



저작권법에서 본 데이터 윤리

김 윤 명*

Abstract

저작권법은 창작자의 권리를 보호하는 동시에 공정한 이용을 보장하여 문화·산업 발전의 균형을 도모한다. AI 기술의 발전으로 인해 대량의 데이터가 활용되면서, 저작권법과 데이터 윤리 사이의 충돌이 심화하고 있다. 특히, 공정이용(Fair Use)과 텍스트·데이터 마이닝(Text and Data Mining, TDM)의 법적 한계는 AI 학습 과정에서 중요한 쟁점으로 부각되고 있다. AI가 방대한 데이터를 학습하는 과정에서 저작권법이 어떻게 적용될 것인가에 대한 논쟁이 지속되고 있으며, 공정이용이 AI 학습 데이터에 적용될 수 있는 범위에 대한 명확한 기준이 요구된다. 또한, 데이터 윤리는 AI 모델의 학습과 결과물의 공정성을 보장하는 중요한 요소로 작용하지만, 윤리적 고려가 법적 규제와 어떻게 조화를 이루어야 하는지는 여전히 불확실하다.

이 글은 저작권법과 데이터 윤리의 관계를 분석하고, AI 학습을 위한 데이터 활용이 법적으로 허용될 수 있는 범위를 검토하며, 법적·윤리적 방안으로써 제시되는 표시, 게시중단, 머신 언러닝 등 여러 제도나 이에 연계된 기술에 대해 검토한다. 이를 통해 AI 기술 발전과 저작권의 접근성 사이의 균형을 유지하면서도, 법적 명확성을 강화하는 방향을 모색하고자 한다.

I. 저작권법이 고민해야 할 윤리

1. 데이터윤리의 필요

데이터(data)는 지구상에 존재하는 모든 것이다. 그 것이 기록된 문자나 이미지, 또는 음성이나 영상이라고 하더라도 데이터로서 존재한다. 데이터인 순간 누구라도, 어떤 용도라도

* 디지털정책연구소 소장, digitallaw@naver.com

쉽게 이용이 가능하다. 데이터의 속성은 저작물성, 개인정보성, 영업비밀성, 사실 정보(fact)와 같은 퍼블릭도메인(public domain) 등 다양하다.¹⁾ 그 다양한 만큼, 이해관계는 복잡하며, 관련된 법률이나 보호체계도 상이하다. 이런 면이 법적 측면에서 데이터에 대한 관계를 보다 복잡하게 하고 있다. 대표적인 관련 법률로는 저작물로서 데이터는 저작권법, 개인정보로서 데이터는 개인정보 보호법, 영업비밀인 경우는 영업비밀보호법을 들 수 있다. 데이터산업 전반적으로는 데이터 산업진흥 및 이용촉진에 관한 기본법(이하, '데이터산업법'이라 함)이, 공공영역의 데이터와 관련해서는 공공데이터법 등 다양한 법률이 관련된다. 다양한 법률이 존재함에도 데이터 수집, 가공, 이용 등 일련의 과정을 제대로 확인할 수 없다는 점에서 데이터의 윤리적인 체계를 고려하지 않을 수 없다. 그러한 고려 없이 데이터를 이용할 경우, 향후 관련 분쟁은 자명할 것이기 때문이다.²⁾ 데이터 윤리는 데이터를 사용함으로써 발생할 수 있는 여러 가지 문제를 포함하여 데이터를 수집, 가공, 이용 등의 처리과정에서 윤리적인 문제에 대응하기 위한 체계로 볼 수 있다.

데이터 윤리가 필요한 이유는 다음과 같다. 데이터 윤리는 개인정보 보호, 데이터 사용, 알고리즘 사용, 데이터 관리 및 공개 등과 같은 데이터와 관련된 윤리적 문제를 다룬다. 데이터는 디지털전환 시대에서 매우 중요한 자산이며, 그것이 우리 삶의 거의 모든 측면을 지배하고 있다. 데이터 윤리는 인공지능 및 기계학습 알고리즘을 개발하고 사용할 때 공정성과 투명성을 유지하는 데 중요하다. 이것은 인공지능이 인간의 편견을 반영하지 않고 공정하게 작동하도록 보장하는 것을 의미한다.

또한, 데이터 윤리는 기업의 책임과도 관련이 있다. 기업은 고객의 데이터를 보호하고 고객의 신뢰를 유지해야 한다. 불법적인 데이터 수집, 유통 및 판매는 기업의 이미지를 손상시킬 수 있기 때문이다. 따라서, 데이터 윤리는 개인, 기업 및 사회적 측면에서 중요한 역할을 한다. 데이터 윤리에 대한 이해와 적절한 적용은 인공지능의 신뢰성을 보장하고, 나아가 기업의 신뢰도와 이미지를 유지하는 데 도움이 될 것이다.

2. 문제제기: 데이터 편향의 문제

인간의 편향성이 AI에도 전이될 수 있다는 점은 AI의 신뢰로 이어질 수 있다는 점에서

1) 데이터에 대한 물건성은 인정되지 않는다. 민법 제98조에서 물건을 유체물 및 전기 기타 관리할 수 있는 자연력으로 정의하고 있기 때문이다. 형법도 마찬가지로, 정보의 물건성을 부정하는 대법원 판례가 있었다. 다만, 민법상 관리가능성 여부에 따라 무형재에 대해서도 물건성을 인정할 수 있다는 견해도 있다.

2) 2024년, 독일에서는 TDM을 긍정하는 판결이 내려지기도 했다. OpenAI와 MS 등에 대해 저작권 침해를 이유로 하는 일련의 소송이 진행 중이다. 반면, 피고는 공정이용을 근거로 반론을 펴고 있다.

해결해야 할 문제이다. 알고리즘은 데이터 없이 작동하기 쉽지 않으며, 데이터를 통한 학습 과정을 거치지 않을 경우, 효율적이고 효과적인 결과를 만들어 내기 어렵다.³⁾ 문제는 기계 학습 과정에서 사용되는 데이터가 현재 상황을 반영한다기 보다는 과거 정보를 담고 있다는 점이다. 이는 과거의 차별적인 관습과 문화가 데이터에 쌓여있기 때문에 학습과정이 객관적이고 공정하다고 하더라도, 결과는 편향될 수밖에 없는 한계를 갖는다.⁴⁾

인공지능 기술은 대규모 데이터에 대한 기계의 자가 학습(self-learning)을 통하여 지속적으로 알고리즘 성능을 강화하므로 데이터를 산업의 새로운 경쟁 원천으로 부각시킴으로써 기존 산업구조의 근본적 변화를 가져오고 있다. 예컨대, 스스로 데이터를 확보할 수 있는 생태계를 구축하고 이를 활용할 수 있는 알고리즘을 보유한 기업이 시장을 주도하고 많은 이윤을 창출할 수 있다. 따라서, 지능정보사회에서는 지능정보기술을 활용한 산업이 보다 많은 사용자를 플랫폼 기반 생태계에 참여하게 함으로써 데이터를 지속적으로 생성, 활용하는 구조가 형성되면서 승자독식의 플랫폼 경쟁을 심화시키는 가운데 새로운 성장의 기회도 창출할 것이다.⁵⁾

다만, 데이터 확보는 플랫폼사업자나 자본력을 가진 경우가 아니라면 쉽지 않기 때문에 데이터 독점에 대한 문제를 해결하기 위한 인류의 공통된 노력이 필요하다. 무엇보다, 플랫폼사업자의 독점화에 따른 이슈로 이용자의 데이터를 기반으로 데이터를 이용하는 사업자들은 데이터 주권이 아닌 데이터 독점을 피하고 있기 때문이다. 공정이용이 아닐 경우, 데이터 접근이 제한될 수 있다. 그러면, 편향에 대한 문제는 더욱 확산될 수 있을 것이다. 기계학습에서의 공정이용의 가용성을 제한하는 것이 사회 정의에 바람직하지 않은 결과를 초래할 수 있다.⁶⁾ 데이터 편향은 지속되고, 안정적이고 예측가능한 제도설계를 위해 입법론에 대한 고민이 필요하다.

저작권법은 창작자의 권리를 보호하면서도, 공정한 이용을 도모하는 것이다. 이로써, 문화 및 산업의 향상발전이라는 사회적 이익을 균형 있게 도모하는 것이다. 저작권법을 포함

3) 따라서, “머신러닝을 통하여 인공지능을 제대로 학습시키기 위해서는 기존 훈련데이터가 특정한 성별, 지위, 단체, 사건유형과 관련하여 편향되거나 오류가 담기지 않는 것이 전제되어야 한다.”, 고유강, 법관업무의 지원을 위한 머신러닝의 발전상황에 대한 소고, LAW&TECHNOLOGY 제15권 제5호(통권 제83호), 2019, 13면.

4) “알고리즘은 자료의 편향성을 판별하지 않기 때문에, 오히려 자료에 내재한 편향성을 그대로 학습하는 것이다. 이러한 문제는 활용하는 자료가 많아진다고 하더라도 해소되지 않는다”고 한다. 허유선, 인공지능에 의한 차별과 그 책임 논의를 위한 예비적 고찰 - 알고리즘의 편향성 학습과 인간 행위자를 중심으로, 한국여성철학 제29권, 2018, 90면.

5) 이원태 외, 지능정보사회의 규범체계 정립을 위한 법제도 연구, 2016, 28~29면.

6) Sobel, Benjamin, Artificial Intelligence's Fair Use Crisis(September 4, 2017). 41 Colum. J.L. & Arts 45(2017), Available at SSRN: <https://ssrn.com/abstract=3032076>, p.95.

한 법률이 윤리적 딜레마를 완벽하게 반영하지는 못한다. 다만, 법적 허용 범위 내에서 윤리적으로 문제가 될 수 있는 행위들이 존재할 수 있다는 점에서 여지를 남겨두는 것이다. 대표적으로 공정이용(Fair Use) 법리는 법적으로는 허용되는 예외로써, 특정 상황에서는 권리침해를 면책한다. AI 윤리나 데이터 윤리가 저작권법에 지나치게 깊이 개입할 필요는 없겠지만, 윤리적 고민이 법의 해석과 향후 인간이 아닌 저작자의 등장에 따른 법개정에 영향을 미칠 가능성은 높다. AI가 저작권 체계를 흔들 가능성도 배제할 수 없는 상황에서 법이 기술 변화에 적응하기 위해서는 윤리적 논의가 필수적이다. 그 논의 방향은 윤리를 배제하는 것이 아니라 법과 함께 균형 있게 고려하는 것이야 한다. 다만, 법적 판단이나 책임이 윤리 뒤로 숨는 모양새는 바람직하지 않다.

저작권법은 데이터 윤리에 대해 어떤 입장인지 확인이 필요하다. 그렇지만, 인공지능 윤리가 강조되는 이유는 간단하다. 윤리는 법적 책임이 없다는 점에서 부담을 완화할 수 있다는 매력적인 수단으로 활용하는 것이다. 데이터 출처, 데이터의 편향성 등 책임에서 어느 정도 자유로울 수 있기 때문이다. 그렇지만, AI윤리는 그 범위나 시간에서 있어서도, 상당히 제한적이다. EU AI 법을 비롯하여, 미국의 주법 및 우리나라의 인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(이하, 'AI 기본법'이라 함) 등 각국에서 입법이 이루어지고 있기 때문이다. EU는 입법 과정에서 저작권 처리에 대해서는 기존 법제의 적용을 기본으로 하고 있다.⁷⁾ 이를 위하여 저작권법에서 규정하는 제반 규정과 실무적으로 논란이 되고 있는 사항에 대해 데이터 윤리와 그에 따른 법적 검토를 통해 AI 모델과 시스템에 대한 신뢰성을 확보할 수 있는 방안에 대해 살펴보고자 한다.

II. 저작권으로서 데이터의 가치

전통적인 생산의 3요소인 토지, 노동, 자본에서 이제는 데이터가 포함된 4요소로 확장하고 있다. 데이터가 그만큼 생산에 있어서 중요한 가치재가 있다는 의미이다. 이하에서는, 저작물 내지 저작물을 가공하여 만들어진 데이터가 어떠한 가치가 내재하는지 살펴보고자 한다. 이는 데이터의 가치평가와는 다르며, 데이터의 가치에 대한 법적 접근을 넘어 윤리적인 고려까지 확장하고자 하는 취지이다.

7) EU AI 법 제53조 (c) 유럽연합 저작권 관련 법 및 관련 권리를 준수하고 특히 최신기술을 통한 것을 포함하여 지침 (EU) 제2019/790호 제4조(3)에 따라 표시된 권리 유보를 확인하고 준수하기 위한 정책을 시행한다.

1. 기술적 상상으로서의 존재와 데이터

가. 기술적 상상과 존재

존재하는 모든 것을 0과 1로 조직화할 수 있다면, 그 모든 것은 데이터(data)로 볼 수 있다. 인터넷에 존재하는 많은 것은 우리가 숨쉬며 살고있는 현실사회와 더불어 2중의 사회이기는 하지만 이미 가상사회는 디지털화된 또다른 사회이기도 하다. 좀 더 정확히 하자면 아바타(avatar)인 나를 중심으로 하는 메타버스(metaverse) 사회이다. 기술적 상상으로서 메타버스 사회에서 존재하는 모든 것은 데이터로 구성된다는 점에서 데이터 기반 사회라고 해도 과언이 아니다. 자아를 나타내는 아바타가 공연을 하거나 다양한 창작활동을 한다.⁸⁾ 이러한 과정에서 또다른 데이터의 집합을 이루게 된다.

인간을 인간답게 하는 상상하기를 나 아닌 타자가 인식할 수 있는 현실로 보여준다면, 가상사회와 현실을 분리할 수 있을까? 뇌 임플란트(brain implant) 기술이 발전과정에 있다.⁹⁾ 사람의 생각과 꿈속의 모습이 바로 디지털로 전환된다면, 사람의 의식적, 무의식적 생각은 데이터로 나타낼 수 있다. 생각하는 모든 것을 현실에서 볼 수 있다는 의미이다. 물론, 아직은 꿈을 뿌옇게 디지털화하는 수준이지만, 언젠가는 선명하게 보게 될 것이다. 그렇다면, 그 꿈은 인간의 창작적 의지가 반영된 것인가? 꿈속의 영상을 디지털화할 경우, 그 영상은 창작물이 될 수 있을까? 아니면 무엇으로 취급할 수 있을까? 꿈속의 나는 또다른 자아로 볼 수 있는 것이 아닌가 하는 생각도 들 수 있다. 또다른 자아로 성장 또는 발전하고 있는 인공지능과는 어떤 차이가 있을까? 의식 속에 존재하는 상상이나 생각이 타자에게 비추어지는 것은 철학적 이슈를 담고 있음은 별론으로 하더라도, 존재하는 것이 아닌 것이 존재하는 것으로 나타나는 현실은 법률적으로 풀어야 할 과제가 될 것이다.

나. 만들어진 데이터

인공지능이 인간을 대신할 수 있는가? 이 질문에 답하기가 쉽지 않다. 역사적으로 볼 때, 상상은 언젠가 기술이 발달함으로써 실현되었다. 인공지능 암흑기(AI winter)가 도래했던

8) 90년대, 경희대학교의 라이언(LION; Love Is On the Net)이라는 사이버 대학생을 중심으로 각 대학에서는 가상의 대학생을 선보인 바 있다. 그렇지만, 그 당시 기술력과 어떻게 활용할 것인지에 대한 상상력 부재로 자취를 감추었다. 이제는 상상력과 이를 구현할 수 있는 기술력의 뒷받침으로 신한라이프의 22살 '로지(ROZY)'와 같은, 또다른 인간이 탄생하고 있다. 24시간 활동할 수 있는 사이버 인간의 활동모습은 인공지능의 모습과도 겹친다.

