

Friction as a Critical Concept for Media Art in the Age of Generative AI: Intersections across the Model, Dataset, and Output Layers

Daechan Heo^{1,2*}

¹Department of Visual Communication Design, Lecturer, Kookmin University, Seoul, Korea

²Editor-in-Chief, Aliceon, Media art & culture channel/Collective, Seoul, Korea

Abstract

Background Contemporary generative artificial intelligence (AI) increasingly produces visual and interactive experiences that appear natural and smooth through automation, prediction, and probabilistic optimization. As users generate texts and images and receive system responses, they tend to encounter outputs without sufficiently recognizing the underlying computational procedures, the ways in which data are configured, or the limitations and biases of the models involved. Recent scholarship at the intersection of art and technology has actively addressed issues such as datasets, synthetic media, prompts, and the politics of technology. However, conceptual frameworks that connect these issues to the critical functions performed by AI-based media art and analyze how artworks formally position their viewers remain comparatively underdeveloped.

Methods This study proposes friction as an analytical framework for media art criticism in the age of generative AI. The study brings Anna Tsing's concept of friction into dialogue with Alexander R. Galloway's theory of the interface and Jacques Rancière's concept of the distribution of the sensible. Drawing additionally on Kate Crawford and Vladan Joler's anatomical approach to AI systems, the study formulates the model, dataset, and output as three analytical layers. Friction is defined here not simply as a metaphor for resistance or technological failure, but as a critical event in which the conditions of reduction, omission, selection, and exclusion within AI systems are made perceptible through, or sutured over by, the formal and interface structures of artworks. These three layers are not intended as a fixed taxonomy that exhaustively explains AI systems. Rather, they function as components of an analytical framework for examining how individual artworks perform their critical operations.

Results Jooyoung Oh's BirthMark juxtaposes original artwork footage, the AI's object-detection process, keyword outputs, and the results of the AI's semantic processing of artists' interpretations of their own works. Through this parallel presentation, the work makes perceptible the gap between different forms of computational interpretation and the experience of art appreciation. Unmake Lab's dataset practices appropriate visual conventions of machine-learning datasets while introducing residual forms that cannot readily be reduced to existing classificatory categories. In doing so, they reveal how datasets construct certain entities as recognizable objects while omitting or leaving others as remainders. By contrast, Refik Anadol's Unsupervised functions as a comparative case in which certain technical conditions are partially disclosed, yet the tensions made perceptible by this disclosure are re-sutured through immersive visual surfaces, an anthropomorphized self-narrative, and the authority of the technical research community. Friction thus becomes partially perceptible before being reabsorbed into a smooth viewing experience.

Conclusions This study proposes friction not as an evaluative criterion for judging the critical value of an artwork, but as a critical framework for analyzing where the conditions of an AI system become perceptible and how they are subsequently sutured over. The study thereby shifts the criticism of AI-based media art away from assessments of technological novelty or the possibilities of artistic-technological convergence and toward an examination of the conditions of reduction, omission, and mediation embedded within AI systems. The proposed analytical framework may also serve as a conceptual resource for design criticism and design education concerned with the critical interpretation of generative-AI-based visual culture, interfaces, and data representation.

Keywords Friction, Generative Artificial Intelligence, AI-based Media Art, Interface, Dataset

*Corresponding author: Daechan Heo (s2016511@kookmin.ac.kr)

Citation: Heo, D. (2026). Friction as a Critical Concept for Media Art in the Age of Generative AI: Intersections across the Model, Dataset, and Output Layers. *Archives of Design Research*, 39(3), 365-384.

<http://dx.doi.org/10.15187/adr.2026.08.39.3.365>

Received : Apr. 24. 2026 ; **Reviewed :** Jul. 17. 2026 ; **Accepted :** Jul. 29. 2026
pISSN 1226-8046
eISSN 2288-2987

Copyright : 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 educational and non-commercial use, provided the original work is properly cited.

1. 서론

생성형 인공지능의 일상화는 기술의 복잡한 작동 과정을 더 잘 보여주기보다, 오히려 그것을 사용자의 경험에서 지워가는 방향으로 진행되고 있다. 사용자는 텍스트를 입력하고 이미지를 생성하며 질문에 대한 응답을 받지만, 그 사이에서 어떤 데이터가 선택되고 모델이 무엇을 계산하며 어떤 결과가 제외되었는지는 거의 의식하지 않는다. 입력과 출력 사이의 복잡한 매개 과정은 짧고 연속적인 상호작용으로 압축되고, 사용자는 그 과정이 아니라 유려하게 완성된 결과물과 마주한다. 이러한 조건은 효율성이나 편의성의 문제에 그치지 않고, 세계를 보고 이해하며 감각하는 방식을 조직하는 기술적·설계적 환경으로 이해될 필요가 있다(Chun, 2006; Crawford, 2021; Crawford & Joler, 2018).

이 논문에서 ‘매끄러움(smoothness)’은 이미지의 표면이 부드럽거나 출력이 자연스럽다는 뜻에 한정되지 않는다. 자연스러움이 결과물의 지각적 개연성이나 인간이 만든 것과의 유사성을 가리킨다면, 매끄러움은 입력과 출력 사이에서 이루어지는 변환, 선택, 삭제, 오류와 불확실성이 사용자 경험에서 감지되지 않도록 조직되는 상태를 뜻한다. 본 논문은 분석의 목적상 이를 유사한 용어인 seamlessness와 frictionlessness로부터 구별한다. Seamlessness가 매개 과정의 단절이 드러나지 않는 경험적 연속성을, frictionlessness가 저항이나 중단이 제거된 것으로 상상되는 상태를 강조한다면, 여기서 smoothness는 시스템 내부의 변환과 비대칭이 출력의 완결성과 인터페이스의 연속성 속에서 감지되지 않도록 조직되는 조건을 가리킨다(Galloway, 2012; Chun, 2006). 따라서 매끄러움은 마찰이 사라진 상태라기보다 시스템의 환원과 비대칭이 다른 작동 과정 속에 감춰지거나 결과의 완결성 안에 흡수된 상태에 가깝다.

그러나 생성형 AI의 매끄러운 작동을 비평적으로 분석하기 위한 개념적 어휘는 아직 충분히 마련되지 않았다. 기존의 AI 비판은 주로 편향과 차별, 저작권, 노동의 자동화, 기술적 위험을 중심으로 전개되어 왔고, 생성형 AI 시대의 미디어아트 논의 역시 기술 활용이나 형식적 새로움을 평가하는 데 집중하는 경우가 많았다. 이러한 논의는 AI가 야기하는 구체적인 문제를 확인하는 데 중요하지만, AI 시스템이 무엇을 생략하고 평균화하며 무엇을 보이거나 보이지 않게 만드는지를 작품의 형식과 관람 경험 차원에서 설명하는 데에는 한계가 있다.

최근 예술·기술 영역에서는 데이터셋, 합성 미디어, 프롬프트, 계산 규모, 기술의 정치성과 같은 의제가 활발히 논의되어 왔다. 질린스카(Zylinska, 2020)는 AI 예술 담론이 계산 기술의 정치적 전제와 창작 조건을 충분히 드러내지 못한다고 비판했고, 슈타이얼(Steyerl, 2023)은 생성 이미지를 확률과 평균의 시각적 렌더링이자 사회적 필터를 통과한 ‘평균 이미지(mean image)’로 분석했다. 국내에서는 노현정·이예승(Ro & Lee, 2026)이 생성형 AI의 창작 활용 패턴을 구조화했으며, 정규리(Chung, 2023)와 박태성(Park, 2024)은 레픽 아나돌의 작업을 중심으로 데이터, 알고리즘, 설명가능성의 문제를 논의했다. 이들 연구는 생성형 AI의 기술적·미학적·윤리적 조건을 각각 조명했지만, 생성형 AI 시대의 미디어아트가 그러한 조건을 작품의 형식과 관람 경험 속에서 어떻게 드러내거나 다시 통합하는지를 분석하기 위한 비평 개념으로까지는 충분히 정식화하지 않았다.

이러한 문제의식에서 이 논문은 생성형 AI 시대 미디어아트의 비평적 작동을 설명하기 위한 개념으로 ‘마찰(friction)’을 제안한다. 여기서 마찰은 단순한 방해나 저항, 기술적 오류를 뜻하지 않는다. 마찰은 AI 시스템의 조건 자체가 아니라, 그 시스템에 내재한 환원과 누락이 작품의 형식과 관람 경험 속에서 감각 가능한 사건으로 출현하는 순간을 가리킨다. 데이터의 선택과 모델의 환원, 출력의 감각적 조직이 더 이상 자명하고 자연스러운 과정으로 유지되지 못하고 다시 지각될 때 마찰이 발생한다.

분석 대상은 생성 모델을 직접 활용한 작품에만 한정하지 않는다. 생성 모델과 객체 인식·분류 모델은 기술적으로 동일한 파이프라인을 이루지 않지만, 세계의 일부를 데이터로 선택하고 계산 가능한 관계와 패턴으로 조직한다는 인식론적 조건을 공유한다. 생성형 AI의 문제는 최종적인 이미지나 문장의 생성에서만

시작되지 않는다. 무엇이 학습 가능한 자료로 구성되고, 어떤 특징과 관계가 계산 대상으로 추출되며, 그 결과가 어떤 인터페이스와 시각적 형식으로 제시되는가에서 이미 형성된다. 이에 따라 생성 모델을 직접 활용한 작업과 함께, 생성형 AI 시대의 시각문화를 구성하는 데이터화·분류·확률적 환원의 조건을 비평적으로 다루는 미디어아트를 분석 범위에 포함한다.

다만 이 범위가 AI를 사용하는 모든 미디어아트를 포괄하는 것은 아니다. AI를 제작 도구로 사용하면서 그 작동 조건을 작품의 문제로 다루지 않는 작업이나 기술에 대한 선언적 찬반에 머무르는 작업은 분석 대상에서 제외한다. 분석의 초점은 모델의 환원, 데이터셋의 구성과 분류, 출력과 인터페이스의 감각적 조직이 작품의 형식과 관람 경험 안에서 어떻게 비평적 질문으로 전환되는가에 있다.

이상의 문제를 구체화하기 위해 다음 두 가지 연구 질문을 설정한다.

연구 질문 1. 생성형 AI 시대의 미디어아트는 모델·데이터셋·출력의 층위에서 발생하는 환원과 누락을 작품의 형식과 관람 경험 속에서 어떻게 드러내거나 다시 봉합하는가?

연구 질문 2. 인터페이스의 구성, 데이터의 표현, 출력의 연출과 같은 디자인적 선택은 마찰의 가시화와 봉합을 어떻게 매개하는가?

첫 번째 질문은 마찰을 AI 시스템의 구체적인 작동 층위와 연결한다. 두 번째 질문은 작품의 비평적 의미가 다루는 기술이나 작가의 선언만으로 결정되지 않고, 그 기술을 감각 가능한 형식으로 조직하는 디자인적 선택에 의해 달라진다는 점을 검토한다. 여기서 디자인은 완성된 결과물의 조형이나 사용성에 한정되지 않는다. 모델이 사용하는 범주, 데이터의 선택과 배열, 인터페이스의 정보 구조, 출력의 감각적 연출을 통해 무엇을 전경화하고 무엇을 감추는가를 조직하는 과정까지 포함한다.

