

생성형 AI 디자인의 사용자 신뢰 형성: 심미적 평가와 출처 고지의 역할

백 은 지¹ · 조 동 민^{2*}

¹전북대학교 산업디자인학과 석사과정

²전북대학교 산업디자인학과 교수

User Trust Formation in Generative AI Design: Roles of Aesthetic Evaluation and AI Disclosure

Eun-Ji Baek¹ · Dong-Min Cho^{2*}

¹Master's Course, Department of Industrial Design, Jeonbuk National University, Jeonju 54896, Korea

²Professor, Department of Industrial Design, Jeonbuk National University, Jeonju 54896, Korea

[요 약]

생성형 AI의 활용이 디자인 분야로 확대되면서 AI 디자인 결과물에 대한 사용자 신뢰 형성이 중요해지고 있다. 본 연구는 생성형 AI 디자인 환경에서 디자인 요소 지각이 사용자 신뢰에 미치는 영향과 심미적 평가의 매개효과, 출처 고지 여부의 조절효과를 검증하였다. 동일한 생성형 AI 포스터를 활용하여 출처 고지 집단과 미고지 집단을 대상으로 온라인 실험을 실시하였으며, 총 150명의 자료를 SPSS로 분석하고, 조절된 매개효과 검증을 위해 PROCESS Macro Model 14를 적용하였다. 분석 결과, 디자인 요소 지각은 심미적 평가에 유의한 정(+)의 영향을 미쳤으며, 심미적 평가는 사용자 신뢰에 유의한 정(+)의 영향을 미쳤다. 또한 디자인 요소 지각이 사용자 신뢰에 미치는 직접효과와 심미적 평가를 통한 간접효과가 모두 유의하여 부분매개가 확인되었다. 반면, 출처 고지 집단과 미고지 집단 간 사용자 신뢰의 차이, 출처 고지 여부의 조절효과 및 조절된 매개효과는 유의하지 않았다. 본 연구는 생성형 AI 디자인에서 디자인 요소에 대한 지각과 심미적 평가가 사용자 신뢰 형성에 중요한 역할을 한다는 점을 실증적으로 확인하였다.

[Abstract]

Given the extensive use of generative AI in design, user trust in AI-generated design outcomes has become increasingly important. This study examined the effects of perception of design elements on user trust, with a focus on the mediating role of aesthetic evaluation and moderating role of AI disclosure. In an online between-subjects experiment, 150 participants viewed the same AI-generated poster with or without AI disclosure. The data were analyzed using SPSS, and PROCESS Macro Model 14 was employed in testing the moderated mediation effect. The results showed that perception of design elements had a significant positive effect on aesthetic evaluation and aesthetic evaluation had a significant positive effect on user trust. Both direct and indirect effects were significant, indicating partial mediation. By contrast, the difference in user trust between AI disclosure and non-disclosure groups, moderating effect of AI disclosure, and moderated mediation effect were not significant. Thus, perception of design elements and aesthetic evaluation play important roles in user trust formation in generative AI design.

색인어 : 생성형 AI 디자인, 출처 고지, 사용자 신뢰, 디자인 요소 지각, 심미적 평가

Keyword : Generative AI Design, AI Disclosure, User Trust, Perception of Design Elements, Aesthetic Evaluation

<http://dx.doi.org/10.9728/dcs.2026.27.8.2253>



This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License (<http://creativecommons.org/licenses/by-nc/3.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

Received 13 July 2026; **Revised** 31 July 2026

Accepted 11 August 2026

***Corresponding Author; Dong-Min Cho**

Tel: +82-63-270-3752

E-mail: mellgipson@daum.net

I. 서 론

생성형 AI 기술의 발전은 시각디자인 분야의 콘텐츠 제작 방식에도 큰 변화를 가져오고 있다. Midjourney, Adobe Firefly, DALL·E와 같은 생성형 AI 기반 이미지 생성 도구는 브랜드 콘텐츠, 광고 시각물, SNS 콘텐츠, 전시회 포스터 등 다양한 시각 커뮤니케이션 분야에서 활용 범위를 빠르게 확대하고 있으며 [1]~[3], 도구별 기술적 특성과 사용자 경험에 따라 활용 양상에도 차이가 나타나는 것으로 보고된다[3]. 이러한 흐름 속에서 생성형 AI는 기존의 디자인 제작 과정을 효율화하는 동시에 새로운 창작 방식으로 자리 잡고 있으며, 사용자는 일상적으로 생성형 AI가 제작한 디자인 결과물을 접하고 있다.

그러나 생성형 AI 기술이 고도화되면서 사용자가 디자인 결과물의 제작 주체를 명확히 인식하기 어려운 사례가 증가하고 있으며, 이는 결과물에 대한 평가와 신뢰 판단의 새로운 쟁점으로 대두되고 있다. 이러한 환경에서는 디자인 결과물 자체의 품질뿐 아니라 누가 제작하였는가에 대한 정보가 결과물에 대한 평가와 신뢰 형성에 영향을 미칠 가능성이 제기된다. 이에 따라 생성형 AI가 제작한 사실을 사용자에게 명시하는 출처 고지(AI disclosure)가 사용자 인식과 신뢰 판단에 어떠한 영향을 미치는지에 대한 연구의 필요성이 커지고 있다.

선행연구에서는 AI 생성 사실의 고지가 창작 결과물에 대한 이용자의 평가에 영향을 미칠 수 있음을 보고하였으며[4], 인간 제작 결과물과의 선호도 비교[5], AI 서비스에 대한 신뢰와 지속사용의도[6] 등에 관한 연구도 활발히 이루어지고 있다. 국내에서도 AI 생성 콘텐츠의 가치 유형과 제시 형식에 따라 출처 공개에 대한 소비자 인식이 다르게 나타남이 보고된 바 있다[7]. 그러나 대부분의 연구는 AI와 인간 제작물을 비교하거나 AI 서비스 자체에 대한 수용 태도를 중심으로 수행되었으며, 동일한 생성형 AI 디자인을 대상으로 출처 고지 여부만을 조작하여 사용자 신뢰 형성 과정을 분석한 연구는 상대적으로 부족한 실정이다. 특히 시각디자인 분야에서 디자인 요소 지각, 심미적 평가 및 사용자 신뢰 간의 관계를 통합적으로 설명하고, 출처 고지의 조절효과를 함께 검증한 연구는 제한적으로 이루어져 왔다. 따라서 본 연구는 동일한 생성형 AI 디자인 결과물을 대상으로 출처 고지 여부만을 실험적으로 조작하여 디자인 요소 지각, 심미적 평가 및 사용자 신뢰의 형성 과정을 통합적으로 검증하였다.

Ngo 등[8]은 균형, 대칭, 통일성 및 단순성 등 화면 구성의 시각적 질서를 정량적으로 평가하는 인터페이스 심미성 모형을 제시하였다. 이러한 논의는 디자인의 시각적 구성과 심미적 평가 간 관계를 검토하는 데 이론적 참고점을 제공하지만, 색상·레이아웃·타이포그래피에 대한 사용자의 지각이 심미적 평가와 사용자 신뢰로 이어지는 경로를 직접 제시한 것은 아니다. 따라서 본 연구는 생성형 AI 디자인 환경에서 디자인 요소 지각, 심미적 평가 및 사용자 신뢰 간의 관계를 별도로 실증 분석하고, 출처 고지 여부가 이러한 관계에 어떠한

영향을 미치는지를 확인하고자 하였다.

