

03

우리는 어떻게 인공지능 기술에 대응해 나갈 것인가? 언론보도와 AI를 둘러싼 법·제도적 문제들

라기원 한국법제연구원 부연구위원



1. 언론·인공지능 기업이 맞닥뜨린 뉴스 학습 규범 전쟁

- 인공지능의 언론 학습데이터화에 대응한 분쟁 현황과 각국의 법제도 동향
- 언론 내부 차원의 인공지능 활용을 위한 지침, 기준 마련 동향
- 인공지능 기본사회에서 언론이 나아가야 할 방향

최근 우리나라에서는 인공지능 학습 과정에서의 저작물 무단 이용 문제가 본격적인 분쟁으로 이어지고 있다. 여러 언론 보도에 따르면, 네이버의 초대규모 언어모델 ‘하이퍼클로바X’가 학습 과정에서 언론사의 뉴스 콘텐츠를 무단 활용했다는 의혹이 제기되면서 공영방송사들이 법적 대응에 나섰다. 한국방송협회 보도자료(2025. 1. 13.)에 따르면 KBS, MBC, SBS 등 공중파 3사는 2024년 1월 13일 네이버와 네이버클라우드를 상대로 각 2억 원, 총 6억 원 규모의 손해배상 청구 소송과 인공지능 추가 학습 금지 청구 소송을 제기했다. 방송협회는 ‘네이버 AI가 학습한 자료 중 뉴스는 13.1%로, 전체 자료에서 상당한 비중을 차지한다고 지적했다.¹⁾ 방송협회가 문제 삼은 기사 비율은 네이버가 공개한 자료는 아니지만, 2021년에 발표된 인공지능 연구 논문에서 나온 수치와 동일하다. 이 논문에 따르면, 네이버가 처음 구축한 인공지능 학습 자료는 인터넷 글·댓글·블로그·뉴스 등 다양한 정보로 이루어져 있는데, 이 전체 자료 중 뉴스 기사가 약 13%를 차지하는 것으로 보고되어 있다. 이에 방송협회는 뉴스가 대량으로 학습에 쓰였다면 언론사의 콘텐츠가 인공지능에 흡수되는 셈이고, 그만큼 언론이 입는 경제적 손실도 크다고 주장하고 있다.²⁾

그리고 지난 11월 6일 KBS·MBC·SBS가 네이버·네이버클라우드를 상대로 제기한 ‘인공지능의 저작권 침해로 인한 손해배상 및 인공지능 추가 학습 금지’ 청구 소송 2차 변론기일이 열렸다. 재판부는 KBS·MBC·SBS(이하 ‘원고 측’)에 “저작권 침해를 주장한다면, 네이버·네이버클라우드(이하 ‘피고 측’)가 학습한 ‘개별 저작물’을 특정해야 한다”며 소송물 특정을 요구하였다. 이에 원고 측은 “기사 및 영상이 대량 데이터 덩치로 학습되므로 특정이 사실상 불가능하다”고 주장했으나, 재판부는 설령 피고 측 인공지능이 원고 측 저작물 수만 건을 학습했더라도 각 저작물에 대한 저작권

1) 신은빈. (2025. 10. 13). ‘뉴스 무단 학습’ 논란 네이버 AI...언론 단체 거액 소송 본격화. 뉴스 1. <https://www.news1.kr/it-science/general-it/5938666>

2) Kim, S., Lee, G., Kim, B., Chung, H. W., Kim, G., Kim, S., & Kim, N. (2021). Billions-scale Korean generative pretrained transformers. In Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing. Association for Computational Linguistics. EMNLP 2021, pp.8483–8501.

침해는 별개의 불법행위를 구성하므로, ‘손해배상청구권도 개별 저작물마다 별개의 소송물’로 판단할 수밖에 없다며 원고 측의 주장을 받아들이지 않았다.

피고 측 역시 원고 측 저작권 주장은 이유 없다고 반박하며, 원고 측이 문제 삼은 사건·사고 보도 상당수는 「저작권법」 제7조제5호에서 정한 바와 같이 ‘사실의 전달에 불과한 시사보도’로서 저작권법의 보호 대상에 해당하지 않기 때문에 애초에 저작권이 성립되지 않는다고 맞섰다. 또 2020년 원고 측과 피고 측 포털 송출 계약에서 ‘새로운 서비스 연구·개발 목적 기사 이용 가능’ 조항이 존재했음을 이유로, 생성형 인공지능 학습도 정당한 계약으로 인한 것이라고 주장했다.

피고 측 포털 송출 계약에 근거한 기사 이용권 주장은 「독점규제 및 공정거래에 관한 법률(이하 ‘공정거래법’)」상 언론사와 플랫폼 간 분쟁과도 연결된다. 한국신문협회는 지난 4월, 피고 측이 디지털 기사를 인공지능 학습용으로 무단 사용하였다는 혐의로 공정거래위원회에 신고했다. 한국신문협회 측 대리인은 “인공지능 학습 때문에 언론 트래픽이 급감하는데도, 포털의 압도적인 지위 때문에 언론사는 사실상 거래거부가 불가능한 상황”이라는 점을 지적했다.

이처럼 생성형 인공지능이 언론 콘텐츠를 대규모로 학습하고 재구성해 제공하는 시대가 되면서 언론계와 포털, 그리고 인공지능 기업 간 충돌이 본격화되고 있지만, 「저작권법」 및 「공정거래법」 등 기존의 법제로는 언론·인공지능 기업의 규범전쟁에 대응하기는 부족한 실정이다.