연구는 이론적 문헌고찰, 개념적 종합, 비교 사례분석으로 진행된다. 먼저 칭(Anna Tsing)의 마찰 개념, 알렉산더 R. 갤러웨이(Alexander R. Galloway)의 인터페이스 이론, 자크 랑시에르(Jacques Rancière)의 감각의 분할 개념을 각각 현실과의 접촉, 시스템 내부의 매개, 감각 가능한 것의 분할이라는 차원에서 교차 검토하여 마찰의 의미를 재구성한다. 이어 크로퍼드와 윌러(Crawford & Joler, 2018)의 AI 시스템 해부학을 참조하여 모델·데이터셋·출력을 세 분석 층위로 설정하고, 그 관계를 3장에서 교차적 분석틀로 정식화한다.

비교 사례로는 오주영의 〈BirthMark〉, 언메이크랩의 데이터셋 실천, 레픽 아나돌(Refik Anadol)의 〈Unsupervised〉를 선정하였다. 세 사례는 AI 미디어아트 전체를 대표하는 표본이 아니라, 객체 인식과 인지 모델, 기계 시각과 데이터셋 구성, 생성 모델과 물입형 출력이라는 서로 다른 기술적 조건에서 마찰이 어떻게 다르게 작동하는지를 비교하기 위해 선택하였다. 작품의 기술적 구성과 작가의 진술, 전시 자료 및 비평적 논의를 함께 확인할 수 있다는 점도 선정 과정에서 고려하였다.

사례분석은 제안된 개념을 실증적으로 검증하거나 작품의 비평적 우열을 판정하기 위한 절차가 아니다. 각 사례에서 모델·데이터셋·출력의 층위가 어떻게 교차하며, 인터페이스와 데이터 표현, 출력의 연출이 마찰을 드러내거나 다시 봉합하는지를 비교함으로써 분석틀의 적용 가능성과 한계를 점검하고 개념을 정교화한다. 이를 통해 생성형 AI 시대의 미디어아트가 기술 활용에 머물지 않고, 기술이 세계와 감각을 조직하는 조건을 다시 드러내고 재배치하는 비평적 실천으로 읽힐 수 있음을 논증하고자 한다.

2. 마찰의 이론적 재구성

2. 1. 보편성과 현실의 접촉으로서의 마찰

안나 칭(Anna Tsing)의 마찰 개념은 글로벌라이제이션(globalization)의 인류학적 분석에서 출발한다. 칭이 문제 삼는 것은 보편적 체계가 지역적 현실과 대립한다는 단순한 구도가 아니다. 보편성은 이미 완성된 형태로 세계 곳곳에 동일하게 적용되는 것이 아니라, 서로 다른 장소와 역사, 문화와 권력관계를 통과하는 과정에서 비로소 현실적인 힘을 얻는다. 그 과정에는 번역과 협상, 오해와 불균등한 협력이 개입하며, 이동하는 체계 역시 접촉 이전과 동일한 상태로 남지 않는다. 칭은 이러한 이질적인 접촉의 질을 마찰이라고 부른다(Tsing, 2005).

따라서 마찰은 연결을 외부에서 방해하는 장애물이나 보편적 체계에 맞서는 저항과 동일하지 않다. 도로의 마찰이 이동을 늦추면서도 이동 자체를 가능하게 하듯이, 사회적·문화적 마찰 역시 연결을 어렵고 불균등하게 만들면서 그 연결이 실제 세계에서 작동하도록 한다. 중요한 것은 마찰이 연결의 실패 이후에 출현하는 예외가 아니라, 연결이 성립하는 과정 안에 처음부터 포함되어 있다는 점이다. 보편적 체계는 마찰을 제거함으로써 현실에 도달하는 것이 아니라, 마찰을 통과하면서 구체적인 형태를 얻는다. 이때 접촉은 보편성을 단순히 적용하는 과정이 아니라 그것을 변형하고 예상하지 못한 결과를 만들어내는 과정이 된다.

이 관점은 생성형 AI가 보편적 지능이나 일반적 문제 해결 장치로 제시되는 방식을 비판적으로 읽게 한다. 생성형 AI의 일반화 능력은 추상적인 계산 체계만으로 성립하지 않는다. 무엇이 학습 자료로 수집되고 어떤 범주와 규범이 사용되는가에 따라 시스템이 다룰 수 있는 세계가 달라진다. 보편적으로 작동하는 것처럼 보이는 시스템도 실제로는 특정한 데이터의 분포와 계산 가능한 형식에 의존한다. 이 과정에서 발생하는 누락과 오분류, 의미의 축소는 시스템 바깥에서 우연히 생겨난 예외만이 아니라, 이질적인 현실을 계산 가능한 관계로 변환하면서 구조적으로 발생하는 결과다. 기술이 정확하고 자연스럽게 작동하는 순간에도 구체적인 대상은 학습된 패턴과 범주 안에서 선택적으로 번역되고 환원된다.

생성형 AI 시대 미디어아트에서 중요한 것은 이러한 접촉이 존재한다는 사실 자체가 아니라, 평소에는 자연화된 접촉의 비대칭이 작품의 형식과 관람 경험 속에서 어떻게 다시 지각되는가다. 칭의 논의가 이 논문에 제공하는 첫 번째 이론적 차원은 바로 이 지점에 있다. 마찰은 보편적 계산 체계가 구체적 현실과 만나는 장소와 그 접촉의 불균등한 조건을 묻게 한다. 다만 접촉이 시스템 내부에서 어떠한 번역과 선택을 거쳐 결과로 조직되는지는 이 개념만으로 충분히 설명되지 않는다. 다음 절에서는 이 문제를 인터페이스의 매개 작용을 통해 살핀다.

2. 2. 인터페이스와 생성적 교란

알렉산더 R. 갤러웨이(Alexander R. Galloway)의 인터페이스 이론은 생성형 AI의 매끄러운 인터페이스가 자기 작동의 조건을 어떻게 감추고 자연화하는지를 분석하기 위한 관점을 제공한다. 그에게 인터페이스는 단순한 화면이나 경계점이 아니라 서로 다른 상태를 매개하는 능동적 문턱이며, 고정된 객체라기보다 결과를 산출하는 과정이다. 그가 말하듯 인터페이스는 “사물이 아니라 결과를 산출하는 과정(not things, but rather processes that effect a result)”이다(Galloway, 2012, p. vii).

이 관점에서 인터페이스의 매끄러움은 매개가 사라졌다는 뜻이 아니라, 매개가 지나치게 잘 작동하여 스스로를 감춘 결과다. 웬디 후이 경(Wendy Hui Kyong Chun)은 네트워크가 한때 “마찰 없는(friction-free)” 자유의 공간처럼 상상되었지만, 실제로는 통제와 프로토콜의 비가시적 작용을 통해 유지된다고 지적한다(Chun, 2006). 생성형 AI 인터페이스 역시 사용자가 단지 질문을 입력하고 결과를 받는다고 느끼게 만들지만, 그 사이에는 수많은 변환·선택·필터링이 작동한다. 인터페이스는 자연스러운 상호작용의 표면처럼 보이지만

실제로는 무엇을 보이게 하고 무엇을 삭제할지를 조정하는 구조다. 이 점에서 생성형 AI 인터페이스는 단순한 사용 환경이 아니라, 무엇을 자연스럽게 투명한 것으로 경험하게 만드는지를 조직하는 인터페이스 정치(interface politics)의 장으로 이해될 수 있다.

이러한 인터페이스 관점은 보편적 계산과 구체적 현실의 접촉이 시스템 내부에서 어떻게 번역되고 조정되는지를 설명한다. 사용자가 마주하는 자연스러운 결과는 매개가 사라진 상태가 아니라, 수많은 변환과 선택이 하나의 연속적인 경험으로 조직된 결과다. 2.1이 계산 체계와 현실의 비대칭적 접촉을 다루었다면, 이 절은 그 접촉이 인터페이스의 선택·필터링·표현을 거쳐 어떻게 특정한 결과로 구성되는지를 두 번째 이론적 차원으로 제시한다.

2. 3. 감성의 분할과 비가시성의 정치학

자크 랑시에르(Jacques Rancière)의 감성의 분할 개념은 생성형 AI의 출력이 무엇을 감각 가능한 것으로 만들고 무엇을 비가시화하는지를 읽기 위한 정치미학적 관점을 제공한다. 그에게 감성의 분할은 무엇이 보이고 들리며 말해질 수 있는지를 조직하는 질서이며, 어떤 존재가 공통의 세계 안에서 지각 가능한 것으로 인정되는가를 결정한다(Rancière, 2004).

이 관점에서 접근할 때 생성형 AI는 단순히 정보를 처리하는 기술이 아니라, 어떤 패턴을 전경화하고 어떤 잔여를 제거하는 감각적 장치로 이해될 수 있다. 생성형 AI는 세계를 있는 그대로 재현하지 않는다. 그것은 이미 어떤 것을 보이게 하고 어떤 것을 보이지 않게 하는 새로운 분할 체계를 구성한다.

따라서 생성형 AI 시대 미디어아트 비평은 출력 결과의 미학적 새로움만을 평가하는 데 머물 수 없다. 어떤 패턴과 존재가 감각 가능한 것으로 전경화되고, 어떤 차이와 잔여가 보이지 않는 것으로 밀려나는지를 함께 읽어야 한다. 랑시에르의 논의는 인터페이스의 매개가 단순한 정보 처리에 그치지 않고, 무엇이 공통의 세계에 나타날 수 있는가를 조직하는 감각적·정치적 작동임을 보여준다. 이로써 보편적 체계와 현실의 접촉, 인터페이스 내부의 매개에 이어 감각 가능한 것의 분할이라는 세 번째 이론적 차원이 마련된다.

2. 4. 생성형 AI 시대 미디어아트 비평 개념으로서의 마찰

칭, 갤러웨이, 랑시에르의 논의는 서로 다른 층위의 질문에 답한다. 칭은 보편적 체계가 어디에서 어떤 비대칭을 수반하며 구체적 현실과 접촉하는지를 묻게 한다. 갤러웨이는 그 접촉이 시스템 내부의 번역과 선택, 필터링을 거쳐 어떻게 결과를 산출하는지를 분석하게 한다. 랑시에르는 그러한 매개의 결과가 무엇을 감각 가능한 것으로 만들고 무엇을 비가시화하는지를 드러낸다. 세 논의는 각각 접촉의 조건, 매개의 작동, 감각의 분할을 설명하며, 이 세 차원이 교차할 때 생성형 AI 시대 미디어아트 비평적 작동을 분석할 수 있는 마찰 개념이 구성된다.

이 논문에서 마찰은 보편화된 계산 체계와 구체적 현실의 비대칭적 접촉이 인터페이스의 매개를 거쳐 작품의 형식과 관람 경험 속에서 감각 가능한 문제로 출현하는 비평적 사건을 뜻한다. 마찰은 AI 시스템의 모든 조건을 포괄하는 총체적 명칭도 아니며, 기술적 오류나 불편 자체를 지시하지도 않는다. 모델의 환원, 데이터셋의 분류와 누락, 출력의 감각적 자연화가 더 이상 자명한 과정으로 유지되지 못하고 관객이 지각하고 해석할 수 있는 긴장으로 전환될 때 마찰이 발생한다.

