

생성형 AI 저작권과 이중구조 보상 모델:

‘초능력 문학소녀’를 위한 새로운 문화 사회계약*

연세대학교 일반대학원 법학과 박사 수료(국제법 전공) **염 규 현****

*논문접수 : 2026. 4. 19. *심사개시 : 2026. 4. 22. *게재확정 : 2026. 5. 6.

〈목 차〉

I. 서론 : TDM 보상 체계가 갖는 구조적 공백	1. 일회성 보상 패러다임이 갖는 문제점
II. 파운데이션 모델과 저작권법의 불일치	2. 이중 구조 보상 모델의 이론적 구성
1. AI 학습의 기술적 메커니즘	IV. 산출물 기반 과세적 보상에 대한 이론적 정당화
2. 학습 데이터와 생성 결과물의 관계	1. 예상 반론에 대한 이론적 재반론
3. TDM 관련 주요국 법제 및 공통적 맹점	2. 제도 실현을 위한 실행 로드맵
III. 보상과 과세 분리 방안 제안 : 이중구조 보상 모델	V. 결론 : 생성형 AI시대와 사회계약의 재구성

I. 서론 : TDM 보상 체계가 갖는 구조적 공백

2022년 11월 챗GPT가 처음 등장한 이후, 생성형 AI는 기하급수적 발전을 거듭하며 급기야 창작과 저작권의 영역을 흔들고 있다. 이는 관련 규범에 대한 전면적인 재검토의 필요성을 제기하고 있다. 거대언어모델(Large Language Model, LLM)은 방대한 텍스트 데이터를 학습하여 번역, 요약부터 소설 창작 등의 글쓰기까지 다양한 텍스트 산출물을 제공한다. 이 과정에서 방대한 분량의 웹 문서를 무단 학습했다는 논란이 제기되면서 전 세계적으로 ‘텍스트 데이터 수집(Text and Data Mining, TDM)에 관한 법적 규제 논의가 이뤄지며, 국제 법무의 중요 의제로 부상하고 있다.¹⁾

* 본 연구는 문화체육관광부 및 한국콘텐츠진흥원의 2026년도 문화체육관광 연구개발사업으로 수행되었음.
(과제명: AI 기반 실시간 콘텐츠 제작 및 글로벌 유통을 위한 실·가상 융합 방송 자동 영상 생성 기술개발,
과제번호 : RS-2024-00349534, 기여율: 100%)

그런데, 현재 진행 중인 주된 논의는 빅테크들이 AI 모델 학습 과정에서 데이터를 무단 사용한 부분에 초점이 맞춰져 있다. 쉽게 비유하자면 AI가 공부를 할 때 ‘참고서 사는 비용’을 지불해야 하는가. 한다면 ‘얼마를 지불해야 하는가’에 대한 문제가 주가 되고 있는 것이다.²⁾

지난해 9월, 미국의 생성형 AI 모델 ‘클로드(Claude)’ 개발사 앤스로픽(Anthropic)은 작가들이 저작권 침해 소송을 제기한 것에 대해 15억 달러, 우리돈 2조 원을 지급하기로 합의한 바 있는데,³⁾ 이 사례를 참고하면 주요 인공지능 테크 기업들도 일단 큰 틀에서는 ‘참고서 비용’을 지불하자는 추세로 가고 있는 것으로 보인다.

그러나, 필자는 본고를 통해 이런 추세에 한 가지 의문을 추가하고자 한다. 바로, ‘참고서 비용’만 내면 과연 끝날 것인가 하는 점이다. 학습이 완료된 이후에도 배분할 정당한 가치에 대한 문제를 짚어보고자 하는 것이다. 만일, 한 문학소녀가 소설책과 문학 참고서를 사서 공부한 뒤 그 지식을 바탕으로 평생 작가로 활동하며 돈을 번다면 참고서 저자에게는 책값 이외에 아무런 대가도 돌아가지 않는 것이 타당하다. 학습은 스스로 한 것이며, 학습 결과물로 얻은 능력도 개인의 것이기 때문이다. 인간 창작자의 경우, ‘참고서 비용’만 지불해도 문제가 되지 않는다.

반면, AI의 경우에는 상황이 다르다. 우선, 인공지능은 인간이 평생동안 습득할 분량을 몇 달 안에 학습할 수 있다. 또, 학습의 결과물이 새로운 대체 콘텐츠로 대규모 생성되어 시장에 배포될 수 있다. 그로 인해, 굉장히 빠른 속도로 기존 창작자들의 시장을 잠식하고

** MBC AI전략자회사 (주)도스트일레븐 경영기획본부장(CSO), 한양대학교 미디어커뮤니케이션학과 겸임교수, 연세대학교 국제경제법 박사과정 (AI 국제규제, EU AI Act 연구)

1) 세계지적재산권기구(WIPO)에서도 11차례에 걸쳐 인공지능과 저작권 문제에 관한 회의를 진행하며, 이에 대한 법적 논의를 이어가고 있다. (WIPO, “Artificial Intelligence and Intellectual Property”, webpage) <https://www.wipo.int/en/web/frontier-technologies/artificial-intelligence/index>. EU의 디지털 단일시장 저작권 지침(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)에서도 TDM 문제에 대해 초점을 맞추고 있다.

2) 단순한 창작 콘텐츠 뿐 아니라 학술적 콘텐츠들까지도 무단 학습 논란이 불거지고 있다. (공인희 편집인, “AI 학습 데이터 논란” 연구자들 “우리 논문 무단 사용됐다”, AI매터스 2024년 8월 30일 보도.) <https://aimatters.co.kr/news-report/ai-report/2916/>

3) 김광연 기자, “클로드 운명사 앤스로픽, ‘저작권’ 소송 작가들에게 2조원 지급기로”, IT조선, 2025년 9월 6일 보도. <http://it.chosun.com/news/articleView.html?idxno=2023092146990>

피해를 끼칠 수 있는 것이다.⁴⁾ 그래서 필자는 이런 AI 특징을 반영해 AI 언어모델을 ‘초능력 문학소녀’라 명명하였다. 즉, 단순한 ‘문학소녀’와 ‘초능력 문학소녀’는 그 대우를 달리해야 한다는 문제의식에서 본고의 중심 생각은 출발한다.

이에 따라, 본고에서는 TDM에 대한 보상은 학습 단계에서 끝나는가, 아니면 학습을 통해 획득한 능력의 지속적 활용에까지 미쳐야 하는지 짚어보고 이러한 구분에 따른 차등적인 규율 체계를 새롭게 제안하고자 한다.

II. 파운데이션 모델과 저작권법의 불일치

1. AI 학습의 기술적 메커니즘

앞서 서론에서 언급한 ‘초능력 문학소녀’의 비유가 적합한지 따져보기 위해서는 AI의 모델 학습 메커니즘에 대해 먼저 살펴볼 필요가 있다. 특히, 텍스트 데이터 수집과 관련하여 문제가 되고 있는 거대언어모델(LLM)이 어떻게 학습하고 생성하는지 그 과정에 대한 정확한 이해가 선행되어야 한다.

(1) 트랜스포머(Transformer) 아키텍처를 통한 대규모 학습

현행 LLM 모델은 트랜스포머(Transformer) 구조를 기반으로 한다. 트랜스포머는 2017년 구글이 발표한 신경망으로 텍스트 내 단어들 간의 관계를 학습하게 되는데⁵⁾, 학습 과정은 크게 사전학습(Pre-training)과 미세조정(fine-tuning) 과정, 두 단계로 나눌 수 있다. 사전학습 단계에서는 방대한 텍스트를 입력받아 다음에 나올 단어에 대한 통계적 예측을 반복적으

4) 챗GPT는 플러스 요금제의 경우, 월 20달러의 구독료를 받고 있으며 비즈니스 요금제는 월 25달러, 프로 요금제는 월 200달러를 받고 있다. 매출이 20억 달러를 돌파하고 있다. 이러한 생성형 AI 산출물은 콘텐츠는 기존 작가, 일러스트레이터, 번역가 등의 시장을 직접적으로 잠식하고 있지만 학습 데이터를 제공한 창작자들은 일회성 보상(또는 무보상) 이후 AI의 지속적인 상업 활동으로부터 어떠한 추가적 보상도 받지 못하고 있다. (이승윤 기자, “오픈AI, 챗GPT 출시 2년 반 만에 연간 매출 100억 달러 돌파”. YTN, 2025년 6월 10일 보도)

5) 트랜스포머 구조는 기존의 AI 학습 방식(RNN, CNN)의 한계를 극복하기 위해 고안된 개념으로 2017년 구글 브레인 소속 연구진에 의해 발표되었고, GPT 시리즈의 기본구조로 채택되었으며, 현재 대부분의 언어모델이 이 구조에 기초하고 있다. (Ashish Vaswani et al, “Attention Is All You Need”, Arxiv, <https://doi.org/10.48550/arXiv.1706.03762>) 실제 GPT도 Generative Pre-trained Transformer의 줄임말로 이름 자체에 트랜스포머 구조를 명문화하였다.

로 수행하게 된다.⁶⁾ 문제는 이런 학습 대상 텍스트에 저작권법의 보호를 받는 수많은 저작물이 포함되어 있다는 것이다. 뉴스 기사와 웹 게시물 등을 비롯해, 논문과 소설 등 단행본까지 사용된 것으로 강하게 추정되며, 대다수는 명시적 허락 없이 진행되었다.⁷⁾ 이후 미세조정 단계에서는 이렇게 사전학습을 거친 모델을 특정 용도에 맞게 다듬게 되는데, 대화형 모델로 만들기 위해 질문과 답변 데이터를 추가시키고, 유해한 답변은 걸러내는 등의 조정을 거치게 된다.⁸⁾

(2) AI 학습의 본래적 특징 : 데이터 저장이 아닌 숨겨진 패턴의 일반화

AI 모델을 상대로 제기되는 저작권 소송이나 데이터 무단 사용 담론을 다루는 과정에서 흔히 'AI가 내 작품을 베꼈다'라고 하거나 'AI가 학습 데이터를 그대로 뺏어 낸다'는 식의 표현이 등장하는 경우가 간혹 있다. 하지만, 이같은 표현은 AI 학습의 특성을 정확히 표현한 말은 아니다. 오히려, 그릇된 표현에 더 가깝다.

왜냐하면 AI가 모델 학습을 할 때, 학습 데이터 자체를 저장하지 않기 때문이다. 언어모델의 경우를 살펴보면, 거대언어모델(LLM)은 방대한 양의 텍스트를 분석하여 단어와 단어의 관계를 매개변수(parameter) 연산을 통한 가중치 형태로 저장하게 되는데,⁹⁾ 이런 가

6) 이 과정은 수십억 개의 문장에 대해 반복적으로 이루어지게 되는데, 예를 들어, “오늘 기분이 매우”라는 문장이 주어졌을 때, 그 다음에 나올 가능성이 높은 단어인 “~좋다.”, “~나쁘다.” 등을 확률적으로 예측하게 된다.

7) 2020년에 출시된 GPT-3는 약 45TB의 텍스트 데이터(필터링 후 약 570GB)를 학습했으며, 여기에는 웹 크롤링 데이터, 디지털 도서 등이 다수 포함된 것으로 추정된다. GPT-4의 경우 개발사인 오픈AI가 데이터 및 파라미터 규모를 공개하지는 않았지만 업계에서는 GPT-3에 사용된 토큰의 20배 정도가 사용됐을 것으로 추정한다.(Josh Howarth, “Number of Parameters in GPT-4 (Latest Data)”, EXPLODING TOPICS, June 17, 2025 : <https://explodingtopics.com/blog/gpt-parameters>)

8) 다만 본고에서 다루고자 하는 사전학습(pre-training)의 일회성과 추론 및 활용의 지속성은 전세계에서 광범위하게 사용되고 있는 상업용 폐쇄형 파운데이션 모델(foundation model)을 기준으로 제시하는 것이다. 주로, 챗GPT, Claude, Gemini 등의 모델이 여기에 해당된다고 할 수 있다. 사용자 피드백을 통한 파인튜닝이나 추론 단계에서 외부 데이터 검색을 활용하는 이른바, 'RAG(Retrieval-Augmented Generation) 모델'은 학습과 활용 경계가 모호한 경우도 발생하고 있다. 전자의 경우인 반복적인 파인튜닝은 개별 단계를 별도의 학습 사이클로 볼 수 있다는 점, 후자의 경우인 RAG의 경우에는 본질적으로 학습이라기보다는 복제권이나 전송권 등 전통적 저작권 개념이 적용된다는 점에서 본 연구에서는 다루지 않았다.

9) GPT-4의 경우 매개변수 숫자는 1.8조개, Gemini의 경우 1조개 이상, Grok3 모델의 경우 비공개이지만 3천억 개 정도로 추정되며 이러한 매개변수는 신경망의 각 연결(connection)에 할당된 가중치 값이다. 다만, 기술 발전으로 인해 동일한 성능을 내는 데 필요한 매개변수의 수는 줄어들면서 효율화하는 추세이다. (변희원 기자, “고성능·저비용·개방형 모델 늘어… 일상 속으로 스며든 AI”, 조선일보, 2025년 4월 9일 보도.)

중치를 통해 흔히 프롬프트라고 일컫는 특정 입력값이 주어졌을 때, 답변을 예측하여 산출물(output)을 만들어 낸다.

이 과정에서 인공지능 모델은 학습 데이터를 그대로 검색하여 내뱉지 않는다. 학습 데이터는 이미 매개변수와 가중치 형태로 전처리되어 있으며, 따라서 생성물이 학습 데이터의 특정 저작물과 실질적, 결과적 유사성을 보이더라도 그것은 직접적 복제에 의한 것이라기보다는 유사한 패턴의 일반화를 통한 결과적 재생성으로 보는 것이 합당하다.¹⁰⁾

이는 인간이 실제로 학습하는 과정과 유사하다고 할 수 있다. 만약에 어떤 ‘문학소녀’가 소설가로 성장하는 과정에서 세계문학전집을 읽어 소설가로서 문장력을 키울 수 있었다고 가정해보자. 이 때, 그 ‘문학소녀’는 수백 권의 소설을 통째로 암기하는 것이 아니다. 많은 책들을 읽으면서, 해당 소설들이 가진 패턴이나, 개념, 특정 작가의 스타일이나 문체, 전형적인 표현 기법 등을 추상화하여 하나의 소설관을 정립하게 되는 것처럼, AI 모델도 개별 문장이 아닌 학습 대상인 문장들이 공통적으로 갖고 있는 통계적 규칙성을 일반화하게 되는 것이다.¹¹⁾ 만약 AI 모델이 원저작물을 ‘기억’하여 재현한다면 이는 복제권 침해가 명확할 수 있다. 그러나, 해당 모델이 ‘일반화(generalization)’를 수행하여 원저작물과 유사성이 없는 새로운 무언가를 생성해 낸다면 복제권 침해를 다투기는 어려운 상황이다.¹²⁾

2. 학습 데이터와 생성 결과물의 관계

그렇다면 AI 모델에 쓰인 학습 데이터와 산출물(output)의 관계는 어떻게 보아야 할까.