2. 우리나라의 입법 동향

가. 최근의 입법 동향

국회는 언론저작물 보호를 위한 입법에 시동을 걸고 있다. 정연욱 국민의힘 의원(부산 수영)은 신문·인터넷신문·뉴스통신 기사를 ‘언론저작물’로 명시하는 「저작권법 일부개정법률안」을 발의하며, “인공지능 시대에 언론의 창작물 가치를 법으로 분명히 해야 한다”고 강조했다.³⁾ 2026년 1월 시행을 앞두고 있는 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법(이하 ‘인공

3) 정연욱의원 대표발의, “저작권법 일부개정법률안”, 2212043, (2025. 8. 7.)

지능기본법」) 또한 개정안이 제안되었다. 2025년 6월 17일 발의된 「인공지능기본법 일부개정법률안」에는 생성형 인공지능 및 고영향 인공지능, 딥페이크 기술 등이 사용된 결과물의 경우 사용자에게 대하여 인공지능 기술이 사용되었다는 사전고지 의무와 결과물 표시(워터마크) 의무를 부과한 기존의 인공지능 투명성 조문(제31조)에 제4항을 신설하는 내용이 포함되었다. 신설된 제4항은 “저작물 등의 권리자가 자신의 저작물 등이 학습용 데이터로 이용되었는지 확인을 요청하는 경우, 인공지능사업자가 이를 확인할 수 있는 절차를 마련하여야 한다”는 내용으로, 일종의 ‘학습데이터 이용 확인권’을 도입한 것이다.⁴⁾ 그동안 저작권·콘텐츠 산업계가 지속적으로 제기해 온 쟁점은 ‘인공지능이 무엇을 학습했는지, 내 저작물이 포함되었는지 여부조차 알 수 없다면 권리 보호가 사실상 불가능하다’는 것이었다. 언론계 또한 인공지능 학습데이터의 13% 이상을 기사가 차지하고 있다는 연구결과에도 불구하고, 인공지능이 어떤 데이터를 학습했는지 정확히 알 수 없다는 문제를 안고 있었다. 인공지능의 학습 저작물 확인 요청권에 관한 법률안이 통과될 경우, 앞서 본 분쟁에서 언론이 인공지능 기업을 상대로 저작권 침해를 원인으로 한 손해배상 청구 대상 및 범위를 특정할 수 있는 근거로 작용할 수 있을 것으로 보인다.

나. 언론과 인공지능 문제를 해결하기 위한 입법 연혁

한편, 인공지능 시대에 대비하여 국회는 언론 분야의 법제도 개선을 활발히 추진해 왔다. 2016년 미국 대선에서 도널드 트럼프 대통령이 주요 언론을 향해 가짜뉴스(Fake News)라는 표현을 연일 사용하기 시작한 이후 국내 언론 또한 신뢰성의 위기를 맞았고, 동시에 SNS, 온라인 플랫폼에서 확산된 허위 정보들이 여론에 중대한 영향을 미칠 수 있다는 문제의식이 대두되었다. 이러한 흐름 속에 제22대 국회의원 선거를 앞둔 2023년 12월 19일, 허위정보 확산에 대응하기 위한 취지의 법률안이 국회에서 다수 발의되었고, 이들을 통합한 「공직선거법 일부개정법률안(대안)」이 제출 하루 만에 본회의를 통과하는 기록을 세웠다.⁵⁾ ‘선거 과정에서 딥페이크 기술이 활용될 경우, 유권자는 진실과 조작된 정보를 식별하기 어려워지고 그 결과

4) 김기현의원 대표발의, “인공지능 발전과 신뢰 기반 조성 등에 관한 기본법 일부개정법률안”, 2210903, (2025. 6. 17.)

5) 정치개혁특별위원회장, “공직선거법 일부개정법률안(대안)”, 2125972, (2023. 12. 19.)



허위·조작된 정보가 가짜뉴스의 형태로 급속히 확산될 수 있다는 우려에서 마련된 해당 개정안은 단순한 기술 규제를 넘어 인공지능 기반 조작·허위 콘텐츠의 유포 및 확산으로부터 선거의 공정성과 민주적 의사형성 과정을 보호하기 위한 예방적 규율로 이해될 수 있다. 이 외에도 가짜뉴스의 생성 가능성을 우려하여, 허위 정보나 타인의 의사에 반하여 사람의 얼굴·신체 또는 음성을 편집·가공한 영상물 또는 음성물의 유통을 금지하고, 위반 시 처벌하도록 하는 취지의 법률안도 제안되었지만, 임기만료로 폐기되었다.⁶⁾ 이는 이후 2024년 12월 26일 제정된 「인공지능기본법」의 제31조에 인공지능 사업자가 생성형 인공지능을 이용한 제품 또는 서비스를 제공하려는 경우 그 결과물이 생성형 인공지능에 의하여 생성되었다는 사실을 표현해야 한다는 취지의 조문으로 입법되었다.

공정성의 관점에서는 2020년 9월 11일, 언론이 생산한 기사를 제공 또는 매개하는 인터넷뉴스서비스사업자에게 기사배열의 기본방침과 책임자, 기사배열에 사용되는 프로그램 등 구체적인 기준을 공개하도록 하는 내용의 「신문 등의 진흥에 관한 법률 일부개정안」이 제안되었다.⁷⁾ 2020년경의 인공지능 기술은 생성형 인공지능 단계까지는 이르지 못하였지만, 인터넷 포털 사이트에서 기사 배열 알고리즘으로 활용되기 시작하였다. 포털의 기사 배열은 국민 여론 형성에 상당한 영향을 미치는 만큼, 알고리즘 기반 배열의

6) 박성중의원 대표발의, “정보통신망 이용촉진 및 정보보호 등에 관한 법률 일부개정법률안”, 2125438, (2023. 11. 15.)