이 정의에서 매끄러움과 마찰은 서로 배타적인 두 상태가 아니다. 매끄러움은 시스템 내부의 변환과 비대칭이 사라진 상태가 아니라, 그것이 사용자의 경험에서 감지되지 않도록 조직된 상태다. 마찰은 그 조직이 중단되거나 실패할 때에만 나타나지 않는다. 작품이 매끄러운 출력의 형성 과정을 분해하거나 서로 다른 해석 체계를 병치하고, 분류되지 않는 잔여를 데이터의 형식 안에 배치할 때에도 발생한다. 반대로 작품이 데이터와 연산의 조건을 일부 드러내면서도 이를 몰입적 출력과 의인화된 서사 속에 다시 흡수할 수 있다. 따라서 이

논문이 제안하는 분석틀은 마찰의 유무만을 판정하지 않고, 마찰이 작품 안에서 어떠한 형식으로 출현하며 다른 장치에 의해 어떻게 다시 통합되는지를 함께 추적한다.

마찰의 개념적 범위를 보다 분명히 하기 위해 이를 매개, 인터페이스 정치, 불투명성, 데이터 정치, 생산적 마찰과 구분할 필요가 있다(Table 1).

Table 1 Comparison of friction and adjacent concepts

개념	대표 논의	주요 분석 대상	마찰과의 관계
Mediation	Galloway	서로 다른 형식과 상태를 연결하고 결과를 산출하는 매개의 작동	마찰은 매개 일반이 아니라, 매개 과정의 번역과 환원이 비대칭적 접촉으로 지각되는 순간을 분석한다.
Interface politics	Galloway, Chun	인터페이스가 선택·삭제·표현을 조직하면서 생산하는 권력 효과	마찰은 인터페이스의 권력 효과가 작품과 관람 경험에서 무엇을 보이게 하거나 비가시화하는지 분석한다.
Opacity	Burrell	머신러닝 시스템의 불투명성과 인간의 해석 능력 사이의 불일치	마찰은 불투명성 자체보다, 그 불일치가 작품의 형식 속에서 감각 가능한 문제로 전환되는 사건에 주목한다.
Data politics	Crawford & Joler, D'Ignazio & Klein	데이터의 수집·분류·표현에 내재한 권력과 배제	마찰은 데이터 정치가 작품의 형식과 관람 경험 속에서 어떻게 표현되고 의식화되는지를 분석한다.
Productive friction	Cox et al.; Mejtoft et al.	자동적 행동을 중단시키고 성찰을 유도하기 위해 의도적으로 삽입된 상호작용	마찰은 사용자 행동을 조절하기 위한 설계 전략이 아니라, 시스템의 환원과 누락이 감각 가능한 사건으로 출현하는 방식을 분석한다.

버렐(Burrell, 2016, pp. 4-5)은 머신러닝 알고리즘의 불투명성(opacity)을 기업 비밀, 기술 문해력의 부족, 머신러닝의 수학적 최적화와 인간의 의미 해석 사이의 근본적인 불일치로 구분한다. 이 가운데 세 번째 불투명성은 이 논문의 문제의식과 직접 연결된다. 그러나 버렐이 이러한 불일치를 분석하고 해명해야 할 사회기술적 문제로 다룬다면, 마찰은 같은 불일치가 작품의 형식과 관람 경험 속에서 감각 가능한 사건으로 변환되는 방식에 주목한다. 마찰은 불투명성을 해소하거나 시스템을 투명하게 만들기 위한 개념이 아니라, 불투명성과 해석의 간극이 미학적·비평적 문제로 출현하는 장면을 읽기 위한 개념이다.

국내 연구에서도 한지윤과 최유미(Han & Choi, 2026)는 생성 모델의 파라미터를 조정하여 의도된 기술적 불완전성을 유도하고, 이를 관객의 미적 경험을 활성화하는 설계 요소로 분석하였다. 이 접근이 불완전성을 경험 설계의 자원으로 회수한다면, 본 연구의 마찰은 불완전성 자체보다 AI 시스템의 환원과 누락이 작품의 형식과 관람 경험 속에서 비평적으로 가시화되는 방식에 초점을 둔다.

HCI와 UX에서 마찰은 대체로 사용자 경험의 흐름을 방해하는 요소로 간주되어 제거의 대상이 되어 왔다. 일부 인터랙션 디자인 연구는 이를 생산적 마찰(productive friction)로 전환하여, 의도적인 지연이나 중단이 사용자의 자동적 행동을 멈추고 반성적 판단을 유도할 수 있다고 본다(Cox et al., 2016; Mejtoft et al., 2019). 이 논문의 마찰은 사용자의 행동을 늦추기 위해 설계자가 삽입하는 상호작용 장치와 다르다. 분석의 대상은 마찰을 의도적으로 설계했는가가 아니라, AI 시스템의 환원과 누락이 작품의 형식과 관람 경험 속에서 어떠한 문제로 조직되는가다. 따라서 마찰은 행동 조절을 위한 설계 전략이 아니라, 기존 작품과 인터페이스에서 시스템의 조건이 어떻게 감각되는지를 분석하기 위한 비평 개념이다.

비판적 디자인과 성찰적 디자인은 기술과 디자인에 내재한 가치와 전제를 질문하고 성찰을 유도하는 제작 및 설계 접근이라는 점에서 이 논문의 문제의식과 접점을 가진다. 그러나 이 논문에서 마찰은 특정한 디자인 장르나 제작 방법론을 지시하지 않는다. 작가나 디자이너가 비판적 의도를 선언했는지보다, 모델·데이터셋·출력의 환원과 누락이 작품의 형식과 관람 경험 속에서 실제로 드러나거나 자연화되는 방식을 분석한다는 점에서 구별된다(Dunne & Raby, 2013; Sengers et al., 2005).

이러한 개념적 구분을 바탕으로 다음 장에서는 마찰이 작품 속에서 발생하거나 통합되는 양상을 모델·데이터셋·출력의 세 층위에서 살펴본다. 세 층위는 칭, 갤러웨이, 랑시에르의 논의와 일대일로 대응하지 않는다. 보편적 체계와 현실의 접촉, 인터페이스의 매개, 감각 가능한 것의 분할은 각 층위에서 서로 다른 비중으로 교차한다. 모델·데이터셋·출력은 이러한 교차가 작품의 형식과 관람 경험 안에서 어떻게 구체화되는지를 추적하기 위한 분석 층위다.

3. 생성형 AI 시스템의 세 분석 층위와 마찰

2장에서 정리한 현실과의 접촉, 인터페이스의 매개, 감각 가능한 것의 분할은 마찰을 구성하는 이론적 차원이다. 이에 비해 모델·데이터셋·출력은 마찰이 AI 시스템과 작품의 형식 안에서 어떻게 구체화되는지를 살피기 위한 분석 층위다. 두 체계는 일대일로 대응하지 않는다. 현실과의 비대칭적 접촉, 인터페이스의 번역과 선택, 감각 가능한 것의 전경화와 비가시화는 세 층위 모두에서 서로 다른 비중으로 작동한다.

각 사례 역시 하나의 층위에 고정되지 않는다. 다만 작품마다 특정 층위의 환원과 누락이 중심으로 전경화되며, 그 비평적 작동은 다른 층위들과의 관계 속에서 형성된다. Figure 1은 세 이론적 차원과 세 분석 층위의 교차, 각 사례의 중심 층위와 연관 층위, 그리고 사례 분석에서 도출된 핵심 내용을 요약한다.

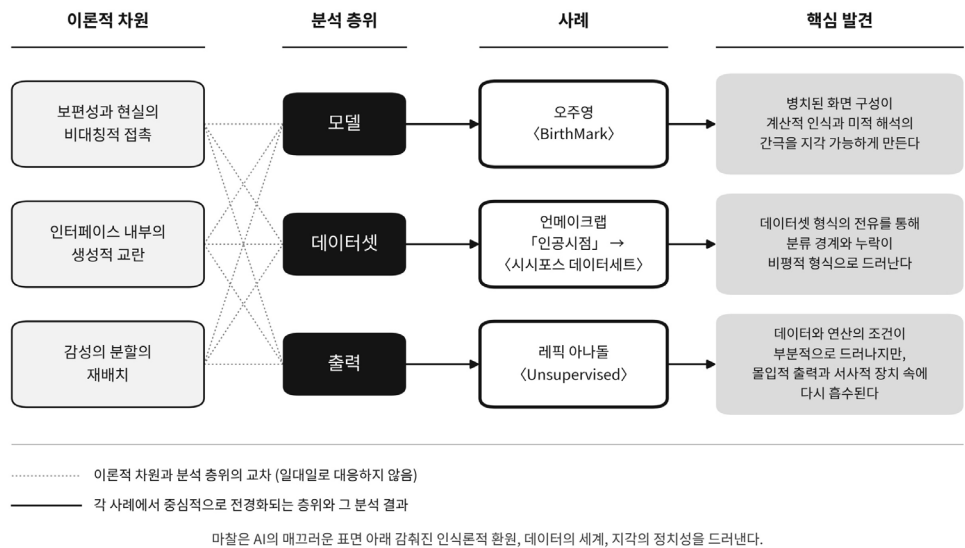


Figure 1 Intersections among theoretical dimensions, analytical layers, cases, and key findings

3. 1. 생성형 AI 시스템의 해부학과 세 분석 층위

모델·데이터셋·출력의 세 층위는 생성형 AI 시스템의 기술적 실재를 완전하게 기술하기 위한 구조도, 비평 작업을 순서대로 수행하기 위한 단계적 방법론도 아니다. 이 논문은 기존 AI 파이프라인 담론을 미디어아트 비평의 관점에서 재독해하여, 세 층위를 작품 분석을 위한 교차적 분석틀로 사용한다. 이 틀은 작품마다 환원과 누락이 어느 층위에서 두드러지게 나타나는지, 그것이 다른 층위와 어떤 관계를 맺으며 작품의 형식과 관람 경험으로 전환되는지를 추적한다. 따라서 개별 작품은 모델·데이터셋·출력 가운데 하나에 고정되지 않는다. 특정 층위가 중심으로 전경화될 수는 있지만, 그 비평적 작동은 세 층위의 연결 속에서 형성된다.

세 분석 층위의 구성에는 크로퍼드와 올레르(Crawford & Joler, 2018)의 AI 시스템 해부학이 주요한 참조점이 된다. 다만 이 논문의 목적은 자원, 노동, 인프라를 포함한 생성형 AI의 물질적·사회경제적 조건 전체를 해부하는 데 있지 않다. 분석의 초점은 AI 시스템이 세계를 인식하고 분류하며 감각 가능한 결과로 조직하는 방식이 작품의 형식 속에서 어떻게 드러나는가에 있다. 이에 AI 시스템 해부학이 제시하는 여러 차원 가운데 인식과 감각의 조직에 직접 관여하는 모델, 데이터셋, 출력의 층위를 선택하여 논의를 전개한다.

파스퀴넬리(Pasquinelli, 2023, 2024)는 AI를 생물학적 지능의 모방이 아니라 노동과 사회적 관계의 지능을 포획하고 자동화하는 역사적 프로젝트로 이해한다. 이 관점은 모델·데이터셋·출력 바깥에도 자원, 노동, 인프라의 사회기술적 조건이 존재한다는 점을 환기한다. 다만 이 논문은 분석 범위를 AI 시스템이 인식과 감각을 조직하는 방식으로 제한하며, 이러한 물질적·사회경제적 조건에 관한 논의는 후속 연구의 과제로 남긴다.