9) 뇌임플란트(brain implant)는 인간의 뇌에 전자 장치를 이식하여 신경 신호를 모니터링하거나 조작하는 기술을 말한다. 주로 신경과학, 생의학, 인공지능(AI) 등의 기술이 결합된 최첨단 분야로, 의료 및 비의료(향상 기술) 목적으로 연구되고 있다.

것도 아이디어에 대한 기술적 뒷받침이 어려웠기 때문이었다. 이제는 인공지능 암흑기를 극복하고 있다는 점에서 기술적 한계는 어느 순간 극복되기 마련임을 알 수 있다.

인간적인 사고와 의식을 인공지능이 가질 수 있을지에 대한 많은 논란이 있는 것도 사실이다. 그렇다면, 인간의 상상력을 인공지능이 대신할 수 있을까? 보다 풍성하게 할 가능성도 높다. 지금까지 상상하는 능력은 생각하는 능력과 연관된 인간의 능력으로만 알고 있었다. 인공지능에 대해 논하기 전에 오래전 읽었던 문덕수 시인의 ‘시쓰는 법’에서 인간만의 능력으로 간주되는 상상이 어떠한 의미를 갖는지 살펴본다.¹⁰⁾

“상상은 과거에 보고 듣고 겪었던 사물의 이미지를 마음속에서 다시 생각해 내는 일이다. 지난날의 일을 회상하는 것을 흔히 기억이라고 하는데, 이 점에서 상상은 기억과 같다. 미국의 심리학자 윌리엄 제임스(William James)는 재생적 상상과 생산적 상상을 구분하는데, 전자는 과거의 기억을 그대로 이미지화하는 것을 말하며, 후자는 지난날에 겪었던 이미지들에서 선택된 여러 가지 요소들을 결합하여 새로운 이미지로 만들어내는 것을 말한다.”

인공지능은 데이터에 기반하고, 데이터는 우리의 역사가 반영된 결과물이다. 그렇기 때문에 인공지능이 데이터를 학습하는 과정은 과거에 ‘만들어진’ 데이터를 바탕으로 새로운 가치를 형성한다는 점에서 앞서 살펴본 상상적 이미지 또는 변형적 이용(transformative use)¹¹⁾과 유사하다. 과거의 데이터이기 때문에 현재와 일치하지 않을 수 있다는 점이 가장 큰 문제이다. 즉, 과거의 기준이 현재에 적용됨에 따라 현실성이 떨어질 수 있다는 점, 과거의 가치와 현재의 사회문화적 가치가 일치하지 않을 수 있다는 점 등으로 인하여, 인공지능이 내보이는 결과가 편향적이라는 것이다. 이에, 데이터는 현재의 데이터라기보다는 과거의 축적된 산물이라는 점에서 인종이나 사회적·문화적 차별을 그대로 담고 있어, 기계학습 과정에서 데이터에 담겨있는 편향을 알고리즘이 학습하는 우를 범할 수 있다. 즉, 데이터 자체도 중립적이지 않거니와 데이터를 학습하는 과정에서 축적된 편견이 무의식적으로 알고리즘에 반영되는 결과를 가져올 수 있기 때문이다. 또한, 학습데이터를 선택하는 과정에서 인간의 편견이 반영될 수 있다. 이처럼 모든 단계에서 인간의 편향이 의도적이든 의도적이지 않든 반영된다는 점에서 데이터의 편향성은 기본적인 문제이다.

아울러, 인간의 영역으로만 여겨졌던 창작활동이 이제는 인공지능이 저작물을 이용하고,

10) 문덕수, 시쓰는 법(重版), 동원출판사, 1984, 118~119면.

11) Campbell v. Acuff-Rose Music, Inc., 510 U.S. 569(1994).

그 결과로써 창작물을 만들어내는 지적 활동으로 이어지고 있다. 인공지능이 보다 개선되고 바람직한 로직과 판단을 내릴 수 있도록 학습시키게 되며, 인공지능의 기계학습(machine learning) 과정에서 무엇보다 중요한 것이 ‘데이터’이며,¹²⁾ 산업화시대의 중요한 자원(resources)이었던 원유에 비교되기도 한다.

2. 혁신을 이끄는 요소로서 데이터

가. 저작권법은 기술혁신을 저해하는가?

법이나 제도가 혁신을 저해하는 경우를 보통 규제라고 한다.¹³⁾ 규제는 그 필요에 따라 이루어진 경우가 적지 않다. 해당 기술이 가져올 수 있는 파급력을 확인할 수 없는 경우에도 이루어지기도 한다. ‘미지의 기술’이 가져올 파장 또한 예측하기 어렵기 때문에 규제라는 방패를 사용한다. 물론, 이러한 상황이나 또는 규제를 회피하거나 극복하는 과정에서 노력이 더해지는 경우, 혁신이 더해지는 것도 사실이다. 그렇기 때문에 법이나 규제가 기술혁신을 저해한다고 단정하기는 어렵다. 다만, 규제라는 것은 정부의 정책적 결단이기 때문에 법적 근거가 있어야 한다. 보통 정부 규제의 일반적인 근거조항은 기본권 제한의 법률유보를 규정하고 있는 헌법 제37조 제2항을 제시하고 있으나 반대로 기본권 보장을 위한 규제의 근거로 작용하기도 한다.¹⁴⁾

인공지능이 세상을 이끌어가는 사회, 그 인공지능의 중심에는 알고리즘이 위치한다.¹⁵⁾ 알고리즘은 문제해결을 위한 방법, 즉 해법(solution)이다. 존재하는 많은 문제들을 해결할 수 있는 방법으로서 알고리즘은 다양한 영역에서 구현되고 있다. 수학 분야에서는 방정식이나 함수와 같은 방법으로 구현돼 왔고, SW 분야에서는 프로그래밍 언어를 통하여 컴퓨터에 적합한 방식으로 구현돼 왔다. 저작권법은 알고리즘, 즉 해법에 대하여 명시적으로 보호범위에 포함되지 않는다고 규정하고 있다.¹⁶⁾ 다만, 알고리즘은 저작권법상 보호범위에서 제외되는 해법이기는 하나 그것이 구체적으로 표현된 결과물이라면 달리 볼 수 있다.

12) Mark A. Lemley, Bryab Casey, Fair Learning(2020), at <http://ssrn.com/abstract=3528447>.

13) 박균성, 정책, 규제와 입법, 박영사, 2022, 37면.

14) 김유환, 행정법과 규제정책, 법문사, 2012, 17면.

15) 알고리즘이 갖는 다양한 법률 문제에 대해서는 김운명, 알고리즘과 법, 정보화진흥원, 2019.

16) 저작권법 제101조의2(보호의 대상) 프로그램을 작성하기 위하여 사용하는 다음 각 호의 사항에는 이 법을 적용하지 아니한다.

1. 프로그램 언어: 프로그램을 표현하는 수단으로서 문자·기호 및 그 체계
2. 규약: 특정한 프로그램에서 프로그램 언어의 용법에 관한 특별한 약속
3. 해법: 프로그램에서 지시·명령의 조합방법

인공지능 기술은 SW 기술과 HW 기술의 발전이 공진화하면서 발전해 왔다. SW 기술이 갖는 특성은 급작스럽기 보다는 점진적이고 누적적이기 때문에 인공지능 또한 그러한 과정을 거쳐 왔다고 해도 과언이 아니다.¹⁷⁾ 물론, 인공지능을 구현하는 대부분이 SW나 HW와 융합된 SW이기 때문에 SW 기술에 대한 법적인 검토와 이해는 인공지능에 대한 이해의 기반이 된다. 전파기술은 방송이라는 정보유통의 혁신을 이끌어왔고, 마그네틱 기술은 정보의 저장과 배포에 또 다른 혁신을 이끌었다. 이 과정에서 공정이용(fair use)이라는 제도적 혁신이 강구되기도 했다.

그동안 기술혁신과 저작권법의 응전은 인간이 중심이었으나 이제는 인간이 매개되거나 인간을 위한 것이 아닌 유형이 대두되고 있다. 즉, 인공지능이 저작권법에 도전하고 있는 상황이다. 인공지능은 산업은 물론, 사회를 변화시키고 있다. 기술혁신은 점진적인 기술발전 과정에서 유의미하게 나타나곤 한다. 그 과정에서 기술발전에 따라 의도하지 않게 직접적인 규제로서 작용하는 법이 저작권법이라는 점은 역설적이다.¹⁸⁾

나. 데이터 확보를 위한 저작권법의 역할론

인공지능의 발전을 위해 필요한 기계학습의 기본이 되는 데이터의 확보에 있어서도 저작권법은 적지 않은 관련을 맺고 있다. 각국은 기계학습에 필요한 데이터 확보를 위해 저작권법을 개정하고 있으며, EU에서는 지침을 제정하여 데이터 확보에 필요한 입법을 정비한 바 있다. 우리나라도 ‘인공지능사회’¹⁹⁾를 이끌어 나가기 위해서라도, 법제정비가 필요한 시점이다.²⁰⁾ 무엇보다, 데이터 확보를 하는 경우나 확보된 데이터를 이용하여 기계학습을 하는 경우에 저작권을 침해하는 것인지 논란이 적지 않기 때문에 이에 대한 입법적 해결이 필요하다. 물론, 공정이용 규정을 통해서도 가능하다. 다만, 좀더 명확하게 그 기준을 제시하는 것도 의미있다. 물론, TDM 규정을 두고 있는 나라는 공정이용에 관한 일반조항을 두고있지 않는 점도 고려될 필요가 있다.²¹⁾ 역사적으로 다양한 기술이 저작권법의 변화를 가

17) 김윤명, 발명의 컴퓨터 구현 보호체계 합리화를 위한 특허제도 개선방안 연구, 특허청, 2014, 5면.

18) 규제와 권리행사는 상대적이다. 서로 다른 입장에서의 차이라고 보는 것이 타당하다.

19) AI 기본법에 따르면 인공지능사회란 인공지능을 통하여 산업·경제, 사회·문화, 행정 등 모든 분야에서 가치를 창출하고 발전을 이끌어가는 사회를 말한다.

20) 20대 국회에서 박정희원이 빅데이터 확보를 위한 발의한 저작권법 개정이유를 “빅데이터를 활용하기 위해서는 저작물 등 기존 자료를 처리하여 필요한 정보를 분석하는 것이 필수적이나 현행법은 저작자의 이익을 부당하게 해치지 않는 범위 내에서 저작물의 공정한 이용이 가능하다고 포괄적으로 규정하고 있을 뿐 정보를 분석하는 과정에 저작물 등을 복제할 수 있는지 명시적으로 규정하고 있지 않아 빅데이터 활용의 활성화가 어려운 실정”이라고 밝히고 있다. 2017.12 발의(의안번호, 2010698)

21) 같은 생각으로는 남형두, 공정이용의 역설, 경인문화사, 2025, 276면.

져왔다. 인공지능사회에서 사회혁신을 이끌기도 한 저작권법을 어떠한 방향으로 개정해야 할 지에 대한 입법자의 고민이 필요한 시점이다.²²⁾

3. 문화 및 산업 발전을 위한 데이터

가. 저작권법의 산업적 역할

저작권법에서 데이터를 고민하는 주된 이유 중 하나는 데이터베이스에 대한 규정에 기인한다. 데이터산업법에서 데이터란 다양한 부가가치 창출을 위하여 관찰, 실험, 조사, 수집 등으로 취득하거나 정보시스템 및 소프트웨어 진흥법 제2조 제1호에 따른 소프트웨어 등을 통하여 생성된 것으로서 광(光) 또는 전자적 방식으로 처리될 수 있는 자료 또는 정보를 말한다. 이러한 데이터 하나하나의 집합물로서 데이터베이스라는 점이다. 데이터라는 소재의 선택과 배열에 있어서 데이터베이스에 관한 것이지만, 실상 데이터의 가치를 이해함에 있어서 데이터베이스 보다는 개별 데이터가 중요하다. 데이터산업법을 관할하는 부처가 데이터에 대한 산업적인 정책을 주도하고 있다는 점도 고려될 필요가 있다. 다만, 정책이라는 것은 관련 부처의 협력을 통해서 시너지를 낼 수 있다는 점에서 부처간 협력은 중요한 과제이다.

저작권법은 인공지능이라는 기계를 학습시키기 위한 데이터, AI 모델 및 이를 활용한 AI 시스템, 그리고 시스템이 만들어내는 결과물의 저작물성 여부 등 AI 생애주기와 관련있다. 이에 저작권법은 산업적인 측면과 문화적인 측면을 아우르는 법률로서 역할을 하며, 인공지능시대에 중요한 법률임에는 틀림없다.

나. 저작권법 목적규정에 대한 검토

저작권법은 문화와 산업의 진흥을 궁극적인 목적으로 하고 있다. 생산의 3요소에 더하여, 데이터를 4요소로 보기도 한다. 저작권법은 AI를 어떤 역할을 할 수 있는가? 기계학습을 통해 구축되는 AI 모델은 다양한 데이터가 학습된 구조이다. 학습이 끝난 AI 모델은 데이터 그 자체는 아니다.

데이터를 이용하여 구축된 AI 모델과 이를 특정 목적과 용도에 맞게 파인튜닝(fine-tuning)한 AI 시스템은 직접적으로 이용자와 관계한다. 누구라도 AI 시스템을 활용할 수 있게 되었다. 개발자의 개발도 민주화하고 있지만, 이용자도 민주화의 혜택을 받고 있다.

22) 2024년 AI 기본법 제정과정에서 문화부는 학습데이터 공개입장을 고수한 바 있다.

산업적 측면에서 경쟁의 제한이 없어 누구라도 창작의 주체로서 위치가 지워질 수 있다. 이러한 점은 저작권 산업에 있어서도 긍정적이다. AI가 인간의 영역을 침탈한다는 주장도 가능하지만, 모든 영역은 아니다. 창작의 영역에서 AI는 인간의 보조적 수단으로써 활용되기 때문이다. 결국, 저작권법이 AI와 그 학습을 위한 데이터에 대한 입장은 그 흐름에 벗어나지 않은 방향이어야 한다는 점이다.

저작권법은 기술의 도래에 따른 긴장관계의 형성에 있었고, 그 궁극적인 목표는 해당 서비스에 대한 섷다운이었다. 그렇지만, 그것이 온전히 사회적인 부담을 덜 수 있을지, 저작권자에게 돌아가는 혜택만큼 이용자에게 돌아갈 수 있는 것인지에 대한 사리분별은 필요하다. 무수한 섷다운이 이루어졌지만, 그 섷다운은 오래가지 못했다. 문화적인 향유라는 것이 의도하든 그렇지 않든, 사회를 움직이는 힘이 있기 때문이다. AI도 사람이 관여하는 기술적 사회적인 산물이다. 합리적인 선택을 이끌어 낼 수 있는 것은 기술이 아닌 사람의 몫이기 때문이다. 이를 뒷받침할 수 있는 제도, 그리고 일관성을 유지할 수 있는 입법은 무엇보다 중요한 사회적 기술이다.

Ⅲ. 데이터 윤리와 저작권법의 접점

1. 데이터 윤리의 역할

데이터 윤리는 데이터의 수집, 가공 및 활용 과정에서 발생하는 윤리적 문제를 다루는 개념으로, 개인정보 보호, 데이터 접근성, 알고리즘 공정성 및 투명성과 같은 다양한 요소를 포함한다. AI 학습에서 데이터 편향성 문제가 심화되면서, 데이터의 윤리적 이용이 더욱 중요한 화두가 되고 있다. 저작권법은 창작자의 권리를 보호하면서도 공정한 이용을 보장하는 법적 틀을 제공한다. 그러나 AI 학습 데이터가 저작권 보호를 받는 경우, 저작권법이 데이터 윤리를 방해하는 요소로 작용할 수 있다. 따라서 AI 학습 과정에서 공정이용이 어느 정도 허용될 수 있는지에 대한 법적 해석이 중요하다.