7) 정희용의원 대표발의, “신문 등의 진흥에 관한 법률 일부개정법률안”, 2103807, (2020. 9. 11.)

중립성과 공정성에 대한 사회적 의문이 제기되었다. 이러한 문제의식 속에서 포털사업자가 뉴스 배열을 임의로 조작하지 못하도록 하는 취지의 법률안이 연이어 발의되었다. 특히 2021년 6월 15일에는 인터넷뉴스서비스사업자가 인공지능 알고리즘에 따라 기사를 배열하는 과정에서 발생할 수 있는 기사 노출의 불공정성에 대응하기 위해, 사업자의 자체적 판단이나 기준에 따른 기사 추천·편집 행위를 제한하려는 내용의 개정안들이 제출되었다.⁸⁾ 그러나 이러한 법안들은 모두 국회 임기 만료로 폐기되었다.

저작권의 관점에서 국회는 2020년 12월 21일, 우리나라 현행법상 인공지능이 생성한 저작물에 관한 명확한 규정이 존재하지 않는다는 점을 이유로, 인간의 창작활동을 보조하거나 스스로 저작물을 생산할 수 있는 인공지능이 제작한 저작물에 대한 저작권을 ‘인공지능 저작물의 저작자’라는 법적 근거를 마련함으로써 보장하고, 나아가 인공지능 저작물의 창작과 이에 대한 투자를 활성화하고자 했다.⁹⁾ 이는 언론, 기고, 영상제작물 등을 인공지능으로 제작하였을 경우 저작권의 소유권한 행사에 대한 근간이 될 수 있는 규정으로 기대받았으나, 인공지능이 생성한 결과물의 저작권 등록 가능성에 대한 명확한 법적 해석은 여전히 부재한 상황이다.

한편, 언론을 보호해야 한다는 관점과는 반대로 언론 등 저작물 활용이 인공지능 및 데이터 산업의 활성화에 필수적이라는 문제의식 하에 데이터마이닝을 정당화하는 취지의 법률안도 제안되었다.¹⁰⁾ 또 컴퓨터 기반 자동화 분석 기술을 활용하여 대량의 정보를 분석하고 새로운 정보나 가치를 생성하는 이른바 ‘데이터마이닝(Data Mining)’ 또는 ‘텍스트·데이터마이닝(Text and Data Mining, TDM)’을 허용하려는 취지의 법률안도 제안된 바 있다. 이러한 규정은 저작권자의 사전 허락 없이도 저작물에 대한 복제 및 전송을 가능하도록 하여, 데이터 산업의 발전을 촉진하려는 데에 목적이 있었다. 영국(2014년), 독일(2017년), 일본(2018년) 등은 이미 「저작권법」을 개정하여 정보분석을 위한 복제를 명시적으로 허용하는 제도를 도입하였고, 이 입법례를 참고하여 국내에서도 유사한 법안이 발의된 것이었다. 그러나 당시 입법 논의 과정에서 여러 우려가 제기되었다. 특히 독일과 일본은 영리 목적의 데이터마이닝까지 허용한 반면, 영국은 비상업적 연구

8) 김의겸의원 대표발의, “신문 등의 진흥에 관한 법률 일부개정법률안”, 2110802, (2021. 6. 15.)

9) 이용호의원 대표발의, “저작권법 일부개정법률안”, 2117990, (2024. 5. 29.)

10) 황보승희의원 대표발의, “저작권법 일부개정법률안”, 2122537, (2023. 6. 8.)

목적에 한정하고 있어, 우리나라가 영리 목적까지 포괄하는 데이터마이닝을 허용할 경우 권리자들의 반발이 커질 것이라는 지적도 제기되었다. 한국 언론진흥재단 또한 개정안이 도입될 경우 언론사의 뉴스저작권이 침해될 소지가 있다는 의견을 제시하였다. 이러한 논란 속에서 법안은 국회 임기 만료로 폐기되었다.

3. 미국, 유럽연합, 일본의 규범동향

가. 미국

미국은 언론사와 인공지능 기업 간 충돌이 가장 먼저 격렬하게 표면화된 국가이다. 2023년 12월 27일, 뉴욕타임스(The New York Times)는 오픈AI(OpenAI)와 마이크로소프트를 상대로 미국 뉴욕 남부연방법원에 저작권 침해 소송을 제기하면서, 두 회사가 구독자 전용 기사를 포함한 자사 기사들을 허락 없이 대규모로 수집해 GPT 계열 모델의 학습에 사용하였고, 그 결과 챗GPT가 뉴욕타임스 기사를 거의 그대로 재현하는 출력물을 여러 차례 생성했다고 주장하였다.¹¹⁾ 2024년에는 시카고 트리뷴(Chicago Tribune), 뉴욕 데일리뉴스(New York Daily News)를 포함한 미국 지역 신문사들이 유사한 취지의 집단 소송에 동참하며, 기업들이 AI의 학습을 위해 언론사의 핵심 자산을 무단으로 이용하고 있다는 문제의식이 본격적인 법정 분쟁의 형태로 확산되었다.¹²⁾