세 층위는 독립된 분석 단위가 아니라 상호 순환하는 시스템의 절단면이다. 모델은 데이터셋의 통계적 관계 위에서 학습되며, 데이터셋은 이전 세대 출력의 누적과 재가공을 포함하고, 출력은 다시 다음 세대 모델의 학습 자료가 될 수 있다. 합성 데이터의 본격적인 활용은 출력이 데이터셋의 외부에 놓인 최종 결과가 아니라, 이후 데이터셋의 일부로 편입되는 구조를 보여준다. 이 지점에서 모델·데이터셋·출력은 선형적 파이프라인이라기보다 자기참조적으로 순환하는 관계로 이해될 필요가 있다.

매끄러움은 이러한 순환이 사용자 경험에서 닫힌 체계처럼 자연화될 때 나타난다. 어느 한 층위의 환원은 다른 층위에서 반복되거나 증폭될 수 있으며, 출력의 완결성은 그 누적 과정을 비가시화한다. 반대로 마찰은 이 순환의 특정 지점이 작품의 형식 속에서 의식화되는 사건이다. 모델의 분류가 데이터셋의 분포와 어떻게 연결되는지, 출력의 매끄러움이 어떤 환원의 누적 위에 형성되었는지를 드러낼 때 마찰이 발생한다. 따라서 교차적 분석들은 마찰이 두드러지는 층위뿐 아니라, 그것이 다른 층위와 연결되거나 작품의 형식 안에서 다시 흡수되는 과정까지 추적한다.

3. 2. 모델의 마찰

모델의 층위에서는 AI가 무엇을 인식하고 어떤 관계를 의미 있는 것으로 계산하는지가 핵심이 된다. 모델은 세계를 있는 그대로 이해하는 것이 아니라, 학습된 범주와 패턴에 따라 대상을 분류하고 다음 항목을 예측 가능한 형태로 환원한다. 대규모 언어 모델의 경우 트랜스포머 아키텍처(Transformer Architecture)와 어텐션 메커니즘(Attention Mechanism)을 바탕으로 문맥적 관계를 학습하지만, 이러한 성능은 의미 그 자체에 대한 이해를 보장하기보다 언어적·기호적 패턴의 정교한 확률 계산에 가깝다(Bender et al., 2021).

에밀리 M. 벤티 외(Bender et al., 2021, p. 617)는 대규모 언어 모델을 “확률적 앵무새(stochastic parrots)”라고 규정하며, 그것이 의미를 이해하기보다 훈련 데이터에서 관찰한 패턴을 확률적으로 조합한다고 비판한다. 이 비판을 서로 다른 AI 모델에 동일한 기술적 방식으로 적용할 수는 없다. 언어 모델과 시각 모델은 구조와 목적에서 차이를 보이지만, 학습 데이터에서 추출한 패턴과 범주를 통해 세계를 계산 가능한 형태로 조직한다는 환원 문제를 공유한다. 따라서 모델의 마찰은 기술의 실패나 오작동뿐 아니라 성공적인 인식, 분류, 생성처럼 보이는 순간에도 드러날 수 있다.

시각 AI는 객체 탐지와 이미지 분류 모델, CLIP과 같은 시각-언어 모델, Stable Diffusion과 같은 생성 모델에 이르기까지 서로 다른 구조와 목적을 가진다. 그러나 이러한 학습 기반 모델의 성능과 출력은 모두 대규모 데이터의 구성과 범주화 방식에 의존한다. 데이터셋은 모델에 투입되는 중립적인 원재료가 아니라, 무엇을 식별 가능한 대상으로 보고 어떤 관계를 학습할 것인지를 미리 조직한다. 크로퍼드와 트레버 페글렌(Crawford & Paglen, 2021)은 *Excavating AI*에서 ImageNet의 사람(Person) 카테고리가 인종·국적·직업·도덕성에 이르는 2,833개의 하위 분류를 포함하며, 그 모든 분류 층위에 정치·이데올로기·편견의 흔적이 스며들어

있다고 지적한다. 이들이 강조하는 것은 데이터셋이 알고리즘을 위한 원재료가 아니라 그 자체가 정치적 개입이며, 그 분류의 인식론적 가정들이 역사적 시각 분류 체계와 연속선상에 놓여 있다는 점이다. 그 위에서 학습된 시각 모델은 그러한 분류 체계의 시선을 다시 재생산한다.

조이 부올람위니와 팀니트 게브루(Buolamwini & Gebru, 2018)의 *Gender Shades* 연구는 상업용 얼굴 분류 시스템에서 어두운 피부 여성에 대한 오분류율이 최대 34.7%에 이르는 반면, 밝은 피부 남성에 대한 오분류율은 0.8% 이하라는 사실을 보여준다. 이 격차는 시각 모델의 인식 정확도가 학습 데이터에 표상된 신체와 형태의 분포에 강하게 의존한다는 점을 드러낸다.

따라서 모델의 마찰은 AI 모델이 인간의 해석과 감상을 대체하지 못한다는 사실만을 가리키지 않는다. 그것은 의미와 맥락, 판단, 시각적 분류를 계산 가능한 패턴으로 압축하는 방식이 작품 안에서 어떻게 드러나는가를 묻는 층위다. 분류 결과, 키워드, 확률값, 화면의 병치와 같은 표현 방식은 이러한 환원을 관객이 감각하도록 만드는 디자인적 조건이 된다.

3. 3. 데이터셋의 마찰

데이터셋의 층위에서는 무엇이 수집되고 분류되며 어떤 존재와 차이가 누락되는지가 핵심이 된다. 데이터셋은 단순한 원재료가 아니라 생성형 AI가 세계를 보는 방식을 미리 조직하는 구조다. 어떤 데이터가 포함되고 어떤 데이터가 배제되는지는 기술적 선택인 동시에 정치적 선택이다. 크로퍼드(2021)는 기계학습 시스템을 훈련하는 데이터셋이 단순한 정보 저장소가 아니라 특정한 세계관을 내장한 분류 체계라고 설명한다.

파스퀴넬리의 논의는 이러한 층위를 사회적·역사적으로 보강한다. 자동화가 단순한 기술적 대체가 아니라 노동과 사회적 관계를 추상화하고 재조직하는 과정이라면, 데이터셋은 어떤 세계가 계산 가능하고 학습 가능한 것으로 선언되는지를 보여주는 핵심 장치가 된다(Pasquinelli, 2023, 2024). 따라서 생성형 AI 비평은 그것이 얼마나 정확한가를 묻는 문제를 넘어서, 무엇을 추출하고 어떤 사회적 능력을 어떤 형태로 재조직하는가를 함께 물어야 한다.

데이터셋의 정치성에 대한 이러한 문제의식은 최근 AI 윤리 담론에서도 활발히 다루어져 왔다. 편향(bias), 공정성(fairness), 대표성(representation)의 개념은 머신러닝 시스템이 특정 인구 집단을 누락하거나 왜곡하는 문제를 식별하고 교정하기 위한 분석 도구로 자리 잡았으며, demographic parity나 equalized odds 같은 정량적 공정성 지표와 더 다양한 데이터셋 구축이 해결책으로 제시돼 왔다. 캐서린 디그나지오와 로런 F. 클라인(D'Ignazio & Klein, 2020)은 *Data Feminism*에서 이러한 기술적 교정만으로는 데이터 시스템의 권력 구조를 충분히 다룰 수 없다고 보면서, 데이터 작업을 권력의 검토와 도전, 맥락의 고려, 노동의 가시화를 포함하는 정치적 실천으로 재정의한다.

그러나 이 논문에서 데이터셋의 마찰은 이러한 윤리적·정치적 비판과 초점이 다르다. AI 윤리 담론과 비판적 데이터 연구가 데이터셋의 편향과 누락을 진단하고 그 교정 또는 권력 구조의 재편을 지향한다면, 이 논문은 데이터셋의 분류와 누락이 작품의 형식 속에서 어떻게 감각 가능한 문제로 전환되는지를 분석한다. 마찰은 더 정확한 데이터셋이나 더 공정한 분류 체계를 처방하지 않는다. 오히려 분류 가능한 것과 분류되지 않는 것의 경계 자체가 작품 경험 속에서 어떻게 의식화되는지를 비평의 대상으로 삼는다. 이때 라벨과 범주, 확률값, 이미지의 배열과 빈칸 같은 표현 요소는 데이터셋의 경계를 감각적으로 조직하는 디자인적 조건이 된다.

3. 4. 출력의 마찰

출력의 층위에서는 계산 결과가 어떻게 감각 가능한 형식으로 제시되고, 그 과정에서 어떤 긴장과 잔여가 매끄러운 시각적 표면 속에 봉합되는지가 중요해진다. 생성형 AI의 힘은 종종 이 층위에서 가장 강하게 체감된다. 이미지·영상·텍스트는 더 이상 기계적 결과처럼 보이지 않고 유려하고 장엄하며 몰입적인 형태를

취한다. 바로 이 지점 때문에 출력의 층위는 마찰이 가장 눈에 띄게 드러나기도 하고, 반대로 가장 효과적으로 흡수되기도 하는 자리다.

슈타이얼(2023)은 머신러닝이 만들어내는 시각적 출력물을 통계적 렌더링(statistical renderings)이라고 부르며, 이러한 이미지가 더 이상 사실성의 지표가 아니라 확률과 평균의 시각적 렌더링이라고 지적한다. 그녀가 말하는 “평균 이미지(mean images)”는 실제 세계의 지표적 재현이라기보다 데이터셋과 사회적 필터를 통과한 평균적 경향과 통계적 규범이 감각적 표면으로 현전하는 결과물이다. 이 점에서 생성형 AI의 출력은 단순한 모델 내부 연산의 최종 결과가 아니라, 사회적 편향과 통계적 평균, 확률적 규범이 시각적으로 응축된 장면으로 이해될 필요가 있다(Steyerl, 2023).

출력의 층위는 오늘날 생성형 AI 시스템에서 점차 일회적 결과의 자리에서 벗어나고 있다. 합성 데이터의 본격적 활용은 출력이 다음 세대 모델의 학습 자료로 재흡수되는 구조를 일상화하며, 출력은 시스템의 종착점이 아니라 다음 세대 데이터셋의 일부가 된다. 동시에 오늘날의 생성형 AI에서 출력은 정적 이미지나 텍스트에 머물지 않고, 사용자의 프롬프트와 시스템의 응답이 실시간으로 교환되는 인터페이스의 형태를 취한다. 이 인터페이스에서 사용자는 출력을 일방적으로 수용하는 위치에 머물지 않으며, 자신의 입력을 통해 출력의 흐름에 다시 개입하는 행위자로 편입된다(Galloway, 2012; Pasquinelli, 2023, 2024).

그러나 바로 이 즉시성이 비평적 거리를 봉합하기도 한다. 사용자는 자신의 입력이 출력에 반영된다는 감각 속에서 시스템 조건 자체를 다시 질문하기 어려워질 수 있다. 출력의 매끄러운 인터페이스는 데이터 선택과 알고리즘 선택의 자연화를 가속하며, 시스템의 환원 구조를 비가시화한다. 이때 출력 층위의 마찰은 기술적 오류가 아니라, 매끄러운 표면이 어떤 데이터 분포와 모델의 환원 위에 형성되었는지를 다시 감각하게 하는 형식적 작동으로 나타난다.

