L}_{perceptual}(z) \quad (4)$$

식 (4)를 최적화 함수를 이용하여 계속해서 최소값을 찾

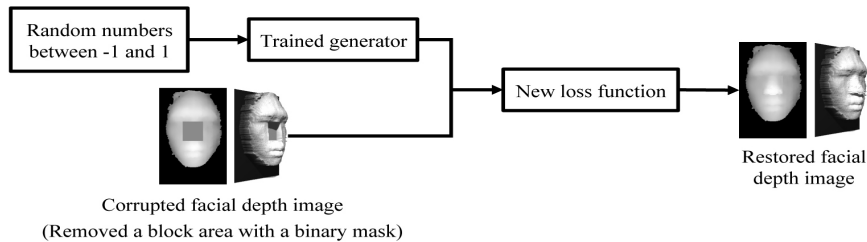


그림 5. 학습된 생성자를 이용한 손실된 깊이 영상 복원
Fig. 5. Depth image restoration using trained generator

으며 그 최소값을 가질 때의 입력 값이 식 (5)에서의 $\hat{z}$ 이다.

$$\hat{z} = \operatorname{argmin} \mathcal{L}(z) \quad (5)$$

최종적으로 손실 값이 최소가 되는 $\hat{z}$ 값과 생성자 G 및 이진 마스크를 이용하여 식 (6)을 통해 손실된 부분을 복원한다.

$$x_{reconstructed}(z) = M \odot y + (1 - M) \odot G(\hat{z}) \quad (6)$$

IV. 실험 결과

본 논문에서의 실험은 TensorFlow 프레임워크 상에서 NVIDIA GeForce GTX 1080 Ti GPU를 사용하여 총 5일간 수행하였으며 DCGAN의 미니맥스 게임 및 손실된 깊이

표 1. 미니맥스 게임 및 손실된 깊이 영상을 복원할 때 사용한 하이퍼 파라미터들
Table 1. Hyper parameters used in minimax game and restoration

Hyper Parameters	Values
Iteration Number	25,000
Batch Size	64
Regularization Weight (λ)	0.1
Learning Rate	0.0001, 0.00012, 0.00015, 0.00017, 0.0002
Optimization Function	Adam optimizer

영상을 복원할 때 사용한 하이퍼 파라미터(hyper parameter)들은 다음 표 1에 정리하였다.

그림 6은 그림 4의 (b)에서 더 나아가 총 25,000번 배치를 가져와 미니맥스 게임을 하며 학습된 생성자가 만들어진 얼굴 깊이 영상이다. 손실된 부분을 복원하는 데에 좀 더 다양한 학습률로 학습된 생성자를 사용하기 위해 미니맥스 게임을 총 다섯 가지의 학습률을 사용하여 진행하였으며 그 결과 학습용으로 사용하였던 데이터 이외의 얼굴 깊이 영상들을 생성하였다.

