

RAPGAN와 RRDB를 이용한 Image-to-Image Translation의 성능 개선

윤동식¹, 객노윤^{2*}

¹고려대학교 전기전자공학과 석사과정, ²백석대학교 컴퓨터공학부 교수

Performance Improvement of Image-to-Image Translation with RAPGAN and RRDB

Dongsik Yoon¹, Noyoon Kwak^{2*}

¹Graduate Student, Department of Electronics and Electrical Engineering, Korea University

²Professor, Division of Computer Engineering, Baekseok University

요약 본 논문은 RAPGAN(Relativistic Average Patch GAN)과 RRDB(Residual in Residual Dense Block)을 이용한 Image-to-Image 변환의 성능 개선에 관한 것이다. 본 논문은 Image-to-Image 변환의 일종인 기존의 pix2pix의 결점을 보완하기 위해 세 가지 측면의 기술적 개선을 통한 성능 향상을 도모함에 그 목적이 있다. 첫째, 기존의 pix2pix 생성자와 달리 입력 이미지를 인코딩하는 부분에서 RRDB를 이용함으로써 더욱 더 깊은 학습을 가능하게 한다. 둘째, RAPGAN 기반의 손실함수를 사용해 원본 이미지가 생성된 이미지에 비해 얼마나 진짜 같은지를 예측하기 때문에 이 두 이미지가 모두 적대적 생성 학습에 영향을 미치게 된다. 마지막으로, 생성자를 사전학습시켜 판별자가 초기에 학습되는 것을 억제하도록 조치한다. 제안된 방법에 따르면, FID 측면에서 기존의 pix2pix보다 평균 13% 이상의 우수한 이미지를 생성할 수 있었다.

주제어 : 적대적 생성 신경망, pix2pix, Conditional GAN, RRDB, Relativistic Average Patch GAN

Abstract This paper is related to performance improvement of Image-to-Image translation using Relativistic Average Patch GAN and Residual in Residual Dense Block. The purpose of this paper is to improve performance through technical improvements in three aspects to compensate for the shortcomings of the previous pix2pix, a type of Image-to-Image translation. First, unlike the previous pix2pix constructor, it enables deeper learning by using Residual in Residual Block in the part of encoding the input image. Second, since we use a loss function based on Relativistic Average Patch GAN to predict how real the original image is compared to the generated image, both of these images affect adversarial generative learning. Finally, the generator is pre-trained to prevent the discriminator from being learned prematurely. According to the proposed method, it was possible to generate images superior to the previous pix2pix by more than 13% on average at the aspect of FID.

Key Words : GAN, pix2pix, Conditional GAN, RRDB, Relativistic Average Patch GAN

1. 서론

적대적 생성 신경망(GAN: Generative Adversarial Network)[1]은 심층 신경망 설계에 변화가 있을 수 있음을 알려준 모델로, 일반적인 방식으로는 학습하기 힘든 생성 모델을 효과적으로 구축할 수 있는 장점이 있다. 이러한 GAN의 색다른 진화의 일종으로 널리 알려진 바 있는 Image-to-Image 변환인 pix2pix[2]는 Conditional GAN[3]을 사용해 이미지들 간의 변환을 수행하는 것으로, 통상의 GAN 모델에서는 해결할 수 없었던 생성된 이미지를 가변적으로 조절할 수 있는 장점이 있다. 예컨대 그림 지도를 위성 지도로 변환하거나 스케치 수준의 이미지를 실사 이미지로 변환하는 등의 다양하고 변화무쌍한 변환이 가능한 편리함이 있는 반면에 생성된 이미지의 해상도와 품질이 낮은 단점이 있다.