모델·데이터셋·출력은 작품을 세 유형으로 나누기 위한 배타적 분류가 아니다. 각 작품은 세 층위 모두와 관계하지만, 특정 층위의 환원과 누락을 더 강하게 전경화한다. 따라서 사례 분석에서는 중심 층위를 출발점으로 삼되, 마찰이 다른 층위와 연결되어 구체화되고 작품의 형식 안에서 다시 흡수되는 과정을 함께 살핀다.

4. 사례를 통한 적용

서론에서 제시한 선정 기준에 따라, 본 장은 AI 시스템의 모델·데이터셋·출력 조건을 작품의 비평적 문제로 구성하는 세 사례를 비교한다. 오주영의 〈BirthMark〉, 언메이크랩의 데이터셋 실천, 레픽 아나들의 〈Unsupervised〉는 각각 모델·데이터셋·출력의 한 층위를 중심으로 전경화한다. 그러나 각 사례의 비평적 작동은 단일 층위에 한정되지 않으며, 다른 두 층위와의 관계 속에서 형성된다.

세 사례가 사용하는 기술도 동일하지 않다. 〈Unsupervised〉가 StyleGAN 계열의 생성 모델을 직접 사용하는 반면, 〈BirthMark〉와 언메이크랩의 작업은 객체 인식과 컴퓨터 비전의 분류 모델을 다룬다. 앞의 두 사례는 생성형 AI 출력에 선행하는 모델과 데이터셋의 환원을 드러내고, 〈Unsupervised〉는 그 조건 위에서 산출되는 출력의 층위를 전경화한다.

세 사례는 비대칭적으로 구성된다. 〈BirthMark〉와 언메이크랩의 실천이 마찰을 작품의 형식 속에서 가시화한다면, 〈Unsupervised〉는 마찰이 물입적인 출력 속에 다시 봉합되는 대조 사례다. 이러한 비대칭은 본 연구의 분석틀이 마찰의 유무나 작품의 비평적 우열을 판정하기보다, 시스템 조건의 가시화와 봉합 방식을 비교하기 위한 것임을 보여준다(Table 2).

Table 2 Analytical alignment of friction layers, key questions, and case studies

분석 층위	핵심 질문	주요 사례	분석 초점
모델	AI는 어떻게 지각과 해석을 계산 가능한 범주로 환원하는가	오주영, 〈BirthMark〉	감상과 계산적 인식 사이의 간극
데이터셋	어떤 세계가 데이터로 선택되고, 분류되며, 누락되고, 다시 구성되는가	언메이크랩, 데이터셋의 실천 궤적	분류 범주의 경계와 부재하는 데이터셋
출력	기술적·정치적 조건은 어떻게 감각적 스펙터를 속에 흡수되는가	Refik Anadol, 〈Unsupervised〉	기술 조건의 부분적 공개와 물입적 표면·서사·기술 권위를 통한 재봉합

4. 1. 오주영의 〈BirthMark〉: 인공지능 관객 모델과 감상의 잔여

작가 오주영의 〈BirthMark〉(2017)는 인공지능 관객 모델을 제안하며 인간의 감상 과정을 위장(camouflage), 해결(solution), 통찰(insight)의 세 단계로 정의한다. 이 작업은 2020년 SIGGRAPH Asia Art Gallery 논문에서 인공지능 관객 모델로 소개되었다(Oh, 2020). 작품에서 AI는 16명 작가의 서로 다른 작품이 담긴 아카이브 영상을 보면서 인간처럼 작품을 감상하려고 시도한다. 이를 위해 YOLO-9000이라는 객체 탐지 시스템(object detection system)이 작품 이미지 속 대상을 추적하고, ACT-R이라는 뇌의 구조를 모방하도록 설계된 인지 아키텍처(cognitive architecture)가 이를 읽고 지각하는 과정을 수행한다.

이때 AI의 인식 과정은 비디오로 시각화되고, AI가 찾아낸 작품의 키워드는 작은 화면에 제시되며, 슬라이드 프로젝트는 작가가 자신의 작업에 대해 제시한 해석을 AI가 의미적으로 처리한 결과를 보여준다. 오주영(Oh, 2020, p. 1)은 AI가 이미지에서 정확하게 분석하는 키워드가 “only 2 to 5 out of 300 words”에 불과하며, 작업이 더 추상적일수록 AI의 이해가능성(intelligibility)이 더 낮아진다고 밝힌다.

그러나 〈BirthMark〉의 비평적 작동은 이러한 작가 진술에 머물지 않으며 작품의 전시 형식 자체에서 발생한다. 〈BirthMark〉의 설치는 네 가지 시각적 층위를 같은 어두운 공간 속에 동시에 펼쳐 놓는다(Figure 2). 중앙의 큰 모니터에는 원본 작품 영상 위에 객체 탐지 박스와 분류 정보가 표시되고, 좌측 보조 모니터에는 AI의 내부 처리 과정이 코드의 흐름으로 나타나며, 우측 보조 모니터에는 그 처리 결과로 추출된 키워드가 제시된다. 그 아래의 슬라이드 프로젝트는 참여 작가들이 자신의 작업에 대해 제시한 해석을 AI가 의미적으로 처리한 결과를 비춘다. 관객은 이 네 층위를 한 번에 본다. 작품은 단일한 해석 화면을 제공하지 않고, 원본 이미지·모델의 탐지 과정·키워드 출력·작가 해석에 대한 AI의 의미 처리 결과를 병렬적으로 배치함으로써 서로 다른 계산적 해석이 일치하지 않는 지점을 비교하게 만든다. 곧 관객의 위치는 동일한 작품이 서로 다른 계산 과정 속에서 각기 다른 정보로 환원되는 양상을 동시에 감각하는 자리로 구성된다.



Figure 2 Oh Jooyoung, *BirthMark*, 2017, installation view
Source: Courtesy of Oh Jooyoung

이 동시적 펼침은 관객 경험에서 모델의 마찰이 발생하는 자리를 만든다. 관객은 AI가 이미지에서 추출한 키워드와 원본 작품 영상, 그리고 작가의 해석을 AI가 의미적으로 처리한 결과를 동시에 비교한다. 이때 슬라이드 프로젝트는 작가의 해석 원본을 직접 제시하는 것이 아니라, 이미지 기반 객체 탐지와는 다른 입력을 모델이 어떻게 의미적으로 처리하는지를 보여주는 또 하나의 계산적 해석 층위를 삽입한다. 관객은 동일한 작품이 객체 탐지·키워드 추출·텍스트 의미 처리라는 서로 다른 과정 안에서 각기 다르게 환원되는 양상을 감각한다. 작품이 드러내는 것은 단순히 AI의 인식률이 낮다는 정량적 사실이 아니라, 작품 감상이 이러한 계산적 처리들의 결합과 완전히 일치하지 않는다는 인식론적 조건이다.

이 모델의 마찰은 다른 두 층위와도 연결된다. 데이터셋 차원에서, <BirthMark>가 사용한 YOLO-9000은 ImageNet과 COCO를 기반으로 학습된 범용 객체 탐지 모델이며(Redmon & Farhadi, 2017), 그 학습 분포는 일상의 사물·인물·동물 카테고리 집중에 있다. 작가의 작품 영상이 더 추상적일수록 AI의 이해가능성이 떨어진다는 관찰은 모델 자체의 한계가 아니라, 모델이 학습한 데이터셋의 분포가 미술 작품의 이미지 세계와 어긋난다는 사실을 가리킨다.

출력 차원에서, AI가 산출하는 키워드는 매끄럽게 다듬어진 시각적 표면이 아니라 빈약하고 거친 텍스트 단편의 형태로 직접 노출된다. 이 점에서 <BirthMark>는 3장에서 제시한 세 층위의 순환을 역방향으로 작동시킨다. 출력의 매끄러움이 모델과 데이터셋의 환원을 자연화하는 대신, 출력의 의도적 거칠음이 모델과 데이터셋의 환원을 감각적으로 노출하는 형식을 취한다.

디자인 비평의 관점에서 이 전시 형식을 다시 읽으면, <BirthMark>의 인터페이스 결정이 생성형 AI 인터페이스 설계의 일반적 최적화 방향과 정확히 반대를 향하고 있음이 드러난다. 일반적인 사용자 경험 설계는 사용자의 인지적 부담을 줄이고 정보의 해석 경로를 명확히 조직하는 방향을 지향한다. 이러한 설계 관점에서라면 코드 흐름을 감추거나, 키워드 출력을 주화면에 통합하거나, 네 층위 사이에 보다 명확한 시각적 위계를 부여하는 선택이 가능했을 것이다. <BirthMark>의 설치 디자인은 이러한 최적화 방향을 따르지 않는다. 네 해석 체계를 단일한 사용자 경험으로 통합하지 않는 이 결정이 바로 모델의 마찰이 감각적으로 발생하는 조건을 만든다. 관객은 순차적 수신자가 아닌 동시적 비교자의 자리로 들어선다.

여기서 디자인 비평의 새로운 물음이 열린다. 생성형 AI 인터페이스를 사용성을 중심으로 최적화하는 과정에서 채택되는 단일하고 매끄러운 출력, 내부 계산의 은폐, 간결한 인터랙션 흐름은 중립적인 설계 결정이 아니다. 그것은 모델의 환원과 데이터셋의 분류를 자연화하는 출력 층위의 통합 장치다. <BirthMark>의 이러한 전시 형식은 그 통합의 메커니즘을 역방향으로 드러내며, 디자인 비평이 AI 인터페이스의 사용성 평가만이 아니라 그 사용성이 무엇을 통합하는지를 함께 읽어야 한다는 점을 작품의 형식 차원에서 입증한다.

따라서 <BirthMark>에서 마찰은 실패한 기술의 결합이 아니다. 감상이라는 행위가 애초에 계산적 인식 구조를 초과한다는 사실이 작품의 다층적 장치와 관객의 동시 비교 경험 속에서 가시화되는 비평적 사건이다. 이 작품은 모델의 마찰을 정량적으로 측정하기보다, 그 마찰이 관객의 감각 속에서 의식화되는 형식 자체를 구축한다.

4. 2. 언메이크랩의 데이터셋 실천: 「인공시점」에서 <시시포스 데이터세트>까지

앞선 오주영의 사례가 단일 작품을 중심으로 모델의 마찰을 드러낸다면, 언메이크랩의 경우 데이터셋 층위의 마찰은 하나의 완결된 결과물보다 기계 시각 연구·분류 체계 비판·데이터셋 구성 행위가 서로 연결된 리서치와 작업의 궤적 속에서 더 선명하게 드러난다. 본 절은 「인공시점」 리서치 프로젝트에서 제기된 기계 시각의 분류 문제를 <시시포스 데이터세트>(2020)의 “부재하는 데이터세트”와 “데이터세-팅하기(dataset-ting)” 개념으로 응축하는 흐름 전체를 하나의 분석 단위로 다룬다.

