

A Study on an AI-Based System for Hair Style and Color Synthesis

Hyeonwoo Choi[†] · Woonghee Lee^{††}

ABSTRACT

Recently, facial image synthesis technology has been actively applied in industries such as beauty and fashion. In particular, the increasing demand for personalized image generation has underscored the need for more sophisticated and natural synthesis techniques. However, existing commercial technologies often rely on predefined settings, limiting their ability to accommodate diverse user preferences. To address these limitations and offer users a wider range of customization options, this study proposes a hairstyle and color synthesis system. Specifically, a convolutional neural network is employed to naturally synthesize user-selected hairstyles, while accurate face parsing techniques are used to recognize and segment facial structures, enabling the seamless integration of the chosen hair color with the original image. Experimental results show that the proposed system achieves realistic synthesis without noticeable artifacts when compared to the original images. Furthermore, user evaluations yielded over 94% positive feedback, confirming the system's effectiveness.

Keywords : Face image synthesis, Hairstyle synthesis, Hair Color Synthesis, GAN, Semantic segmentation

인공지능 기반 머리 스타일 및 머리 색 합성 시스템 연구

최 현 우[†] · 이 웅 희^{††}

요 약

최근 얼굴 이미지 합성 기술은 뷰티 및 패션 산업 등에서 활발하게 활용되고 있다. 특히 사용자 맞춤형 이미지 생성에 대한 수요가 증가하면서, 더욱 정교하고 자연스러운 합성 기술의 필요성이 대두되고 있다. 그러나 기존의 상용화된 기술은 미리 정해진 설정값을 바탕으로 얼굴 이미지를 합성하므로 사용자의 다양한 요구사항을 충분히 수용하지 못하는 한계가 있다. 본 연구에서는 이러한 한계를 극복하고 사용자에게 보다 폭넓은 선택지를 제공하기 위해 머리 스타일과 머리 색 합성 시스템을 제안한다. 구체적으로, 합성곱 신경망을 활용하여 사용자가 선택한 머리 스타일을 자연스럽게 합성하고, 정교한 얼굴 이미지 파싱(parsing) 기술을 기반으로 얼굴의 세부 구조를 정확히 인식 및 분할하여 선택된 머리 색상과 이미지를 조화롭게 결합한다. 실험 결과, 원본 이미지와 비교했을 때 이질감 없는 자연스러운 합성이 이루어졌으며, 사용자 평가에서도 94% 이상의 긍정적인 피드백을 얻어 시스템의 우수성을 확인하였다.

키워드 : 얼굴 이미지 합성, 머리 스타일 합성, 머리 색 합성, GAN, 의미론적 분할

1. 서 론

최근 인공지능(Artificial Intelligence, AI) 기술은 우리 삶의 중요한 영역으로 자리 잡고 있다. 이러한 추세에 따라 AI를 활용하여 사진 속 인물의 외모를 개선하려는 시도가 증가하고 있다. 기존의 외모 개선을 위한 이미지 편집 작업은 주로 수작업으로 이루어졌으나, 최근에는 AI를 활용하여 보다 빠르고 간편하게 자연스러운 결과를 얻으려는 연구들이 활발히 진행

중이다[1].

그러나 현재 상용화된 외모 개선용 이미지 편집 기술은 몇 가지 한계점을 가지고 있어 사용자들의 다양한 요구를 모두 충족하기에는 아직 부족하다. 첫째, 이러한 시스템들은 미리 학습된 설정값을 기반으로 이미지를 편집하므로 사용자가 원하는 새로운 스타일이나 창의적이고 독창적인 변화를 구현하는 데 어려움이 있으며, 정해진 범위 내에서만 편집이 가능하다는 제약이 있다. 이로 인해 사용자가 창의적인 표현을 시도하거나 기존 패턴을 벗어나고자 할 때 제한사항으로 작용한다. 둘째, 기존 기술은 각 기능이 독립적으로 작동하고 서로 유기적으로 연동되지 않는 구조를 가지고 있어 사용자의 선택 옵션이 제한된다. 기능 간 상호작용이 부족해 사용자는 여러 단계를 거쳐야 하며, 이 과정에서 직관성과 사용성도 저하된

※ This research was financially supported by Hansung University.