미국의 기업들은 소송과 여론의 압박을 계기로, 법적 리스크를 관리함과 동시에 고품질 뉴스 데이터의 안정적 확보를 위해 언론사와의 대규모 라이선스 계약에 적극 나서는 방향으로 전략을 선회했다. 구글(Google)은 월스트리트저널(The Wall Street Journal)의 모회사인 뉴스 코프(News Corp)와 인공지능 관련 콘텐츠 및 제품 개발을 위한 계약을 체결하였고,¹³⁾

11) The New York Times Company v. Microsoft Corporation and OpenAI Inc., Complaint, No. 1:23-cv-11195, U.S. District Court for the Southern District of New York, filed (2023, December 27), https://nytco-assets.nytimes.com/2023/12/NYT_Complaint_Dec2023.pdf

12) Cat Zakrzewski, (2023, December 28), New York Times sues OpenAI, Microsoft for using articles to train ChatGPT, The Washington Post, <https://www.washingtonpost.com/technology/2023/12/27/new-york-times-sues-openai-chatgpt/>

13) Laurie Sullivan, (2024, April 30), Google Reportedly Strikes Deal With News Corp. To Develop AI Content, MediaPost, <https://www.mediapost.com/publications/article/395651/google-reportedly-strikes-deal-with-news-corp-to-develop-ai-content.html?edition=>

오픈AI는 파이낸셜타임스(Financial Times)와의 전략적 파트너십 및 라이선스 계약을 통해 기사 이용 및 모델 학습에 뉴스 기사를 활용하는 방안을 공식화하였다.¹⁴⁾ 또한 오픈AI는 비즈니스 인사이더(Business Insider)와 폴리티코(Politico) 등을 보유한 독일 언론 그룹 악셀 스프링어(Axel Springer)와도 인공지능 학습 및 생성형 답변에 뉴스 콘텐츠를 활용하기 위한 장기 라이선스 계약을 체결하고, 유료 구독으로 접근이 제한되는 기사(paywall)에 대해서도 대가를 지급하기로 합의하였다.¹⁵⁾

위와 같은 일련의 계약은 단순히 뉴스 저작물에 대한 이용료를 지급하는 차원을 넘어, ‘뉴스 콘텐츠는 인공지능의 성능을 끌어올리는 가장 고급의 훈련 데이터’라는 인식을 전제로 한 거래라고 평가된다.¹⁶⁾ 미국에서는 입법이나 포괄적 규제체계가 먼저 정비되는 것이 아니라, 시장에서 뉴스 데이터에 대한 대가 지급 관행이 형성되고, 그 관행이 새로운 표준으로 부상하는 방식으로 질서가 세워지고 있다. 인공지능 기업은 소송과 규제의 불확실성을 줄이는 대신, 라이선스 비용을 지불하고 고품질 데이터를 확보하는 길을 택했고, 언론사는 자신의 뉴스가 인공지능 시대에 어떤 가치로 평가되며 얼마만큼의 보상을 받을 수 있는지를 협상 테이블에서 직접 다루기 시작했다.

나. 유럽연합

유럽연합은 생성형 인공지능과 플랫폼 기반 뉴스 소비가 민주주의와 언론 다양성에 미치는 파급효과를 단순한 기술 발전이 아니라 정치·사회적 질서를 재구성하는 문제로 받아들이고 있다. 유럽연합은 이른 시점부터 인공지능 기반 추천 시스템을 「디지털서비스법(Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC, Digital Services Act, ‘DSA’)」 안으로 끌어들여 규율하기 시작했다.

14) Financial Times. (2024, April 29). Financial Times announces strategic partnership with OpenAI [Press Release]. https://aboutus.ft.com/press_release/openai

15) OpenAI Official Blog. (2023, December 13). Partnership with Axel Springer to deepen beneficial use of AI in journalism. <https://openai.com/index/axel-springer-partnership/>

16) Diana Kwon, (2024, December 9). Publishers are selling papers to train AIs — and making millions of dollars. Nature. <https://www.nature.com/articles/d41586-024-04018-5>



그 핵심은 바로 「디지털서비스법」 제27조의 ‘추천 알고리즘 투명성 규정’이다. 이 조항에 따라 플랫폼은 뉴스를 노출할 때 어떤 데이터 신호(data signals)를 입력값으로 활용했는지, 그 신호들에 어느 정도의 내부 가중치(parameter)를 부여해 추천 순위를 결정했는지에 관한 정보를 공개하여, 알고리즘의 작동 원리를 일반 이용자가 이해할 수 있는 언어로 설명해야 한다. 여기서 말하는 ‘데이터 신호’란 조회수, 체류 시간, 이용자의 관심사·행동 패턴, 정치적 성향 추정치, 지역·언어 정보, 광고 최적화 신호 등 추천 시스템의 입력값을 의미하며, 가중치(parameter)는 이러한 신호가 알고리즘 내부에서 어떻게 조합되고 어떤 비율로 반영되는지를 결정하는 내부 설정값을 뜻한다. 유럽연합은 플랫폼이 이 두 가지를 모두 설명하도록 의무화함으로써, 뉴스 유통 과정에서 알고리즘의 작동 원리가 더 이상 기업의 블랙박스로 남지 않도록 하고 있다. 또한 이용자는 추천 시스템의 구조를 이해한 뒤, 원한다면 개인맞춤형이 아닌 일반 추천 방식으로 언제든지 전환할 권리를 가진다. 이는 플랫폼이 만들어낸 개인화 흐름에 이용자가 일방적으로 끌려가지 않도록 하기 위한 제도적 안전장치이다.