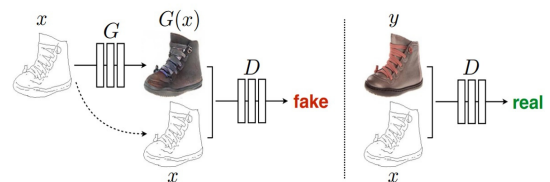
기존의 고해상도의 이미지를 사용해 이미지들 간의 변환을 수행하는 연구[4]가 있었지만 이 방법은 대량의 고해상도의 이미지 데이터셋을 확보해야 하는 불편함이 있다. 본 논문의 연구진은 이러한 문제점을 해소하기 위해 낮은 해상도 이미지를 사용하여 높은 품질을 갖는 이미지를 생성하는 모델을 제안하고자 한다. 본 논문은 기존의 pix2pix의 결점을 보완하기 위해 세 가지 측면의 기술적 개선을 통한 성능 향상을 도모함에 그 목적이 있다. 첫째, 기존의 pix2pix 생성자 층이 상대적으로 깊지 않은 것을 보완하기 위해 RRDB(Residual in Residual Block)[5]를 적층 구조로 결합해 생성자 네트워크를 구성한다. RRDB를 사용해 더욱 더 깊은 학습을 가능하게 함으로써 진짜와 유사한 고해상도의 이미지를 생성하고자 한다. 둘째, RAPGAN(Relativistic Average Patch GAN) 기반의 손실함수를 사용함으로써 원본 이미지가 적대적 학습에 영향을 주지 못해 파생됐던 문제를 해결하고자 한다. RAPGAN은 전체 영역을 패치 단위로 나눠 진짜 이미지가 생성된 가짜 이미지보다 얼마나 더 진짜와 유사한지를 패치 단위로 예측한다. 마지막으로, 통상 판별자(discriminator)의 학습 속도가 생성자(generator)보다 빠른 것을 감안해, 생성자와 판별자 간의 적대적 학습이 진행되기 전에 생성자만을 사전 학습시킴으로써 학습 초기에 판별자와 생성자가 수렴하여 모델이 훈련되지 않는 문제를 해결하고자 한다. 결과적으로 기존의 pix2pix보다 정량적이나 정성적으로 뛰어난 이미지를 생성할 수 있을 것으로 기대된다.

본 논문의 구성은 다음과 같다. 2장에서는 관련 연구에 대해 살펴보고, 3장에서는 제안된 pix2pix 네트워크

구조와 학습 방법에 대해 설명한다. 4장에서는 시뮬레이션 결과 및 고찰을 서술하고, 마지막으로 5장에서는 결론을 맺는다.

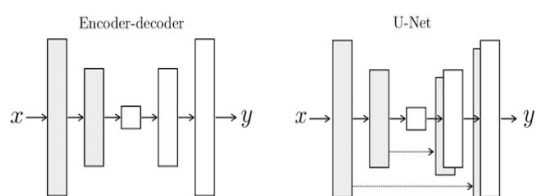
2. 관련 연구

딥러닝 기술의 비약적이고 급속한 발전에 힘입어 Image-to-Image 변환 분야[2]도 빠르고 독창적인 기술들[6, 7, 8, 9, 10, 11, 12, 13]이 다수 제안되어 있다. 특히 BicycleGAN[12]은 Image-to-Image 변환 문제에서 다양한 출력을 생성하기 위해 저장된 잠재 벡터와 Bijective Consistency를 활용한 방법을 통해 하나의 입력 이미지에 대해 다양한 시각적 특성을 가진 출력 이미지를 생성하는 것이 특징이다. 또한 UNIT[13]은 비지도학습 기반의 Image-to-Image 변환 모델로서 Coupled GANs 기반의 공유 잠재 공간을 활용한 공변량 변화 문제를 해결함으로써 주목할만한 성능을 제공하고 있다. 본 논문은 Image-to-Image 변환 모델들 중에서 구조가 직관적이며 성능이 좋은 pix2pix 모델[2]을 개선함에 그 목적이 있다. 기존의 pix2pix는 Conditional GAN을 사용하여 이미지 간의 변환을 수행하는데, [Fig. 1]은 pix2pix에서 Conditional GAN을 이용하여 네트워크를 학습시키는 과정을 나타낸 것이다. [Fig. 1]과 같이 우선, pix2pix를 학습시키기 위해선 원본 이미지 x 와 변환될 목표 이미지 y 를 짝 지은 데이터로 구성된 입력 데이터셋을 사전에 준비할 필요가 있다. 이 입력 데이터셋에 토대로 생성자 G 는 판별자 D 를 속이기 위해 원본 이미지 x 를 목표 이미지 y 에 가까운 가짜 이미지 $G(x)$ 를 만들어내도록 학습한다. 이후 판별자 G 는 목표 이미지 y 를 포함한 한쌍의 데이터는 진짜로 분류하고, 생성자가 생성한 가짜 이미지 $G(x)$ 를 포함한 한쌍의 데이터는 가짜로 분류하도록 학습한다.