가. 인공지능 윤리

인공지능을 선의로 활용하기 위해서는 인공지능 자체, 인공지능의 개발, 기계학습에 필요한 데이터셋(data set) 등 관련 분야에서 사회적 가치와 충돌할 수 있는 사항들을 합리적으로 처리할 필요가 있다. 쉽게 생각하는 것이 규제일 수 있으나, 획일적인 규제는 기술진보를 저해하기 마련이다. 기술을 살리고, 원래 의도대로 사용하려면 인공지능에 대한 윤리에

집중할 필요가 있다. 윤리는 최소한의 도덕으로서 법이 가져올 확실성을 대신하여 유연한 연성법(soft law)으로 대응할 수 있다는 긍정적인 의미도 있기 때문이다.

윤리의 한계도 명확하다. 인공지능이 인식해야 할 윤리가 무엇인지 의문이기 때문이다. 인간에게도 어려운 윤리를 인공지능에 학습시킬 수 있을 것인가? 자율주행차를 중심으로 논의되고 있는 인공지능 윤리는 어떠한 모습이어야 하는가? 한 단계 높은 철학적 논의에서 말하는 정의란 무엇인가? 이러한 질문을 포함하여, 인류가 최고선이라고 하는 윤리적 가치도 최고 수준의 가치라고 단정하기 어렵다. 예를 들면, 법적으로 살인은 금지되나 전쟁에서, 또는 생명을 위협하는 경우에는 정당방위로서 허용되기도 한다. 만약, 로봇이나 인공지능이 가져야 할 윤리가 수립된다고 가정하더라도, 인공지능만이 윤리적이어야 하는가? 이를 기획하고, 개발하고 때론 이용하는 인간은 어떠해야 하는가? 국가는 어떠해야 하는가? 정답은 없다. 다만, 지금 내리는 선택이 우리의 미래세대에 영향을 미칠 것이라는 점이다.

인공지능의 안전성, 공정성, 투명성 확보를 위해 윤리가 대안으로서 제시되지만, ‘어떻게?’라는 물음에 명확히 답하기는 쉽지 않다. 설계 단계에서 어떻게 윤리적이어야 하는지, 공정성을 담보할 수 있을 지에 대한 가이드라인 수준의 논의가 이루어지고 있기는 하지만, 단순하게 ‘로봇공학 3원칙’이 아닌 구체적이고 프로그래밍 가능한 또는 학습 가능한 가이드라인이 제시될 필요가 있다. 무엇보다, 인공지능이 가져야 할 기본적인 명제는 ‘인간을 위하여 존재해야 한다’는 점이다. 인공지능이 인간의 규범의 적용 없이 별개로 활동하는 것은 공존이 아닌 충돌을 가져올 수밖에 없기 때문이다. 이러한 명제를 달성하기 위해서는 윤리적인 원칙도 중요하지만, 이를 강제해야 할 법적인 수단도 필요하다.

나. AI 윤리와 구별되는 데이터 윤리

AI 윤리와 데이터 윤리는 서로 밀접한 관련이 있지만, 목적과 범위가 다르기 때문에 구별될 수 있다. 데이터 윤리는 주로 데이터 수집, 저장, 처리 및 공유와 같은 데이터 활동에서 발생하는 윤리적 문제에 중점을 둔다. 이는 데이터의 정확성, 투명성, 개인정보 보호, 접근성 및 미디어의 독점적인 소유와 같은 문제를 다룬다. 예를 들어, 데이터 윤리의 한 예로는 대용량 데이터베이스를 사용하여 개인의 인종, 성별, 출신국가, 종교, 성적 취향 등을 추측하는 것과 같은 문제가 있다. 이러한 데이터 수집과 처리는 인권 침해 및 사회적 공정성 문제를 야기할 수 있다.

반면, AI 윤리는 인공지능 시스템의 설계, 개발, 구현, 운영 및 사용에 대한 윤리적 문제에 중점을 둔다. 이는 인공지능의 효과성, 공정성, 투명성, 안전성 및 책임성과 같은 문제

를 다룬다. 예를 들어, AI 윤리의 한 예로는 자율주행차가 인간의 생명을 위협할 수 있는 상황에서 누구에게 책임이 있는지와 같은 문제가 있다. 이러한 문제는 인공지능 시스템의 사용과 구현에 대한 윤리적 고민을 유발한다.²³⁾ AI 윤리와 데이터 윤리는 서로 관련이 있지만 각각 다른 범위와 목적을 가지고 있다. AI와 데이터를 사용하는 모든 사람은 둘 다 고려해야 할 중요한 윤리적 문제이며, 이를 고려하여 적절한 결정을 내리는 것이 중요하다.

데이터 처리(수집, 가공, 분석) 및 이용과정의 문제로서 나타나는 것은 데이터의 편향이다. 데이터 편향성은 데이터 자체의 편향성을 포함하여, 데이터를 어떤 의도에 따라 수집하거나 가공, 분석, 분류하거나 이렇게 처리된 데이터를 재처리하는 과정에서도 발생할 수 있는 문제이다. 처리과정의 복잡성에 따라 그 과정에서 의식적이거나 무의식적으로 처리자의 의도가 반영될 수 있기 때문에 데이터의 수집과정에서 나타나거나 또는 발생할 수 있는 문제에 대한 정리가 필요하다.

데이터 수집과정에서는 데이터의 수집행위가 공정한지 여부가 문제될 수 있으며, 타인의 권리를 침해하는 것은 지양되어야 한다. 데이터 학습을 위해 저작물을 사용하는 행위가 저작권법상 공정이용에 해당하는지 여부도 검토될 필요가 있다.²⁴⁾ 이와 별개로, 일본, 영국, 독일 등 각국은 데이터 확보를 위한 경우에는 저작권재산권의 제한규정을 도입하여 해결하고 있다.

데이터 가공 과정에서는 비정형 데이터를 정형화하여 데이터베이스를 제작하는 것과 같은 행태로 보이며, 저작권법상 데이터베이스 제작에 상당한 투자가 이루어진 경우에는 별도 법적 보호가 발생할 수 있는 과정이다. 문제는 데이터 가공 과정에서 의도성이 반영될 수밖에 없으며, 이 과정에서 차별이나 공정성을 해하는 경우가 나타날 가능성이 크다는 점이다. 이처럼, 처리과정에서 데이터의 편향성을 의도하는 것은 데이터가 갖는 영향력을 확대 재생산하는 것으로, 개인에 대한 데이터 편향을 확대하는 것으로 이해할 수 있다.²⁵⁾ 따

23) 트롤리딜레마에 대한 다양한 논의가 이루어지고 있으나, 이러한 상황자체를 만들지 않도록 자율주행차의 윤리적 설계가 중요하다. 기계가 통신이 가능하다면 인간의 주행통제보다 훨씬 안전한 상황에서 자율주행이 가능할 것이기 때문이다.

24) “훈련데이터 자체에 내재된 편향성뿐만 아니라, 데이터를 수집하고 인공지능을 활용한 예측모델을 설계하는 사람의 주관적인 판단(가령 어떠한 요소에 가중치를 얼마나 둘지)이나 막대한 양의 훈련데이터를 통계적으로 분석하기 위하여 가공하는 과정에서 개별 데이터가 가지는 특징이 소실되어 왜곡될 우려 등도 경계해야 한다.”고 지적된다. 고유강, 법관업무의 지원을 위한 머신러닝의 발전상황에 대한 소고, LAW&TECHNOLOGY 제15권 제5호(통권 제83호), 2019, 13면.

25) 이러한 이유는 “알고리즘 설계자의 편향성이 의도적 또는 무의식적으로 반영되어 데이터 마이닝의 결과가 왜곡될 수 있기 때문이다. 예를 들어, 개인정보에 관한 빅데이터가 알고리즘 설계자가 갖고 있는 편향성에 따라 분석됨으로써 특정 개인에 대한 프로파일링이 왜곡될 수 있는 것이다. 그러므로 자신의 개인정보와 관련을 맺는 데이터 마이닝

라서, 데이터 분석 과정에서는 데이터가 갖는 속성, 의미 등을 추려냄으로써 데이터의 통계적 사용을 하게 되는 것으로 데이터 전문가(data scientist)의 역할이 무엇보다도 중요하다.

이처럼, 데이터 윤리는 AI 학습을 위한 데이터의 수집, 저장 및 이용 과정에서 발생하는 윤리적 문제를 포함한다. AI 윤리는 기술 자체의 윤리성을 논의하는 반면, 데이터 윤리는 데이터 활용 방식의 윤리성을 논의하는 것이다. AI가 학습하는 데이터에 윤리적 문제가 내재되어 있다면, AI의 결과물도 공정성을 확보하기 어려울 것이다.

다. 저작권법과 데이터 윤리

법률은 사회현상의 모든 것을 담아내기 어렵다. 사회현상이라는 것이 가변적이기 때문에 입법적으로 해결하는 것이 때론 시기를 놓치거나 유연한 대응을 막을 수 있다. 그렇기 때문에 일반조항을 통해, 보편타당한 방안을 제시하고자 하는 것이 입법적 결단이기도 하다. 일례로 공정이용의 요건을 갖출 경우, 저작권 침해임에도 불구하고 면책하겠다는 것이다. 물론, 판단이전에 수범자에게 대략적인 기준을 제시함으로써 이러한 요건에 부합한 행위인지에 대해 고려할 수 있도록 하겠다는 의도가 담겨있다. 법이 의도하는 목적이기도 하지만, 윤리의 기본적인 역할이기도 하다. 따라서, 법과 윤리는 명확하게 구분될 수 있기보다는 둘의 상호작용을 통해 사회가 추구하는 가치를 쉬이 달성할 수 있도록 하는 역할을 하게 된다. AI 시대에 있어서, 윤리를 중요하게 여기는 이유도 마찬가지다.

또한, 저작권법에서 윤리적인 고민을 해야 하는 것은 저작권법이 앞서 나가는 것이 때로는 저작권법이 추구하는 가치와 충돌할 수 있기 때문이다. 저작권법은 궁극적으로 문화향상과 산업발전을 꾀하고 있기 때문이다. AI는 문화적 요소와 산업적 요소를 모두 관련있다. 따라서, 저작권법이 보호와 이용이라는 전통적인 의미의 법에서 나아가, 인공지능 사회에 대한 대응체계를 갖추어나가는 것도 필요하다. AI 데이터 확보를 위해 공정이용 법리가 인정될 수 있다면, 데이터 확보가 되지 않아 나타날 수 있는 편향이나 차별, 또는 독점에 따른 사회적 법익의 침해를 극복할 수 있기 때문이다.²⁶⁾ 앞으로 살펴야 할 사항은 로봇 배제 표준(robot.txt)의 적용, 계약과 데이터윤리, 출처표시 요건, 공정이용 등을 들 수 있다.

AI 기본법에 규정된 사업자의 의무를 대체적으로 노력의무로 제한된다. 강제적인 처벌규정이 없기 때문에 형사책임은 묻기 어렵겠지만, 노력의무 위반이 인정될 경우라면 어느정

이 정확하게 이루어지도록 요구할 수 있는 권리를 정보주체에게 부여해야 한다”고 한다. 양천수 외, 현대 빅데이터 사회와 새로운 인권 구상, 안암법학 57권, 2018, 26면.

26) Sobel, Benjamin, Artificial Intelligence's Fair Use Crisis(September 4, 2017). 41 Colum. J.L. & Arts 45(2017), Available at SSRN: <https://ssrn.com/abstract=3032076>. p.95.

도 민사책임은 가능할 것이다. 만약, 사업자가 AI 기본법에서 정한 노력의무를 완전히 무시하여 이용자에게 피해가 발생했다면, 과실을 이유로 민사상 손해배상책임이 성립할 수 있다. 또한, AI 서비스를 제공하는 계약 관계에서, 사업자가 AI 기본법상 요구되는 안전성, 신뢰성 등의 기준을 충족하려는 노력을 하지 않았다면 계약상 의무 위반이 될 수 있다. 특히, AI 서비스 약관이나 계약서에 일정한 기준을 준수하겠다고 명시되어 있다면, 그 기준을 지키지 않은 경우 손해배상 책임이 인정될 가능성이 크다. 물론, 이용약관 등에서 사업자들은 면책관련 규정을 보다 구체적으로 열거할 가능성도 배제하기 어렵다.

2. 데이터 저작물 수집과 윤리

가. 크롤링은 윤리적인가?

크롤링은 데이터 윤리 관점에서 보면, 크롤러가 수집하는 데이터가 다른 사람들의 개인정보와 같은 민감한 정보인 경우 문제가 발생할 수 있다. 크롤러가 이러한 정보를 수집하여 사용하는 경우, 개인정보 보호법과 같은 법적인 문제가 발생할 수 있다. 크롤러가 데이터를 수집할 때 사이트의 로봇 배제 표준(robots.txt)을 무시하거나, 서비스 제공자의 이용약관을 무시하는 경우도 있다. 이러한 행위는 법적으로 문제가 될 수 있을 뿐만 아니라, 윤리적인 문제도 제기될 수 있다. 그러나, 크롤링 자체가 모든 경우에 윤리적으로 부적절하다고 말할 수는 없다. 예를 들어, 검색 엔진은 웹 페이지를 크롤링하여 정보를 수집한다. 이 경우, 크롤러는 검색 엔진의 서비스를 위해 정보를 수집하고 있으며, 이용약관과 로봇 배제 표준을 준수하고 있다. 따라서, 크롤링이 데이터 윤리적인 측면에서 문제가 되는지는 크롤러가 어떤 데이터를 수집하고, 그것을 어떻게 사용하느냐에 따라 다르게 판단된다.

허락없이 크롤링 하는 것도 윤리적이지는 않다. 저작권을 허락 없이 크롤링하는 것은 부당한 데이터 수집 방법 중 하나이다. 크롤링된 데이터는 원래 소유주의 동의를 받은 경우를 제외하면 사용할 수 없다. 예를 들어, 어떤 웹사이트에서 정보를 크롤링하여 다른 목적으로 사용하거나, 해당 웹사이트의 콘텐츠를 무단으로 복제하여 사용하는 것은 부적절하다. 크롤링할 때는 반드시 저작권법 및 개인정보 보호법을 준수해야 한다. 따라서, 면책을 주장하기 위해서는 크롤링을 통해 수집한 데이터의 사용 목적과 방법도 중요하다. 예를 들어, 민감한 개인정보를 수집하거나, 인종, 종교, 성별 등의 개인적 특성을 포함한 데이터를 부적절하게 사용하는 것은 데이터 윤리상 문제가 될 수 있다. 따라서, 데이터 수집 시 개인정보 보호 및 공정한 사용 등을 고려하는 것이 바람직하다.

나. 로봇배제원칙은 법적 효력이 있는가?

(1) 로봇배제원칙

로봇배제원칙은 검색 엔진 로봇이 웹사이트를 크롤링할 때 따라야 하는 지침을 제공하는 파일이다. 이 파일은 웹사이트 소유자가 자신의 웹사이트에서 수집되는 데이터의 양과 종류를 제어할 수 있는 방법을 제공한다. 로봇배제원칙이 담긴 robots.txt 파일은 법적 효력이 있는 파일은 아니지만, 검색 엔진에서 권장되는 규약이며, 대부분의 검색엔진은 이 원칙을 따른다. 하지만, robots.txt 파일을 무시하고 웹사이트를 크롤링하는 것은 법적으로 문제가 될 수 있다. 특히, 웹사이트 소유자가 robots.txt 파일에서 허용하지 않은 부분을 크롤링하거나, 크롤링을 금지한 경우에는 법적 분쟁이 발생할 수 있다. 또한, 크롤링을 통해 수집된 데이터의 사용에 대한 법적 문제가 발생할 수 있으므로, 항상 데이터 윤리와 관련된 법과 규제를 준수해야 한다.