- 10) 일반적으로 LLM에서 생성된 의사 표현은 저작권을 침해하지 않는다. 이러한 모델은 학습 데이터 내에서 잠재 특징(feature)과 연관성을 학습(learning)하기 때문이다.(Matthew Sag, “Copyright safety for generative AI”. *Houston Law Review*, 61(2), pp. 95-347. p. 302)
- 11) AI 모델 학습이 완료된 후, 모델은 추론(inference) 과정을 통해 텍스트를 생성한다. 사용자가 프롬프트를 입력하면, 모델은 학습된 파라미터를 기반으로 다음에 올 단어를 확률적으로 예측하며, 이를 반복하여 문장을 완성하며 각 단계에서 가장 가능성이 높은 단어를 선택하게 된다.
- 12) 다만, 지난해 발표된 연구에 따르면 이런 일반화 과정에서도 데이터를 복제하는 듯한 결과가 나오는 것이 확인됐는데, 모델의 크기나 문맥의 길이, 학습 데이터의 중복 빈도가 높을수록 데이터가 재현될 수 있음이 확인되었다. 하지만, 이것이 복사 및 붙여넣기 형태의 복제와는 다르며, 전체 출력 중 일부 데이터에 해당되며 프롬프트로 입력된 데이터의 반복이나 서식을 복사하는 과정에서의 이른바 ‘유사 중복’효과가 있다고 분석된 바 있다. 본질적으로는 LLM 등 AI 모델이 만들어내는 생성물은 확률적 조합에 의한 것이라고 할 수 있다.(Kiyomaru, H. et al, “A comprehensive analysis of memorization in large language models.” *In Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing(EMNLP 2024)*, pp. 584-592.)

만약, 산출물이 학습 데이터를 복제한 것이라면 복제권 침해가 명확하겠지만 앞서 분석한 대로 AI의 산출물은 학습된 일반적 패턴에 따른 재창작이기 때문에 학습 데이터로 제공되었다는 사실 만으로 저작권법상 복제권 침해를 구성한다고 볼 수는 없다.¹³⁾

저작권법상 복제권 침해는 크게 두 가지 요건으로 구성된다. 우선, 원저작물에 대한 접근이 있어야 하고, 그 다음으로 그에 따른 원저작물과의 실질적 유사성이 있어야 한다.¹⁴⁾ 이 조건에 비추어 보면 AI 산출물은 학습 과정에서 특정 저작물을 입력받아 학습하는 것이 일반적이기 때문에 일단 원저작물에 대한 접근 자체는 명백히 성립한다고 할 수 있다. 그러나, AI 산출물은 프롬프트에 따라 사용자의 목적에 맞는 콘텐츠를 생성하기 때문에 사안에 따라 달리 판단할 수밖에 없는데, 생성형 AI는 특정 작품을 모방하는데 주로 쓰이기보다는 새로운 창작에 동원되는 경우가 상대적으로 더 많기 때문에 실질적 유사성이 인정될 수 있는 여지는 매우 드물다고 할 수 있다.¹⁵⁾

예를 들어, A라는 특정 소설을 재료로 학습한 모델이 만약 B라는 새로운 소설을 생성해 냈다면, A와 B는 서로 다른 소설이기 때문에 실질적 유사성이 발생할 개연성은 낮을 것이다. 또, 특정 작가의 문체를 흉내 내도록 프롬프트를 입력받아 새로운 텍스트를 생성해 냈다 하더라도 현행 저작권법상 문체나 스타일은 일반적인 아이디어 영역으로 간주되어 보호받지 못할 수 있다.¹⁶⁾ 결론적으로 현재 동의를 받지 않고 다량의 데이터를 활용해 학습한 AI 모델이 생성하는 산출물에 대해 데이터 소유자가 복제권 침해를 주장하기는 어렵다

13) 생성형 AI의 산출물은 데이터 특징을 알고리즘을 통해 분석하여 확률에 따라 재조합하기 때문에 이를 원 데이터의 '화학적 변화'라고 언급하는 학자도 있다.(이지연, 김원오, “저작물 학습자로서 생성형 AI의 저작권 침해 책임”, 인하대학교 법학연구소 IP & Data 法, 제4권 제1호, 2024년 6월 30일, 75~120쪽, 82쪽.)

14) “저작권법이 보호하는 복제권의 침해가 있다고 하기 위해서는 침해되었다고 주장되는 기존의 저작물과 대비되는 저작물 사이에 실질적 유사성이 있다는 점 외에도, 그 저작물이 기존 저작물에 의거하여 작성되었다는 점이 인정되어야 한다.” (대법원 2024.1.28. 선고 2013다21782 판결.)

15) 김병일, “인공지능 학습데이터의 이용과 저작권 제한”, 『AI법제 Insight』, 25-1호, 118-157면, 125면.

16) 미국 저작권법은 “아이디어와 표현의 이분법(idea-expression dichotomy)”의 원칙에 따라, 아이디어는 보호하지 않고 구체적 표현만 보호하고 있다. (17 U.S.C. § 102(b); Baker v. Selden, 101 U.S. 99 (1879); Ninth Circuit, 2015 F.3d Decision). 문체나 스타일은 일반적으로 “아이디어” 영역으로 간주되어 보호받지 못한다. 따라서 헤밍웨이 스타일의 소설을 쓴다고 해서 저작권 침해가 되지 않는다. AI가 특정 작가의 스타일을 학습하여 유사한 텍스트를 생성하는 것도 마찬가지로 비침해로 볼 여지가 있다. (“In no case does copyright protection for an original work of authorship extend to any idea, procedure, process, system, method of operation, concept, principle, or discovery, regardless of the form in which it is described, explained, illustrated, or embodied in such work.”, U.S. Copyright Act of 1976, Title 17 U.S.C. § 102(b))

는 것이다.¹⁷⁾

게다가, 구체적인 책임을 묻는 과정보다 복잡할 수밖에 없는데, AI 모델이 기본적으로 수천억 단위, 조 단위의 매개변수를 통해 가중치를 학습하기 때문이다. 따라서, 특정 산출물이 어느 한 저작물과 유사한 점을 발견했다손 치더라도, 그것이 특정 저작물의 영향을 받은 것인지 구분하기 어려울 수밖에 없다.¹⁸⁾

3. TDM 관련 주요국 법제 및 공통적 맹점

앞서 분석한 기술적 특성과 법리 적용의 어려움에 대응하여 전 세계 주요국들은 텍스트 데이터 수집(TDM)에 관한 법적 규율을 시도하고 있다. 본 절에서는 이같은 규율의 형태를 살펴보고, 현행 TDM 규제들이 갖는 맹점에 대해 짚어보고자 한다.

먼저, 유럽연합(EU)의 경우 2019년 디지털 단일시장(DSM) 저작권 지침을 통해 세계 최초로 TDM 명시적 예외 규정을 마련했다.¹⁹⁾ 해당 지침 제2조 제2항에서는 TDM을 “패턴 및 트렌드 상관관계를 포함하지만 이에 국한되지 않는 정보를 생성하기 위해 디지털 형태로 텍스트와 데이터를 분석하는 것을 목표로 한 자동화된 분석 기법”으로 정의하고 있으며, 제3조는 “연구기관과 문화유산 기관”이 행하는 복제 및 추출행위를 면제할 것을 규정하고 있다. 즉, 연구 목적 TDM은 강제적으로 허용하되 상업적 목적은 배제하는 원칙을 채택하고 있으며, 제4조에서는 모든 주체가 일반 TDM을 수행할 수 있으나 저작권리자가 이를 기계가 읽을 수 있는 방식으로 거부 의사를 표시하면 학습을 금지하도록 하는 이른바 ‘사전 거부(opt-out)’ 방식을 채택했다.²⁰⁾

17) 이 장에서 언급된 복제권 침해는 데이터 학습 단계에서 발생하는 복제권 침해를 의미한 것이 아니라 최종 산출물과의 관계만에 주목하여 판단한 것이다.

18) 예를 들어 모델이 판타지 소설을 생성했다고 가정했을 때, 이것이 해리포터나 반지의 제왕의 영향을 받은 것인지 아니면 해당 작품들을 포함한 수천 편의 판타지 소설들로부터 학습한 장르적 관습의 재현인지 가려내기가 어려울 수밖에 없다. 전통적인 저작권법은 하나의 원저작물에서 하나의 복제물이 나오는 일대일 구조를 상정하고 있으나 생성형 AI는 수백만, 수천만 수억의 원저작물에서 하나의 생성물이 나오는 다대일 관계를 만들어내기 때문이다.

19) 류시원, “저작권법상 텍스트·데이터 마이닝(TDM) 면책규정 도입 방향의 검토”, 『선진상사법률연구』, No.101, 347-390면, 357-360면.

20) P. Bernt Hugenholtz, “The New Copyright Directive: Text and Data Mining (Articles 3 and 4)”, Kluwer Copyright Blog, *Institute for Information Law (IViR)*, July 24, 2019 (<https://legalblogs.wolterskluwer.com/copyright-blog/the-new-copyright-directive-text-and-data-mining-articles-3->

일견 공익 보호라는 가치와 저작권자의 동의에 기반한 균형 잡힌 논리적 구성으로 보이지만 이 같은 접근 방식에는 핵심적인 문제를 내재하고 있다. 해당 지침은 대규모 데이터를 학습하는 불가피하게 생기는 복제 행위에 대해, 전통적인 저작권 침해의 관점에서 1회적으로 면제하는 것이지, AI 모델 학습이 완료된 이후 해당 모델을 상업적으로 활용하는 것까지 면제하는 것은 아니다.²¹⁾ 만약, 챗GPT가 해당 조항인 제4조에 따라 합법적으로 유럽에서 수백만 개의 특정 텍스트 저작물을 활용해 인공지능 거대언어모델 학습을 진행했다고 했을 때, 해당 모델로부터 발생하는 추가적인 수익이 학습 데이터를 제공한 유럽 저작자들에게 배분되지는 않는다.

영국도 사정은 비슷하다. 영국 정부는 지난 2014년 저작권법을 개정해 TDM 예외를 도입했으나 계약으로 이를 배제할 수 있도록 하고, Section 29A를 통해 연산 분석(computational analysis) 목적의 복제를 허용하고 있다.²²⁾ 비영리 연구 목적에 한해 허용한다고 규정하고 있지만 정작 분석이 완료된 이후 모델의 상업적 활용(commercial use)에 대해서는 아무런 규정이 없다. 당시 개정 과정에서 영국 정부는 AI 산업 경쟁력 강화를 명시적 목표로 제시하면서 창작자 보호보다는 AI 기업 지원에 방점을 두는 듯한 모습을 보인 바 있는데, 아직 명확한 결론이 나지는 않은 상태이다.²³⁾

나아가, 일본은 2018년 저작권법 제30조의4를 신설해 세계에서 가장 광범위한 TDM 예

and-4/)

21) 박희영, “인공지능 학습과 저작권 제한 -유럽 저작권지침 및 독일 저작권법의 텍스트 및 데이터마이닝(TDM)”, 『Copyright Issue Report』, 한국저작권위원회, 5면.

22) Copyright, Designs and Patents Act 1988, 29A (legislation.gov.uk) Copies for text and data analysis for non-commercial research

(1) The making of a copy of a work by a person who has lawful access to the work does not infringe copyright in the work provided that—

(a) the copy is made in order that a person who has lawful access to the work may carry out a computational analysis of anything recorded in the work for the sole purpose of research for a non-commercial purpose, and

(b) the copy is accompanied by a sufficient acknowledgement (unless this would be impossible for reasons of practicality or otherwise).

23) 영국 정부의 경우, 2021년 10월부터 2022년 1월까지 저작권 디자인 및 특허법(CDPA) 개정 가능성에 대한 협의를 마친 뒤, TDM 예외 범위를 확대하려는 계획을 일단 취소하기도 하였다. (Chelsea Chung, “UK Halts Its Expansion Of Existing TDM Exception For Copyright Infringements”. *The London Financial*, Feb 17, 2023)

의를 도입하고 있다.²⁴⁾ 저작물을 향유(enjoyment)하지 않는 목적의 이용을 포괄적으로 허용한다. 비향유 개념은 AI 학습이 저작물의 본래적 가치를 소비하는 것이 아니라 데이터로 분석하는 것이므로 저작권 침해가 아니라는 논리인데, 이는 앞서 분석한 기술적 메커니즘과는 정확히 부합한다고 할 수 있다.²⁵⁾

그러나, 이러한 관대한 규제에 의해 일본의 모든 저작물을 AI가 자유롭게 학습할 수 있으며 저작권자는 어떠한 대가도 받지 못하는 결과로 이어질 수 있다. 앞서 언급한 제30조의4는 복제 등의 행위를 허용하는 조항이지 학습 완료 후 모델의 지속적 상업적 활용에 대한 규율은 아니다. 만약, 챗GPT가 일본어 저작물을 대량 학습하여 일본 시장에서 수익을 올려도 일본 저작권자들은 어떠한 보상도 받지 못하는 것이다.²⁶⁾

미국은 TDM에 관한 명시적 입법 없이 공정이용(fair use) 법리를 통해 사안별로 판례를 통해 판단하고 있는데, Google Books Case 판결에서는 수백만 권을 스캔해 검색 가능하게 한 것을 변형적이고 시장 대체 효과가 없다는 이유로 공정이용으로 인정한 바 있다.²⁷⁾ 그렇다면 이런 판례가 AI 학습(training) 문제에도 적용될 수 있을지 따져볼 필요가 있다. 앞선 사건에서 구글은 스캔한 책으로 새로운 책을 생성하지 않고, 단지 검색과 발견을 가능하게 했을 뿐이었지만 생성형 AI는 학습한 데이터로 원저작물의 검색을 돕는 것이 아니라 전혀 새로운 제2, 제3의 콘텐츠를 생산하도록 함으로써 동일한 시장에서 경쟁하는 새로운 콘텐츠를 양산해 내고 있다. 즉, Google Books Case 판결의 법리가 AI 학습 문제에 준용된다면 저작자들의 피해는 오히려 더 커질 수 있는 상황인 것이다.²⁸⁾

24) 김성원, 최상민, 이수원, “텍스트·데이터 마이닝 과정의 저작물 이용 면책 규정 신설안에 대한 소고-적대적 생성신경망에의 적용-”, 『지식재산연구』, 2023 vol.18, no.1, pp.147-180, 153-156면.

25) Hugh Stephens, “Japan’s Text and Data Mining (TDM) Copyright Exception for AI Training: A Needed and Welcome Clarification from the Responsible Agency”, hughstephensblog, March 10, 2024.
(<https://hughstephensblog.net/2024/03/10/japans-text-and-data-mining-tdm-copyright-exception-for-ai-training-a-needed-and-welcome-clarification-from-the-responsible-agency/>)

26) 일본 정부는 AI 산업 진흥을 최우선 목표로 삼았다고 밝혔는데, 산업 진흥 관점에서는 긍정적 효과가 있을 수 있겠지만 그렇게 거두게 될 성과가 창작자들의 희생 위에서 달성됐다는 비판을 피하기 어려워 보인다.