그림 7은 깊이 카메라로 취득한 얼굴 깊이 영상을 손상시

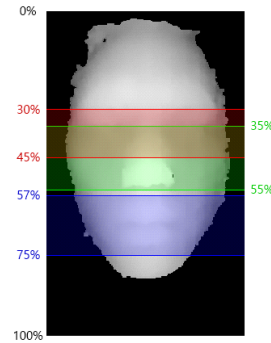


그림 7. 취득한 얼굴 깊이 영상에서 합성 실험 영상을 얻기 위해 이진 마스크를 수행할 눈, 코, 입 영역

Fig. 7. Eye, nose, mouth region of obtained facial depth image

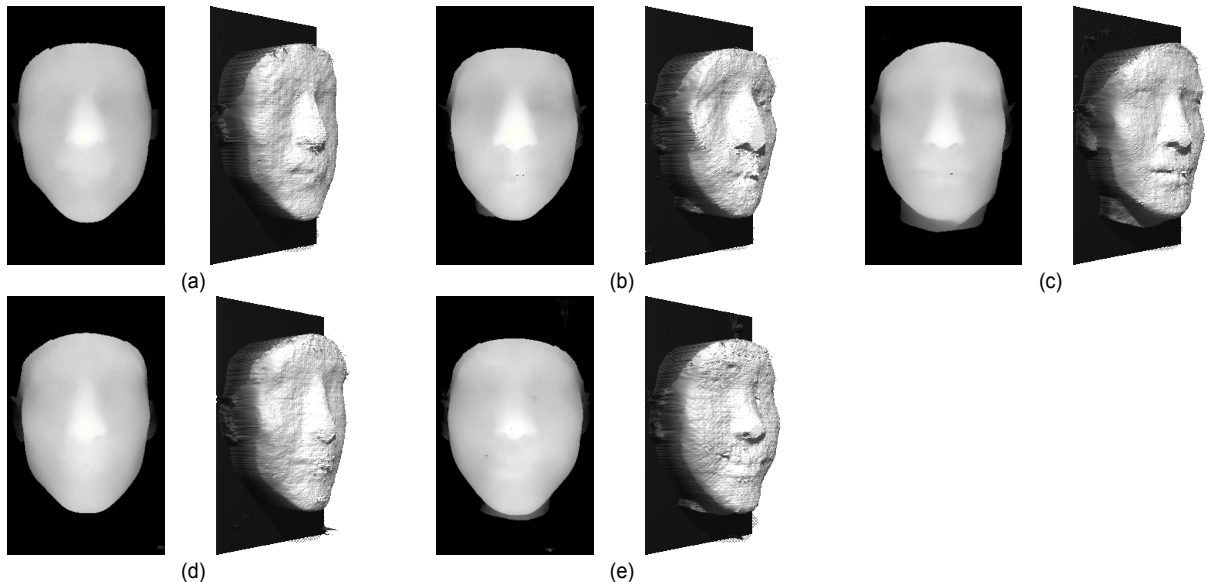


그림 6. 다섯 가지 학습률을 이용하여 학습된 생성자로 만들어진 얼굴 깊이 영상. (a) 0.0001, (b) 0.00012, (c) 0.00015, (d) 0.00017, (e) 0.0002

Fig. 6. Generated facial depth image by trained generator using five different learning rates. (a) 0.0001, (b) 0.00012, (c) 0.00015, (d) 0.00017, (e) 0.0002

켜 합성 실험 영상을 얻기 위해 눈, 코, 입에 이진 마스크를 어떻게 적용했는 지를 나타낸다. 각 부위는 얼굴 영역의 맨 윗부분을 기준으로 눈은 30% ~ 45% (빨간 부분), 코는 35% ~ 55% (초록 부분), 입은 57% ~ 75% (파란 부분)으로

지정하였으며 각 부위에 대해 이진 마스크의 높이를 고정시키고 너비를 변화시켜가며 손실된 깊이 영상을 생성하였다.