이에 본 연구는 동일한 생성형 AI 디자인 자극물을 제시하고 출처 고지 여부만을 조작하는 단일요인 피험자 간 (between-subjects) 실험설계를 적용하여 디자인 요소 지각이 심미적 평가를 거쳐 사용자 신뢰에 미치는 영향을 검증하였다. 또한 출처 고지 여부가 심미적 평가와 사용자 신뢰 간의 관계를 조절하는지를 분석하였다. 이를 바탕으로 생성형 AI 디자인의 사용자 신뢰 형성 과정을 살펴보고자 한다.

II. 이론적 배경

2-1 생성형 AI 디자인

생성형 AI 디자인은 텍스트 프롬프트를 입력으로 받아 이미지, 형상, 색상 조합 등 시각 결과물을 자동으로 산출하는 AI 기반 디자인 방식을 의미한다. 박인평·한정엽[2]은 Midjourney, Stable Diffusion, Adobe Firefly, DALL·E 등의 이미지 생성 AI 도구를 유형별로 분류하고, 각 도구의 생성 방식과 이미지 표현 특성을 분석하였다. 전수진·김승인[3]은 이 가운데 Midjourney, DALL·E, Adobe Firefly의 기술적 특성과 창의적 활용 가능성을 사용자 경험을 중심으로 비교하여, 도구별로 생성 품질과 활용 양상에 차이가 있음을 보고하였다. 이들 도구는 대규모 이미지-텍스트 데이터를 학습한 확산 모형(diffusion model) 기반으로[9], 사용자의 프롬프트 입력만으로 상업적 수준의 시각 결과물을 산출할 수 있다는 공통적 특징을 지닌다.

생성형 AI 디자인의 실무적 활용 범위는 브랜드 콘텐츠, 광고 시각물, SNS 이미지, 전시 포스터 등으로 확대되고 있다 [2],[3]. 안성민[1]은 디자이너를 대상으로 생성형 디자인 프로그램에 대한 수용 요인을 분석하고, 인지된 유용성과 사용 용이성이 수용 의도에 유의한 영향을 미침을 보고하였다. 이는 생성형 AI 디자인이 실무에서 활용 가능한 설계 도구로 기능할 수 있음을 보여준다.

관련 선행연구는 크게 두 방향으로 구분된다. 첫째, Bellaiche 등[5]과 같이 AI 생성 결과물과 인간 제작 결과물 간의 선호도 차이를 비교한 연구이며, 둘째, 김소연 외[6]와 같이 생성형 AI 서비스 전반에 대한 신뢰와 수용 의도를 분석한 연구이다. 그러나 AI 생성 디자인 결과물을 분석 단위로 설정하고, 색상·레이아웃·타이포그래피와 같은 개별 디자인 요소 지각이 심미적 평가를 거쳐 신뢰 판단에 이르는 구조적 경로를 검증한 연구는 충분히 이루어지지 않았다. 특히 출처 고지 여부를 함께 고려하여 이러한 관계를 통합적으로 검증한 연구는 더욱 제한적이다.

2-2 디자인 요소 지각과 심미적 평가

시각 디자인의 구성 요소는 색상(color), 레이아웃(layout),

타이포그래피(typography)로 구분된다. 색상은 색조, 채도, 명도의 조합을 통해 시각적 인상과 정서 반응을 형성하며, 레이아웃은 구성 요소의 공간적 배치와 위계 구조를 통해 정보 전달의 효율성과 균형감을 결정하는 역할을 한다. 타이포그래피는 서체 선택, 크기, 행간, 자간 등의 속성을 통해 가독성과 디자인 전체의 시각적 통일감에 영향을 미친다.

한편 Tractinsky, Katz, Ikar[10]는 현금자동입출금기 인터페이스를 대상으로 사용자가 지각한 심미성과 지각된 사용성 간에 높은 관련성이 나타나며, 이러한 관계가 실제 사용 이후에도 유지됨을 확인하였다. 이는 심미적 평가가 대상의 다른 속성에 대한 후속 판단과 관련될 수 있음을 보여준다. 다만 해당 연구는 심미성과 사용성 간의 관계를 검증한 것으로, 사용자 신뢰에 미치는 영향을 직접 검증한 것은 아니다.

Lavie와 Tractinsky[11]는 웹사이트에 대한 지각된 시각적 심미성 척도를 개발하고, 심미성이 정돈성과 명료성을 중심으로 하는 고전적 심미성과 창의성 및 독창성을 중심으로 하는 표현적 심미성으로 구분됨을 제시하였다. 이는 심미적 평가가 개별 시각요소에 대한 평가와 구분되는 대상 전반의 지각적 판단으로 측정될 수 있음을 보여준다. 본 연구는 이러한 논의를 바탕으로 심미적 평가를 생성형 AI 디자인 결과물에 대한 전반적인 미적 판단으로 정의하였다. 본 연구에서는 이러한 논의를 생성형 AI 디자인의 신뢰 판단 맥락에 적용하여 심미적 평가와 사용자 신뢰 간의 관계를 검증하였다.

Ngo, Teo, Byrne[8]은 스크린 레이아웃의 심미성을 정량화하기 위하여 균형, 대칭, 응집성, 통일성, 비례, 단순성 등 14개의 수학적 측정지표를 제안하였다. 해당 모형은 화면 구성의 시각적 질서와 복잡성을 평가하기 위한 것으로, 색상·레이아웃·타이포그래피에 대한 지각이 심미적 평가로 이어지는 경로를 직접 제시한 것은 아니다.

본 연구에서는 Ngo 등[8]의 모형이 제시한 시각적 구성과 질서에 관한 논의를 참고하되, 생성형 AI 디자인의 특성을 고려하여 색상, 레이아웃 및 타이포그래피에 대한 사용자의 지각을 통합적인 디자인 요소 지각으로 재구성하여 적용하였다. 따라서 본 연구의 디자인 요소 지각은 Ngo 등의 원모형을 그대로 적용한 것이 아니라, 생성형 AI 디자인 맥락에 맞게 확장하여 구성한 개념이다.

심미적으로 긍정적인 평가를 받은 디자인은 전문성, 완성도 및 신뢰성에 대한 긍정적인 인식을 형성할 수 있다. 따라서 심미적 평가는 생성형 AI 디자인에 대한 사용자 신뢰 형성과 관련된 중요한 선행요인으로 볼 수 있다. 그러나 이러한 심미적 평가가 항상 동일한 수준의 신뢰로 이어지는 것은 아니며, AI 출처 고지와 같은 출처 정보가 심미적 평가와 사용자 신뢰 간의 관계에 영향을 미칠 가능성이 있다.