27) 박경신, “[2015-22 미국] 항소법원, 구글의 도서 디지털 변환 사업인 구글 북스는 공정 이용에 해당 한다”, 한국저작권위원회 저작권 동향, 2025년 11월 6일.
<https://www.copyright.or.kr/information-materials/trend/the-copyright/view.do?brdctsn=17863&list.do%3FpageIndex=1&brdclasscode=&nationcode=&searchText=&servicecode=06&searchTarget=ALL&brdctstatecode=>

28) 상기 판례에서 재판부는 “Google Books의 검색 기능과 미리보기 기능이 변형성이 인정되며 상당한 대체

이 때문에 이 같은 인식을 바탕으로 추가 판례를 통해 판례법을 만들어갈 필요가 있지만 최근 벌어지고 있는 *Authors Guild v. OpenAI*, *New York Times v. OpenAI* 등 소송들을 살펴보면 쟁점은 결국, 학습 과정에서의 복제가 공정이용(fair use)인지에 집중돼 있다. 서론에서 소개한 Claude의 보상 사례 역시 학습 데이터의 대가 지불에 초점이 맞춰져 있다.²⁹⁾ 즉, 미국 내에서도 학습 과정에서의 TDM이 법적 쟁점이 되고 있을 뿐 학습 완료 후 인공지능 언어모델의 지속적 상업적 운영에 대해서는 제대로 된 조명이 이뤄지지 않고 있는 것으로 보인다.

이러한 경향은 최근 관련 소송의 절차적 변화 과정에서도 일부 엿볼 수 있다. 지난해 4월, 뉴욕 남부연방법원은 작가 길드(*Authors Guild*)와 뉴욕타임스 등이 Open AI 및 마이크로소프트를 상대로 제기한 개별 저작권 소송들을 하나의 재판부로 병합하여 심리하기로 결정한 바 있는데,³⁰⁾ 법원의 이 같은 결정은 일견 소송 경제를 명분으로 삼고 있지만, 결과적으로 AI 저작권 분쟁의 핵심을 ‘학습 단계의 공정 이용’이라는 단일한 프레임에 가두고 단순화하는 효과를 불러올 우려가 있다.³¹⁾

끝으로, 우리나라의 경우에는 TDM에 관한 명시적 조항이 없어 저작권법 제35조의5에 명시된 공정한 이용 조항에 의존해야 한다.³²⁾ 일부 학자들은 AI 학습이 기술 개발 목적이

적 경쟁 관계가 없는 상황에서 구글의 영리 목적은 상당한 변형적 목적을 압도하지 못한다.”고 판시한 바 있는데, 만약 구글이 스캔한 책들의 내용을 바탕으로 새로운 책들을 자동 생성하여 판매하고자 했다면 분명 결론이 달랐을 것이다. 작금의 생성형 AI는 그런 형태의 일을 하고 있는 상황이다.

29) 최승재, “인공지능 학습에 대한 미국 법원들의 공정이용 판단”, 『Copyright Trends Report』, 한국저작권위원회, 2025년 7월.

30) Ella Creamer, “US authors’ copyright lawsuits against OpenAI and Microsoft combined in New York with newspaper actions.”, *The Guardian*, 4 April 2025.
(<https://www.theguardian.com/books/2025/apr/04/us-authors-copyright-lawsuits-against-openai-and-microsoft-combined-in-new-york-with-newspaper-actions>)

31) 즉, 인공지능 모델이 실제로 시장에서 운영되면서 발생하는 여러 경제적, 법적 효과들이 ‘데이터 수집 절차’에만 매몰되는 우를 범할 수 있음을 시사한다.

32) 제35조의5(저작물의 공정한 이용) ① 제23조부터 제35조의4까지, 제101조의3부터 제101조의5까지의 경우 외에 저작물의 일반적인 이용 방법과 충돌하지 아니하고 저작자의 정당한 이익을 부당하게 해치지 아니하는 경우에는 저작물을 이용할 수 있다. <개정 2016. 3. 22., 2019. 11. 26., 2023. 8. 8.>

② 저작물 이용 행위가 제1항에 해당하는지를 판단할 때에는 다음 각 호의 사항 등을 고려하여야 한다. <개정 2016. 3. 22.>

1. 이용의 목적 및 성격
2. 저작물의 종류 및 용도
3. 이용된 부분이 저작물 전체에서 차지하는 비중과 그 중요성

고 비항유적이며 변형적이므로 정당화될 수 있다고 주장하고 있으나, 우리나라 내의 논의 역시 학습 행위의 적법성에만 집중하고 있다.³³⁾ 우리나라 역시 인공지능과 저작권 문제에 대한 공론화 단계에 접어들었으나 학습 단계 그 자체에 주로 주목할 뿐 학습 완료 후 모델의 지속적 상업 활용 단계에 대한 대응에 대해서는 적극적인 논의가 나오지는 않고 있다.³⁴⁾

<표 1> 주요국 TDM 법제 비교

구분	주요 특징	활용 단계 규율
EU	DSM 지침으로 명시적 TDM 예외 도입, 권리자 opt-out 방식	없음.
영국	저작권법 Section 29A로 비영리 연구 목적에 한정해 허용	
일본	저작권법 제30조의4로 비항유 목적 이용을 광범위하게 허용	
미국	명시적 입법 없이 공정이용(fair use) 법리로 사안별 판단	
한국	명시적 TDM 조항 부재, 저작권법 제35조의5(공정한 이용) 의존	

이상의 비교법적 분석을 통해 내릴 수 있는 결론은 주요국들의 TDM에 대한 접근법은 모두 동일한 맹점을 공유한다는 것이다. 학습 시점의 일회성 사용의 문제에만 집중할 뿐, 학습 완료 후 AI가 그 능력을 지속적으로 활용하여 벌어들이는 수익에 대한 이익 배분의 문제에 대해서는 고려하지 않고 있다는 것이다. 이에 본고에서는 이런 현황과 문제 의식을 바탕으로 인공지능의 개발과 활용 단계를 구분하여 각 단계별로 상응하는 규제 정책을 제안해 보고자 한다.

III. 보상과 과세 분리 방안 제안 : 이중 구조 보상 모델

1. 일회성 보상 패러다임이 갖는 문제점

앞서 서술한 바와 같이 주요 국가들의 현행 TDM 규제 관련 법제는 AI 모델 학습 단계의 데이터 무단 사용에만 주로 초점을 맞추고 있다. 이렇다 보니 자연스럽게 모델 학습 시

4. 저작물의 이용이 그 저작물의 현재 시장 또는 가치나 잠재적인 시장 또는 가치에 미치는 영향

33) 장민권 기자, “AI 학습시 ‘저작권 침해 책임’ 면제해야”, 파이낸셜뉴스, 2025년 9월 15일 보도.

<https://www.fnnews.com/news/202509151816376363>

34) 국내에서 제시된 입법 논의를 봐도 정보 분석의 대상이 되는 저작물에 적법하게 접근할 것을 요구하면서도 이에 대한 구체적 의미가 없다는 학계의 지적도 있다. (강미영, “생성형AI의 활용에 따른 저작권 침해와 형사 정책적 개선 방안”, 부산대학교 법학연구 제66권, 제1호통권123호, 2025년 2월, 193~218쪽, 205-206쪽.)

점에 저작권 침해가 발생했는지, 혹시 발생했다면 이것을 예외로 인정할지, 만약 필요하다면 그에 상응하는 어떤 보상을 지급할 것인지의 문제로 귀결된다. 바꿔 말하면, 이 모든 논의가 일회성 보상에 국한하고 있는 것이다.

앞서 비유한 대로, 학생이 서점에서 학습 목적으로 참고서를 사면 저자와의 관계는 책 값을 지불하는 단계에서 그 순간 종료된다. 그 이후에 학생이 그렇게 구입한 서적으로 스스로 공부해서 작가가 되어 돈을 벌든, 시험에 합격하든 저자와의 후속 관계가 발생하지는 않는다. 지식의 전파와 활용은 자유로워야 하고, 학습의 기본권이 보장되는 상황에서 이는 당연한 결론이다.³⁵⁾

그러나, AI 모델 학습의 경우에는 이 같은 논리 구조를 그대로 적용할 수 없다는 게 필자의 견해이다. 그 논리적 근거로 우선 시간적 차이를 들 수 있다. 인간이 교과서든 참고서든, 소설이든 특정 텍스트로 학습하여 자신의 학술적 혹은 문학적 세계를 구축해 나가는 시간은 작게는 수년에서 수십 년이 소요된다.³⁶⁾ 그렇다 보니, 그걸 토대로 한 사람이 평생 생산해 낼 수 있는 창작물의 수도 제한적일 수밖에 없다.

반면, 인공지능의 경우, 학습에 필요한 시간은 GPU 자원이 넉넉할 경우, 수주에서 수개월의 시간이면 충분하고 대규모 데이터로 학습을 마친 즉시 해당 모델을 활용해 산출물을 대규모로 생산할 수 있다. 학습하는 데 걸리는 시간이 압축적으로 짧아질 수 있기 때문에 앞선 인간의 사례와 비교하면 시간적 격차가 커지게 된다. 학습을 마친 이후, 산출물을 생산하는 속도도 압도적으로 빨라 초당 수천 개의 텍스트를 생성할 수 있으며, 이 같은 콘텐츠는 디지털 형태로 생산되어 빠르게 복제될 수도 있다. 예를 들어, 한 소설가가 평생 서른 권 정도의 책을 썼다고 쳤을 때, 그 소설가가 읽은 참고서나 비평서, 문학 작품은 수백 권 혹은 수천 권에 이를 수 있다. 그러나 그 소설가가 쓴 서른 권의 책이 기존 문학 시장을 직접 대체한다고 보지는 않는다. 평생에 걸쳐 점진적으로 시장에 진입하며 오히려, 문학 생태계의 저변을 넓히는 기여로 인식되기도 한다. 그러나, 챗GPT로 대표되는 거대언어 모델의 경우, 수백만 권 혹은 수억 권의 책을 학습한 이후에 하루 만에도 수십만 개, 수천

35) 예를 들어, 노벨문학상을 받은 한강 작가가 세계문학전집 등을 읽고 소설가의 꿈을 키우며, 자신의 문학적 역량을 키워 소설가가 되었다고 해서 그의 작품이 세계문학을 표절했다거나 베꼈다는 의심을 받지 않는다. 실제로 한강 작가는 어릴 때 문인 집안에서 성장하면서 많은 소설을 읽은 것으로 전해진다. (김호영 기자, “한강, 캄캄한 방에서 공상 즐기던 소녀”, 채널A, 2024년 10월 11일 보도.)

36) 심지어, 평생에 걸쳐 과업을 완성하는 작가도 있다.

만 개, 수억 개의 텍스트를 단시간에 쏟아내면서 기존 창작물 시장을 단시간에 교란시킬 수 있는 상황이다.³⁷⁾

둘째로 생각해볼 수 있는 논리적 근거는, 비용과 수익의 불균형 문제와 관련이 있다. 작가가 문학적 소양을 쌓기 위해 구입하는 책의 가격은 수백만 원에서 수천만 원이다. 그리고, 그걸로 작품 활동을 하면 보통 책 판매 가격의 10% 내외를 인세 수입으로 거둬들인다. 책 판매 기간이 평생에 걸쳐 이뤄지면 이 같은 수입도 생애 주기에 맞춰 순차적으로 발생한다.³⁸⁾ 그러나, 인공지능 거대언어모델의 경우, 데이터 수집과 처리, 학습에 일정 비용을 지불하게 되고 이로 인해 파운데이션 모델을 완성하면 그 모델은 다중 사용자에게 개방된다. 오픈AI의 지난해 연간 매출이 20억 달러 이상으로 추정되고 있는데, 데이터 수집 비용 대비 큰 수익이라고 할 수 있다.³⁹⁾ 물론, AI 테크 기업들의 경우, 현재까지 수익을 내지 못한 상태이며, 여전히 적자 국면이라고 항변하지만 적자에 들어간 비용은 대부분 GPU 등의 하드웨어 공급비용이며, 데이터 수집 및 학습에 대한 대가 지불은 미미하거나 거의 없는 상황이다. 그마저도, 초기의 인공지능 모델들은 사실상 무단 크롤링 과정을 통해 대가 없이 학습한 것으로 알려져 있다. 설령, 추후에 학습 비용을 지불한다 하더라도 학습 데이터 제공자들은 일회성 보상 이후 배제되는데 이 경우 비용과 수익의 불균형 문제가 발생할 수 있는 것이다.

셋째는, 인공지능은 단순히 텍스트 생성을 넘어 특정 창작자의 스타일 학습까지 가능하다는 점이다. 얼마 전, 지브리 스튜디오풍 이미지 생성 유행이 생겨나면서 인공지능의 스타일 모방 능력이 널리 회자된 바 있지만 실제로, 단순히 인공지능은 산출물만 만드는 게 아니라 특정 작가를 모방하는 것도 가능하기 때문에 또 다른 2차 피해를 유발할 수 있고 특정 작가의 수요를 다른 형태로 잠식하는 것도 가능하다.⁴⁰⁾ 단순히 ‘학습했다.’고만 표현

37) 단순히 텍스트 뿐 아니라 이미지나 영상 생성 모델은 창작자를 넘어 배우들의 경쟁 환경까지 교란 시키고 있는 실정이다. (방성훈 기자, “AI 여배우 출현에 헐리우드 ‘발칵’ ... 유명 배우들 “인간 창작 모욕”, 이데일리 2025년 10월 1일 보도)

<https://www.edaily.co.kr/News/Read?newsId=03847446642328656&mediaCodeNo=257&OutLnkChk=Y>

38) 작가노조준비위원회 저자들이 박한 인세 수입의 문제를 지적하는 내용의 책을 출간하면서 작가들의 수입에 관한 실상을 공개한 적이 있다. (송광호 기자, “강연료가 원고료 초과하는 씽씽한 작가 현실... ‘작가노동선언’, 연합뉴스, 2025년 4월 29일 보도)

<https://www.yna.co.kr/view/AKR20250429071100005?input=1195m>

39) 고계연 기자, “챗GPT 앱, 누적수익 2.7조원... “AI 앱 시장 지배””, 초이스경제, 2025년 8월 16일 보도.

<https://www.choicenews.co.kr/news/articleView.html?idxno=152329>

할 것이 아니라, 수십만, 수백만 명의 창작자가 평생에 걸쳐 축적한 창작 능력을 압축적으로 내재화하면서 차원이 다른 결과를 낳게 되는 것이다.