그림 8은 위 그림 7에서 설명한 눈, 코, 입 부분에 각각 이진 마스크를 여러 비율에 따라 적용시켜 합성 실험 영상

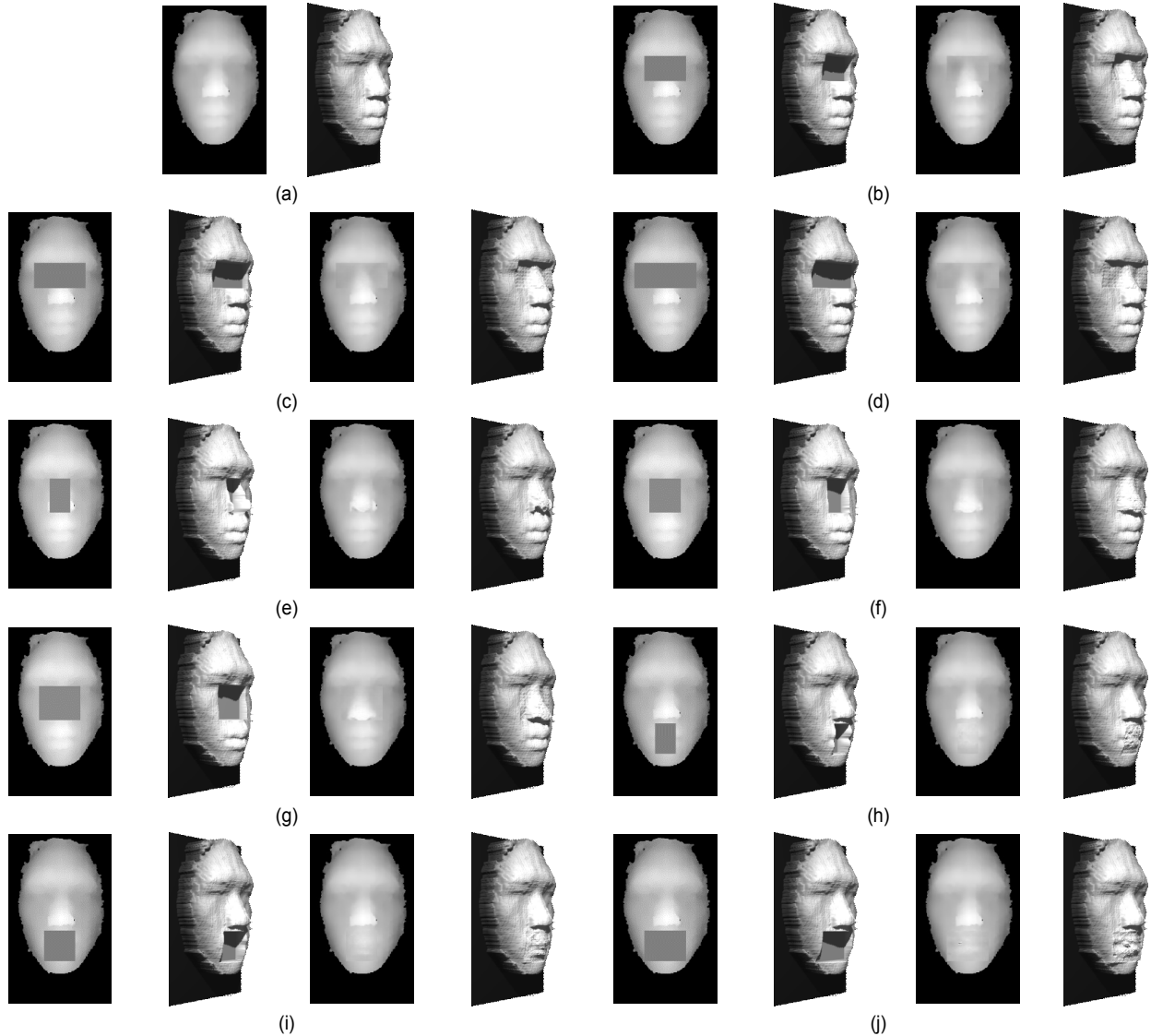


그림 8. 실제 깊이 영상에 각 부위 별 여러 비율의 이진 마스크를 적용하여 영상을 손상시킨 후 다시 복원한 결과. (a) 원본 깊이 영상, (b), (c), (d) 눈 부위에 각각 너비의 40%, 50%, 60%를 이진 마스크 삭제한 뒤 복원한 결과, (e), (f), (g) 코 부위에 각각 너비의 20%, 30%, 40%를 이진 마스크로 삭제한 뒤 복원한 결과, (h), (i), (j) 입 부위에 각각 너비의 20%, 30%, 40%를 이진 마스크로 삭제한 뒤 복원한 결과 (미니맥스 게임에 사용된 학습률: 0.00017, 복원 과정에 사용된 학습률: 0.00015)

Fig. 8. Results of depth image restoration after corrupting the original depth image with binary masks with different area. (a) Original depth image, (b), (c), (d) Restoration result after corrupting 40%, 50%, and 60% of eyes part width, (e), (f), (g) Restoration result after corrupting 20%, 30%, and 40% on nose part width, (h), (i), (j) Restoration result after corrupting 20%, 30%, and 40% on mouth part width (Learning rate used in minimax game: 0.00017, learning rate used in restoration: 0.00015)