2-3 AI 출처 고지와 사용자 신뢰

AI 출처 고지(AI disclosure)는 콘텐츠나 디자인 결과물이 AI에 의해 생성되었음을 사용자에게 명시적으로 알리는 행위

를 의미한다. 이는 단순한 정보 제공을 넘어, 사용자가 결과물을 해석하고 판단하는 인지적 맥락에 영향을 미칠 수 있는 요인으로 이해된다. Raj 등[4]은 AI 생성 사실의 고지가 창작 결과물에 대한 평가를 저하시킬 수 있으며, 이러한 효과가 진정성 지각을 통해 나타남을 보고하였다.

본 연구에서 사용자 신뢰는 생성형 AI 디자인 결과물을 신뢰할 수 있다고 판단하는 정도를 의미하며, 해당 결과물의 제작 주체에 대한 신뢰를 포함한다. 김소연 외[6]는 생성형 AI 서비스의 신뢰도에 영향을 미치는 요인을 탐색적으로 분석하였으며, AI 생성 결과물에 대한 신뢰는 서비스 수용 및 지속 사용 의도와 직접적으로 연결된다고 보고하였다.

AI 출처 고지와 사용자 신뢰의 관계를 구체적으로 분석한 연구는 최근 점차 축적되고 있다. Schilke와 Reimann[12]은 AI를 활용한 행위자에 대한 신뢰가 AI 사용 사실의 공개 이후 저하될 수 있음을 보고하였다. 다만 해당 연구는 디자인 결과물 자체에 대한 신뢰가 아니라 AI를 사용한 행위자에 대한 신뢰를 대상으로 했다는 점에서 본 연구와 차이가 있다. 반면 국내에서는 AI 생성 콘텐츠의 가치 유형과 형식에 따라 출처 공개에 대한 소비자 인식이 달라질 수 있음이 보고되었다[7]. 이러한 결과는 출처 고지의 효과가 콘텐츠의 특성과 연구 맥락에 따라 달라질 가능성을 보여준다. 이에 본 연구에서는 출처 고지 효과의 방향을 특정하지 않고 조절효과를 검증하였다. 그러나 기존 연구는 출처 고지의 영향을 신뢰, 선호도, 평가 등의 종속변수와 직접 연결하는 데 주로 초점을 두었으며, 심미적 평가가 사용자 신뢰로 이어지는 과정에서 출처 고지가 어떠한 역할을 하는지에 대해서는 충분히 검증되지 않았다.

따라서 본 연구에서는 출처 고지 여부를 집단 구분을 위한 실험요인으로 조작하는 동시에, 심미적 평가와 사용자 신뢰 간 관계를 조절하는 변수로 설정하였다. 즉, 출처 고지 여부는 디자인 자체의 시각적 특성을 변화시키기보다는, 심미적으로 평가된 결과물을 신뢰로 해석하는 사용자의 인지적 판단 과정에 영향을 미치는 조절요인으로 기능할 가능성이 있다.

III. 연구모형 및 연구가설

3-1 연구모형

본 연구는 생성형 AI 디자인에 대한 사용자 신뢰 형성 과정을 디자인 요소 지각, 심미적 평가 및 출처 고지의 관계를 중심으로 설명하고자 하였으며, 연구모형은 그림 1과 같다. Ngo 등[8]은 균형, 대칭, 통일성 및 단순성 등 화면 구성의 시각적 질서를 바탕으로 인터페이스 심미성을 평가하는 모형을 제시하였다. 본 연구는 이러한 시각적 구성과 심미성에 관한 논의를 참고하여, 생성형 AI 디자인에서 색상, 레이아웃 및 타이포그래피에 대한 통합적 지각이 심미적 평가 및 사용

자 신뢰와 어떠한 관계를 형성하는지를 실증적으로 검증하고자 하였다.

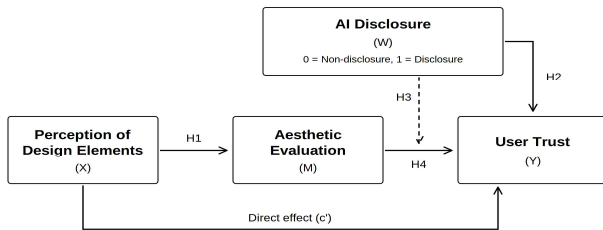


그림 1. 조절된 매개효과 연구모형

Fig. 1. Moderated mediation research model

본 연구에서는 생성형 AI 디자인의 출처 고지 효과를 검증하기 위하여 단일요인 피험자 간(between-subjects) 설계를 적용하였다.

선행연구에서는 AI 생성 사실의 고지가 이용자의 평가와 사용자 신뢰에 영향을 미칠 수 있음이 보고되었다[4],[12]. 이에 출처 고지 여부에 따라 생성형 AI 디자인 결과물에 대한 사용자 신뢰 수준에 직접적인 차이가 나타날 가능성을 검증하고자 하였다. 동일한 디자인 결과물이라 하더라도 출처가 고지되었는지 여부에 따라 심미적 평가가 사용자 신뢰로 이어지는 정도는 달라질 수 있다. 이에 본 연구에서는 출처 고지 여부를 심미적 평가와 사용자 신뢰 간의 관계를 조절하는 변수로 설정하였다.

이에 따라 본 연구모형은 디자인 요소 지각을 독립변수(X), 심미적 평가를 매개변수(M), 사용자 신뢰를 종속변수(Y)로 설정하고, 출처 고지 여부를 심미적 평가와 사용자 신뢰 간의 관계를 조절하는 변수(W)로 설정한 조절된 매개모형(moderated mediation model)으로 구성하였다.

3-2 연구가설

H1. 디자인 요소 지각은 심미적 평가에 정(+)의 영향을 미칠 것이다.

H2. 출처 고지 집단과 미고지 집단은 사용자 신뢰 수준에 차이가 있을 것이다.

H3. 출처 고지 여부는 심미적 평가와 사용자 신뢰 간의 관계를 조절할 것이다.

H4. 심미적 평가는 사용자 신뢰에 정(+)의 영향을 미칠 것이다.

IV. 연구방법

4-1 연구설계

본 연구는 생성형 AI 디자인의 출처 고지 여부가 사용자

신뢰 형성 과정에서 수행하는 조절적 역할을 검증하기 위하여 단일요인 두 집단 피험자 간(between-subjects) 설계를 채택하였다. 연구모형에서 디자인 요소 지각은 독립변수(X), 심미적 평가는 매개변수(M), 사용자 신뢰는 종속변수(Y), 출처 고지 여부는 심미적 평가와 사용자 신뢰 간 관계를 조절하는 변수(W)로 설정하였다.

실험 자극물은 Adobe Firefly를 활용하여 제작한 생성형 AI 기반 문화예술 전시 포스터로 구성하였다. 모든 참여자에게 동일한 디자인의 포스터를 제시하되, 고지 집단에는 생성형 AI로 제작되었음을 알리는 문구를 포스터 상단에 추가하였고, 미고지 집단에는 해당 문구를 제시하지 않았다. 따라서 두 조건 간 차이는 생성형 AI 제작 사실을 알리는 고지 문구의 유무로 한정하였으며, 그 외의 이미지, 텍스트, 레이아웃 및 규격은 동일하게 유지하였다. 이를 통해 자극물의 시각적 차이에서 발생할 수 있는 혼입효과를 최소화하였다. 실험에 사용된 자극물은 그림 2와 같다.