유럽연합은 이러한 장치를 통해 플랫폼의 추천 시스템을 사기업의 내부 알고리즘이 아니라 민주주의에 영향을 미치는 공적 인프라로 재정의하고, 여론 형성에 대한 책임을 플랫폼에게 법적으로 귀속시키고 있다. 특히 유럽 연합 국가의 월간 이용자 수가 4,500만 명이 넘는 매우 큰 온라인 플랫폼(Very Large Online Platforms, VLOPs), 예를 들어 유튜브, 위키피디아, 아마존닷컴, 알리익스프레스, 구글, 마이크로소프트 등의 경우, 위험평가

와 위험완화 조치 의무까지 부과되며, 제27조를 위반할 경우 전 세계 연간 총매출액의 최대 6%에 이르는 과징금과 시정명령은 물론, 필요하다면 유럽 연합 역내 서비스 제한까지 가능하다는 점에서 그 규범력이 상당히 강하다.

인공지능과 언론의 두 번째 규범은 「디지털 단일시장 저작권 지침(DSM Directive, 2019/790)」이다. 유럽연합은 생성형 인공지능이 전 세계적 화두가 되기 전인 2019년에 이미 텍스트·데이터 마이닝(TDM)을 둘러싼 저작권 규율을 준비해 두었다. 그 가운데 제4조는 인공지능 학습과 언론사의 권리 보호를 직접적으로 연결하는 핵심 조항으로 평가된다. 인터넷에 공개된 기사라 하더라도, 언론사가 로봇 배제 표준(robot.txt)이나 메타데이터를 통해 '이 콘텐츠는 TDM·인공지능 학습에 사용하지 말라'는 권리 유보 의사를 표시하면, 인공지능 기업은 그 순간부터 더 이상 해당 기사를 학습데이터로 사용할 수 없게 된다. 이 규범 구조 속에서 언론사는 단순한 데이터 제공자가 아니라, 자신이 생산한 뉴스가 인공지능 학습에 사용될지 여부를 결정할 수 있는 통제권과 협상력을 보장받게 된다.

2024년 8월 발효된 「EU 인공지능법(EU AI Act, Regulation (EU) 2024/1689)」은 처음으로 “인공지능이 어떤 데이터를 학습했고 어떤 결과물을 만들어내는지”를 모델 개발자와 제공자의 법적 책임으로 명확히 규정했다. 특히 「EU 인공지능법」 제53조는 범용 인공지능(General-Purpose AI) 제공자에게 「EU 저작권법」 준수를 위한 내부 정책 마련을 요구하고, 학습에 사용된 콘텐츠의 종류와 특성에 대한 공개 및 EU 인공지능청(EU AI Office)에의 제출 의무를 부과하였으며, 콘텐츠 제공자가 학습을 허락했는지 또는 금지했는지 식별하고 이를 따르기 위한 절차 구축을 의무화하고 있다. 이 조항으로 인해 인공지능 학습데이터는 더 이상 “기술이 알아서 처리하는 입력값”이 아니라, 투명성과 책임을 요구받는 규제 대상으로 자리 잡게 되었다. 다만 「EU 인공지능법」의 본문 규정은 주로 기본 원칙과 규율 구조를 제시하는 수준에 머물러 있어, 이러한 의무가 실제 기술 설계나 조직 운영에서 어떻게 구현되어야 하는지에 대해서는 추가적인 하위 지침과 실무 기준이 마련될 필요가 있다.

유럽연합은 2025년 7월 10일 범용 목적 인공지능 제공자를 대상으로 하는 「실무준칙(Code of Practice)」을 마련했다.¹⁷⁾ 이 준칙은 2025년 8월 2

17) European Commission, (2025, July 10). The General-Purpose AI Code of Practice.

일부러 적용되는 「EU 인공지능법」 규정 중 투명성(transparency), 저작권(copyright) 의무를 준수하기 위한 실질적인 지침으로 설계되었다. 「실무준칙」은 「EU 인공지능법」과 「디지털 단일시장 저작권 지침」이 선언한 추상적 규범을 실제 학습데이터 수집·책임 배분에 적용하는 역할을 수행하며, 법에 준하는 준헌법적 기술 규범이라는 평가를 받고 있다.¹⁸⁾ 2025년 8월 1일 「실무준칙」 서명 참여 기업 목록이 공개된 이후 2025년 11월 현재까지 40여 개 이상의 글로벌 AI 제공자가 이 코드에 서명하면서 해당 지침은 사실상 업계 표준으로 자리 잡는 흐름을 보이고 있다.¹⁹⁾ 요컨대, 유럽연합은 「실무준칙」의 적용을 통해 모든 범용 인공지능 모델 제공자는 자사의 모든 범용 인공지능 모델에 대해 「EU 저작권법」 준수를 위한 내부 정책을 문서화해야 한다는 점, 권리자의 허락 또는 거부 의사를 식별하고 준수하는 절차를 갖추어야 한다는 점, 그리고 이러한 정책을 집행·감독할 책임자를 반드시 지정해야 한다는 점을 규정하였다. 또한 저작권 준수 여부를 대외적으로 공개할 것을 권고하고 있다.