[†] 비 회 원 : 한성대학교 AI융용학과 학사과정

^{††} 정 회 원 : 한성대학교 AI융용학과 조교수

Manuscript Received : June 30, 2025

First Revision : August 18, 2025

Accepted : September 8, 2025

* Corresponding Author : Woonghee Lee(whlee@hansung.ac.kr)

다. 이는 시스템의 효율성과 사용자 경험을 저하시킬 수 있다.

이러한 기존 연구의 한계를 극복하기 위해, 본 연구에서는 인물 이미지 편집 분야에서 특히 머리 스타일과 머리 색에 초점을 맞추어, 사용자가 원하는 머리 스타일 이미지와 얼굴 이미지를 결합하고 원하는 색으로 머리 색을 변경할 수 있는 시스템을 제안한다. 이를 위해 합성곱 신경망(Convolutional Neural Network, CNN)기반의 실시간 얼굴 파싱(Face Parsing) 모델인 BiSeNet[2]을 활용하여 특정 영역에 대한 색상 편집을 수행하였다. 또한 생성적 적대 신경망(Generative Adversarial Network, GAN) 기반의 생성 모델인 Style-Your-Hair[3]를 통해 사용자가 선택한 머리 스타일과 얼굴 이미지를 자연스럽게 합성하였다. 본 연구에서 활용하는 최신 이미지 합성 모델들은 대부분 GAN을 기반으로 하는데, GAN의 핵심 구성 요소로 이미지의 특징을 효과적으로 추출하는 CNN이 활용된다. 또한, 얼굴 파싱은 이미지의 모든 픽셀을 의미론적 단위로 분류하는 의미론적 분할(Semantic Segmentation) 기술을 얼굴 영역에 특화시킨 분야로, 머리카락, 피부 등 세부 요소를 정교하게 분리하는 데 사용된다[2]. 따라서 본 연구는 CNN을 근간으로 하는 GAN 모델과 얼굴에 특화된 의미론적 분할 기술을 활용하여 제안하는 시스템을 구현한다. 또한, 제안된 시스템을 실제 사진에 적용하여 결과를 평가하고, 사용자 의견을 수집하여 본 시스템이 실사용자들에게 긍정적인 평가를 받았음을 검증하였다.

본 논문의 구성은 다음과 같다. 섹션 2에서는 인물 이미지 편집과 관련된 기존의 다양한 연구들을 소개한다. 섹션 3에서는 제안한 시스템의 구성을 상세히 설명한다. 섹션 4에서는 제안된 시스템의 구현 및 평가 결과를 설명하며, 섹션 5에서는 결론을 기술한다.

2. 관련 연구

본 섹션에서는 제안하는 시스템의 기반이 되는 주요 기술들을 설명한다. 먼저, 이미지의 픽셀 단위 분할을 위한 Semantic Segmentation 기술과, 이를 얼굴 영역에 특화시킨 얼굴 파싱에 대해 알아본다. 다음으로, 사실적인 이미지를 생성하고 변형하는 데 사용되는 GAN 기술과 관련 연구들을 살펴본다.

고품질의 얼굴 이미지 생성 및 편집 연구는 주로 GAN을 기반으로 발전해왔으며, 특히 StyleGAN[4]과 그 후속 모델들은 이 분야의 기반 기술로 자리 잡았다[1,3,4]. 이 모델들은 매우 사실적인 이미지를 생성할 수 있었지만, 초기에는 생성된 이미지의 특정 속성을 사용자가 원하는 대로 정교하게 제어하는 데에는 한계가 있었다. 이러한 문제를 해결하기 위해, 모델의 잠재 공간(Latent space)을 최적화하여 사용자의 의도를 더 잘 반영하려는 연구들이 등장했다[5]. 예를 들어, 자세 변화에 관계없이 일관된 스타일을 적용하거나, 이미지의 특정

속성만을 정교하게 변경하는 기술들이 제안되었다[6]. 하지만 이러한 연구들 역시, 미리 학습된 데이터의 범주를 벗어나거나 사용자가 완전히 새로운 스타일을 창의적으로 적용하는 데에는 여전히 어려움이 있었다. 즉, 대부분 기술이 미리 정의된 스타일의 조합이나 보간에 머물러, 사용자의 창의적 요구를 충족시키기엔 부족했다[7].