[Fig. 1] Training process of pix2pix using Conditional GAN

이상과 같은 적대적 학습의 반복을 통해 생성자와 판별자는 각각 생성력과 판별력을 향상시킨다. 그 결과 생성자는 진짜 이미지와 유사한 가짜 이미지를 생성할 수 있게 되고, 또한 판별자는 판별자가 생성한 가짜 이미지에 대한 판별 능력이 50% 수준으로 수렴하게 된다. [Fig. 2]는 기존의 pix2pix 생성자의 네트워크 구조이다. pix2pix 생성자의 네트워크에서 채택하고 있는 Encoder-decoder는 구조상 이미지의 크기를 줄이고 늘리는 과정에서 형상(shape)은 유지되지만 정보 손실이 발생해 이미지가 블러링되는 문제가 있다.



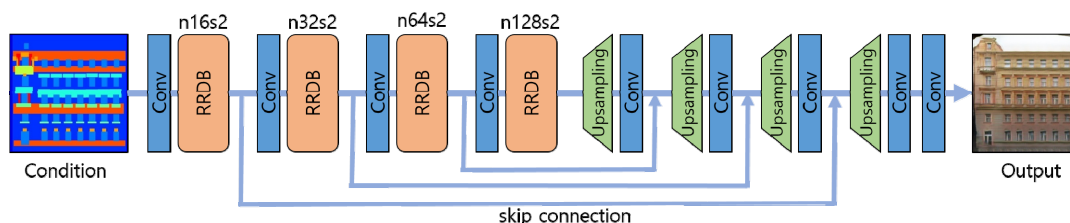
[Fig. 2] U-Net architecture of the original pix2pix generator

이를 개선하기 위해 채택하고 있는 pix2pix 생성자 네트워크의 U-Net 구조는 Skip Connection을 통해 이러한 문제를 회피함으로써 이미지 생성 과정에서 기존

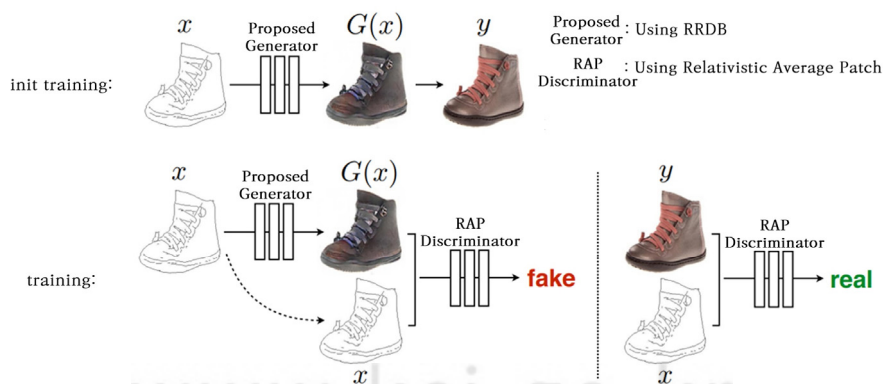
이미지의 정보를 직접 전달받아 선명한 이미지를 만드는 장점이 있다. 판별자 네트워크는 Patch GAN을 사용하여 이미지를 전체 영역이 아니라 특정 크기의 패치 단위로 이미지의 진위 여부를 판별한다. 이상과 같이 pix2pix의 생성자가 최종적으로 생성한 이미지의 품질을 높이기 위해 다양한 조치를 취하고 있음에도 불구하고 여전히 생성된 이미지의 해상도와 품질이 낮은 한계가 있다[2].