다. 일본

일본은 2018년 「저작권법(著作権法)」 개정을 통해, 텍스트·데이터 마이닝(TDM)에 관한 한 세계에서 가장 과감한 예외 규정을 도입한 국가로 꼽힌다. 2018년 5월 법률 제30호로 도입된 저작권법 제30조의4는 ‘정보 분석을 위한 이용에 대해서는 상업적 목적과 비상업적 목적을 구별하지 않고 폭넓게 허용하는 내용을 담고 있다.’²⁰⁾ 이 규정의 도입으로 인해, 연구 목적의 TDM뿐 아니라 상업적 인공지능 학습의 상당 부분도 원칙적으로는 허용된다는 해석이 일반화되었고, 실제 일본 지적재산전략본부는 2019년 보고서에서 일본을 AI 학습에 관한 한 세계에서 가장 이용범위가 넓은 국가 중 하나라고 평가하였다.²¹⁾ 이처럼 일본은 한때 AI 학습하기 좋은 나라, 인공지능 학습의 천국이라고 불릴 만큼, 데이터 분석과 학습에 필요한 저작물

18) Nicola Kerr-Shaw, David A. Simon, Stuart D. Levi, Susanne Werry & Aleksander J. Aleksiev, (2025, August 14), EEU's General-Purpose AI Obligations Are Now in Force, With New Guidance, Skadden, <https://www.skadden.com/insights/publications/2025/08/eus-general-purpose-ai-obligations>

19) European Commission, (2025, July 10), General-Purpose AI Code of Practice now available [Press Release], IP/25/1787.

20) 日本国, 「著作権法」, (文化庁) 第30条の4(情報解析のための利用).

21) 内閣府 知的財産戦略本部, 「AI・データ時代における知財制度の在り方」, (2019).

이용을 넓게 허용하는 방향으로 저작권 법제를 개정해왔다. 하지만 Kenji Tosaki, Takahiro Hatori 그리고 Tomoki Abe(2023)에 따르면 이와 같은 광범위한 TDM 예외는 주로 머신러닝을 위한 데이터 분석 단계에 한정되는 것으로, 이미지 생성과 같이 학습·생성 전 과정에서 저작물을 광범위하게 활용하는 생성형 AI 기술에 그대로 적용하기 어렵다. 즉, 일본의 저작권 법제는 생성형 AI까지 완전히 포섭하는 구조가 아니므로 생성형 AI에서는 여전히 저작권 침해 위험이 존재한다.²²⁾

최근 들어 이러한 문제의식이 일본 내부에서도 점차 확산되고 있다. 일본 신문협회(日本新聞協會), 일본민영방송연맹(日本民間放送連盟), 일본작가회의(日本作家會議), 만화가 단체 등 다양한 창작자·언론 단체는, TDM이 폭넓게 허용된다 하더라도 언론 기사나 고부가가치 창작물을 제한 없이 학습 데이터로 활용하는 것에는 정당한 보상이 따라야 한다는 문제를 제기하기 시작했다. 특히 언론계는 인공지능이 학습 과정에서 활용하는 뉴스 데이터가 단순한 텍스트 축적물이 아니라, 언론사가 수십 년간 축적해 온 브랜드 가치, 취재 역량, 편집의 노하우 등 무형의 편집·보도 자산을 체계적으로 흡수하는 과정이라는 점을 강조하며, 이에 상응하는 보상·협상 구조의 제도화가 필요하다는 문제의식을 제기하고 있다.²³⁾

이런 가운데 2025년 8월 26일, 일본의 주요 일간지인 일본경제신문(日本經濟新聞)과 아사히신문(朝日新聞)은 인공지능 기반 검색엔진을 운영하는 Perplexity AI를 상대로 저작권 침해 소송을 제기했다.²⁴⁾ 두 신문사는 Perplexity AI가 유료 기사(paywalled articles)를 포함한 자사 콘텐츠를 무단 크롤링하여 학습데이터로 활용했다고 주장하며, 이는 일본 저작권법의 허용 범위를 벗어난 불법적 이용이라고 지적하였다. 이는 2025년 8월 7일 일본 최대 발행부수를 가진 요미우리신문(讀賣新聞)이 Perplexity AI를 상대로 소송을 제기한 사건과도 연결된다. 이 사건을 계기로 일본 언론계 전반에서 “AI 기업의 뉴스 데이터 무단 학습은 언론사의 핵심 자산을

22) Kenji Tosaki, Takahiro Hatori & Tomoki Abe. (2023, July). Japan a Paradise for Machine Learning, Not Generative AI. 知的財産法 ニュースレター. Nagashima Ohno & Tsunematsu. https://www.nagashima.com/wp-content/uploads/2023/07/ip_en_no3.pdf

23) 日本民間放送連盟. (2024). 「生成AIの報道利用に関する検討」.

24) Nikkei staff writers. (2025, August 26). Japanese newspapers Nikkei and Asahi sue Perplexity AI over copyright. Nikkei Asia. <https://asia.nikkei.com/business/technology/artificial-intelligence/japanese-newspapers-nikkei-and-asahi-sue-perplexity-ai-over-copyright>

침해한다”는 문제의식이 공고해지고 있다.²⁵⁾

라. 소결

앞서 살펴본 세 나라의 접근방식을 비교해 보면, 각 국가가 서로 다른 규범 설계 철학을 어떻게 구현하고 있는지 뚜렷하게 드러난다. 미국은 우선 소송과 시장을 통해 인공지능 기업이 데이터 제공자인 언론에게 대가를 지불해야 한다는 구조를 현실에서 먼저 형성한 이후 규범화나 입법을 검토하고 있다. 유럽연합은 「디지털서비스법」, 「디지털 단일시장 저작권 지침」, 「EU 인공지능법」과 「실무준칙」에 이르기까지, 학습·생성·유통 전 단계에 걸쳐 규범, 보상 구조, 투명성 의무, 기술적 준칙을 모두 함께 설계하는 방식을 취함으로써, 인공지능과 언론의 관계를 단순한 계약관계나 개별 소송의 문제를 넘어 종합적인 법제도 체계로 규율하고 있다. 일본은 저작권법상 TDM 예외를 폭넓게 인정하여 인공지능의 분석과 학습을 원칙적으로 허용하였으나, 언론 데이터와 고품질 창작물에 대해서는 별도의 보상과 규율을 논의해야 한다는 방향으로 사회적 논의가 이동하고 있다.