머리 색상을 자연스럽게 변경하기 위해서는 이미지 내에서 머리카락 영역을 정확하게 분리하는 의미론적 분할(Semantic segmentation) 기술이 선행되어야 한다[8]. 최근에는 실시간 처리가 가능하면서도 높은 정확도를 유지하는 양방향 분할 네트워크(Bilateral Segmentation Network, BiSeNet)와 같은 효율적인 모델들이 제안되었다[9,10]. BiSeNet[2]은 CNN 기반 얼굴 파싱 모델의 대표적인 예시로 이미지의 세밀한 경계 정보와 문맥 정보를 동시에 활용하여 머리카락과 같은 복잡한 구조의 객체를 정교하게 분할할 수 있다. 그러나 이 연구들은 주로 특정 작업(분할, 색상 적용)의 정확도 향상에 집중해, 스타일 변화와 결합되는 통합적 경험을 제공하지는 못했다.

종합하면, 기존 연구들은 머리 스타일 합성[3,4,6,12]이나 색상 변경을 위한 영역 분할[2,8,9,10]이라는 개별 작업에서는 높은 기술적 완성도를 보였다. 그러나 대부분의 연구가 스타일 합성과 색상 편집이라는 두 가지 핵심 기능을 독립적으로 다루고 있어, 사용자가 이 둘을 하나의 시스템 안에서 유기적으로 결합하여 동시에 제어하는 통합적인 프레임워크에 대한 논의는 부족했다. 이로 인해 사용자는 여러 단계를 거쳐야 했으며, 기능 간 상호작용 부족으로 직관성과 사용성이 저하되는 문제가 있었다. 본 연구는 이 문제를 해결하기 위해 스타일 합성 모델과 정교한 분할 모델을 결합하여, 사용자가 머리 스타일과 색상을 동시에, 그리고 직관적으로 제어할 수 있는 통합 시스템을 제안함으로써 기존 연구의 한계를 극복하고자 한다.

3. 시스템 설계

본 섹션에서는 제안한 시스템 전반 및 세부 과정들에 대해 설명한다.

3.1 전체 시스템 구조

Fig. 1은 제안하는 시스템의 전체적인 구조를 나타낸다. 본 시스템은 서론에서 지적한 기능 간의 단절 문제를 해결하기 위해, 머리 스타일 합성과 색상 변경이라는 두 가지 핵심 기능을 하나의 유기적인 파이프라인으로 통합하였다. 사용자가 입력한 머리 스타일 및 얼굴 이미지는 먼저 Hair synthesis 단계를 거쳐 새로운 스타일이 적용된 합성 이미지로 생성된다. 이 부분은 Algorithm 1에서 line 2~3에 해당된다. 이 단계의 결과물은 독립적으로 존재하는 것이 아니라, 곧바로 다음 Face parsing 및 Coloring 단계의 필수적인 입력으로 전달된다. 즉,

게 정렬하여, 두 정보가 자연스럽게 결합되도록 조정한다. 이러한 과정을 통해 모델은 얼굴 디테일을 유지하며 머리 스타일 형태 변화를 자연스럽게 반영한다. 이 과정의 결과물은 Fig. 1의 "Final Output"으로 표시된, 1024x1024 해상도의 머리 스타일 합성 이미지이다. 이 과정은 Algorithm 2의 line 5~6에 해당된다.

〈Algorithm 2〉 Hairstyle Generation

```

Input:
* norm_face_img: 정규화된 얼굴 이미지
* norm_hair_img: 정규화된 머리 스타일 이미지
Output:
* synthesized_img: 머리 스타일이 합성된 이미지
---
1. SynthesizeHair( ):
// StyleGAN2의 W+ 잠재 공간으로 임베딩
2. W_src ← EmbedToLatent(norm_face_img)
   W_trg ← EmbedToLatent(norm_hair_img)
// 얼굴 자세 정렬 (Pose Alignment)
3. W_aligned ← AlignPose(source_latent=W_src*, target_latent=W_trg*)
// 잠재 벡터 블렌딩을 통한 스타일 합성
4. W_final ← BlendLatents(source_latent=W_src*, aligned_target=W_aligned*)
// 최종 잠재 벡터로부터 이미지 생성
5. synthesized_img ← GenerateImageFromLatent(W_final)
6. return synthesized_img
---
```

3.4 영역추출 및 색상 합성

머리 스타일 합성이 완료된 후, Fig. 1의 가운데 블록인 "Face parsing" 단계가 진행된다. 전체 과정은 Algorithm 3에 기술되어있다. 합성된 이미지는 BiSeNetv2 모델[9]을 사용하여 얼굴 파싱이 수행된다. 이 모델은 두 가지 주요 경로로 구성된다. 첫 번째는 "Backbone"을 통해 특징을 추출하여 "Booster"로 전달하는 경로이다. 두 번째는 "Aggregation Layer"를 통해 이 특징들을 결합하고, 최종적으로 얼굴 영역을 분할하는 "Seg Head"로 전달하는 경로이다.