넷째로, 위와 같은 특징으로 인해 비가역적 피해가 발생한다. 실제로, 생성형 AI 도입 이후, 많은 영역에서 창작자 생태계 붕괴가 실제로 목격되고 있으며, 많은 VFX 스튜디오들이 디자이너들을 해고하고 있으며, IT 회사들의 경우에도 개발자들의 대량 해고로 이어지고 있다.⁴¹⁾ 더 큰 문제는 창작 의욕 감소이다. AI가 일회성 보상을 통해 학습을 마친 뒤 대규모 창작 활동을 수행하게 되면 창작자들의 창작 의욕이 줄어 전반적으로 돌이킬 수 없는 피해를 낼 수 있는 것이다.⁴²⁾

위와 같은 이유로, 현재의 일회성 보상 체계에는 문제가 있다는 것이 필자의 주장이다. AI 기술의 특성상 활용 단계에서는 복제권 이슈가 발생하는 것은 아니지만 그 이후에 생기는 실질적 타격을 고려해야 하기 때문이다. 비유하자면 ‘문학소녀’에게는 ‘책값’만 받아도 되지만 ‘초능력 문학소녀’에게는 책값뿐 아니라 일종의 ‘초능력세’를 거두어야 한다는 것이다. 개울물에 규제가 필요 없다고 해서 쓰나미에도 규제가 필요 없는 것은 아닐 것이다. 규모와 속도의 질적 전환은 그 사회적 효과도 달라지게 할 수 있는 만큼 TDM 문제 해결을 위해 학습 단계 보상만으로는 부족하며, 해당 모델의 지속적 활용에 대한 별도의 규율, 즉 활용 단계의 과세적 접근이 필수적이다.

2. 이중 구조 보상 모델의 이론적 구성

앞 절에서는 TDM 법제 논의 과정에서 일회성 보상에 그치는 것이 왜 근본적 한계를 가

40) 권은주 기자, “챗GPT發 ‘지브리 프사’ 열풍... “애니 작가 밥그릇까지 빼앗나””, 이코노미뉴스, 2025년 4월 3일 보도. (<https://www.m-economynews.com/news/article.html?no=52869>)

41) 개발자들이 쓰는 프로그래밍 언어 또한 언어의 일종으로 거대언어모델이 매우 잘 이해하고, 생성하는 분야이기도 하다. 이런 특성에 기반해 언어모델이 개발자나 디자이너 자리를 대체하면서 글로벌 기업들을 중심으로 작가는 수만 명에서 수십만 명에 달하는 해고 움직임이 일고 있다. (장유미 기자, “[AI는 지금] “무려 60만 명 해고”... AI·로봇 직원 등장에 인간 일자리 사라진다”, 지디넷코리아, 2025년 10월 24일 보도)

42) Basbanes v. Microsoft Corporation and OpenAI (filed January 5, 2024) 사건에서 사건 원고들은 Open AI와 MS가 원고들의 저작물을 의도적으로 이용했고, 이에 대한 비용 지불도 제대로 하지 않았으며, 2차 저작물로 높은 이득을 취하며, 창작 의욕을 상실케 하고 시장을 파괴한다고 주장한바 있다. (이지연, 김원오, 앞의 논문, 91쪽.)

지는지 살펴보았다. 이같은 한계 극복을 위한 대안으로서 본고에서는 활용 단계에서 일종의 과세를 부과하는 ‘이중 구조 보상 모델’을(Two-Tier Compensation Model)을 제안하고자 한다.

‘이중 구조 보상 모델’의 구조는 단순하다. 인공지능 모델의 보상 단계를 특정 저작물을 수집하고, 전처리하여 학습 알고리즘을 적용해 매개변수(parameter)를 최적화하는 1) 학습 단계(Input stage)와, 이렇게 개발된 AI 파운데이션 모델을 활용하여 텍스트나 이미지 등을 생성하여 그렇게 만들어진 산출물이 시장에 투입되는 과정인 2) 활용 단계(Output stage)로 나누어 각 단계별 규율 방식을 적용하자는 것이다.

일단 두 단계로 나누어 접근하고자 하는 건 저작권이나 시장에 미치는 영향력의 범위나 규율의 초점이 다를 수 있기 때문이다. 우선, 학습 단계는 그 자체로 과거의 일회성 행위인 반면에, 활용 단계는 현재와 미래에 걸쳐 계속 지속될 수 있다. 또, 학습 단계는 저작물을 직접 활용하게 되어 그 자체로 복제권 문제가 직접적으로 제기될 수 있으나, 활용 단계는 복제권은 관여하지 않고, 시장의 대체 효과라는 간접적 영향이 주된 쟁점이 된다. 그렇다 보니 학습 단계는 단순히 학습에 쓰인 저작물과 인공지능 모델 사이에 각각 일대일(一對一) 대응이 성립하게 되지만, 활용 단계에서는 인공지능 모델이 산업 전체, 나아가 창작 생태계 전반을 교란할 수 있는 일대다(一對多)의 문제로 치환된다. 끝으로, 학습 데이터는 학습에 쓰였는지 여부를 AI 모델 개발사가 공개적으로 공표한다는 것을 전제로 상대적으로 간명하게 확인할 수 있지만,⁴³⁾ 후자의 시장 영향은 그 직간접적 영향과 정도를 가늠하기 어렵다.

필자가 핵심적으로 주장하고자 하는 요지는, 이처럼 단일한 법적 규율만으로는 앞서 상술한 두 단계인 학습 단계와 활용 단계를 모두 포섭하는 것이 구조적으로 어렵다는 것이

43) 학습 데이터를 모두 공개한 대표적 모델로 스위스 정부가 개발한 Apertus를 들 수 있는데, 스위스 정부는 지난 9월, 최초로 자체 국영 LLM ‘Apertus’를 공개하면서 이름처럼(라틴어로 Apertus는 ‘열린’이라는 뜻) 모델 구조, 가중치, 학습 데이터까지 모두 공개했다. 투명성과 신뢰성을 전면에 내세운 프로젝트라는 점이 핵심이다. OpenAI나 Google과 같은 소위 빅테크 기업의 최신 LLM들은 여전히 학습 데이터나 설계 과정을 외부에 공개하지 않고 있는데, EU AI Act가 본격 시행되면, 기업과 연구자는 모델 학습 데이터와 과정에 대한 투명한 설명과 문서화를 요구 받게 된다. 이때, 빅테크의 불투명한 모델에 의존하는 것은 법적·윤리적 위험으로 이어질 수 있어 스위스는 이런 리스크를 선제적으로 피해가기 위해 처음부터 규제 친화적인 모델을 구축한 것으로 보인다.(Miriam Partington, “Switzerland unveils national LLM in tech sovereign push”, SIFTED, September 2, 2025)

다. 이에 따라, 본고에서는 이 두 단계를 명확히 구분하여 AI 모델 학습 단계에서는 전통적인 저작권 법리에 따라 복제권 침해 논의와 보상의 구조로 다루고, 해당 모델을 상업적으로 활용하는 단계에서는 저작권법의 렌즈를 넘어, 시장 질서 정립이라든지 창작 생태계 유지 등 창작 생태계 시장 보호를 보호 법익으로 삼는 조세적 차원의 다층적 접근을 제안하고자 하는 것이다.

이 같은 구조를 설계하기 위해 필자는 해당 조치에 대응하는 법적 근거도 차등적으로 접근할 것을 제안한다. 즉, 1단계에서는 저작권이라는 사적 권리 영역을 보호하는 법리를 검토해야 하지만, 2단계인 활용 단계에서는 경제법과 조세법 같은 공적 규제 관점에서의 접근이 필요하다는 것이다.

(1) 1단계: 학습 데이터 보상 체계 (Compensation for Training Data)

1단계의 학습 데이터 보상은 인공지능 모델 학습 과정에서 발생하는 저작물 이용에 대한 것으로, 전통적 저작권 법리 적용의 연장선상에 있다. TDM 문제를 두고 다양한 학설과 견해, 대안들이 제기되고 있지만 큰 틀에서 법적 정당화 근거는 크게 두 가지로 접근해 살펴볼 수 있는데, 일설은 전통적 복제권 침해의 법리 관점에서 살펴볼 수 있다는 견해, 그리고 다른 일설은 새로운 형태의 예외적 사례로 볼 수 있다는 견해로 나뉘볼 수 있다.

먼저, 전자의 경우는 이른바 '복제권 침해 보상설'에 해당하는 데, 일단 특정 저작물을 인공지능이 학습하기 위해서는 이를 스토리지 서버에 저장하는 행위가 수반된다. 물론 AI 모델에 데이터가 그 자체로 남지는 않지만 학습(model training) 목적으로 옮겨 담는 과정에서 저작권법상 복제에 해당하기 때문에 이에 대한 학습 데이터 보상이 필요하다는 논지를 펼 수 있다.⁴⁴⁾

후자의 경우에는 TDM 행위를 새로운 예외로 간주하자는 것으로 본고에서 주장하고자 하는 개념이다. AI 학습 자체는 공익적 목적, 과학 기술 발전의 목적으로 보고 우선 합법으로 인정하고 이에 대한 일종의 '특별한 희생'을 고려해 보상 청구권을 인정하자는 것이다.⁴⁵⁾ 이는 국내외 학계에서도 제안된 바 있는 견해인데 토지 수용 등에 대한 손실보상과

44) 이지연, 김원오, 앞의 논문, 108면.

45) 공정이용(fair use)의 법리에 따라 TDM을 허용하더라도 보상 체계가 필요하다는 견해도 제시된 바 있다. (이상미, "AI 학습데이터 무단 사용에 대한 저작권 보호 방안", 『계간 저작권』, 37(2), 79-121면, 95-96면.)

유사한 법리를 적용하자는 것이다.⁴⁶⁾ 즉, TDM 행위가 합법이라 하더라도 상응하는 정당한 보상(Just compensation)을 하자는 것이다.⁴⁷⁾

AI 모델이 특정 데이터를 학습한다고 해서 해당 데이터를 그대로 저장했다가 출력하는 형태가 아니라는 점, 딥러닝(deep learning) 방식의 새로운 과학적 기법에 의한 개발이라는 점, 인공지능이 사회 전반의 효율성 향상에 기여할 수 있다는 점을 고려할 때 단순한 복제권 침해로 보긴 어려운 측면이 크기 때문이다.⁴⁸⁾

이처럼 1단계 보상의 이론적 근거를 마련한 뒤에는 데이터 유형별 차등적 설계가 필요하다. 저작권 보호의 강도가 다른 만큼 그에 따른 보호 가치도 달라지고 이에 따라 특별한 희생의 정도도 달라질 수 있기 때문이다.⁴⁹⁾

<표 2> 데이터 유형별 보상 체계

구분	주요 특징
Tier 1 (무보상 데이터)	- 퍼블릭 도메인 (저작권 보호기간 만료) - 오픈 라이선스 저작물 (CC0, CC BY 등) - 정부저작물 등 공공저작물 - 창작자가 AI 학습을 명시 허용한 자발적 기부 데이터
Tier 2 (최소 보상 데이터)	- 웹상 일반 보호저작물(블로그, 뉴스 기사, 온라인 소설 등) - 학술DB, 뉴스 아카이브 등 데이터베이스 저작물

46) 국내에서는 미국에서 제안되고 있는 의 이른바 ‘공용 수용’ 논의를 비판적으로 수용하면서 행정법적 접근이 아니라 한국 민법의 ‘주위지(周圍地) 통행권 이론’을 원용하는 학설이 제시된 바 있는데, 결국 국내외 학계에서 주목하는 핵심은 기존 저작권자의 권리를 사후적으로 보상해야 한다는 당위성 측면에서는 공통적인 견해를 취하고 있는 것으로 보인다. (남형두, 『공정이용의 역설 -시소에 올라탄 거인, 균형의 복원』, 경인문화사, 2025년 2월, 168-192면.)

47) 미국의 경우 행정법적 근거가 아닌, 수정헌법 제5조 (정부의 권한 남용에 대한 보호)를 근거로 이른바, 정당한 보상(Just compensation) 개념을 유추하고자 하는 시도가 제안되고 있다. (Molly McInnis, “Advancing Computation Of Just Compensation: An AI-Based Approach To Property Valuation In Eminent Domain”, *University of Cincinnati Law Review* Vol.94, Oct 12, 2025)

48) 이런 인공지능의 특징 때문에 심지어 산출물이 학습데이터와 유사하게 출력된다고 하더라도 동 저작물에 대해 저작권법 의거성을 인정할 수 없다는 견해도 있다. (손영화, “생성형 AI에 의한 창작물과 저작권”, 『법과 정책연구』, 한국법정책학회, 제23권 제3호, 2023. 9. 357-389면. 372면.)

49) 해외 연구에서는 학습 대상 데이터를 구분하면서 1) 강력한 보호(Strong Protection) 2) 적격 보호 (Qualified Protection) 3) 보호 면제 (Exclusion from protection)로 나눈 접근 방식이 제시된 바 있으나, 본고에서는 보상 관점에서 재해석하여 술어를 보상이 필요없는 단계부터 보상이 많이 필요한 단계로 역으로 정리하였다. (Zichun Xu, Zhilang Xu. “Tiered copyrightability for generative artificial intelligence: An empirical analysis of China and the United States judicial practices”, WILEY, open access, 22 August 2025 <https://doi.org/10.1002/aaai.70018>)

구분	주요 특징
Tier 3 (개별 협상 데이터)	- 베스트셀러, 전문학술서, 유료DB 등 - 의료·법률 등 특수목적 데이터 - 집중관리 거부 및 개별협상 선택 데이터

우선 가장 낮은 보호 단계의 데이터로 무보상 데이터(No Compensation Data)를 상정할 수 있다. Tier 1에 속하는 무보상 데이터는 저작권 자체를 주장하기 어려운 데이터를 말하는 데, 예를 들어 저작권 보호 기간이 이미 만료된 저작물이나, 오픈 라이선스 저작물 또는 정부나 공공기관이 공적 목적으로 생산한 저작물, 창작자가 자발적으로 AI 학습을 허용한 저작물 등이 이에 속할 수 있다. 이 경우, 1단계 학습 과정에서 TDM으로 활용되었다고 하더라도 보상 체계가 작동할 근거는 낮다고 할 수 있다.

다음으로 Tier 2에 해당하는 최소 보상 데이터(Minimal Compensation Data)이다. 이 데이터는 사실상 가장 논란이 될 수 있는 영역으로 웹사이트상에는 공개돼 있으나 명시적으로 라이선스 표시나 권리 침해 요소가 적은 데이터가 여기에 속한다. 블로그에 공개한 글이나 뉴스 기사, 온라인에 게시된 공개 댓글 등이 여기에 포함될 수 있는데 명시적인 라이선스 표시가 없더라도 저작권자의 동의 없이 사용했다면 당연히 특별한 희생에 의한 보상 논의가 이 단계부터 발생한다고 할 수 있을 것이다.