*The original Korean disclosure statement was retained to preserve the integrity of the experimental stimulus.

그림 2. 실험에 사용된 자극물(AI 미고지 조건과 AI 고지 조건)

Fig. 2. Experimental stimuli (AI non-disclosure condition vs. AI disclosure condition)

본 연구는 출처 고지 여부 이외의 시각적 차이를 통제하기 위해 단일 자극물을 사용하였다. 이는 자극물 간 차이로 인한 혼입효과를 최소화하고 내적 타당도를 확보하기 위한 선택이지만, 특정 자극물의 특성이 연구 결과에 영향을 미칠 가능성이 있으므로 자극물 유형에 따른 결과의 일반화에는 제한이 있을 수 있다. 향후 연구에서는 다양한 주제와 시각적 특성을 지닌 복수의 자극물을 활용하여 연구 결과의 일반화 가능성을 검증할 필요가 있다.

설문에 앞서 참여자에게 연구 목적과 조사 절차를 안내하고 자발적 참여 동의를 얻었다. 이후 온라인 설문 플랫폼을 통해 참여자를 고지 집단 또는 미고지 집단에 무작위로 배정하였다. 출처 고지 여부는 분석을 위해 미고지 집단을 0, 고지 집단을 1로 더미코딩하였다.

참여자는 실험 자극물을 충분히 확인한 후 디자인 요소 지각, 심미적 평가 및 사용자 신뢰 문항에 순차적으로 응답하였다. 출처 고지 조작에 대한 직접적인 주의 환기가 주요 측정 변수 응답에 영향을 미치는 것을 최소화하기 위해 조작점점 문항은 설문 마지막에 배치하였다.

4-2 연구대상 및 표집

본 연구의 모집단은 생성형 AI 기반 디자인 결과물에 노출될 가능성이 있는 20~40대 성인으로 설정하였다. 조사 참여자는 온라인 패널을 활용하여 성별과 연령대를 기준으로 할당된 비율을 할당표집 방식으로 모집하였다. 성별은 동일한 비율로 구성하고, 연령대는 20대·30대·40대가 동일한 비율로 포함되도록 할당하였다. 모집된 참여자는 출처 고지 집단과 미고지 집단에 각각 75명씩 무작위로 배정하였으며, 최종 분석에는 총 150명의 응답을 활용하였다.

표본 규모는 G*Power 3.1을 이용하여 산정하였다. 연구 모형 중 예측변수가 가장 많은 사용자 신뢰 회귀모형을 기준으로 중간 효과크기($f^2=.15$), 유의수준 $\alpha=.05$, 검정력(1-

β)=.80을 적용한 결과 최소 필요 표본수는 85명으로 산출되었다. 본 연구의 최종 표본수는 150명으로 최소 필요 표본수를 충족하였다.

또한 무작위 배정 이후 두 집단의 기초적 동질성을 확인하기 위해 성별, 연령 및 디자인 관련 경험에 대한 교차분석과 카이제곱 검정을 실시하였다.

4-3 변수의 조작적 정의

본 연구에서 사용한 주요 변수의 조작적 정의와 측정방법은 표 1에 제시하였다. 디자인 요소 지각은 생성형 AI 포스터의 색상, 레이아웃 및 타이포그래피에 대한 지각으로 정의하였으며, Ngo 등[8]의 인터페이스 심미성에 관한 논의를 참고하여 9개 문항으로 재구성하였다. 탐색적 요인분석 결과 교차적재가 확인된 1개 문항을 제외하고 최종 8개 문항을 분석에 활용하였다. 심미적 평가는 생성형 AI 디자인 결과물에 대한 전반적인 미적 평가로 정의하였으며, Lavie와 Tractinsky[11]의 심미성 논의를 참고하여 4개 문항으로 구성하였다. 사용자 신뢰는 생성형 AI 디자인 결과물에 대한 신뢰를 중심으로, 해당 결과

표 1. 변수의 조작적 정의

Table 1. Operational definitions of variables

Variable	Operational Definition	Measurement
Perception of Design Elements (X)	Respondents' perception of color, layout, and typography	8 items, 5-point Likert scale
Aesthetic Evaluation (M)	Overall aesthetic evaluation of the design outcome	4 items, 5-point Likert scale
User Trust (Y)	Trust toward the AI-generated design outcome and the entity responsible for its creation	5 items, 5-point Likert scale
AI Disclosure (W)	Presence or absence of information indicating that the design was generated using AI	Experimental manipulation: 0=non-disclosure, 1=disclosure
AI Disclosure Recognition	Degree to which participants recognized the AI-generated origin of the stimulus	2 items, 5-point Likert scale

표 2. 측정문항 및 출처

Table 2. Measurement items and sources

Variable	Measurement Items	Source
Perception of Design Elements (8 items)	X1. The color combination of this poster is harmonious. X2. Its colors match the intended theme. X3. Its colors provide visual stability. X4. Its information layout is easy to view. X5. Its elements are arranged in a balanced manner. X6. Important information is easy to identify. X7. Its typeface is highly readable. X8. Its font size and placement are appropriate.	Developed by the authors with reference to Ngo et al. [8]
Aesthetic Evaluation (4 items)	M1. This poster is aesthetically pleasing overall. M2. This poster is visually attractive. M3. This poster is pleasing to look at. M4. I am satisfied with the overall design of this poster.	Developed by the authors based on Lavie and Tractinsky [11]
User Trust (5 items)	Y1. I can trust this design outcome. Y2. I consider this design outcome to be reliable. Y3. I believe that this type of design outcome can be trusted and used in other contexts. Y4. I judge this design outcome to be trustworthy. Y5. I can trust the entity responsible for creating this design.	Developed by the authors based on Kim et al. [6] and Schilke and Reimann [12]
AI Disclosure Recognition (2 items)	C1. I think this poster was created by AI. C2. I was aware that this poster was created using AI.	Developed by the authors

Note. All items were measured on a 5-point Likert scale. X9, "The typeface of this poster matches its overall atmosphere," was excluded due to cross-loading in the EFA. The items were administered in Korean, and their English translations are presented in this table.

물의 제작 주체에 대한 신뢰를 포함하는 개념으로 정의하였다. 생성형 AI 서비스와 AI 활용 행위자에 대한 신뢰를 다룬 선행 연구[6],[12]를 참고하여, 디자인 결과물에 대한 신뢰 4개 문항과 제작 주체에 대한 신뢰 1개 문항으로 수정·구성하였다. 출처 고지 여부는 생성형 AI 제작 사실을 알리는 문구의 유무로 실험적으로 조작하였으며, 미고지 집단은 0, 고지 집단은 1로 코딩하였다. 출처 인식도는 해당 조작이 참여자에게 인지되었는지를 확인하기 위한 조작점검 변수로 활용하였다. 출처 고지 여부를 제외한 주요 구성개념은 5점 Likert 척도로 측정하였으며, 구체적인 측정문항과 출처는 표 2에 제시하였다.