이 "Seg Head"는 "Detail Branch"와 "Semantic Branch"로 구성된 출력을 생성한다. Fig. 1에서의 파란색 레이어들로 표시된 "Detail Branch"는 얼굴 윤곽선과 같은 저수준 세부 정보를 담당하며, 초록색 레이어들로 표시된 "Semantic Branch"는 얼굴의 주요 구성 요소를 분류하고 고수준 의미 정보를 제공한다. 이 두 브랜치의 출력을 결합하여 얼굴의 세부 구조와 의미 정보를 통합한 특징 표현을 생성한다. 최종적으로, 이 "Seg Head"의 출력은 얼굴의 다양한 영역을 색상으로 구분하여 나타내는 분할 마스크 이미지로 표시된다. 이처럼 BiSeNetv2를 통해 얼굴 전체의 분할 마스크를 생성하고, 그로부터 머리카락 영역만 추출하는 과정은 Algorithm 3의 line 2~3에 해당된다.

마지막으로, Fig. 1의 하단 블록인 "Coloring" 단계가 진행

된다. "Hair synthesis"를 통해 생성된 이미지와 "Face parsing"을 통해 얻은 분할 마스크를 사용하여 색상 합성이 이루어진다. 이 단계에서는 OpenCV를 사용하여 이미지의 색 공간을 HSV로 변환한 후, Hue 채널에 사용자가 원하는 색상을 적용하여 머리 색상만 변경한다. 이때 채도와 명도는 유지된다. 특히, 이미지의 자연스러움을 위해 명도를 낮춘 후 전체 이미지의 색 변환을 진행하며, 변환이 완료된 이미지에서 머리 영역만 추출하여 기본 이미지에 합성하는 방식으로 처리한다. 이 과정은 Algorithm 3의 line 4~7에 해당된다. 최종적으로 합성된 이미지는 원본 머리 영역에만 적용된다. Fig. 1에서 이 과정은 합성된 얼굴 이미지와 분할 마스크가 합쳐져서 최종적으로 머리 색상이 변경된 "Output" 이미지를 생성하는 것으로 나타난다. 이 최종 합성 과정은 Algorithm 3의 line 10에 명시되어 있다.

〈Algorithm 3〉 Hair Colorization

```

Input:
* synthesized_img: 스타일이 합성된 이미지
* target_color: 적용할 머리 색상
Output:
* colored_output_img: 최종 색상이 적용된 이미지
---
1. ApplyHairColor( ):
// BiSeNetv2를 이용한 얼굴 영역 분할
2. seg_mask ← ParseFace(synthesized_img)
3. hair_mask ← ExtractHairRegion(seg_mask) // 분할된 마스크에서 머리카락 영역 추출
// 색상 적용
4. hsv_img ← ConvertToHSV(synthesized_img)
5. for each pixel p in hsv_img:
6. if hair_mask at p is True:
7. hsv_img.Hue at p ← target_color.Hue
// 자연스러움을 위해 명도(Value)를 소폭 조정
8. hsv_img.Value at p ← hsv_img.Value at p × 0.9
9. rgb_img ← ConvertToRGB(hsv_img)
// 원본 이미지와 머리 영역만 합성하여 자연스러움 유지
10. colored_output_img ← Combine(original_img=synthesized_img*, colored_hair_img=rgb_img*, mask=hair_mask*)
11. return colored_output_img
---
```

4. 구현 및 성능 평가

본 섹션에서는 시스템의 구현 내용과 성능 평가 결과를 설명한다.

4.1 구현

1) 구현 환경

본 연구에서는 RTX3060 12GB GPU, 16GB RAM, AMD 5600X @3.7GHz CPU를 탑재한 기기를 사용하여 구현 및 실험을 수행하였다. 실험 환경에서 사용된 소프트웨어는

Python 3.11.9, PyTorch 2.3.0, OpenCV 4.10.0, CUDA 12.1 버전이다.