3. 제안된 pix2pix 네트워크 구조와 학습 방법

본 논문은 기존의 pix2pix의 픽셀간의 변환을 수행하는 개념은 유지하면서 세 가지 측면의 기술적 개선을 통해 성능향상을 도모한다. 따라서 기존의 pix2pix와 다른 네트워크 구조와 학습 과정으로 구성된다. [Fig. 3]은 RRDB를 적용한 제안된 생성자의 구조를 나타낸 것이고, [Fig. 4]는 제안된 pix2pix 네트워크의 전체 학습 과정을 나타낸 것이다. 본 논문에서는 판별자의 학습 속도와 생성자의 학습속도가 달라 초기에 모델이 수렴하는 문제를 해결하기 위해 제안된 RRDB를 적용한 생성자만을 사전 학습시킨



[Fig. 3] The architecture of the proposed generator layering the RRDB



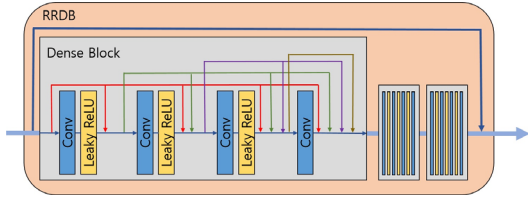
[Fig. 4] Training process of the proposed pix2pix

다. 이때 생성된 이미지와 원본 이미지의 픽셀 차이를 최소화시키는 방향으로 훈련을 진행한다. 그 후 기존의 pix2pix와 같이 제안된 생성자와 판별자를 사용한 적대적 학습을 수행한다.

3.1 RRDB(Residual in Residual Dense Block)

본 논문에서는 통상적으로, 층과 연결의 적절한 증대가 성능 향상에 기여한다는 연구 결과들[14, 15]과 기존의 pix2pix 생성자의 네트워크 구조가 Residual Block을 사용하지 않은, 얇은 구조라는 점에 착안해 RRDB를 사용한 구조적 개선을 통해 성능 향상을 도모한다.

타 연구진의 연구 성과에 따르면, 이미지 초해상화를 수행하는 네트워크[5]가 Upsampling 이전 단계에서 RRDB를 사용하여 뛰어난 성능을 낸다는 사실을 감안해, 본 논문에서는 [Fig. 3]와 같이 생성자의 Encoding 과정에서 RRDB를 적용하여 양질의 latent code를 얻고자 한다. RRDB는 [Fig. 5]와 같이 BN(Batch Normalization)층이 없는 것이 특징이다.



[Fig. 5] Architecture of RRDB

BN층은 배치의 평균과 분산을 사용해 특징들을 정규화하여 훈련하고 테스트 시엔 전체 데이터셋의 평균과 분산을 사용하기 때문에 훈련과 테스트 데이터셋의 통계가 크게 다른 경우, 시각적 거부감을 야기하고 일반화 성능을 저하시키는 경향이 있다. 기존의 이미지 초해상화 네트워크[5]에서는 이를 제외함으로써 훈련을 안정화시키면서 일반화 성능은 높이고 계산 복잡도를 낮추었다. 본 논문에서는 이 RRDB를 적절하게 적용함으로써 성능 향상을 도모하고 있다.

3.2 제안된 RAPGAN

본 논문에서는 생성자의 구조만 새로 제시한 것이 아니라 Relativistic GAN[16]과 Patch GAN[2]을 효과적으로 결합한 RAPGAN(Relativistic Average Patch GAN)을 제안한다. 기존의 pix2pix 판별자는 입력 이미지 x 가 얼마나 진짜 같은지에 대한 가능성을 판별하는

것이고, Relativistic GAN 판별자는 진짜 이미지 x_r 이, 생성된 가짜 이미지 x_f 보다 얼마나 더 진짜 같은지를 예측한다. 제안된 RAPGAN은 진짜 이미지 x_r 의 전체 영역이 아니라 특정 크기의 패치 단위의 평균이, 생성된 가짜 이미지 x_f 의 특정 패치 평균보다 얼마나 더 진짜 같은지를 예측한다. 본 논문에서 사용한 판별자는 식 (1)과 같이 나타낼 수 있다. 여기서 σ 는 시그모이드 함수이고 $C(x_r)$ 은 진짜 이미지의 특정 패치를 통과시킨 변환되지 않은 판별자의 출력을 의미한다.