로봇배제원칙이 법적 강제성이 없다는 점이 논란이다. 강제성이 없다는 점에서 이를 위반한다고 하더라도 법적 책임을 묻기가 어렵다. 로봇배제원칙은 저작권법과 직접적인 관련이 있는 것은 아니지만, 웹사이트의 콘텐츠가 저작권 보호 대상인 경우 이를 크롤링하고 활용하는 것은 저작권 침해로 이어질 수 있다. 웹사이트의 이용약관에 크롤링 금지를 명시하고 있다면, 이를 무시한 크롤링은 계약 위반(Breach of Contract)에 해당할 수 있다. 회원으로 가입하지 않은 경우에는 이용약관의 적용을 받지 않는다는 점도 고려되어야 한다.²⁷⁾ 이처럼, 단순히 로봇배제원칙만으로 법적 효력이 발생하는 것은 아니므로, 크롤링이 불법인지 여부는 개별 사안에 따라 판단될 수밖에 없다.

(2) 법원의 판단

여기어때 사건에서 대법원은 크롤링과 관련하여 무죄판결을 내렸다.²⁸⁾ 데이터베이스제작자는 그의 데이터베이스의 전부 또는 상당한 부분을 복제·배포·방송 또는 전송할 권리를 가지고(저작권법 제93조 제1항), 데이터베이스의 개별 소재는 데이터베이스의 상당한 부분으로 간주되지 않지만, 개별 소재의 복제 등이라 하더라도 반복적이거나 특정한 목적을 위하여 체계적으로 함으로써 해당 데이터베이스의 통상적인 이용과 충돌하거나 데이터베이스 제작자의 이익을 부당하게 해치는 경우에는 해당 데이터베이스의 상당한 부분의 복제 등으로 본다(저작권법 제93조 제2항). 이는 지식정보사회의 진전으로 데이터베이스에 대한 수요가 급증함에 따라 창작성의 유무를 구분하지 않고 데이터베이스를 제작하거나 그 갱신·

27) 대법원 2022. 5. 12. 선고 2021도1533 판결.

28) 대법원 2022. 5. 12. 선고 2021도1533 판결.

검증 또는 보충을 위하여 상당한 투자를 한 자에 대하여는 일정기간 해당 데이터베이스의 복제 등 권리를 부여하면서도, 그로 인해 정보공유를 저해하여 정보화 사회에 역행하고 경쟁을 오히려 제한하게 되는 부정적 측면을 방지하기 위하여 단순히 데이터베이스의 개별 소재의 복제 등이나 상당한 부분에 이르지 못한 부분의 복제 등만으로는 데이터베이스제작자의 권리가 침해되지 않는다고 규정한 것이다. 데이터베이스제작자의 권리가 침해되었다고 하기 위해서는 데이터베이스제작자의 허락 없이 데이터베이스의 전부 또는 상당한 부분의 복제 등이 되어야 하는데, 여기서 상당한 부분의 복제 등에 해당하는지를 판단할 때는 양적인 측면만이 아니라 질적인 측면도 함께 고려하여야 한다. 양적으로 상당한 부분인지 여부는 복제 등이 된 부분을 전체 데이터베이스의 규모와 비교하여 판단하여야 하며, 질적으로 상당한 부분인지 여부는 복제 등이 된 부분에 포함되어 있는 개별 소재 자체의 가치나 그 개별 소재의 생산에 들어간 투자가 아니라 데이터베이스제작자가 그 복제 등이 된 부분의 제작 또는 그 소재의 갱신·검증 또는 보충에 인적 또는 물적으로 상당한 투자를 하였는지를 기준으로 제반 사정에 비추어 판단하여야 한다. 또한 앞서 본 규정의 취지에 비추어 보면, 데이터베이스의 개별 소재 또는 상당한 부분에 이르지 못하는 부분의 반복적이거나 특정한 목적을 위한 체계적 복제 등에 의한 데이터베이스제작자의 권리 침해는 데이터베이스의 개별 소재 또는 상당하지 않은 부분에 대한 반복적이고 체계적인 복제 등으로 결국 상당한 부분의 복제 등을 한 것과 같은 결과를 발생하게 한 경우에 한하여 인정함이 타당하다.

저작권법 위반과 마찬가지로, 네트워크 침입 여부에 대해서도 무죄판단을 하였다. 정보통신망법 제48조는 “누구든지 정보통신망에 불법적으로 침입하여 정보를 삭제·변경 또는 이용하거나 정보통신망의 정상적인 운영을 방해한 자는 5년 이하의 징역 또는 5천만 원 이하의 벌금에 처한다”라고 규정하고 있다(법 제71조 제1항 제9호). 위 규정은 이용자의 신뢰 내지 그의 이익을 보호하기 위한 규정이 아니라 정보통신망 자체의 안정성과 그 정보의 신뢰성을 보호하기 위한 것이므로, 위 규정에서 접근권한을 부여하거나 허용되는 범위를 설정하는 주체는 서비스제공자이다. 따라서 서비스제공자로부터 권한을 부여받은 이용자가 아닌 제3자가 정보통신망에 접속한 경우 그에게 접근권한이 있는지 여부는 서비스제공자가 부여한 접근권한을 기준으로 판단하여야 한다. 그리고 정보통신망에 대하여 서비스제공자가 접근권한을 제한하고 있는지 여부는 보호조치나 이용약관 등 객관적으로 드러난 여러 사정을 종합적으로 고려하여 신중하게 판단하여야 한다.

다. 데이터 윤리와의 저촉

robots.txt는 법적으로 크롤링을 차단하는 강제력 있는 수단이 아니다. 다만, 이를 무시하고 크롤링하는 것은 윤리적으로 논란이 될 수 있으며, 이용약관과 결합될 경우 계약법적으로 문제될 수 있다. 저작권 보호 여부는 robots.txt와 무관하게 독립적으로 판단된다. 즉, robots.txt가 존재하더라도 크롤링된 콘텐츠가 저작권 보호 대상이라면 무단 사용 시 저작권 침해가 될 수 있다. 다만, 크롤링을 원천적으로 막는 것은 또다른 문제가 될 수 있다. 특히, 공공데이터를 차단하겠다는 것은 문제이다. 공공데이터법, 정보접근성 등 데이터의 공공성을 위한 입법과 제도가 수립되고 있기 때문이다.

3. 데이터 저작물 이용과 윤리

가. 데이터 공정이용

인공지능의 맥락에서 데이터를 더 잘 활용하기 위해 규정을 개선하기 위한 지속적인 논의와 이니셔티브가 필요하다. 예를 들어, EU는 현재 EU 전체에서 데이터 사용을 위한 보다 조화되고 일관된 법적 프레임워크를 만드는 것을 목표로 하는 새로운 데이터 법에 대한 제안을 준비하고 있다. EU의 DSM 지침²⁹⁾을 통해 개정된 데이터베이스 지침은 TDM을 포함하여 데이터 액세스, 공유 및 재사용과 관련된 문제를 해결할 것으로 예상된다.

나. TDM을 위한 저작권법 개정

인공지능의 맥락에서 데이터를 더 잘 활용하기 위해 규정을 개선하기 위한 지속적인 논의와 이니셔티브가 필요하다. 예를 들어, EU는 현재 EU 전체에서 데이터 사용을 위한 보다 조화되고 일관된 법적 프레임워크를 만드는 것을 목표로 하는 새로운 데이터 법에 대한 제안을 준비하고 있다. EU의 DSM 지침을 통해 개정된 데이터베이스 지침은 TDM을 포함하여 데이터 액세스, 공유 및 재사용과 관련된 문제를 해결할 것으로 예상된다.³⁰⁾

29) Directive(EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on copyright and related rights in the Digital Single Market and amending Directives 96/9/EC and 2001/29/EC(Text with EEA relevance.)

30) DSM 지침 제3조 과학적 연구 목적의 텍스트/데이터 마이닝

(1) 회원국은 연구기관과 문화유산기관이 과학적 연구 목적으로 그들이 합법적인 접근 권한을 가지는 저작물이나 그 밖의 보호대상을 텍스트/데이터 마이닝을 수행하기 위해 복제하고 추출하는 것에 대해 데이터베이스보호지침 제5조 (a)와 제7조 제1항, 정보사회저작권지침 제2조 및 이 지침 제15조 제1항에 규정된 권리에 대한 예외를 규정하여야 한다.

(2) 제1항에 따라 만들어진, 저작물이나 그 밖의 보호대상의 복제물은 적절한 수준의 보안을 갖춰 저장하여야 하고,

또한 다양한 국가와 지역에서 데이터 거버넌스를 개선하고 데이터 공유를 촉진하는 방법을 모색하고 있다. 예를 들어, 일본은 최근 AI 및 기타 혁신적인 애플리케이션을 위한 데이터 사용을 촉진하는 것을 목표로 하는 데이터 활용에 관한 새로운 법률을 도입했다. 이 법은 데이터 공유를 촉진하고 책임 있는 데이터 사용에 대한 지침을 설정한다.

저작권법 개정안(이용호 의원안)에 따른 TDM 규정

제35조의5(정보분석을 위한 복제·전송 등) ① 컴퓨터를 이용한 자동화 분석 기술을 통하여 추가적인 정보 또는 가치를 생성하기 위한 목적으로 다수의 저작물을 포함한 대량의 정보를 분석(규칙, 구조, 경향 및 상관관계 등의 정보를 추출하는 경우를 말한다. 이하 이 조에서 "정보분석"이라 한다)하는 것으로 다음 각 호의 요건을 모두 충족하는 경우에는 필요한 범위 안에서 저작물을 복제·전송할 수 있다.

1. 저작물에 표현된 사상이나 감정을 향유하지 아니할 것
2. 정보분석의 대상이 되는 해당 저작물에 적법하게 접근할 것
- ② 제1항에 따라 저작물을 복제하는 자는 정보분석을 위하여 필요한 한도 안에서 복제물을 보관할 수 있다. 이 경우 저작권 및 그 밖에 이 법에 따라 보호되는 권리의 침해를 방지하기 위하여 복제방지조치 등 대통령령으로 정하는 필요한 조치를 하여야 한다.
- ③ 정보분석의 결과물에 대하여 다음 각 호의 어느 하나에 해당하는 목적으로 적법하게 접근하는 경우에는 「부정경쟁방지 및 영업비밀보호에 관한 법률」, 「데이터 산업진흥 및 이용촉진에 관한 기본법」, 「산업 디지털 전환 촉진법」 및 그 밖의 데이터 보호에 관한 다른 법률의 규정에도 불구하고 해당 결과물을 이용할 수 있다. 다만, 정보분석을 위하여 정당한 권리자로부터 저작물의 복제·전송에 대한 이용의 허락을 받은 경우에는 그러하지 아니하다.
 1. 교육·조사·연구 등 비상업적 목적
 2. 저작물의 창작 목적

연구결과의 검증 등 과학적 연구 목적을 위해 유지할 수 있다.

- (3) 권리자들은 저작물이나 그 밖의 보호대상이 호스팅되는 네트워크와 데이터베이스의 보안과 무결성을 보장하기 위한 조치를 적용하는 것이 허용되어야 한다. 그러한 조치는 그 목적을 달성하는 데에 필요한 수준을 넘어서서는 안 된다.
- (4) 회원국은 권리자, 연구기관 그리고 문화유산기관이 제2항과 제3항에 각각 언급된 의무와 조치의 적용에 관하여 통상적으로 합의된 최적 관행을 정의하도록 권장하여야 한다.

제4조 텍스트/데이터 마이닝을 위한 예외와 제한

- (1) 회원국은 텍스트/데이터 마이닝을 목적으로 합법적으로 접근 가능한 저작물과 그 밖의 보호대상의 복제와 추출을 위해 데이터베이스지침 제5조(a)와 제7조제1항, 정보사회저작권지침 제2조, 컴퓨터프로그램지침 제4조 제1항(a)와 (b) 그리고 이 지침 제15조 제1항에 규정된 권리에 대한 예외와 제한을 규정하여야 한다.
- (2) 제1항에 따라 만들어진 복제물과 추출물은 텍스트/데이터 마이닝 목적으로 필요한 한 보관될 수 있다.
- (3) 제1항에 규정된 예외와 제한은 제1항에 언급된 저작물과 그 밖의 보호대상의 이용이 권리자에 의해, 콘텐츠가 온라인으로 공중에게 이용 제공되는 경우에 기계가독형 수단 등, 적절한 방법으로 명시적으로 표명되지 않았다는 것을 조건으로 적용되어야 한다.
- (4) 이 조항은 이 지침 제3조의 적용에 영향을 미쳐서는 안 된다.

전반적으로 AI 개발을 위한 데이터의 중요성에 대한 인식이 높아지고 있으며 데이터 액세스 및 혁신의 필요성과 개인정보 보호 및 기타 고려 사항의 균형을 맞추는 규제 틀을 만들기 위한 노력이 이루어지고 있다.

다. TDM 관련 법원 판결

최근 TDM과 공정이용 관련 중요한 판결이 내려졌다. EU의 경우와 미국의 경우라는 점에서 각각의 진영에 대한 법원의 판결과 시장에 미치는 영향을 살펴본다. 입법적으로 해결한 EU는 소송결과는 당연한 것으로 받아들여질 수 있다. 반면, 미국은 공정이용이 기업이 오랫동안 주장해 온 대응논리였다는 점에서 부정적으로 작용할 가능성이 커졌다. 다만, 이번 판결은 경쟁사업자의 데이터 이용에 관한 것으로 생성형 AI에 적용되는 사건은 아니라는 점에서 전세계적으로 진행 중인 AI 데이터 소송에 일반화하는 것은 경계할 필요가 있다.

(1) 독일

독일 함부르크 지방법원은 2024년 9월 27일 크네쉬케 대 라이온(Kneschke vs. LAION) 사건에서 데이터셋의 제작 과정에서 발생하는 저작권 침해와 이러한 침해가 텍스트 및 데이터 마이닝(TDM) 제한 규정으로 정당화될 수 있는지에 대하여 최초로 판결했다.³¹⁾

주요 판결 내용은 다음과 같다: LAION의 데이터셋 생성 활동은 독일 저작권법 제60d조에 따른 비상업적 과학 연구 목적의 TDM에 해당한다. LAION이 상업 기업과 협력 관계에 있다는 사실만으로는 비상업적 성격이 부정되지 않는다. 웹사이트 이용약관에 명시된 자연어로 된 TDM 금지 문구도 '기계가 읽을 수 있는 형식'의 opt-out으로 간주될 수 있다. 법원은 LAION의 이미지 데이터셋 구축 활동을 비상업적 과학 연구 목적의 TDM으로 인정했다. 이는 학계와 산업계 간의 협력 연구에 대한 법적 보호를 강화할 수 있다. 순수 학술 연구뿐만 아니라 산학 협력 프로젝트도 TDM 예외의 혜택을 받을 수 있다는 점을 명확히 하였다. 이로써, AI 연구와 개발에 더 넓은 자유를 제공할 수 있을 것이다.

(2) 미국

톰슨 로이터는 대형 법률 검색 플랫폼인 Westlaw를 운영하며, 판례, 법령, 법학 논문 등을 제공한다. 또한 법적 논점을 정리한 헤드노트(headnotes) 및 Key Number System을 보유하고 있으며, 이에 대한 저작권을 주장하고 있다. ROSS는 AI 기반 법률 검색 엔진을 개발하는 기업으로, AI를 훈련하기 위한 법률 데이터를 필요로 했다. ROSS는 Westlaw의

31) LG Hamburg Urteil vom 27.09.2024 – 310 O 227/23.

