
인공지능 윤리와 저작권의 규제체계 연구

신용우(연구책임자)

2024. 9. 28.



국회입법조사처
NATIONAL ASSEMBLY RESEARCH SERVICE

인공지능 윤리와 저작권의 규제체계 연구

2024. 9. 28.

연구책임자 : 신용우 [법무법인(유한) 지평]

참여연구원 : 최정규 [법무법인(유한) 지평]

허 중 [법무법인(유한) 지평]

이혜온 [법무법인(유한) 지평]

이재명 [법무법인(유한) 지평]

김재훈 [법무법인(유한) 지평]

전준우 [법무법인(유한) 지평]

조영록 [법무법인(유한) 지평]

법무법인(유한) 지평

이 보고서는 국회입법조사처의 정책연구용역사업으로 수행된 것으로서,
보고서의 내용은 연구용역사업을 수행한 연구자의 의견이며,
국회입법조사처의 공식 견해가 아님을 밝힙니다.

제 출 문

국회입법조사처장 귀하

본 보고서를 「인공지능 윤리와 저작권의 규제체계 연구」 연구용역의 최종
보고서로 제출합니다.

2024년 9월

신용우[법무법인(유) 지평]

요 약

1. 인공지능 규제 연구의 방향

- ☐ 인공지능(AI)은 다양한 분야의 혁신을 주도하고 있지만, 윤리적 문제가 대두되며 법적 규제의 필요성이 제기되고 있음
- ☐ 인공지능 기술 개발과 활용 전반에 공정성, 투명성, 책임성 등 윤리적 가치를 준수해야 한다는 주장이 제기되었으며, 특히 허위조작정보로 인한 피해는 당장의 현실적 문제가 되고 있음
- ☐ 생성형 인공지능의 활용이 확대되면서 다양한 저작권 이슈가 제기되고 있고, 특히 저작물 학습 시 저작권 면책 여부가 문제되며 다수의 저작권 분쟁이 진행 중임
- ☐ 본 연구는 인공지능 기술로 인한 윤리와 저작권 분야의 규제체계 마련을 위한 국내외 동향을 분석하고 입법 방향을 제시함

2. 인공지능 윤리 규제체계 현황 및 입법 방향

- ☐ 인공지능 윤리 원칙의 규범력과 강제력 확보를 위해 정책이나 입법이 필요하며, 본 연구는 인공지능 윤리와 관련하여 법으로써 규율할 수 있는 규제체계에 중점을 두었음

- 최근 발효된 EU 인공지능법은 세계 최초의 통합적이고 옴니버스 형태를 가진 인공지능 규제로서 주목하여 살펴볼 필요가 있음
 - EU 인공지능법은 위험 기반 접근법을 채택하였으며, 인공지능 기술 개발 및 활용과 관련한 이해관계자를 제공자, 배포자, 수입업자 및 유통업자로 구분하고 역할별로 의무사항을 상세히 규정하였음
 - 범용 인공지능 기반 모델 규제, 강한 행정규제, EU와 회원국 간 다층적 거버넌스, 혁신지원 제도 신설 등도 특징임
 - EU의 통합적인 접근방식은 전 분야에 걸쳐 인공지능의 위험성을 사전에 예방한다는 장점이 있지만, 여러 분야를 한데 묶어 규율하다 보니 법률이 복잡하고 아직 명확하지 않은 규제 영역이 많아 향후 법 시행 상황을 지켜볼 필요가 있음
 - EU 인공지능법을 무조건적으로 수용하기보다는 그 특징과 장점을 잘 살리되 미국 등 다른 국가의 사례와 우리나라의 특수성 등을 고려하여 우리나라의 고유한 제도로 발전시키는 방안을 모색할 필요가 있음
- 미국은 연방정부 차원의 포괄적인 규제 법률은 제정되어 있지 않지만, 행정명령을 통해 공공분야에서부터 인공지능 정책을 추진하고 있으며, 캘리포니아 주(州), 콜로라도 주 등 개별 주에서 상당한 강한 규제가 담긴 법률이 제정되고 있음
- 규제방식은 EU 방식의 수평적 규제와 미국 방식의 수직적 규제로 구분할 수 있으며, 유연하게 혼합하여 적용하는 방안을 고려할 필요가 있음
 - EU와 캐나다는 상대적으로 중앙집권적인 집행이 가능한 수평적 규제, 즉 대부분의 분야와 애플리케이션에 걸쳐 적용되도록 설계된 규제를 도입하는 접근 방식을 취함

- 미국, 영국, 중국은 상대적으로 분산된 접근 방식을 취하고 있으며, 여러 규제 기관에서 시행하는 수직적 규제, 즉 인공지능의 특정 애플리케이션이나 특정 부문에만 적용되는 규제에 대한 의존도가 더 높음
 - 미국의 접근 방식은 구속력이 없는 원칙, 위험 관리에 대한 자발적 지침, 연방 차원에서 새로운 인공지능 관련 법률을 개발하기보다는 기존의 부문별 법률을 적용하는 것이 특징임
 - EU와 미국 방식 모두 장단점을 갖고 있으며, 어느 한 방식을 규범적으로 선택하여 적용하기보다는, 국민의 권리를 보호하면서도 인공지능 기술·산업 발전을 촉진할 수 있는 방안을 유연하게 혼합하여 적용하는 방안을 고려할 필요가 있음
- 허위조작정보(딥페이크)는 최근 사회적 문제로 대두되고 있으며, 허위 영상물에 대한 금지는 강하게 추진하고, 인공지능 생성물 표시 의무는 현실적이고 합리적인 방식으로 설계할 필요가 있음
- 국내외에서 허위 영상물을 금지하거나 인공지능이 생성한 영상물·이미지라는 점을 표시하도록 하는 규제가 입법화되거나 추진 중에 있음
 - 허위조작정보로 인한 피해가 현실에서 발생하고 있으므로, 분야별 규제가 필요하며, 그 외에 다양하게 발생하는 문제점들을 효과적으로 해결하기 위하여 인공지능 기본법 또는 특별법에서 이를 예방·방지하는 시책을 추진하도록 하는 내용을 담을 수 있음
 - 인공지능으로 만든 가상 정보에 대하여 표시를 하도록 하는 의무는 허위조작정보 금지뿐만 아니라 저작권 침해 예방에도 효과적일 것으로 전망되며, 인공지능 전반을 규율하는 기본법에 규정할 수 있음

- 그러나 가상 정보 표시에 대한 기술적 회피·조작 가능성, 실제 사실에 대한 거짓 표시 가능성 등이 있으므로, 현 시점에서는 적용 콘텐츠나 수범자를 좁히고 미준수 시 제재는 최소한으로 하되, 관련 기술 발전 상황을 지켜보면서 규제를 넓혀가는 방안을 고려할 필요가 있음
- 인공지능 윤리 규제체계 입법 방향으로서 EU와 같은 옴니버스식 규제 도입을 고려하되 규제 수준은 낮추고, 규제 집행의 실효성 확보와 인공지능 리스크 책임 분배를 위해 컨트롤타워 정립이 필요함
- EU의 수평적·하향식 규제체계의 장점을 도입하되, 규제의 대상이 되는 분야를 한정적으로 열거하고 사업자에게 최소한의 의무만을 부여하며, 낮은 수준의 제재를 부과하는 방안을 고려할 수 있음
- 이후 인공지능 기술의 발전상을 따라가며, 인공지능으로 인하여 피해를 입은 구체적 사례를 분석하고 이를 축적하여 적절한 규제 범위와 제재 수준을 정해가는 것이 바람직함
- 위험 기반 규제 방식 도입을 고려하되, 현 단계에서는 공공분야를 중심으로 국민의 권리·의무에 상당한 영향을 줄 수 있는 분야를 엄선하고, 사업자 부담을 줄이는 것이 적절함
- 생성형(범용) 인공지능의 범용성과 고성능을 고려하여 별도 규제를 통해 위험을 저감할 필요가 있으며, 다만, 생성형 인공지능의 연산능력에 따른 기준을 법률에서 정하는 것은 적절하지 않음
- 제공자·배포자 등 인공지능 생태계에서 활동하는 각 주체별로 그에 적합한 의무를 부과하는 방안을 고려할 수 있음
- 인공지능 관련 정책과 감독을 전담하는 독립적인 기구는 업무수행의 독립성, 공정성, 전문성을 확보해야 하고 각 부처별, 인공지능 관련 기구들과의 원활한 협력관계를 구축해야 함

- 정책을 자문하는 전문가 집단의 구성을 고려할 필요가 있으며, 단순 자문을 넘어 실질적인 영향력과 구속력을 가질 수 있는 체계를 고려할 필요가 있음

3. 인공지능 저작권 규제체계 현황 및 입법 방향

- 본 연구에서는 인공지능이 저작물을 학습하는 경우 저작권이 면책되는지 여부에 관하여 중점적으로 살펴봄
 - 인공지능 자체에 저작권을 부여하기 어려우며, 사람인 저작자가 인공지능을 도구로서 사용한 경우에는 기존의 법리에 따라 저작권 인정이 가능함
 - 인공지능이 생성한 결과물이 다른 저작권을 침해하는지 여부는 기존의 저작권 침해 법리를 따르면 됨
 - 생성형 인공지능 기반 모델이 저작물을 학습할 때 저작권 침해가 면책될 수 있는지, 즉 저작재산권이 제한될 수 있는지가 주요 쟁점이며, 해외에서 다수 소송이 진행 중임
- 미국은 포괄적 공정이용 조항을 두고 있으며, 생성형 인공지능의 학습시 공정이용이 인정될지 여부는 명확하지 않음
 - 인공지능 학습은 비표현적이고 상업적 목적이 없다고 해석될 수 있으며, 저작권이 있는 자료를 인공지능 훈련에 사용하는 것은 혁신성을 인정받아 공정이용으로 판단될 가능성이 있음
 - 그러나 생성형 인공지능 모델은 상업적 성격이 충분히 인정될 수 있으며, 인공지능 생성물이 학습데이터인 저작물에 대한 수요를 대체할 수 있다고 볼 여지가 있음

- 최근 미국 연방대법원 판결에 따르면, 생성형 인공지능 결과물의 목적이 원저작물과 동일하다면 공정이용 가능성이 낮아질 것이며, 반대로 인공지능 개발자는 인공지능 결과물과 원저작물 간 목적이 실질적으로 다르다는 점을 입증하면 공정이용이 인정될 여지가 있을 것임

□ EU, 영국, 일본, 독일, 싱가포르 등은 TDM(Text-Data Mining) 면책 규정을 도입하였으며, TDM 생태계 내 다양한 이해관계를 조정하기 위해 각기 다른 접근 방식을 취하고 있음

- EU의 ‘디지털 단일시장(DSM) 저작권 지침’(DSM 지침)은 학술연구 목적의 TDM과 일반적인 TDM을 구분하여 이원적으로 규율하고 있으며, 학술연구 목적의 TDM에 대해서는 계약이나 기술적 수단으로 면책을 배제하지 못하도록 강제하는 반면, 일반 TDM에 대해서는 옵트아웃(opt-out, 데이터 수집 등에 대한 거부의를 표시할 경우 정보수집 금지)을 허용함으로써 균형을 맞추고 있음
- 영국과 싱가포르는 상업적/비상업적 목적의 구분 없이 옵트아웃을 허용하지 않는 강력한 규정을 도입함으로써 EU보다 더 넓은 범위의 면책을 제공하고 있음
- 일본은 EU DSM 지침과 유사하게 TDM 면책 규정을 이원화하여 정보해석 그 자체를 위한 비향수적 이용과, 정보해석에 부수하는 경미한 이용을 구분하고 있음
- 우리나라에서는 2021년 TDM 면책 규정을 저작권법에 도입하는 내용의 저작권법 전부개정법률안이 발의되었지만 저작권자들의 반대 등으로 통과되지 못하고 임기만료로 폐기되었음

- 인공지능 저작권 규제체계 입법 방향과 관련하여, TDM 면책 규정을 도입하되, 저작권자의 옵트아웃 권리, 저작물에 대한 적법한 접근의 범위, 학습데이터 출처 표시 방법, TDM 면책 시 보상 의무 등을 종합적으로 고려해야 함
 - 포괄적 공정이용 조항만으로는 수범자들의 예측가능성이 저해될 수 있으며 TDM 면책 조항과 그 입법 취지와 요건도 다르므로, TDM 면책 조항을 별도로 도입하는 것이 적절함
 - 생성형 인공지능 개발자(제공자)·배포자는 통상 영리 목적으로 개발·이용하며, 영리 목적과 비영리 목적의 구분이 불분명할 수 있으므로, 영리 목적인 경우에도 TDM 면책을 인정하되, 저작권자의 다른 권리 강화를 통해 균형을 도모할 수 있음
 - 저작권자에 대한 옵트아웃 권리 부여를 고려할 수 있지만 인공지능이 학습할 수 있는 저작물의 범위가 상당히 제한될 수 있으므로, 옵트아웃 표시 방법을 제한하는 방안, 저작권자가 옵트아웃을 하더라도 정당한 보상을 하면 면책되도록 하는 방안 등을 함께 고려할 수 있음
 - TDM 면책은 저작재산권을 제한하는 규정으로서 저작권자에게 불측의 손해가 발생할 우려가 있는 점을 고려하면 복제·전송 외에 그 범위를 확대하는 것은 바람직하지 않을 수 있음
 - 저작물에 대한 적법한 접근이라고 볼 수 있는지 여부는 옵트아웃 등 저작권자의 권리를 함께 고려하여 유연하게 제도를 설계할 필요가 있음
 - 학습데이터 출처 표시와 관련하여, 인공지능이 학습한 데이터의 출처에 대한 요약을 표시하도록 하는 방안을 고려할 수 있지만, 어느 정도로 상세하게 표시해야 할지 이해관계자의 사정을 종합적으로 고려해야 함

- TDM 면책 시 보상 의무를 입법화할 것인지 여부에 대하여, 주요국 현황을 고려할 필요가 있으며, 옵트아웃 권리 인정 여부, 영리 목적 이용 시 면책 여부 등 여러 쟁점들을 종합적으로 고려해야 함

4. 시사점

- 본 연구는 인공지능 기술과 산업 발전을 견인하면서 기존 이해관계자와 국민의 권리를 보호하는 규제체계 마련에 도움이 될 것으로 기대함
- 인공지능이 갖고 있는 내재적인 불확실성·오류와 인공지능의 강력한 성능에 대한 오·남용을 최소화함으로써 인공지능 기술에 대한 신뢰성을 제고하고 사회적 수용을 높일 수 있음
- 이해관계자들 간 갈등을 줄여 기술·산업 발전을 저해하지 않으면서도 기존 산업 관계자와 국민의 권리를 보호할 수 있음
- 국제적 규제체계 정합성을 통해 우리나라 기업의 글로벌 진출도 지원할 수 있음
- 인공지능 분야는 강한 규제가 산업 혁신을 저해할 우려가 크지만, EU도 규제샌드박스과 같은 혁신 지향적 조항을 두고 있고, 미국의 연방 법률안과 각 주의 법률에서도 EU 방식의 강한 규제가 일부 포함되어 있는 것으로 나타나 규제와 혁신 사이의 균형이 인공지능 입법의 궁극적인 지향점이 될 것으로 보임
- 아직 성숙하지 않은 인공지능 산업에 대하여 과도한 규제를 지양함으로써 궁극적으로 기술과 산업의 발전을 견인할 수 있음

차 례

□ 요약

I. 서론 / 1

1. 연구배경과 목적	1
2. 연구범위	2
가. 인공지능 윤리 규제체계 연구	3
나. 인공지능 저작권 규제체계 연구	3
3. 연구의 기대효과 및 활용방안	5

II. 인공지능 윤리 및 규제체계 / 6

1. 인공지능 윤리 및 규제체계	6
가. 인공지능 윤리	6
나. 우리나라 정부 및 공공기관 정책 동향	7
다. 윤리와 법	8
라. 인공지능 윤리 규제체계 연구 방향	9
2. 주요국 규제체계 현황	10
가. EU	10
나. 미국	29
다. 영국	46
라. 일본	50
마. EU와 미국 간 규제체계 접근 방식 비교	53
3. 국내 규제체계 진행 상황	57
가. 제21대 국회 발의 경과	57

나. 제22대 국회 발의 현황	58
4. 허위조작정보(딥페이크) 규제체계 논의 현황	78
가. 문제점	78
나. 국내외 입법 동향	79
다. 입법 방향	97
5. 인공지능 윤리 및 규제체계 입법 방향	100
가. 총론	100
나. 위험 기반 규제 방식 도입 여부	101
다. 생성형(범용) 인공지능 규제 여부	102
라. 연산능력에 따른 규제 여부	103
마. 제공자와 배포자 등 구분 여부	104
바. 거버넌스	104

Ⅲ. 인공지능 저작권 규제체계 / 107

1. 인공지능과 저작권	107
가. 저작물 등의 정의	107
나. 우리나라 정부 및 공공기관 정책 현황	107
다. 인공지능 저작권 관련 분쟁 현황	108
2. 학습데이터로서의 저작물 이용	115
가. 저작권법상 보호대상인 학습데이터	115
나. 인공지능의 저작물 학습 관련 특징	116
다. 학습데이터 이용에 따른 저작권 침해	117
라. 학습데이터 이용의 면책 여부 검토	117
마. 학습데이터 이용 시 저작권자 보상	165
바. 학습데이터 공개 관련	167
3. 인공지능 산출물의 보호	169

가. 인공지능 산출물 보호 관련 국내외 동향	169
나. 인공지능 산출물의 저작자 결정	172
4. 인공지능 산출물에 의한 저작권 침해	173
가. 일반 법리	173
나. 침해 인정 사례	175
5. 인공지능 저작권 규제체계 입법 방향	177
가. 총론	177
나. 포괄적 공정이용 조항 외에 TDM 면책 조항 필요 여부	177
다. 영리 목적 인정 여부	178
라. 저작권자의 옵트아웃 권리 인정 여부	178
마. 대상 권리 범위	179
바. 적법한 접근 관련	179
사. 학습데이터 출처 표시 여부	179
아. 복제물의 보관기간 관련	180
자. 보상 여부	180

IV. 결론 / 181

1. 연구 요약	181
가. 인공지능 윤리 규제체계 현황	181
나. 인공지능 저작권 규제체계 현황	183
2. 입법 방향	185
가. 인공지능 윤리 입법 방향	185
나. 인공지능 저작권 입법 방향	187
3. 시사점	189

□ 참고문헌

표 차례

[표 1] 주요 연구내용	4
[표 2] EU 인공지능법의 구조 및 개관	12
[표 3] 애플리케이션 유형별 EU와 미국의 AI 위험 관리 비교	54
[표 4] 제22대 국회 발의안 (2024.8.28. 기준)	58
[표 5] 법안별 규율체계 비교	59
[표 6] 인공지능 법안의 공통 내용 및 일부 법안에 포함된 내용	60
[표 7] 고위험영역 인공지능 규제 체계	68
[표 8] 제22대 국회 발의 인공지능 법안별 조문 체계 비교	71
[표 9] 허위조작정보 규제(콘텐츠산업진흥법 개정안)에 대한 이해관계자 의견 ·	81
[표 10] 허위조작정보 규제(정보통신망법 개정안)에 대한 이해관계자 의견 ·	83
[표 11] 인공지능 저작권 관련 주요 소송	108
[표 12] 주요국의 저작권 면책 관련 입법 현황	118
[표 13] 제21대 국회 TDM 면책 조항 발의안	131
[표 14] 제21대 국회 발의 법안의 TDM 정의 비교	135
[표 15] 제21대 국회 발의 법안별 TDM 규정 주요 내용 비교 분석 ···	136
[표 16] TDM 법안에 관한 저작권자(협회) 의견	141
[표 17] 주요국 저작권 법제 현황	163
[표 18] 주요국 TDM 정의	164

그림 차례

[그림 1] 사람이 중심이 되는 인공지능 윤리 기준	6
[그림 2] 영국 AI Regulation Landscape	48
[그림 3] 일본 AI 사업자 가이드라인 주요 내용	52
[그림 4] 생성형 인공지능에 의해 복제된 이미지 사례	117
[그림 5] Goldsmith의 사진(원저작물)과 Warhol의 작품 ‘Orange Prince’ (새로운 저작물)	123
[그림 6] 스테판 탈러의 ‘A Recent Entrance to Paradise’	170

I. 서론

1. 연구배경과 목적

인공지능(AI)은 사회 전반에 걸쳐 확산되고 있으며, 기술의 발전과 함께 그 영향력도 급격히 확대되고 있다. 인공지능은 의료, 금융, 교육, 교통 등 다양한 분야에서 혁신을 이끌어 내고 있으며, 그로 인해 인간의 삶의 질을 향상시키는 중요한 역할을 하고 있다. 특히, 최근 활발하게 개발·이용되고 있는 생성형 인공지능의 경우 텍스트, 이미지, 음악, 코드 등 다양한 형식의 콘텐츠를 자동으로 생성할 수 있어 다양한 산업에 혁신을 가져오고 있다.

그러나 인공지능의 발전과 보급이 확산됨에 따라 윤리적 문제와 법적 규제의 필요성이 대두되고 있다. 인공지능 관련 윤리적·법적 문제를 개발과 이용 단계로 나누어 살펴보면, 인공지능을 개발하는 단계에서는 인공지능이 저작물을 학습할 때 저작권 침해가 발생하는지 여부가 큰 화두이며, 인공지능을 이용하는 측면에서는 인공지능의 환각(hallucination) 현상으로 잘못된 정보를 제공하거나 인공지능의 강력한 성능을 오·남용함으로써 발생하는 윤리적 이슈에 대한 규제 논의가 한창이다.

즉, 인공지능 이용 시 기술의 내재적 불확실성과 불투명성으로 인하여 오동작 등 안전에 관한 우려가 있으며, 환각 현상으로 잘못된 정보를 제공하기도 한다. 특히 최근 생성형 인공지능을 활용한 허위조작정보(딥페이크, Deepfake)가 문제되고 있으며, 악성코드나 피싱 이메일 생성 등 사이버범죄에 악용되는 사례도 발생하고 있다.

아울러, 인공지능 학습에 양질의 대규모 데이터가 필요한데, 저작물을 학습할 경우 저작권 침해가 발생하는지, 공정이용으로 인정될 가능성은 있는지 명확하지 않은 상황이다. 인공지능 개발자·배포자와 저작권자 모두 법적 권리

와 책임이 불분명하며, 법적 분쟁의 가능성이 높아지고 있고 실제로 해외에서는 다수의 소송이 진행 중이다. 이처럼 저작권 규제체계가 정립되지 않아 문화 창달과 기술 혁신 모두가 저해될 수 있다.

즉, 인공지능 기술의 발전 및 사회적 수용성 제고를 위하여 현재 입법적으로 논의할 사항으로서, 인공지능 개발 단계에서는 데이터 학습 시 저작권 침해와 관련하여, 이용 단계에서는 윤리적 문제에 대하여 규제체계 마련이 중요하고 시급하다고 할 수 있다.

이에, 본 연구는 인공지능 기술의 발전이 윤리와 저작권에 미치는 영향에 주목하고 그 대응방안으로서 올바른 규제체계를 마련하기 위한 입법적 방향성을 모색한다. 인공지능 기술이 가져올 수 있는 잠재적 위험을 최소화하고 신뢰를 높이며 인간 중심의 기술 발전을 도모할 수 있는 윤리 규제체계를 제안하고자 한다. 아울러 저작권 보호와 인공지능 기술 발전 간 균형과 조화를 추구하는 저작권 규제체계를 제안하고자 한다.

2. 연구범위

본 연구는 윤리 규제체계와 저작권 규제체계에 대한 연구로 구분된다. 인공지능 기술의 발전과 확산으로 인하여 발생하는 법적·사회적 쟁점들을 규율하는 체계를 마련한다는 점에서는 공통된 측면이 있지만, 인공지능 윤리 규제체계는 인공지능 신뢰성 확보, 위험 관리, 거버넌스 구축 등 포괄적이고 거시적인 관점에서의 규제를 포함하고 있는 반면, 인공지능 저작권 규제체계는 인공지능의 저작물 학습 시 면책 여부, 인공지능 생성물의 저작권 인정 여부, 인공지능 생성물의 저작권 침해 등 저작권법을 둘러싼 쟁점들에 대한 연구가 주된 내용이라는 차이점이 있다.

가. 인공지능 윤리 규제체계 연구

인공지능 개발과 활용에서의 윤리원칙으로서 책임성, 안전성, 신뢰성 등이 제시되고 있는데, 이러한 윤리 원칙은 일반적으로 지향해야 할 중요한 가치들이지만, 그 자체만으로는 규범력과 강제력을 갖지 못하므로 정책이나 입법으로서 효과를 얻기 어려운 한계가 있다.

본 연구의 목적은 인공지능 윤리와 관련된 입법적 쟁점과 국내외 입법 현황 및 정책 사례 연구이므로, 법으로써 규율할 수 있는 규제체계와 관련된 쟁점에 한정하여 논의한다.

이에 따라, 주요국 규제체계 현황을 살펴보고 그 중 중요하게 다루어지고 있는 EU와 미국 사례를 중점적으로 분석하고 비교한다. 아울러 국내에서 발의된 여러 법안을 비교·분석하고 주요 입법 쟁점별로 방향을 제시한다.

한편, 최근 생성형 인공지능의 이용한 허위조작정보(딥페이크)가 사회적 문제로 대두되고 있어 이에 대한 국내외 규제 동향을 살펴보고 입법방향을 제시한다.

나. 인공지능 저작권 규제체계 연구

인공지능 기술의 개발 및 이용 확산으로 저작권법상 다양한 쟁점이 문제되고 있으며, 특히 생성형 인공지능의 출현은 저작권 산업 및 제도에 상당한 파급력을 미치고 있다.

그 중 입법적으로 논의가 중점적으로 필요한 사항은 인공지능의 저작물 학습 시 저작권 면책 여부이다. 인공지능이 대량의 데이터 학습 시 상당수는 저작물이 포함되어 있으며, 공개된 저작물인 경우에도 이를 이용하려면 원칙적으로 저작권자의 동의를 받아야 한다. 그러나 현실적으로 수많은 저작물의 저작권자로부터 동의를 받는 것은 불가능에 가깝다. 이에 따라 인공지능의 저

작물 학습이 포괄적 공정이용(저작권법 제35조의5)에 따라 면책이 가능한지, 저작권 면책을 위해 별도의 입법이 필요한지에 대한 논의가 활발히 진행되고 있다.

본 연구에서는 인공지능 저작물 학습 시 포괄적 공정이용 적용이 가능한지에 대하여 국내외 연구 및 판례를 살펴본다. 아울러 국내외에서 입법이 추진되거나 완료된 TDM(Text-Data Mining) 면책 조항이 인공지능 학습에 적용될 수 있는지, 적절한 입법방안은 무엇인지 살펴본다.