$\mathbb{E}_{x_f}[C(x_f)]$ 는 미니 배치에 있는 모든 가짜 데이터의

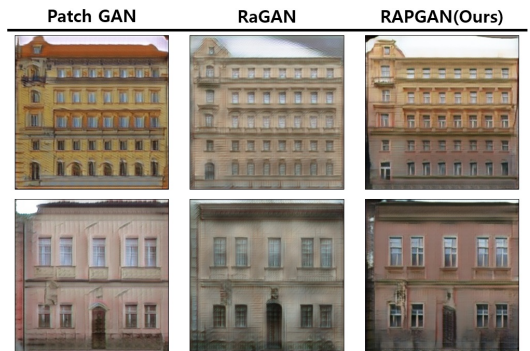
특정 패치 평균을 의미한다. 이를 사용한 손실함수는 식 (2) 및 식 (3)과 같으며 판별자와 생성자가 대칭적인 형태를 이룬다.

$$D_{Ra}(x_r, x_f) = \sigma(C(x_r) - \mathbb{E}_{x_f}[C(x_f)]) \quad (1)$$

$$L_D^{Ra} = -\mathbb{E}_{x_r}[\log(D_{Ra}(x_r, x_f))] - \mathbb{E}_{x_f}[\log(1 - D_{Ra}(x_f, x_r))] \quad (2)$$

$$L_G^{Ra} = -\mathbb{E}_{x_r}[\log(1 - D_{Ra}(x_r, x_f))] - \mathbb{E}_{x_f}[\log(D_{Ra}(x_f, x_r))] \quad (3)$$

여기서 $x_f = G(x_i)$ 이고 x_i 는 Condition 이미지를 의미한다. 식 (3)에서와 같이 생성자의 손실함수는 x_r 과 x_f 모두를 포함하기 때문에 기존의 pix2pix와 달리 생성된 이미지와 원본 이미지 모두 적대적 학습에 영향을 미친다.



[Fig. 6] Simulation results using the proposed RRDB stacked generator and several discriminators

[Fig. 6]은 제안된 RRDB를 적용한 생성자와 여러 판별자들을 사용해 시뮬레이션한 결과이다. [Fig. 6]의 가

장 우측의 그림과 같이 제안된 생성자 네트워크를 구성할 경우, 기존의 Patch GAN이나 Relativistic Average GAN보다 제안된 RAPGAN의 성능이 가장 우수한 것을 시각적으로 확인할 수 있다.

3.3 생성자의 사전 학습

본 논문에서는 기존 판별자의 학습속도가 생성자보다 빠른 것을 감안해 생성자를 사전 학습시켜 학습 초기에 판별자와 생성자가 수렴하여 모델이 혼란되지 않는 문제를 해결하고자 한다. 식 (4)는 생성자의 사전 학습 손실함수로, 생성된 이미지 $G(x_i)$ 와 원본 이미지 y 의 차이를 구한 것이다.

$$L_1 = \mathbb{E}_{x_i} \| G(x_i) - y \| \quad (4)$$

3.4 제안된 생성자와 판별자의 손실함수

제안된 생성자와 판별자의 손실함수는 각각 식 (5)와 식 (6)과 같다.

$$L_G = L_{percep} + \lambda_1 L_G^{Ra} + \lambda_2 L_{cG}^* + \lambda_3 L_1 \quad (5)$$

$$L_D = \lambda_1 L_D^{Ra} + \lambda_2 L_{cG}^* \quad (6)$$

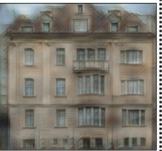
$$L_{cG}^* = \arg \min_G \max_D L_{cGAN}(G, D) \quad (7)$$

$$L_{cGAN}(G, D) = \mathbb{E}_{x_i, x_r} [\log D(x_i, x_r)] + \mathbb{E}_{x_i} [\log (1 - D(x_i, G(x_i)))] \quad (8)$$

식 (5)와 식 (6)에서 λ 는 다른 손실과의 균형을 맞추기 위한 계수이다. 식 (5)의 L_{percep} 은 생성 이미지와 원본 이미지의 VGG19 네트워크를 통과시킨 손실이고 L_{cG}^* 는 식 (7)과 식 (8)에서 알 수 있듯이 기존의 Conditional GAN의 손실함수이다. 그림 7은 생성자의 손실함수에 따른 시뮬레이션 결과를 비교한 것이다.