콘텐츠 사용을 요청했지만, 경쟁 관계로 인해 톰슨 로이터는 이를 거절했다. 이에 따라 ROSS는 LegalEase라는 제3자 기업과 계약을 맺고, Bulk Memo라는 법률 질문 및 답변 데이터를 확보했다. 이러한 Bulk Memo는 Westlaw의 헤드노트를 기반으로 생성되었으며, 결과적으로 ROSS의 AI가 Westlaw의 저작물을 학습하는 데 사용되었다. 이에 톰슨 로이터는 ROSS를 저작권 침해로 고소했다.

델라웨어 연방지방법원은 ROSS의 공정이용 항변을 다음과 같은 이유로 기각하였다.³²⁾ ROSS는 AI 훈련을 위한 중간 과정에서만 헤드노트를 사용했으므로 공정이용에 해당한다고 주장했다. 그러나 법원은 공정이용을 인정하지 않았다. 첫째, 이용 목적 및 성격(Purpose and Character of Use)에 있어서 ROSS는 상업적 목적을 가지고 Westlaw의 헤드노트를 이용했으며, AI 훈련을 위한 중간 단계에서 사용되었다고 하더라도 변형적 이용(transformative use)으로 볼 수 없다고 판결했다. 둘째, 원저작물의 성격(Nature of the Copyrighted Work)에 있어서 헤드노트는 단순한 사실이 아니라 창작적 요소가 포함된 편집물로, 공정이용이 적용되기 어렵다. 셋째, 사용된 분량 및 중요성(Amount and Substantiality of the Portion Used)에 있어서 ROSS는 수천 개의 헤드노트를 복제했으며, 이는 Westlaw의 핵심 콘텐츠를 이용한 것이므로 공정이용이 성립하지 않는다. 넷째, 시장 대체 효과(Market Effect) Ross의 검색엔진은 Westlaw와 경쟁하는 제품으로, Westlaw의 시장에 직접적인 영향을 미치므로 공정이용으로 볼 수 없다.

법원은 Ross가 2,243개의 헤드노트를 무단으로 복제하여 저작권을 침해했다고 판단하며, 공정이용 항변을 인정하지 않았다. 그렇지만, 이 판례에서 고려해야 할 사항은 ROSS가 직접적으로 기계학습에 이용했는지 확인되지 않았다. 또한, 우리 저작권법의 일시적 복제와 같은 ‘중간복사’를 인정하지 않았다. 데이터를 이용하기 위하여 복제가 이루어질 수밖에 없다는 점을 인정한 것으로 볼 수 있다.

무엇보다, Westlaw와 피고는 경쟁관계라는 점이다. 그렇기 때문에 공정이용의 4번째 요건인 시장대체성 요건을 충족하지 못한다고 본 것이다. 만약, 경쟁관계가 아닌 범용 AI 모델을 구축하는 경우라는 법원은 다르게 판단하였을 것이다. 다만, 이 사건은 많은 소송의 이슈가 되고 있는 생성형 AI 모델과 관련된 사례는 아니다. 따라서, 생성형 AI 소송에 적용되는 일반적인 사건은 아니라는 점이다. 따라서, 이 사건이 모든 생성형 AI 소송에 일반화되는 것은 아니라고 할 것이다.

32) Thomson Reuters Enter. Centre GmbH et al v. ROSS Intelligence Inc., Case No. 1:20-cv-00613-SB(D. Del. Feb. 11, 2025).

4. 생성된 데이터의 윤리

AI가 스스로 생성하는 것은 기술적으로 어렵다. 다부스(Dabus) 사건에서도 법원은 그 점을 명시하고 있다.³³⁾ 설령, 기술적으로 그렇다고 하더라도 입법적 뒷받침 없이 권리주체가 되지 못한다. 이에 대한 법체계 전반적인 검토가 요구된다. 따라서, AI를 도구적으로 생성한 경우로 한정한다. 생성물은 저작물성을 인정받을 수 있는지에 대한 논란은 여전하다. 획일적으로 인정하는 것은 어렵다고 하더라도, 경우에 따라 살펴보는 것이 필요하다. 창작적 기여를 인정할 수 있느냐의 여부에 따라 달라질 수 있을 것이다. 2025년 발간된 미국저작권청 보고서에도 명확하게 기준을 제시한 것으로 보기 어렵다.³⁴⁾ 경우에 따라서 달리 판단할 수 있다는 것으로 이해되기 때문이다. 결국은 최소한의 창작성이라는 기준에 따를 수밖에 없을 것이다.

AI 시스템을 통해 이용자가 생성하는 결과물이 원래 저작권과 유사할 수 있다. 이는 여러 가지 요인을 생각할 수 있다. 먼저, 복제되는 경우이다. 이는 중복된 학습데이터 출처가 다른 동일한 데이터가 중복적으로 학습데이터로 정제된 경우로 볼 수 있다. 텍스트-이미지 모델의 맥락에서 모델이 동일한 작품의 여러 복제본에서 학습되고, 이미지가 고유한 텍스트 설명과 연관되고, 모델 크기와 학습 데이터의 비율이 비교적 클 때 학습 데이터를 암기할 가능성이 더 높다고 한다.³⁵⁾

모델이 저작권이 있는 작품을 모방하도록 의도적으로 구축되지 않았더라도 결국 저작권을 침해하는 수준으로 모방될 수 있다. 기계학습 모델은 때때로 입력 데이터의 특성이 과적합됨으로서, 동일한 결과가 나올 수 있기 때문이다.³⁶⁾

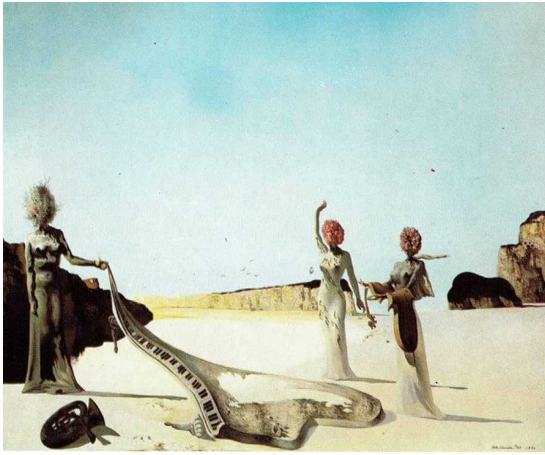
33) 서울고등법원 2024.5.16. 선고 2023누52088 판결.

34) U.S. Copyright Office, Copyright and Artificial Intelligence Part 2, 2025.

35) Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593>, p.302.

36) Sobel, Benjamin, Artificial Intelligence's Fair Use Crisis(September 4, 2017). 41 Colum. J.L. & Arts 45(2017), Available at SSRN: <https://ssrn.com/abstract=3032076>, p.64.

표 1 원작과 AI 생성 결과

구분	원작 ³⁷⁾	코파일럿 생성 ³⁸⁾
살바도로 달리		
해리포터		

* 출처: 좌(인터넷 검색), 우상(연구자 via 코파일럿), 우하(인터넷 검색)³⁹⁾

기계학습 모델에 충분히 동질적인 학습데이터 세트가 제공되었거나 해당 학습데이터 세트에 과적합되었을 경우, 입력 데이터와 실질적으로 유사하거나 완전히 복제된 출력이 생성될 수 있다.⁴⁰⁾ 이처럼, LLM은 너무 커서 훈련 데이터에서 특정 작품을 본질적으로 기억

37) 살바도르 달리가 1936년에 그린 Three Young Surrealist Women Holding in Their Arms the Skins of an Orchestra이다.

38) Three Young Surrealist Women Holding in Their Arms the Skins of an Orchestra을 프롬프트로 하여, 코파일럿을 통해 생성한 이미지이다.

39) 코파일럿, ChatGPT 등 생성형 AI는 저작권 침해를 이유로 해리포터와 관련된 프롬프트를 이용한 생성을 거부하였다. 이에 따라, 인터넷 검색을 통해 찾은 이미지로 비교하였다.

40) Sobel, Benjamin, Artificial Intelligence's Fair Use Crisis(September 4, 2017). 41 Colum. J.L. & Arts 45(2017), Available at SSRN: <https://ssrn.com/abstract=3032076>. p.64.

할 수 있다는 점을 인식할 필요가 있다. 모델이 학습데이터를 기억하는 경우, 출력을 통해 학습데이터의 원래 표현을 전달할 수 있기 때문이다.⁴¹⁾

AI 모델 학습에 상상을 초월한 데이터가 입력되기 때문에 데이터에 내재된 문제를 다 정제하기는 어려운 상황이다. 방대한 분량과 중복되는 데이터, 그리고 그 출처가 다를 수 있다는 점이 더 오랜 기억의 형태로 출력될 수 있다. 물론, 학습데이터로 이용된 일부는 공정 이용을 통해서 허용가능한 데이터가 사용되었을 수 있다. 또한, 개인 블로그에 올려진 폐쇄형 공간의 데이터를 크롤링했을 가능성도 있다. 그렇기 때문에 학습데이터의 출처를 확인하지 않고서는 중복여부를 판단하기는 쉽지 않을 것이다. 이는 기업의 데이터 거버넌스가 제대로 작동하지 못한다는 증거로 볼 수도 있다.

둘째는 AI 모델의 버그이다. MLL은 데이터 자체를 기억하지 않도록 설계되었지만, 알 수 없는 사정으로 인하여 동일하거나 유사한 데이터가 생성될 수 있다는 점이다. 이는 블랙박스의 특성상 확인이 쉽지 않다는 점에서 버그로 볼 수 있는 것이다.

셋째, AI 모델이 학습하여 나름대로의 규칙이나 방식에 따라 생성한 것이 우연찮게 일치하는 경우이다. 인간의 창작에서도 동일한 결과물을 생성하기도 한다. 인간의 뇌구조를 모방한 AI에서도 같은 현상이 발생할 수 있을 것이다. 이는 과학적으로 확인할 수 없으나, 저작권 소송에 있어서도 우연의 일치로 판단된 경우도 자못 있다는 점이 이를 반증할 수 있다고 본다. 기계가 수많은 데이터를 통해 그 특징을 학습하였고, 그 특징들이 통계적으로 생성된다는 점에서 충분히 가능성이 있다고 생각된다.

이처럼, 유사한 결과물이 출력되는 것은 추론에 따른 것으로 중복된 데이터 또는 학습한 AI 모델이 우연적으로 생성할 수 있다는 점이 명확하게 밝혀지는 것은 아니다. AI 생성물의 성격에 대해 암기, 의도, 무의식적 표현이나, 우연의 일치이나 등 논란이 예상된다. 저작권법은 고의에 의한 침해와 달리 무의식적 침해에 대해서는 고의성이 인정되지 않기 때문에 형사책임에서는 면책될 것으로 보인다.

5. 소결

LLM이 기존 저작물과 유사한 결과를 생성하는 것은 여러 가지 요인이 있을 것이나, 확인가능한 것은 동일한 데이터를 학습데이터로 사용한 경우에는 그 가능성이 높다는 것이

41) Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593> or <http://dx.doi.org/10.2139/ssrn.4438593>, p.312.

다. 또한, 헤리포터와 같은 대중적인 데이터를 활용하는 경우에도 전문적인 데이터보다 가능성이 높다.⁴²⁾ 데이터 윤리와 저작권법의 관계 설정을 위해 인공지능 윤리에 더하여, 기계학습의 기반이 되는 데이터 윤리에 대한 관심도 필요하다. 아무리 인공지능이 윤리적이고 하더라도, 학습에 필요한 데이터셋이 왜곡된 경우라면 알고리즘이나 소스코드를 공개 하더라도, 그 원인을 찾기가 쉽지 않기 때문이다. 윤리적 적용이 어렵게 되면, 규제단계를 넘어 금지유형으로 형법 체계에서 인공지능이 다루어질 가능성이 크다. 인공지능 윤리는 인공지능과의 싸움 이전에, 규제기관과의 충돌을 피하기 위한 방안이라는 점을 인식할 필요가 있다.

IV. 데이터 윤리를 위한 관련 제도 검토

AI 윤리가 저작권법에 영향을 미칠 수 있는 내용은 앞서 살펴본 바와 같이 다양하다. 다만, 보다 구체적으로 제도화에 대한 논의가 저작권법 또는 AI 기본법 등 관련 법령을 대상으로 이루어지고 있기 때문에 관련하여 살펴본다.

1. 표시제도

가. 표시제도의 의의

인공지능은 단순한 창작/복제의 도구에 그치지 않고 기존의 창작방식과 전혀 다른 새로운 방식으로 엄청나게 빠른 속도로 대량의 콘텐츠를 생산하고 소비하는 전혀 새로운 저작권 생태계를 만들어 나갈 것이다.⁴³⁾ 이러한 창작활동은 인간과 비교할 수 없을 정도이기 때문에 AI 생성물이 확산하면서 인간의 창작과 구별이 필요하다.⁴⁴⁾ 인공지능에 의해 만들어진 것이 인간을 감동시킨다면 인간이 만든 것과 다를 바 없다.⁴⁵⁾ 이러한 상황에서, 저작권 제

42) 헤리포터 시리즈는 다른 일반적인 책보다 훨씬 대중적이고 다양한 언어로 번역되었기 때문에 그 가능성이 높다고 할 것이다. Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593>, p.336.

43) 정상조, 인공지능시대의 저작권법 과제, 계간저작권 통권 122호, 2018, 67면.

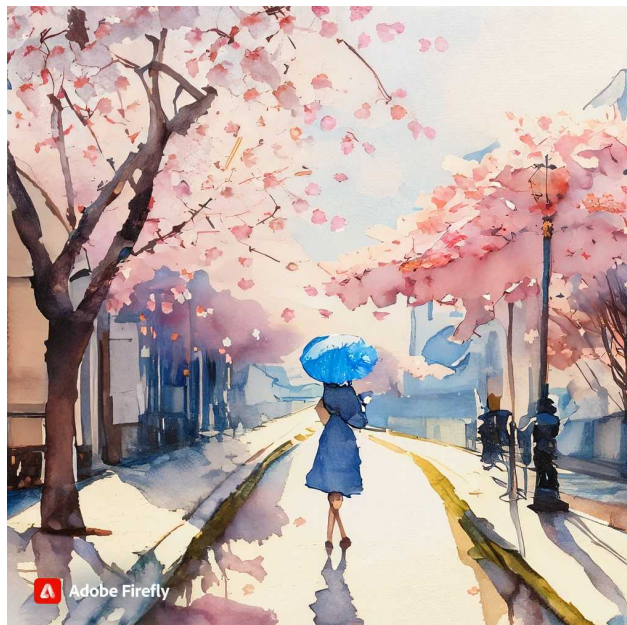
44) AI 창작물에 대해 권리를 부여하는 것에 대해 우려도 있다. 즉, AI가 생성한 콘텐츠에 법적 효를 제공하면 인간의 저작이 억제될 것이라는 우려를 표명하였으며, 법적으로 보호되는 AI가 생성한 결과물이 급증하면 창의성이 억제되어 인간이 창작할 의욕이 저하되어 대중이 볼 수 있는 인간이 만든 작품이 전반적으로 감소할 것이고 주장했다고 한다. U.S. Copyright Office, Copyright and Artificial Intelligence Part 2, 2025, p.34.

45) Annemarie Bridy, "Coding Creativity: Copyright and the Artificially Intelligent Author", 2012 Stan.