이 외에 인공지능 저작권과 관련된 쟁점으로 인공지능이 생성한 저작물에 저작권이 부여될 수 있는지 여부가 쟁점이 될 수 있는데, 인공지능에 권리능력이나 행위능력 등을 인정하기 어려워 순수하게 인공지능 자체에 저작권을 부여하기 어려우며, 사람인 저작자가 인공지능을 도구로서 사용한 경우에는 기존의 법리에 따라 창작성이 인정된다면 저작권이 인정될 수 있으므로 본 연구에서는 국내외 동향을 간략히만 살펴본다.

그리고 인공지능이 생성한 결과물이 다른 저작권을 침해하는지 여부도 쟁점이 될 수 있는데, 이는 기존의 저작권 침해 법리를 따르면 될 것으로 보여 이 또한 간략히 살펴본다.

[표 1] 주요 연구내용

분야	세부 주제	연구 방법 및 전략
인공지능 윤리 규제체계	- 주요국 규제체계 현황 - 국내 입법발의 현황 - 허위조작정보 규제 현황	• 문헌 검토를 통한 자료 수집 • 보고서, 논문, 정책 자료, 신문기사, 국내외 주요 사이트 등 검토
인공지능 저작권 규제체계	- 인공지능 저작물 학습 - 인공지능 생성물의 저작권 - 인공지능 생성물의 저작권 침해	

3. 연구의 기대효과 및 활용방안

본 연구를 통해 기대할 수 있는 첫 번째 효과는 윤리적인 인공지능 개발의 촉진이다. 본 연구를 통해 도출되는 규제체계는 인공지능 기술의 개발과 활용 과정에서 인간의 기본권을 보호하고 위험을 관리하여, 신뢰할 수 있고 책임성과 공정성 등을 갖춘 인공지능이 개발될 수 있도록 유도할 것이다.

아울러 본 연구를 통해 마련된 규제체계는 기술 발전을 저해하지 않으면서도, 그에 따르는 윤리적, 법적 문제를 해결하는 방향으로 작용할 것이다. 이는 기술 발전과 법적 안정성 간의 조화를 이루는 데 중요한 역할을 할 것으로 기대된다.

한편, 인공지능 개발 시 상당량의 저작물을 학습하는데, 이해관계자들 간 균형 있는 규제체계를 제안하여 이와 관련된 법적 불확실성을 줄임으로써 인공지능 기술의 개발을 촉진하고 저작권자의 권리도 보호할 수 있을 것으로 기대된다.

본 연구의 결과는 인공지능 윤리와 저작권 규제체계 정립을 위한 입법정책에 직접적으로 활용될 수 있다. 안전하고 신뢰할 수 있는 인공지능 기술 확보를 위한 규제의 근거 자료로 활용될 수 있고, 저작권자와 인공지능 개발자·배포자들 간 이해관계를 균형 있게 조정하는 법률안을 제시하는 데 활용될 수 있다.

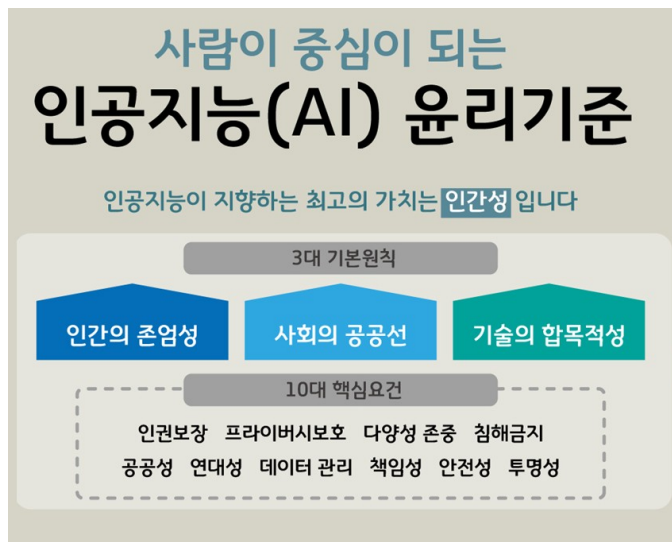
II. 인공지능 윤리 및 규제체계

1. 인공지능 윤리 및 규제체계

가. 인공지능 윤리

인공지능 윤리의 개념과 관련하여, 정부에서 2020년 발표한 사람이 중심이 되는 「인공지능(AI) 윤리기준」은 ‘인간성(Humanity)’을 위한 3대 기본원칙과 10대 핵심요건을 제시하고 있으며, 10대 핵심요건에 ‘책임성’, ‘투명성’, ‘다양성 존중(비차별성, 공정성)’, ‘프라이버시 보호’를 포함하고 있다.

[그림 1] 사람이 중심이 되는 인공지능 윤리 기준



자료 : 과학기술정보통신부·4차산업혁명위원회

아울러 소프트웨어정책연구소(SPRi)는 신뢰성 개념에 프라이버시 보호, 공정성, 투명성, 책임성을 포함하고 있다.¹⁾

나. 우리나라 정부 및 공공기관 정책 동향

우리나라 정부는 2019년 인공지능 강국을 비전으로 한 ‘인공지능 국가전략’을 수립하였고,²⁾ 인공지능 산업 진흥과 활용 기반을 강화하면서 역기능을 방지하기 위해 ‘인공지능 법·제도 규제 정비 로드맵’을 발표하였다.³⁾

이후, 각 부처별로 인공지능 정책을 수립하고 있다. 금융위원회는 2021년 소비자 및 금융 사업·시장 건전성 보호를 위해 인공지능 운영의 방향성을 제시한 ‘금융 분야 AI 가이드라인’을 발표하였고, 국가인권위원회는 2022년 인공지능 개발과 활용 과정에서 발생할 수 있는 인권침해와 차별을 방지하기 위해 ‘인공지능 개발과 활용에 관한 인권 가이드라인’을 마련하여 국무총리와 관련 부처 장관에게 권고한 바 있다. 개인정보보호위원회는 2023년 개인정보 침해 위험은 최소화하면서 인공지능 혁신 생태계 발전에 꼭 필요한 데이터는 안전하게 활용할 수 있도록 하는 ‘인공지능 시대 안전한 개인정보 활용 정책 방향’을 발표하였다.⁴⁾

과학기술정보통신부는 2023년 인공지능을 비롯하여 보편적 디지털 질서 규범을 위한 현장인 ‘디지털 권리장전’을 발표하였다.⁵⁾ 디지털 권리장전은 디지털 심화 시대에 맞는 국가적 차원의 기준과 원칙을 제시하고, 글로벌을 리

- 1) 조원영 외, 「인공지능 신뢰체계 정립방안 연구」, 연구보고서 RE-123, 소프트웨어정책연구소, 2023; 유재홍 외, 「글로벌 AI 신뢰성 정책 동향 연구」, 연구보고서 RE-150, 소프트웨어정책연구소, 2023.
- 2) 과학기술정보통신부 보도자료, “범정부 역량을 결집하여 AI 시대 미래 비전과 전략을 담은 ‘AI 국가전략’ 발표”, 2019. 12. 17.
- 3) 과학기술정보통신부 보도자료, “인공지능 시대를 준비하는 법·제도·규제 정비 로드맵 마련”, 2020. 12. 23.
- 4) 법제처 미래법제혁신기획단, 「인공지능(AI) 관련 국내외 법제 동향」, 법제소식, 2024. 7.
- 5) 과학기술정보통신부 보도자료, “대한민국이 새로운 디지털 규범질서를 전 세계에 제시합니다!”, 2023. 9. 25.

드릴 수 있는 보편적 디지털 질서 규범의 기본방향을 담은 헌장(憲章)으로서 배경과 목적을 담은 전문과 함께 총 6장, 28개조가 담긴 본문으로 구성된다. ‘디지털 공동번영사회’ 구현을 위한 기본원칙으로서 제1장에서 ❶ 디지털 환경에서의 자유와 권리 보장, ❷ 디지털에 대한 공정한 접근과 기회의 균등, ❸ 안전하고 신뢰할 수 있는 디지털 사회, ❹ 자율과 창의 기반의 디지털 혁신의 촉진, ❺ 인류 후생의 증진의 5가지를 제시하였으며, 이후 각 원칙의 세부 원칙을 규정하였다.

정부출연연구소, 국책연구기관들도 인공지능 윤리 정책을 심도있게 연구하고 있다. 한국법제연구원은 2021년 「AI 윤리 관련 법제화 방안 연구」를 통해 국제적으로 다양한 단체 또는 기관에서 배포한 AI 윤리 원칙에 대한 비교 분석을 통해 향후 우리나라의 AI 관련 윤리 이슈의 법제화 방안 제시하였다.⁶⁾

소프트웨어정책연구소는 인공지능 윤리 및 신뢰에 관한 연구를 이어오고 있으며, 2023년 「글로벌 AI 신뢰성 정책 동향 연구」를 통해 인공지능 신뢰성 확보를 위한 국내외 정책 동향을 조사하고 분석하여 국내 인공지능 정책 고도화를 위한 시사점을 제공하였다.⁷⁾

다. 윤리와 법

윤리⁸⁾와 법의 구분에 관하여, 칸트는 법과 도덕의 구분은 ‘강제’의 유무에서 차이가 있다고 하였으며, 엘리네크는 법을 ‘도덕의 최소한’이라고 표현한 바 있다.

거트(Gert)는 도덕 또는 윤리가 비형식적인 공적 체계임에 반하여 법은 형

6) 최경진·이기평, 「AI 윤리 관련 법제화 방안 연구」, 글로벌법제전략 연구, 2021. 10.

7) 위 유재홍 외 보고서

8) ‘윤리’와 ‘도덕’ 용어 간 차이가 있다는 견해도 있지만, 본 보고서에서는 통용하여 사용하기로 한다.

식적·공적 체계라고 하였다. 즉, 법은 모든 국민에게 적용되는 행위의 규범뿐만 아니라 행위가 실현되기 위해 따라야 하는 절차를 규정함으로써 법익을 보호하고자 한다. 그리고 이런 목표의 달성은 형식적인 체계에 의한 것으로 해당 행위와 절차를 명문으로 규정하고, 규정의 위반에 따른 처리 절차 역시 형식적인 체계에 부합하도록 운영하는 특성을 지닌다. 따라서 이러한 형식적인 체계에 부합하거나 형식적인 체계에 의해 관리될 수 없는 것은 법률의 내용으로 삼기 어렵고, 설사 이러한 내용이 법률에 포함된다 하더라도 실질적인 효과를 가져오기 어렵다.⁹⁾

라. 인공지능 윤리 규제체계 연구 방향

2010년대 인공지능 기술이 비약적으로 발전하고 2016년 알파고와 이세돌의 바둑대국 이후 관심이 높아지면서 인공지능 기술의 개발 및 이용에 관한 윤리원칙 연구가 다수 진행되었다. 인공지능 신뢰성 확보를 위해 프라이버시 보호, 공정성, 투명성, 책임성 등의 가치들이 추구되어야 한다는 것이다.

이러한 윤리 원칙을 정립하고 이를 준수하도록 함으로써 인공지능 기술에 대한 신뢰가 높아질 수 있고 기술에 대한 막연한 불안을 불식시켜 사회적 수용성을 높일 수 있다. 그러나 윤리원칙 그 자체만으로는 규범력과 강제력을 갖지 못하므로 정책이나 입법으로서 효과를 야기할 수 없는 한계가 있다.

이에 따라, 본 연구에서는 법으로써 규율할 수 있는 규제체계와 관련된 입법적 쟁점에 한정하여 논의한다.

먼저 주요국 규제체계 동향을 살펴본다. EU 인공지능법은 2024년 8월 1일에 발효되었고 2025년 2월 2일부터 단계적으로 예외적 적용이 개시될 예정이다. 이 법은 세계 최초로 제정된 포괄적이고 전분야에 걸친 규제로서 전 세

9) 최경석, 「생명윤리에서의 법, 도덕 및 윤리의 역할과 한계」, 이화여자대학교 법학논집 제15권 제4호, 2011. 6.

계에서 주목하고 있으므로 상세히 살펴본다. 아울러 미국의 인공지능 행정명령과 각 주(州)의 주요 규제를 살펴본다. 그리고 우리나라 제22대 국회에서 발의된 인공지능 관련 법안을 비교·분석한다. 허위조작정보 관련 국내외 규제 동향을 살펴보고, 인공지능 윤리 규제체계 입법 방향에 대하여 논의한다.

2. 주요국 규제체계 현황

가. EU

(1) EU 인공지능법의 주요 내용과 특성

(가) 도입 연혁 및 개관

1) 도입 연혁

인공지능 기술의 발전과 경쟁이 심화되고 있다. 그만큼 인공지능 규제에 관한 논쟁도 뜨겁다. 유럽의회는 2024년 3월 세계 최초로 ‘EU 인공지능법(EU AI Act)’을 통과시켰다. 적용 지역은 27개 회원국이 포함된 EU 전역이며, 발효 6개월 후부터 순차적으로 시행될 예정이다. 인공지능 기술개발과 활용에 대한 강력한 규제와 처벌 조항이 포함되어 있어 여러 논의가 이루어지고 있다.

EU는 인공지능 기술에 대한 규제를 마련하기 위해 점진적으로 시도해왔다. EU 이사회는 2017년 10월 19일 EU 집행위원회에 인공지능에 관한 유럽의 접근 전략을 제안할 것을 권고하였고, EU 집행위원회는 2018년 4월 25일에 유럽을 위한 인공지능 정책 추진안을 발표하였다. 인공지능 기술에 대한 구속력 있는 규범을 제정하려는 본격적인 시도는 EU 집행위원회 위원장인 폰 테어 라이엔(Ursula von der Leyen)의 공약에서 시작되었다. 2019년 EU 집행위원회 위원장에 취임한 폰 테어 라이엔은 정책 공약집에서 인공지능에 대한

EU의 조직적 접근 계획을 모색하기 위한 법안을 제출할 것을 언급하였고¹⁰⁾, 취임 직후 ‘인간의 안전과 권리를 존중하는 새로운 인공지능 규범’을 마련한다는 내용을 포함한 디지털 혁신 전략 계획을 발표하였다.¹¹⁾ 이러한 논의를 바탕으로, EU 집행위원회는 인공지능 기술에 대한 비구속적인 윤리 원칙인 ‘신뢰할 수 있는 인공지능에 대한 윤리지침’을 제정하였다.¹²⁾

EU 집행위원회는 이후 2020년 2월 19일 인공지능 백서를 발표하고¹³⁾ 2021년 4월 21일 새로운 인공지능법안을 제안하였다.¹⁴⁾ 유럽 의회는 2023년 6월 14일 의견 수렴을 거친 EU 집행위원회의 제안에 대한 수정안을 제안하였으며 유럽 이사회가 2024년 2월 13일 수정안에 합의하였다. 이후 유럽 의회가 2024년 3월 13일 합의안을 최종 채택하고 이어서 2024년 5월 21일 유럽 이사회가 이를 최종 승인함으로써 세계 최초의 인공지능법안이 시행되게 되었다. EU 인공지능법은 2024년 8월 1일 공식 발효되었고, 발효일로부터 24개월 이후 시행된다. 다만 일부 조항은 발효일로부터 6개월, 12개월 또는 36개월 이후 단계적으로 시행된다.

10) Ursula von der Leyen, “A Union that strives for more, My agenda for Europe By Candidate for President of the European Commission”, Political Guidelines for the Next European Commission 2019-2024, 2018, p. 13.; 김건우 편저, 「인공지능 규제 거버넌스의 현재와 미래」, 파이돈, 2023, 151면.

11) European Commission, “A Union that strives for more: the first 100 days”, 2020. 3. 6., <https://neighbourhood-enlargement.ec.europa.eu/news/union-strives-more-first-100-days-2020-03-06_en> 최종방문일 2024. 8. 29.>

12) High-Level Expert group on AI, “Ethics guidelines for trustworthy AI”, 2019.

13) European Commission, “White Paper on Artificial Intelligence: a European approach to excellence and trust”, 2020. 2. 19.

14) European Commission, “Proposal for a REGULATION OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL LAYING DOWN HARMONISED RULES ON ARTIFICIAL INTELLIGENCE(ARTIFICIAL INTELLIGENCE ACT) AND AMENDING CERTAIN UNION LEGISLATIVE ACTS”, Brussels, 2021. 4. 21., COM(2021) 206 final.

2) EU 인공지능법의 구조 및 개관

EU 인공지능법은 총 113개의 규정과 13개의 부속서로 구성되어 있다. 각 장의 제목과 내용은 다음과 같다.¹⁵⁾

[표 2] EU 인공지능법의 구조 및 개관

장	제목		조항	시행일
제1장	총칙		제1~4조	6개월 후
제2장	금지된 AI 행위		제5조	
제3장	고위험 AI 시스템	제1절 고위험 AI 시스템의 분류	제6~7조	24개월 후 다만, 제6조제1항 (고위험 AI 시스템 분류, 의무)은 36개월 후
		제2절 고위험 AI 시스템에 대한 요건	제8~15조	24개월 후
		제3절 고위험 AI 시스템 제공자, 배포자 및 기타 당사자의 의무	제16~27조	
		제4절 통보기관 및 인증기관	제28~39조	12개월 후
		제5절 표준, 적합성 평가, 인증서, 등록	제40~49조	24개월 후
제4장	특정 AI 시스템 제공자 및 배포자에 대한 투명성 의무		제50조	
제5장	범용 AI 모델	제1절 분류 규칙	제51~52조	12개월 후
		제2절 범용 AI 모델 제공자의 의무	제53~54조	
		제3절 시스템적 위험이 있는 범용 AI 모델 제공자의 의무	제55조	
		제4절 실천 강령	제56조	
제6장	혁신 지원 방안		제57~63조	24개월 후
제7장	거버넌스	제1절 EU 수준에서의 거버넌스	제64~69조	12개월 후
		제2절 국가 관할기관	제70조	
제8장	고위험 AI 시스템을 위한 EU 데이터베이스		제71조	24개월 후

15) 한국법제연구원(2024), “EU AI ACT 번역본” 참고

장	제목		조항	시행일
제9장	출시 후 모니터링, 정보 공유, 시장감독	제1절 출시 후 모니터링	제72조	24개월 후
		제2절 중대 사건에 관한 정보의 공유	제73조	
		제3절 집행	제74~84조	
		제4절 구제수단	제85~87조	
		제5절 범용 AI 모델 제공자에 대한 감독, 조사, 집행 및 감시	제88~94조	
제10장	행동강령 및 지침		제95~96조	
제11장	권한의 위임 및 위원회의 절차		제97~98조	
제12장	벌칙		제99~101조	제99~100조는 12개월 후 제101조(범용 AI 모델 제공자에 대한 벌금)는 24개월 후
제13장	최종 조항		제102~113조	24개월 후
13개의 부속서(I~XIII)				부속서 언급 조항에 따른 시행

EU 인공지능법은 AI 시스템을 위험도에 따라 ① 금지된 AI 행위, ② 고 위험 AI 시스템, ③ 제한적 위험 AI 시스템¹⁶⁾, ④ 최소 위험 AI 시스템으로 분류하고 그 수범자들에게 차등적으로 의무를 부과한다.

사람의 잠재의식 조작, 착취, 사회통제 등 EU가 보호하는 기본권, 가치에 위배되는 방식으로 AI 시스템을 사용하는 것은 금지된다. EU 내 개인의 보건, 안전, 기본권에 중대하고도 유해한 영향을 미치는 고위험 AI 시스템에는 적합성 평가와 시판 후 모니터링 등 기본권 영향평가 의무 등이 부과된다. 사람과

16) 법 최종본에서는 ‘특정 AI 시스템’으로 규정되어 있으며, 본 보고서에서는 ‘제한적 위험 AI 시스템’과 ‘특정 AI 시스템’ 용어를 혼용하여 사용한다.

상호작용하는 AI 시스템은 기만, 조작 등의 문제가 발생할 수 있으므로 이러한 제한적 위험 AI 시스템 내지 특정 AI 시스템에는 AI 생성 콘텐츠임을 표시하는 등의 투명성 의무가 부과된다. 최소 위험 AI 시스템이란 게임에서 제공되는 AI 보조 기능 등 일반적인 AI 시스템으로, EU 인공지능법에서 특별히 규제하지 않는다.¹⁷⁾

EU 인공지능법은 또한 다양한 AI 시스템 개발의 기초로 쓰일 수 있는 파운데이션 모델을 범용 AI 모델로 규정하고 그에 관한 규제 조항을 별도로 두어 범용 AI 모델 제공자가 준수해야 할 사항들을 정하고 있다.

그밖에 EU 인공지능법은 AI 규제 샌드박스, 신생기업을 포함한 중소기업에 대한 지원 등의 내용을 담은 AI 산업, 기술 진흥을 지원할 근거를 마련하고 있다.

(나) EU 인공지능법의 적용 범위

EU 인공지능법은 ① EU의 역내 혹은 역외에서 설립되었거나 소재하고 있는지를 불문하고 EU에서 AI 시스템을 시장에 출시하거나 서비스를 제공하는 제공자(Provider), ② EU 내에서 설립되었거나 소재하고 있으며 자신의 권한 하에 AI 시스템을 사용하는 AI 시스템의 배포자(Deployer), ③ (AI 시스템의 개발 장소가 EU의 역외에 위치하더라도) AI 시스템의 산출물을 EU에서 사용하는 경우, EU의 역외에서 설립되었거나 소재하고 있는 제공자 및 배포자, ④ AI 시스템의 수입업자 및 유통업자(importers and distributors) 등에게 적용된다.¹⁸⁾¹⁹⁾ EU 인공지능법은 AI 시스템을 위험성에 기반하여 분류한 후 그에

17) 라기원, 「유럽연합 인공지능법(EU AI ACT)의 특성과 쟁점-우리나라 인공지능 입법에 대한 시사점」, 한국법제연구원, 2024, 13쪽.

18) EU 인공지능법 제2조

19) 심소연, 「규제 중심의 유럽연합 인공지능법(EU AI ACT)」, 국회도서관 최신 외국입법정보, 2024, 4쪽

따라 위와 같은 수범자들에게 여러 제한과 의무를 부과한다.

EU 인공지능법은 EU 내에서 AI 시스템을 출시하거나 서비스를 제공하거나 사용하는 행위에 적용된다. 다만 회원국의 국가안보, 군사, 국방 또는 국가안보의 목적인 경우, 국제기구가 EU 또는 하나 이상의 회원국과 국제협력 또는 법 집행 및 사법 협력을 위한 협약에 따른 경우, 오직 연구·개발만을 목적으로 개발되고 서비스를 제공하는 모델 또는 시스템, 시장에 출시되거나 서비스를 제공하기 전의 모델 또는 시스템, 오직 개인적이고 비전문적인 활동 과정에서 AI 시스템을 사용하는 자연인 등에 대해서는 EU 인공지능법이 적용되지 않는다.²⁰⁾

(다) AI 시스템의 분류와 관계자의 의무

1) 개관 - 위험 기반 규제

EU 인공지능법은 범용 AI 모델 중 시스템적 위험이 있는 범용 AI 모델을 분류하는 한편²¹⁾, 이를 기반으로 만든 AI 시스템을 위험성에 따라 분류하여 규제 수준을 차등화하는 위험 기반 접근법을 취한다(Risk-based approach).²²⁾ 이에 따라 EU 인공지능법은 ① 특정한 AI 행위를 금지하거나, ② 고위험 AI 시스템, ③ 제한적 위험 AI 시스템, ④ 최소 위험 AI 시스템 등으로 분류하여 각각에 대해 서로 다른 제한과 의무를 규정하는 등 차등적으로 규제하는 방식을 취하고 있다.

20) EU 인공지능법 제2조 제3항 내지 제12항

21) EU 인공지능법 제51조

22) 라기원, 위 논문 12쪽.

2) 범용 AI 모델

범용 AI 모델이란 대규모 자기지도(Self-supervised learning)를 이용하여 다량의 데이터로 학습된 AI 모델로서 상당한 일반성을 보이고, 모델이 시장에 출시되는 방식과 관계없이 다양한 고유 작업을 능숙하게 수행하며, 다양한 하위 시스템이나 응용프로그램(어플리케이션)에 통합될 수 있는 AI 모델을 말한다.²³⁾ EU 인공지능법은 범용 AI 모델이 일반성과 광범위한 업무를 능숙하게 수행할 수 있는 능력을 갖는다는 점에서 일반적인 AI 시스템과 구별되어야 한다는 점을 명시하고 있다.²⁴⁾

EU 인공지능법은 범용 AI 모델 제공자의 의무를 다음과 같이 규정하고 있다²⁵⁾.

범용 AI 모델 제공자의 의무 ²⁶⁾
• 부속서 XI에 규정된 정보를 포함하는 학습 및 시험 과정과 평가 결과를 포함한 기술문서의 작성 및 업데이트 • EU 법령 및 회원국의 국내법에 따라 지식재산권과 기밀 사업 정보 또는 영업비밀을 침해하지 아니하는 범위에서, 범용 AI 모델을 AI 시스템에 통합하려는 AI 시스템 제공자에게 최신 상태의 정보와 문서를 제공 • EU 저작권법을 준수하고 지침(EU) 제2019/790호²⁷⁾ 제4조 제3항에 따라 표시된 권리의 유보를 확인하고 준수하기 위한 정책을 시행 • 인공지능청이 제공하는 양식에 따라 범용 AI 모델의 학습에 사용된 콘텐츠에 관하여 충분히 상세한 요약서를 작성하여 공개

23) EU 인공지능법 제3조 제63항

24) EU 인공지능법 전문 97

25) EU 인공지능법 제53조

26) EU 인공지능법 제53조

27) 디지털 단일 시장에서 저작권 및 관련 권리에 관한 지침(EU) 2019/790

EU 인공지능법은 범용 AI 모델 중 고성능 모델을 다시 시스템적 위험(systemic risk)이 있는 범용 AI 모델로 분류한다.²⁸⁾ 적절한 기술적 도구와 방법론을 기반으로 평가할 때 고(高)영향력이 있다고 인정되거나, EU 집행위원회가 직권으로 결정하는 경우 등에는 시스템적 위험이 있는 범용 AI 모델로 분류되며²⁹⁾, 훈련에 사용된 누적 연산량이 10^{25} FLOPs(Floating point Operations Per Second, 부동소수점 연산) 이상인 경우 ‘고영향력’이 있는 것으로 추정된다.³⁰⁾ 시스템적 위험이 있는 범용 AI 모델의 제공자는 일반적인 범용 AI 모델의 제공자에 비해 추가적인 의무를 부담하게 된다.³¹⁾

시스템적 위험이 있는 범용 AI 모델 제공자의 의무³²⁾

- 기술문서 제출 의무 강화(부속서 11 제2절)
- 시스템적 위험을 확인 및 완화하기 위한 적대적 시험을 수행하고 문서화하는 등 최신 기술을 반영하는 표준화된 강령 및 도구에 따라 모델 평가를 수행
- 시스템적 위험이 있는 범용 AI 모델의 개발, 시장 출시 또는 사용으로 인해 발생할 수 있는 EU 수준의 시스템적 위험과 그 원인을 평가하고 완화하는 조치
- 중대한 사건에 관한 관련 정보와 이를 해결하기 위한 가능한 시정 조치를 추적하고, 문서화하고, 과도한 지체 없이 인공지능청에 보고하며 적절하게 국가 관할기관에 보고
- 해당 모델의 물리적 기반시설에 대한 적절한 수준의 사이버보안 보호를 보장
- EU 통합표준 공표 전까지 제56조에 따른 실천 강령 준수

28) EU 인공지능법 제51조 제1항

29) EU 인공지능법 제51조 제1항

30) EU 인공지능법 제51조 제2항

31) EU 인공지능법 제55조

32) EU 인공지능법 제55조

3) 금지된 AI 행위

EU 인공지능법은 인간의 존엄성, 자유, 평등, 민주, 법치, 기본권 등에 위협을 발생시킬 수 있는 금지된 AI 행위 내지는 사용례들(Prohibited AI Practices)을 규정하고 있다.³³⁾

이와 같은 금지된 AI 행위의 유형으로는 ① 사람의 잠재의식을 조작하거나 특정 집단의 취약성을 이용하여 상당한 피해를 유발할 수 있는 경우³⁴⁾, ② 사회적 행동이나 개인의 특성에 기반하여 생성·수집된 정보를 기반으로 개인이나 집단의 사회적 점수(social score)를 도출하여 불리한 처우를 초래할 수 있는 경우³⁵⁾, ③ 프로파일링 또는 개인의 성격 특성만을 토대로 개인의 범죄가능성을 예측 및 평가하는 경우³⁶⁾, ④ 인터넷이나 CCTV 영상의 얼굴 이미지를 임의로 수집하여 안면 인식 데이터베이스를 생성하는 경우³⁷⁾, ⑤ 직장 또는 교육기관에서 사람의 감정을 추론하기 위해 활용되는 경우³⁸⁾, ⑥ 생체인식데이터를 기반으로 인종, 정치적 견해, 종교적 또는 철학적 신념, 성적 지향성, 성생활 또는 성적 지향을 추론하는 경우³⁹⁾, ⑦ 공공장소에서 실시간 원격으로 생체정보를 인식하는 경우⁴⁰⁾ 등이 포함된다.