4. 시뮬레이션 결과 및 고찰

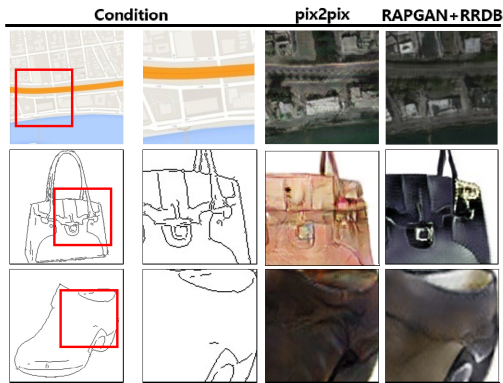
본 논문에서는 Jupyter Notebook과 Nvidia GTX 1070Ti GPU 환경에서 시뮬레이션을 수행하였다. <Table. 1>은 기존의 pix2pix 방법과 제안된 방법을 적용한 결과 이미지의 FID(Fréchet Inception Distance) 성능 비교를 나타낸 것이다. 제안된 방법은 16×16 과 32×32 의 패치 크기를 사용하는데, FID 측면에서 기존의 pix2pix에 비해 정량적으로 우수함을 확인할 수 있었다. [Fig. 8]은 시뮬레이션 결과를 국부적으로 확대한 결과 이미지이다. [Fig. 8]과 같이 제안된 방법을 사용해 생성한 결과가 기존의 pix2pix보다 더 실제 같은 이미지를 생성함을 볼 수 있다. [Fig. 9]는 여러 생성자와 GAN을 번갈아 사용하여 테스트한 시뮬레이션 결과로, 가장 우측에 있는 RRDB를 적용한 생성자와 RAPGAN을 같이 사용한 테스트의 시뮬레이션 결과가 가장 좋은 것을 볼 수 있다.

L_G	Proposed				
L_G^{Ra}	O	O	O	X	O
L_{cG}^*	O	O	X	O	O
L_1	O	X	O	O	O
L_{percep}	X	O	O	O	O
Simulation Results (Facades)					
					

[Fig. 7] Comparison of the loss functions in the proposed generator

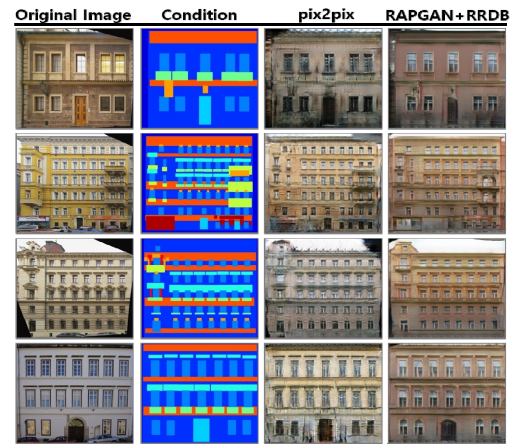
〈Table 1〉 Comparison of FID performance between the previous pix2pix and the proposed method

데이터셋	방법	가존의 pix2pix	RAPGAN+RRDB (patch: 16×16)	RAPGAN+RRDB (patch: 32×32)
Facades		11.38	10.72	9.77
Edge2Handbags		9.55	8.53	8.78
Edge2Shoes		10.13	8.05	8.20
Maps		12.41	10.38	10.23
Cityscapes		9.32	8.98	8.56



[Fig. 8] Simulation results using the previous pix2pix and proposed discriminators

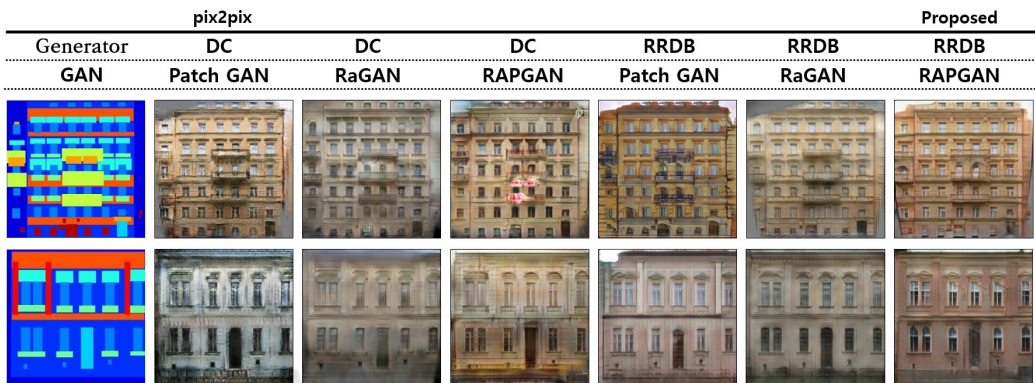
[Fig. 10]은 기존의 pix2pix 방법과 제안된 방법의 결과를 facades 데이터셋을 사용하여 비교한 것으로, 기존 pix2pix의 결과 이미지가 경계선이 모호하고 시각적 거부감을 주는 반면에 제안된 방법으로 생성된 결과는 우수한 품질의 이미지를 생성한다. RRDB를 적층한 구조를 통해 상대적으로 깊은 학습을 수행하고, 혹은 생성자를 사전 학습시켜 판별자가 조기에 학습되는 것을 억제하는 등의 제안된 방법의 조치가 효과가 있음을 시각적으로 확인할 수 있었다.