도에 대한 근본적인 변화가 필요하다. 다만, 저작권법 개정 논의보다는 AI가 가져오는 여러 가지 한계점을 극복하기 위한 제도적 방안이 제안되고 있다. 대표적으로, 표시제도를 들 수 있다. AI 창작물에의 표시를 통하여, 인간과 AI 창작을 구별함으로써 신뢰관계를 형성할 수 있을 것이다. 즉, AI가 생성한 것을 인간의 것으로 오인하지 않도록 하는 것이다. 저작물과 인공지능 창작물 간의 오인과 혼동을 예방하고, 표절과 저작물의 불법이용을 방지하기 위한 수단이 될 것이라는 점이다.⁴⁶⁾

표시제도는 사업자에게 규제적인 속성으로 작용할 수 있다. AI 시스템에 표시를 해야 한다는 점에서 이용자의 반발이 예상될 수 있을 것이고, 오히려 저작자로서 오인될 가능성도 있기 때문이다. 예를 들면, 아래 그림의 ‘Adobe firefly’라는 로고는 어도비(adobe) 사의 파이어플라이(firefly)라는 생성형 AI를 이용한 경우에 생성된다. 그렇지만, 이용자는 자신이 창작한 것에 대해 제3자의 표시를 의무화하는 것을 수용하지 않을 수 있다는 점도 고려될 필요가 있다.⁴⁷⁾

그림 1 Adobe firefly 워터마크가 표시된 이미지



* 출처: 연구자(adobe firefly)

Tech. L. Rev. 5.(2012) 정상조, 인공지능시대의 저작권법 과제, 계간저작권 통권 122호, 2018, 43면에서 재인용.

46) 소병수, 미디어 영역에서 인공지능(AI) 기술의 확산과 제문제, 원광법학 제39권 제3호, 2023, 39면.

47) AI 기본법에서는 표시 의무화이외에 구체적인 방법 등에 대해서는 시행령에 위임하고 있기 때문에 시행령에 이러한 이용자의 요구사항이 반영될 것으로 기대한다.

생성물표시는 저작권법의 적용 여부와는 크게 상관이 없다. 저작권법은 등록표시 등 표시요건이 법적 효과는 없기 때문이다. 저작권법에 식별표지로서 의무화하는 것과는 별개이지만, 저작물성이 없는 경우에 대해서까지 저작권법이 관여하는 것은 입법취지와 목적에 부합하지 않기 때문에 위법성 논란이 예상된다.⁴⁸⁾ 따라서, 저작권법이 아닌 다른 법률 예를 들면 콘텐츠산업 진흥법이나 AI 기본법이 타당하다. 다행히, AI 기본법에는 표시요건이 사업자의 의무로 규정되어있기 때문에 이러한 논란은 해결될 것이다. 다만, 저작권법 개정안에 표시요건이 들어가 있기 때문에 이에 대해서는 살펴볼 필요가 있다.

표시제도는 AI가 생성했다는 것을 이용자에게 제시하도록 하는 것이다. 인간 창작과의 구분이 가능하도록 하자는 의도이다. 이는 이용자 혼동 방지나 콘텐츠의 신뢰성 확보를 기대할 수 있는 정책적 판단에 따른 것이다. AI 기본법에서도 신뢰성을 확보하기 위하여 생성형 AI로 생성한 결과물에 대해서는 표시하도록 하고 있다.

이러한 방법으로써, 표시요건을 고안한 것이다. 물론, 표시요건은 새로운 것은 아니다. 이미 콘텐츠산업 진흥법에서도 표시를 손해배상의 요건으로 하고 있다. 또한, 제조물책임법에서 표시요건을 흠결한 경우에도 결함으로 보고 있는 등 다양한 법제에서 표시를 요건으로 하고 있다.

나. 표시 효과

AI 기본법에서는 표시가 권리발생 등의 효과를 부여하는 것은 아니다. 사업자의 의무로써 생성형 AI로 창작한 결과물에 대해서는 표시하도록 하고 있을 뿐이다. 이와 관련하여 참고할 수 있는 법률로는 콘텐츠산업 진흥법을 들 수 있다. 동 법은 누구든지 정당한 권한 없이 콘텐츠제작자가 상당한 노력으로 제작하여 대통령령으로 정하는 방법에 따라 콘텐츠 또는 그 포장에 제작연월일, 제작자명 및 이 법에 따라 보호받는다라는 사실을 표시한 콘텐츠의 전부 또는 상당한 부분을 복제·배포·방송 또는 전송함으로써 콘텐츠제작자의 영업에 관한 이익을 침해하여서는 아니된다. 다만, 콘텐츠를 최초로 제작한 날부터 5년이 지났을 때에는 그러하지 아니하다. 또한, 누구든지 정당한 권한 없이 콘텐츠제작자나 그로부터 허락을 받은 자가 제1항 본문의 침해행위를 효과적으로 방지하기 위하여 콘텐츠에 적용한 기

48) “생성형 인공지능(글·소리·그림·영상 등의 결과물을 다양한 자율성의 수준에서 생성하도록 만들어진 인공지능을 말한다) 기술을 이용하여 인간의 사상 또는 감정을 표현한 저작물을 창작한 자는 해당 저작물이 생성형 인공지능 기술을 이용하여 제작되었다는 사실을 표시하여야 한다.”고 규정하며, 제2항에서는 표시의 내용과 방법에 필요한 사항을 대통령령에 위임한다. 위 신설 조항은 표시 대상을 “저작물”로 표현하고 있어서, 저작물로 인정되지 않는 인공지능 산출물은 표시 의무의 범위에서 제외되는 중대한 문제가 있다고 지적된다. 류시원, 인공지능 산출물 표시 규제 방안의 검토, 저스티스 통권 제205호, 2024, 552면.

술적보호조치를 회피·제거 또는 변경(이하 “무력화”라 한다)하는 것을 주된 목적으로 하는 기술·서비스·장치 또는 그 주요 부품을 제공·수입·제조·양도·대여 또는 전송하거나 이를 양도·대여하기 위하여 전시하는 행위를 하여서는 아니된다. 다만, 기술적보호조치의 연구·개발을 위하여 기술적보호조치를 무력화하는 장치 또는 부품을 제조하는 경우에는 그러하지 아니하다. 더욱이, 콘텐츠제작자가 표시사항을 거짓으로 표시하거나 변경하여 복제·배포·방송 또는 전송한 경우에는 처음부터 표시가 없었던 것으로 본다(제37조). 행위규제 형태로 표시된 콘텐츠의 불법적인 이용을 금지하고 있다. 더욱이, 이렇게 표시된 콘텐츠를 위법하게 이용할 경우에는 손해배상 청구권을 부여하고 있다. 금지되는 행위로 인하여 자신의 영업에 관한 이익이 침해되거나 침해될 우려가 있는 자는 그 위반행위의 중지 또는 예방 및 그 위반행위로 인한 손해의 배상을 청구할 수 있다. 다만, 위반하는 행위에 대하여 콘텐츠제작자가 같은 항의 표시사항을 콘텐츠에 표시하지 아니한 경우에는 그러하지 아니하다. 한편, 법원은 손해의 발생은 인정되나 손해액을 산정하기 곤란한 경우에는 변론의 취지 및 증거조사 결과를 고려하여 상당한 손해액을 인정할 수 있다(제38조). 따라서, 표시를 하지 않은 경우에는 손해배상 청구권이 인정되지 않는다.⁴⁹⁾

AI 기본법은 사업자의 의무로써 표시를 요구하고 있다. 인공지능사업자는 고영향 인공지능이나 생성형 인공지능을 이용한 제품 또는 서비스를 제공하려는 경우 제품 또는 서비스가 해당 인공지능에 기반하여 운용된다는 사실을 이용자에게 사전에 고지하여야 한다. 또한, 인공지능사업자는 생성형 인공지능 또는 이를 이용한 제품 또는 서비스를 제공하는 경우 그 결과물이 생성형 인공지능에 의하여 생성되었다는 사실을 표시하여야 한다. 예술적인 목적으로 창작된 경우에는 예외를 두고 있다. 즉, 인공지능사업자는 인공지능시스템을 이용하여 실제와 구분하기 어려운 가상의 음향, 이미지 또는 영상 등의 결과물을 제공하는 경우 해당 결과물이 인공지능시스템에 의하여 생성되었다는 사실을 이용자가 명확하게 인식할 수 있는 방식으로 고지 또는 표시하여야 한다. 이 경우 해당 결과물이 예술적·창의적 표현물에 해당하거나 그 일부를 구성하는 경우에는 전시 또는 향유 등을 저해하지 아니하는 방식으로 고지 또는 표시할 수 있다. 다만, 사전고지, 표시, 고지 또는 표시의 방법 및 그 예외 등에 관하여 필요한 사항은 대통령령으로 위임하고 있다(제37조).

AI 기본법에서 요구하는 표시요건과 콘텐츠산업 진흥법에서의 표시요건의 성격에 차이가 있으나, 전자는 AI 사업자의 의무이고 후자는 콘텐츠사업자의 권리요건이라는 점에서 양자의 중복적인 적용여부도 검토될 필요가 있다.

49) 대법원 2005.7.14. 선고 2004도7962 판결.

2. 저작권법상 권리관리정보로서 워터마크

가. 권리관리정보

저작권법상 권리관리정보란 다음 각 목의 어느 하나에 해당하는 정보나 그 정보를 나타내는 숫자 또는 부호로서 각 정보가 저작권, 그 밖에 이 법에 따라 보호되는 권리에 의하여 보호되는 저작물등의 원본이나 그 복제물에 붙여지거나 그 공연·실행 또는 공중송신에 수반되는 것을 말한다.

가. 저작물등을 식별하기 위한 정보

나. 저작권, 그 밖에 이 법에 따라 보호되는 권리를 가진 자를 식별하기 위한 정보

다. 저작물등의 이용 방법 및 조건에 관한 정보

이러한 권리관리정보와 관련하여, 누구든지 정당한 권한 없이 저작권, 그 밖에 이 법에 따라 보호되는 권리의 침해를 유발 또는 은닉한다는 사실을 알거나 과실로 알지 못하고 다음 각 호의 어느 하나에 해당하는 행위를 하여서는 아니된다. 다만, 국가의 법집행, 합법적인 정보수집 또는 안전보장 등을 위하여 필요한 경우에는 적용하지 아니한다(제104조의3).

1. 권리관리정보를 고의로 제거·변경하거나 거짓으로 부가하는 행위
2. 권리관리정보가 정당한 권한 없이 제거 또는 변경되었다는 사실을 알면서 그 권리관리정보를 배포하거나 배포할 목적으로 수입하는 행위
3. 권리관리정보가 정당한 권한 없이 제거·변경되거나 거짓으로 부가된 사실을 알면서 해당 저작물등의 원본이나 그 복제물을 배포·공연 또는 공중송신하거나 배포를 목적으로 수입하는 행위

이처럼 권리관리정보는 주로 불법복제 등의 사후적 발견이나 적법한 이용을 위해 필요한 권리처리의 수행을 용이하게 하는 데 관련된 것이다.⁵⁰⁾ 실상 표시형태로 콘텐츠에 부가하는 것은 기술적 보호조치라기 보다는 권리관리정보로 볼 수 있다. 다만, 권리관리정보로 인정받기 위해서는 권리주체 등에 대한 정보를 표시할 필요가 있다.

나. 권리관리정보로서 워터마크

워터마크는 13세기 이탈리아에서 종이의 제조업체를 식별하기 위해 처음 사용되었으며, 이후 디지털 콘텐츠로 활용 영역을 넓히며 사용되어 왔다. 이미지 등의 매체에 로고나 텍

50) 박성호, 저작권법(제3판), 박영사, 2023, 705면.

스트 형태의 식별자를 추가하거나 사람이 인지할 수 없는 패턴을 삽입하여 콘텐츠의 출처 및 소유권 확인에 사용된다. 워터마크는 사용자 시각을 기준으로 인지 여부에 따라 인지 가능 워터마크 및 인지 불가능 워터마크로 분류된다. 특히, 인지 가능 워터마크는 생성물에 적용된 워터마크를 육안으로 구분 가능한 방법으로, 예를 들어 원본 이미지에 로고나 텍스트 등의 표시를 추가하게 된다.⁵¹⁾

디지털 형태의 워터마크는 AI 생성물에 표시하는 식별자로서 역할을 한다. 기술적 표시를 통해, 저작권 정보를 확인토록 하거나 불법적인 이용을 제한하는 역할을 하게 된다. 그렇지만, 워터마크는 삭제 삭제가 가능하다. 이미지에 표시된 워터마크는 해당 부분을 쉽게 지울 수 있기 때문이다.⁵²⁾ 삭제된 워터마크는 원래 소유자를 확인할 수 없는 상태가 되기 때문에 이는 워터마크라는 기술이 갖는 한계로 볼 수 있다. 이를 극복하기 위하여, 법적으로 워터마크를 삭제하거나 하는 경우에는 처벌토록 규정하고 있다. 즉, 워터마크가 저작권을 확인할 수 있는 권리관리정보로서 역할을 한다면, 해당 워터마크를 권리관리정보로 볼 수 있기 때문이다.

다. 법적 보완

저작권법은 권리관리정보를 개변할 경우에는 형사처벌토록 하고 있다. 저작물을 공정한 유통을 방해는 것이기 때문에 이에 대해 제재토록 하고 있다. 따라서, 워터마크를 임의로 개변하거나 삭제한다면 이는 저작권법 위반에 해당한다. 생성형 AI에 표시된 워터마크를 삭제하는 경우 저작권법에 따른 처벌이 이루어질 수 있다. 다만, 저작권법은 저작물에 해당하기 때문에 생성형 AI로 만들어진 결과물이 저작물이 아니라면, 이는 표시에 대한 개변조가 저작권법 위반에 해당하는 것으로 보기 어렵다.

AI 기본법이나 콘텐츠산업 진흥법 등 저작물성에 대한 논란이 없는 경우라면 저작물성에 대한 요건없이도 위법성을 따질 수는 있을 것이다. 현재 AI 기본법이나 콘텐츠산업 진흥법에서는 이러한 경우에는 처벌할 수 있는 규정이 없기 때문에 법적으로 보완되기 어렵다. 다만, 사업자의 의무규정으로 되어있기 때문에 표시만 할 수 있을 뿐이다. 향후, 입법을 통해 표시를 해제하는 경우에는 처벌할 수 있는 근거를 마련하는 것이 필요하다.

51) 신준호 외, 인공지능(AI) 워터마크 기술 동향 보고서, 한국정보통신기술협회, 2024, 16면.

52) 워터마크 정보의 경우 파일 내에 콘텐츠 데이터의 일부로 포함되므로 메타데이터와 같은 방법으로는 제거할 수 없다. 그러나 사이즈 축소, 잘라내기, 회전, 압축, 밝기 조정, 왜곡, 필터 적용 등의 이미지 편집 과정에서 워터마크 정보가 손상되거나 제거될 수 있다고 한다. 류시원, 인공지능 산출물 표시 규제 방안의 검토, 저스티스 통권 제205호, 2024, 527면.