본 연구는 색상, 레이아웃 및 타이포그래피 각각의 개별 효과를 비교하기보다, 디자인 요소 전반에 대한 통합적 지각이 심미적 평가와 사용자 신뢰에 미치는 영향을 분석하는 데 목적이 있다. 디자인 요소 지각은 최초 9개 문항으로 측정하였으나, 탐색적 요인분석 결과 “이 포스터의 글자체는 전체적인 분위기와 어울린다” 문항(X9)에서 교차적재가 확인되어 해당 문항을 제외하고 최종 8개 문항의 평균값을 분석에 활용하였다.

심미적 평가는 제시된 디자인 결과물에 대한 전반적인 미적·시각적 평가를 의미하며, 사용자 신뢰는 해당 디자인 결과물과 그 제작 주체에 대해 형성된 신뢰 수준을 의미한다. 두 구성개념 간 관련성이 높을 가능성을 고려하여, 심미적 평가 4개 문항과 사용자 신뢰 5개 문항을 대상으로 2요인 탐색적 요인분석을 추가로 실시하여 두 구성개념의 요인구조가 구분되는지를 확인하였다.

출처 고지 여부는 미고지 조건과 고지 조건으로 구분하여 실험적으로 조작하였으며, 조작점검 문항은 주요 종속변수 측정 이후 설문 마지막에 배치하였다.

4-4 분석방법

수집된 자료는 SPSS 27.0과 PROCESS Macro 5.0을 활용하여 분석하였다. 먼저 기술통계분석으로 응답자 특성과 주요 변수의 평균 및 표준편차를 확인하고, 집단 간 기초적 동질성 검증을 위해 성별, 연령 및 디자인 관련 경험에 대한 교차분석과 카이제곱 검정을 실시하였다. 이어 Cronbach's α 로 내적 일관성을, 탐색적 요인분석으로 요인구조를 확인하였으며, 심미적 평가와 사용자 신뢰의 구분 가능성은 두 변수의 총 9개 문항에 주축요인추출법과 직접 오블리민 회전을 적용한 2요인 분석으로 검토하였다. 또한 Pearson 상관분석, 조작점검 문항별 교차분석 및 통합점수 t-검정, 출처 고지 여부에 따른 사용자 신뢰 차이 검증을 위한 독립표본 t-검정을 순차적으로 실시하였다.

연구가설 검증에는 PROCESS Macro Model 14를 적용하였다. 디자인 요소 지각을 독립변수(X), 심미적 평가를 매개변수(M), 사용자 신뢰를 종속변수(Y), 출처 고지 여부를 조절변수(W)로 투입하였으며, 출처 고지 여부는 미고지 집단을 0, 고지 집단을 1로 더미코딩하고 상호작용항을 구성하는 심미적 평가는 평균중심화하였다. 조건부 간접효과와 조절된 매개

효과의 유의성은 부트스트래핑 5,000회를 통해 산출한 95% 백분위 부트스트랩 신뢰구간을 기준으로 판단하였으며, 신뢰구간에 0이 포함되지 않는 경우 유의한 것으로 해석하였다.

아울러 Model 14의 결과가 직접경로의 조절 가능성을 고려한 경우에도 유지되는지를 확인하기 위하여 PROCESS Macro Model 15를 강건성 검토로 추가 실시하였다. 이는 연구모형을 대체하기 위한 것이 아니라 주요 매개경로와 조절효과의 안정성을 보완적으로 확인하기 위한 분석이다.

V. 연구결과

5-1 응답자 특성

본 연구는 생성형 AI 디자인의 출처 고지 여부가 사용자 신뢰 형성에 미치는 영향을 검증하기 위하여 총 150명을 대상으로 온라인 설문조사를 실시하였다. 참여자는 출처 고지 집단(n=75)과 미고지 집단(n=75)에 무작위로 배정되었다. 성별은 남성과 여성이 각각 75명(50.0%)으로 동일하게 구성되었으며, 연령은 20대, 30대, 40대가 각각 50명(33.3%)씩 포함되었다. 또한 디자인 관련 경험은 경험 있음 25명(16.7%), 경험 없음 125명(83.3%)으로 나타났다.

표 3. 응답자 특성

Table 3. Demographic characteristics of respondents

Variable	Category	Discl. n(%)	Non-discl. n(%)	Total n
Gender	Male	37(49.3)	38(50.7)	75
	Female	38(50.7)	37(49.3)	75
Age	20s	25(33.3)	25(33.3)	50
	30s	25(33.3)	25(33.3)	50
	40s	25(33.3)	25(33.3)	50
Design Experience	Yes	14(18.7)	11(14.7)	25
	No	61(81.3)	64(85.3)	125

참여자는 출처 고지 집단과 미고지 집단에 각각 75명씩 배정되었다. 집단별 성별, 연령 및 디자인 관련 경험의 분포는 표 3과 같다. 집단 간 기초적 동질성을 확인하기 위한 카이제곱 검정 결과, 성별($\chi^2=.027$, $p=.870$), 연령($\chi^2=.000$, $p=1.000$), 디자인 관련 경험($\chi^2=.432$, $p=.511$)에서 모두 유의한 차이가 나타나지 않았다. 따라서 성별, 연령 및 디자인 관련 경험에서 두 집단 간 통계적으로 유의한 차이는 확인되지 않았다.

5-2 신뢰도·타당도

측정도구의 신뢰도와 구성타당도 분석 결과는 표 4와 같다. 신뢰도 분석 결과, 디자인 요소 지각은 Cronbach's α

표 4. 신뢰도 및 타당도 분석 결과

Table 4. Reliability and validity analysis

Variable	Items	α	KMO	Explained Variance (%)	Factor Loadings
Perception of Design Elements	8	.864	.815	51.535	.506~.819
Aesthetic Evaluation	4	.894	.827	75.953	.842~.887
User Trust	5	.932	.898	78.625	.869~.926

Note. α = Cronbach's alpha.

=.864, 심미적 평가는 α =.894, 사용자 신뢰는 α =.932로 나타나 모두 일반적인 기준인 .70 이상을 충족하여 높은 내적 일관성을 확보하였다.

구성타당도 검증을 위하여 각 척도별 탐색적 요인분석을 실시하였다. 디자인 요소 지각은 색상, 레이아웃 및 타이포그래피를 포함하는 통합적 구성개념으로 정의하였다. 최초 9개 문항을 대상으로 요인 수를 지정하지 않은 탐색적 요인분석을 실시한 결과, 제1요인의 고유값은 4.617로 전체 분산의 51.300%를 설명하였으며, 제2요인의 고유값이 1을 다소 상회하였으나 스크리 도표에서 제1요인 이후 고유값이 급격히 감소하였고, 디자인 요소 지각을 색상·레이아웃·타이포그래피에 대한 통합적 지각으로 설정한 이론적 정의를 함께 고려하여 단일요인 구조가 적합한 것으로 판단하였다. 다만 이 과정에서 X9 문항의 교차적재가 확인되어 해당 문항을 제외하고, 최종 8개 문항을 대상으로 요인 수를 1로 지정한 탐색적 요인분석을 실시하였다. 분석 결과, 디자인 요소 지각은 KMO=.815, Bartlett의 구형성 검정은 $\chi^2=548.287(df=28, p<.001)$ 로 나타났으며, 단일요인은 전체 분산의 51.535%를 설명하였다. 심미적 평가는 KMO=.827, Bartlett의 구형성 검정은 $\chi^2=349.359(df=6, p<.001)$ 로 나타났으며, 단일요인은 전체 분산의 75.953%를 설명하였다. 사용자 신뢰는 KMO=.898, Bartlett의 구형성 검정은 $\chi^2=592.860(df=10, p<.001)$ 로 나타났으며, 단일요인은 전체 분산의 78.625%를 설명하였다. 또한 각 척도 내 모든 문항의 요인적재량은 .50 이상(.506~.926)으로 나타나 각 측정척도가 해당 구성개념을 일관되게 반영하는 것으로 확인되었다.