2) 전처리

본 연구에서는 효과적인 얼굴 인식을 위해 전처리 단계에서 OpenCV 라이브러리를 사용하였다. 사용자가 입력한 이미지는 Haarcascade 얼굴 검출기(haarcascade_frontalface_alt.xml)를 이용하여 얼굴 영역을 정확히 검출한다. 검출된 얼굴 영역의 중심을 기준으로 각 방향으로 512픽셀씩 확장하여, 총 1024×1024 픽셀 크기로 이미지를 잘라 정규화를 진행한다.

3) 머리 스타일 변경

머리 스타일 변경 모델 구현을 구현하기 위해 PyTorch 1.8.0, TorchVision 0.9.0, Torchaudio 0.8.0, cuDNN 11.1, face_alignment, face-recognition, 그리고 matplotlib 라이브러리가 활용되었다. 주요 구현 내용으로 pytorch를 중심으로 모델과 GPU 가속 연산을 처리하였으며 StyleGAN2를 기반으로 W+ 및 특징 공간에서 임베딩 변환, 스타일 손실과 키포인트 정렬 손실 최적화를 통한 정렬 및 블렌딩 과정이 포함되어있다. 정렬 과정에서는 얼굴과 머리 구조를 보정하고 블렌딩 과정에서는 합성된 이미지와 원본 이미지 간의 시각적 스타일 일관성을 유지하기 위해 그램 행렬 손실(Gram matrix loss)를 활용하였다. 결과적으로, Fig. 2와 같이 머리 스타일 이미지와 얼굴 이미지를 모델에 입력하면, 얼굴 이미지의 특징이 반영되고 머리 스타일 이미지가 합성된 이미지가 출력된다.



Fig. 2. Example results of hair synthesis

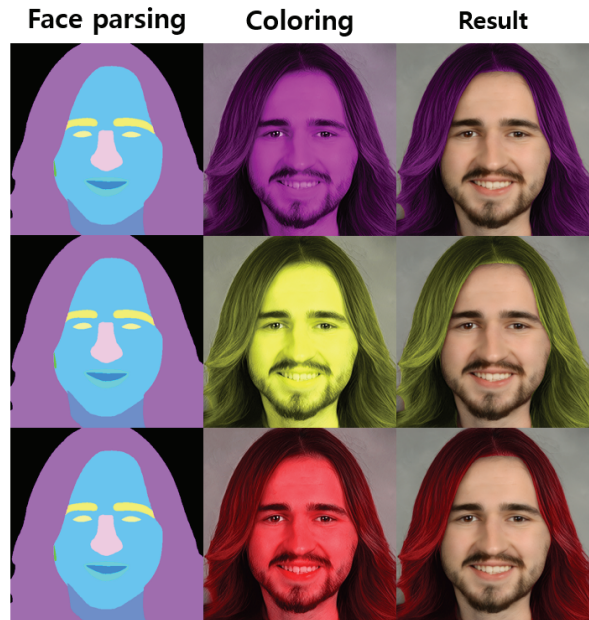


Fig. 3. Example results of coloring

4) 머리 색 변경

본 구현에서 머리 색 변경은 OpenCV, NumPy, Scikit-Image를 활용하여 진행된다. 입력 이미지는 OpenCV로 로드된 후, BiSeNet 모델(PyTorch 기반, CelebAMask-HQ 데이터셋으로 학습)을 통해 Fig. 3의 Face parsing과 같이 19개 클래스로 분할된다. 색상 변경은 Fig. 3의 Coloring에서처럼 HSV 색 공간 변환 및 RGB 조작을 통해 머리 영역에 사용자가 지정한 색상을 적용한다. 또한, Scikit-Image를 활용한 Gaussian Blur 및 샤프닝 효과를 통해 머리카락을 더욱 선명하게 표현하였다.

4.2 성능 평가

본 연구에서는 제안하는 시스템의 성능을 종합적으로 검증하기 위해, 36명의 참가자를 대상으로 정량적 지표 분석과 정성적 사용자 만족도 평가를 병행하였다. 평가의 객관성과 재현성을 확보하고자, 조명, 구도 등 통제되지 않는 변수가 많은 개인 사진 대신, 해당 분야의 표준 데이터셋인 CelebA-HQ (Celeb-Amateur High-Quality)[13]에서 인종, 성별, 머리 스타일의 다양성을 고려한 5개의 이미지를 선정하여 사용하였다.