라. 데이터 포이즈닝과 자력구제

데이터 포이즈닝(Data Poisoning)이란 데이터에 정당한 방법으로 학습데이터로 이용하지 않는 기계학습을 통해 생성된 AI 모델을 제대로 작동하지 않도록 데이터에 악성코드를 삽입하는 것이다. 기계학습 과정에 악의적인 데이터를 삽입하여 모델이 잘못된 패턴을 학습하게 만듦으로써, 구축된 AI 모델이 실제 환경에서 오작동하도록 유도하는 방식이다. 일종의 적대적 공격(Adversaria Attack)으로 볼 수 있다. 선해하자면, 저작자로서 자신의 저작물이 불법적으로 이용되는 것에 대해 스스로 저작권을 지키고자 하는 것으로 볼 수 있다.⁵³⁾

데이터 포이즈닝에 대해서는 2가지를 고려해야 할 것이다. 먼저, 데이터에 악성코드를 심는 것이 허용되는 행위인지 여부이다. 앞서 살펴본 권리관리정보로서 역할을 한다면 어느 정도 법에서 허용한 방식으로 볼 수 있다. 그렇지만, 권리관리정보로 보기 어렵다면 이는 문제가 될 수 있다. 또한, 자력구제 형식으로 데이터가 AI 모델을 해킹하거나 제대로 작동하지 못하도록 한다면 이는 업무방해에 해당할 수 있다. 이 경우에는 형법상 컴퓨터업무방해죄에 해당하여, 형사처벌을 받을 수 있다는 점도 고려해야 한다. 또한, 이러한 방식을 기술적 자력구제라고 하지만, 자력구제는 민법이나 형법에서 엄격하게 금지하는 사적구제 행위이다. 사적구제를 금지하는 것은 사법체계에 부정적 영향을 줄 수 있기 때문이다.

데이터포이즈닝된 데이터를 제공하는 것은 데이터 자체를 이용하도록 하기 때문에 적극적인 이용허락, 최소한 묵시적 이용허락을 했다고 볼 수도 있다. 그러한 신뢰의 관계에서 이용한 경우에 대한 하자는 책임이 없다. 그렇지만, 고의로 데이터에 오염원을 내재시키고 이를 선의의 제3자가 이용토록 한 것이다. 문제가 될 것이라는 것을 사전에 인지하고 있으며, 예견가능성을 알고 있었다고 이해된다. 이는 신뢰이익을 저버린 것으로 업무방해 및 손해배상 책임을 질 수 있다.

3. 데이터 및 출처의 공개

가. 의의

데이터의 공개는 저작권에 대한 윤리적인 이용여부를 확인할 수 있는 방법이다. 생성 AI가 저작권법을 위반하여 의사 표현을 만드는 데 얼마나 자주 사용되는지 추정하기 어렵다. 관련 데이터가 공개되지 않았기 때문이다. 관련 정보가 공개된 경우에도 훈련 데이터의

53) Visger, Mark, Garbage In, Garbage Out: Data Poisoning Attacks and their Legal Implications(October 27, 2022). Big Data and International Humanitarian Law(published by the Lieber Institute), Forthcoming, Available at SSRN: <https://ssrn.com/abstract=4260456>.

어떤 것과 너무 유사한지 여부를 판단하기 위해 출력을 분석해야 하며 그렇다면 그럴 가능성이 있는지 여부를 판단해야 한다.⁵⁴⁾ 데이터를 공개한다고 하더라도, 이를 분석하는 절차가 더욱 복잡할 것이다. 따라서, 데이터 공개보다 중요한 것은 데이터에 대한 구체적인 정보를 기재토록 하는 것이다. EU AI법에서는 범용인공지능모델의 학습에 사용하는 콘텐츠에 관한 충분히 상세한 요약물 작성하고 공개토록 요구하고 있다. 뒤에서 설명하겠지만, 데이터의 공개보다는 실제 데이터의 구성에 관한 정보를 요구하는 것이 보다 합리적일 수 있다고 본다.

또한, 데이터 공개의 의의는 신뢰성, 재현가능성을 확보하기 위한 방법이다. 이러한 공개를 통하여 데이터에 대한 투명성을 신뢰할 수 있게될 것이다. 만약, 데이터가 윤리적인 방법으로 확보되지 않을 경우 해당 데이터로 학습한 AI 모델에 대한 법적 위험으로 인하여 서비스 안전성이 저해될 수 있기 때문이다.

나. 데이터 공개 논의

EU AI법에서는 데이터 공개를 요구하는 것은 아니지만, 학습데이터 위험을 관리하기 위해 데이터 관리 및 학습데이터 요약본 공개 의무를 규정하고, 기술적 위험 관리를 위해 기술문서 작성과 적합성 평가 등을 의무화하고 있다. 저작권에서 보호하는 텍스트 및 데이터를 포함한 범용인공지능모델의 사전학습 및 학습에 사용되는 데이터에 관한 투명성을 높이기 위하여 해당 모델의 제공자는 범용인공지능모델의 학습에 사용된 콘텐츠에 관한 충분히 상세한 요약물 작성하고 공개하는 것이 적절하다고 밝히고 있다.⁵⁵⁾ 범용인공지능모델 제공자는 인공지능사무국이 제공하는 양식에 따라 범용인공지능모델의 학습에 사용하는 콘텐츠에 관한 충분히 상세한 요약(a sufficiently detailed summary)을 작성하고 공개하여야 한다(제53조 제1항 (d)). 고위험 AI 시스템의 공급자는 데이터 관리, 기술문서 작성, 배포자에 이용지침 제공, 적합성 평가 수행 등의 의무를 이행해야 하고, 배포자는 이용지침을 준수하고 기본권 영향평가 등을 수행해야 한다. 생성형 AI를 포함하는 범용 AI(GPAI)의 공급자는 기술문서와 사용설명서를 제공, 저작권 지침 준수, 학습에 사용된 데이터의 요약본 공개 등의 의무를 이행해야 한다. AI 기본법에서는 데이터 공개에 대한 규정을 두고 있지 않다. 다만, 논의과정에서 문화부는 EU AI법⁵⁶⁾과 같이 학습 데이터 공개에 관한 규정을 포함

54) Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593>. p.336.

55) EU AI법 recital 107.

할 것을 요구하였으나 반영되지 않았다.⁵⁷⁾

다. 데이터 공개의 기대효과

기계학습 과정에서 데이터윤리에 위배되는 행위는 다양하다. 무엇보다, 저작권자의 허락 없이 데이터를 수집하여 기계학습에 이용하는 것이다. 물론, 이 과정에서 공정이용이 될 가능성도 상당히 높다. 다만, 독일이나 미국의 판례가 서로 다른 결과를 가져오고 있지만, 이는 고민이 필요하다. TDM이 허용되고 있지 않지만, 공정이용 규정을 두고 있는 우리 저작권법의 해석도 상이하다. 만약, 저작권법을 개정하여 출처 등을 공개토록 할 경우에 나타날 수 있는 우려는 없는지도 검토되어야 한다. EU AI법 제정과정이나 우리 AI 기본법 제정과정에서도 데이터의 공개가 요구된 바 있다. EU는 규제기관이 서버에 접근하여 해당 내용을 충분히 파악할 수 있는 안이 제안된 바 있다. 물론, 이는 충분한 내용을 공개하는 것으로 변경된 바 있다.

따라서, 데이터를 공개하는 등의 제도가 도입될 경우에는 그로 인하여 소송이 확산할 수 있다는 점을 고려할 필요가 있다. 소송이 확산하지 않도록 제한규정을 두거나 공개에 따른 인센티브를 사업자에게 제공하는 방안도 검토될 필요가 있다.⁵⁸⁾

4. 게시중단의 한계와 머신 언러닝

가. 저작권법의 요구로써 데이터 또는 게시물의 삭제

기계학습 과정에서 다양한 데이터가 사용되는 것은 사실이고, 경우에 따라서는 데이터를 삭제해야할 상황이 발생할 수도 있다. 개인정보나 저작권 침해의 경우에 해당 데이터의 삭제가 요구될 수 있기 때문이다. 인간 피드백 기반 강화학습(HLRF) 방식은 모델의 출력을 사후적으로 필터링하는 것보다 효과적이다.⁵⁹⁾ 이외에 게시중단(notice & take down)은 특

56) EU AI 법 제53조(범용 인공지능 모델 제공자의 의무) 1. 범용 인공지능 모델 제공자는 다음 사항에 따른다.(d) 인공지능사무국이 제공하는 양식에 따라 범용 인공지능 모델의 학습에 사용하는 콘텐츠에 관한 충분히 상세한 요약물 작성하고 공개한다.

57) 문화체육관광부는 저작권 생태계 및 창작자 권리 보호와 학습데이터의 투명성 확보를 위해, EU 인공지능법의 경우와 같이 인공지능사업자에 대한 학습데이터 목록 공개 의무 부과가 필요하다는 의견을 제출하였다. 이에 대하여 과학기술정보통신부는 EU 인공지능법에서 작성하도록 하는 상세한 요약본의 내용 및 서식 등이 공개되지 않은 상황에서 우리나라가 선제적으로 구체적인 조치를 하는 것은 곤란하며, EU 외 다른 국가의 동향을 함께 고려하여 단계적으로 보완 입법하는 것이 바람직하다는 입장이다. 김건오, 인공지능 산업 육성 및 신뢰 확보에 관한 법률안 등 검토보고, 국회과학기술정보방송통신위원회, 2024.8. 30면.

58) 김윤명, AI 관련 발명의 공개와 데이터 기탁제도, *홍익법학* vol.24, no.2, 2023, 291~324면.

정 URL이 있어야 가능하다. 포괄적인 키워드 형태의 게시중단은 저작권법의 요건을 갖춘 것으로 보기 어렵다.

게시중단은 URL을 특정해야 진행될 수 있지만, 생성형 AI가 창작한 결과물은 설명 저작권 침해물이 생성되었다고 하더라도 이를 특정하기 어렵다. 이러한 한계 때문에 학습된 모델의 일부를 되돌리기(roll back)하는 방식의 언러닝(unlearning) 학습방법이 제안되고 있다.⁶⁰⁾ 좀더 구체적으로 한다면, 특정 학습데이터의 흔적을 지우는 방식으로 이해할 수 있다. 보다 많은 연구가 이루어져야 하겠지만 특정 영역의 데이터를 지울 경우, 다른 영역의 기억에 대해서 까지 영향을 미칠 수 있을 것으로 보인다. 다만, 몇몇 연구결과에 따르면 큰 영향을 미칠 것은 아니라고 한다.

나. 머신 언러닝 방식

인스턴스 언러닝은 신생 연구 분야이며 삭제된 데이터 포인트 없이 모델을 재교육하는 것은 비용이 많이 들 수 있다. 단기적으로 가장 가능성 있는 접근 방식은 기초 모델이 쉽게 삭제할 수 있는 영구 콘텐츠를 호스팅하는 대신 사용자에게 대한 수요에 따라 콘텐츠를 생성하기 때문에 필터링 메커니즘을 통해 저작권이 있는 작품과 너무 유사한 모델 출력을 삭제하는 것이다. 비교적 새로운 환경에서 삭제 요청을 처리하기 위한 새롭고 개선된 메커니즘을 식별하기 위한 새로운 연구가 필요하다.⁶¹⁾

AI 모델은 학습된 상태이기 때문에 이를 되돌리기는 쉽지 않다. 인간도 학습한 내용을 기억에 남겨지기 때문에 무의식적으로도 해당 내용이 생각나기 마련이다. 기계도 다르지 않을 것이다. 따라서, 해당 기억을 지우는 것은 언러닝 모델을 통해서 이루어질 수 있다고 하더라도, 아직은 정확도 높은 학습으로 보기 어렵다. 설령 가능하더라도, 비용이 소요될 수 있을 것이며 경우에 따라서는 삭제가 불가능하거나 비용이 소요되는 경우에 해당할 수 있을 것이다. 이는 저작권법에서 허용하는 온라인 서비스 제공자(OPS)의 면책에 해당할 수

59) Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593>, p.339.

60) Bensaoud, Ahmed and Kalita, Jugal, Advancing Machine Unlearning: A Novel Approach for Efficient and Effective Data Removal. Available at SSRN: <https://ssrn.com/abstract=5023189>; Hine, Emmie et al., Supporting Trustworthy AI Through Machine Unlearning(November 24, 2023). Science and Engineering Ethics, volume 30, issue 5, 2024, Available at SSRN: <https://ssrn.com/abstract=4643518>.

61) Henderson, Peter et al., Foundation Models and Fair Use (March 27, 2023). Stanford Law and Economics Olin Working Paper No. 584, Available at SSRN: <https://ssrn.com/abstract=4404340>, p.19.

있다. 즉, OSP에게 게시중단을 요청하는 것이 기술적으로 불가능한 경우에는 게시중단 자체도 어렵기 때문이다.⁶²⁾

언러닝은 신생 연구 분야이며 삭제된 데이터 포인트 없이 모델을 재교육하는 것은 엄청나게 비용이 많이 들 수 있다. 단기적으로 가장 가능성 있는 접근 방식은 기초 모델이 쉽게 삭제할 수 있는 영구 콘텐츠를 호스팅하는 대신 사용자에게 대한 수요에 따라 콘텐츠를 생성하기 때문에 필터링 메커니즘을 통해 저작권이 있는 작품과 너무 유사한 모델 출력을 삭제하는 것이다. 그러나 이 비교적 새로운 환경에서 삭제 요청을 처리하기 위한 새롭고 개선된 메커니즘을 식별하기 위한 새로운 연구가 필요하다.⁶³⁾

이러한 방식 이외에 모델이 이미 저작권이 있는 자료에 대해 학습되었고 배포될 예정이라고 가정할 때, 학습 데이터가 재생산되는 것을 방지하기 위한 간단한 아이디어 중 하나는 모델 추론 중에 필터를 적용하여 학습 데이터를 미러링하는 모든 출력을 감지할 수 있도록 하는 것이다.⁶⁴⁾ 필터링을 통해 생성되는 데이터에 대한 내용을 통제함으로써 윤리적으로 지 않은 결과물이 생성되지 않도록 하는 것이다.

문제는 데이터가 만들어내는 다양성이 제한될 수 있다. 실령 학습과정에서 다양한 데이터를 바탕으로 학습이 이루어졌다고 하더라도, 출력단에서 다양성이 훼손되는 필터링이 이루어질 경우라면 AI 모델의 편향성은 확장되고 고착화할 수 있기 때문이다. 따라서, 출력 필터링은 입력 데이터의 다양성 만큼이나 중요한 영역이다. 물론, 필터링이 윤리적인 가치를 위한 것이라는 것과 명확한 기준을 통해 이루어질 경우라면 어느정도 신뢰성을 가질 수는 있을 것이다. 그렇지만, 여전히 데이터의 다양성은 필요하다. 이용자가 이용하는 과정에서 데이터의 편향보다는 편향을 수 있기 때문이다.

다. 게시중단의 기술적인 한계

불법적인 게시물 등에 대해 저작권법은 URL을 특정하여, 게시중단을 요청할 경우 OSP

62) 저작권법 제102조(온라인서비스제공자의 책임 제한) ② 제1항에도 불구하고 온라인서비스제공자가 제1항에 따른 조치를 취하는 것이 기술적으로 불가능한 경우에는 다른 사람에 의한 저작물등의 복제·전송으로 인한 저작권, 그 밖에 이 법에 따라 보호되는 권리의 침해에 대하여 책임을 지지 아니한다. <개정 2011. 6. 30.>

③ 제1항에 따른 책임 제한과 관련하여 온라인서비스제공자는 자신의 서비스 안에서 침해행위가 일어나는지를 모니터링하거나 그 침해행위에 관하여 적극적으로 조사할 의무를 지지 아니한다. <신설 2011. 6. 30.>

63) Henderson, Peter et al., Foundation Models and Fair Use(March 27, 2023), Stanford Law and Economics Olin Working Paper No. 584, Available at SSRN: <https://ssrn.com/abstract=4404340>, p.19.

64) Henderson, Peter et al., Foundation Models and Fair Use(March 27, 2023), Stanford Law and Economics Olin Working Paper No. 584, Available at SSRN: <https://ssrn.com/abstract=4404340>, p.22.