[Fig. 10] Simulation results of the previous pix2pix and proposed method using Facades datasets

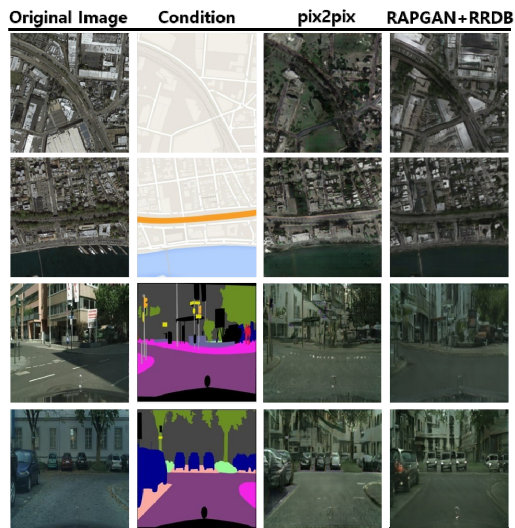


[Fig. 11] Simulation results of the previous pix2pix and the proposed method using Edges2photo



[Fig. 9] Simulation results of multiple generators and GAN

[Fig. 11]은 Edges2photo 데이터셋을 사용해 기존의 pix2pix 방법과 제안된 방법 결과를 비교한 것이다. [Fig. 11]에서 살펴볼 수 있듯이 제안된 방법은 주어진 정보가 부족한 신발의 코 부분과 가방 표면 등의 질감이나 명암 표현이 보다 더 현실적으로 구현된 것을 볼 수 있다. 또한 [Fig. 12]는 Maps와 Cityscapes 데이터셋을 활용해 결과영상을 비교한 것이다. 이전의 다른 데이터셋으로 비교한 결과와 같이 입력 이미지로 주는 조건의 정보가 부족한 경우에도 기존의 pix2pix와 달리 실사와 같은 결과 이미지를 만들어 내는 것을 확인할 수 있다.



[Fig. 12] Simulation results of the previous pix2pix and the proposed method using Maps and cityscapes datasets

5. 결론

본 논문에서는 기존의 pix2pix 성능을 개선하기 위해 RAPGAN과 RRDB를 이용한 개선된 네트워크 구조와 학습 방법을 제안하였다. 제안된 방법은 기존의 pix2pix 생성자보다 많은 층과 연결을 통해 생성자의 성능을 개선하였다. 또한 RAPGAN을 사용하여 진짜 이미지와 생성된 이미지 모두를 가울기에 포함시켜 더욱 안정적인 학습이 가능하게 하였다. 제안된 방법으로 생성한 개별 이미지마다의 성능은 기존의 이미지 변환 GAN보다 훨씬 안정적이지만, 생성되는 이미지들의 색채가 편향되는 한계점이 있다. 그럼에도 불구하고 제안된 방법이 저해상도의 이미지 데이터셋으로 높은 품질의 결과 이미지를

제공하는 것을 감안하여 추가적인 개선이 이뤄진다면, 제안된 네트워크 구조에 사용자가 원하는 조건을 추가하여 고품질의 영상 편집 및 이미지 생성이 가능할 것으로 기대된다.