먼저, 시스템의 성능을 객관적인 수치로 비교하기 위해 이미지 품질 평가 지표인 PSNR과 SSIM을 측정하여 다수의 상용 애플리케이션들과 비교하였다. 단, 상용 애플리케이션들은 지정된 프리셋 방식으로만 작동하여 머리 스타일 및 색상 선택이 한정되는 한계가 있었다. 따라서 본 평가에서 상용 애플리케이션은 얼굴 이미지만을 입력으로 사용하였고, 제안 시스템은 동일한 얼굴 이미지에 머리 이미지를 추가로 입력하여 처리하는 조건으로 진행되었으며, 그 결과는 Table 1과 같다.

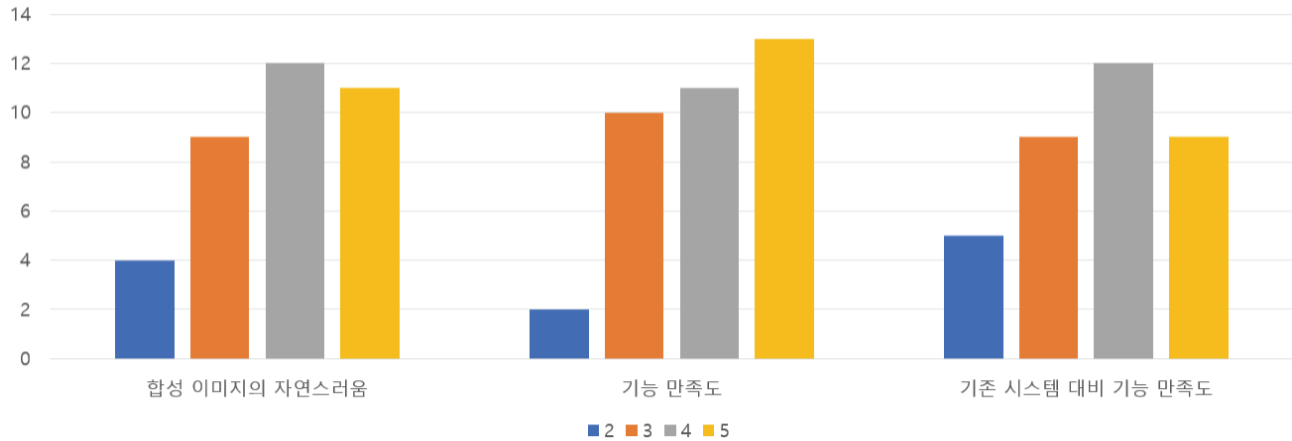


Fig. 4. Results of user satisfaction evaluation

Table 1. Quantitative Performance Comparison of the Proposed System and Commercial Applications

Model	PSNR	SSIM
Our System	11.87 dB	0.5192
PhotoDirector	15.64 dB	0.7381
YouCam	13.07 dB	0.5812
EPIK	19.22 dB	0.8162
SNOW	18.91 dB	0.7037

Table 1에서, 제안 시스템은 PSNR 및 SSIM 수치가 기존 상용 애플리케이션보다 낮게 기록되었다. 그러나 이는 제안 시스템의 성능적 결함이 아닌, 평가 지표가 가진 본질적 한계와 시스템이 추구하는 목표 간의 근본적인 불일치에서 기인한다. PSNR과 SSIM은 원본과 결과물 간의 픽셀 단위 유사도를 측정하는 지표로, 정답이 정해진 복원 작업에 최적화되어 있다. 상용 애플리케이션의 높은 수치는 오히려 그들의 기능적 한계를 시사한다. 이 애플리케이션들은 개발사가 사전에 정의하고 최적화한 수십 개의 프리셋 내에서만 작동하는 닫힌 시스템이다. 이 구조는 프리셋 적용 시 예측 가능한 결과를 내므로 높은 수치 기록에 유리하지만, 선택 옵션에 없는 새로운 스타일 생성은 원천적으로 불가능하다. 이는 사용자의 창의성을 제한하고, 결과물의 다양성을 저해하는 명백한 단점으로 작용한다.