는 주의의무를 규정하고 있다. 대체적으로 일반적인 게시물은 URL을 특정할 수 있다. 따라서, OSP는 해당 주소의 게시물을 처리할 수 있다. 그렇지만, AI 시스템을 이용하는 과정에서 게시 중단에 해당하는 결과물이 생성될 경우에 이를 특정하기가 쉽지 않다. AI 모델에 특징점의 형태로 저장된 정보를 삭제하는 것은 불가능에 가깝다. 우리 저작권법도 게시 중단이 기술적으로 불가능할 경우에 OSP는 다른 사람에 의한 저작물등의 복제·전송으로 인한 저작권, 그 밖에 이 법에 따라 보호되는 권리의 침해에 대하여 책임을 지지 아니한다(제102조 제2항). OSP 면책규정에서 기술적으로 불가능한 경우에는 면책이 이루어지고 있으나, 향후 머신 언리닝 방식이 실효성있는 수준으로 발전할 경우에는 기술적으로 불가능한 경우로 보기 어려운 상황에 이를 수 있다고 본다. 이러한 상황에서는 면책요건으로써 기술적인 불가능성은 적용되기 어려울 것이다.

5. 보상체계로서 보상청구권

가. 데이터 윤리와 보상

학습데이터에 대한 정당한 보상이 이루어진다면 법적이거나 윤리적으로 문제될 것은 많지 않다. 다만, 그 이용이 제대로 된 보상이 없이 이루어지고 있다는 점에서 윤리적인 측면에서 비난받을 수 있다는 점이다. 저작권법의 위배 여부에 대해서는 대법원의 결론이 내려지지 않았기 때문에 위법하다고 단정할 수 없는 것이다. 만약, 저작권자에게 적절한 보상이 이루어질 수 있다면 소송이라는 비용을 들이지 않고 법적인 위험이 없는 상태에서 AI 모델을 구축할 수 있을 것이다.

데이터 보상에서 고려할 수 있는 권리자는 2가지 유형이다. 사업자형 권리자, 개인형 권리자로서, 전자는 신탁관리단체 등 전문적인 저작권자를 들 수 있다. 후자는 SNS에 올려지는 수많은 데이터에 대한 권리를 갖고 있는 개인들이다. 이들은 개인이지만, 양적인 측면에서는 사업자형 권리자보다 많을 수 있다. 그렇지만, 교섭력이 없기 때문에 제대로 된 보상이 이루어질 가능성은 크지 않다. 그렇지만, 이들은 가장 정제된 데이터를 생성하고 있다는 점에서 무임승차의 영역에서 제외되어서는 곤란하다고 본다. 예를 들어, X에 올려지는 사진은 가장 정제성이 높은 데이터라고 한다. 별도 가공이 이루어질 필요가 없다. 특히, 사진과 그 사진에 대한 태그(tag)나 설명은 가장 정확한 레이블링이 될 수 있기 때문이다. 다만, 어떻게 구현할 것인지에 대한 논란이 남아있다. 이를 해결하기 위해서는 제도적 기술적 방안이 강구되어야 할 것이다.

나. 보상체계의 강구

방송사업자들이 음악저작물을 사용하는 경우에 보상청구권이 권리자에게 인정되고 있다. 참고로, 보상청구권의 일 유형인 사적복제보상은 공매체(CD, DVD 등)에 수수료를 징수하여 사적복제에 따른 잠재적 손실에 대한 보상으로 저작권자에게 배분하는 제도이다. 이 시스템은 저작권 보유자에게 일종의 보상을 제공할 수 있지만 공정이용 또는 특정 사용에 대한 예외 문제를 다루지는 않는다.

공정이용은 교육, 비평, 논평 및 뉴스 보도와 같이 저작권 소유자의 허가 없이 저작물의 특정 사용을 허용하므로 저작권법의 중요한 측면이다. 공정한 사용이 없으면 개인과 기업이 저작권 침해에 대한 두려움 때문에 콘텐츠를 사용하거나 만드는 것을 주저할 수 있으므로 창의성과 혁신의 잠재력이 제한될 수 있다. 더욱이 사적 복사에 대한 보상에만 초점을 맞추는 법률이 반드시 저작권과 관련된 광범위한 경제 및 산업 발전 문제를 다루는 것은 아니다. 예를 들어, 공정이용은 저작물을 새롭고 혁신적인 방식으로 사용하기 위한 틀을 제공함으로써 저널리즘, 교육 및 예술과 같은 산업의 성장을 장려할 수 있다. 전반적으로 저작권 소유자에게 일종의 보상을 제공할 수 있지만 공정이용 및 기타 저작권 예외의 광범위한 사회적 혜택을 고려하는 것도 중요하다.

V. 결론 및 시사점

1. 결 론

저작권법은 창작자의 권리를 보호하는 동시에, 공정한 이용을 보장하여 사회적 이익과 문화·산업 발전의 균형을 도모한다. 그러나 AI 기술의 발전으로 인해 대량의 데이터 활용이 필수가 되면서, 저작권법과 데이터 윤리 사이의 충돌이 심화되고 있다. 특히, 공정이용(Fair Use)과 텍스트·데이터 마이닝(TDM)의 법적 한계는 AI 학습 과정에서 중요한 쟁점으로 부각되고 있다. AI가 방대한 데이터를 학습하는 과정에서 저작권법이 어떻게 적용될 것인가에 대한 논쟁이 지속되고 있으며, 공정이용이 AI 학습 데이터에 적용될 수 있는 범위에 대한 명확한 기준이 요구된다. 또한, 데이터 윤리는 AI 모델의 학습과 결과물의 공정성을 보장하는 중요한 요소로 작용하지만, 윤리적 고려가 법적 규제와 어떻게 조화를 이루어야 하는지는 여전히 불확실하다.

AI 학습을 위한 데이터 수집 과정에서 저작권법과 데이터 윤리는 밀접하게 연결된다. 빅데이터 처리와 AI 학습 과정에서는 데이터의 복제, 분석 등이 필수적으로 발생하며, 이를

무분별하게 활용할 경우 저작권 위반 행위가 될 수 있다. 특히, TDM의 특성상 불특정하게 수집된 데이터 안에 저작물이 포함될 가능성이 높아 저작권 보호제도와 충돌할 가능성이 크다. 이러한 문제는 단순한 법적 갈등을 넘어서 데이터 수집의 윤리성, 개인정보 보호, 그리고 창작자의 권리 존중과 같은 윤리적 고려까지 포함된다.

공정이용 원칙은 AI 학습을 위한 데이터 활용에서 중요한 역할을 하지만, 법원의 판례를 통해 구체적인 적용 범위가 확립되어야 한다. 공정이용은 법적인 근거를 갖지만, 일반 조항으로서 명확한 기준을 제시하는 것이 아니라 법원이 개별 사안에 따라 해석하는 방식으로 적용된다. 따라서, 공정이용이 AI 학습에 적용될 수 있는 구체적인 요건과 한계를 명확히 하는 것이 필요하다.

AI 윤리나 데이터 윤리가 저작권법에 지나치게 깊이 개입할 필요는 없겠지만, 법 해석과 법 개정에 미치는 영향은 점점 커질 것으로 보인다. 특히, AI가 창작 활동에 본격적으로 개입하면서 ‘비인간 저작자’의 개념이 등장할 가능성이 있으며, 이는 기존 저작권 체계에 도전하는 요소가 될 수 있다. 따라서 AI와 저작권법의 관계를 설정하는 과정에서 윤리적 논의를 병행하는 것이 필수적이며, 윤리와 법이 조화를 이루는 방식으로 접근해야 한다.

데이터 윤리와 저작권법은 기술 혁신과 개인의 권리 보호 사이에서 균형을 유지하는 방향으로 발전해야 한다. AI의 학습 데이터가 인터넷에서 크롤링된 정보로 구성될 경우, 이용 허락 조건을 위반하는 방식으로 활용된다면 저작권 침해 문제에서 자유로울 수 없다. 따라서 AI 기술 발전이라는 사회적 가치와 창작자의 권리 보호라는 가치 사이의 균형점을 찾는 것이 중요한 과제가 된다. AI 관련 법제는 기존 저작권법의 해석과 적용뿐만 아니라, 기술 변화에 유연하게 대응할 수 있도록 사회적 합의를 반영한 윤리적 고려까지 포함해야 한다. AI 모델의 고도화를 위해 지속적인 데이터 공급이 필요하지만, 수집된 데이터의 오남용으로 인한 특정 집단에 대한 편견과 차별, 사회적 불평등 심화 등의 문제를 해결하기 위한 법적·윤리적 책임도 함께 논의되어야 한다.

2. 시사점

데이터 윤리가 확립될 경우 AI 시스템에 대한 신뢰를 확보할 수 있으며, 이는 데이터의 수집 및 이용 과정에서 저작권법을 준수하는 것이 필수적임을 의미한다. 법적 규정을 준수하는 데이터 활용 방식은 AI 개발 기업과 서비스 제공자가 장기적으로 법적 분쟁을 예방하는 데 기여할 수 있으며, 이용자의 신뢰를 얻는 기반이 된다. 또한, 윤리적인 데이터 활용은 AI의 성능 향상에도 기여할 가능성이 높다.

저작권이 있는 데이터의 윤리적 이용은 데이터의 질적 가치를 높이는 효과를 가져올 수 있으며, 데이터 처리 과정에서 적절한 정제 과정을 거치게 함으로써 윤리적 문제를 최소화할 수 있다. 특히, 데이터 편향성 문제를 해결하는 데 기여할 수 있다. AI 모델이 편향된 데이터를 학습하면 알고리즘이 편향된 결과를 내놓을 가능성이 크며, 이는 사회적 불평등을 심화하는 결과로 이어질 수 있다.

데이터 편향 문제를 해결하기 위한 방안으로는 가치중립적인 키워드와 변수를 설정하여 보다 공정한 데이터 선별과 이용이 가능하도록 하는 접근이 있다. 또한, AI 모델이 학습하는 데이터의 다양성을 확보하여 편향성을 줄이는 방법도 고려해야 한다. 예를 들어, 성별, 연령, 국적 등 다양한 인구통계학적 특성을 반영한 데이터 세트를 구축하는 것이 바람직하다. AI 모델이 설명 가능성(Explainability)을 갖추도록 연구하는 것도 데이터 윤리 실현의 중요한 방향이 될 수 있다.

데이터 윤리를 보장하기 위해서는 법적 규제와 함께 정책적 대응도 중요하다. AI 개발 기업이 데이터 윤리를 준수하도록 유도하기 위해 자율규제 모델을 도입할 수도 있으며, 정부 차원의 데이터 윤리 가이드라인을 강화하는 방안도 고려할 수 있다. 데이터 편향성이 사회적 차별을 고착화하거나 그러한 가능성이 높은 경우, 해당 알고리즘을 개발·운영하는 기업에 대한 강력한 규제가 필요할 것이다. 반면, 의도적이지 않은 데이터 편향에 대해서는 기술적·정책적 유연성을 유지하면서 개선 방안을 마련하는 것이 중요하다.

마지막으로, AI 모델의 공정성과 신뢰성을 확보하기 위해서는 다양한 데이터를 학습하도록 하는 것이 필수적이다. 특정 집단이나 특정 속성만을 반영한 데이터로 AI를 훈련할 경우, 모델이 갖는 편향성은 더욱 심화될 수밖에 없다. 따라서 다양한 관점과 경험을 반영하는 데이터 세트를 구축하는 것이 필요하며, 데이터 접근성과 품질을 동시에 고려한 법적·윤리적 기준을 마련해야 한다. 이러한 조치를 통해 AI 기술 발전과 저작권법의 조화를 이루고, 기술혁신과 사회적 가치를 동시에 실현할 수 있는 기반을 마련할 수 있을 것이다.

참고문헌

- 고유강, 법관업무의 지원을 위한 머신러닝의 발전상황에 대한 소고, LAW&TECHNOLOGY 제15권 제5호(통권 제83호), 2019.
- 김진오, 인공지능 산업 육성 및 신뢰 확보에 관한 법률안 등 검토보고, 국회과학기술정보방송통신위원회, 2024. 8.
- 김유환, 행정법과 규제정책, 법문사, 2012.
- 김윤명, AI관련 발명의 공개와 데이터 기탁제도, 홍익법학 vol.24, no.2, 2023.
- 김윤명, 발명의 컴퓨터 구현 보호체계 합리화를 위한 특허제도 개선방안 연구, 특허청, 2014.
- 김윤명, 알고리즘과 법, 정보화진흥원, 2019.
- 김윤명, 알고리즘과 법: 알고리즘 차별에 따른 공정성 확보방안, 정보화진흥원, 2019.
- 김윤명, 현재의 선택과 선의의 기술로서 AI, SW중심사회 no.69, 2020.
- 김진형, AI최강의 수업, 매일경제신문사, 2020, 270면.
- 남형두, 공정이용의 역설, 경인문화사, 2025.
- 류시원, 인공지능 산출물 표시 규제 방안의 검토, 저스티스 통권 제205호, 2024.
- 문덕수, 시쓰는 법(重版), 동원출판사, 1984.
- 박균성, 정책, 규제와 입법, 박영사, 2022.
- 박성호, 저작권법(제3판), 박영사, 2023.
- 소병수, 미디어 영역에서 인공지능(AI) 기술의 확산과 제문제, 원광법학 제39권 제3호, 2023.
- 신준호 외, 인공지능(AI) 워터마크 기술 동향 보고서, 한국정보통신기술협회, 2024.
- 양천수 외, 현대 빅데이터 사회와 새로운 인권 구상, 안암법학 57권, 2018.
- 이원태 외, 지능정보사회의 규범체계 정립을 위한 법제도 연구, 2016.
- 정상조, 인공지능시대의 저작권법 과제, 계간저작권 통권 122호, 2018.
- 허유선, 인공지능에 의한 차별과 그 책임 논의를 위한 예비적 고찰 - 알고리즘의 편향성 학습과 인간 행위자를 중심으로, 한국여성철학 제29권, 2018.
- Annemarie Bridy, Coding Creativity: Copyright and the Artificially Intelligent Author, 2012 Stan. Tech. L. Rev. 5.(2012).
- Bensaoud, Ahmed and Kalita, Jugal, Advancing Machine Unlearning: A Novel Approach for Efficient and Effective Data Removal. Available at SSRN: <https://ssrn.com/abstract=5023189>.
- Henderson, Peter et.al., Foundation Models and Fair Use(March 27, 2023). Stanford Law and Economics Olin Working Paper No. 584, Available at SSRN: <https://ssrn.com/abstract=4404340>. p.22.

Hine, Emmie et al., Supporting Trustworthy AI Through Machine Unlearning(November 24, 2023). Science and Engineering Ethics, volume 30, issue 5, 2024, Available at SSRN: Mark A. Lemley, Bryab Casey, Fair Learning(2020), at <http://ssrn.com/abstract=3528447>.

Safiya Umoja Noble, 구글은 어떻게 여성을 차별하는가, 한스미디어, 2019.

Sag, Matthew, Copyright Safety for Generative AI(December 3, 2023). Forthcoming in the Houston Law Review, Houston Law Review, Vol. 61, No. 2, 2023, Available at SSRN: <https://ssrn.com/abstract=4438593>.

Sobel, Benjamin, Artificial Intelligence's Fair Use Crisis(September 4, 2017). 41 Colum. J.L. & Arts 45(2017), Available at SSRN: <https://ssrn.com/abstract=3032076>.
<https://ssrn.com/abstract=4643518>.

U.S. Copyright Office, Copyright and Artificial Intelligence Part 2. 2025.

Visger, Mark, Garbage In, Garbage Out: Data Poisoning Attacks and their Legal Implications(October 27, 2022). Big Data and International Humanitarian Law (published by the Lieber Institute), Forthcoming, Available at SSRN: <https://ssrn.com/abstract=4260456>.