다만, Tier 2의 경우에는 웹상에서 크롤링하여 학습하게 될 경우에 수십만 건, 수백만 건 나아가 수천만 건에 대해 일대일 대응 방식으로 보상하는 것은 한계가 있기 때문에 주로 해당 사이트나 플랫폼 단위로 보상 체계를 협의하고 저작권자별 기여도 알고리즘에 따라 추정 분배하는 것이 합리적인 접근법일 것으로 사료 된다.⁵⁰⁾

끝으로, Tier 3으로는 개별 협상 데이터(Individually Negotiated Data)로 가장 저작권 보호 수준이 높은 데이터가 여기에 속한다. 베스트셀러 소설이나 전문 학술 서적, 유료 방송 데이터 등 고부가가치 콘텐츠가 여기에 해당할 수 있다. 단순히, 금전적 가치가 높은 데이터

50) 기존의 저작권 보호 접근법과 다르게 AI 학습 과정에서 데이터 제공자의 기여도를 수학적으로 분석하여 경제적 보상을 제공하는 시스템을 설계하자는 연구(Jiachen T. Wang et al, "An Economic Solution to Copyright Challenges of Generative AI", arxiv, 2024.04, <https://arxiv.org/abs/2404.13964>)도 제안되고 있는데, 이 경우 인공지능 모델의 성능을 저해하지 않으면서도 공정한 보상을 모색할 수 있다. (한국저작권위원회, "[이슈 브리프] 생성형 AI 생성물의 저작권 문제 해결을 위한 기술적 접근: 3가지 최신 연구 분석"(2025년 2-2호), 2025년 2월19일).

뿐만 아니라 의료 기록이나 법률 문서처럼 개인 정보가 포함된 민감 데이터 등도 이 범주에 속한다고 할 수 있을 것이다.

또, Tier 2에 속한다고 하더라도 저작권자가 강하게 학습에 동의하지 않은 경우라면 개별 협상 데이터로 분류할 수 있을 것이다. 이 항목에 속하는 데이터의 경우, 금전적 가치로 환산하여 추산하는 것이 비교적 쉽기 때문에 시장 가격에 따라 특별한 희생에 대한 보상 범위를 달리 협의해야 할 것으로 생각된다.⁵¹⁾

(2) 2단계: 산출물 기반 과세적 보상 (Output-based Taxational Compensation)

앞서 언급한 1단계인 AI 모델 학습 단계의 경우에는 일회적인 모델 학습 행위에 대해 일종의 ‘특별한 희생’ 관점에서 데이터 유형별 차등적인 보상이 가능할 수 있는 법적 근거에 대해 살펴보았다. 여기에 더해, 본 고의 핵심 제안이 바로 활용 단계의 산출물 기반 과세적 보상 제도인 이른바 2단계의 보상적 기여 방안이다.

다만, 개념상 구분해야 할 점은 1단계에 제안하는 보상 체계와 달리 2단계 과세적 보상 단계는 저작권 침해에 대한 반대급부 성격의 보상 개념이 아니라는 점이다. 본고에서 제안하고자 하는 2단계 산출물 기반 과세적 보상제도는 생성형 AI모델을 직접 활용하는 단계에서 납부하는 일종의 ‘AI 창작 생태계 기여금’에 해당할 수 있다.

이는 필자가 입법론적 관점에서 AI 관련 국제법무 체계 중 하나로 새로 창설할 것을 제안하는 개념으로, 조세법의 렌즈에 비추본다면 특정 산업에 대한 일종의 목적세 성격을 지닌다고도 할 수 있을 것이다. 핵심은 일단 AI 모델 개발을 위해 저작권 사용에 대한 보상을 지불하고 난 뒤, 그 데이터를 정당하게 학습하여 개발된 AI 모델에 대해서도 실제 활용할 때 일회적 보상을 넘어 활용 행위로 발생하는 부가가치에 대해 산출물 기반의 조세 성격의 추가 기여를 AI 모델 보유기업 등에 부과해야 한다는 것이다.⁵²⁾

51) 시장 가치가 높은 데이터는 개별 협상이 훨씬 더 효율적일 수 있으며, 권리자의 선택권을 존중해야 한다는 점에서 합당하다고 생각된다.

52) 경제법적 관점에서는 부정적 외부효과(negative externality)의 내재화를 위한 일종의 피구세(Pigovian Tax) 관점이라고 할 수 있다. 생산 효율성을 높이는 AI는 긍정적 외부효과로서 집단적 문화 자산을 무상 흡수하는 동시에, 시장 진입을 통해 창작자 수익을 잠식하는 부정적 외부효과도 가지고 있다. 결국, 효익과 실익을 상계하여 적정 수준의 과세가 필요하다는 개념으로 결국 핵심은 AI 기업이 개발 및 활용 행위로 인해 발생하는 사회적 비용을 부담해야 한다는 것이다. 결과적으로 조세 개념을 좀 더 완화한다면 문화 정책적 관점의 기여금 성격으로 볼 여지도 있을 것이다.

2단계의 활용 단계 과세를 저작권법 밖에 두어야 하는 이유는 다음과 같다. 우선, 일단 AI 모델이 학습을 거치고 나면 내부에 데이터 자체가 남지 않기 때문이다. AI 모델 내부에는 매개변수와 가중치 형태로 저장된 거대한 알고리즘 체계가 남아있게 된다.⁵³⁾ 이 경우, AI 모델은 매 프롬프트마다 개별 추론을 거쳐 산출물(Output)을 생성하게 된다. 이 경우, 인과관계를 추론하기 힘들다. 특정 산출물이 어느 저작물로부터 나왔는지 추정하기 어렵는데 저작권법을 매개로 대가를 요구하는 것은 논리적으로 성립하기 어렵다고 할 수 있다.

두 번째로, 한번 학습을 마치면 대규모 생산이 단시간에 가능한 AI 모델의 특성은 특정 개인에게 피해를 준다고보다 생태계 전체에 대해 피해를 준다는 점에서 개별 저작권 대응이 어려울 수 있다. 2단계에서 발생하는 침해의 성격 자체가 공동체 집단에 대한 대항성을 지닌다고 할 수 있을 것이기 때문에 국제법에서 등장하는 개념인 대세적 의무(duties erga omnes)의 법리를 적용하여 창작 공동체 전체에 대한 책임 개념으로 접근할 필요가 있다.⁵⁴⁾

끝으로, 기존 국제 저작권 법무체계와의 정합성 문제를 생각하더라도 2단계의 활용단계 과세는 저작권법과 별개로 고려되어야 한다. 베른 협약과 TRIPS 협정에서는 회원국이 저작권의 최소 보호수준(minimum standards)을 준수할 것을 규정하고 있는데,⁵⁵⁾ 만일 2단계

53) Adam Buick, "Copyright and AI Training Data—Transparency to the rescue?", *Journal of Intellectual Property Law & Practice* 20(3): pp.182-204, 12 Dec 2024, p.186

54) UC 버클리대학의 교수인 Haimo Schack은 지식재산권 자체가 본래부터 대세적 효력(operating erga omnes)을 지닌 절대적 권리(absolute right)라고 주장한다. 즉, 지재권을 침해한 자가 1인에게 피해를 가한 것이 아니라 해당 권리를 향하는 공동체 전체에 가해행위를 한 것으로 인식한다. 이런 이론적 논의를 바탕으로 AI 산출물의 활용 단계에서 벌어지는 침해행위에 대한 공동 대응성격의 입법론을 제시할 수 있을 것으로 생각된다. (Schack, H., "The Law Applicable to Unregistered IP Rights After Rome II.", *Ritsumeikan Law Review*, No. 26, 2009, pp.129-144, p.130)

55) 베른 협약(Berne Convention for the Protection of Literary and Artistic Works)
"Article 5

(1) Authors shall enjoy, in respect of works for which they are protected under this Convention, in countries of the Union other than the country of origin, the rights which their respective laws do now or may hereafter grant to their nationals, as well as the rights specially granted by this Convention.

(2) The enjoyment and the exercise of these rights shall not be subject to any formality."

(3) Protection in the country of origin is governed by domestic law.

무역관련 지식재산권협정(TRIPS Agreement, 1994)

"Article 1 —Minimum Standards of Protection) "Members shall give effect to the provisions of this Agreement. Members may, but shall not be obliged to, implement in their law more extensive protection than is required by this Agreement."

활용 단계에서 산출물에 기반하여 청구하는 대가를 기존 저작권법에서 파생되는 새로운 권리로 규정하게 될 경우, 기존 국제 규범과의 양립 가능성을 두고 해석상 논란이 생길 여지도 있다.⁵⁶⁾ 기존 국제 저작권법 체계 안에서도 복제권 침해 등을 구성할 때 복제의 사실 관계 등 이를 해당 법문의 구성요건에 맞게 입증하는 과정이 필요한데, AI 모델은 앞서 설명한 대로 학습 데이터와 모델 사이, 학습 데이터와 산출물 사이의 인과관계를 정밀하게 추론하기 어렵기 때문에 기존 협약 등에 명시되지 않는 새로운 권리나 의무 관계를 창설할 경우 국제 저작권 체계의 기존 구조를 흔들어 오히려 혼란을 부추길 가능성이 있다.⁵⁷⁾

이런 논란을 피하면서도 합당한 대가를 이전하기 위해 2단계인 ‘활용 단계’에서 부과하는 금원을 특별한 조세나 문화 정책적 기여금 방식을 부과하는 방식으로 접근하게 되면 정책 목적을 달성하면서도 기존 저작권법 밖에서 효과적으로 규율하는 것이 가능할 것이다. 또, 시간적으로도 기존의 법체계를 수정하거나 새로운 해석적 접근을 협의하는 것보다 상대적으로 짧은 시간에 신규 규범 창설을 통해 창작 생태계에 주는 피해에 대한 적절한 보상 방안을 강구할 수 있을 것으로 사료된다.⁵⁸⁾ 이런 접근이 전례가 없었던 것도 아니다. 실제로, 많은 나라들은 자국의 문화산업 정책이나 조세정책의 일환으로 특정 산업에 대한 목적세나 부담금을 부과하고 있으며, 정책적 재량도 보유하고 있다.⁵⁹⁾

다만, 이 같은 논의는 AI 산출물 자체의 저작권을 어떻게 인정할 것인가에 따라 논리 전개가 달라질 수 있어 이에 대한 사전 논의도 중요할 수 있다. 대체로, 아직까지는 AI 산출

(Article 9(1) — Berne Incorporation Clause) “Members shall comply with Articles 1 through 21 of the Berne Convention (1971) and the Appendix thereto, except for Article 6bis on moral rights.”

56) 이는 베른협약 제5조의 내국민대우 원칙 및 TRIPS 협정의 3단계 테스트(three-step test)와의 해석상 논란을 야기할 수 있다

57) 반면, 조세법적 접근 또는 문화정책적 기여금 방식은 저작권법 체계 밖에서 작동하므로 베른협약이나 TRIPS 협정의 직접적인 규율 대상이 되지 않는다. 각국은 자국의 문화산업 정책이나 조세정책의 일환으로 특정 산업에 대한 목적세나 부담금을 부과할 수 있는 정책적 재량을 보유하고 있으며, 이는 저작권의 국제적 최소 보호 기준과는 별개의 영역으로 볼 수 있다.

58) 이와 같은 조세법적 접근은 국내에서도 제안된 바 있다. 남형두(2025)는 “공정이용 판단이라는 사전 단계에서 이를 막아내지 못한 경우 사후에 조세적 해법으로 치유하는 것을 고려해 볼 수 있다.”고 주장한 바 있는데, (남형두, 앞의 책, 448면.) 필자는 이 같은 접근에 동의하면서 조세적 접근의 차별적 적용을 주장하고자 한다. TDM 법제를 단일 과정으로 규정하고, 보상이나 과세 개념을 적용할 경우 생길 수 있는 논리적 모순점 등을 피하기 위해 AI 파운데이션 모델의 개발 및 보급 과정을 1단계(학습 단계)와 2단계(활용 단계)로 나누어, 이중 구조로 설계한 뒤, 2단계에서 조세법적 접근을 적용하자고 제안하는 것이다.

59) 실제로 프랑스의 구글세(Google Tax), 각국의 디지털세 논의 등은 저작권법이 아닌 조세법 체계를 통해 디지털 경제의 새로운 과세 문제를 해결하려는 시도로 볼 수 있다.

물에 대한 저작권성을 인정하지 않고 있는 추세이나 일부 인정하는 움직임도 있는 만큼 향후, 인정 여부에 따라 후속 조치도 달라질 수 있다.⁶⁰⁾ 저작권성이 인정될 경우 상업적 활용에 대한 경제적 대가 지불의 명분은 더 강해질 것이고, 만약 저작권성이 인정되지 않는다 하더라도 해당 산출물을 통해 수익활동을 하는 경우, 주인이 없는 무주물(無主物)을 통해 얻은 수익의 공익적 환원 개념도 상정할 수 있을 것이다. 또, 인정 여부에 관계없이 AI 모델의 사용료를 받는 빅테크 기업의 부가가치에는 정당한 대가를 요구할 수 있는 논리를 갖출 수 있을 것으로 생각된다.

이런 점에 비춰볼 때, 2단계인 인공지능 모델 활용 단계에서 부과하는 대가 청구행위를 '저작권 보상'이 아닌 'AI 창작 생태계 기여금'이라는 별도의 조세적 성격으로 구성할 경우, 기존 체계와의 모순이나 충돌을 피하면서도 AI 시대의 창작 생태계 보호라는 실질적 목적을 달성할 수 있을 것으로 기대된다.

<표 3> AI 모델 활용 단계 과세 논의를 저작권법으로 규율하기 어려운 이유

구분	주요 내용 및 함의
인과관계 입증의 어려움	- AI 모델에는 원본 데이터가 남아있지 않으며 알고리즘만 존재. - 특정 산출물에 특정 저작물이 기여했는지 인과 추정 불가능. → 저작권법상 인과관계 요건을 충족하는 것이 어려울 수 있음.
집단적으로 발생하는 피해의 성격	- AI 모델 학습 이후, 대규모 산출물 생산이 단시간에 가능. - 특정 저작자만 피해를 입는 게 아니라 산업 생태계 전체 영향. → 저작권들의 개별적인 권리행사로는 체계적 대응이 불가능.
기존 국제 저작권 법무 체계와의 정합성	- 베른 협약 등 기존 국제 저작권 법무체계로는 규율이 어려움. - 새로운 권리 창설이나 해석 확장 등 국제적 합의에 시간 소요 → 기존 국제 규범과 충돌하지 않는 신규 트랙 창설이 필요함.

이런 이해를 바탕으로 2단계 인공지능 모델 활용 단계에서 일종의 '준과세 체계'를 구축하게 될 경우, 과세 대상은 누구로 할지가 문제될 수 있는데, 본고에서는 크게 생산자와 소비자 두 가지 측면으로 나누어 살펴보고자 한다.

우선, 크게 빅테크에 직접 부과하는 세금과 빅테크 사용자에게 부과하는 세금으로 나누

60) 예컨대, 미국 저작권청은 2023년 *Zarya of the Dawn* 사건에서 AI가 활용된 만화책에 대해 AI 생성한 부분에 대해서는 선택적으로 저작권 등록을 거부한 바 있으며, 우리나라 저작권 규제 역시 AI 산출물의 저작물성은 명시적으로 인정하지 않고 있다. 다만, 중국 베이징인터넷법원은 2023년 *Li v. Liu* 사건에서, 프롬프트 설계 등 사용자가 창작적 기여를 했다고 인정되는 경우, 저작물성을 인정하기도 하였다.