반면, 본 연구가 제안하는 시스템은 사용자가 입력하는 임의의 이미지를 기반으로 새로운 스타일을 실시간으로 합성하는 열린 시스템으로서, 상용 앱이 제공하지 못하는 독창적인 기능적 우위성을 확보하였다. 예를 들어, 특정 인물의 독창적인 머리 스타일을 자신의 사진에 적용하고자 할 때, 상용 앱은 해당 스타일과 유사한 프리셋이 존재하지 않는 한 이를 구현할 수 없다. 하지만 제안 시스템은 해당 이미지를 직접 입력하여 스타일을 합성함으로써, 높은 자율성을 사용자에게 제공한다. 이처럼 본 시스템은 단순한 필터 적용을 넘어 새로운 결과

물을 창조하는 것을 목표로 하므로, 원본과의 단순 비교를 전제하는 정량적 지표는 그 핵심 가치를 온전히 평가하기에 부적합하다.

이러한 정량적 지표의 한계를 보완하고 시스템의 실질적 효용성을 입증하기 위해, 36명의 참가자를 대상으로 사용자 만족도 조사를 수행하였다. 참가자들은 기능 설명을 읽은 후, 무작위로 제시된 합성 이미지 5장을 보고 3개 항목(결과 자연스러움, 기능 만족도, 기존 시스템 대비 개선 여부)을 5점 척도로 평가했다. 그 결과는 Fig. 4와 같다. 전반적으로 모든 항목에서 긍정적인 평가가 주를 이루었으며, 특히 기능 만족도 항목에서 5점(매우 만족)을 선택한 응답자가 13명으로 가장 높은 평가를 받았다. 합성 이미지의 자연스러움 항목과 기존 시스템 대비 기능 만족도 항목에서도 각각 21명의 참가자가 4점(만족) 이상으로 평가하며 높은 만족도를 보였다. 구체적으로, 3점 이상의 긍정적인 평가를 한 참가자의 비율은 기능 만족도에서 94% (34명), 기존 시스템 대비 기능 만족도에서

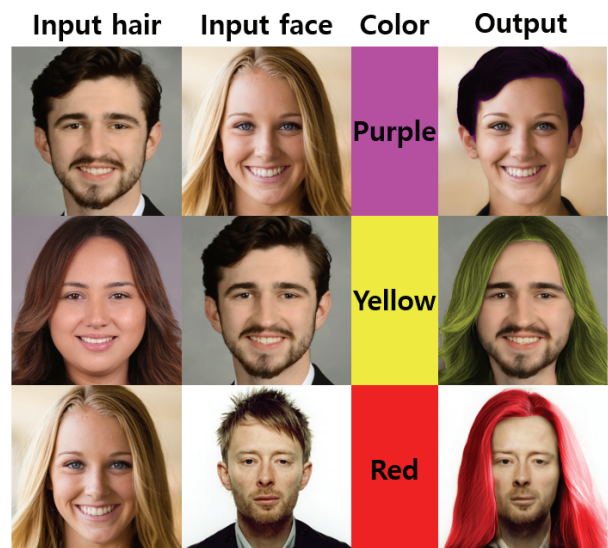


Fig. 5. Example results of the proposed system

86% (31명), 합성 이미지의 자연스러움에서 89% (32명)에 달했다. 특히, Fig. 5에서 볼 수 있듯이 색상에 구애받지 않고 다양한 머리 스타일이 자연스럽게 적용되는 점이 두드러졌다. 결론적으로, 제안 시스템의 가치는 수치적 비교를 넘어, 사용자에게 자율성을 부여했다는 정성적 평가를 통해 입증된다. 이러한 결과는 제안 시스템이 기존 상용 시스템의 한계를 극복하고, 사용자에게 높은 기능적 만족도와 자율성을 제공함을 보여준다.

5. 결 론

본 연구는 기존 얼굴 이미지 편집 시스템이 가진 두 가지 주요 한계, 즉 미리 정의된 설정값으로 인한 사용자의 창의성 제한과 스타일 합성과 색상 변경 기능의 단절로 인한 비유기적인 사용자 경험을 극복하는 것을 목표로 하였다. 이를 위해, StyleGAN2와 BiSeNetv2 모델을 기반으로, 두 핵심 기능을 하나의 유기적인 파이프라인으로 통합한 새로운 시스템을 제안하였다. 제안하는 시스템은 사용자가 원하는 어떤 얼굴과 머리 스타일 이미지도 실시간으로 조합하고, 그 결과물에 자유롭게 색상을 적용할 수 있는 높은 자유도를 제공한다. 사용자 만족도 조사 결과, 시스템의 핵심 기능은 94%의 긍정적 평가를 획득하여 제안된 방법론의 높은 기능적 효용성을 실증적으로 보여주었다. 또한, 생성된 이미지의 시각적 자연스러움에 대한 긍정적 평가는 89%에 달해, 이는 합성 결과물의 품질이 사용자의 수용 기준을 충족함을 의미한다. 더 나아가, 기존 상용 시스템 대비 경험의 개선도를 묻는 항목에서 86%의 사용자가 향상되었다고 응답하여, 본 연구가 제안하는 통합적 파이프라인이 기존 기술의 한계를 성공적으로 극복했음을 보여주었다.