심미적 평가와 사용자 신뢰가 서로 구분되는 구성개념인지를 추가로 확인하기 위하여 심미적 평가 4개 문항과 사용자 신뢰 5개 문항을 대상으로 2요인 탐색적 요인분석을 실시하였다. 주축요인추출법과 직접 오블리민 회전을 적용한 결과, KMO는 .924였으며, Bartlett의 구형성 검정은 $\chi^2=1060.875(df=36, p<.001)$ 로 나타나 요인분석에 적합한 것으로 확인되었다. 분석 결과, 사용자 신뢰 문항은 제1요인에 .684~.941, 심미적 평가 문항은 제2요인에 .605~.931의 요인적재량을 보였으며, .40 이상의 교차적재는 나타나지 않았다. 두 요인 간 상관은 .740으로 나타났다. 따라서 심미적 평가와 사용자 신뢰는 높은 관련성을 보이면서도 탐색적 수준에서 서로 구

표 5. 심미적 평가와 사용자 신뢰의 2요인 탐색적 요인분석 결과

Table 5. Two-Factor exploratory factor analysis of aesthetic evaluation and user trust

Construct	Number of Items	Factor Loading Range	Cross-loading
User Trust	5	.684~.941	None
Aesthetic Evaluation	4	.605~.931	None

Note. Principal axis factoring with direct oblimin rotation; KMO=.924, Bartlett's $\chi^2=1060.875(df=36, p<.001)$, factor correlation=.740.

분되는 요인구조를 형성하는 것으로 나타났으며, 구체적인 결과는 표 5와 같다.

5-3 기술통계·상관분석

주요 변수의 기술통계와 Pearson 상관분석을 실시한 결과는 표 6과 같다. 디자인 요소 지각의 평균은 3.332(SD=.681), 심미적 평가는 3.258(SD=.806), 사용자 신뢰는 3.235(SD=.748)로 나타났다.

상관분석 결과, 디자인 요소 지각은 심미적 평가($r=.685, p<.01$) 및 사용자 신뢰($r=.648, p<.01$)와 모두 유의한 정적 상관관계를 보였다. 또한 심미적 평가는 사용자 신뢰와 가장 높은 정적 상관관계($r=.715, p<.01$)를 나타냈다. 모든 상관계수는 .80 미만으로 나타나 변수 간 상관이 과도하게 높은 수준은 아닌 것으로 확인되었다.

표 6. 기술통계 및 상관분석

Table 6. Descriptive Statistics and Correlations

Variable	M	SD	1	2	3
Perception of Design Elements	3.332	.681	1		
Aesthetic Evaluation	3.258	.806	.685**	1	
User Trust	3.235	.748	.648**	.715**	1

Note. M = Mean; SD = Standard Deviation. Values in columns 1-3 are Pearson correlation coefficients. ** $p < .01$.

5-4 조작점검

조작점검은 출처 인식도 2개 문항(C1, C2)으로 측정하였으며, 문항별 응답분포는 표 7과 같다.

표 7. 조작점검 문항별 응답분포

Table 7. Response distributions for manipulation check items

Response	Item 1 Discl.	Item 1 Non-discl.	Item 2 Discl.	Item 2 Non-discl.
1	0(0.0)	2(2.7)	5(6.7)	15(20.0)
2	11(14.7)	18(24.0)	17(22.7)	11(14.7)
3	24(32.0)	24(32.0)	17(22.7)	25(33.3)
4	30(40.0)	24(32.0)	27(36.0)	16(21.3)
5	10(13.3)	7(9.3)	9(12.0)	8(10.7)

Note. Item 1: $\chi^2=4.886, p=.299$; Item 2: $\chi^2=10.682, p=.030$.

표 8. 조작점검 결과

Table 8. Manipulation check results

Variable	Discl. (n=75)	Non-discl. (n=75)	t	p	Cohen's d
AI Disclosure Recognition	3.380 (.834)	3.047 (1.024)	2.186*	.030	.357

Note. Discl. = disclosure condition; Non-discl. = non-disclosure condition. Values are presented as M (SD). *p < .05.

문항별 교차분석 결과, 첫 번째 문항에서는 고지 집단과 미고지 집단 간 응답분포의 차이가 유의하지 않았으나($\chi^2=4.886$, $p=.299$), 두 번째 문항에서는 집단 간 응답분포의 차이가 유의하게 나타났다($\chi^2=10.682$, $p=.030$). 단, 첫 번째 문항의 경우 일부 셀에서 기대빈도가 5 미만으로 나타나 해석에 주의가 필요하다.

또한 두 문항의 평균값을 산출한 통합점수 분석 결과는 표 8과 같으며, 고지 집단($M=3.38$, $SD=.83$)이 미고지 집단($M=3.05$, $SD=1.02$)보다 출처 인식도가 유의하게 높게 나타났다($t(148)=2.186$, $p=.030$, Cohen's $d=.357$). 다만 효과 크기는 작은 효과와 중간 효과 사이의 수준에 그쳤으며, 문항별 점검 결과도 일관되지 않았다는 점을 고려할 때 출처 고지 조작의 효과는 제한적으로 확인되었다.

5-5 출처 고지 여부에 따른 사용자 신뢰 차이 및 H2 검증

출처 고지 여부에 따라 사용자 신뢰 수준에 차이가 나타나 있는지를 확인하기 위하여 독립표본 t-검정을 실시하였다. 분석 결과, 고지 집단의 사용자 신뢰는 평균 3.155($SD=.720$), 미고지 집단은 평균 3.315($SD=.772$)로 나타났다. 두 집단 간 평균 차이는 통계적으로 유의하지 않았다($t(148)=-1.313$, $p=.191$, 평균차 $=-.160$, 95% CI $[-.401, .081]$, Cohen's $d=-.214$). 따라서 출처 고지 집단과 미고지 집단의 사용자 신뢰 수준에 차이가 있을 것이라는 H2는 지지되지 않았다.