이 흐름의 출발점은 기계 시각이 대상을 객관적으로 인식한다는 통념에 대한 비판이다. 언메이크랩은 「인공시점」 프로젝트 중 “포즈에 따른 AI 시각 연구”에서 포즈를 바꾸는 동일한 인물이 인공지능 시각 앞에서 “사람”, “댄서”, “골퍼”, “플레이어”, 심지어 “매력적(attractive)”과 같은 값으로 확률적으로 반환되는 장면을 제시하고, 짧은 머리의 여성을 “여성(woman, female)”으로 제대로 분류하지 못하는 사례를 보여준다. 이러한 결과는 단순한 기술적 오류라기보다, 어떤 데이터가 학습되었고 어떤 범주가 누락되었는지를 역으로 노출하는 징후다. 안전고깔(traffic cone)이 위상 변화에 따라 립스틱·램프·지우개로 분류되는 사례 역시 같은 방향을 가리킨다. 컴퓨터 비전의 분류 결과는 사물의 객관적 인식이 아니라 학습 데이터셋의 분포와 분류 체계의 한계가 작품 표면에 드러나는 자리다. 이러한 오분류 결과는 기계 시각이 세계를 어떤 범주와 확률 분포 안에서 조직하는지를 관객에게 시각적으로 드러낸다(Unmake Lab, 2021).

이 문제의식은 〈시시포스 데이터셋〉에서 작품의 형식 자체로 응축된다(Figure 3). 작품은 등글지만 부서진 외곽을 가진 돌 이미지들을 흰색 배경 위에 균일한 조명과 그림자 처리로 촬영하고 격자 형태로 배열한다. 이 시각적 형식은 우연이 아니다. 흰 배경, 객체의 분리, 균일한 조명과 격자 배열은 대상을 개별 표본으로 분리해 제시하는 머신러닝 데이터셋의 시각적 관습을 차용한다. 그러나 그 형식 안에 놓인 돌들은 기존의 분류 범주에 쉽게 환원되지 않는 잔여적 형태들이다. 이들은 자연 표본처럼 보이지만 동시에 인간의 추출과 개발 과정에서 변형된 흔적을 가진 객체들이기도 하다. 작품은 데이터셋의 형식을 빌리면서 그 형식이 무엇을 분류 가능한 것으로 받아들이고 무엇을 받아들이지 않는지를 동시에 노출한다.



Figure 3 Unmake Lab, Sisyphus Dataset, 2020

Source: Courtesy of Unmake Lab

언메이크랩은 이 작업을 둘러싸고 “부재하는 데이터셋”과 “데이터세-팅하기(dataset-ting)”라는 두 개념어를 직접 제시한다. 작가는 〈시시포스 데이터셋〉을 구성하는 행위가 부재하는 데이터셋을 구성하는 행위를 통해 데이터셋 만들기를 과거의 기억에 대한 비판적 유추와 재구성의 질문으로 만들어 미래로 되먹임하는 과정이 되며, 이것이 곧 데이터세-팅하기의 과정이라고 설명한다(Unmake Lab, 2021).

이 개념적 배치에는 제프리 코헨(Jeffrey Jerome Cohen)의 『Stone: An Ecology of the Inhuman』(2015)이 중요한 참조점으로 기능하며, 둘은 인간 중심의 서사와 욕망 안에서 끊임없이 변형되는 일반자연의 표본이자 그 변형의 폭력을 되묻는 질문의 단위로 다뤄진다. 이때 부재하는 데이터셋은 더 정확한 분류로 채워져야 할 누락이 아니라, 분류 체계 자체가 무엇을 인식 가능한 것으로 미리 결정하는지를 노출하는 음각적 형식이다. 데이터세-팅하기는 데이터셋을 더 잘 만드는 기법이 아니라, 데이터셋을 만드는 행위 자체를 비평적 질문의 형식으로 전환하는 방법론이다.

이 점에서 〈시시포스 데이터셋〉은 AI 윤리 담론이 데이터셋 문제에 대해 제시하는 처방과 분명히 구별된다. 데이터셋 편향에 관한 많은 실천적·정책적 논의가 누락을 보완하고 보다 공정한 분류 체계를 구축하는 데 초점을 둔다면, 〈시시포스 데이터셋〉은 분류 가능한 카테고리를 확장하지 않으며 더 정확하거나 더 공정한 데이터셋을 처방하지도 않는다. 작품은 분류 가능한 것과 분류되지 않는 것의 경계 자체가 작품 경험 속에서 어떻게 의식화되는지를 형식적으로 구축한다. 데이터셋의 마찰은 바로 이 자리에서 발생한다. 그것은 더 나은 데이터셋의 진단이 아니라, 데이터셋이라는 형식 자체가 무엇을 받아들이고 무엇을 누락하는지를 작품의 시각적 형식 속에서 다시 의식화하는 사건이다.

데이터 표현 설계의 관점에서 보면, 〈시시포스 데이터셋〉가 차용한 흰 배경, 객체의 분리, 균일한 조명과 격자 배열은 대상을 기계 판독 가능한 표본으로 제시하는 설계된 표현 형식이다. 이 형식은 세계가 분리되고 분류 가능한 범주들로 조직될 수 있다는 전제를 시각적으로 자연화한다. 작품은 이러한 관습을 차용하면서 그 안에 기존 범주로 쉽게 환원되지 않는 둘의 형상을 배치함으로써, 데이터 표현이 무엇을 인식 가능한 대상으로 구성하고 무엇을 잔여로 남기는지를 드러낸다. 이는 데이터셋의 미적 관습이 기술적 필연이 아니라 특정한 인식론을 내장한 디자인 결정임을 보여준다.

이 데이터셋의 마찰은 모델과 출력 차원과의 연결된다. 모델 차원에서, 「인공시점」 리서치가 노출한 분류 오류는 모델 자체의 한계가 아니라 모델이 학습한 데이터셋의 분류 체계가 어떤 세계를 인식 가능한 것으로 미리 조직했는가의 문제로 귀결된다. 출력 차원에서, 〈시시포스 데이터셋〉가 보여주는 것은 매끄러운 생성 출력이 아니라 데이터셋 그 자체의 형식으로 응축된 출력이다. 작품은 생성형 AI 시스템이 산출할 매끄러운 출력을 우회하고 그 출력의 전제가 되는 데이터셋의 형식으로 직접 들어가, 출력의 매끄러움이 어떤 환원 위에 서 있는지를 데이터셋 차원에서 미리 노출한다. 이 점에서 언메이크업의 실천은 생성 출력의 표면으로 나아가기보다, 그 전제가 되는 데이터셋의 형식과 구성 행위 자체를 비평적 사유의 대상으로 전환한다.

4. 3. 레픽 아나돌(Refik Anadol)의 〈Unsupervised〉: 마찰의 흡수와 매끄러운 출력

MoMA가 2021-2023년 작품으로 수록하고 있는 〈Unsupervised—Machine Hallucinations—MoMA〉는 본 논문의 세 사례 가운데 가장 대조적인 위치를 차지한다(Museum of Modern Art, n.d.-b; Figure 4). 이 작업은 MoMA의 공개 컬렉션 데이터와 디지털 이미지 아카이브를 바탕으로 구축되었다. Refik Anadol Studio(n.d.)에 따르면, 작품은 MoMA 컬렉션의 이미지 138,151건을 처리하고, StyleGAN2 ADA 모델을 활용해 생성한 잠재 공간(latent space)을 자체 개발한 Latent Space Browser로 탐색하는 방식으로 제작되었다. 이렇게 산출된 시각적·청각적 흐름은 MoMA 건드 로비(Gund Lobby)의 빛·움직임·음향과 외부 날씨 변화에 반응하며 실시간으로 변화한다(Museum of Modern Art, n.d.-a). Refik Anadol Studio는 작품을 설명하며 “기계의 마음이 MoMA의 방대한 컬렉션을 ‘본’ 뒤에 무엇을 꿈꾸는가(What would a machine mind dream of after ‘seeing’ the vast collection)”라는 질문을 전면에 내세우고, AI의 작동을 “환각(hallucinations)”, “자기 재생성하는 놀라움(self-regenerating element of surprise)”, “사이버네틱 우연을 통한 새로운 감각적 자율성(new form of sensorial autonomy via cybernetic serendipity)”과 같은 어휘로 호명한다(Refik Anadol Studio, n.d.).

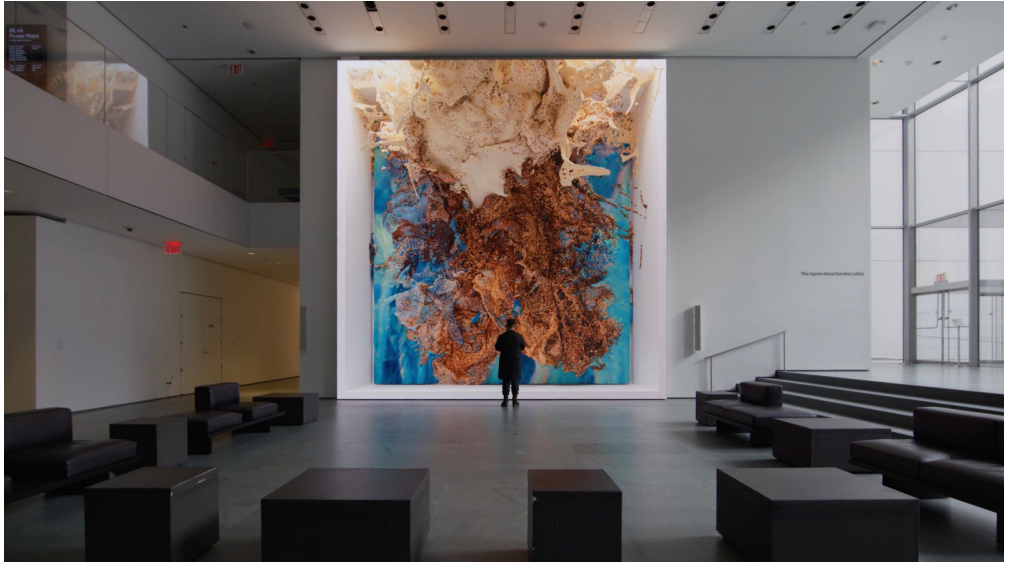


Figure 4 Refik Anadol, *Unsupervised — Machine Hallucinations* — MoMA, 2021–2023, installation view at MoMA, 2022

Source: Refik Anadol Studio, retrieved May 20, 2026

이 어휘들은 단순한 마케팅 수사가 아니라 작품의 작동 방식을 규정하는 수사적 환경을 구성한다. 작품의 매끄러운 출력 표면에 대한 의문—어떤 데이터가 어떻게 분류되었는지, 어떤 미학적·정치적 선택이 그 출력을 가능하게 했는지, 출력의 어떤 가능성이 전경화되고 어떤 가능성이 배제되었는지—은 AI를 의인화된 창작 주체로 호명하는 어휘 안에서 미리 봉합된다. “기계가 꿈꾼다”는 명제는 출력의 책임 소재를 인간 행위자들의 선택과 결정에서 의인화된 시스템 자체로 이전하며, 이 이전을 통해 출력에 대한 비평적 질문은 시스템의 “자율성”과 “환각”이라는 어휘 속으로 흡수된다. 출력의 매끄러움은 기술적 효과이면서 동시에 수사적 효과다.

이러한 수사적 봉합은 작품에 대한 기술 연구 공동체의 평가 발언이 공식 페이지에 배치되면서 한 번 더 강화된다. 〈Unsupervised〉 공식 페이지는 NVIDIA Research의 수석 연구원 야코 레티넨(Jaakko Lehtinen)과 UMAP 개발자 릴랜드 매키니스(Leland McInnes)의 평가 발언을 나란히 제시하며, 작품의 미학적 성취에 기술 연구 공동체의 권위를 부여한다(Refik Anadol Studio, n.d.).