어 접근할 수 있다. 빅테크에 직접 부과하는 세금의 경우에는 직접적으로 구독료를 상정할 수 있다. 챗GPT나 Claude, Midjourney나 Kling 같은 유료 플랜 구독 모델들의 경우, 구독료 수입의 일부가 과세 표준이 될 수 있으며 해당 구독료는 B2C 형태를 통한 직접 구독 모델과 API 형태의 B2B 형태의 구독 모델이 모두 포함된다.

다음으로, AI 모델을 활용한 사용자(user)에게 부과하는 경우로는 AI 모델 산출물의 직접 판매 수익 등을 들 수 있다. 예를 들어, AI로 만든 소설이나 만화, 이미지 등을 직접 판매할 때 수익이 발생하게 되는데 이 경우 사용자가 거두는 수익의 일부도 2단계 활용 과세의 일종의 과세 표준에 포함할 수 있다.⁶¹⁾

IV. 산출물 기반 과세적 보상에 대한 이론적 정당화

1. 예상 반론에 대한 이론적 재반론

본고에서 제안한 이중 구조 보상 모델의 핵심은 AI 모델을 학습 단계와 활용 단계로 나누어 규제를 하자는 것을 골자로 한다. 그 중에서도 2단계인 AI 모델활용 단계에서의 산출물 기반 과세적 보상 체계의 신규 도입은 직관적으로 이해할 수 있는 제안이기는 하지만 이론적, 실무적 관점에서 다양한 비판과 우려, 또는 반론이 제기될 수 있다. 본고의 제안안이 기존 국제 법무 체계에는 없던 새로운 제도인 데다, 조세 제도의 본질상 그 형태를 막론하고 침익적인 재정적 처분에 대해서는 조세 저항에 직면할 수 있기 때문이다. 본 절에서는 필자가 제시한 이중 구조 보상 모델에 대해 예상되는 반론들을 검토하고, 이에 대한 이론적인 재반론을 통해 논지를 강화하고자 한다.

(1) 이중 과세 반론에 대한 검토

가장 먼저 제기될 수 있는 반론은 본고에서 제안하는 새로운 산출물 기반 부담 체계가

61) 텍스트 기반의 콘텐츠를 유통하는 웹소설 플랫폼의 경우, 빠르게 상업화를 진행하면서 플랫폼에 입점한 작가들의 수익이 연간 3천만 원을 넘어섰다는 조사 결과도 있다. 만약 이런 수익화 과정에서 생성형 AI를 도구를 활용할 경우 이에 대한 판매 수익에 대해 간접적으로 과세를 검토할 수 있을 것으로 생각된다. (김옥조 기자, “‘웹소설’ 산업 규모 1조 원 돌파·이용자 수 6백만 명 육박”, KBC뉴스, 2023년 9월 7일 보도.) 다만, AI 모델에 붙는 광고 수익 등의 간접 수익의 처리 문제 등은 과세 체계를 가다듬을 때 고려 요소가 될 수 있고, 개인에게 과세하는 경우에도 비상업적 사용이나 학술적, 교육적 목적의 사용 등에 대해서는 공제 혜택을 부여하는 보완적 방안이 검토되어야 할 것으로 생각된다.

사실상의 이중과세 내지는 이중 부담 문제라는 주장이다. 특히, AI 모델을 개발하기 위해 이미 거금을 주고 학습 데이터를 구매해 보상을 마친 경우라면 이미 대가를 지급했는데 모델 활용 단계에서 추가로 금전적 부담을 진다면 이를 사실상의 이중과세로 간주하고 부담하다는 주장을 펼 수 있을 것으로 예상된다.⁶²⁾

그러나, 이러한 주장에는 그릇된 전제가 내포돼 있다. 즉, 1단계의 인공지능 모델 학습 행위와 2단계에서 학습을 마친 인공지능 모델의 산출물 생성 행위를 '동일한 행위'로 간주한다는 전제가 깔려 있다고 볼 수 있다. 그러나, 앞서 필자가 분석한 대로 1단계 모델 학습단계와 2단계 모델 활용 단계는 시간적 측면에서, 법적 측면에서, 또 실제 생태계에 미치는 경제 효과 측면에서도 분명히 구분되는 별개의 행위이다.

첫째로, 1단계에서는 학습을 위해 데이터를 복제하는 그 행위 자체를 규율하는 것이며 사용자의 동의 없이 전송, 복제하는 행위에 대한 대가를 지불하도록 하는 것이다. 이는 전통적인 저작권법이 주목해 온 영역에 속한다. 반면, 2단계는 활용 시점에서의 상업적 경쟁에 영향을 미치는 행위에 대해 규율하는 것이다. AI 모델이 산출물을 만들어 낼 때는 기존의 저작물을 참고하지 않으며, 빅데이터를 통해 구축된 확률적 분포를 통해 추론하게 된다. 이를 통해 지속적으로 경제 활동을 하면서 생태계 전체에 영향을 주게 된다. 즉, 복제와 경쟁은 보호하는 법익이 다르며 서로 영향을 주지 않는다고 반론을 제시할 수 있다.

또, 1단계에서의 보상은 사적 권리인 저작권 침해에 대한 일대일 보상이다. 이는 민사적 차원의 사적 보상 체계에 가깝다고 할 수 있다. 반면, 2단계 보상은 공적 권리를 보호하는데 목적이 크다고 할 수 있다. 저작권 생태계 전체에 대한 대세적 의무, 즉 일대다(一對多) 보상이다. 경제학에서 말하는 외부효과를 내재화하는 수단으로, 이는 공법 영역의 보상 성격을 지니게 되어 보호 법익에서 차이가 발생한다.⁶³⁾

62) 미국의 비영리단체인 납세자보호연맹(Taxpayers Protection Alliance)는 인공지능 과세에 대한 반대 입장을 밝히면서 알고리즘, 인공지능 모델, 이로 인한 산출물 등에 부과하는 세금을 차별적 세금(discriminatory tax)이라고 규정하기도 했는데, 기술 발전 장려를 명분으로 내세운 조세 저항이 이미 나타나고 있다. (TPA, "Taxing Bias: ALEC's AI Equity Framework", July 18, 2025)

63) 1단계에서의 보상은 사법(私法)상 보상이고, 2단계는 공법(公法)상 조세인데 양자를 어떻게 단일한 보상 모델로 포섭할 수 있는가 하는 의문이 제기될 수 있으나, 본고가 보상(Compensation)이라는 용어를 채택한 것은, 두 단계 모두 그 형식과 명칭이 무엇이든 그 종국적 효과가 창작자에 대한 가치 환류(feedback)라는 점에서 공통되기 때문이다.

끝으로, 해당 규범이 적용되는 관할권의 시간적 범위 또한 차이가 있다. 1단계 인공지능 모델 학습 단계에서 이뤄지는 보상의 경우, 과거에 이미 지나간 데이터 복제행위, 그 일회성 행위에 대한 보상인 반면에, 2단계 인공지능 모델 활용 단계에서 이뤄지는 과세 처분의 경우에는 과거가 아니라 현재와 미래 시점에 대해서도 지속적으로 효력이 미치는 행위이다.

이 세 가지 특징들을 종합해보면 이중과세나 이중 부담이라고 주장하는 반론에 대해서는 과세 성격과 보호 법익, 해당 처분의 시간적 관할권 측면 명확히 구분되는 별개의 행위라는 점을 분명히 할 수 있다. 다른 행위에 대해 다르게 별도로 처분을 하는 것이다. 인간이 참고서를 구입해 스스로 학습을 할 때도 참고서를 구입할 때 이미 돈을 냈다고 해서 학습의 결과로 발생한 소득에 대해 소득세를 면제해 주지는 않듯이 인공지능 모델도 학습시 보상을 마쳤으므로 한 번만 데이터 소유주 측에 보상금을 내고 무제한으로 수익을 창출할 수 있다는 주장을 펼친다면 오히려 형평을 해치게 되는 결과로 이어진다고 할 수 있다. 그러므로 본고에서 주장하는 ‘이중 구조 접근법’은 이중 과세가 아닌 AI의 특수성에 기반해 법률상(de jure) 형평뿐 아니라 사실상(de facto)형평을 보장하는 적합하고, 명확한 과세라고 주장할 수 있다.

(2) 기업 혁신을 저해할 거란 반론에 대한 검토

두 번째 예상 반론은 AI 관련 규제가 지나칠 경우, 기업들의 기술적 혁신을 저해할 수 있다는 우려이다. 지난 2022년 말, 챗GPT가 전세계적으로 부상한 이후, AI 모델 개발업체들, 이른바 빅테크들은 관련 기술을 급속도로 발전시키며 경쟁적으로 신규 모델들을 출시하고 있다. 기술이 가져올 변화가 파괴적이다 보니 인공지능 파운데이션 모델은 국가의 전략적 기술(strategic technology)로도 간주되며, 국가별 경쟁도 격화하고 있으며, 이른바 ‘소버린 AI’ 개발 지원 사업으로 대표되는 국가적 차원의 개발 지원도 이뤄지고 있는 상황이다.⁶⁴⁾ 이런 상황에서 필자가 제시한 이중 구조 보상 모델은 가뜩이나 수익화가 먼 AI 기술 기업들의 투자 유인을 떨어뜨리고, 개발 속도를 늦춰 혁신을 해칠 수 있다는 반론에 직면할 수 있다. 특히, AI 후발 주자인 한국이나 유럽의 입장에서 보면 기술이 무르익기 전 선도적 규제가 등장하면 국내 관련 산업이 위축되고 경쟁력이 약화될 수 있다는 주장

64) 실제로 인공지능 기술은 인류에게 많은 편익을 제공하며 업무 효율성을 높이고, 의료, 교육, 콘텐츠 제작 등 다양한 분야의 혁신도 이끌어 내면서 긍정적 영향도 많이 끼치고 있는 것이 사실이다.

이 나올 가능성이 더 높다.⁶⁵⁾

이 비판은 실제 기업에 규제가 한층 추가된다는 면에서는 일견 타당하다. 그러나, 그렇다고 해서 아무런 규제 없이 이뤄지는 무제한의 TDM 허용도 앞서 살펴본 여러 부작용을 낳을 수 있다. 이런 반론에 대한 재반론과 대안을 제시하고자 한다.

우선, 적정 세율 설정을 통해 혁신 저해를 최소화할 수 있는 방안을 모색해 볼 수 있다. 이중 구조 보상 모델은 기업의 수익성에 치명적인 수준의 징벌적 과세를 하자는 논의가 아니라 모델 학습과 활용에 대한 생태계 균형을 추구하자는 것으로, 수익 기반 과세를 추구하고 있어 기업의 부담을 최소화하면서도 정책 목적을 달성할 수 있는 세율을 설정할 수 있다. 스타트업이나 중소기업의 경우 세제 혜택을 부여하는 식으로도 공존할 수 있는 보완책이 가능하다.

두 번째는, 비용 감소 효과도 고려해야 한다. 많은 AI 기업들이 규제가 늘어나면 비용이 늘어날 거라고 생각하지만 규제가 없어도 오히려 그 불확실성에 대한 비용이 늘어나게 된다. 신산업인 AI에 대한 특별 규제가 존재하지 않게 되면, 기업 입장에서는 제대로 포섭할 수 없는 기존 전통적 법률(conventional law)에 기초해 의사결정과 리스크 분석을 할 수밖에 없으며, 이 과정에서는 컴플라이언스 비용이 오히려 증가할 수 있다. 즉, 규제가 없으면 모든 규제와의 상충을 검토하는 과정에서 기업들이 혁신을 오히려 저해받을 요소도 있는 것이다. 실제로 오픈AI를 비롯한 미국 주요 빅테크들이 단일 규제를 창설해달라고 요청하기도 했는데,⁶⁶⁾ 이는 컴플라이언스 비용을 낮추고자 하는 동인으로 해석할 수 있다. 이런 관점에서 보면 명확한 1단계 보상과 2단계 과세 체계를 갖추게 되면 기업들의 법적 예측 가능성이 높아지고, 불확실성이 해소되는 효과도 가지게 되는 것이다.

65) 2025년 9월에 개최된 대통령 주재의 제1차 '핵심 규제 합리화 전략회의'에서는 AI 학습 과정에서 발생하는 TDM 문제와 관련해 저작권 책임을 면제해야 한다는 주장이 기술 업계를 중심으로 제기되었다. 현재 국내외를 막론하고 세계적으로 TDM 면제 요청의 주요 근거가 되는 주장이다. (장민권 기자, "AI 학습 저작권 책임 면제해야"···"TDM 도입 등 AI 규제혁신 목소리 봇물", 파이낸셜뉴스, 2025년 9월 15일 보도.)

66) OpenAI는 AI 시스템이 안전하고 광범위하게 유익하도록 유지하기 위해 규제(regulation)가 필요하다는 입장을 공개적으로 밝히고 있으며, (Open AI, "Our approach to AI safety", April 5, 2023 : <https://openai.com/index/our-approach-to-ai-safety/>) 오픈AI 창업자인 샘 알트만도 인터뷰를 통해 적절한 규제가 필요하다고 밝힌 바 있다. (Ed Osmond, "OpenAI CEO says possible to get regulation wrong, but should not fear it", Reuters, September 26, 2023)

<표 4> 예상되는 반론에 대한 본고의 이론적 재반론

구분	예상되는 반론	재반론 요지와 논거
(1) 이중과세	AI 모델학습을 할 때 이미 데이터에 대해 보상을 마친 상황에서, 다시 모델 활용 시 과세를 하면 이중과세 일 수 있다는 반론.	- 1단계와 2단계는 행위와 보호 법익 시간적 관할권 측면에서 구분됨. 1) 행위 구분 : 복제 이슈 vs 경쟁 이슈 2) 법익 구분 : 민사적 vs 공법적 3) 시간 구분 : 사전적 vs 사후적
(2) 혁신저해	AI 기술 개발이 경쟁적으로 이뤄지는 상황에서 선도적으로 규제가 생기면 기업들의 투자 유인을 약화시켜 기술 경쟁력과 혁신을 저해할 거라는 주장.	- 초기에 규제를 마련하면 법적 안정성과 기술의 예측 가능성을 높일 수 있음. - 적절한 세율 설정을 통해 조절 가능. - 규제가 불확실성 해소 측면도 있음. - 이를 통해 기업들의 규제 준수비용은 오히려 줄어들 수 있음.
기대효과 : 제도의 정당성과 생태계의 지속성을 높여 공정한 경쟁과 완전한 형평 추구		

끝으로, 생태계 존립과 장기적 득실 분석을 통해 기업 혁신을 저해할 수 있다는 반론에 대해 재반론을 제시할 수 있다. 물론, 기존에 존재하지 않던 과세 체계 등이 새로 창설되면 기업 입장에서는 단기적으로 부담으로 인식할 수 있다. 그러나, AI 기업 관점에서도 길게 보면 본고에서 제안하는 수익 환류 체계가 기업에게도 이익이 될 수 있는 측면이 있다. 마치 황금알을 낳는 거위의 배를 가르듯이 창작 생태계가 일시에 붕괴해 버리게 된다면 AI 모델이 추후 추가 학습할 데이터도 함께 사라질 수 있기 때문이다.⁶⁷⁾ 따라서, 단기적으로 창작 생태계에 기여하는 과세는 조세적 성격을 띄게 되어 기업들에게 부담을 줄 수 있지만 장기적 관점에서 보면 기업들에게 오히려 이익이 될 수 있는 측면도 있다. 양질의 학습용 데이터를 안정적으로 확보할 수 있는 환경을 만들 수도 있는 것이다.