5-6 가설검증

PROCESS Macro Model 14 분석 결과는 표 9와 같으며, 디자인 요소 지각은 심미적 평가에 유의한 정(+)의 영향을 미치는 것으로 나타났다($B=.810$, $SE=.071$, $t=11.424$, $p<.001$). 또한 심미적 평가는 사용자 신뢰에 유의한 정(+)의 영향을 미치는 것으로 나타났다($B=.523$, $SE=.087$, $t=6.022$, $p<.001$). 또한 디자인 요소 지각은 사용자 신뢰에 직접적으로도 유의한 영향을 미쳤다($B=.326$, $SE=.083$, $t=3.953$, $p<.001$). 따라서 H1과 H4는 지지되었다. 반면, 출처 고지 여부의 조절효과는 유의하지 않았으며, 이에 대한 조절된 매개 효과 분석 결과는 5-7절에 제시하였다.

5-7 조절된 매개효과

분석 결과, 출처 고지 여부의 회귀계수는 유의하지 않았으며($B=-.119$, $SE=.082$, $t=-1.457$, $p=.147$), 심미적 평가와

표 9. PROCESS Macro Model 14 분석 결과

Table 9. Regression results of PROCESS Macro Model 14

Outcome	Predictor	B	SE	t	p
Aesthetic Evaluation	Perception of Design Elements	.810	.071	11.424	<.001
User Trust	Perception of Design Elements	.326	.083	3.953	<.001
User Trust	Aesthetic Evaluation	.523	.087	6.022	<.001
User Trust	AI Disclosure	-.119	.082	-1.457	.147
User Trust	Aesthetic Evaluation × AI Disclosure	-.109	.102	-1.069	.287

Note. AI Disclosure: 0=Non-disclosure, 1=Disclosure. Aesthetic evaluation was mean-centered. Mediator model: $R^2=.469$, $F(1,148)=130.512$, $p<.001$; User trust model: $R^2=.568$, $F(4,145)=47.732$, $p<.001$.

출처 고지 여부의 상호작용 효과도 유의하지 않았다($B=-.109$, $SE=.102$, $t=-1.069$, $p=.287$). 따라서 출처 고지 여부가 심미적 평가와 사용자 신뢰 간 관계를 조절할 것이라는 H3은 지지되지 않았다.

조건부 간접효과와 조절된 매개효과 분석 결과는 표 10과 같으며, 디자인 요소 지각이 심미적 평가를 거쳐 사용자 신뢰에 미치는 간접효과는 미고지 집단($B=.423$, $BootSE=.084$, 95% Boot CI $[.276, .599]$)과 고지 집단($B=.335$, $BootSE=.121$, 95% Boot CI $[.090, .563]$)에서 모두 유의한 것으로 나타났다. 그러나 조절된 매개지수는 $-.088$ 이었으며, $BootSE=.109$, 95% 부트스트랩 신뢰구간 $[-.318, .109]$ 에 0이 포함되어 조절된 매개효과는 유의하지 않았다. 이는 디자인 요소 지각이 심미적 평가를 거쳐 사용자 신뢰로 이어지는 간접경로가 두 집단 모두에서 유의하였으나, 그 효과의 크기가 출처 고지 여부에 따라 통계적으로 달라지지는 않았음을 의미한다.

직접효과와 간접효과가 모두 유의하여 심미적 평가는 디자인 요소 지각과 사용자 신뢰 간 관계를 부분적으로 매개하는 것으로 나타났다.

표 10. 조절된 매개효과 분석 결과

Table 10. Moderated mediation analysis results

Conditional Indirect Effects				
AI Disclosure	Effect	BootSE	BootLLCI	BootULCI
Non-disclosure	.423	.084	.276	.599
Disclosure	.335	.121	.090	.563
Index of Moderated Mediation				
Statistic	Effect	BootSE	BootLLCI	BootULCI
Index of Moderated Mediation	-.088	.109	-.318	.109

Note. Bootstrap samples = 5,000; 95% percentile confidence intervals.

5-8 추가분석

Model 14의 결과가 직접경로의 조절 가능성을 추가로 고려한 경우에도 유지되는지를 확인하기 위하여 PROCESS Macro Model 15를 적용하였다. 분석 결과, 디자인 요소 지각과 출처 고지 여부의 상호작용항은 유의하지 않았으며 ($B=-.211, p=.211$), 심미적 평가와 출처 고지 여부의 상호작용항 역시 유의하지 않았다($B=.011, p=.939$). 또한 조절된 매개지수는 .009였으며, 95% 부트스트랩 신뢰구간 $[-.375, .264]$ 에 0이 포함되어 유의하지 않았다. 따라서 직접경로의 조절 가능성을 추가로 고려한 경우에도 출처 고지의 조절효과와 조절된 매개효과는 유의하지 않은 것으로 나타났다.

VI. 결 론

생성형 AI 기술의 활용이 디자인 분야 전반으로 확대되면서 생성형 AI 디자인 결과물에 대한 사용자 신뢰 형성 과정을 이해하는 것이 중요한 과제로 대두되고 있다. 그러나 생성형 AI 디자인에서 어떠한 요인이 사용자 신뢰 형성에 영향을 미치며, 그 과정이 어떠한 구조로 이루어지는지에 관한 실증적 연구는 아직 부족한 실정이다. 이에 본 연구는 생성형 AI 디자인 환경에서 디자인 요소 지각이 사용자 신뢰에 미치는 영향과 심미적 평가의 매개효과, 그리고 출처 고지 여부의 조절효과를 실증적으로 검증하고자 하였다.

분석 결과, 디자인 요소 지각은 심미적 평가에 유의한 정(+)의 영향을 미쳤으며, 심미적 평가는 사용자 신뢰에 유의한 정(+)의 영향을 미치는 것으로 나타났다. 또한 디자인 요소 지각이 심미적 평가를 통해 사용자 신뢰에 영향을 미치는 간접효과는 출처 고지 집단과 미고지 집단에서 모두 유의하게 나타나 심미적 평가의 매개역할이 확인되었다. 특히 디자인 요소 지각의 직접효과와 간접효과가 모두 유의하여 심미적 평가는 디자인 요소 지각과 사용자 신뢰 간 관계를 부분적으로 매개하는 것으로 확인되었다. 반면, 출처 고지 집단과 미고지 집단 간 사용자 신뢰의 직접적인 차이는 유의하지 않았으며, 출처 고지 여부가 심미적 평가와 사용자 신뢰 간의 관계를 조절하는 효과와 조절된 매개효과 역시 유의하지 않았다. 직접경로의 조절 가능성을 고려한 Model 15에서도 출처 고지의 조절효과와 조절된 매개효과는 유의하지 않았다. 이러한 결과는 본 연구의 실험 조건에서 출처 고지 여부에 따른 사용자 신뢰의 차이는 확인되지 않은 반면, 디자인 요소에 대한 지각과 심미적 평가를 통한 사용자 신뢰 형성 경로는 일관되게 나타났음을 보여준다.