이 배치는 작품을 머신러닝 연구의 예술적 응용 사례로 정당화하는 동시에, 데이터 선택과 알고리즘 구성, 시각화 결정에 대한 비평적 질문을 기술적 혁신의 서사 안에 위치시킨다. 마찰이 흡수되는 자리는 출력의 시각적 표면 하나가 아니라, 출력과 그 출력을 둘러싸는 서사적·제도적 장치들의 결합이다.

R. H. 로신(R. H. Lossin)은 e-flux Criticism에 발표한 글에서 이 작품을 “기술 부스터리즘의 전형(tech boosterism at its best)”이라고 비판한다(Lossin, 2023). 그에 따르면 이 작업은 예술사를 사회적 의미를 지닌 상징 실천의 아카이브가 아니라 기계가 추출한 시각적 특징들을 재배열하고 결합할 수 있는 요소들의 집합으로 다룬다. 그 결과 기술 부스터리즘, 감시 체계, 환경적으로 고비용인 계산 과정은 장엄한 시각성 속에서 봉합된다.

이 비판은 본 연구의 분석틀에 따라 다음과 같이 해석할 수 있다. 〈Unsupervised〉는 모델과 데이터셋의 기술적 조건을 공식 설명 자료에서 부분적으로 가시화하지만, 이러한 공개가 매끄러운 시각적 표면, 의인화된 자기서사, 기술 연구 공동체의 권위 부여와 결합하면서 다시 봉합되는 양가적 사례다. 로신의 비판이 작품의 장엄한 시각성과 기술 우월주의, 그 배후의 제도적·정치적 조건을 전면화한다면, 본 연구는 이를 기술 조건의 부분적 가시화와 재봉합이 동시에 작동하는 메커니즘으로 다시 조직한다.

이러한 통합은 3장에서 논의한 출력의 즉시성이 실시간 반응성의 형식으로 구현되는 장면에서도 확인된다. 〈Unsupervised〉는 정적인 이미지가 아니라 로비의 빛·움직임·음향과 외부 날씨 변화에 반응하며 끊임없이 변하는 시각적·청각적 흐름으로 제시된다. 이 구성은 현장의 환경과 출력이 직접 연결되어 있다는 감각을 형성하고, 관객을 분석적 관찰자보다 변화하는 화면에 둘러싸인 몰입적 관람자의 위치에 놓는다. 관객은 출력에 영향을 미치는 조건의 일부로 자신을 인식할 수 있지만, 이러한 참여의 형식이 데이터와 알고리즘의 구성 조건을 함께 드러내지는 않는다. 오히려 출력의 연속성과 반응성이 전면화되면서, 그 출력을 가능하게 한 선택과 환원은 배경으로 물러난다.

이러한 출력의 통합은 모델과 데이터셋 차원의 통합을 동시에 수행한다. 모델 차원에서, 어떤 이미지 하위집합이 학습에 사용됐고 그 이미지들이 어떤 잠재 공간의 구조로 변환되었는지, 그 과정에서 어떤 미학적 가능성이 전경화되었는지는 유동적인 출력의 장엄함 속에서 질문되기 어렵다. 데이터셋 차원에서, MoMA 컬렉션이 형성되어 온 수집·분류·디지털화의 역사와 그 안의 불균형은 “기계의 꿈”이라는 수사 뒤로 물러난다. 컬렉션은 특정한 제도적 선택의 결과라기보다 기계가 자유롭게 탐색하는 보편적 시각 기억의 저장소처럼 제시된다.

〈Unsupervised〉 분석은 마찰의 통합이 출력의 시각적 표면에만 한정되지 않으며, 그 출력을 둘러싸는 의인화된 서사와 기술 연구 공동체의 권위 부여를 통해 다층적으로 강화된다는 점을 보여준다. 따라서 출력 층위의 분석은 화면에 나타난 결과뿐 아니라 이를 정당화하는 수사적·제도적 장치까지 포함해야 한다.

국내의 아나돌 연구는 이러한 긴장을 다른 방식으로 포착해 왔다. 정규리(Chung, 2023)는 아나돌의 작업을 데이터 아트(Data Art)와 데이터 프레젠테이션(Data Presentation)의 차이를 보여주는 융합예술 형식으로 논의하고, 박태성(Park, 2024)은 방대한 학습 데이터와 알고리즘의 구성, 설명 가능한 인공지능의 필요성을 중심으로 분석한다. 이러한 연구가 아나돌 작업의 융합적 가능성과 기술적 설명 가능성에 초점을 둔다면, 본 연구는 출력의 시각성과 이를 둘러싼 수사적·제도적 장치가 기술적 조건을 어떻게 자연화하는지에 주목한다.

이 사례는 모든 AI 기반 미디어아트가 자동적으로 비평적 마찰을 생성하는 것은 아니며, 마찰의 가시화 여부가 기술 사용 자체보다 그 사용을 조직하는 감각의 질서와 수사적·제도적 장치에 달려 있음을 보여준다.

5. 결론

본 연구는 생성형 AI 시대의 AI 기반 미디어아트 비평을 위한 분석틀로서 마찰 개념을 정식화했다. 생성형 AI가 자동화·예측·확률적 최적화를 통해 매끄러운 경험을 생산하는 조건에서, 그러한 표면은 계산적 전제와 데이터의 선택 및 배제, 감각의 정치적 재배치를 자연화한다. 또한 이러한 조건 속에서 마찰을, 보편화하는 시스템과 구체적 현실의 접촉, 인터페이스 내부의 매개 작동, 감성의 분할이 교차하면서 시스템의 조건이 감각적으로 드러나는 비평적 사건으로 재정의했다. 그리고 이를 모델·데이터셋·출력의 세 층위로 구조화하여, 마찰이 어디에서 발생하고 어떻게 가시화되거나 통합되는지를 분석하기 위한 틀을 제안했다.

세 사례 분석은 마찰이 작품의 형식 안에서 서로 다른 방식으로 가시화되거나 통합되는 양상을 보여준다. 오주영의 〈BirthMark〉는 원본 작품 영상, AI의 객체 탐지 과정, 키워드 출력, 작가의 해석을 AI가 의미적으로 처리한 결과를 동시에 펼치는 다층적 장치를 통해 관객이 서로 다른 계산적 해석과 감상 사이의 간극을 시지각적으로 감각하는 자리를 구축했다. 언메이크업의 데이터셋 실전은 머신러닝 데이터셋의 시각적 관습을 차용하면서 그 안에 기존 분류 범주로 쉽게 환원되지 않는 잔여적 형상을 배치함으로써 분류 체계의 경계를 작품의 형식으로 드러냈다. 반면 아나돌의 〈Unsupervised〉는 공식 설명을 통해 모델과 데이터셋의 일부 기술

조건을 공개하면서도, 매끄러운 시각적 표면, 의인화된 자기서사, 기술 연구 공동체의 권위 부여가 결합하면서 그 조건을 다시 배경으로 밀어내는 양가적 사례로 기능했다. 이는 기술 조건의 공개 자체가 비평적 마찰을 보장하지 않으며, 그 공개가 어떠한 인터페이스·서사·제도적 장치와 결합하는가에 따라 다시 봉합될 수 있음을 보여준다.

이 점에서 본 연구가 제안하는 마찰 개념의 의의는 새로운 비평 어휘를 추가하는 데 그치지 않는다. 마찰은 미디어아트 비평의 초점을 기술 사용의 참신성에 대한 평가나 융합예술의 가능성 진단에서 이동시켜, 동시대 AI 기반 미디어아트가 무엇을 비평하는 실천일 수 있는가를 다시 묻게 한다. 데이터셋 편향에 관한 많은 실천적·정책적 논의가 누락을 보완하고 분류의 공정성을 높이는 데 초점을 둔다면, 본 연구의 마찰 개념은 그러한 선택과 배제의 조건이 작품의 감각적 표면에서 다시 사유되는 자리를 가리킨다. 융합예술 연구가 기술과 예술의 결합 가능성을 평가한다면, 본 연구의 분석틀은 그 결합이 시스템 조건을 가시화하는지, 아니면 매끄러운 관람 경험 안에 다시 봉합하는지를 구분한다.

연구를 진행하며 검토한 국내 선행연구가 생성형 AI의 창작 활용 패턴(Ro & Lee, 2026), 아나돌 작업의 융합예술적 가능성(Chung, 2023), 데이터와 알고리즘의 기술적 구성 및 설명 가능성(Park, 2024)에 주목했다면, 본 연구는 이와 다른 방향에서 AI 기반 미디어아트의 비평적 기능을 설명하기 위한 개념을 정식화했다. 이는 기존 연구를 대체하기보다, 기술 활용과 작품 형식의 분석을 시스템 조건의 가시화와 봉합이라는 비평적 질문에 연결하는 시도다.

이러한 분석틀은 미디어아트 비평을 넘어 디자인 연구로 확장되는 접점을 갖는다. 본 논문이 2장에서 다룬 HCI·인터랙션 디자인의 생산적 마찰은 사용자의 행동을 늦추거나 성찰을 유도하기 위해 의도적으로 삽입되는 인터랙션 장치를 가리킨다. 이와 달리 본 연구의 마찰은 행동을 설계하는 전략이 아니라, 시스템의 조건이 감각적·인식론적으로 가시화되는 비평적 사건을 의미한다. 두 개념은 생성형 AI의 매끄러운 표면에 비판적으로 대응한다는 문제의식을 공유하며, 이 접점에서 본 연구의 분석틀은 생성형 AI 기반 인터페이스와 시각문화를 읽기 위한 디자인 비평의 개념적 자원으로 기능할 수 있다.

〈BirthMark〉와 〈시시포스 데이터세트〉의 분석은 이러한 접점이 어떻게 작동하는지를 구체적으로 보여준다. 〈BirthMark〉에서는 인터페이스 최적화의 관습을 거스르는 다층적 정보 배치가 모델의 환원 과정을 감각적으로 노출한다. 이는 사용성이나 정보 전달의 효율만으로는 충분히 설명되기 어려운 층위다. 〈시시포스 데이터세트〉에서는 흰 배경, 객체의 분리, 균일한 조명과 격자 배열 같은 데이터 표현의 시각적 관습이 특정한 인식론적 전제를 내장한다는 점이 드러난다. 이는 데이터의 품질이나 공정성 기준만으로는 충분히 포착되기 어려운 층위다. 본 연구는 새로운 UX 방법론이나 인터랙션 설계 지침을 제안하지 않는다. 이 연구의 디자인 연구적 기여는 미디어아트 사례에서 출발해, 생성형 AI 기반 시각문화와 인터페이스를 비판적으로 읽기 위한 개념적 틀을 제안한 데 있다.

다만 본 연구에는 한계가 있다. 개념 정식화에 초점을 둔 만큼 사례 분석은 세 사례에 한정되며, 각 작품에 대한 관객 반응을 경험적으로 조사하지 않았다. 따라서 관객 경험에 관한 논의는 작품의 설치 구성과 인터페이스가 관객을 위치시키는 방식에 대한 해석에 한정된다. 또한 모델·데이터셋·출력의 세 층위는 AI 시스템을 완결적으로 설명하는 고정된 분류체계가 아니라, 작품의 비평적 작동 지점을 구분하기 위한 잠정적 분석틀로 이해되어야 한다. 후속 연구에서는 이 틀을 더 다양한 AI 기반 미디어아트 사례에 적용하여 그 비평적 유효성과 한계를 검토할 필요가 있다. 이를 디자인 교육과 비평 교육의 맥락으로 확장하여, AI 시대의 이미지와 인터페이스, 데이터셋을 어떻게 읽고 가르칠 것인지 탐구하는 작업도 중요한 과제로 남는다.