결론적으로, 2단계 과세 체계가 포함된 이중 구조 보상 모델이 혁신을 저해한다는 우려는 오히려 기업의 예측 가능성을 높이고, 건전한 데이터 생태계를 유지할 수 있다는 점에서 AI의 지속가능성을 보장한다는 취지에서 반론이 가능하다. EU가 제정한 데이터보호규정인 GDPR이 처음 도입됐을 때도, EU 기업들만 불리하다는 우려가 있었으나 결과적으로는 선도 입법이 되어 많은 글로벌 기업들이 관련 규범을 준수하는 것을 오늘날 목도하고 있다.⁶⁸⁾ EU는 인공지능법도 세계 최초로 제정하면서 관련 규범 질서를 만들어 가고 있는

67) AI 모델이 생성한 콘텐츠를 다시 AI가 모델 학습에 사용할 경우, 모델 품질이 저하될 수 있다는 모델 붕괴현상에 대한 연구도 있다. (Shumailov, I. et, "AI models collapse when trained on recursively generated data", *Nature*, 631(8031), pp. 755-759. <https://doi.org/10.1038/s41586-024-07566-y>)

데, 규제로 인한 혁신 저해에 대한 우려보다는 규제 선점의 관점에서 적극 모색할 때 얻을 수 있는 이점도 크다고 생각된다.

2. 제도 실현을 위한 실행 로드맵

이중 구조 보상 모델은 앞서 살펴본 대로 AI 모델의 속성과 이론적 근거를 바탕으로 살펴볼 때, AI 모델에 대한 과세 체계 도입의 법적 근거는 설득력 있게 제시할 수 있다. 다만, 이런 제도를 실제 구현하는 과정에서의 현실적 고려 사항들을 짚어볼 때 몇 가지 고려할 사항들이 존재한다.

(1) 국제법무 협력의 필요성

현재 지구상에서 주로 소비되는 대부분의 AI 파운데이션 모델은 미국과 중국이 주도하고 있다. 생성형 AI의 시작을 알린 오픈AI를 비롯해서 구글의 Gemini, 엔트로픽의 Claude 까지. 모두 미국에서 개발된 모델이다. 그리고 이를 맹렬하게 추격하는 모델들은 주로 중국 기업들에 편중돼 있다. 특히 중국의 경우, 이미지와 영상 생성 모델 분야에서 두각을 나타내고 있는데, 이렇다 보니 전 세계 사용자들이 역외 서비스 이용 형태로 AI 모델을 사용하고 있는 실정이다.

이처럼 현재까지 개발된 AI 모델이 소비자에게 전달되는 경로는 국경을 넘어 역외로 제공되는 경우가 많아 주권 국가가 입법 관할권을 가지고 있어도 이를 실효적으로 집행하기는 어려울 수 있는 것이다. 즉, 효율적인 집행관할권 확보가 한계로 지적된다.

이 같은 문제를 해결하기 위해서는 국제법 및 국제경제법적 접근과 연구가 필요하다. 과거 통상규범은 디지털 재화에 대해서는 규율하는 데 한계가 명확했는데, AI 모델은 단순한 디지털 서비스를 넘어 복합적으로 작용하는 재화로서 이를 규율하기 위한 국제 사회의 규범 창설이 반드시 필요하다고 생각된다. 우리나라가 국내법 제정의 관점에서 독자적으로 규제를 설계하고 제안하는 데 그칠 것이 아니라 본고의 제안을 비롯하여 선도적 입법 모델을 제시하고, 이를 국제기구 차원의 논의로 의제를 이끌어 갈 필요가 있다.⁶⁹⁾

68) 당초, GDPR이 처음 도입될 당시 유럽 기업들이 경쟁력이 떨어질 것이란 우려도 나왔었지만 규제의 세계화를 촉진하는 이른바 '브뤼셀 효과'를 나타내며 각국으로 확산하여 글로벌 규범화하였다. Diba Aalami Harandi, "The Extra-Territorial Impact of the European Union's General Data Protection Regulation on United States Businesses", University of Denver, April 25, 2025.

필요하다면 양자 협정 등을 통한 해법도 가능하다. 개별 국가와 양자 협정을 통해 AI 관련 규제의 상호 인정과 집행을 보장할 수 있으며, 이 과정에서 EU, 일본 등 상대적으로 미국, 중국에 비해 경쟁에 뒤쳐진 후발주자 국가들과 연대하여 규제 집행을 담보하는 방식의 국제 협력 사례를 늘려나가는 방식을 생각해 볼 수 있다.

그렇다면, 미국처럼 AI 기술 선도 국가와의 협의는 어려울 있다는 우려가 나올 수 있는데, 이같은 AI 선진국들과는 현지의 작가 조합이나 배우 조합 같은 창작자 중심의 이익집단과 연대하여 미국 정부를 설득하는 방식으로 관련 규제를 도입하기 위한 노력에 나서는 방안을 생각해 볼 수 있을 것이다.

또, 이런 국가 차원의 협력 뿐 아니라 글로벌 플랫폼과의 협력 및 규제 적용도 필요할 것으로 생각된다. 구글과 애플 같은 주요 플랫폼들은 이미 각국의 과세 당국과 접촉면이 있으며, 현행 과세 체계에 포섭된 주체이다. 특히, 많은 AI 관련 앱들이 구글 플레이스토어나 애플 앱 스토어 등을 통해 시장에 진입하고 있다는 점을 감안하면 이러한 플랫폼을 통해 원천 징수 형태의 납세 의무를 부과함으로써 과세 체계를 조정할 수 있을 것으로 예상된다.

(2) 특별법 제정을 통한 흠결 보완

이중 구조 보상 모델은 현행 법 체계로는 불가능한 만큼 새로운 규범 창설이 반드시 필요하다. 이 과정에서 기존 저작권법을 수정하는 방안과 특별법을 제정하는 방안을 생각해 볼 수 있는데, 본고에서는 단계별 구분을 통한 입법론을 제안하고자 한다.

우선은 전술하였다시피 1단계는 기존의 데이터를 복제권 관점에서 저작권법과 매개해 보상하는 단계이다. 그러므로, 이 영역은 기존의 저작권법의 틀에서 소화할 수 있다. 이 때문에, 1단계 TDM 과정 자체와 관련한 입법은 기존 저작권법 내에 TDM 예외와 보상 체계를 규정하여 보완하면 기존 저작권법과 양립하는 규제 체계를 창설 할 수 있을 것으로 예상된다. 이 경우, 기존 저작권법에서 보호하는 관리 체계나 분쟁 조정 절차, 구제 수단 등을 그대로 활용할 수 있어 저작자 입장에서 유리하다고 할 수 있을 것이다.

69) 대표적으로 OECD나 WIPO 등의 국제기구를 들 수 있는데, OECD에서는 디지털 플랫폼에 대한 과세에 관해 논의한 적이 있으며, 자금 세탁 방지 등도 국제적 연대를 통한 규율에 나서고 있다. OECD/G20 Base Erosion and Profit Shifting Project Outcome Statement on the Two-Pillar Solution to Address the Tax Challenges Arising from the Digitalisation of the Economy

그러나, 2단계의 경우에는 새로운 형태의 과세 체계를 창설하는 것이다. 그러나, 개념상 과세 체계에 가깝다는 것이지 세법을 매개로 하는 전통적인 조세 개념과는 약간 차이가 있다. 따라서, 2단계 모델 활용 단계에서의 과세는 특별법 제정을 통해 규율하는 것이 바람직하다고 생각된다.

앞서 언급했던 저작권법 민사 관계 규정이 본질이라면 2단계는 공적 의무를 부과하는 단계인 만큼 새로운 법적 의무를 창설할 수 있는 전용 법률이 필요하다. 우리 법에서는 자연인(自然人)과 법인(法人)의 구분만 있을 뿐 인공지능(人工人)에 대한 규율이 없는 만큼 AI의 특성을 충분히 반영한 새로운 법을 제정할 필요가 있다. 이렇게 제정하는 새로운 법에는 앞 절에서 살펴본 국제법 규범과의 조화에 필요한 이행 법률적 요소도 포함되는 것이 바람직할 것이다.

이를 통해, 단순한 언어모델(language model)뿐 아니라 음악, 음성, 이미지, 영상까지 다양한 유형의 생성형 AI 모델을 포괄할 수 있도록 유연한 제도 설계도 가능할 수 있을 것이다. 일종의 디지털 기여금 형태를 부과하는 법적 근거를 창설하면서도 기술적 유연성과 중립성을 확보한 포괄적 규제를 설계하는 것이 중요할 것이다. 우리나라의 입법 실정을 되짚어보면 서로 파편화되어 있는 이익집단의 이해(interest)를 조율하는 과정이 순탄치 않은 만큼 본 고에서 제안한 특별법 제정 담론은 사회적 합의 도출을 위한 지난한 조율 과정을 요하고, 적잖은 사회적 비용을 지불할 수 있을 걸로 예상된다. 따라서, 입법 과정에서 기술적 가변성을 수용할 수 있는 유연한 설계 못지 않게 이해관계자 간의 소통을 가능케 할 '상설 거버넌스' 체계도 필요할 것이다.

V. 결론 : 생성형 AI 시대 사회계약의 재구성

본 연구는 인공지능 기업들이 대량의 데이터를 활용해 생성형 AI 모델 학습을 진행하면서 생기는 이른바 TDM 이슈와 관련해, 활용되는 저작물에 관한 보상 체계에 대한 새로운 대안을 입체적으로 제시하고자 하였다. 현행 TDM 관련 법제는 AI 모델학습 과정에서의 일회성 보상에만 주로 집중할 뿐 AI 기술이 창작 생태계 전반에 미치는 파괴적 영향은 주목을 받지 못하고 있는 실정이다. 따라서, 필자는 이를 보완하기 위해 기존의 TDM 논의에 더해 산출물 기반의 과세적 보상 제도를 추가한 '이중 구조 보상 모델'을 제안하였다.

본 모델의 핵심은 AI를 모델 학습 단계와 산출물 생성 및 활용 단계로 나누어 각 단계별 법익을 별도로 분석해 차등적 규율 체계를 확립하고자 하는 데 있다. 1단계인 학습 단계에서는 기존의 전통적 저작권법 틀을 확대 발전시켜, 데이터 유형별로 차등적 보상안을 설계하여 제안하였고, 2단계인 AI 모델 활용단계에서는 저작권법 밖에서 조세법 관점을 더해 일종의 ‘창작 생태계 기여금’을 추가하는 이원적 구조를 제시하고자 하였다. 복제권이라는 사적 이익과 창작 생태계 보호라는 공적 이익을 조화롭게 추구하고, 기존 저작권법의 한계도 극복할 수 있는 대안이 될 것으로 확신한다.

일각에서는 단순히 세금을 추가한 정도로 단순화해 이 모델을 평가할 수도 있겠지만 필자는 이 같은 시도가 인류가 쌓은 문화적 자산을 통해 학습한 AI와 인간 창작자들이 어떻게 공존할 수 있는지 중요한 물을 제정하는 계기가 될 것이라고 확신한다. AI 대전환기에 제도적 첫 단추를 어떻게 끼우느냐에 따라 그 효익과 피해가 극명하게 갈라질 수 있기 때문이다. 1,2차 산업혁명이 태동하면서 노동시간 문제, 안전 문제, 나아가 환경 문제까지 다양한 사회 계약이 파생되었듯이. 또, 인터넷으로 대표되는 3차 산업혁명을 거치면서 개인 정보 보호 문제, 사이버 범죄 문제 등의 사회 계약도 파생되었듯이. AI와 로봇으로 대표되는 4차 산업 혁명은 창작 생태계와 저작권이라는 문화 사회계약을 재정의 할 것을 강하게 요구하고 있다. 이는 18세기 저작권법이 출현한 이래 가장 큰 패러다임 전환이라고도 볼 수 있을 것이다.

물론, 이런 구조를 설계하는 것은 단일 국가의 노력만으로는 어려울 수 있다. AI 서비스가 미국, 중국 중심의 일부 선도국들에 의해 주도된다는 점, 고도의 AI 서비스를 제공하기 위해서는 막대한 하드웨어 인프라가 필요하고 이를 위해서는 데이터가 역외 전송 처리될 수밖에 없는 점 등을 고려할 때 국제 협력과 관련 국제법 제정, 이에 대한 국내 이행입법 제정 등 후속 과제도 만만치 않은 과정을 요구한다. 그러나, 이런 어려운 여건을 딛고 우리가 선도적으로 규제를 제안하고 주도할 수 있다면 EU의 GDPR이나 AI ACT가 글로벌 벤치마크 규범으로 자리 잡은 것처럼 오히려 규범의 리더십을 확보할 수 있을 것이라고 생각된다. 특히, 글로벌 콘텐츠 시장에서 트렌드를 주도한 우리 한국이 콘텐츠 창작 생태계를 보호할 수 있는 형태의 새로운 문화 사회계약을 제안한다면 이는 국제적으로 강한 신호를 보낼 수 있고 이를 규범 권력화 할 수 있을 것이라고 생각된다. 특히, 본고에서 규정한 2단계 활용 단계에서 적용 받게 되는 특별법은 AI의 법적 지위를 명확히 하고, 자연

인과 법인을 넘어선 '인공인'에 대한 규율 체계를 최초로 정립할 수 있다는 점에서 법제사적인 의의도 크다고 할 수 있으며, 이 경우 단순히 저작권법의 한계를 극복하는 것을 넘어 AI 시대 법제 전반의 재구성을 촉진하는 계기를 마련할 수 있을 것이다.

단, 본 연구에서도 AI 생태계에서 발생할 수 있는 모든 경우의 수를 포섭하는 완결성을 가진 세부 규제 체계까지는 제시하지 못하였다. 앞서 언급한 대로, 과세 개념을 도입한다면 적정 세율은 얼마가 좋을지, 구체적인 징수 메커니즘은 어떻게 가져가야 할지 제도의 세부 설계에 대해서는 후속 연구가 절실하다. 또, 자칫 이같은 조세적 접근이 오히려 빅테크들에게 정당성을 부여할 수 있을 거란 우려도 상존하고 있다.⁷⁰⁾ 세법이나 경제법, 경쟁법 등 관련 전문 분야와의 통섭(統攝)이 중요할 것으로 생각된다.