이러한 금지된 AI 행위를 목적으로 하는 AI 시스템은 EU 인공지능법에서 명문으로 규정된 특별한 예외가 없는 한 시장에 출시되거나 서비스가 제공되는 것이 말 그대로 금지된다.

33) EU 인공지능법 제5조

34) EU 인공지능법 제5조 제1항 (a), (b)

35) EU 인공지능법 제5조 제1항 (c)

36) EU 인공지능법 제5조 제1항 (d)

37) EU 인공지능법 제5조 제1항 (e)

38) EU 인공지능법 제5조 제1항 (f)

39) EU 인공지능법 제5조 제1항 (g)

40) EU 인공지능법 제5조 제1항 (h)

4) 고위험 AI 시스템

고위험(high risk) AI 시스템이란 자연인의 기본권 또는 보건과 안전에 위해를 가할 중대한 위험이 있고 시스템의 사용이 최종적인 결정에 영향을 미치는 AI 시스템으로이다.⁴¹⁾ 구체적으로 첫째, <부속서 I>에 열거된 EU 조정법⁴²⁾(Union harmonisation legislation)의 적용을 받는 제품⁴³⁾의 안전 구성 요소이거나 제품 그 자체이면서, 해당 EU 조정법이 제품의 출시 또는 서비스 개시를 위해 제품에 대한 제3자 적합성 평가를 요구하는 경우이다. 둘째, <부속서 III>에 명시된 8가지 목적⁴⁴⁾을 위해 사용되는 AI 시스템을 말한다.⁴⁵⁾ 단, <부속서 III>에 따른 AI 시스템이 자연인의 건강, 안전 또는 기본권에 중대한 위해를 가할 위험이 없는 경우⁴⁶⁾ 고위험으로 간주되지 않지만 그 경우에도 해당 AI 시스템이 자연인에 대한 프로파일링을 수행하는 경우에는 항상 고위험으로 간주된다.⁴⁷⁾

41) EU 인공지능법 제6조 제3항

42) 혹은 EU 통합법률 등으로 지칭됨

43) 기계류, 장난감, 엘리베이터, 케이블카, 의료기기, 수상기기, 항공기기, 철도시스템, 자동차 등을 포함한다.

44) ① 금지되지 않는 생체 인식·분류,
 ② 주요 사회기반시설의 관리·운영,
 ③ 교육 및 직업훈련,
 ④ 고용, 인사 관리, 개인의 업무 성과 평가 및 감독,
 ⑤ 필수 민간 서비스 및 공공 서비스에 대한 접근 및 혜택의 향유,
 ⑥ EU법 또는 회원국 법에 의해 허용되는 법 집행,
 ⑦ 이민, 망명 및 국경 통제 관리,
 ⑧ 사법행정 및 선거 또는 국민투표 등 민주적 절차

45) EU 인공지능법 제6조

46) AI 시스템이 (a) 좁은 절차적 작업을 수행하기 위한 경우, (b) 이전에 완료된 인간 활동의 결과를 개선하기 위한 경우, (c) 의사결정 패턴 또는 이전의 의사결정 패턴으로부터 이탈을 탐지하기 위한 것이며, 적절한 인적 검토 없이 이전에 완료된 인적 평가를 대체하거나 이에 영향을 미치기 위한 것이 아닌 경우, 부속서 III에 열거된 배포 사례의 목적과 관련된 평가에 대한 준비 업무를 수행하기 위한 경우

EU 인공지능법은 고위험 AI 시스템의 시장 출시 단계부터 활용에 따른 효과 분석까지 시스템의 전 주기에 걸친 적법 요건, 요구사항을 규정하여 AI 시스템이 사람의 기본권, 보건과 안전에 미치는 피해를 최소화하고자 한다.⁴⁸⁾ EU 인공지능법은 고위험 AI 시스템의 이해관계자별 준수사항을 다음과 같이 상세히 규정하고 있다.

첫째, 고위험 AI 시스템 ‘제공자’의 의무는 다음과 같다.

고위험 AI 시스템 제공자의 의무 ⁴⁹⁾
• 시스템상에 이름, 상호·상표, 주소 등을 표시 또는 동봉 문서에 표시 • 품질관리 시스템 마련 • 기술문서, 품질관리 문서, 인증기관의 변경 승인 문서, 인증기관이 발급한 결정문, EU 적합성 선언서 등 문서보관 의무 • 고위험 AI 시스템에서 자동으로 생성된 로그를 보관 • EU 시장 출시 또는 서비스 투입 전 제43조에 따른 적합성 평가 절차 이행 및 EU 적합성 선언서 작성 • CE 마크⁵⁰⁾를 부착하거나 포장 또는 동봉 문서에 적합성 표시 • EU 집행위원회가 설치·관리하는 EU 데이터베이스에 등록 • 시정조치 및 정보제공 의무 • 지침(EU) 2016/2102⁵¹⁾ 및 지침(EU) 2019/882⁵²⁾에 따른 접근성 요구사항 준수 확인 • 시장 출시 이후 모니터링 • 중대한 사건 보고의무

47) EU 인공지능법 제6조 제3항 (d)

48) 라기원, 위 논문 15쪽.

49) EU 인공지능법 제16조

50) 프랑스어 'Conformité Européenne'의 약자로, CE 마크를 획득한 제품은 소비자의 안전과 건강, 위생, 환경보호 등에 관한 유럽의 기준에 부합한다는 의미의 안전 인증 마크

51) 공공기관의 웹사이트와 모바일 앱 접근성에 관한 지침(EU) 2016/2102

52) 제품과 서비스에 대한 접근성 요구조건에 관한 지침(EU) 2019/882

둘째, 고위험 AI 시스템 ‘배포자’의 의무는 다음과 같다.

고위험 AI 시스템 배포자의 의무 ⁵³⁾	
• 적절한 기술적·조직적 조치 및 필요한 역량, 학습, 권한과 자원을 갖춘 자연인에게 인적 감독 책임을 부여 • AI 시스템의 설계 목적을 고려하여 입력 데이터의 충분한 관련성, 대표성이 있는지 확인 • 고위험 AI 시스템 모니터링, 시판 후 모니터링 결과를 제공자에게 고지 • AI 시스템이 위험을 유발할 수 있는 경우 사용을 중단하고 제공자, 유통업자 및 시장감독기관에 통지 • 고위험 AI 시스템에서 자동으로 생성된 로그를 최소 6개월 보관(단, EU 법률이나 회원국 법률, 특히 개인 정보 보호에 관한 EU 법률에 달리 명시된 경우는 제외) • 사업주인 배포자는 근로자 대표와 영향을 받는 근로자에게 고위험 AI 시스템 사용 대상이 될 것임을 고지 • 배포자가 공공기관 또는 EU의 기관인 경우 EU 데이터베이스에 시스템 등록 의무 준수 • 범죄 수사 관련 원격 생체 정보 인식용 고위험 AI 시스템 배포자는 사법당국에 사전에 또는 48시간 이내에 승인 요청. 승인이 거부되는 경우 관련된 원격 생체 인식용 시스템의 사용을 중단하고 관련된 데이터 삭제 • 관련 시장감시기관 및 국가 데이터 보호 기관에 원격 생체 인식용 시스템의 사용에 관한 연례보고서를 제출 • (배포자가 공공기관이거나 공공서비스를 제공하는 경우) 기본권 영향평가 의무 부과(제27조)	

셋째, 둘째, 고위험 AI 시스템 ‘수입업체와 유통업체’의 의무는 다음과 같다.

고위험 AI 시스템 수입업체와 유통업체의 의무 ⁵⁴⁾		
구분	수입업체	유통업체
공통 의무	• CE 마크 표시 여부, EU 적합성 선언 및 사용 지침의 제공 여부 확인 및 보장	

53) EU 인공지능법 제26조

54) EU 인공지능법 제23조, 제24조

고위험 AI 시스템 수입업체와 유통업체의 의무 ⁵⁴⁾		
	• 해당 시스템이 규정 및 요건에 부합하지 않는 것으로 판단할 충분한 이유가 있는 경우 출시 금지 • 해당 시스템이 인간의 건강, 안전 또는 기본권에 위험을 초래하는 경우 시스템 제공자 및 시장감독기관에 통지 • 고위험 AI 시스템의 위험 감소, 완화를 위해 당국과 협력	
개별 의무	• 기술문서 작성 여부 확인 및 보장 • 자신의 이름, 등록 상호 또는 상표, 연락할 수 있는 주소를 해당 시스템과 포장 또는 동봉된 문서에 표시 • 인증기관의 인증서, 사용 설명서, EU 적합성 선언서 등 관련 서류 10년간 보관	• 이미 시장에 출시된 이후 요건에 부합하지 않는 부분이 발견된 경우 리콜 조치 등 시정조치

5) 제한적 위험 AI 시스템

제한적 위험(limited risk) AI 시스템이란 사람과 상호작용하는 AI 시스템 중 의사결정 결과에 실질적인 영향을 주지 않고 건강, 안전 또는 기본권에 중대한 위해를 가할 위험이 없지만⁵⁵⁾ 딥페이크, 합성 콘텐츠를 생성하는 등으로 기만, 조작 등의 문제를 일으킬 수 있는 AI 시스템을 말한다.⁵⁶⁾ 이러한 제한적 위험 AI 시스템의 제공자와 배포자는 투명성 의무를 부담한다.

제공자는 AI 시스템과 사용자가 상호작용한다는 사실을 비롯하여 해당 AI 시스템의 작동 방식, 기능과 한계 등을 이해할 수 있도록 설계하고, 그러한 정보를 사용자에게 제공해야 한다. 배포자도 해당 AI 시스템을 통해 인위적으로 생성·조작된 콘텐츠라는 사실을 알리고 출처를 표시해야 한다.⁵⁷⁾

55) 심소연, 위 논문 6쪽

56) EU 인공지능법 제50조

57) 심소연, 위 논문 6쪽

6) 최소 위험 AI 시스템

최소 위험(minimal risk) AI 시스템이란 금지된 AI, 고위험 AI, 제한적 위험 AI에 속하지 않는 그 밖의 AI 시스템을 말한다. 이러한 최소 위험 AI 시스템에는 별도의 규제가 없다.

(라) 거버넌스

1) EU 차원의 거버넌스

첫째, AI 청이다. 2024년 1월 24일 EU 집행위원회의 결정으로 EU AI 청이 설치되었다. AI 청은 AI 시스템 및 AI 거버넌스의 구현, 점검을 지원하고 감독한다.⁵⁸⁾ AI 청은 EU 전체의 AI 분야 전문성과 역량 증진을 지원하는 역할 등을 수행한다.⁵⁹⁾

둘째, AI 위원회이다. EU AI 위원회는 회원국 당 1명의 대표로 구성되며, 각국의 대표들이 참석하여 의결권을 행사한다. 유럽데이터보호감독기구가 참관인으로 참가하며, AI 청은 의결권 없이 참여한다.⁶⁰⁾ AI 위원회는 EU 집행위원회와 회원국에 대한 자문 및 기술 지원, 전문지식과 모범사례의 수집 및 공유, 범용 AI 모델과 관련한 EU 인공지능법 자문, EU 인공지능법의 시행 및 일관되고 효과적인 적용과 관련한 문제에 대한 권고 및 의견을 제출하는 등의 활동을 할 수 있다.⁶¹⁾

셋째, 자문단이다. EU 인공지능법은 산업, 신생기업, 중소기업, 시민사회 및 학계를 포함한 이해관계자들로 구성된 자문단의 설치 근거를 두고 있다.⁶²⁾

58) EU 인공지능법 제3조 제47항

59) EU 인공지능법 제64조

60) EU 인공지능법 제65조

61) EU 인공지능법 제66조

62) EU 인공지능법 제67조

자문단은 AI 위원회나 EU 집행위원회의 요청이 있는 경우에는 권고 및 의견을 제출하는 등의 활동을 할 수 있으며 특정한 분야의 문제를 심의하기 위해 상임 또는 소위원회를 둘 수 있다.⁶³⁾

넷째, 과학 패널이다. 자문단 외에도 과학·기술 분야 지원을 위해 인공지능 분야에서 전문지식을 갖추고 AI 시스템 또는 범용 AI 모델의 제공자로부터 독립된 전문가로 구성된 독립 전문가 과학 패널의 설치 근거를 두고 있다.⁶⁴⁾ 인공지능 분야에 대한 전문지식을 갖춘 과학 패널은 범용 AI 모델 및 시스템에 관한 EU 인공지능법의 시행 및 집행 지원에 관한 조언, 인공지능청에 범용 AI 모델의 위험성 경고, 범용 AI 모델 및 시스템의 성능을 평가하는 도구 및 방법의 개발 지원 등의 활동을 할 수 있다.⁶⁵⁾

2) 회원국의 거버넌스

EU의 각 회원국은 관할기관을 지정하고 연락 창구를 두어야 한다. EU 인공지능법은 법의 적용 및 준수 여부를 감독하기 위해 적합성 평가 기구를 지정·모니터링·감독하는 통보기관⁶⁶⁾과 AI 시스템의 제공자 등의 적법요건 준수 여부를 감독하는 시장 감시기관⁶⁷⁾을 하나 이상 두어야 한다.⁶⁸⁾ 회원국은 관할기관이 효과적으로 업무를 수행할 수 있도록 기술, 금융, 인적 자원을 지원하고 특히 인공지능 기술, 데이터 및 데이터 계산, 개인정보 보호, 사이버 보안, 기본권, 건강 및 안전 위험에 대한 깊은 이해와 법률적 지식을 포함한 전문지식이 있는 충분한 수의 인력을 상시적으로 갖추어야 한다.⁶⁹⁾

63) EU 인공지능법 제67조

64) EU 인공지능법 제68조 제1항, 제2항

65) EU 인공지능법 제68조 제3항

66) EU 인공지능법 제3조 제29항

67) EU 인공지능법 제3조 제26항

68) EU 인공지능법 제70조

(마) 제재 수단

1) 시장감시기관의 시정명령

회원국의 시장감시기관은 법 제48조를 위반하는 방식으로 ‘CE 마크’를 부착하거나 이를 부착하지 아니한 경우, 적합성 선언서를 작성하지 않거나 정확하지 않게 작성한 경우 등에 대하여 시정을 요구할 수 있다.⁷⁰⁾ 시정 요구에도 불구하고 위와 같은 미준수가 계속되는 경우 회원국의 시장감시기관은 고위험 AI 시스템의 제공을 제한 또는 금지하거나 지체 없이 시장에서 회수 또는 철수하도록 하기 위한 적절하고 비례성 있는 조치를 취할 수 있다.⁷¹⁾ 시장감독 기관의 조치에 대해 상대방은 이의를 제기할 수 있다.⁷²⁾

2) 금전적 제재

가) AI 시스템 관련 의무 위반에 대한 제재

금지된 AI 시스템 관련 의무를 위반한 경우, 최대 3,500만 유로의 금전적인 제재를 받을 수 있다. 위반자가 기업인 경우 위 금액과 직전 회계연도의 전 세계 총 연간 매출액의 최대 7퍼센트 이하 금액 중 더 높은 금액의 제재금이 부과된다.

고위험 AI 시스템, 특정 AI 시스템 관련 의무를 위반한 경우로서 (a) 제16조에 따른 제공자의 의무, (b) 제22조에 따른 수권대리인의 의무, (c) 제23조에 따른 수입업자의 의무, (d) 제24조에 따른 유통업자의 의무, (e) 제26조에 따른 배포자의 의무, (f) 제31조, 제33조 제1항, 제33조 제3항, 제33조 제4항 또는 제34조에 따른 인증기관의 요건 및 의무, (g) 제50조에 따른 제공자 및 사용자가 준수해야 하는 투명성 의무를 위반한 경우에는 최대 1,500만 유로의

69) EU 인공지능법 제70조

70) EU 인공지능법 제83조

71) EU 인공지능법 제83조

72) EU 인공지능법 제85조

금전적인 제재를 받을 수 있다. 위반자가 기업인 경우 위 금액과 직전 회계연도의 전 세계 총 연간 매출액의 최대 3퍼센트 이하 금액 중 더 높은 금액의 제재금이 부과된다.⁷³⁾

인증기관이나 국가 관할기관의 요청에 부정확, 불완전 또는 허위 정보를 제공한 경우 최대 750만 유로의 금전적인 제재를 받을 수 있다. 위반자가 기업인 경우 위 금액과 직전 회계연도의 전 세계 총 연간 매출액의 최대 1퍼센트 이하 중 더 높은 금액의 제재금이 부과된다.⁷⁴⁾

단, 신생기업을 포함한 중소기업의 경우, 제99조 제3항, 제4항, 제5항에 언급된 비율 또는 금액 중 낮은 금액이 제재금이 부과된다.⁷⁵⁾

위와 같은 금전적 제재 부과 여부와 그 금액을 결정할 때는 (a) AI 시스템의 목적, 위반의 내용 및 그로 인한 결과의 성질, 심각성, 기간, 피침해자의 수와 손해의 정도, (b) 동일한 위반 행위에 대해 하나 이상의 회원국 또는 다른 시장감시기관이 금전적 제재를 부과하였는지 여부, (c) 동일한 위반 행위에 대해 다른 EU 법령 또는 회원국 법에 근거하여 다른 기관이 금전적 제재를 부과하였는지 여부, (d) 위반행위를 한 자의 규모, 연간 매출액 및 시장점유율, (e) 위반으로 인하여 직간접적으로 얻은 이익이나 회피한 손실 등 가중·감경 요소, (f) 위반을 치유하고 부정적 영향을 완화하기 위한 국가 관할기관과의 협력 정도, (g) 위반자가 실행한 기술적, 조직적 조치를 고려한 책임 정도, (h) 위반 사실이 국가 관할기관에게 알려진 방식 및 특히 위반자가 위반을 통지하였는지 여부와 그 범위, (i) 위반의 고의 또는 과실 정도 (j) 피침해자의 피해를 완화하기 위하여 위반자가 취한 조치 등 제반 정황을 모두 고려한다.⁷⁶⁾

73) EU 인공지능법 제99조 제4항

74) EU 인공지능법 제99조 제5항

75) EU 인공지능법 제99조 제6항

76) EU 인공지능법 제99조 제7항

나) 범용 AI 모델 관련 의무 위반에 대한 제재

범용 AI 모델 관련 의무 위반에 대하여는 AI 시스템 관련 의무 위반과 별도의 제재 규정을 두고 있다. EU 집행위원회는 (a) 범용 AI 모델의 제공자가 EU 인공지능법 관련 규정을 위반하거나, (b) 제91조에 따른 문서 또는 정보 요청을 준수하지 않거나 부정확, 불완전 또는 허위 정보를 제공한 경우, (c) 제93조에 따라 요구되는 조치를 준수하지 않은 경우, (d) EU 집행위원회가 제92조에 따른 평가를 실시하기 위해 필요한 시스템적 위험이 있는 범용 AI 모델 또는 범용 AI 모델에 대한 접근 권한을 제공하지 않는 경우에 대하여 위반자의 고의 또는 과실이 있다고 판단하는 경우 1천 5백만 유로 또는 직전 회계연도 전 세계 총 매출액의 3% 중 높은 금액의 금전적 제재를 부과할 수 있다.⁷⁷⁾

EU 집행위원회가 금전적 제재를 확정할 때는 비례성, 적절성의 원칙과 위반행위의 성질, 심각성, 기간을 고려하여야 한다.⁷⁸⁾ EU 집행위원회는 금전적 제재를 결정하기 전에 범용 AI 모델 또는 시스템적 위험이 있는 범용 AI 모델 제공자에게 예비조사결과를 통지하고 의견을 제시할 기회를 제공해야 한다.⁷⁹⁾ EU 재판소는 위와 같은 EU 집행위원회의 결정에 대해 제한 없는 심사권을 가지며, 부과된 금전적 제재를 취소, 감액 또는 증액할 수 있다.⁸⁰⁾

(바) AI 혁신을 지원하기 위한 제도적 장치

EU 인공지능법은 AI 산업 진흥 내지는 AI 기술 혁신을 지원하기 위해 다음과 같은 제도적 장치에 대한 근거를 마련하고 있다.

77) EU 인공지능법 제101조 제1항

78) EU 인공지능법 제101조 제1항

79) EU 인공지능법 제101조 제2항

80) EU 인공지능법 제101조 제5항

1) AI 규제 샌드박스

EU 인공지능법은 AI 규제 샌드박스 제도에 관한 근거를 마련하고 있다.⁸¹⁾ AI 규제 샌드박스는 AI 시스템의 제공자 또는 장래의 제공자에게 한시적으로 관할 기관의 감독 아래 샌드박스계획에 따라 혁신적인 AI 시스템을 개발, 훈련, 검증 및 시험할 수 있도록 하는 것이다.⁸²⁾ EU 집행위원회는 AI 규제 샌드박스의 설립 및 운영을 위한 기술 지원, 자문 및 도구를 제공할 수 있다. 또한, 유럽데이터보호감독관은 EU 기관, 단체, 사무소 및 에이전시를 위한 AI 규제 샌드박스를 설립할 수 있고 관할 당국으로서 역할과 업무를 수행할 수 있다.⁸³⁾ AI 규제 샌드박스 외에도 회원국은 실제(real world) 조건에서 고위험 AI 시스템 시험의 실시를 승인할 수 있다.⁸⁴⁾

2) 신생기업을 포함한 중소기업에 대한 혁신 지원 방안

EU 인공지능법은 신생기업을 포함한 중소기업에 대한 혁신 지원 방안을 마련하고 있다.⁸⁵⁾ EU 회원국은 EU에 등록된 사무소 또는 지점을 둔 신생기업을 포함한 중소기업에 AI 규제 샌드박스 우선 접근권을 제공, EU 인공지능법에 대한 인식을 제고 및 훈련, EU 인공지능법에 관한 자문을 받을 수 있는 소통 채널 마련, 표준 개발과정에서 참여 촉진 등의 조치를 취해야 한다.⁸⁶⁾ 그밖에 제43조에 따른 적합성평가를 받는 경우 수수료를 할인하며⁸⁷⁾ EC 권고 제2003/361호에서 정하는 영세기업⁸⁸⁾은 제17조에 따른 품질관리를 간소화된 방식으로 이행할 수 있도록 하고 있다.⁸⁹⁾

81) EU 인공지능법 제57조 내지 제61조

82) EU 인공지능법 제3조 제55항

83) EU 인공지능법 제57조

84) EU 인공지능법 제60조

85) EU 인공지능법 제1조 제2항 (g), 제62조

86) EU 인공지능법 제62조

87) EU 인공지능법 제62조

88) Commission Recommendation of 6 May 2023 concerning the definition of micro, small and medium-sized enterprises(2003/361/EC).

(2) 시사점

EU 인공지능법은 위험 기반 접근법, 범용AI 모델 규제, 강한 행정규제, EU/회원국 간 다층적 거버넌스, 혁신지원 제도 등 인공지능 규제를 위한 다양한 제도적 도구를 담고 있으며, 세계 최초로 통합적이고 옴니버스 형태의 인공지능 규제를 담고 있으며, 글로벌 논의를 주도하고 있다. GDPR의 사례와 같이 전 세계에 영향을 줄 것이라는 전망도 있으며, 실제로 우리나라를 비롯한 여러 국가, 미국의 일부 주에서 EU 인공지능법과 유사한 형태의 입법을 추진하고 있다.

EU와 같은 통합적인 접근방식은 인공지능의 위험성을 사전에 예방한다는 장점이 있지만, 한편으로 이질적인 분야를 한데 묶어 규율하다 보니 법률이 복잡하고 아직 명확하지 않은 규제 영역이 많아 그 실효성에 의문을 갖는 경우가 있으며, EU의 규제가 전 세계에 영향을 주는 소위 ‘브뤼셀 효과’가 제한적일 것이라는 전망도 있다.⁹⁰⁾

EU 인공지능법은 다양한 논의를 거쳐 좋은 제도를 규정하고 있지만, 무조건적으로 수용하기보다는 그 특징과 장점을 잘 살리되 미국 등 다른 국가의 사례와 우리나라의 특수성 등을 고려하여 우리나라의 고유한 제도로 발전시키는 방안을 모색할 필요가 있다.

나. 미국

(1) 미국의 AI관련 입법 동향

미국 의회는 포괄적인 연방 법률로 「국가 인공지능 이니셔티브법(National

89) EU 인공지능법 제63조

90) Alex Engler, The EU AI Act will have global impact, but a limited Brussels Effect, Brookings Research, 2022. 6. 8.

Artificial Intelligence Initiative Act of 2020)», 「정부인공지능법(AI in Government Act of 2020)», 「미국 인공지능진흥법(Advancing American AI Act)», 「인공지능 교육법(AI Training Act)」등을 제정하였다. 미국 연방정부 차원에서는 연방 정부 사업에 관한 법률들은 제정되어 있으나, EU의 인공지능법에 상응하는 포괄적인 규제 법률이 제정되어 있지는 않다.⁹¹⁾

반면, 개별 주(州)에서는 상당히 강한 규제가 담긴 법률이 제정되고 있다. 대표적으로 캘리포니아 주, 콜로라도 주의 법률을 들 수 있다.