을 생성하고 손상된 영역을 복원한 실험 결과이다. 이진 마스크의 비율은 깊이 영상의 중앙을 기준으로 눈은 깊이 영상의 너비의 40%, 50%, 60%, 코와 입은 20%, 30%, 40%로 변경하며 실험하였다. 복원 과정에서의 손실함수를 위한 학습률 또한 미니맥스 게임과 동일하게 총 다섯 가지를 사용하였으며 그 중 미니맥스 게임에서 사용한 학습률이 0.00017, 복원 과정에서 사용한 학습률이 0.00015일 때 결과가 육안으로 봤을 때 가장 우수한 결과를 생성하였다.

제안하는 알고리즘의 성능을 정량적으로 분석하기 위하여, 그림 8에서 각 부위 별로 복원한 부분과 원본 깊이 영상에서 동일한 부분을 비교하여 PSNR 값을 계산한 결과를 표 2에 제시하였다. 실험 결과 PSNR과 이진 마스크의 비율은 큰 관련이 없었으며 코 부위에서의 PSNR이 다른 부위에 비해 좀 더 높게 측정되었다. 이는 눈과 입 부위에 비해 코 부위의 복원 정도가 더 뛰어나는 나타내며 코 부위에서 깊이 값의 변화가 더욱 크기 때문에 학습된 생성자가 복원하기 쉬웠을 것이라 추정하였다.

그림 9는 실제로 깊이 카메라를 사용하여 취득할 때 손실된 부분을 위에서 학습시킨 생성자를 이용하여 복원한 결과이다. 손실된 부분을 생성하기 위해 깊이 범위를 30cm에서 40cm 사이로 제한하여 얼굴을 손으로 가려 촬영하였고 이진 마스크는 취득한 깊이 영상에서 값이 0인 부분만을

지정하여 복원해주었다.

V. 결 론

본 연구를 통해 얼굴 깊이 영상을 데이터를 이용하여 미니맥스 게임기법을 통해 생성자와 구별자를 학습시켜 얼굴 깊이 영상에서 완전히 손실된 부분이 발생하였을 때 이를 얼굴 형상에 맞게 복원하였으며 기존 DCGAN에서 사용된 구별자의 손실함수가 아닌 Wasserstein distance를 이용하여 좀 더 개선된 결과를 도출하였다.

따라서 앞으로 이를 더욱 강화하고 최적화한다면 얼굴 깊이 영상에서 뿐만 아니라 깊이 카메라로 다른 일반적인 깊이 영상을 취득했을 때 손실된 부분을 깊이 카메라 자체에서 적응적으로 복원할 수 있고 기존 알고리즘과 차별화된 손실 복원을 해줄 수 있을 것임을 제시하였다.

참 고 문 헌 (References)

- [1] K. Xu, J. Zhou, and Z. Wang, "A method of hole-filling for the depth map generated by Kinect with moving objects detection," *Proceeding of IEEE international Symposium on Broadband Multimedia Systems and Broadcasting*, pp. 1-5, June 2012.

표 2. 각 부위 별 복원한 결과와 원본 깊이 영상과의 비교 결과

Table 2. Comparison of restoration results and original depth images for each part

	Eyes			Nose			Mouth		
Image Corruption Ratio (%)	40	50	60	20	30	40	20	30	40
PSNR (dB)	29.34	30.06	30.11	31.83	31.7	30.39	29.03	29.15	29.6



(a)



(b)



(c)

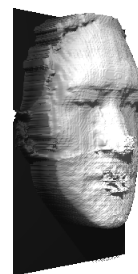


그림 9. 영상에서 발생한 손실 부분을 복원한 결과. (a) RGB 영상, (b) 손실과 함께 취득된 얼굴 깊이 영상, (c) 손실된 부분을 복원한 결과