그리고 나아가 이렇게 취득한 새로운 공동체의 재원을 어떻게 분배하고 공유할 수 있을지, 이것을 이른바 기본소득(Universal Basic Income) 형태로 가져갈지를 포함한 보다 더 큰 틀에서의 사회계약이 필요하다. 이는 빅테크의 성장이 개인의 '인지잉여(Cognitive Surplus)'와 법의 지원에 기반하고 있음을 지적하며, 그 성장의 과실을 일종의 부당 이득(Unjust enrichment)로 규정하며 개발 이익 환수 차원에서 공적으로 회수하여 기본소득 재원 마련의 단초로 삼아야 한다는 논의와 궤를 같이 한다.⁷¹⁾ 다만, 구체적인 분배 메커니즘에 대한 설계는 본 연구의 한계로 남겨두며 후속 연구의 과제로 돌리고자 한다.

결국, 본고에서 제안한 새로운 보상 모델이 만들어지고, 국제 법무체계에 도입되는 과정은 AI 대전환기 새로운 형태의 문화 사회계약이라고도 할 수 있을 것이다. 기술 혁신을 저해하지 않는 선에서 창작자의 정당한 권리를 보장하고, 공익의 균형을 이룰 수 있는 사회적 계약을 만들어 낼 수 있다면, AI 기술의 발전과 지속 가능성을 담보할 수 있을 것으로 예상된다.

70) 남형두, 앞의 책, 456면.

71) 남형두, “잉여(剩餘) —빅테크와 양봉업자—”, 『법철학연구』 제25권 제2호, 한국법철학회, 2022, 272-273면, 279-281면.

참고문헌

1. 국내 문헌

- 남형두, 『공정이용의 역설 —시소에 올라탄 거인, 균형의 복원』, 경인문화사, 2025년 2월.
- 강미영, “생성형AI의 활용에 따른 저작권 침해와 형사 정책적 개선 방안”, 부산대학교 법학연구 제66권, 제1호통권123호, 2025년 2월, 193~218쪽.
- 김병일, “인공지능 학습데이터의 이용과 저작권 제한”, 『AI법제 Insight』, 25-1호, 118-157면.
- 김성원, 최상민, 이수원, “텍스트·데이터 마이닝 과정의 저작물 이용 면책규정 신설안에 대한 소고 —적대적 생성신경망에의 적용—”, 『지식재산연구』, 2023 vol.18, no.1, 147-180면.
- 남형두, “잉여(剩餘) —빅테크와 양봉업자—”, 『법철학연구』, 제25권 제2호, 한국법철학회, 2022, 217-286면.
- 류시원, “저작권법상 텍스트·데이터 마이닝(TDM) 면책규정 도입 방향의 검토”, 『선진 상사법률연구』, No.101, 347-390면.
- 박경신, “[2015-22 미국] 항소법원, 구글의 도서 디지털 변환 사업인 구글 북스는 공정 이용에 해당 한다”, 한국저작권위원회 저작권 동향, 2025년 11월 6일.
- 박희영, “인공지능 학습과 저작권 제한 —유럽 저작권지침 및 독일 저작권법의 텍스트 및 데이터마이닝(TDM)”, 『Copyright Issue Report』, 한국저작권위원회
- 손영화, “생성형 AI에 의한 창작물과 저작권”, 『법과 정책연구』, 한국법정책학회, 제23권 제3호, 2023. 9. 357-389면.
- 이상미, “AI 학습데이터 무단 사용에 대한 저작권 보호 방안”, 『계간 저작권』, 37(2), 79-121면.
- 이지연, 김원오, “저작물 학습자로서 생성형 AI의 저작권 침해 책임”, 인하대학교 법학 연구소, IP & Data 法, 제4권 제1호, 2024년 6월 30일, 75~120면, 108면.

최승재, “인공지능 학습에 대한 미국 법원들의 공정이용 판단”, 『Copyright Trends Report』, 한국저작권위원회, 2025년 7월.

한국저작권위원회, “[이슈 브리프] 생성형 AI 생성물의 저작권 문제 해결을 위한 기술적 접근: 3가지 최신 연구 분석” (2025년 2-2호), 2025년 2월 19일.

고계연 기자, “챗GPT 앱, 누적수익 2.7조원...“AI 앱 시장 지배”, 초이스경제, 2025년 8월 16일 보도. <https://www.choicenews.co.kr/news/articleView.html?idxno=152329>

공인희 편집인, “AI 학습 데이터 논란... 연구자들 “우리 논문 무단 사용됐다”, AI매티스, 2024년 8월 30일 보도. <https://aimatters.co.kr/news-report/ai-report/2916/>

권은주 기자, “챗GPT發 ‘지브리 프사’ 열풍... “애니 작가 밥그릇까지 빼앗나”, 이코노미뉴스, 2025년 4월 3일 보도. <https://www.m-economynews.com/news/article.html?no=52869>

김광연 기자, “클로드 운영사 엔스로픽, ‘저작권’ 소송 작가들에게 2조원 지급키로”, IT조선, 2025년 9월 6일 보도. <http://it.chosun.com/news/articleView.html?idxno=2023092146990>

김옥조 기자, “‘웹소설’ 산업 규모 1조 원 돌파·이용자 수 6백만 명 육박”, KBC뉴스, 2023년 9월 7일 보도.

김호영 기자, “한강, 캄캄한 방에서 공상 즐기던 소녀”, 채널A, 2024년 10월 11일 보도.

방성훈 기자, “AI 여배우 출현에 헐리우드 ‘발칵’...유명 배우들 ‘인간 창작 모욕’”, 이데일리, 2025년 10월 1일 보도.

변희원 기자, “고성능·저비용·개방형 모델 늘어... 일상 속으로 스며든 AI”, 조선일보, 2025년 4월 9일 보도.

송광호 기자, “강연료가 원고료 초과하는 씹쓸한 작가 현실...‘작가노동 선언’”, 연합뉴스, 2025년 4월 29일 보도.

<https://www.yna.co.kr/view/AKR20250429071100005?input=1195m>

이승윤 기자, “오픈AI, 챗GPT 출시 2년 반 만에 연간 매출 100억 달러 돌파”, YTN, 2025년 6월 10일 보도.

장민권 기자, “AI 학습시 ‘저작권 침해 책임’ 면제해야”, 파이낸셜뉴스, 2025년 9월 15일 보도. <https://www.fnnews.com/news/202509151816376363>

장민권 기자, “AI 학습 저작권 책임 면제해야”...TDM 도입 등 AI 규제혁신 목소리 봇물”, 파이낸셜뉴스, 2025년 9월 15일 보도.

장유미 기자, “[AI는 지금] “무려 60만 명 해고”...AI·로봇 직원 등장에 인간 일자리 사라진다”, 지디넷코리아, 2025년 10월 24일 보도.

2. 외국 문헌

Adam Buick, “Copyright and AI Training Data—Transparency to the rescue?”, *Journal of Intellectual Property Law & Practice* 20(3): pp.182–204, 12 Dec 2024, p.186

Ashish Vaswani et al., “Attention Is All You Need”, *Arxiv*, <https://doi.org/10.48550/arXiv.1706.03762>

Chelsea Chung, “UK Halts Its Expansion Of Existing TDM Exception For Copyright Infringements”, *The London Financial*, Feb 17, 2023

Diba Aalami Harandi, “The Extra-Territorial Impact of the European Union’s General Data Protection Regulation on United States Businesses”, *University of Denver*, April 25, 2025.

Ed Osmond, “OpenAI CEO says possible to get regulation wrong, but should not fear it”, *Reuters*, September 26, 2023

Ella Creamer, “US authors’ copyright lawsuits against OpenAI and Microsoft combined in New York with newspaper actions.”, *The Guardian*, 4 April 2025
(<https://www.theguardian.com/books/2025/apr/04/us-authors-copyright-law-suits-against-openai-and-microsoft-combined-in-new-york-with-newspaper-actions>)

Hugh Stephens, “Japan’s Text and Data Mining (TDM) Copyright Exception for AI Training: A Needed and Welcome Clarification from the Responsible Agency”, *hughstephensblog*, March 10, 2024.

<https://hughstephensblog.net/2024/03/10/japans-text-and-data-mining-tdm-copyright-exception-for-ai-training-a-needed-and-welcome-clarification-from-the-responsible-agency/>

Jiachen T. Wang et al., “An Economic Solution to Copyright Challenges of Generative AI”, arxiv, 2024.04, <https://arxiv.org/abs/2404.13964>

Josh Howarth, “Number of Parameters in GPT-4 (Latest Data)”, EXPLODING TOPICS, June 17, 2025. <https://explodingtopics.com/blog/gpt-parameters>

Kiyomaru, H. et al., “A comprehensive analysis of memorization in large language models”, In Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing (EMNLP 2024), pp. 584-592.

Matthew Sag, “Copyright safety for generative AI”, Houston Law Review, 61(2), pp.95-347, p. 302

Miriam Partington, “Switzerland unveils national LLM in tech sovereign push”, SIFTED, September 2, 2025

Molly McInnis, “Advancing Computation Of Just Compensation: An AI-Based Approach To Property Valuation In Eminent Domain”, University of Cincinnati Law Review Vol.94, Oct 12, 2025

P. Bernt Hugenholtz, “The New Copyright Directive: Text and Data Mining (Articles 3 and 4)”, Kluwer Copyright Blog, Institute for Information Law (IViR), July 24, 2019. <https://legalblogs.wolterskluwer.com/copyright-blog/the-new-copyright-directive-text-and-data-mining-articles-3-and-4/>

Schack, H., “The Law Applicable to Unregistered IP Rights After Rome II”, Ritsumeikan Law Review, No. 26, 2009, pp.129-144, p.130

Shumailov, I. et al., “AI models collapse when trained on recursively generated data”, Nature, 631(8031), pp.755-759. <https://doi.org/10.1038/s41586-024-07566-y>

TPA, “Taxing Bias: ALEC’s AI Equity Framework”, July 18, 2025

Zichun Xu, Zhilang Xu, “Tiered copyrightability for generative artificial intelligence: An empirical analysis of China and the United States judicial practices”, WILEY, open access, 22 August 2025. <https://doi.org/10.1002/aaai.70018>

3. 온라인 자료

Copyright, Designs and Patents Act 1988, 29A (legislation.gov.uk) - Copies for text and data analysis for non-commercial research.

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.

OECD/G20 Base Erosion and Profit Shifting Project Outcome Statement on the Two-Pillar Solution to Address the Tax Challenges Arising from the Digitalisation of the Economy.

Open AI, “Our approach to AI safety”, April 5, 2023.
<https://openai.com/index/our-approach-to-ai-safety/>

TRIPS Agreement, 1994, Article 1 – Minimum Standards of Protection, Article 9(1) – Berne Incorporation Clause.

U.S. Copyright Act of 1976, Title 17 U.S.C. § 102(b).

WIPO, “Artificial Intelligence and Intellectual Property”, webpage.
<https://www.wipo.int/en/web/frontier-technologies/artificial-intelligence/index>

대한민국 저작권법 제35조의5 (저작물의 공정한 이용).

베른협약 (Berne Convention for the Protection of Literary and Artistic Works), Article 5.

대법원 2014. 1. 28. 선고 2013다21782 판결.

[국문초록]

생성형 AI 저작권과 이중구조 보상 모델:

‘초능력 문학소녀’를 위한 새로운 문화 사회계약

염규현

본 논문은 대량의 데이터를 활용해 생성형 AI 모델을 학습하는 과정에서 무분별하게 이루어지는 텍스트 수집 행위인, TDM에 대한 새로운 접근법을 제시하고자 하였다. 기존 논의가 전통적인 저작권법에 기초하여 이른바 ‘학습 단계’에서의 일회적 복제와 이용에 대해서만 집중한 나머지, 실제 학습 이후에 광범위하게 상용화되는 과정의 보상 체계를 제대로 반영하지 못하고 있다는 점을 비판적으로 조명하고자 하였다.

이를 위해 거대언어모델(LLM)을 일종의 ‘초능력 문학소녀’로 비유하고, 생성형 인공지능이 가진 ‘초능력성’에 대해 별도의 과세적 개념의 규율이 필요함을 역설하였다. 인공지능 기술이 가진 비대칭적 특성을 반영해 1) 학습 단계와 2) 활용 단계로 구분하고, 학습 단계에 대해서는 저작권법에 기초한 적절한 보상을, 2) 활용 단계에 대해서는 수익에 따른 과세 체계를 적용하는 이른바 ‘이중 구조 보상 모델(Two-tier Compensation Model)’을 도입할 것을 제안하였다.

1단계에서는 TDM을 전통적인 복제권 침해 내지는 일종의 합법적 수용 형태로 간주하여 데이터 유형별 차등 보상 체계를 제안하였고, 2단계에서는 모델 내부에는 원본 데이터가 남지 않는다는 점에 착안하여 저작권법상의 개별 권리 구제 형태가 아닌 ‘창작 생태계’ 전반에 기여하는 기여금 형태의 목적세 도입을 주장하였다.

또, 이런 정책을 도입하는 과정에서 발생할 수 있는 이중 과세 논란, 혁신 저해 주장 등에 대한 반론을 모색하였고, 국제 협력 및 국내법 개정 등 실행 로드맵도 제시하여 생성형 AI 시대 새로운 문화 사회계약의 모델을 제안하였다.

주제어: 생성형 인공지능, 텍스트·데이터 마이닝(TDM), 저작권 제한·면책, 이단계 보상 모형(Two-Tier Compensation Model), 목적세

[Abstract]

Generative AI Copyright and a Two-Tier Compensation Model: Toward a New Cultural Social Contract for the 'Super-powered Literary Girl'

Yum, Kyuhyun

This study proposes a new approach to Text and Data Mining (TDM), which has been indiscriminately conducted in the training process of generative AI. It critically examines how existing discussions, based on copyright law and focused solely on one-time reproduction and use at the 'learning stage,' fail to adequately reflect the compensation structure for the commercialization process that occurs extensively after learning.

To this end, LLMs are likened to a 'super-powered literary girl,' emphasizing the need for a separate taxation-based regulatory framework for the 'super-power' possessed by generative AI. Reflecting the asymmetric characteristics of AI technology, this study distinguishes between 1) the training stage and 2) the utilization stage, proposing the introduction of a 'Two-Tier Compensation Model' that applies appropriate compensation based on copyright law for the training stage and a taxation system based on revenue for the utilization stage.

At Stage 1, TDM is regarded as either traditional copyright infringement or a form of legitimate appropriation, and a differentiated compensation system by data type is proposed. At Stage 2, noting that original data does not remain within the model, the study advocates for introducing a purpose tax in the form of contributions to the overall 'creative ecosystem,' rather than individual rights remedies under copyright law.

Furthermore, this study addresses potential controversies such as double taxation and innovation deterrence that may arise during policy implementation, and presents an implementation roadmap including international cooperation and domestic law amendments, thereby proposing a model for a new cultural social contract in the era of generative AI.

Keyword: Generative Artificial Intelligence, Text and Data Mining (TDM), Limitations and Exceptions to Copyright, Two-Tier Compensation Model, Purpose Tax