(2) 행정명령 제14110호

(가) 배경

바이든 행정부는 2023년 10월에 안전하고 보안이 보장되며 신뢰할 수 있는 인공지능의 개발과 사용(Safe, Secure and Trustworthy Development and Use of Artificial Intelligence)에 관한 행정명령 제14110호(이하 ‘행정명령 제14110호’)를 발표하였다.⁹²⁾

이에 앞서 트럼프 행정부는 2019년 ‘AI분야 미국의 리더십 유지를 위한 명령(Maintaining American Leadership in Artificial Intelligence)’(행정명령 제13859호)에 서명하고,⁹³⁾ 2020년에는 ‘연방정부에서 신뢰할 수 있는 인공지능 사용 촉진에 관한 행정명령(Promoting the Use of Trustworthy Artificial Intelligence

91) 오유빈, 「미국의 미국의 인공지능 입법 현황과 바이든 행정부의 행정명령」, 국회도서관, 2024. 1.

92) 미국 백악관 보도자료 <<https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>> 최종방문일 2024. 9. 11.>

87) 미국 연방관보 <<https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence>> 최종방문일 2024. 9. 11.>

in the Federal Government)’(행정명령 제13960호)를 발표하였다. 행정명령 제13859호는 인공지능 분야에서 미국의 기술 우위를 공고화하여 경제성장과 국가안보를 강화하기 위한 조치로 해석되며, 트럼프 행정부 당시 미국 정부의 접근방식은 시장을 통제하고 규율하기 보다는 정부가 시장의 현황을 정확히 파악하여 시장에 대한 이해를 높이고 인공지능 활용을 최대화할 수 있는 예산 및 규정 정비에 강조점을 두었다고 이해된다.

반면 바이든 행정부의 행정명령 제14110호는 AI의 잠재성은 극대화하면서 서도, AI의 위험성을 규제하는 내용들을 담고 있다. 트럼프 행정부의 행정명령 제13859호와 바이든 행정부의 행정명령 제14110호가 선언하고 있는 목표를 비교하여 보면, “AI연구개발과 적용에서 미국의 선도적 지위 유지와 강화”와 “막대한 위험 완화”에 각기 강조점이 있는 행정명령의 방향성 차이를 확인할 수 있다.

행정명령 제13859호 제1조(정책 및 원칙)	행정명령 제14110호 제1조(정책 및 원칙)
인공지능은 미국 경제의 성장을 촉진하고 경제 및 국가안보를 강화하며 삶의 질을 개선할 것을 약속한다. 미국은 AI 연구개발 및 적용 분야에서 세계적인 선두 주자이다. AI 분야에서 미국의 지속적인 선도적 지위는 미국의 경제 및 국가안보를 유지하고 미국의 가치, 정책 및 우선순위에 부합하는 방식으로 AI의 세계적 진화를 형성하는 데 가장 중요하다. 연방정부는 AI 연구개발을 촉진하고 AI 관련 기술의 개발과 적용에 대한 미국 국민의 신뢰를 촉진하며 직장에서 AI를 사용할 수 있는 인력을 교육하고 전략적 경쟁자와 적대적인 국가들이 획득을 시도하지 않도록 미국 AI 기술 기반을 보호하는 데 중요한 역할을 수행한다. AI분야에서 미국의 선도적 지위를 유지하려면 기술과 혁신의 발전을 촉진함과 동시에 미국의 기술, 경제와 국가안보, 시민의 자유, 사생활 및 미국의 가치를 보호하고 외국 동반자나 동맹국과의 국제협력과 산업협력을 강화하기 위해서는 공동의 노력이 필요하다. 미국 정부가 5대 원칙에 따라 인도된 연방정부 전략인 미국 AI 선도계획을 통해 마연구개발과 적용에서 미국의 과학적, 기술적 및 경제적 선도적 지위 위치를 유지하고 강화하는 것이 미국 정부의 정책이다.	인공지능은 엄청난 잠재력과 동시에 위험을 내재하고 있다. AI는 책임감 있게 사용한다면 시급한 과제를 해결하고 동시에 세계를 더 번영하고 생산적, 혁신적이며 안전한 곳으로 만들 수 있는 잠재력을 가지고 있다. 반면에 AI를 무책임하게 사용한다면 사기, 차별, 편견 및 허위 정보와 같은 사회적 피해를 악화시키고 근로자를 대체하며 그 권한을 박탈하고 경쟁을 억제하며 국가 안보에 위협을 초래할 수 있다. AI를 선의의 목적으로 활용하고 그에 따른 다양한 이점을 실현하기 위하여는 그에 따른 막대한 위험을 완화하여야 한다. 이를 위해 정부, 민간 부문, 학계 및 시민사회를 포함한 사회 전체의 노력이 필요하다. 우리 행정부는 안전하게 책임감 있는 AI 개발과 활용을 관리하는 데 최우선순위를 두고 있으며 따라서 이를 위해 연방정부 전반에 걸친 조정된 접근방식을 추진하고 있다. AI역량이 급속하게 발전함에 따라 미국은 안보, 경제, 사회를 위해 주도권을 가져야 한다. AI는 AI개발자, 사용자 및 구축된 데이터의 원칙을 반영한다. 본인은 급변의 시대에도 미국이 반영할 수 있었던 데에는 이상의 힘, 우리 사회의 근간, 우리 국민의 창의성, 다양성과 품위가 큰 역할을 했다고 단언한다. 이는 우리가 오늘날에도 다시 한 번 성공할 수 있는 이유이기도 하다. 우리는 모든 사람을 위한 정의, 안보 및 모든 기회를 보장하기 위하여 AI를 활용할 수 있는 그 이상의 능력을 갖추고 있다.

(나) 개괄

행정명령 제14110호는 8개 정책영역을 구분하고 미국 연방 기관이 준수하여야 하는 구체적인 세부지침을 명시하고 있다. 63쪽 분량의 행정명령 제

14110호의 8개의 지도 원칙은 ① 안전 및 보안의 보장(제4조), ② 혁신 및 경쟁의 촉진(제5조), ③ 근로자에 대한 지원(제6조), ④ 형평성 및 시민권의 진전(제7조), ⑤ 소비자, 환자, 승객 및 학생에 대한 보호(제8조), ⑥ 프라이버시 보호(제9조), ⑦ 연방기관의 AI사용의 발전(제10조), ⑧ 해외에서 미국 선도자로서의 강화(제11조)로 구체화된다.

행정명령 제14110호는 AI 개발자들이 안전 테스트 결과와 중요한 기술 정보들을 제품 출시 전 미국 정부와 사전 공유하도록 하고, 국가안보·경제·공공보건 등 주요 분야에서 AI 모델을 개발하는 회사는 안전테스트 결과를 정부에 보고하도록 하였다. 또한 기업들은 개인정보 규제를 위한 지침을 만들고, 노동 시장에 미칠 영향에 대해 보고서를 작성하고 AI로 일자리를 잃는 근로자를 지원할 방법을 연구해야 하며, 형평성 있고 책임감 있는 AI 사용법을 연구하도록 하였다. 행정명령은 행정기관 및 그 소속 공무원에 대하여 구속력이 발생하고 일반 사인에 대해 직접 효력은 없으나, 행정기관이 제시하는 가이드라인이나 행정기관에 발주하는 업체에 대한 의무 사항은 민간업체에 영향을 미치게 되므로, 행정명령은 미국의 AI 규제체계의 기본적인 틀로 이해할 수 있다. 행정명령 제14110호는 직접적인 금지 규정을 두고 있는 EU의 규제 체계와는 달리 정부가 가이드라인을 제시하고 인공지능의 시스템 결합이나 취약점은 실증적 접근을 통하여 확인하도록 하는 체계를 취하고 있어 EU보다 시장주의적인 규제 방식을 취하고 있다고 평가된다.

(다) 주요 세부내용

1) 안전 및 보안의 보장

명령 시행일로부터 270일 이내에 AI 안전 및 보안을 위한 지침, 모범사례 확립 및 AI 개발자 특히 ‘이중 용도 기초 모델’⁹⁴⁾ 개발자가 AI 레드티밍⁹⁵⁾ 시험을 수행할 수 있도록 적절한 절차와 지침을 수립해야 한다.

이중용도 인공지능 모델을 개발하고자 하는 기업은 이중용도 기초 모델의 훈련, 개발, 생산과 관련한 활동, 레드팀 시험 결과 및 안전 개선 조치 등과 관련한 보고서 등을 연방정부에 제공해야 한다.

잠재적인 대규모 컴퓨팅 클러스터를 획득, 개발, 소유하는 회사는 클러스터의 존재, 위치 등을 보고해야 한다.

명령 시행일로부터 240일 이내에 상무부 장관은 관련 기관장과 협의하여 합성 콘텐츠 발견, 합성 콘텐츠 라벨링, 콘텐츠 인증 및 출처 추적을 위한 보고서를 제출하고 이후 180일 이내에 디지털 콘텐츠 인증 및 합성 콘텐츠 발견 조치를 위한 지침을 개발해야 한다. 또한 지침 개발 180일 이내에 미국 정부의 공식 디지털 콘텐츠 무결성에 대한 대중의 신뢰를 강화하기 위하여 해당 기관이 생산하거나 게시하는 콘텐츠에 라벨링하고 인증하기 위한 지침을 발행해야 한다.

2) 혁신 및 경쟁의 촉진

AI 및 기타 중요한 신규 기술 분야 연구자 등 인재 확보를 위하여 비자 및 이민 제도를 개선해야 한다. 아울러, AI 시장의 공정경쟁을 보장하고 소비자, 근로자 보호를 위한 규범 및 규칙을 마련해야 한다.

94) 이중 용도 기초 모델(dual-use information model)은 국가 경제안보, 국가 공공 보건 또는 안전에 심각한 위험을 초래하는 작업에서 높은 수준의 성능을 발휘하거나 쉽게 수정될 수 있는 모델을 의미하며, 구체적으로 ① 대량살상무기의 설계, 합성, 획득 또는 사용에 대한 비전문가의 진입 장벽을 실질적으로 낮추는 것, ② 자동화된 취약성 식별과 악용을 통해 잠재적 사이버 공격의 대상에 대한 강력한 공격적 사이버 작전을 가능하도록 하는 것, ③ 기만이나 은폐 수단을 통해 인간의 통제나 감독을 회피할 수 있도록 허용하는 것을 말한다[제3조(k) 참조].

95) AI 레드티밍(AI red-teaming)은 AI시스템의 결함과 취약성을 식별하기 위한 구조화된 시험을 말하며, AI 시스템의 유해하거나 차별적인 출력, 예상하지 못하거나 바람직하지 아니한 시스템 작동, 한계 또는 시스템의 오용과 관련된 잠재적 위험 등과 같이 AI시스템의 결함과 취약점을 식별한다[제3조(d)참조]

명령 시행일로부터 120일 이내에 미국특허청장은 생성형 인공지능을 발명 과정에서 사용하는 것과 관련한 발명자 자격 지침을 발표하여야 한다. 명령 시행일로부터 270일 또는 미국의회도서관 저작권청이 AI 연구를 공표한 후 180일 이내 중으로 저작권과 관련하여 AI를 사용하여 제작된 저작물의 보호 범위와 AI 훈련시 저작권 저작물의 처리를 포함하여 저작권과 관련한 문제를 다루는 내용으로 대통령에게 권고사항을 발행하여야 한다.

3) 근로자에 대한 지원

경제자문위원회장은 AI가 노동시장에 미치는 영향에 관한 보고서를 준비하여 대통령에게 제출하여야 하고, 노동부 장관은 AI 도입으로 대체된 근로자를 지원하는 기관의 능력을 분석한 보고서를 대통령에게 제출해야 한다.

노동부 장관은 명령 시행일로부터 180일 이내에 피고용인의 복지에 대한 AI의 잠재적 피해를 완화하고 잠재적 이점을 극대화할 수 있는 원칙과 모범 사례를 고용주를 위하여 배포하고 공표해야 한다.

4) 형평성 및 시민권의 진전

명령 시행일로부터 365일 이내에 법무부 장관은 형사사법제도에서 AI 사용을 다루는 보고서를 대통령에게 제출해야 한다.

명령 시행일로부터 265일 이내에 노동부 장관은 고용에서 AI로 인한 불법적 차별을 방지하기 위해 고용시 차별 금지에 관한 연방 계약 업체에 대한 지침을 공표해야 한다.

5) 소비자, 환자, 승객 및 학생에 대한 보호

규제기관은 사기, 차별, 개인정보 보호 위협으로부터 미국 소비자를 보호하고 AI 사용으로 인한 위험을 해결하기 위하여 규칙 제정을 고려하고 기존 규정과 지침이 AI에 적용되는 부분을 강조하거나 명확히 하도록 권장한다.

6) 프라이버시 보호

행정관리예산국장은 기관이 조달한 상업이용가능정보(commercially available information, CAI), 특히 개인식별정보가 포함된 상업이용가능정보, 데이터 중 개상으로부터 조달한 상업이용가능정보 및 판매자를 통해 간접적으로 조달되고 처리된 상업이용가능정보를 평가하고 식별하기 위한 조치를 실시해야 한다.

개인식별정보가 포함된 상업이용가능정보의 수집, 처리, 유지, 사용, 공유, 전파, 처리와 관련된 기관 표준 및 절차를 평가해야 한다.

(라) 시사점

행정명령 제14110호는 연방정부를 중심으로 사회의 전 분야에서 AI의 위험과 기회에 대응하기 위하여 계획을 수립하고 지침을 마련하도록 하고 있으며, 구체적인 이행기간을 못 박고 있다는 점도 주목할 만하다. 행정명령 이행 주체는 행정기관이지만, 이중용도 인공지능 모델 개발사와 같이 국가 경제, 안보, 공공보건, 안전에 위험을 크게 초래하는 분야의 경우 보고 의무를 부담하게 하였으며, 특히 전 세계적으로 문제되고 있는 합성콘텐츠와 관련하여서도 디지털 콘텐츠 무결성과 관련한 미국 정부의 라벨링 인증 지침이 확정되면 민간 영역에도 큰 영향을 미칠 것으로 예상된다.

(마) 후속 조치⁹⁶⁾

미국 백악관 예산관리국(Office of Management and Budget, 이하 ‘OMB’)은 2024년 3월에 ‘Advancing Governance, Innovation, and Risk Management for Agency Use of Artificial Intelligence’라는 제목의 메모랜덤을 발표하여 각

96) 채은선, 「美, AI 행정명령(2023.10.30.)의 주요 내용 및 이행 현황」, 디지털 법제 Brief, 2024. 7.

정부기관이 준수하여야 할 정책을 지시했다.⁹⁷⁾ 이에 따르면 정부기관은 ① 60일 이내에 AI최고 책임자(Chief AI Officer, CAIO)를 임명하고, ② 최소 연 1회 AI 사용 사례를 조사하여 목록을 작성하여 OMB에 제출하고 공개해야 하고, ③ AI 사용 사례가 ‘안전에 영향을 미치는 AI’(Safety-impacting AI)와 ‘권리에 영향을 미치는 AI’(Rights-impacting AI)에 해당하는지 식별하고 이러한 사용이 초래하는 위험에 대한 세부 정보와 위험 관리 방안을 보고해야 하는 등의 내용을 담고 있다.

‘안전에 영향을 미치는 AI’란, 그 결과물이 행동을 유발하거나 인간의 생명과 웰빙, 기후나 환경(돌이킬 수 없거나 중대한 환경 피해 포함), 주요 기반 시설, 전략적 자산이나 자원의 안전에 중대한 영향을 미칠 수 있는 결정의 주요 근거로 작용하는 AI를 지칭한다. ‘권리에 영향을 미치는 AI’는 그 결과물이 특정 개인이나 단체에 관한 결정이나 행동의 주요 근거로 작용하며 개인이나 단체의 시민권, 시민의 자유, 프라이버시, 평등한 기회, 중요한 정부 자원이나 서비스에 대한 접근 또는 신청 능력과 관련하여 중요한 영향을 미치는 AI를 의미한다. AI를 ‘안전’과 ‘권리’에 영향을 미치는 AI에 해당하는지 평가하고, ‘안전’과 ‘권리’에 영향을 미치는 경우 영향평가를 의무화하고 테스트와 모니터링 의무를 부여하는 것으로, 규제 대상인 AI의 분류 기준이나 규제의 정도와 관련하여 시사점을 제공한다.

91) 미국 백악관 성명 <<https://www.whitehouse.gov/briefing-room/statements-releases/2024/03/28/fact-sheet-vice-president-harris-announces-omb-policy-to-advance-governance-innovation-and-risk-management-in-federal-agencies-use-of-artificial-intelligence/>> 최중방문일 2024. 9. 11.>

(3) 콜로라도 인공지능법(Colorado Artificial Intelligence Act, CAIA, SB 24-205⁹⁸⁾)

(가) 개요

콜로라도 주는 2026년 2월 1일부터 발효되는 콜로라도 인공지능법(CAIA)을 제정하여 특정 인공지능(AI) 시스템의 개발 및 배포에 대한 포괄적인 규제를 마련했다. 이 법은 특히 고위험 AI 시스템에 중점을 두고 있으며, 교육, 고용, 금융, 주택, 의료, 법률 서비스와 관련된 중요한 결정에 영향을 미치는 AI 시스템을 규제한다. CAIA는 AI 시스템이 보호된 특성을 기반으로 개인 또는 집단을 차별하지 않도록 하는 알고리즘 차별 방지를 목표로 하고 있다. 알고리즘 차별은 연령, 피부색, 장애, 민족, 성별, 군필 여부 등 보호된 특성을 기준으로 개인이나 그룹을 불리하게 대우하는 행위로 정의된다.

CAIA는 고위험 AI 시스템의 개발자와 배포자에게 일련의 의무를 부과한다. ‘개발자’란 콜로라도에서 사업을 하는 개인 또는 단체로서 고위험 AI 시스템을 개발하거나 의도적으로 실질적으로 수정하는 사람을 말한다. ‘배포자’는 콜로라도에서 사업을 하는 개인 또는 단체로 고위험 AI 시스템을 배포하는 사람을 말한다. 개발자와 배포자 모두 고위험 AI 시스템의 사용으로 인해 발생하는 알고리즘 차별에 따른 알려진 또는 합리적으로 예측 가능한 위험으로부터 소비자를 보호하기 위해 합리적인 주의를 기울여야 한다.

98) Stuart D. Levi et al, “Colorado’s Landmark AI Act: What Companies Need To Know”, Skadden Publication, 2024. 6. 24.; Colorado General Assembly <<https://leg.colorado.gov/bills/sb24-205> 최종 방문일 2024. 9. 1.>

(나) 주요 내용

1) 고위험 AI 시스템의 정의

고위험 AI 시스템은 배포 시 ‘결과적 결정(consequential decisions)’, 즉 다음 항목들의 제공 또는 비용에 중대한 영향을 미치는 결정을 내리거나 결정에 중요한 요소가 되는 시스템으로 정의한다.

- 교육 등록 또는 교육 기회
- 고용 또는 고용 기회
- 금융 또는 대출 서비스
- 필수 정부 서비스
- 의료 서비스
- 주택
- 보험
- 법률 서비스.

2) 개발자의 책임

개발자는 고위험 AI 시스템의 사용과 관련된 예측 가능한 위험과 그 용도를 설명하는 문서를 제공해야 한다. 이 문서에는 다음의 내용이 포함된다.

- 시스템의 성능을 평가한 방법
- 알고리즘 차별의 영향을 완화하기 위해 취한 조치
- 데이터 거버넌스 조치(데이터 소스의 적합성, 편향 가능성 및 적절한 완화를 검토하는 데 사용된 조치 포함)
- 시스템의 의도된 결과물
- 개인이 시스템을 사용하거나 사용하지 않고 모니터링하는 방법

아울러 개발자는 자신이 개발한 고위험 AI 시스템에 대한 공개 사항을 명확하게 표시하고, 알고리즘 차별의 위험을 관리하는 방법을 설명해야 한다.

또한, 알고리즘 차별을 유발할 가능성이 높다는 사실을 발견하면 90일 이내에 법무부 장관과 관련 배포자에게 이를 통보해야 한다.

3) 배포자의 책임

고위험 AI 시스템이 소비자에 대한 중요한 결정을 내리는 경우, 배포자는 소비자에게 AI 시스템의 사용 여부를 사전에 통지해야 한다. 또한, 불리한 결정을 내린 경우, 그 이유와 사용된 데이터, 그리고 소비자가 잘못된 데이터를 수정할 기회를 제공해야 한다.

또한 고위험 AI 시스템이 소비자에게 불리한 결정에 도달하는 경우, 배포자는 소비자에게 다음 사항에 대한 설명을 제공해야 한다.

- 결과적 결정의 이유
- 고위험 AI 시스템이 해당 결정에 기여한 정도.
- 시스템에서 처리한 데이터의 유형 및 해당 데이터의 출처

소비자에게 사용된 잘못된 개인 데이터를 수정할 수 있는 기회와 불리한 결정에 이의를 제기하고 사람의 검토를 요청할 수 있는 기회를 제공해야 한다.

배포자는 웹사이트에서 모든 정보를 명확하고 쉽게 이용할 수 있도록 다음의 사항들을 공개해야 한다.

- 현재 배포된 고위험 AI 시스템의 유형
- 발생할 수 있는 알고리즘 차별의 알려진 또는 합리적으로 예측 가능한 위험을 관리하는 방법
- 배포자가 AI 시스템과 관련하여 수집하고 사용하는 정보의 성격, 출처 및 범위

소비자와 상호 작용하는 AI 시스템을 제공하는 배포자는 (고위험이 아니더라도) 합리적인 사람이 알 수 있는 경우가 아니라면 소비자에게 AI 시스템과 상호 작용하고 있다는 사실을 공개해야 한다.

법무부 장관이 고위험 AI 시스템 배포자에 대해 집행 조치를 취하는 경우, 배포자가 다음 기준을 충족하는 경우 CAIA에서 요구하는 합리적인 주의를 기울인 것으로 반박할 수 있는 추정이 인정된다.

- 배포자가 알고리즘 차별의 위험을 식별, 문서화 및 완화하는 데 사용하는 원칙, 프로세스 및 인력을 명시하는 최신 위험 관리 정책 및 프로그램을 보유하고 있다.
 - CAIA는 필수 위험 관리 프로그램의 벤치마크로 미국 국립표준기술연구소(NIST)의 “인공 지능 위험 관리 프레임워크”를 인용하지만, 다른 국내 또는 국제 프레임워크(예: ISO/IEC 42001) 또는 콜로라도 주에서 지정한 기타 프레임워크도 허용한다.
 - 또한 CAIA는 배포자의 규모와 복잡성, 고위험 AI 시스템의 성격과 범위, 처리되는 데이터의 민감도 및 양을 고려하여 정책이 합리적이어야 한다고 지적한다.
- 적어도 매년, 그리고 고위험 AI 시스템을 크게 수정한 후 90일 이내에 재평가하는 영향 평가를 수행한다.
 - 평가에는 다음 사항이 포함되어야 한다.
 - 고위험 AI 시스템의 목적, 의도된 사용 사례 및 이점에 대한 설명
 - 시스템이 알고리즘 차별의 알려진 또는 합리적으로 예측 가능한 위험을 초래하는지 여부와 그러한 위험이 있는 경우 그러한 위험의 성격 및 위험 완화를 위해 취한 조치에 대한 분석
 - 시스템이 입력으로 처리하고 출력으로 생성하는 데이터 범주에 대한 설명
 - 시스템의 성능과 알려진 한계를 평가하는 데 사용되는 지표.
 - 고위험 AI 시스템에 대한 연례 검토를 수행하여 알고리즘 차별을 유발하지 않는지 확인한다. 따라서 배포자는 법무부 조치 이전에도 이러한 정책과 절차를 갖추고 있는지 확인해야 한다.

아울러, 개발자와 마찬가지로 고위험 AI 시스템의 배포자는 알고리즘 차별을 발견한 날로부터 90일 이내에 법무부 장관에게 알려야 한다.

4) 면제

다음 기준을 충족하는 배포자는 고위험 AI 시스템을 사용하여 소비자에 대한 결과적 결정이 내려졌다는 사실을 소비자에게 통지해야 하는 의무를 제외하고는 CAIA의 적용을 면제받는다.

- 배포자의 직원 수가 50명 미만
- 배포자가 자체 데이터를 사용하여 시스템을 학습시키지 않음
- 배포된 시스템이 의도된 목적(개발자가 지정하고 CAIA에 따라 요구되는)에 맞게 사용됨
- 배포자는 개발자가 제공한 영향 평가를 소비자에게 제공함

또한 연방 기관의 승인을 받은 시스템, 건강보험 이동성 및 책임에 관한 법률의 적용을 받는 기관 및 기존 법률과 규정의 적용을 받는 기타 특정 기관에 대해서는 예외가 적용된다.

5) 법무부 장관의 집행

법무부 장관은 CAIA를 집행할 독점적 권한을 가진다. 개발자 또는 배포자는 CAIA에 명시된 요건이 충족되었음을 입증할 책임을 부담한다.

집행 조치를 받는 개발자와 배포자는 다음과 같은 경우 적극적 방어권을 갖는다.

- 자체 내부 검토 또는 “레드팀”(즉, 위험을 발견하기 위한 내부 프로세스에 따라) 또는 외부 피드백의 결과로 위반을 해결한 경우
- NIST AI 위험 관리 프레임워크의 최신 버전, 기타 국내 또는 국제적으로 인정받는 AI 위험 관리 프레임워크 또는 법무부 장관이 선택한 프레임워크를 준수한 경우

한편, 법무부 장관은 CAIA에 명시된 규정을 시행하고 집행하기 위해 필요에 따라 규칙을 공포할 권리가 있지만 의무는 아니다. 이러한 규정에는 다음과 같은 사항이 포함될 수 있다.

- 개발자에게 요구되는 세부 문서
- 소비자 고지 및 공시의 내용
- 위험 관리 정책 및 영향 평가의 내용

(4) 캘리포니아 인공지능법(The Safe and Secure Innovation for Frontier Artificial Intelligence Models Act, SB-1047)⁹⁹⁾

(가) 개요

캘리포니아 주 하원은 2024년 8월 28일에 AI 규제법인 “The Safe and Secure Innovation for Frontier Artificial Intelligence Models Act(첨단AI모델을 위한 안전 및 보안 혁신법, 이하 “AI법)”을 찬성41표, 반대9표로 통과시켰다. SB-1047이라고 불리는 이 법안은 캘리포니아 상원의원인 스콧위너(Scott Wiener, 민주당)가 발의하여 하원으로 넘어온 것으로서, 상원을 통과해야 하지만 통과될 것으로 예상되며, 현재 캘리포니아 주지사인 개빈뉴섬(Gavin Newsom)의 서명만을 남겨놓고 있다.

이 법안은 거대AI모델에 대한 강력한 규제를 담고 있어 이미 미국 내 AI 사업자, 연구진, 여러 이해관계자들은 법안에 찬반이 갈려 있다.