Fig. 9. Restoration result of real corrupted facial depth image. (a) RGB image, (b) Obtained corrupted facial depth image, (c) Restoration result on loss part

- [2] L. Feng, L.-M. Po, X. Xu, K.-H. Ng, C.-H. Cheung, and K.-w. Cheung, "An adaptive background biased depth map hole-filling method for Kinect," *Proceeding of IEEE Industrial Electronics Society*, pp. 2366 - 2371, November 2013.
- [3] S. Ikehata, J. Cho, and K. Aizawa, "Depth map inpainting and super-resolution based on internal statistics of geometry and appearance," *Proceeding of IEEE International Conference on Image Processing*, pp. 938-942, September 2013.
- [4] A. Radford, L. Metz, and S. Chintala, "Unsupervised representation learning with deep convolutional generative adversarial networks," *Proceeding of International Conference on Learning Representations*, May 2016.
- [5] R. A. Yeh, C. Chen, T. Yian Lim, A. G. Schwing, M. Hasegawa-Johnson, and M. N. Do, "Semantic image inpainting with deep generative models," *Proceeding of IEEE Conference on Computer Vision and Pattern Recognition*, July 2017.
- [6] M. Arjovsky, S. Chintala, and L. Bottou, "Wasserstein GAN," *Proceeding of International Conference on Machine Learning*, vol. 70, pp. 214-223, August. 2017.
- [7] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial networks," *Proceeding of Advances in Neural Information Processing Systems*, December 2014.
- [8] A. T. Tran, T. Hassner, I. Masi, and G. Medioni, "Regressing robust and discriminative 3D morphable models with a very deep neural net-work," *Proceeding of IEEE Conference on Computer Vision and Pattern Recognition*, July 2017.
- [9] C. Cao, Y. Weng, S. Zhou, Y. Tong, and K. Zhou, "FaceWarehouse: a 3D facial expression database for visual computing," *IEEE Transaction on Visualization and Computer Graphics*, vol. 20, no. 3, pp. 413-425, March 2014.
- [10] P. Paysan, R. Knothe, B. Amberg, S. Romdhani, and T. Vetter, "A 3D face model for pose and illumination invariant face recognition," *Proceeding of IEEE International Conference on Advanced Video and Signal Based Surveillance*, October 2009.
- [11] Intel® RealSense™ Camera SR300, <https://software.intel.com/sites/default/files/managed/0c/ec/realsense-sr300-product-data-sheet-rev-1-0.pdf> (accessed August 13, 2018).

저 자 소 개

나 준 엽



- 2018년 2월 : 인하대학교 정보통신공학과 학사
- 2018년 3월 ~ 현재 : 인하대학교 정보통신공학과 석사과정
- ORCID : <http://orcid.org/0000-0002-2576-5240>
- 주관심분야 : 컴퓨터비전, deep learning, GPGPU

심 창 훈



- 2018년 2월 : 인하대학교 정보통신공학과 학사
- ORCID : <http://orcid.org/0000-0002-1367-1896>
- 주관심분야 : 컴퓨터비전, deep learning, GPGPU

박 인 규



- 1995년 2월 : 서울대학교 제어계측공학과 학사
- 1997년 2월 : 서울대학교 제어계측공학과 석사
- 2001년 8월 : 서울대학교 전기컴퓨터공학부 박사
- 2001년 9월 ~ 2004년 3월 : 삼성중합기술원 멀티미디어랩 전문연구원
- 2007년 1월 ~ 2008년 2월 : Mitsubishi Electric Research Laboratories (MERL) 방문연구원
- 2014년 9월 ~ 2015년 8월 : MIT Media Lab 방문부교수
- 2018년 7월 ~ 2019년 2월 : University of California San Diego (UCSD) 방문연구원
- 2004년 3월 ~ 현재 : 인하대학교 정보통신공학과 교수
- ORCID : <http://orcid.org/0000-0003-4774-7841>
- 주관심분야 : 컴퓨터 그래픽스 및 비전 (영상기반 3차원 형상 복원, 증강현실, computational photography), GPGPU