4. 언론의 인공지능 활용 규범화 동향

오늘날 인공지능은 언론을 위협하는 존재이지만, 동시에 언론이 스스로 인공지능을 활용하여 언론의 일부 역할을 대체하려는 시도도 활발히 진행되고 있다. 2025년 3월, BBC는 기존 뉴스룸 구조를 재편하며 BBC 뉴스 성장·혁신·인공지능(BBC News Growth, Innovation and AI) 부서를 신설한다고 발표했다. 이 부서는 인공지능 기술을 활용해 사용자 개인화 뉴스 콘텐츠를 제공하려는 목적으로 신설되었다. BBC는 젊은 층의 소셜미디어 중심 뉴스 소비 패턴 변화에 대응하기 위해 AI 기반 맞춤형 뉴스 제공과 콘텐츠 추천, 사용자 경험 개인화 등을 추진할 예정이라고 밝혔다.²⁶⁾ 나아가 BBC는 생성형 인공지능을 뉴스 제작에 활용하며, 긴 기사의 핵심을

²⁵⁾ Nikkei staff writer, (2025, August 7). Japanese newspaper Yomiuri sues Perplexity AI over copyright, Nikkei Asia. <https://asia.nikkei.com/business/technology/artificial-intelligence/japanese-newspaper-yomiuri-sues-perplexity-ai-over-copyright>

²⁶⁾ Michael Savage, (2025, March 6). BBC News is to create a new department that will use AI to give the public more personalised content, The Guardian. <https://www.theguardian.com/media/2025/mar/06/bbc-news-ai-artificial-intelligence-department-personalised-content>

요약해 주는 ‘한눈에 보기(At a glance)’기능, 지역 뉴스 기사를 BBC 고유 스타일로 재편집해주는 ‘BBC 스타일 어시스트(BBC Style Assist)’ 기능을 도입한다고 밝혔다.²⁷⁾ 우리나라 방송 및 기사에서도 인공지능이 편집·제작한 결과물을 볼 수 있는 경우가 흔해지고 있다.

이처럼 언론의 인공지능 활용이 빠르게 확산되는 상황에서, 인공지능 기업이 언론 저작물을 이용하는 문제와는 별개로 언론이 인공지능을 활용하는 방식 자체도 새로운 규범적 쟁점으로 부상하고 있다. 유럽공영방송 연합(European Broadcasting Union, EBU)은 2025년 발표한 국제조사 보고서에서, 생성형 AI의 등장으로 뉴스 제작·편집·배포의 구조가 근본적으로 변화하고 있음에도 불구하고, 편집권(editorial integrity)과 뉴스 신뢰성은 인간이 책임져야 한다는 원칙을 분명히 했다.²⁸⁾ 해당 보고서는 ChatGPT, Gemini, Copilot, Perplexity 등 주요 AI 어시스턴트가 생성한 뉴스 응답의 45%가 사실 오류·출처 오류·맥락 왜곡을 포함한다는 조사 결과를 제시하며, 신뢰할 수 있는 정보 공급자로서 공영방송의 역할이 더욱 중요해졌다고 강조했다. 이에 따라, EBU는 투명성, 인간 중심 편집권, 인공지능 윤리 기준, 저작권 보호를 공영방송의 핵심 원칙으로 제시한다. 또 인공지능이 뉴스 추천, 요약, 자동 생성에 관여할 수는 있지만, 최종적으로 어떤 뉴스를 어떻게 배치하고 어떤 맥락으로 제공할 것인지는 인간 편집자가 책임져야 한다는 점을 강조하였다.²⁹⁾ 즉, AI가 뉴스 업무 전 과정에 걸쳐 보조적 역할을 수행할 수 있지만, 공영방송과 주요 언론사들은 일관되게 인간 책임 원칙을 핵심 가치로 두고 있음을 보여준다는 점에서 의미가 있다.

한편, 유럽의 주요 언론사들은 AI가 취재·기사 작성·편집 과정에까지 깊숙이 개입하는 현실에 대응하기 위해, 내부 윤리 강령과 업무 규정을 재정비하는 흐름을 보이고 있다. 이 과정에서 자주 언급되는 주체가 바로 공영방송(Public Service Media, PSM)이다. BBC·ARD·ZDF·France Télévisions과 같은 공영 언론기관은 국가로부터 공적 책임을 부여받은 만큼, 상업 매체보다 공공성·중립성·책임성이 더 강하게 요구되며, 이에 따라

27) Rhodri Talfan Davies. (2025, June 27). BBC to launch new Generative AI pilots to support news production, BBC Media Center, <https://www.bbc.co.uk/mediacentre/2025/articles/bbc-to-launch-new-generative-ai-pilots-to-support-news-production/>

28) EBU·BBC. (2025, October). News Integrity in AI Assistants: An International PSM Study.

29) Ibid.; EBU. (2025). News Integrity in AI Assistants Toolkit.

가장 먼저 규범 정비 논의를 시작하고 있다.