99) California Legislative Information<https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240SB1047 최종 방문일 2024. 9. 1.>

(나) 주요 내용

1) 대상AI 모델(Covered AI models)

2027년 1월 1일 이전에는 ① 1억 달러 초과 개발비용 소모 및 전체 10^{26} FLOP 등 초과 컴퓨팅 연산량 이용 또는 ② 1천만 달러 초과 비용 소모 및 전체 $10^{25} \times 3$ FLOP 등 컴퓨팅 연산량 초과 파인튜닝 이용하는 AI 모델이 대상AI 모델이 된다.

2027년 1월 1일 이후에는 비용은 위와 동일하고(인플레이션 조정 가능), 컴퓨팅 연산량 기준은 캘리포니아주의 Government Operations Agency에서 추후 결정할 예정이다.

2) 안전평가 관련

AI로 인한 핵심 피해(Critical harm)를 ① 대량학살을 일으킬 수 있는 화학, 생물학, 방사능 또는 핵무기의 제작 또는 이용, ② 핵심기반 시설 목표 사이버공격을 통한 대량학살 또는 최소 5억 달러 이상의 피해, ③ AI모델의 관여로 다음 두 가지를 일으키는 행위[(i) 인간의 감독, 개입등의 제한, (ii) 사망, 중상해, 재산상피해, (인간이라면) 고의·중과실로 인해 형사처벌되는 행위, 그러한 행위의 교사·방조]를 통해 발생한 대량학살 또는 최소 5억 달러 이상의 피해, ④ 이와 유사한 공공안전 및 보안에 대한 막대한 피해로 규정한다.

이러한 핵심 피해를 막기 위해 개발자들은 문서화된 안전 및 보안프로토콜을 수립해야 한다.

대상AI 모델을 이용 또는 공개하기 전에 개발자들은 모델이 핵심 피해를 일으킬 가능성을 평가하고, 평가결과 기록을 보관하며, 적절한 안전조치를 도입해야 한다.

3) 클러스터 정책

컴퓨팅 클러스터란 데이터센터의 100GB/초를 초과하는 네트워크로 연결되고, 이론적으로 최소 10^{20} 이상의 초당 FLOP 등의 최대 컴퓨팅 연산량을 가지며, AI학습을 위해 이용될 수 있는 기계집합을 의미한다.

컴퓨팅 클러스터 운영자는 고객이 대상AI 모델을 학습하기 위해 충분한 컴퓨팅 자원을 이용하는 경우에 대비한 정책 및 절차를 구비해야 한다.

해당 정책 및 절차는 컴퓨팅 클러스터 운영자가 잠재고객의 기본 인적정보와 사업목적 등의 정보를 확보하고, 잠재적 고객의 클러스터 이용의도가 대상AI 모델 학습용인지를 평가하며, 고객의 통제하에 모델을 학습 또는 작동시키는 전체 자원의 즉각중단 기능 시행 등의 내용을 포함한다.

4) 중단 기능

대상AI 모델학습 전에 개발자들은 학습을 포함하여 모든 모델기능을 중단시키는 전체 중단기능(full shutdown)을 구현해야 한다.

5) 법무 장관 소송

이 법안 내용의 위반으로 사망, 중상해, 재산상피해, 도난, 재산착복 또는 임박한 공공안전 위협 등에 대해 법무 장관(Attorney General)이 민사소송 제기를 할 수 있다.

법무 장관은 징벌적 손해배상을 포함한 금전적 손해배상, 민사페널티, 금지명령 구제, 선언적 구제 등을 구할 수 있다.

6) 감사 및 보고

2026년부터 대상AI 모델 개발자들은 매년 준법감시를 위한 독립적 감사를 위해 제3의 기관으로부터 감사를 진행해야 한다.

대상AI 모델 개발자들은 민감정보를 삭제한 안전 및 보안프로토콜과 감사 보고서를 공개하며, 법무 장관 요청시 법무 장관에게 무삭제본을 제공해야 한다.

대상AI 모델 개발자들은 준수 성명서(compliance statements)를 매년 제출하고, 안전사고 발생시 법무 장관에게 72시간 내 보고해야 한다.

7) 내부고발자 보호

내부 직원의 법률 미준수 고발 또는 AI 모델의 비합리적 위험폭로에 대한 대상AI 모델 개발자들의 저지 또는 보복을 금지한다.

다. 영국

(1) AI 백서

영국 과학혁신기술부(DSIT)는 2023년 3월에 ‘AI규제에 대한 혁신적 접근(A pro-innovation approach to AI regulation)’이라는 AI 백서를 발표하여, ‘허용할 수 없는 위험(unacceptable risk)’을 가진 AI는 엄격히 금지하고, ‘고위험 AI(high risk)’에 대하여는 엄격한 요구사항을 부과하는 EU의 강력하고 포괄적인 규제와는 차별화되는 유연한 규제 체계를 제시하였다.¹⁰⁰⁾

위 AI 백서는 ① AI 주도 국가로서의 영국의 입지 강화, ② 책임있는 혁신 도모 및 규제 불확실성 감소를 통한 AI의 성장과 번영, ③ 위험 규율 및 기본가치 보호를 통한 AI에 대한 대중의 신뢰 제고라는 3가지 목표를 추구한다. 이를 위해, 혁신을 억제할 수 있는 AI에 대한 고정 규칙이나 법안 도입은 최소화하고, (새로운 단일 규제기관을 신설하여 전권을 부여하는 것이 아니라) 기존 관련 기관이 부문별·상황별 지침을 채택하는 유연한 규제방식을 취한다

100) 영국 정부 홈페이지 <<https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper>> 최종 방문일 2024. 9. 1.>

는 방침을 내세우고 있다. 특히, 위 기관들로 하여금 AI가 활용되는 개별적 맥락을 중요시하면서도, (1) 안전·보안·견고성, (2) 투명성·설명 가능성, (3) 공정성, (4) 책임·거버넌스, (5) 이의제기·시정이라는 5가지 공통적 원칙을 제시하여, 우선순위를 정해 규제를 적용하고, 성장과 혁신을 모두 지원할 수 있는 비례적 접근을 취할 것을 제안하였다.

영국은 위 백서에서, AI 규제의 초기에는 법적 구속력 없는 원칙을 발표하고, 규제기관은 자체 권한 내에서 이를 적용하는 방식을 취하되, 차후 실효성 확보를 위하여 규제기관의 추가 개입 및 법률 변경이 이루어질 수 있다고 설명하였다. 이와 같은 배경에서 영국 정부는 상기 규제에 대한 의견청취 정부 답변에서 영국의 경쟁시장청(CMA)·광고기준청(ASA)·통신청(Ofcom) 등의 규제기관이 AI 원칙을 자신의 권한 내에서 적극적으로 이행하는 현 상황을 긍정적으로 보고 있다고 밝혔으며, 정보위원회(ICO)·의약품규제청(MHRA) 등도 각자의 이행계획을 진행 중이라고 기재하였다. 또한 상기 답변에서 정부는 새로운 AI 규제가 기존 법령과 모순되거나 중복되지 않도록 주의해야 한다는 의견을 고려하고 있으며, 규제기관의 기존 법적 권한과 현행 입법 체계 내에서의 잠재적 격차나 마찰을 계속 평가할 예정이라고 밝혔다.

(2) AI 규제 환경 청사진

영국 정부는 2024년 2월 6일에 위 백서에 대한 여러 유관기관의 질문에 대한 답변을 취합·정리한 정부 답변(Government response)을 발간하였고 이를 통하여 새로운 AI 규제 환경의 청사진(규제 프레임워크)을 제시하였다.¹⁰¹⁾

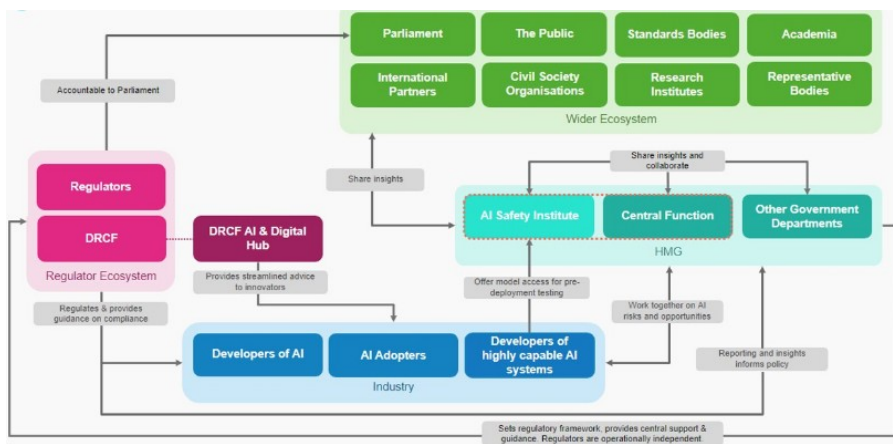
거버넌스 측면에서, 디지털 정책 관련 규제 조정을 위해 경쟁시장청

101) 영국 정부 홈페이지 <<https://www.gov.uk/government/consultations/ai-regulation-a-pro-innovation-approach-policy-proposals/outcome/a-pro-innovation-approach-to-ai-regulation-government-response>> 최종 방문일 2024. 9. 1.>

(CMA)·금융행위감독청(FCA)·정보위원회(ICO)·통신청(Ofcom)이 구성원이 되어 설립한 디지털 규제 공동 포럼(DRCF)과 그 내부에 설치된 AI & Digital Hub 및 다른 규제기관들이 규제기관 생태계(Regulator Ecosystem)을 구성하도록 하였다. 규제기관은 독립적으로 운영되면서 AI 산업계를 규제 및 지도하고, 정부는 AI 안전 연구소(AI Safety Institute)·핵심 기능·여러 정부 부서를 중심으로 규제기관들을 지원하고 업무 관련 보고를 받으며 AI 산업계와 협력하도록 하였다. 자문기구로서 영국 의회와 학계·시민사회·연구기관·국제 파트너 등은 광의의 생태계(Wider Ecosystem)을 구축하여 정부와 의견을 교류할 예정이다. 이 과정에서, 영국 정부는 개별 규제기관이 이미 AI 규제를 도입한 경우 기관 간 규제 내용을 공유하고, 일관된 지침을 제공하는 것을 목표로 하고 있다.

새로운 규제 프레임워크는 혁신 촉진 및 규제 불확실성 해소와 AI에 대한 대중의 신뢰 구축, AI 분야 영국의 글로벌 리더십 강화를 목표로, ① 혁신 친화성, ② 균형성, ③ 신뢰성, ④ 적응성, ⑤ 명확성, ⑥ 상호보완성 6가지 핵심 원칙을 제시하였다.

[그림 2] 영국 AI Regulation Landscape



자료 : 영국 정부 홈페이지

규제 방식은, AI를 폭넓게 정의하여 규제의 유연성을 확보하고, 기술 자체가 아닌 기술 활용 규제를 통한 맥락 기반으로 접근하고, 각 규제기관 간 정보공유 등을 통해 기관별 개별 원칙·특수성을 고려하는데 초점을 두었다. 특히 이와 관련하여 정부는 규제기관의 혁신적인 태도를 위해 위험 평가·교육 등 여러 핵심 기능을 지원하도록 하였다.

한편, LLM의 막대한 영향력을 고려할 때, 규제의 필요성은 인정되나 자칫 혁신 저해로 이어질 수 있으므로 현 단계에서는 리스크 예측과 지속적인 모니터링만 진행하며, 추후 특정 지침을 발표할 수 있도록 하였다.

범정부 차원 AI 위험 관리를 위해 DSIT 내 AI 유관부서가 관련 위험 모니터링 및 평가 수행 중이며, 금년 중 경제 영역 전반의 위험 평가를 위한 AI 위험 등록부 시행 방안 의견 수렴 예정이다. 중복되거나 모순되는 지침을 피하기 위해 기관 간 조정 기능을 마련할 예정이며, 정부대표와 주요 규제기관으로 구성된 AI 거버넌스 조정위원회를 설치 예정이다.

(3) 시사점

영국의 AI 규제 거버넌스는 AI가 지닌 다양한 산업 분야에서의 잠재력을 인정하고 그 토대 위에서 영국이 AI 혁신을 선도하기 위한 기관별 협력 구조와 각종 조치를 담고 있다. 현행 영국의 AI 규제는 별도의 AI 법률 없이 기존의 법적 프레임워크를 통하여 이루어지고 있으며, 차후 AI 규제 백서에서 제시한 ‘원칙 기반 접근법’에 따라 안전, 보안, 공정성, 개인정보 보호, 인권 등의 영역에서 AI로 인한 부작용과 악영향을 완화하면서도 산업의 발전을 저해하지 않는 각종 방안이 마련될 예정이다.¹⁰²⁾

102) 이근우, 「해외 주요국의 AI 규제 거버넌스 현황 및 시사점 - 영국의 AI 규제 거버넌스 현황」, 지능정보사회 법제도 이슈리포트(2024-01), 2024. 8.

라. 일본¹⁰³⁾

(1) 배경

일본 아베 정부는 AI를 자국 성장 전략의 핵심으로 설정하며 2018년 9월 ‘통합 혁신전략회의’를 설치하였고, 2019년 6월에는 국가 차원의 첫 종합 전략인 「AI 전략 2019」과 「인간 중심의 인공지능(AI) 사회원칙」을 발표하였다. 총무성과 문부과학성, 경제산업성 등 3개 부처를 중심으로 AI 개발 촉진을 위해 법적 강제력이 있는 규제 대신, ‘인간 중심의 인공지능 사회원칙’을 바탕으로 기업의 자율을 존중하는 정책 추진을 하였다.

이후 코로나 팬데믹 발생에 따라, 일본 정부는 디지털 인프라 확충을 위해 새로운 목표를 추가한 「AI 전략 2022」를 발표하였고, 그 일환으로 국가 AI 컨트롤타워 역할을 수행할 ‘AI 전략회의’를 설치하였다. ‘AI 전략회의’는 2023년 5월 1차 회의를 시작으로 AI 기술의 확대에 발생할 잠재력과 리스크 등을 검토한 뒤, 2024년 4월 법적 구속력이 없는 자율규제 형태의 「AI 사업자 가이드라인」을 공표하였다.

(2) 거버넌스

거버넌스 측면에서, 일본 정부는 산·학·연의 AI 지식 결집·강화와 부처간 장벽을 배제하기 위해 2016년 첫 국가 컨트롤타워 ‘인공지능기술전략회의’를 설치하였다. ‘인공지능기술전략회의’ 산하의 ‘연구연계회의’는 문부과학성, 경제산업성, 총무성을 연계해 연구 총괄 및 조정, 목표 등을 구체화하였고, ‘산업연계회의’는 연구개발과 산업의 연계 총괄을 조정해 왔다.

103) 정지은, 「해외 주요국의 AI 규제 거버넌스 현황 및 시사점 - 일본의 AI 규제 거버넌스 현황」, 지능정보사회 법제도 이슈리포트(2024-01), 2024. 8.; KoTRA, 일본의 AI 정책과 실제 사례, 2024. 4. 4.

코로나 팬데믹 당시, ‘대규모 재해에 대한 대처’와 ‘사회 실증’을 추가한 「AI 전략 2022」을 발표 하였으며, 내각부 산하에 국가의 AI 컨트롤타워 역할을 해줄 ‘AI전략회의’를 설치하였다. AI 기술 확대로 발생할 잠재력과 동시에 AI 리스크 등에 대한 검토를 거친 뒤, 2024년 4월 「AI 사업자 가이드라인」을 발표하였다.

정부는 AI 정책을 정부 차원에서 추진할 수 있도록 총무성과 문부과학성, 경제산업성을 주축으로 회의를 구성해 AI 정책 추진하고 있다.

자문기구로서 2023년 5월 기시다 후미오 일본 총리는 내각부 산하에 전문가와 관계부처 담당자로 구성된 ‘AI 전략회의’ 설치하고 ‘AI 전략회의’하에 내각총리보좌관이 관계부처 합동의 ‘AI 전략팀’의 팀장을 맡아, AI 관련 문제를 신속히 검토하고 정부 방침을 마련하고자 하였다. 회의 좌장을 맡은 도쿄대학원 마츠오 유타카 교수는 AI 산업 경쟁력 강화 방안과 보안·프라이버시·저작권 등의 리스크를 논의할 계획을 밝혔다.

(3) AI 사업자 가이드라인

「AI 사업자 가이드라인」은 AI를 개발하고 활용하는 사업자를 위한 것으로, 안전성·공정성·투명성 등 10가지 항목에 대한 31개의 지침을 제시하며 AI의 적절한 활용을 촉진하는 것을 목적으로 한다.

AI를 통해 나타날 수 있는 사회적 편견, 가치관, 거짓정보, 성별이나 인종 등에 관한 민감한 속성, 보안, 독성 데이터에 의한 공격, 데이터 보호 등에 대해 다루었고 G7에서 추진하는 히로시마 AI 프로세스를 비롯한 해외 주요국의 AI규제 동향이 포괄적으로 담겨있다.

[그림 3] 일본 AI 사업자 가이드라인 주요 내용

기본 이념	
'인간 중심의 AI 사회 원칙'에 따른 ① 인간의 존엄성 존중받는 사회 ② 다양한 배경의 사람들이 다양한 행복을 추구할 수 있는 사회 ③ 지속가능한 사회 실현	
AI 관련자별 사항	
① AI 개발자: AI 제공·활용될 때 미칠 영향을 사전에 최대한 검토 및 대응책 마련 중요 ② AI 제공자: AI 운영·적절한 활용을 전제로 한 AI 시스템 및 서비스 제공 실현 중요 ③ AI 이용자: AI 제공자가 의도한 범위 내에서 적절하게 이용, 필요한 지식 습득 중요	
AI를 활용한 모든 사업자를 대상으로 한 10대 원칙	
① 인간중심 인감 존엄과 개인 자유 존중 ② 안전성 인간에 의한 컨트롤 확보 ③ 공정성 부당한 차별 최소화 ④ 프라이버시보호 개인정보보호법 기초 대응 ⑤ 시류리터화보 시스템 기밀성 유지	⑥ 투명성 데이터 수집 방법 등 대외 공개 ⑦ 설명책임 AI에 대한 이념, 사상 공표 ⑧ 교육·리터러시 올바른 지식 보급 ⑨ 공정강장확보 이해관계자에 유리하지 않게 ⑩ 이노베이션 사회전체의 기술혁신에 공헌

자료 : KoTRA, 일본의 AI 정책과 실제 사례(2024. 4. 4.); 정지은, 일본의 AI 규제 거버넌스 현황(2024. 8.)

일본은 AI에 대하여 법률적 규제 대신 AI 법적 구속력이 없는 연성규범 방식을 활용하며 기존 법으로 AI를 관리하였고, AI 개발 촉진을 중시해 기업 자율을 강조하는 정책을 추진해 왔다. 「AI사업자 가이드라인」도 AI 관련 사업자가 자율적으로 준수할 것을 권고한 수준이다.

주요 리스크로 보고 있는 저작권과 관련해서도, 일본 정부는 TDM 면책 규정을 마련해 저작권자의 이익을 부당하게 침해하지 않으면 영리·비영리 목적과 관계없이 저작물을 AI 학습용 데이터로 사용할 수 있도록 하였고, 올해 1월 'AI와 저작권'에 대한 해석을 정리한 '인공지능과 저작권에 관한 고찰안' 발표했을 뿐 실질적인 규제는 도입하지 않았다.

(4) 입법 전망

일본은 현재 AI 규제를 위한 특정 법률은 없으며 전문가 및 정부관계자들이 참여하는 ‘AI 전략회의’를 통해 AI 규제·활용에 대한 정책방향을 논의 중이다. 그러나 최근 AI 규제를 강화하는 국제 흐름에 따라, 일본도 지난 5월 22일 ‘AI전략 회의’를 개최해 AI 규제 기본방침과 AI 안전성 확보를 위한 법률 규제를 검토하겠다고 발표했지만, 실제 법 시행까지 속도를 조절해 내년 정기 국회에 법안을 제출하고 2026년 전면 시행을 목표로 검토할 계획임을 밝혔다.¹⁰⁴⁾

마. EU와 미국 간 규제체계 접근 방식 비교

EU와 캐나다는 상대적으로 중앙집권적인 집행이 가능한 수평적 규제, 즉 대부분의 분야와 애플리케이션에 걸쳐 적용되도록 설계된 규제를 도입하는 접근 방식을 취한다. 반면, 미국, 영국, 중국은 보다 분산된 접근 방식을 취하고 있으며, 여러 규제 기관에서 시행하는 수직적 규제, 즉 AI의 특정 애플리케이션이나 특정 부문에만 적용되는 규제에 대한 의존도가 더 높다. 미국의 접근 방식은 구속력이 없는 원칙, 위험 관리에 대한 자발적 지침, 연방 차원에서 새로운 AI 관련 법률을 개발하기보다는 기존의 부문별 법률을 적용하는 것이 특징이다.¹⁰⁵⁾

EU와 미국 모두 AI 규제에 대해 대체로 위험 기반 접근 방식을 지지하며, 신뢰할 수 있는 AI가 작동해야 하는 방식에 대해 유사한 원칙을 설명하고 있다. 미국의 최신 지침 문서(AIBoR 및 NIST AI RMF)와 EU 인공지능법의 원칙을 살펴보면, 모두 정확성과 견고성, 안전성, 비차별, 보안, 투명성과 책임성, 설명 가능성 및 해석 가능성, 데이터 프라이버시를 지지하며 약간의 차이

104) Kotra, “일본 AI 법 규제 본격 논의 시작...일본 기업의 대응은?”, 2024. 8. 27.

105) Huw Roberts et al, 「A comparative framework for ai regulatory policy」, 2023.

만 있을 뿐이다. 또한 EU와 미국 모두 정부 및 국제기구인 표준화 단체가 AI에 대한 가드레일(guardrail)을 설정하는 데 중요한 역할을 할 것으로 기대하고 있다.¹⁰⁶⁾

이러한 광범위한 개념적 일치에도 불구하고 AI 위험 관리에는 상이한 점이 훨씬 많다. EU의 접근 방식은 전체적으로 볼 때 더 많은 애플리케이션을 포함하고 각 애플리케이션에 대해 더 구속력 있는 규칙을 공표한다는 측면에서 미국보다 훨씬 더 중앙에서 조정되고 포괄적인 규제를 적용하고 있다. 미국 기관들이 본격적으로 가이드라인을 작성하고 해당 영역 내 AI 애플리케이션에 대한 규칙 제정을 고려하기 시작했지만, 이러한 규칙을 집행할 수 있는 능력은 아직 불분명하다. 미국 기관은 이러한 규칙을 시행하기 위해 알고리즘을 규제할 명시적인 법적 권한이 없는 새로운 소송을 추진해야 할 수도 있다. 일반적으로 EU 규제 당국은 명확한 조사 권한과 규정 미준수 시 상당한 벌금을 부과하여 AI 애플리케이션에 대한 규칙을 제정할 수 있다.

[표 3] 애플리케이션 유형별 EU와 미국의 AI 위험 관리 비교

애플리케이션	예시	EU 정책 개발	미국 정책 개발
인간 프로세스를 위한 AI/사회경제적 의사 결정	채용, 교육 접근성, 금융 서비스 승인에 활용되는 AI	GDPR은 중요한 의사 결정에 사람이 참여하도록 요구한다. EU 인공지능법 부록 III의 고위험 AI 애플리케이션은 품질 표준을 충족하고 위험 관리 시스템을 구현하며 적합성 평가를 수행해야 한다.	AI 권리장전 및 관련 연방 기관의 조치에 따라 이러한 애플리케이션 중 일부에 대한 패치 워크 감독이 이루어지고 있다.
소비재 분야의 AI	의료 기기, 부분 자율 주행 차량 및 비행기의 AI	EU AI 법은 이미 EU 법률에 따라 규제되고 있는 제품 내에서 구현되는 AI를 고위험으로 간주하고, 나아가 현행 규제 절차에 새로운 AI 표준을 통합할 것이다.	의료 기기의 경우 FDA, 자율 주행차의 경우 DOT, 소비자 제품의 경우 CPSC와 같은 개별 연방 기관에서 적용한다.

106) Alex Engler, 「The EU and U.S. diverge on AI regulation: A transatlantic comparison and steps to alignment」, Brookings Research, 2023. 4. 25.

애플리케이션	예시	EU 정책 개발	미국 정책 개발
챗봇	상업용 웹사이트의 판매 또는 고객 서비스 챗봇	EU AI 법에 따라 챗봇이 사람이 아닌 AI라는 사실을 공개해야 한다.	N/A
소셜 미디어 추천 및 모더레이션 시스템	TikTok, Twitter, Facebook, Instagram의 뉴스피드 및 그룹 추천 시스템	EU DSA는 이러한 AI 시스템에 대한 투명성 요건을 마련하고, 독립적인 연구와 분석을 가능하게 한다.	N/A
이커머스 플랫폼의 알고리즘	Amazon 또는 Shopify에서 제품 및 공급업체를 검색하거나 추천하는 알고리즘	EU DMA는 디지털 시장에서 자체 선호(self-referencing) 알고리즘을 제한할 것이다[전자상거래 알고리즘 및 플랫폼 설계에서 자체 선호를 줄이기 위한 개별 반독점 조치 포함(예: 아마존 및 Google 쇼핑에 대한 조치)].	N/A
파운데이션 모델/제너레이티브 AI	안정성 AI의 안정적 확산과 OpenAI의 GPT-3	EU AI 법은 품질 및 위험 관리 요구 사항을 고려한다.	N/A
얼굴 인식	클리어뷰 AI, 핼라이즈, 아마존 레코그니션	EU AI 법에는 원격 얼굴 인식 및 생체 인식에 대한 제한이 포함된다. EU 데이터 보호 당국은 GDPR에 따라 얼굴 인식 회사에 벌금을 부과했다.	NIST의 AI 얼굴 인식 벤더 테스트 프로그램은 얼굴 인식 소프트웨어 시장에 효율성과 공정성 정보를 제공한다.
타겟팅 광고	웹사이트 및 휴대폰 애플리케이션의 알고리즘 기반 타겟팅 광고	GDPR은 행동 광고에 개인 사용자 데이터를 사용한 Meta에 벌금을 부과했다. DSA는 어린이를 대상으로 한 타겟 광고와 특정 유형의 프로파일링(예: 성적 취향)을 금지하고 있다. 이 법에 따르면 타겟팅 광고에 설명이 포함되어야 하며, 사용자가 어떤 광고가 표시되는지 제어할 수 있어야 합니다.	개별 연방 기관의 소송으로 인해 일부 타겟팅 광고가 약간 축소되었다. 여기에는 차별적인 주택 광고에 대한 Meta의 소송과 타겟팅 광고에 보안 데이터를 사용한 트위터에 대한 FTC의 과징금 부과에 성공한 법무부와 HUD가 포함된다.

자료 : Alex Engler, The EU and U.S. diverge on AI regulation: A transatlantic comparison and steps to alignment, Brookings Research, 2023. 4. 25.