REFERENCES

- [1] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative Adversarial Nets," Proceedings of Conference on Neural Information Processing Systems, 2014.
- [2] P. Isola, J. Zhu, T. Zhou, and A. A. Efros, "Image-to-Image Translation with Conditional Adversarial Networks," Proceedings of Conference on Computer Vision and Pattern Recognition, 2017.
- [3] M. Mirza and S. Osindero, "Conditional Generative Adversarial Nets," arXiv: 1411.1784, Nov. 2014.
- [4] T. Wang, M. Liu, J. Zhu, A. Tao, J. Kautz, and B. Catanzaro, "High-Resolution Image Synthesis and Semantic Manipulation with Conditional GANs," Proceedings of Conference on Computer Vision and Pattern Recognition, 2018.
- [5] X. Wang, K. Yu, S. Wum J. Gu, Y. Liu, C. Dong, Y. Qiao, and C. C. Loy, "ESRGAN: Enhanced Super-resolution Generative Adversarial Networks," Proceedings of European Conference on Computer Vision, pp.63-79, 2018.
- [6] D. Bashkurova, B. Usman, and K. Saenko, "Adversarial Self-Defense for Cycle-Consistent GANs," Proceedings of Conference on Neural Information Processing Systems, 2019.
- [7] T. Bepler, E. D. Zhong, K. Kelley, E. Brignole, and B. Berger, "Explicitly Disentangling Image Content from Translation and Rotation with Spatial-VAE," Proceedings of Conference on Neural Information Processing Systems, 2019.
- [8] Y. Alharbi, N. Smith, and P. Wonka, "Latent Filter Scaling for Multimodal Unsupervised Image-to-Image Translation," Proceedings of Conference on Computer Vision and Pattern Recognition, 2019.
- [9] Tycho F.A. van der Ouderaa and D. E. Worrall, "Reversible GANs for Memory-efficient Image-to-Image Translation," Proceedings of Conference on Computer Vision and Pattern Recognition, 2019.
- [10] K. Frans and C. Cheng, "Unsupervised Image to Sequence Translation with Canvas-Drawer Networks," arXiv preprint arXiv:1809.08340, 2018.
- [11] Z. Murez, S. Kolouri, D. Kriegman, R. Ramamoorthi, and K. Kim, "Image to Image Translation for Domain Adaptation," Proceedings of Conference on Computer

Vision and Pattern Recognition, 2018.

- [12] J. Zhu, R. Zhang, D. Pathak, T. Darrell, A. A. Efros, O. Wang, and E. Shechtman, "Toward Multimodal Image-to-Image Translation," Proceedings of Conference on Neural Information Processing Systems, 2017.
- [13] M. Liu, T. Breuel, and J. Kautz, "Unsupervised Image-to-Image Translation Networks," Proceedings of Conference on Neural Information Processing Systems, 2017.
- [14] B. Lim, S. Son, H. Kim, S. Nah, and K. M. Lee, "Enhanced Deep Residual Networks for Single Image Super-resolution," Proceedings of Conference on Computer Vision and Pattern Recognition, 2017.
- [15] Y. Zhang, Y. Tian, Y. Kong, B. Zhong, and Y. Fu, "Residual Dense Network for Image Super-resolution," Proceedings of Conference on Computer Vision and Pattern Recognition, 2018.
- [16] A. Jolicoeur-Martineau, "The Relativistic Discriminator: A Key Element Missing from Standard GAN," arXiv preprint arXiv:1807.00734, 2018.

윤 동 식(Dongsik Yoon)

[준회원]



- 2021년 2월 : 백석대학교 ICT학부 (공학사)
- 2023년 2월 : 고려대학교 대학원 전기전자공학과 (공학석사)

〈관심분야〉

딥러닝 기반 영상처리 및 컴퓨터비전, 딥러닝 기반 객체 분할, 초해상화 및 인페인팅, 스타일 트랜스퍼

곽 노 윤(Noyoon Kwak)

[종신회원]



- 1994년 2월 : 한국항공대학교 항공전자공학과 (공학사)
- 1996년 2월 : 한국항공대학교 대학원 항공전자공학과 (공학석사)
- 2000년 2월 : 한국항공대학교 대학원 항공전자공학과 (공학박사)

■ 2000년 3월 ~ 현재 : 백석대학교 컴퓨터공학부 교수

〈관심분야〉

딥러닝 기반 영상처리 및 컴퓨터비전, 얼굴 및 시선 인식, 객체 트래킹, 3D 재구성, 제스처 기반 UI/UX, 인공지능