최근 학계와 언론계에서 비공식적으로 언급되는 이른바 ‘Article 12 가이드라인’ 또한 이러한 논의의 연장선에 있다. ‘Article 12’는 유럽연합의 특정 법률 조항을 의미하는 것이 아니라, 여러 공영방송과 해외 언론사들이 내부 윤리강령을 개정하는 과정에서 공통적으로 포함시키는 조항(예: 편집 책임, 투명성 의무 등)을 학계가 편의적으로 부르는 표현으로, 공식 규범은 아니지만 다수의 해외 언론사들이 논의하는 내부 윤리기준을 총칭하는 표현으로 사용된다.³⁰⁾ 이 가이드라인이 강조하는 핵심은 명확하다. 첫 번째는 투명성 원칙으로, 기자가 취재나 기사 작성에 AI 도구를 사용했다면, 그 사실과 사용 범위를 편집국과 독자에게 투명하게 공개해야 한다. 두 번째는 인간의 책임 원칙으로, AI가 리서치·요약·번역 등 보조적 역할을 수행할 수는 있으나, 최종 편집 판단과 책임은 인간 편집자에게 있다는 점을 명확히 해야 한다. 세 번째는 고지의무로, 자동 생성 기사나 개인 맞춤형 뉴스가 제공되는 경우, AI의 개입 여부를 독자가 명확히 인식할 수 있도록 표시해야 한다. 네 번째는 출처표시의무로 텍스트·이미지 등 AI가 생성한 콘텐츠를 기사에 포함할 때에는 가능한 범위 내에서 생성 방식과 원출처를 밝혀야 한다.

이러한 원칙은 단순한 도덕적 권고가 아니라, AI가 언론의 일상적 업무에 깊이 들어온 시대에 언론사가 스스로 마련해야 하는 최소한의 자기 규율(self-regulation)이다. 다시 말해 언론이 어떻게 기술을 책임 있게 통제하고, 그 과정에서 신뢰·편집권·투명성을 지켜낼 방법은 무엇인가라는 질문에 답하기 위한 최소한의 자기 규율이 형성되고 있는 것이다.

5. 재편되는 언론 규범: AI가 바꾼 전선

2025년 현재, 인공지능과 언론의 충돌은 더 이상 미래의 가능성이 아닌 현실의 문제로 뚜렷하게 부상하고 있다. 근본적인 문제는 인공지능이 언론을 학습하는 순간, 언론과 법·제도는 더 이상 과거의 규범 위에 머무를 수 없다는 것이다. 뉴스는 더 이상 발행 후 소비되고 사라지는 것이 아니라, 인공지능 모델 학습을 통해 무한히 재조합·재생산되는 구조적 자원이 된다.

30) Tomas Dodds, Wang Ngai Yeung & Claudia Mellado, (2025), On Controlled Change: Generative AI's Impact on Professional Authority in Journalism, arXiv preprint, <https://arxiv.org/abs/2510.19792>

그렇기에 언론과 인공지능 기업 사이의 갈등은 저작권 분쟁을 넘어 데이터 통제권·보상 구조·책임 귀속을 둘러싼 규범 전쟁으로 비화하고 있다.

우리나라에서는 네이버와 방송 3사의 소송과 학습데이터 확인권 도입 문제가 쟁점이 되고 있으며, 미국에서는 인공지능 기업과 미디어 기업 간 프리미엄 뉴스 라이선스 계약 활성화가 화두로 떠올랐다. EU는 이미 「디지털 서비스법」, 「디지털 단일시장 저작권 지침」, 「EU 인공지능법」 등 여러 법과 지침을 통해 인공지능과 뉴스의 관계를 규율하는 체계를 갖추는 중이다. 일본은 한때 AI 학습을 거의 전면 허용하던 법제를 마련하였으나, 최근 언론계의 강한 반발 이후 학습 허용 범위를 다시 조정하려는 흐름이 나타나고 있다. 이처럼 인공지능 기업과 언론 간 분쟁의 양상은 국가별로 차이가 있지만, 인공지능 학습에 있어 뉴스는 ‘고급 데이터이자 자원’이며 뉴스를 사용하는 방식에 새로운 규칙과 협상이 필요하다는 인식에는 공통점을 보이고 있다.

동시에 인공지능은 언론이 감시하고 견제해야 할 대상이면서도, 언론이 스스로 활용하지 않으면 뒤처지게 되는 필수 도구이기도 하다. BBC가 인공지능만을 다루는 전담 부서를 만든 일, 유럽공영방송연합(EBU)이 최종 편집과 책임은 인간에게 있어야 한다고 밝힌 원칙, 유럽 주요 언론사들이 AI 사용 기준을 내부 규정으로 정비하는 흐름은 ‘언론은 어떻게 인공지능을 활용하면서도 저널리즘의 윤리와 신뢰를 지켜낼 것인가?’라는 질문에 대한 해답을 모색하는 과정이다.

이제 각국 정부, 기업, 언론은 인공지능 시대라는 거대한 전환기 속에서 서로의 권리와 책임이 충돌하는 새로운 전장에 서 있다. 그러나 이 갈등은 단순한 힘겨루기로 끝날 수 없다. 데이터·저작권·책임·투명성을 놓고 충돌하던 이해관계자들은 결국 한 테이블에 마주 앉아야 한다. 인공지능 시대의 정보 질서는 전면전이 아니라, 각국 정부·기업·언론이 함께 체결해야 할 ‘언론과 인공지능의 평화협정’이라는 새로운 형태로 완성될 수밖에 없다. 서로의 권리와 책임을 다시 배치하는 조정의 과정이, 곧 미래 공론장이 될 우리 사회의 안정과 회복을 좌우할 것이다. 