EU와 미국은 채용, 교육 접근성, 금융 서비스 등 사회경제적 결정에 영향을 미치는 AI에 대해 서로 다른 규제 접근 방식을 취하고 있다. EU의 접근 방식은 더 넓은 범위의 애플리케이션과 이러한 AI 애플리케이션에 대한 더 광범위한 규칙을 모두 포함하고 있다. 미국의 접근 방식은 훨씬 더 제한적이며, 현행 기관의 규제 권한을 조정하여 AI를 처리하려는 방식으로 EU보다 좁게 축소되어 있다. 일부 기관은 위 작업을 본격적으로 시작하여 직관적으로 많은 EU 기관보다 앞서 있지만, EU 회원국 기관과 잠재적인 EU AI 위원회는 EU법의 강력한 권한, 새로운 권한 및 자금 지원으로 인해 따라잡을 수 있을 것으로 예상된다. EU와 미국 간의 불균등한 권한과 AI 규제 시기의 불확실성으로 인해 조율에 상당한 어려움이 있을 수 있다.

EU-미국 무역기술위원회(The Trade and Technology Council, TTC)는 미국과 EU가 주요 글로벌 무역, 경제 및 기술 문제에 대한 접근 방식을 조율하고 이러한 공유 가치를 바탕으로 대서양 횡단 무역 및 경제 관계를 심화하기 위한 포럼 역할을 한다. 2021년 6월 15일 브뤼셀에서 열린 EU-미국 정상회담에서 설립되었다. 안전하고 신뢰할 수 있는 전 세계 AI에 대한 공동 비전을 달성하기 위해 공유된 민주적 가치를 기반으로 AI 거버넌스 프레임워크 간의 상호 운용성을 발전시키고 강화할 것으로 기대된다.¹⁰⁷⁾

EU와 비교했을 때 미국의 허가 없는 임의적 혁신 접근 방식은 실용적이고 민첩하며 반복적이고 문제 기반이지만 과편화되어 있다. 이는 종종 불충분한 것으로 여겨지며, 특히 의료 분야에서의 AI의 가능성과 함정에 대해서는 더욱 그렇다. 하지만 혁신을 더 쉽게 실현할 수 있다는 장점도 있다. 일부에서는 GDPR과 EU 인공지능법이 취약한 스타트업과 스케일업에 위축 효과를 가져와 EU 출신의 의료 혁신 유니콘 기업이 탄생할 기회를 줄인다고 주장한다.

107) 미국 백악관 보도자료, “U.S-EU Joint Statement of the Trade and Technology Council”, 2024. 4. 5.

미국의 접근 방식은 지나치게 세분화되고 자유로운 기업을 의미하는 반면, EU의 접근 방식은 법적, 윤리적, 사회경제적 측면에서 지나치게 예방적이라는 견해가 있다.

의료 산업에 필요한 것은 합리적이고(환자 안전과 건전한 기술에 중점을 두고), 실용적이며(이해와 실행이 쉽고), 해당 부문의 특정 요구에 맞는 규제다. 임상시험 비용과 같은 경제적 현실과 생산 및 시장 라이선스와 같은 기존 법적 구조는 다른 산업/부문과 다르므로 규제 당국은 이를 고려해야 한다. 이 두 가지 측면을 제대로 고려하지 않으면 규제가 구체성이 부족하거나 진정으로 중요한 규제 주제를 다루지 못하여 금방 무용지물이 되고 실효성이 떨어진 다. 혁신 기업과 환자 모두에게 진정으로 도움이 되는 규제 환경을 조성하기 위해서는 사전 예방적 혁신과 무허가 혁신의 장점을 혼합하여 의료 분야의 특성에 맞는 실행 가능한 중간 지점을 마련해야 한다.¹⁰⁸⁾

즉, EU와 미국 방식 중 어느 한 방식을 규범적으로 선택하여 적용하기 보다는, 국민의 권리를 보호하면서도 인공지능 기술·산업 발전을 촉진할 수 있는 방안을 유연하게 혼합하여 적용하는 방안을 고려할 필요가 있다.

3. 국내 규제체계 진행 상황

가. 제21대 국회 발의 경과

우리나라 제21대 국회에서는 2020년 7월 최초로 인공지능 법안을 발의한 이후, 개별 발의된 7개 법안을 병합한 위원회 대안이 상임위 법안소위까지 통과되는 성과가 있었지만 결국에는 제21대 국회 만료로 최종 폐기되었다.

108) Suzan Slijpen, Mauritz Kop & I. Glenn Cohen, 「EU and US Regulatory Challenges Facing AI Health Care Innovator Firms」, 2024. 4. 6.

나. 제22대 국회 발의 현황¹⁰⁹⁾

(1) 법안 발의 현황

제22대 국회에 발의된 인공지능 관련 법안은 2024년 8월 28일 기준으로 다음의 총 9개가 계류 중이다.

[표 4] 제22대 국회 발의안 (2024.8.28. 기준)

제안일자	의안번호	의안명	약칭
2024.05.31	제2200053호	인공지능 산업 육성 및 신뢰 확보에 관한 법률안(안철수의원 등 12인)	안철수의원안
2024.06.17	제2200543호	인공지능 발전과 신뢰 기반 조성 등에 관한 법률안(정점식의원 등 108인)	정점식의원안
2024.06.19	제2200673호	인공지능산업 육성 및 신뢰 확보에 관한 법률안(조인철의원 등 19인)	조인철의원안
2024.06.19	제2200675호	인공지능산업 육성 및 신뢰 확보에 관한 법률안(김성원의원 등 11인)	김성원의원안
2024.06.28	제2201158호	인공지능기술 기본법안(민형배의원 등 13인)	민형배의원안
2024.07.04	제2201399호	인공지능 개발 및 이용 등에 관한 법률안(권철승의원 등 15인)	권철승의원안
2024.08.22	제2203072호	인공지능 기본법안(한민수의원 등 10인)	한민수의원안
2024.08.27	제2203235호	인공지능책임법안(황희의원 등 10인)	황희의원안
2024-08-28	제2203297호	인공지능 발전 진흥과 사회적 책임에 관한 법률안(배준영의원 등 10인)	배준영의원안

109) 정점식의원안, 조인철의원안, 김성원의원안, 민형배의원안, 권철승의원안에 대한 검토는 국회 과학기술정보통신위원회 「인공지능 산업 육성 및 신뢰 확보에 관한 법률안 등 검토보고서」(2024. 8.)의 내용을 주로 참고하였으며, 안철수의원안, 한민수의원안, 황희의원안, 배준영의원안은 추가로 검토하였다.

(2) 법안의 규율체계 및 주요 내용 분석

인공지능 관련 제정법안(이하 ‘인공지능 법안’)은 대체로 총칙 및 추진체계, 인공지능 기술 개발·산업육성, 인공지능 윤리·신뢰성 확보, 보칙·벌칙의 체계로 구성되어 있다. 법안별 규율체계를 살펴보면 다음과 같다.

[표 5] 법안별 규율체계 비교

의안	총칙	정책 수립 및 추진체계	개발 및 산업 육성	인공지능 윤리 및 신뢰성 확보	분쟁조정	보칙·벌칙
안철수의원안	제1장	제2장	제3장	제4장	-	제5장
정점식의원안	제1장	제2장	제3장	제4장	-	제5장
조인철의원안	제1장	-	제2장	제3장	-	제4장
김성원의원안	제1장	제2장	제3장	제4장	-	제5장
민형배의원안	제1장	제2장	제3장	제4장	-	제5장
권철승의원안	제1장	제2장	제3장, 제4장	제5장	-	제6장, 제7장
한민수의원안	제1장	제2장	제3장	제4장	-	제5장
황희의원안	제1장	-	제2장	제3장	제4장	제5장
배준영의원안	제1장	제2장	제3장	제4장	-	제5장

인공지능 법안들은 2023년 2월 14일에 제21대 국회 과학기술정보방송통신위원회 정보통신방송법안심사소위원회에서 대안으로 제안하기로 의결한 안 및 안철수의원안과 체계·내용과 대체로 유사한 방식과 내용으로 구성되어 있으며, 이와 유사하게 큰 틀에서 반복적으로 규율하고 있는 내용과 일부 법안에 담긴 내용을 정리하자면 다음과 같다.¹¹⁰⁾

110) [표 5] 인공지능 법안의 일반적 구성에 황희의원안의 규율 체계는 기존의 체계와 대부분의 내용이 상이하여, 황희의원안만이 규율하는 법안 구성요소는 배제함.

[표 6] 인공지능 법안의 공통 내용 및 일부 법안에 포함된 내용

구분	공통 내용	일부 법안의 내용
총칙	• 목적 • 정의 • 기본원칙 • 다른 법률과의 관계	• 국가·지방자치단체 등의 책무(권찰승의원안, 황희의원안) • 국외행위에 대한 적용(정점식의원안)
정책 수립 및 추진체계	• 인공지능 기본계획의 수립 • 인공지능위원회, 위원회의 기능, 위원의 제척·기피 및 회피, 전문위원회 등 • 국가인공지능센터	• 인공지능 실행계획(권찰승의원안) • 지역 인공지능 종합계획의 수립, 지역계획의 심의와 조정, 지역인공지능위원회의 구성 등(민형배의원안) • 인공지능안전연구소(안철수의원안, 정점식의원안)
개발 및 산업 육성	• 인공지능기술 개발 및 안전한 이용 지원 • 인공지능기술의 표준화 • 인공지능 학습용데이터 관련 시책의 수립 등 • 기업의 인공지능기술 도입·활용 지원 • 창업의 활성화 • 인공지능 융합의 촉진 • 제도개선 등 • 전문인력의 확보 • 국제협력 및 해외시장 진출의 지원 • 인공지능집적단지 지정 등	• 우선허용·사후규제 원칙(김성원의원안) • 인공지능 실증 규제특례(조인철의원안) • 중소기업자를 위한 특별지원(한민수의원안, 황희의원안) • 대한인공지능진흥협회의 설립(안철수의원안, 조인철의원안, 김성원의원안)
인공지능 윤리 및 신뢰성 확보	• 인공지능 윤리원칙 등 • 신뢰할 수 있는 인공지능 • 인공지능 신뢰성 검·인증 지원 등 • 고위험영역 인공지능의 확인 • 고위험영역 인공지능 고지의무 • 고위험영역 인공지능과 관련한 사업자의 책무 • 생성형 인공지능 고지 및 표시 • 민간자율인공지능윤리위원회의 설치 등	• 금지된 인공지능의 개발·이용 제한(권찰승의원안) • 인공지능제품의 비상정지(조인철의원안) • 생성형 인공지능 안전 확보 의무(정점식의원안)

황희의원안의 구성에 관하여는 아래 본문에 서술한다.

구분	공동 내용	일부 법안의 내용
보칙	• 인공지능산업의 진흥을 위한 재원의 확충 등 • 실태조사, 통계 및 지표의 작성 • 권한의 위임 및 업무의 위탁	
벌칙	• 벌칙 적용에서 공무원 의제 • 벌칙(직무상비밀누설)	• 벌칙(금지된 인공지능을 개발·이용)(권철승의원안) • 벌칙(검·인증등을 받지 않고 고위험 인공지능을 이용하여 제품·서비스를 제공)(권철승의원안) • 벌칙(고위험 인공지능 운용 사실을 고지하지 아니한 자)(권철승의원안)

인공지능 법안은 대체로 세부항목을 제외하고 유사한 내용과 형식으로 인공지능 등 용어를 정의하고 있으며, 구체적 거버넌스, 운영 등에 차이는 있으나 정부와 민간이 참여하는 위원회를 도입하도록 규정하고 있다. 그 외 산업 발전이나 윤리·신뢰성에 대한 규제체계도 전체적으로 유사한 체계로 구성되어 있으며 법안별 세부내용(고위험 인공지능 관련 책무 수준, 벌칙 규정의 종류 등)에서 차이가 있다. 인공지능 법안은 대체로 기본적인 법안 구조의 동일성을 유지하면서, 큰 틀에서의 규율 체계에 대한 이견은 없는 것으로 보인다.

다만, 각 법안 간 세부 내용에 대한 차이점이 존재하는데, ① ‘생성형 인공지능’을 정의하고 이에 대한 규제를 포함할지 여부(정점식·민형배의원안) 및 고위험 인공지능에 대한 정의를 법률의 형식으로 정할 것인지 대통령령에 위임하는 방식으로 정할 것인지 여부(권철승의원안), ② ‘법의 적용범위’에 관하여 국외에서 이루어진 행위라도 국내 시장·이용자에게 영향을 미치는 경우 인공지능법을 적용할 수 있도록 명문에 규정을 해야 할 필요성(정점식의원안), ③ ‘인공지능위원회’를 대통령 소속으로 규정할지(정점식·조인철의원안), 국무총리 소속으로 규정할지(김성원·권철승의원안) 여부,¹¹¹⁾ ④ ‘국가인공지능센

111) 한편, 민형배의원안의 경우 위원회를 국무총리 소속의 국가인공지능위원회와 지자체의 지역인공지능위원회로 구분하고 있다.

터’를 설치할지(김성원·권철승의원안) 또는 ‘인공지능정책센터’를 설치할지(민형배의원안) 여부 및 나아가 ‘인공지능안전연구소’를 별도로 설치할지(정점식의원안) 여부, ⑤ 사업자를 개발사업자와 이용사업자로 나누어 규정할 것인지 여부(안철수·권철승의원안), ⑥ 규제 샌드박스를 도입 여부 및 인공지능제품에 비상정지 기능 적용 여부(조인철의원안), 우선허용·사후규제 원칙 도입 여부(김성원의원안), ⑦ ‘금지된 인공지능’의 개발과 이용을 금지하고, 금지된 인공지능 및 고위험 인공지능 관련 의무 위반 시 처벌하는 벌칙 규정 등을 도입할지 여부(권철승의원안) 등에 관한 논의가 진행되어야 할 것으로 보인다.

반면, 황희의원안¹¹²⁾의 경우 다른 인공지능 법안들과 다른 독자적 규율체계로 구성되어 있다. 황희의원안은 위 기존 법안들과는 정책 수립 및 추진체계, 개발 및 산업 육성에 관하여는 상대적으로 포괄적인 근거 규정을 두고 있는 반면, 고위험 인공지능의 개발 및 이용 과정에서 사업자의 책임 및 이용자의 권리를 세분화하여 규정하면서 이용자의 권리로서 설명요구권 등을 별도로

112) 황희의원안의 주요내용은 다음과 같다.

- 가. 이 법은 안전하고 신뢰할 수 있는 인공지능 기술의 개발 및 이용, 관련 산업의 육성을 위한 사회적 기반을 마련하는데 필요한 사항을 정함으로써 경제 발전과 국민의 삶의 질 향상에 이바지함을 목적으로 함(안 제1조).
- 나. “인공지능”, “고위험인공지능”, “인공지능사업자” 등에 대하여 정의함(안 제2조).
- 다. 인공지능의 개발 및 이용의 기본원칙이 인류의 발전과 편의 도모를 위함임을 명시하고, 인공지능사업자로 하여금 사업자책임위원회를 운영하도록 함(안 제3조 및 제5조).
- 라. 인공지능 기술의 개발, 기술기준의 마련, 표준화 및 실용화·사업화 등 인공지능 산업의 진흥을 위한 정부의 역할을 규정하고, 인공지능에 대한 규제 원칙을 정함(안 제7조부터 제17조까지).
- 마. 고위험인공지능으로부터 이용자 보호를 위한 정부의 역할, 사업자 책무 그리고 이용자의 설명요구권, 이의제기권 및 책임의 일반원칙 등을 규정함(안 제18조부터 제22조까지).
- 바. 인공지능에 관한 분쟁 조정을 위하여 인공지능분쟁조정위원회를 설치하고 관련 절차를 마련함(안 제23조부터 제32조까지).

규정하고, 분쟁조정위원회의 설립 등 분쟁조정제도를 별도의 장으로 구성하여 도입하는 등 기존의 법안들과 전혀 다른 규율체계를 도입하고 있다.

황희위원안과 관련하여 다음과 같은 정부부처의 의견을 참고할 수 있다. 방송통신위원회는 인공지능 법안에 관하여 이용자 설명요구권, 분쟁조정제도 등 이용자 보호를 위한 추가적인 제도 도입을 검토할 필요가 있다는 입장으로 황희위원안만이 독자적으로 규율하고 있는 요소들과 부합하는 면이 있다. 반면, 과학기술정보통신부는 인공지능에 관한 규율체계는 인공지능 산업 발전 수준을 고려하여 단계적으로 확립해 나가는 것이 바람직하며 이용자 설명요구권, 분쟁조정제도 등 이용자 보호 관련 규정은 주요국에서도 아직 필요성에 대한 논의가 진행 중이고 실제 도입한 사례를 찾아보기 어렵다는 점을 고려하여, 충분한 검토 후 필요하다고 인정될 경우 도입하는 것이 바람직하다는 입장이다.¹¹³⁾

현재 제22대 국회에 계류 중인 인공지능 법안은 기본적으로 인공지능을 정의하고, 인공지능 기술 및 산업을 육성하기 위한 진흥규정을 제공하며, 인공지능 윤리 및 신뢰성을 확보하기 위한 규제규정을 동시에 포함함으로써 국가 경쟁력을 강화하고 국민의 삶의 질을 향상시키는 동시에 국민의 권익과 존엄성을 보호하고자 한다는 점에서 공통점이 있다. 인공지능 기본법의 입법 방향은 인공지능 산업의 진흥과 규제 사이의 균형을 확보하는 것이 핵심이며, 규제의 글로벌 정합성을 고려하면서도 국가경쟁력에 도움이 되는 인공지능윤리 규범을 담은 법안의 조속한 통과가 필요한 상황이다.

(3) 규제 체계 관련 주요 내용

(가) 정의 조항

인공지능 법안의 규제체계에 있어서 인공지능에 관한 정의를 어떻게 할

113) 위 국회 과학기술정보통신위원회 검토보고서(각주 103) 28쪽 참조

것인지는 규율 범위를 구체화하고 법적안정성과 예측가능성을 확보한다는 점에서 그 내용과 규율 방식을 합의할 필요가 있다.

먼저, ‘인공지능’에 대한 정의에 대하여는 대체로 유사하게 규정하고 있으나, 알고리즘, 생성형 인공지능 등 세부 용어¹¹⁴⁾를 구분하여 정의할지 여부에 관하여는 차이점이 있다.

인공지능의 정의	
안철수의원안	“인공지능”이란 자율성을 가지고 외부의 환경 또는 입력에 적응하여 학습, 추론, 지각, 판단, 언어의 이해 등 인간이 가진 지적 능력을 전자적 방법으로 구현한 것을 말한다
그외 8개 법안	“인공지능”이란 학습, 추론, 지각, 판단, 언어의 이해 등 인간이 가진 지적 능력을 전자적 방법으로 구현한 것을 말한다

‘금지된 인공지능’, ‘고위험영역 인공지능’ 및 ‘생성형 인공지능’은 제정안에 따른 규제의 대상 및 범위와 밀접하게 연관이 있으므로, 이하의 관계부처 의견 등을 참조하여 가능한 한 구체적이고 명확하게 규정할 필요가 있다.

안철수·정점식·민형배·한민수·배준영의원안은 새로운 기술동향을 반영하여 ‘생성형 인공지능’을 “글, 소리, 그림, 영상 등의 결과물을 다양한 자율성의 수준에서 생성하도록 만들어진 인공지능을 말한다”고 정의하고 있는데, ‘자율성’의 수준에 대한 판단기준이 모호하여 사실상 글, 소리, 그림, 영상 등을 제작하는 것과 관련된 모든 제품 및 서비스가 사실상 생성형 인공지능에 해당할 우려가 있다는 지적이 있다.

‘고위험영역 인공지능’은 인공지능의 위험도에 따라 차등화된 규제를 적용하기 위한 것으로, 인공지능 법안의 주요 규율 대상이 된다는 점에서 그 내용과 규율방식에 관한 충분한 논의가 필요하다.

114) 안철수의원안 등, “알고리즘”이란 문제를 해결, 업무의 수행 또는 장비·장치·기기 등의 운용을 위하여 기술(記述)된 연산, 규칙과 절차, 명령 또는 논리의 집합으로 이루어진 체계를 말한다.

대부분의 법안이 규제 대상인 고위험인공지능을 법률의 형식으로 규정하고 있다. 안철수·정점식·조인철·김성원·민형배·한민수·배준영의원안은 고위험영역 인공지능의 범위를 구체적으로 열거하고 있으며 세부내용의 차이를 제외하고는 그 내용 또한 유사하다. 황희의원안은 법률이 직접 열거하여 규정하고 있다는 점은 같으나 규율의 내용이 다른 법안에 비해 상대적으로 포괄적이라는 차이가 있다. 반면, 권철승의원안은 고위험 인공지능의 구체적인 기준 및 유형을 대통령령으로 정하도록 하고 있다. 특히, ‘금지된 인공지능’을 별도로 정의하여 고위험 인공지능보다 더 규제의 필요성이 높은 인공지능의 유형을 대통령령으로 별도로 규정할 수 있도록 정의조항을 두고 있다. 인공지능 기술의 개발 및 환경변화의 속도가 빠른 만큼, 대통령령의 형식으로 규율하여 시의성을 확보할 수 있다는 점에서 위임형식의 의의를 제고할 수 있다.¹¹⁵⁾

다만, 권철승의원안의 경우 벌칙조항을 확대하여 규제의 실효성을 강화하고 있는데 이러한 경우 벌칙조항의 구성요건 요소가 되는 인공지능의 기준 및 유형을 대통령령에 위임하는 방식이 죄형법정주의의 관점에서의 위헌 여부를 검토해 볼 필요가 있다.

(나) 인공지능 실증규제 특례 및 우선허용·사후규제 원칙

조인철의원안은 인공지능 실증 규제특례(규제 샌드박스)를, 김성원의원안은 우선허용·사후규제 원칙을 규정하고 있다. 실증규제 특례조항과 우선허용·사후규제 원칙은 기술 및 산업 환경의 변화 속도가 매우 빠른 인공지능 분야의 특성을 고려할 때, 혁신을 촉진하기 위한 규정으로 규제를 완화하는 취지이다.

정보통신기술에 대한 규제 샌드박스과 우선허용·사후규제 원칙은 「정보

115) 한국정보산업연합회, 「인공지능기본법 입법 추진현황 및 산업진흥 측면에서 본 이슈」, 2024. 7., 7쪽

통신 진흥 및 융합 활성화 등에 관한 특별법」 제3조의2 및 제38조의2에 이미 규정되어 있으나, 인공지능기술 및 인공지능제품·서비스의 개발에도 위 조항이 적용된다는 점을 명시적으로 규정하여 확인한다는 의미가 있다.

다만, 인공지능 관련 법안은 진흥과 규제 측면의 균형을 맞추는 것이 법안의 핵심이며 신뢰성과 윤리성에 관한 엄격한 규제 제도의 도입이 예정되는 만큼, 우선허용·사후규제 원칙을 법안에 포함시킬지에 관하여는 신중히 판단해야 할 것으로 보인다.

법안의 요지	
실증규제 특례 (조인철의원안)	혁신적인 인공지능시스템이 시장에 출시되거나 서비스에 투입되기 전에 제한된 기간동안 개발, 테스트 및 검증을 용이하게 하기 위한 인공지능 실증 규제특례를 운영할 수 있도록 함.
우선허용·사후 규제 원칙 (김성원의원안)	국가와 지방자치단체는 인공지능기술의 연구·개발 및 인공지능제품·서비스의 출시를 허용하는 것을 원칙으로 하고, 다만 국민의 생명·안전·권익에 위해가 되거나 공공의 안전 보장, 질서 유지 및 복리 증진을 현저히 저해하는 경우 이를 제한할 수 있음.

(다) 금지된 인공지능에 대한 규제

다른 법안과 달리 권철승의원안은 금지된 인공지능을 별도로 유형화하여, 원칙적 개발·이용의 금지, 예외적 허용을 규정하고 있다.

권철승의원안 제21조(금지된 인공지능의 개발·이용 제한)
① 누구든지 제2조제2호의 금지된 인공지능을 개발하거나 이용하여서는 아니 된다. ② 제1항에도 불구하고 대통령령으로 정하는 바에 따라 다음 각 호의 어느 하나에 해당하는 경우에는 위원회의 심의를 거쳐 금지된 인공지능의 제한적인 개발과 이용을 허용할 수 있다. 1. 범죄피해자 및 「실종아동등의 보호 및 지원에 관한 법률」 제2조제2호에

- 다른 실종아동등에 대한 수색
2. 특정인의 생명이나 신체적 안전 또는 테러 공격에 대한 실질적이고 임박한 위협의 예방
 3. 수사기관 등의 요청에 따라 범인 추적 등 범인 검거를 위한 활동에 반드시 필요한 경우
- ③ 제2항에 따라 금지된 인공지능의 개발과 이용을 제한적으로 허용하는 경우에도 그로 인하여 발생하는 위험과 피해를 최소화할 수 있는 장치를 마련하여야 한다.

다만, 법안의 ‘금지된 인공지능’을 판단하는 기준 및 유형을 대통령령에서 정하도록 위임하고 있어, 법문언만으로 어떠한 인공지능이 금지된 인공지능에 해당하는지 예측하기 어려운 면이 있다. 따라서 금지된 인공지능을 규제하기 이전에 금지된 인공지능을 특정하는 기준과 절차를 법률에 보다 더 구체적으로 규정해야 할 필요가 있다.

(라) 고위험영역 인공지능에 대한 규제

인공지능 법안은 공통적으로 사람의 생명·신체의 안전과 기본권 보호에 중대한 영향을 미칠 우려가 있는 인공지능을 ‘고위험영역 인공지능(고위험 영역에서 활용되는 인공지능, 고위험 인공지능)’으로 정의하고 있으며, 고위험 인공지능의 규제는 대체로 다음과 같이 3개의 내용으로 구성되어 있다.

[표 7] 고위험영역 인공지능 규제 체계

의안	고위험 인공지능 확인	고위험 인공지능 고지의무	고위험 인공지능 사업자의 책무
안철수의원안	제26조	제27조	제28조
정점식의원안	제27조	제26조	제28조
조인철의원안	제25조	제26조	제27조
김성원의원안	제26조	제27조	-
민형배의원안	제28조	제27조	제29조
권칠승의원안	제22조	제23조	제23조
한민수의원안	제26조	제25조	제27조
배준영의원안	제25조	제24조	제26조

법안은 고위험영역 인공지능 규제로서 △ 고위험 인공지능 확인, △ 고위험 인공지능 고지의무, △ 고위험 인공지능 사업자의 책무 등을 규정하고 있다.

사업자의 의무에는 이용자에 대한 사전 고지 의무와, 위험관리방안 및 이용자 보호 방안, 인공지능이 도출한 최종결과에 대한 설명 방안의 수립·운영 등 인공지능의 신뢰성과 안전성을 확보하기 위한 의무가 포함되며, 권칠승의원안과 같이 사업자에게 검·인증 등의 의무를 부과하고 의무 위반에 대한 벌칙을 부과하려는 경우, 검·인증 범위 및 절차 등 의무의 내용을 법률에 구체적으로 규정할 필요가 있다.