그럼에도 본 연구는 하나의 분명한 제안을 남긴다. 생성형 AI가 점점 더 자연스럽고 유려한 표면을 통해 세계를 조직할수록, 미디어아트의 비평적 역할은 그 표면을 단순히 거부하거나 깨뜨리는 데만 있지 않다. 그 역할은

표면 아래에 봉합된 이질성과 잔여, 누락과 배제를 감각 가능한 문제로 다시 드러내는 데 있다. 마찰은 이러한 비평적 드러내기를 분석하기 위한 개념이며, AI 기반 미디어아트가 무엇을 비평하는 실천일 수 있는지를 다시 사유하게 하는 출발점이다.

References

1. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). ACM. <https://doi.org/10.1145/3442188.3445922>
2. Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. *Proceedings of Machine Learning Research*, 81, 77–91.
3. Burrell, J. (2016). How the machine "thinks" Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12. <https://doi.org/10.1177/2053951715622512>
4. Chun, W. H. K. (2006). *Control and freedom: Power and paranoia in the age of fiber optics*. MIT Press.
5. Chung, C. (2023). 인공지능 시대의 융합 예술: 레픽 아나돌(Refik Anadol)의 작품을 중심으로 [Convergence Art in the Age of Artificial Intelligence: Focusing on the Works of Refik Anadol]. *The Korean Society of Science & Art*, 41(3), 383–394. <https://doi.org/10.17548/ksaf.2023.06.30.383>
6. Cohen, J. J. (2015). *Stone: An ecology of the inhuman*. University of Minnesota Press.
7. Cox, A. L., Gould, S. J. J., Cecchinato, M. E., Iacovides, I., & Renfree, I. (2016). Design frictions for mindful interactions: The case for microboundaries. In *Proceedings of the 2016 CHI Conference Extended Abstracts on Human Factors in Computing Systems* (pp. 1389–1397). ACM. <https://doi.org/10.1145/2851581.2892410>
8. Crawford, K. (2021). *Atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. Yale University Press.
9. Crawford, K., & Joler, V. (2018). *Anatomy of an AI system: The Amazon Echo as an anatomical map of human labor, data and planetary resources*. AI Now Institute & Share Lab. <https://anatomyof.ai>
10. Crawford, K., & Paglen, T. (2021). Excavating AI: The politics of images in machine learning training sets. *AI & Society*, 36(4), 1105–1116. <https://doi.org/10.1007/s00146-021-01162-8>
11. D'Ignazio, C., & Klein, L. F. (2020). *Data feminism*. MIT Press.
12. Dunne, A., & Raby, F. (2013). *Speculative everything: Design, fiction, and social dreaming*. MIT Press.
13. Galloway, A. R. (2012). *The interface effect*. Polity.
14. Han, J., & Choi, Y. (2026). AI 기반 인터랙티브 미디어아트에서 의도된 기술적 불완전성으로 인한 관객의 미적 경험 연구 [Exploring Aesthetic Experience through Intended Technical Imperfection in AI-Based Interactive Media Art]. *Archives of Design Research*, 39(1), 375–392. <https://doi.org/10.15187/adr.2026.02.39.1.375>
15. Lossin, R. H. (2023, March). Refik Anadol's "Unsupervised." *e-flux Criticism*. <https://www.e-flux.com/criticism/527236/refik-anadol-s-unsupervised>
16. Mejtoft, T., Hale, S., & Söderström, U. (2019). Design friction. In *Proceedings of the 31st European Conference on Cognitive Ergonomics* (pp. 41–44). ACM. <https://doi.org/10.1145/3335082.3335106>
17. Museum of Modern Art. (n.d.-a). *Refik Anadol: Unsupervised*. Retrieved May 20, 2026, from <https://www.moma.org/calendar/exhibitions/5535>
18. Museum of Modern Art. (n.d.-b). *Refik Anadol. Unsupervised – Machine Hallucinations – MoMA. 2021–2023*. Retrieved May 20, 2026, from <https://www.moma.org/collection/works/442077>
19. Oh, J. (2020). BirthMark: An artificial viewer for appreciation of digital surrogates of art. In *Proceedings of the SIGGRAPH Asia 2020 Art Gallery* (Art. 8, pp. 1–1). ACM. <https://doi.org/10.1145/3414686.3427146>

20. Park, T. (2024). 데이터와 인공지능이 만든 융합예술: 레픽 아나돌(Refik Anadol)의 작품 분석 [Convergence Art Created by Data and AI: A Study of Refik Anadol's Works]. *The Korean Society of Science & Art*, 42(5), 151–162. <https://doi.org/10.17548/ksaf.2024.12.30.151>
21. Pasquinelli, M. (2023). *The eye of the master: A social history of artificial intelligence*. Verso.
22. Pasquinelli, M. (2024). Theories of automation from the industrial factory to AI platforms: An overview of political economy and history of science and technology. *Tecnoscienza: Italian Journal of Science & Technology Studies*, 15(1), 99–112. <https://doi.org/10.6092/issn.2038-3460/20010>
23. Rancière, J. (2004). *The politics of aesthetics: The distribution of the sensible* (G. Rockhill, Trans.). Continuum.
24. Redmon, J., & Farhadi, A. (2017). YOLO9000: Better, faster, stronger. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 6517–6525). <https://doi.org/10.1109/CVPR.2017.690>
25. Refik Anadol Studio. (n.d.). *Unsupervised – Machine Hallucinations – MoMA*. Retrieved May 20, 2026, from <https://refikanadol.com/works/unsupervised/>
26. Ro, H., & Lee, Y. (2026). 국내 미디어아트 창작에서의 생성형 AI 활용 패턴 연구 [Utilization Patterns of Generative AI in Media Art Creation in South Korea]. *Archives of Design Research*, 39(1), 395–410. <https://doi.org/10.15187/adr.2026.02.39.1.395>
27. Sengers, P., Boehner, K., David, S., & Kaye, J. (2005). Reflective design. In *Proceedings of the 4th Decennial Conference on Critical Computing: Between Sense and Sensibility* (pp. 49–58). ACM. <https://doi.org/10.1145/1094562.1094569>
28. Steyerl, H. (2023). Mean images. *New Left Review*, 140/141, 82–97. <https://doi.org/10.64590/uhm>
29. Tsing, A. L. (2005). *Friction: An ethnography of global connection*. Princeton University Press.
30. Unmake Lab. (2021). *아포페니아, 그리고 시시포스 데이터셋 [Apophenia and Sisyphus Dataset]*. In *미술관 없는 사회, 어디에나 있는 미술관 [Living in the Postdigital, Reliving the Museum]* (NJP Reader #10, pp. 103–120). Nam June Paik Art Center.
31. Zylinska, J. (2020). *AI art: Machine visions and warped dreams*. Open Humanities Press.

생성형 AI 시대 미디어아트의 비평을 위한 마찰 개념: 모델·데이터셋·출력 층위의 교차적 분석

허대찬^{1,2*}

¹국민대학교 조형대학 시각디자인과, 강사, 서울, 대한민국

²미디어 문화&예술 채널/컬렉티브 엘리스온 편집장, 서울, 대한민국

초록

연구배경 오늘날 생성형 인공지능은 자동화·예측·확률적 최적화를 통해 점점 더 자연스럽고 매끄러운 시각적·상호작용적 경험을 생산한다. 사용자는 텍스트와 이미지를 생성하고 시스템의 응답을 받아들이는 과정에서 내부의 계산 절차, 데이터의 구성 방식, 모델의 한계와 편향을 충분히 의식하지 않은 채 출력 결과와 마주하기 쉽다. 최근 예술·기술 연구에서는 데이터셋, 합성 미디어, 프롬프트, 기술의 정치성과 같은 의제가 활발히 논의되어 왔지만, 이를 AI 기반 미디어아트가 수행하는 비평적 기능과 연결하여 작품의 형식과 관객의 위치를 분석하는 개념적 틀은 상대적으로 부족하다.

연구방법 본 연구는 생성형 AI 시대의 미디어아트 비평을 위한 분석틀로서 마찰(friction) 개념을 제안한다. 이를 위해 안나 칭의 마찰 개념, 알렉산더 R. 겔러웨이의 인터페이스 이론, 자크 랑시에르의 감성의 분할 개념을 교차적으로 재구성하고, 크로퍼드와 윌러의 AI 시스템에 대한 해부학적 관점을 참조하여 모델·데이터셋·출력을 세 분석 층위로 정식화한다. 이때 마찰은 단순한 저항이나 기술적 실패의 은유가 아니라, AI 시스템의 환원과 누락, 선택과 배제의 조건이 작품의 형식과 인터페이스를 통해 감각적으로 가시화되거나 다시 봉합되는 비평적 사건으로 정의된다. 세 층위는 AI 시스템을 완결적으로 설명하는 고정된 분류체계가 아니라, 각 작품에서 비평적 작동이 두드러지는 지점을 구분하기 위한 분석 층위로 사용된다.

연구결과 오주영의 <BirthMark>는 원본 작품 영상, AI의 객체 탐지 과정, 키워드 출력, 작가의 작품 해석을 AI가 의미적으로 처리한 결과를 병렬적으로 제시함으로써 서로 다른 계산적 해석과 예술 감상 사이의 간극을 시지각적으로 드러낸다. 언메이크랩의 데이터셋 실천은 머신러닝 데이터셋의 시각적 관습을 차용하고 그 안에 기존 분류 범주로 쉽게 환원되지 않는 잔여적 형상을 배치함으로써, 데이터셋이 무엇을 인식 가능한 대상으로 구성하고 무엇을 누락하는지를 가시화한다. 반면 레픽 아나돌의 <Unsupervised>는 공식 설명을 통해 모델과 데이터셋의 일부 기술 조건을 공개하면서도, 몰입적인 시각적 표면, 의인화된 자기서사, 기술 연구 공동체의 권위가 결합된 출력 형식을 통해 그 조건을 다시 배경으로 밀어내고 마찰을 매끄러운 관람 경험 속에 재봉합하는 양가적 사례로 기능한다.

결론 본 연구는 마찰을 작품의 비평적 가치를 판정하는 평가 기준이 아니라, AI 시스템의 조건이 어디서 가시화되고 어떻게 봉합되는지를 분석하기 위한 비평적 분석틀로 제안한다. 이를 통해 AI 기반 미디어아트 비평을 기술 사용의 참신성이나 융합 가능성에 대한 평가에서 시스템의 환원·누락·매개 조건에 대한 분석으로 확장한다. 또한 본 연구의 분석틀은 생성형 AI 기반 시각문화와 인터페이스, 데이터 표현을 비판적으로 읽기 위한 디자인 비평과 디자인 교육의 개념적 자원으로 활용될 수 있다.

주제어 마찰, 생성형 인공지능, AI 기반 미디어아트, 인터페이스, 데이터셋

*교신저자 : 허대찬 (s2016511@kookmin.ac.kr)