본 연구는 스타일과 색상 편집을 유기적으로 통합한 프레임워크를 제안했다는 점에서 학술적 기여를 가진다. 향후에는 본 시스템을 표정, 액세서리 등 다양한 얼굴 구성 요소로 확장하여 사용자에게 폭넓은 선택지를 제공할 계획이다.

References

- [1] T. Karras, M. Aittala, J. Hellsten, S. Laine, J. Lehtinen, and T. Aila, "Training generative adversarial networks with limited data," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020.
- [2] C. Yu, J. Wang, C. Peng, C. Gao, G. Yu, and N. Sang, "BiSeNet: Bilateral segmentation network for real-time semantic segmentation," in *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018.
- [3] T. Karras, S. Laine, M. Aittala, J. Hellsten, J. Lehtinen, and T. Aila, "Analyzing and improving the image quality of StyleGAN," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020.
- [4] T. Karras, S. Laine, and T. Aila, "A style-based generator architecture for generative adversarial networks," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019.
- [5] H. Härkönen, A. Härkönen, J. Lehtinen, and J. Y. A. Karras, "GANSpace: Discovering interpretable GAN controls," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- [6] T. Kim, C. Chung, Y. Kim, S. Park, K. Kim, and J. Choo, "Style your hair: Latent optimization for pose-invariant hairstyle transfer via local-style-aware hair alignment," *arXiv preprint*, arXiv:2208.07765, 2022.
- [7] X. Pan, A. Tewari, T. Leimkühler, L. Liu, A. Meka, and C. Theobalt, "Drag your GAN: Interactive point-based manipulation of the generative image manifold," in *ACM SIGGRAPH 2023 Conference Proceedings*, 2023.
- [8] L. C. Chen, Y. Zhu, G. Papandreou, F. Schroff, and H. Adam, "Encoder-decoder with atrous separable convolution for semantic image segmentation," in *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018.
- [9] C. Yu, C. Gao, J. Wang, G. Yu, C. Shen, and N. Sang, "BiSeNet V2: Bilateral network with guided aggregation for real-time semantic segmentation," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020.
- [10] D. Yaman, H. K. Ekenel, and A. Waibel, "Alpha matte generation from single input for portrait matting," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 2022.
- [11] P. Viola and M. J. Jones, "Rapid object detection using a boosted cascade of simple features," in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR)*, 2001.
- [12] Taeu, Style-Your-Hair [Internet], <https://github.com/Taeu/Style-Your-Hair>.
- [13] CelebAMask-HQ dataset [Internet], <https://github.com/switchablenorms/CelebAMask-HQ?tab=readme-ov-file>.
- [14] C. H. Lee, Z. Liu, L. Wu, and P. Luo, "MaskGAN: Towards diverse and interactive facial image manipulation," in *Proceedings of the AAAI Conference on Artificial Intelligence (AAAI)*, 2020.
- [15] OpenCV-Python [Internet], <https://github.com/opencv/opencv-python>.



최 현 우

<https://orcid.org/0009-0004-6826-4149>
e-mail : apinball91@hansung.ac.kr
2021년 3월~현재 한성대학교 사이버보안
트랙, AI응용학과 학사과정
관심분야 : AI, Deep/Machine learning,
Generative AI, Computer
Vision



이 응 희

<https://orcid.org/0000-0003-0856-6415>
e-mail : whlee@hansung.ac.kr
2013년 2월 고려대학교 전기전자전파공학부
학사
2019년 2월 고려대학교 전기전자공학과 박사
2019년 5월 LG유플러스 선임 연구원
2019년 8월~2021년 2월 삼성전자 Staff Engineer
2021년 3월~현재 한성대학교 AI응용학과 조교수
관심분야 : Internet of Things, AI, Deep/Machine learning