구분	주요 내용
고위험 인공지능 확인	인공지능 제품·서비스를 개발·활용·제공하려는 자는 해당 인공지능이 고위험영역 인공지능에 해당하는지에 대한 확인을 과기정통부장관에게 요청할 수 있고, 과기정통부장관은 요청이 있는 경우 이에 대한 확인을 하여야 함.
고위험 인공지능 고지의무	고위험영역 인공지능을 이용하여 제품·서비스를 제공하려는 자는 해당 제품·서비스가 고위험영역 인공지능에 기반하여 운용된다는 사실을 이용자에게 사전에 고지하여야 함.

구분	주요 내용
고위험 인공지능 사업자의 책무	• 고위험영역 인공지능을 개발하거나 이를 이용하여 제품·서비스를 제공하는 자는 인공지능의 신뢰성과 안전성을 확보하기 위한 조치를 하여야 함. • 고위험영역 인공지능을 이용하여 제품·서비스를 제공하려는 자는 인공지능등의 안전성 및 신뢰성 확보를 위한 검·인증을 받아야 함(권칠승의원안).

한편, 황희의원안은 고위험 인공지능의 개발 및 이용을 별도의 장으로 구성하여, 기본계획 수립에 관한 규정을 별도로 두고, 개발사업자의 책무와 이용사업자의 책무를 별도의 조항으로 구분하여 규정하는 한편, 사업자 책임의 일반원칙 규정을 두어 고위험 인공지능에 관한 규제 조항의 체계를 더욱 세분화하고 있다.

나아가 다른 법안에서는 도입되지 않은 이용자의 권리에 관한 사항을 규정하고 있는데, 인공지능 산업 발전 수준을 고려하여 이용자 설명요구권, 분쟁조정제도 등 이용자 보호 관련 규정은 주요국에서도 아직 필요성에 대한 논의가 진행 중이고 실제 도입한 사례를 찾아보기 어렵다는 점을 고려해 충분한 검토 후 필요하다고 인정될 경우 도입하는 것이 바람직하다는 견해가 있다.

황희의원안	
제3장 고위험인공지능의 개발 및 이용	제18조(고위험인공지능 기본계획 수립) 제19조(고위험인공지능개발사업자의 책무) 제20조(고위험인공지능이용사업자의 책무) 제21조(고위험인공지능 이용자의 권리) 제22조(고위험인공지능사업자 책임의 일반원칙)

(마) 생성형 인공지능에 대한 규제

안철수·정점식·민형배·한민수의원안은 생성형 인공지능을 이용하여 제품 또는 서비스를 제공하려는 자에 대하여 고지 및 표시 의무를 부여하는 규

정을 포함하고 있으며, 특히 정점식의원안의 경우 학습에 사용된 누적 연산량이 일정 기준 이상인 경우에 안전 확보 의무를 부과하는 규정을 별도로 두고 있다.

구분	주요 내용
생성형 인공지능 고지 및 표시	생성형 인공지능을 이용하여 제품·서비스를 제공하려는 자는 해당 제품·서비스가 생성형 인공지능에 기반하여 운용된다는 사실을 이용자에게 사전에 고지하고, 결과물이 생성형 인공지능에 의하여 생성되었다는 사실을 표시하여야 함.
생성형 인공지능 안전 확보 의무	학습에 사용된 누적 연산량이 일정 기준 이상인 생성형 인공지능을 이용하여 제품·서비스를 제공하려는 자는 인공지능안전을 확보하기 위한 위험 식별·평가·완화 및 위험관리체계 구축을 하여야 함.

정점식의원안의 안전 확보 의무의 경우, 사람의 생명·신체 등에 전혀 영향을 미치지 않음에도 단순히 누적 연산량이 일정 규모 이상이라고 하여 안전 확보 의무를 부담하도록 하는 것은 과잉금지원칙 및 비례원칙에 반할 우려가 있다는 의견이 있다.

(바) 벌칙

인공지능 법안들은 공통적으로 ‘벌칙 적용에서 공무원 의제’ 조항을 두고 있으며, 위원회, 전문위원회의 위원 중 공무원이 아닌 위원 및 동법에 의해 위탁받은 업무에 종사하는 공공기관 임직원에 대하여 형법 제129조부터 제132조까지의 규정에 따른 벌칙을 적용할 때 공무원으로 의제하는 내용을 포함하고 있다.

‘벌칙’에 관하여 대체로 법안들은 ‘위원회(전문위원회 등 포함)와 관련하여 직무상 알게 된 비밀을 타인에게 누설하거나 직무상 목적 외의 용도로 사

용한 자에 대해 3년 이하 징역 또는 3천만원 이하 벌금형을 부과하는 처벌 규정을 두고 있다(황희의원안 제외). 권철승의원안은 다른 제정안들과 달리 규제의 실효성을 확대하는 차원에서 금지된·고위험 인공지능과 관련된 의무를 위반한 사업자에 대하여도 벌칙을 부과하는 규정을 두고 있다.

한민수·배준영의원안은 ‘과학기술정보통신부장관은 인공지능정책센터 또는 이와 유사한 명칭을 사용한 자에게 500만 원 이하의 과태료를 부과·징수한다’는 과태료 규정을 두고 있다.

[표 8] 제22대 국회 발의 인공지능 법안별 조문 체계 비교

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙	제1장 총칙
제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)	제1조(목적)
제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)	제2조(정의)
제3조 (기본원칙)	제3조 (기본원칙)	제3조 (기본원칙)	제3조 (기본원칙)	제3조 (기본원칙)	제3조 (인공지능 개발 및 이용의 기본원칙)	제3조 (기본원칙)	제3조 (인공지능 개발 및 이용의 기본원칙)	제3조 (기본원칙)
-	-	-	-	-	제4조 (국가·지방자 치단체 등의 책무)	-	제4조 (국가·지방자 치단체 등의 책무)	-
-	제4조 (국외행위에 대한 적용)	-	-	-	-	-	-	-
-	-	-	-	-	-	-	제5조(인공 지능 사업자의 의무 등)	-
제4조 (다른 법률과의 관계)	제5조 (다른 법률과의 관계)	제4조 (다른 법률과의 관계)	제4 조(다른 법률과의 관계)	제4조 (다른 법률과의 관계)	제5조 (다른 법률과의 관계)	제4조 (다른 법률과의 관계)	제6조 (다른 법률과의 관계)	제4조 (다른 법률과의 관계)
제2장 인공지능산업 육성 및 신뢰확보를 위한 추진체계	제2장 인공지능의 건전한 발전과 신뢰 기반 조성을 위한 추진체계		제2장 인공지능산업 육성 및 신뢰확보를 위한 추진체계	제2장 인공지능의 건전한 발전과 신뢰 기반 조성을 위한 추진체계	제2장 인공지능 정책의 수립 및 추진체계	제2장 인공지능의 건전한 발전과 신뢰 기반 조성을 위한 추진체계		제2장 인공지능의 건전한 발전과 신뢰 기반 조성을 위한 추진체계

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
제5조 (인공지능 기본계획의 수립)	제6조 (인공지능 기본계획의 수립)	제5조 (인공지능 기본계획의 수립)	제5조 (인공지능 기본계획의 수립)	제10조 (국가 인공지능 기본계획의 수립)	제6조 (인공지능 기본계획)	제5조 (인공지능 기본계획의 수립)	-	제5조 (인공지능 기본계획의 수립)
-	-	-	-	-	제7조 (인공지능 실행계획)	-	-	-
-	-	-	-	-	-	-	-	-
-	-	-	-	제11조 (지역 인공지능 종합계획의 수립)	-	-	-	-
-	-	-	-	제12조 (지역계획의 심의와 조정)	-	-	-	-
제6조 (인공지능 위원회)	제7조 (국가 인공지능 위원회)	제6조 (국가 인공지능 위원회)	제6조 (인공지능 위원회)	제5조 (국가 인공지능 위원회)	제8조 (국가 인공지능 위원회)	제6조 (인공지능 위원회)	-	제6조 (인공지능 위원회)
제7조 (위원회의 기능)	제8조 (위원회의 기능)	제7조 (위원회의 기능)	제7조 (위원회의 기능)	제6조 (국가 위원회의 기능)	-	제7조 (위원회의 기능)	-	제7조 (위원회의 기능)
제8조 (위원의 제척·기피 및 회피)	제9조 (위원의 제척·기피 및 회피)	제8조 (위원의 제척·기피 및 회피)	제8조 (위원의 제척·기피 및 회피)	제7조 (위원의 제척·기피 및 회피)	-	제8조 (위원의 제척·기피 및 회피)	-	제8조 (위원의 제척·기피 및 회피)
제9조 (전문위원회 등)	제10조 (분과위원회 등)	제9조 (전문위원회 등)	제9조 (전문위원회 등)	제8조 (전문위원회 등)	-	제9조 (전문위원회 등)	-	제9조 (전문위원회 등)
-	-	-	-	제9조 (지역인공지능 위원회의 구성 등)	-	-	-	-
제10조 (국가인공지능 센터)	제11조 (국가인공지능 센터)	-	제10조 (국가인공지능 센터)	제13조 (인공지능 정책센터)	제9조 (국가인공지능 센터)	제10조 (인공지능 정책센터)	-	제10조 (인공지능 정책센터)
제11조 (인공지능안전 연구소)	제12조 (인공지능안전 연구소)	-	-	-	-	-	-	-
제3장 인공지능 기술 개발 및 산업 육성	제3장 인공지능 기술 개발 및 산업 육성	제2장 인공지능 기술 개발 및 산업 육성	제3장 인공지능 기술 개발 및 산업 육성	제3장 인공지능 기술 개발 및 산업 육성	제3장 인공지능 기술 개발 및 활용기반 조성	제3장 인공지능 기술 개발 및 산업 육성	제2장 인공지능의 진흥	제3장 인공지능 기술 개발 및 산업 육성

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
제1절 인공지능 산업기반 조성	제1절 인공지능 산업기반 조성	제1절 인공지능 산업기반 조성	제1절 인공지능 산업기반 조성	제1절 인공지능 산업기반 조성		제1절 인공지능 산업 기반 조성		제1절 인공지능 산업기반 조성
-	-	-	제11조 (우선허용· 사후규제 원칙)	-	-	-	-	-
제12조 (인공지능기술 개발 및 안전한 이용 지원)	제13조 (인공지능기술 개발 및 안전한 이용 지원)	제10조 (인공지능기술 개발 및 안전한 이용 지원)	제12조 (인공지능기술 개발 및 안전한 이용 지원)	제14조 (인공지능기술 개발 및 안전한 이용 지원)	제10조 (인공지능기술 개발 및 안전한 이용 촉진)	제11조 (인공지능기술 개발 및 안전한 이용 지원)	-	제11조 (인공지능기술 개발 및 안전한 이용 지원)
-	-	-	-	-	-	-	제7조 (인공지능 기술의 개발)	-
-	-	-	-	-	-	-	제8조 (기술기준)	-
제13조 (인공지능 기술의 표준화)	제14조 (인공지능 기술의 표준화)	제11조 (인공지능 기술의 표준화)	제13조 (인공지능 기술의 표준화)	제15조 (인공지능 기술의 표준화)	제11조 (인공지능 기술의 표준화)	제12조 (인공지능 기술의 표준화)	제9조 (인공지능 기술의 표준화)	제12조 (인공지능 기술의 표준화)
제14조 (인공지능 학습용데이터 관련 시책의 수립 등)	제15조 (인공지능 학습용데이터 관련 시책의 수립 등)	제12조 (인공지능 학습용데이터 관련 시책의 수립 등)	제14조 (인공지능 학습용데이터 관련 시책의 수립 등)	제16조 (인공지능 학습용데이터 관련 시책의 수립 등)	제12조 (인공지능 학습용데이터 관련 시책의 수립 등)	제13조 (인공지능 학습용데이터 관련 시책의 수립 등)	-	제13조 (인공지능 학습용데이터 관련 시책의 수립 등)
제2절 인공지능기술 개발 및 인공지능산업 활성화	제2절 인공지능기술 개발 및 인공지능산업 활성화	제2절 인공지능기술 개발 및 인공지능산업 활성화	제2절 인공지능기술 개발 및 인공지능산업 활성화	제2절 인공지능기술 개발 및 인공지능산업 활성화	제4장 인공지능산업 육성 및 활성화 지원	제2절 인공지능기술 개발 및 인공지능산업 활성화	-	제2절 인공지능기술 개발 및 인공지능산업 활성화
제15조 (기업의 인공지능기술 도입·활용 지원)	제16조 (기업의 인공지능기술 도입·활용 지원)	제13조 (기업의 인공지능기술 도입·활용 지원)	제15조 (기업의 인공지능기술 도입·활용 지원)	제17조 (기업의 인공지능기술 도입·활용 지원)	제13조 (기업의 인공지능등 도입·활용 지원)	제14조 (기업의 인공지능기술 도입·활용 지원)	-	제14조 (기업의 인공지능기술 도입·활용 지원)
-	-	제14조 (인공지능 실증 규제특례)	-	-	-	-	-	-
-	-	-	-	-	-	-	제11조 (선도사업의 추진과 지원)	-
-	-	-	-	-	-	-	제12조 (인공 지능의 실용화· 사업화 지원)	-

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
-	-	-	-	-	-	제15조 (중소기업자를 위한 특별지원)	제13조 (중소기업에 대한 지원)	-
-	-	-	-	-	-	-	제14조 (정보분석을 위한 복제·전송)	-
제16조 (창업의 활성화)	제17조 (창업의 활성화)	제15조 (창업의 활성화)	제16조 (창업의 활성화)	제18조 (창업의 활성화)	제14조 (창업의 활성화)	제16조 (창업의 활성화)	-	제15조 (창업의 활성화)
제17조 (인공지능 융합의 촉진)	제18조 (인공지능 융합의 촉진)	제16조 (인공지능 융합의 촉진)	제17조 (인공지능 융합의 촉진)	제19조 (인공지능 융합의 촉진)	제15조 (인공지능 융합 등의 촉진)	제17조 (인공지능 융합의 촉진)	-	제16조 (인공지능 융합의 촉진)
제18조 (제도개선 등)	제19조 (제도개선 등)	제17조 (제도개선 등)	제18조 (제도개선 등)	제20조 (제도개선 등)	-	제18조 (제도개선 등)	-	제17조 (제도개선 등)
제19조 (전문인력의 확보)	제20조 (전문인력의 확보)	제18조 (전문인력의 확보)	제19조 (전문인력의 확보)	제21조 (전문인력의 확보)	제16조 (전문인력의 양성 등)	제19조 (전문인력의 확보)	제10조 (전문인력의 양성)	제18조 (전문인력의 확보)
제20조 (국제협력 및 해외시장 진출의 지원)	제21조 (국제협력 및 해외시장 진출의 지원)	제19조 (국제협력 및 해외시장 진출의 지원)	제20조 (국제협력 및 해외시장 진출의 지원)	제22조 (국제협력 및 해외시장 진출의 지원)	제17조 (국제협력 및 해외시장 진출의 지원)	제20조 (국제협력 및 해외시장 진출의 지원)	-	제19조 (국제협력 및 해외시장 진출의 지원)
제21조 (인공지능집적 단지 지정 등)	제22조 (인공지능집적 단지 지정 등)	제20조 (인공지능집적 단지 지정 등)	제21조 (인공지능집적 단지 지정 등)	제23조 (인공지능집적 단지 지정 등)	-	제21조 (인공지능집적 단지 지정 등)	-	제20조 (인공지능집적 단지 지정 등)
제22조 (대한인공지능 진흥협회의 설립)	-	제21조 (대한인공지능 협회의 설립)	제22조 (대한인공지능 협회의 설립)	-	-	-	-	-
제4장 인공지능 윤리 및 신뢰성 확보	제4장 인공지능 윤리 및 신뢰성 확보	제3장 인공지능 신뢰 확보	제4장 인공지능 신뢰 확보	제4장 인공지능 윤리 및 신뢰성 확보	제5장 인공지능 윤리 및 안전 확보	제4장 인공지능 윤리 및 신뢰성 확보	-	제4장 인공지능 윤리 및 신뢰성 확보
-	-	-	-	-	-	-	제16조 (인공지능 윤리 교육의 실시)	-
-	-	-	-	-	-	-	제17조 (규제의 원칙)	-
제23조 (인공지능 윤리성·신뢰성 원칙 등)	제23조 (인공지능 윤리원칙 등)	제22조 (인공지능 윤리원칙 등)	제23조 (인공지능 윤리원칙 등)	제24조 (인공지능 윤리원칙 등)	제18조 (인공지능 윤리원칙의 확립)	제22조 (인공지능 윤리원칙 등)	-	제21조 (인공지능 윤리에 관한 원칙 등)
제24조 (신뢰할 수 있는 인공지능)	제24조 (신뢰할 수 있는 인공지능)	제23조 (신뢰할 수 있는 인공지능)	제24조 (신뢰할 수 있는 인공지능)	제25조 (신뢰할 수 있는 인공지능)	제19조 (인공지능등의 안전성· 신뢰성 확보)	제23조 (신뢰할 수 있는 인공지능)	-	제22조 (신뢰할 수 있는 인공지능)

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
제25조 (인공지능 신뢰성 검·인증 지원 등)	제25조 (인공지능 신뢰성 검·인증 지원 등)	제24조 (인공지능 신뢰성 검·인증 지원 등)	제25조 (인공지능 신뢰성 검·인증 지원 등)	제26조 (인공지능 신뢰성 검·인증 지원 등)	제20조 (인공지능등의 신뢰성 등 검·인증 지원)	제24조 (인공지능 신뢰성 검·인증 지원 등)	-	제23조 (인공지능 신뢰성 검·인증 지원 등)
							제3장 고위험인공지 능의 개발 및 이용	
-	-	-	-	-	제21조 (금지된 인공지능의 개발·이용 제한)	-	-	-
-	-	-	-	-	-	-	제18조 (고위험인공지 능 기본계획 수립)	-
제26조 (고위험영역 인공지능의 확인)	제27조 (고위험영역 인공지능의 확인)	제25조 (고위험영역 인공지능의 확인)	제26조 (고위험 영역에서 활용되는 인공지능의 확인)	제28조 (고위험영역 인공지능의 확인)	제22조 (고위험 인공지능의 확인)	제26조 (고위험영역 인공지능의 확인)	-	제25조 (고위험영역 인공지능의 확인)
제27조 (고위험영역 인공지능 고지의무)	제26조 (고위험영역 인공지능 고지의무)	제26조 (고위험영역 인공지능의 고지의무)	제27조 (고위험 영역에서 활용되는 인공지능의 고지 의무)	제27조 (고위험영역 인공지능 고지 의무)	제23조 (고위험 인공지능 제공자의 의무)	제25조 (고위험영역 인공지능 고지의무)	-	제24조 (고위험영역 인공지능 고지의무)
제28조 (고위험영역 인공지능과 관련한 사업자의 책임)	제28조 (고위험영역 인공지능과 관련한 사업자의 책임)	제27조 (고위험영역 인공지능과 관련한 사업자의 책임)	-	제29조 (고위험영역 인공지능과 관련한 사업자의 책임)	제23조 (고위험 인공지능 제공자의 의무)	제27조 (고위험영역 인공지능과 관련한 사업자의 책임)	제19조 (고위험인공지 능 개발사업자의 책임) 제20조 (고위험인공지 능 이용사업자의 책임)	제26조 (고위험영역 인공지능과 관련한 사업자의 책임)
-	-	-	-	-	-	-	제21 조 (고위험 인공지능 이용자의 권리)	-
-	-	-	-	-	-	-	제22조 (고위험 인공지능 사업자 책임의 일반원칙)	-

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
-	-	제28조 (인공지능제품의 비상정지)	-	-	-	-	-	-
제29조 (생성형 인공지능 고지 및 표시)	제29조 (생성형 인공지능 고지 및 표시)	-	-	제30조 (생성형 인공지능 고지 및 표시)	-	제28조 (생성형 인공지능 고지 및 표시)	-	제27조 (생성형 인공지능 고지 및 표시)
-	제30조(생성형 인공지능 안전 확보 의무)	-	-	-	-	-	-	-
-	-	-	-	-	-	제29조 (국내대리인 지정)	-	-
제30조 (민간자율 인공지능 윤리위원회의 설치 등)	제31조 (민간자율 인공지능 윤리위원회의 설치 등)	제29조 (민간자율 인공지능 윤리위원회의 설치 등)	-	제31조 (민간자율 인공지능 윤리위원회의 설치 등)	-	제30조 (민간자율 인공지능 윤리위원회의 설치 등)	-	제28조 (민간자율 인공지능 윤리위원회의 설치 등)
							제4장 인공지능 관련 분쟁조정	
-	-	-	-	-	-	-	제23조 (인공지능분쟁 조정위원회)	-
-	-	-	-	-	-	-	제24조 (위원의 신분보장)	-
-	-	-	-	-	-	-	제25조 (위원의 제척·기피· 회피)	-
-	-	-	-	-	-	-	제26조 (조정 신청 등)	-
-	-	-	-	-	-	-	제27조 (처리기간)	-
-	-	-	-	-	-	-	제28조 (자료의 요청 등)	-
-	-	-	-	-	-	-	제29조 (조정 전 합의 권고)	-
-	-	-	-	-	-	-	제30조 (분쟁의 조정)	-

안철수 의원안	정점식 의원안	조인철 의원안	김성원 의원안	민형배 의원안	권철승 의원안	한민수 의원안	황희 의원안	배준영 의원안
-	-	-	-	-	-	-	제31조 (조정의 거부 및 중지)	-
-	-	-	-	-	-	-	제32조 (조정절차 등)	-
제5장 보칙	제5장 보칙	제4장 보칙	제5장 보칙	제5장 보칙	제6장 보칙	제5장 보칙	제5장 보칙	제5장 보칙
제31조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	제32조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	제30조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	제28조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	제32조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	-	제31조 (인공지능 산업의 진흥을 위한 재원의 확충 등)	-	제29조 (인공지능 산업의 진흥을 위한 재원의 확충 등)
제32조 (실태조사, 통계 및 지표의 작성)	제33조 (실태조사, 통계 및 지표의 작성)	제31조 (실태조사, 통계 및 지표의 작성)	제29조 (실태조사, 통계 및 지표의 작성)	제33조 (실태조사, 통계 및 지표의 작성)	제24조 (실태조사, 통계 및 지표의 작성 등)	제24조 (실태조사, 통계 및 지표의 작성 등)	-	제30조 (실태조사, 통계 및 지표의 작성)
제33조 (권한의 위임 및 업무의 위탁)	제34조 (권한의 위임 및 업무의 위탁)	제32조 (권한의 위임 및 업무의 위탁)	제30조 (권한의 위임 및 업무의 위탁)	제34조 (권한의 위임 및 업무의 위탁)	제25조 (업무의 위임 및 위탁)	제33조 (권한의 위임 및 업무의 위탁)	제33조 (권한의 위임·위탁)	제31조 (권한의 위임 및 업무의 위탁)
					제7장 벌칙			
제34조 (벌칙 적용에서 공무원 의제)	제35조 (벌칙 적용에서 공무원 의제)	제33조 (벌칙 적용에서 공무원 의제)	제31조 (벌칙 적용에서 공무원 의제)	제35조 (벌칙 적용에서 공무원 의제)	제26조 (벌칙 적용에서 공무원 의제)	제34조 (벌칙 적용에서 공무원 의제)	제34조 (벌칙 적용에서 공무원 의제)	제32조 (벌칙 적용에서 공무원 의제)
제35조(벌칙)	제36조(벌칙)	제34조(벌칙)	제32조(벌칙)	제36조(벌칙)	제27조(벌칙)	제35조(벌칙)	-	제33조(벌칙)
-	-	-	-	-	-	제34조 (과태료)	-	제34조 (과태료)

4. 허위조작정보(딥페이크) 규제체계 논의 현황

가. 문제점

딥페이크(Deepfake)는 딥러닝(Deep learning)과 가짜(fake)의 합성어로, 머신러닝과 딥러닝을 포함한 인공지능 기술을 사용하여 제작되어 실제처럼 보이지만 실제 사람이 말하거나 행동하지 않은 것처럼 보이는 조작 또는 합성된 오디오 또는 시각 미디어를 의미한다.¹¹⁶⁾

딥페이크는 합성 콘텐츠임에도 불구하고 실제와 구분이 어렵기 때문에 허위 사실을 유포하는데 용이하다. 실제와 구분이 어렵다는 딥페이크의 기능을 활용하여 테일러 스위프트의 딥페이크 성착취물 영상, 로맨스 스캠과 같은 금융사기도 증가하고 있으며, 딥페이크 음원을 만든 뒤 원곡자의 동의를 받지 않고 배포하여 저작권을 침해하고, 바이든 대통령의 목소리로 정치적 목적의 딥페이크 영상을 만들어 민주주의를 위협하는 등 지속적으로 문제가 대두된다. 위와 같이 딥페이크로 인하여 일반공중이 딥페이크로 만든 영상이 실제로 존재하는 것으로 오인·기망을 당할 수 있으며, 딥페이크 영상을 제작하는 과정에서 타인의 성명, 이미지, 목소리, 유사성 등을 무단 사용할 수 있다는 점에서 문제가 된다.

최근 국내에서는 여학생들의 얼굴을 도용해 딥페이크 기술로 성착취물을 만들고 텔레그램을 이용하여 판매하고, 전국적으로 지인 얼굴 사진들을 합성하여 딥페이크 음란물을 제작하는 사례가 늘어나면서 국내에서도 많은 문제점이 드러나고 있다. 국내에서는 2020년 성폭력처벌법을 개정하여 딥페이크 성범죄에 관한 처벌 규정이 생겼으나, 딥페이크 음란물의 단순 시청 또는 반포할 목적 없는 제작 등의 경우에는 처벌 규정이 없어 처벌 공백이 있다는 지적이 이어져 오고 있다.

116) European Parliamentary Research Service, 「Tackling deepfakes in European policy」, 2021. 7., p.2.

나. 국내외 입법 동향

(1) 국내

(가) 법제도 현황

현재 인공지능을 전반적으로 포괄하는 통합법은 없으며, 딥페이크를 악용한 행위 등에 대해서는 「형법」, 「정보통신망 이용촉진 및 정보보호 등에 관한 법률」로 대응해 왔다. 이에 2020년 3월에 「성폭력범죄의처벌 등에 관한 특례법」을 개정하여 별도의 처벌규정¹¹⁷⁾을 마련하였고, 2023년 12월에 「공직선거법」을 일부 개정¹¹⁸⁾하여 딥페이크 영상을 활용한 선거운동을 금지하였다.

117) 성폭력범죄의 처벌 등에 관한 특례법 제14조의2(허위영상물 등의 반포등) ① 반포등을 할 목적으로 사람의 얼굴·신체 또는 음성을 대상으로 한 촬영물·영상물 또는 음성물(이하 이 조에서 “영상물등”이라 한다)을 영상물등의 대상자의 의사에 반하여 성적 욕망 또는 수치심을 유발할 수 있는 형태로 편집·합성 또는 가공(이하 이 조에서 “편집등”이라 한다)한 자는 5년 이하의 징역 또는 5천만원 이하의 벌금에 처한다.