본 연구의 조작점검 결과, 통합 출처 인식도는 고지 집단에서 유의하게 높았으나 효과크기가 크지 않았으며, 문항별 검정 결과도 일관되지 않아 출처 고지 조작의 효과는 제한적인 수준으로 확인되었다. 특히 미고지 집단에서도 생성형 AI 제작 사실을 일정 수준 인지한 것으로 나타나, 출처 고지 문구

의 유무만으로 두 집단의 인식을 명확하게 구분하는 데 한계가 있었을 가능성이 있다. 또한 텍스트 기반 창작물을 대상으로 한 선행연구[4]와 달리 본 연구는 시각디자인 결과물을 대상으로 하였으므로, 콘텐츠 유형과 제시 맥락의 차이가 출처 고지 효과에 영향을 미쳤을 가능성도 고려할 수 있다. 출처 고지의 조절효과가 유의하지 않았다는 결과는 출처 고지가 모든 생성형 AI 디자인 맥락에서 일률적으로 사용자 신뢰를 변화시키는 것은 아닐 가능성을 보여준다. 따라서 출처 고지 효과를 해석할 때에는 콘텐츠 유형, 고지 방식 및 제시 맥락을 함께 고려할 필요가 있다. 다만 본 연구의 조작 효과가 제한적으로 나타났다는 점에서 이를 출처 고지의 일반적인 무효과로 확대해석해서는 안 된다. 향후 연구에서는 이러한 조건을 세분화하여 체계적으로 검증할 필요가 있다.

본 연구는 기존 연구에서 주로 생성형 AI 서비스의 수용이나 AI 생성 여부 자체에 초점을 두었던 사용자 신뢰 연구를 확장하여, 동일한 생성형 AI 디자인 결과물을 활용한 실험을 통해 디자인 요소 지각, 심미적 평가 및 출처 고지 간의 구조적 관계를 통합적으로 검증하였다는 점에서 학문적 의의가 있다. 또한 생성형 AI 디자인을 활용하는 실무에서는 출처 고지 여부만을 고려하기보다 사용자가 긍정적으로 지각할 수 있는 디자인 요소와 높은 심미적 품질을 함께 확보하는 것이 사용자 신뢰 형성에 중요할 수 있음을 시사한다.

다만 본 연구는 하나의 생성형 AI 포스터를 활용하여 실험을 진행하였으므로 다양한 디자인 유형과 콘텐츠 맥락으로 연구결과를 일반화하는 데에는 한계가 있다. 또한 하나의 디자인 결과물을 기반으로 출처 고지 여부의 효과를 검증하였기 때문에 디자인 품질 수준이나 콘텐츠 특성에 따른 차이를 함께 살펴보지 못하였다. 향후 연구에서는 다양한 디자인 유형과 품질 수준뿐 아니라 자극물의 복잡도와 디자인 스타일을 체계적으로 구분하여, 고지 문구의 위치·크기·표현 방식 및 시각적 강조 수준에 따른 출처 고지 효과를 추가적으로 검증할 필요가 있다.

참고문헌

- [1] S. An, "An Identification of Antecedents to Designer Acceptance of Generative Design Tool," *Journal of Korea Design Forum*, Vol. 29, No. 3, pp. 125-138, August 2024. <https://doi.org/10.21326/ksdt.2024.29.3.011>
- [2] I. P. Park and J. Y. Han, "Types of Image Generative AI and Image Analysis: Focus on Representative Cases of Image Transfer, Translation and Extension," *Journal of Korea Institute of Spatial Design*, Vol. 18, No. 8, pp. 647-656, December 2023. <https://doi.org/10.35216/kisd.2023.18.8.647>
- [3] S. J. Jeon and S. I. Kim, "A Study on the Comparison of Technological Characteristics and Creative Utilization of AI-Based Image Generation Tools: Focusing on Midjourney, DALL-E, and Adobe Firefly," *Industry*

Promotion Research, Vol. 10, No. 3, pp. 121-130, July 2025. <https://doi.org/10.21186/IPR.2025.10.3.121>

- [4] M. Raj, J. M. Berg, and R. Seamans, "The Artificial Intelligence Disclosure Penalty: Humans Persistently Devalue AI-Generated Creative Writing," *Journal of Experimental Psychology: General*, Vol. 155, No. 4, pp. 896-915, April 2026. <https://doi.org/10.1037/xge0001889>
- [5] L. Bellaiche, R. Shahi, M. H. Turpin, A. Ragnhildstveit, S. Sprockett, N. Barr, ... and P. Seli, "Humans Versus AI: Whether and Why We Prefer Human-Created Compared to AI-Created Artwork," *Cognitive Research: Principles and Implications*, Vol. 8, No. 1, 42, July 2023. <https://doi.org/10.1186/s41235-023-00499-6>
- [6] S. Kim, J. Y. Cho, and B. G. Lee, "An Exploratory Study on the Trustworthiness Analysis of Generative AI," *Journal of Internet Computing and Services*, Vol. 25, No. 1, pp. 79-90, February 2024. <https://doi.org/10.7472/jksii.2024.25.1.79>
- [7] S. Han, S. Hong, J. Choi, and Y. Jung, "Consumer Perception of Artificial Intelligence Generated Content and Artificial Intelligence Disclaimer: Focusing on Content Value and Format," *Information Systems Review*, Vol. 26, No. 4, pp. 187-216, November 2024. <https://doi.org/10.14329/isr.2024.26.4.187>
- [8] D. C. L. Ngo, L. S. Teo, and J. G. Byrne, "Modelling Interface Aesthetics," *Information Sciences*, Vol. 152, pp. 25-46, June 2003. [https://doi.org/10.1016/S0020-0255\(02\)00404-8](https://doi.org/10.1016/S0020-0255(02)00404-8)
- [9] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer, "High-Resolution Image Synthesis With Latent Diffusion Models," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, New Orleans: LA, pp. 10674-10685, June 2022. <https://doi.org/10.1109/CVPR52688.2022.01042>
- [10] N. Tractinsky, A. S. Katz, and D. Ikar, "What Is Beautiful Is Usable," *Interacting with Computers*, Vol. 13, No. 2, pp. 127-145, December 2000. [https://doi.org/10.1016/S0953-5438\(00\)00031-X](https://doi.org/10.1016/S0953-5438(00)00031-X)
- [11] T. Lavie and N. Tractinsky, "Assessing Dimensions of Perceived Visual Aesthetics of Web Sites," *International Journal of Human-Computer Studies*, Vol. 60, No. 3, pp. 269-298, March 2004. <https://doi.org/10.1016/j.ijhcs.2003.09.002>
- [12] O. Schilke and M. Reimann, "The Transparency Dilemma: How AI Disclosure Erodes Trust," *Organizational Behavior and Human Decision Processes*, Vol. 188, 104405, May 2025. <https://doi.org/10.1016/j.obhdp.2025.104405>



백은지(Eun-Ji Baek)

2024년~현 재: 전북대학교 산업디자인학과 석사과정

※ 관심분야 : 생성형 AI(Generative AI), 산업디자인(Industrial Design), 디자인 연구(Design Research), 사용자 신뢰(User Trust) 등



조동민(Dong-Min Cho)

2004년~2006년: AAU, San Francisco, USA MFA

2008년~2009년: 서강대학교 게임교육연구원 전임강사

2009년~현 재: 전북대학교 산업디자인학과 교수

※ 관심분야 : 게임 디자인(Game Design), 시각디자인(Visual Design), 산업디자인(Industrial Design), 디자인 연구(Design Research) 등