② 제1항에 따른 편집물·합성물·가공물(이하 이 항에서 “편집물등”이라 한다) 또는 복제물(복제물의 복제물을 포함한다. 이하 이 항에서 같다)을 반포등을 한 자 또는 제1항의 편집등을 할 당시에는 영상물등의 대상자의 의사에 반하지 아니한 경우에도 사후에 그 편집물등 또는 복제물을 영상물등의 대상자의 의사에 반하여 반포등을 한 자는 5년 이하의 징역 또는 5천만원 이하의 벌금에 처한다.

③ 영리를 목적으로 영상물등의 대상자의 의사에 반하여 정보통신망을 이용하여 제2항의 죄를 범한 자는 7년 이하의 징역에 처한다.

④ 상습으로 제1항부터 제3항까지의 죄를 범한 때에는 그 죄에 정한 형의 2분의 1까지 가중한다.

118) 공직선거법 제82조의8(딥페이크영상등을 이용한 선거운동) ① 누구든지 선거일 전 90일부터 선거일까지 선거운동을 위하여 인공지능 기술 등을 이용하여 만든 실제와 구분하기 어려운 가상의 음향, 이미지 또는 영상 등(이하 “딥페이크영상등”이라 한다)을 제작·편집·유포·상영 또는 게시하는 행위를 하여서는 아니 된다.

(나) 관련 입법 추진 동향

1) 인공지능 발전과 신뢰 기반 조성 등에 관한 법률안(제2200543호, 정점식의원 대표발의)¹¹⁹⁾

생성형 인공지능 사업자가 생성형 인공지능 운용 사실 고지 및 표시를 하도록 하고, 일정 성능 이상의 생성형 인공지능의 경우 안전을 확보하기 위한 조치를 이행하도록 하였다(안 제29조 및 제30조).

2) 인공지능산업 육성 및 신뢰 확보에 관한 법률안(제2200675호, 김성원의원 대표발의)

고위험 영역에서 활용되는 인공지능을 이용하여 제품 또는 서비스를 제공하려는 자는 해당 제품 또는 서비스가 고위험 영역에서 활용되는 인공지능에 기반하여 운용된다는 사실을 이용자에게 사전에 고지하도록 하였다(안 제27조).

3) 인공지능기술 기본법안(제2201158호, 민형배의원 대표발의)

고위험영역 인공지능을 이용하여 제품 또는 서비스를 제공하려는 자는 이용자에게 사전에 고지하도록 하고, 과학기술정보통신부장관은 요청이 있으면 고위험영역 인공지능의 해당 여부에 관한 확인을 하도록 하였다(안 제27조 및 제28조).

4) 인공지능서비스 이용자 보호에 관한 법률 제정 추진(방송통신위원회)

AI 서비스의 신뢰성을 보장하고 역기능으로부터 이용자를 보호하기 위하여 AI 생성한 콘텐츠를 게시할 경우 AI 생성물 표시를 의무화하였다.¹²⁰⁾

- ② 누구든지 제1항의 기간이 아닌 때에 선거운동을 위하여 딥페이크영상등을 제작·편집·유포·상영 또는 게시하는 경우에는 해당 정보가 인공지능 기술 등을 이용하여 만든 가상의 정보라는 사실을 명확하게 인식할 수 있도록 중앙선거관리위원회규칙으로 정하는 바에 따라 해당 사항을 딥페이크영상등에 표시하여야 한다.

119) 국회 과학기술정보방송통신위원회, 「인공지능 발전과 신뢰 기반 조성 등에 관한 법률안(정점식의원 대표발의, 의안번호 제2200543호) 검토보고서」, 2024. 8.

5) 콘텐츠산업진흥법 일부개정법률안 발의

AI 생성물에 대하여 표시(labeling)하도록 콘텐츠산업 진흥법 일부 개정법률안이 발의된 상태이다. 콘텐츠산업 진흥법 개정안은 콘텐츠제작자에게 대통령령으로 정하는 AI기술을 이용하여 콘텐츠를 제작한 경우에는 해당 콘텐츠가 인공지능 기술을 이용하여 제작된 콘텐츠라는 사실을 표시하도록 하고 표시의 내용과 방법에 필요한 사항은 대통령령으로 정하도록 하였다.¹²¹⁾

이 법안에 대한 이해관계자들의 의견은 다음과 같다.¹²²⁾

[표 9] 허위조작정보 규제(콘텐츠산업진흥법 개정안)에 대한 이해관계자 의견

이해관계자	의견
과학기술정보통신부	• 개정안의 “표시의무 부과”는 AI 사업자가 아닌 불특정 다수의 일반인을 포함한 콘텐츠제작자를 규제하는 사항으로 신중한 접근 필요 • AI 기본법안 제정 후 단계적 마련이 바람직
방송통신위원회	• AI 기술 및 콘텐츠 유형, 표시 내용과 방법 등 핵심 사항들을 모두 시행령으로 위임하고 있어 규제 예측성이 낮고 과도한 포괄입법 우려 존재 • AI콘텐츠의 ‘제작’뿐만 아니라 ‘유통·확산’ 과정에서의 이용자의 알권리를 보호하기 위한 방안도 종합적으로 고려하여 의무를 신설하는 것이 타당
한국음악저작권협회	• 저작권신탁관리제도에 따른 저작권사용료의 정확한 징수 및 분배를 위해 인공지능 활용 콘텐츠의 구별이 필수적 • 콘텐츠 제작자의 자율성에 의존하는 것만으로는 실효성이 없으며, 이미 저작권 분야에서 발생하고 있는 혼란과 피해의 추가 확산 방지가 시급

120) 문화체육관광부 보도자료(2024.2.19.).

121) 콘텐츠산업진흥법 일부개정법률안[제21대 국회에서 이상헌 의원이 대표발의(의안번호 2122180, 2023. 5. 22.)하였고, 2024. 5. 29. 임기만료 폐지되었으나, 제22대 국회에서 강유정 의원이 또다시 대표발의(의안번호 220048, 2024. 5. 31.)].

122) 국회 문화체육관광위원회, 콘텐츠산업 진흥법 일부개정법률안(강유정의원 대표발의, 의안번호 제2200048호) 검토보고서, 2024. 8.

이해관계자	의견
	• 해당 법안은 인공지능의 활용을 제한하는 것이 아니라, 오히려 인공지능 생성물임이 정확히 표기된다면 해당 콘텐츠를 이용함에 따라 추후 발생할 수 있는 각종 비용 부담 및 법적 위험성을 고려하지 않아도 되어 인공지능 생성물에 대한 이용이 촉진될 가능성이 높음
방송협회	• 인공지능에 대한 정의가 모호하며, 명확한 정의 없는 선부른 규제는 법 취지와 달리 산업 생산성을 저해함 • 콘텐츠 제작 방식, 기여도 등 AI를 활용하는 콘텐츠 산업 실정에 맞게 검토 필요
TV조선	• 인공지능으로 만든 콘텐츠의 범위 불명확 • 국내에서 인공지능으로 만들어진 콘텐츠로 인해 사용자들에게 미치는 부작용 사례는 아직 대두된바 없음 • 사업자 자율적으로 관련 사항을 고지하여 시청자에게 정확한 정보를 제공할 수 있는 환경을 조성하는 것이 필요
채널A	• 방송법 등에 출처 명시를 강조하는 조항들이 이미 다수 존재 • 개정안은 방송사들의 부담을 키우는 규제 • 방송제작과 보도의 자유는 방송에 있어 중요한 가치로서, 관련 규제와 권고 기준 마련 등은 최소화 필요
MBN	• ‘인공지능 기술을 이용하여 콘텐츠를 제작하는 경우’라는 기준이 모호해 사업자들의 혼선 우려 • 각 방송사업자별로 자율적인 운영에 맡기는 방안을 포함하여 본 개정안에 대한 신중한 재검토 필요
JTBC	• 규제 대상이 되는 ‘인공지능 기술을 이용한 콘텐츠’ 불명확 • 규제의 본질인 ‘표시내용 및 방법’을 모두 대통령령에 위임하고 있어 ‘포괄위임금지의 원칙’에 반함

6) 정보통신망법 일부개정법률안 발의

제21대에 발의되었던 정보통신망법 일부개정법률안에 따르면 정보제공자가 AI기술을 이용하여 만든 실체와 구분하기 어려운 가상의 음화·화상·영상

등의 정보 중 대통령령으로 정하는 정보를 제공하려는 경우에는 해당 정보가 인공지능 기술을 이용하여 만든 가상의 정보라는 사실을 명확하게 인식할 수 있도록 표시하고, 표시의 방법 등에 필요한 사항은 대통령령으로 정하고, 표시 방법을 지키지 아니하는 AI기술을 이용하여 만들어 게재·전시되어 있는 실제와 구분하기 어려운 가상의 음향·화상·영상 등의 정보를 지체없이 삭제하고, AI 기술을 이용하여 만든 실제와 구분하기 어려운 가상의 정보임을 표시하지 아니한 자에 대해서는 과태료를 부과하도록 하였다.¹²³⁾

제22대 국회에서 발의된 정보통신망법 일부개정법률안은 다음과 같다.

[표 10] 허위조작정보 규제(정보통신망법 개정안)에 대한 이해관계자 의견

순번	대표발의	주요내용
1	김승수 의원 2024. 6. 20. (의안번호 2200713)	(AI생성물 표시 의무) 정보를 제공하려는 자에게 정보인공지능 기술을 이용하여 만든 가상의 정보 표시의무 부과 및 표시의무 불이행시 정보 삭제의무 부과
2	박용갑 의원 2024. 8. 27. (의안번호 2203253)	(딥페이크 콘텐츠 금지, 시책 추진) △ 딥페이크 기술을 악용한 촬영물·영상물 또는 음성물을 편집·합성·가공한 정보의 유통을 금지하고 위반 시 처벌, △ 합성 영상으로 인한 명예훼손 등의 피해 실태 파악, 합성 영상 유통 실태 및 관련 국내외 기술 동향 파악, 유통 방지를 위한 기술 개발의 촉진, 교육·홍보 등의 시책 추진
3	우재준 의원 2024. 8. 28. (의안번호 2203275)	(딥페이크 콘텐츠 금지) 타인에게 손해를 끼칠 목적으로 공공연하게 거짓의 사실을 드러내는 허위 내용의 정보 및 타인의 의사에 반하여 사람의 얼굴·신체 또는 음성을 대상으로 한 촬영물·영상물 또는 음성물을 편집·합성·가공한 정보의 유통을 금지하고 위반 시 처벌

123) 정보통신망법 일부 개정법률안[제21대 국회에서 김승수 의원이 대표발의(의안번호 2126048, 2023. 12. 22.)하였고, 2024. 5. 29. 임기만료 폐지되었으나, 제22대 국회에서 김승수 의원이 또다시 대표발의(의안번호2200713, 2024. 6. 20.)]

순번	대표발의	주요내용
4	김장겸 의원 2024. 8. 29. (의안번호 2203323)	(AI생성물 표시 의무) 정보통신서비스 제공자에게 인공지능 기술을 이용하여 만든 가상의 정보라는 사실을 표시할 수 있는 기능을 마련하도록 하고, 해당 표시방법을 지키지 아니한 정보를 탐지하기 위한 노력 및 삭제 등의 유통 방지 의무를 부과하며, 정보를 게재하려는 자에게는 표시의무 부과와 함께 정당한 이유 없이 표시를 제거하거나 변경하는 것을 금지
5	김용민 의원 2024. 9. 2. (의안번호 2203479)	(딥페이크 콘텐츠 금지) 정보통신서비스 제공자로 하여금 불법촬영물과 같은 수준으로 딥페이크 허위영상물에 대해서도 즉각적인 조치를 취하도록 하고, 의무 위반 시 징벌적 손해배상 도입
6	박충권 의원 2024. 9. 2. (의안번호 2203542)	(딥페이크 콘텐츠 금지, 시책 추진) △ 인공지능 기술을 이용하여 사람의 얼굴·신체 또는 음성을 대상으로 한 촬영물·영상물 또는 음성물을 대상자의 의사에 반하여 편집·합성 또는 가공한 정보의 유통을 금지하고 위반 시 처벌, △ 딥페이크로 인한 명예훼손 등의 피해 실태 파악, 딥페이크 영상 유통 실태 및 관련 국내외 기술 동향 파악, 유통 방지를 위한 기술 개발의 촉진, 교육·홍보 등의 시책 추진
7	김현정 의원 2024. 9. 3. (의안번호 2203592)	(시책 추진) 인공지능 기술을 이용하여 만든 거짓의 음향·화상 또는 영상 등 정보 전반으로 인한 명예훼손 등의 피해 실태 파악, 인공지능 기술을 이용하여 만든 거짓의 음향·화상 또는 영상 등 정보 전반 유통 실태 및 관련 국내외 기술 동향 파악, 유통 방지를 위한 기술 개발 촉진, 교육·홍보 등의 시책 추진
8	조인철 의원 2024. 9. 4. (의안번호 2203627)	(AI생성물 표시의무) 정보통신망에서 유통이 금지되는 정보의 범위에 허위조작정보 및 「성폭력범죄의 처벌 등에 관한 특례법」 제14조의2에 따른 편집물·합성물·가공물 또는 복제물이 포함됨을 명확히 하고, 인공지능 기술 등을 이용하여 만든 실제와 구분하기 어려운 가상의 음향·화상 또는 영상 등의 정보를 제작·편집·유포·상영 또는 게시하는 경우 해당 정보가 인공지능 기술 등을 이용하여 만든 가상의 정보라는 사실을 표시

(2) 미국

(가) 연방 차원

1) 딥페이크 책임법안¹²⁴⁾

“딥페이크 책임법안(DEEPFAKES Accountability Act)”은 딥페이크 기술로 인한 위협으로부터 국가안보를 보호하고 유해한 딥페이크의 피해자들에게 법적 구제 수단을 제공하기 위한 목적으로 발의되었다. 이 법안은 첨단 기술의 허위 사칭 기록(advanced technological false personation record) 제작자에 대한 공개(disclose) 의무, 공개 의무 미준수 또는 공개 내용 변경시의 민·형사상 처벌, 딥페이크 피해자 지원, 딥페이크 감지(Detection of deepfakes) 등에 관한 내용을 포함하고 있다.

2) 딥페이크 신용사기 방지법안¹²⁵⁾

“딥페이크 신용사기 방지법안(Preventing Deep Fake Scams Act)”은 금융서비스 부문의 인공지능 관련 문제에 관해 의회에 보고할 금융서비스 부문 인공지능 태스크포스를 설립하고, 기타 목적을 달성하기 위한 목적으로 발의되었다.

3) AI 행정명령

미국은 「인공지능의 안전하고 보안적이며 신뢰할 수 있는 개발 및 활용에 관한 행정명령(Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence)」(이하 ‘AI 행정명령’)을 2023년 10월 30일

124) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/house-bill/5586/text>> 최종 방문일 2024. 9. 11.>

125) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/house-bill/5808/text>> 최종 방문일 2024. 9. 11.>

발령하였다.¹²⁶⁾

AI 행정명령은 ① AI를 위한 새로운 안전 및 보안기준 마련, ② 개인정보 보호, ③ 형평과 시민권 증진, ④ 소비자 보호, ⑤ 노동자 지원, ⑥ 혁신과 경쟁 촉진, ⑦ 국제 파트너와의 협력, ⑧ 연방정부의 AI 사용과 조달을 위한 지침 개발 총 8개 부분으로 구성되어 있다. 딥페이크와 관련하여 “합성 콘텐츠”에 대한 정의 및 합성 콘텐츠로 인한 위험 감소를 위한 내용이 포함되어 있다.

4) 딥페이크 신고법안(Deepfake Report Act of 2019)¹²⁷⁾

국토안보부의 과학기술국이 디지털 콘텐츠 위조 기술의 상태에 대해 지정된 간격으로 보고하도록 요구하고 있다. 디지털 콘텐츠 위조는 인공지능 및 머신러닝 기술을 포함한 새로운 기술을 사용하여 오도하려는 의도로 오디오, 시각 또는 텍스트 콘텐츠를 제작하거나 조작하는 것을 말한다.

5) 디피언스법(DEFIANCE Act of 2024)¹²⁸⁾

동의 없이 이루어진 사적인 디지털 위조 행위로 피해를 입은 개인의 권리 구제 방안 등을 개선하기 위해 발의된 법안이다.

126) 미국 백악관 보도자료 <<https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>> 최종 방문일 2024. 9. 11.>

127) 미국 의회 홈페이지 <<https://www.congress.gov/bill/116th-congress/senate-bill/2065#:~:text=This%20bill%20requires%20the%20Science,of%20digital%20content%20forgery%20technology>> 최종 방문일 2024. 9. 11.>

128) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/senate-bill/3696/text>> 최종 방문일 2024. 9. 11.>

6) 소비자를 기만적인 AI로부터 보호하는 법안(Protecting Consumers from Deceptive AI Act)¹²⁹⁾

국가표준기술원에서 생성형 인공지능에 의해 생성된 콘텐츠의 식별과 관련된 기술 표준 및 지침의 개발을 촉진하고 정보를 제공하기 위한 태스크포스를 설립하며, 생성형 인공지능에 의해 생성되거나 실질적으로 수정된 오디오 또는 시각적 콘텐츠에 대해서는 콘텐츠를 생성한 인공지능의 출처를 공개하는 내용을 포함하도록 요구하고 있다.

7) AI Labeling Act of 2023¹³⁰⁾

생성형 인공지능 시스템, 개발자 및 이용자로 구분하고 생성형 인공지능 시스템에 대해서 텍스트, 이미지, 동영상, 멀티미디어로 구분하여 표시와 관련된 일정한 의무를 부과하며, 개발자 및 이용자에 대해서 후속 이용에서도 표시가 이뤄지도록 하는 의무를 부과하고, 표시의무 불이행에 대한 집행에 대하여 규정하고 있다.

8) AI Disclosure Act of 2023¹³¹⁾

AI 생성물임을 표시하는 것과 표시의무 위반에 대하여 연방거래위원회(FTC)가 집행하는 내용을 규정하고 있다.

129) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/house-bill/7766/text>> 최종 방문일 2024. 9. 11.>

130) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/senate-bill/2691/text>> 최종 방문일 2024. 9. 11.>

131) 미국 의회 홈페이지 <<https://www.congress.gov/bill/118th-congress/house-bill/3831/text>> 최종 방문일 2024. 9. 11.>

(나) 주 차원

1) 캘리포니아주

가) AB 2839

캘리포니아 의원들은 2024년 8월 30일에 투표 절차, 선거 장비, 후보자에 대한 오해의 소지가 있는 정보를 묘사하는 특정 AI 생성 또는 조작된 정치적 커뮤니케이션의 배포를 금지하도록 민사소송법 제35조를 개정하고, 선거법에 제20012조를 추가하는 법안을 최종 승인하였다.¹³²⁾

나) AB 602¹³³⁾

AB 602에 따르면 ‘묘사된 개인이 동의하지 않았다는 사실을 알고 있거나 합리적으로 알았어야 할 경우 성적으로 노골적인 자료를 만들고 의도적으로 공개하는 경우’ 및 ‘묘사된 인물이 동의하지 않았다는 사실을 알고 있으면서도, 본인이 만들지 않은 성적 노골적인 자료를 의도적으로 공개하는 경우’ 개인 소송권을 제공하도록 하였다.

캘리포니아 민법에서 ‘묘사된 개인(Depicted individual)’이란 ‘디지털화의 결과로 실제로 수행하지 않은 공연을 하거나 변경된 묘사로 공연하는 것처럼 보이는 개인’을 말한다(Cal. Civ. Code § 1708.86 1708.86(a)(4)). 이에 따라 딥페이크 영상을 공개하는 행위 등이 ‘악의로 저질러진 경우’ 이익 환수가 가능하며, 경제적 및 비경제적 손해에 대하여 최대 15만 달러의 법정 손해배상을 받을 수 있다(Cal. Civ. Code § 1708.86(e)).

132) 미국 캘리포니아 의회 홈페이지 <https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240AB2839> 최종 방문일 2024. 9. 11.>

133) 미국 캘리포니아 의회 홈페이지 <https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=201920200AB602> 최종 방문일 2024. 9. 11.>

2) 텍사스주

가) TX SB 751¹³⁴⁾

딥페이크 영상을 제작하고, 딥페이크 영상을 공개하거나 선거일 30일 이내에 배포할 경우 처벌하도록 선거법을 개정하였다.

나) TX SB 1361¹³⁵⁾

어떤 사람이 묘사된 것처럼 보이는 사람의 효과적인 동의 없이 그 사람의 사적인 부위가 노출되거나 성적 행위에 참여하는 것처럼 보이는 딥페이크 영상을 전자적 수단을 통해 의도적으로 제작하거나 배포할 경우 범죄를 저지른 것으로 간주하도록 형법을 개정하였다.

3) 버지니아주

강요, 괴롭힘 또는 협박의 의도를 가지고 나체 상태 등의 타인을 묘사하는 비디오 그래픽 또는 정지 이미지를 어떠한 방법으로든 제작하여 악의적으로 유포 또는 판매하는 경우로서 그러한 권한 등이 없다는 것을 알거나 알 수 있는 경우에는 1급 경범죄에 해당하도록 하였다. 이 경우 ‘타인’은 실존하는 사람을 묘사할 목적으로 비디오 그래픽 등을 제작, 각색 또는 수정(creating, adapting or modifying)하는 데 그의 이미지가 사용된 사람 및 얼굴·유사성 등으로 실제 개인으로 인지할 수 있는 실존 인물을 포함한다. 이는 보복성 음란물에 관한 주법의 적용범위가 딥페이크에 의한 성적 콘텐츠까지 확장된 것으로 평가받는다(Code of Virginia § 18.2-386.2. Unlawful dissemination or sale of images of another; penalty).¹³⁶⁾

134) 미국 텍사스 주 홈페이지 <<https://capitol.texas.gov/tlodocs/86R/billtext/html/SB00751F.htm>> 최종 방문일 2024. 9. 11.>

135) 미국 텍사스 주 홈페이지 <<https://capitol.texas.gov/BillLookup/History.aspx?LegSession=88R&Bill=SB1361>> 최종 방문일 2024. 9. 11.>

4) 플로리다주

SB 1798은 성행위에 참여하는 식별 가능한 미성년자를 묘사하기 위해 전자적, 기계적 또는 기타 수단을 통해 생성, 변경, 각색 또는 수정된 이미지를 범죄로 규정하고 있다.¹³⁷⁾

5) 루이지애나주

Louisiana Act 457은 미성년자의 성적 행위에 관련된 딥페이크를 범죄로 규정하고 있다.¹³⁸⁾

6) 사우스다코타주

SB 79는 아동 포르노 소지, 배포 및 제작과 관련된 법률을 개정하여 컴퓨터 생성 아동 포르노를 포함하도록 하였다. 컴퓨터 생성 아동 포르노는 미성년자가 금지된 성행위에 참여하는 모습을 묘사하기 위해 제작·각색 또는 수정된 실제 미성년자의 시각적 묘사, 미성년자가 금지된 성행위에 참여하는 모습을 묘사하기 위해 제작·각색 또는 수정된 실제 성인, 또는 특정 데이터 입력을 처리 및 해석하여 시각적 묘사를 생성할 수 있는 AI 또는 기타 컴퓨터 기술을 사용하여 생성된 실제 미성년자와 구별할 수 없는 개인을 포함하고 있다.¹³⁹⁾

136) 미국 버지니아 주 홈페이지 <<https://law.lis.virginia.gov/vacode/18.2-386.2/>> 최종 방문일 2024. 9. 11.>

137) 미국 플로리다 주 상원 홈페이지 <<https://www.flsenate.gov/Committees/BillSummaries/2022/html/2658>> 최종 방문일 2024. 9. 11.>

138) 미국 루이지애나 주 홈페이지 <<https://legis.la.gov/legis/ViewDocument.aspx?d=1333325>> 최종 방문일 2024. 9. 11.>

139) 미국 사우스다코다 주 홈페이지 <<https://sdlegislature.gov/Session/Bill/24991/264651>> 최종 방문일 2024. 9. 11.>

7) 뉴멕시코주

HB 182호 법안은 뉴멕시코주의 선거운동 보고법의 조항을 개정하여 실질적으로 기만적인 미디어를 포함하는 광고에 대한 면책 조항 요건을 추가하고, 실질적으로 기만적인 미디어를 배포하거나 배포하기로 타인과 계약을 맺는 행위를 범죄로 규정하고 있다.¹⁴⁰⁾

8) 인디애나주

HB 1133은 조작된 미디어를 포함하는 특정 선거 캠페인 커뮤니케이션에 면책 조항을 포함하도록 요구한다. 또한 필요한 면책 조항이 포함되지 않은 조작된 미디어에 묘사된 후보가 지정된 사람에 대해 민사 소송을 제기할 수 있도록 허용하고 있다.¹⁴¹⁾

9) 워싱턴주

HB 1999은 성적으로 노골적인 딥페이크 피해자에게 민사 및 형사적 법적 구제책을 마련하고 있다.¹⁴²⁾

10) 테네시주

유사성, 음성 및 이미지 보안 보장법(ELVIS, Ensuring Likeness Voice and Image Security Act)은 개인의 이름, 사진, 음성 또는 초상을 보호하기 위해 주의 1984년 개인 권리 보호법을 업데이트하고 대체하였다. 이 법은 사람의 사진, 음성 또는 초상을 허가 없이 제작 및 배포하는 것과 관련된 활동에 대한

140) 미국 뉴멕시코 주 홈페이지 <<https://www.nmlegis.gov/Sessions/24%20Regular/final/HB0182.pdf>> 최종 방문일 2024. 9. 11.>

141) 미국 인디애나 주 홈페이지 <<https://legiscan.com/IN/bill/HB1133/2024>> 최종 방문일 2024. 9. 11.>

142) 미국 워싱턴 주 홈페이지 <<https://legiscan.com/WA/drafts/HB1999/2023>> 최종 방문일 2024. 9. 11.>

민사소송에서의 책임을 규정하고 있다. 사람의 사진, 음성 또는 초상을 허가 없이 사용하는 것을 주된 목적으로 기술을 배포, 전송 또는 기타 방법으로 제공하는 사람에 대한 책임을 포함한다.¹⁴³⁾

11) 오리건주

SB 1571은 선거 캠페인 커뮤니케이션에서 합성 미디어 사용에 대한 공개를 요구하도록 하고 있다.¹⁴⁴⁾

12) 미시시피주

SB 2577은 디지털화의 불법적 배포에 대한 형사 처벌을 신설하였다. 디지털화란 묘사된 사람이 아닌 사람의 이미지나 오디오를 사용하거나 딥페이크라고 일반적으로 불리는 컴퓨터 생성 이미지나 오디오를 활용하여 이미지나 오디오를 현실적인 방식으로 변경하는 것으로 정의하고 있다. 여기에는 소프트웨어, 머신러닝 AI 또는 기타 컴퓨터 생성 또는 기술적 수단을 사용하여 이미지나 오디오를 만드는 것도 해당된다.¹⁴⁵⁾

(다) AI 기업들의 AI 안전 서약¹⁴⁶⁾

오픈AI, 구글, 마이크로소프트 등 미국의 AI 기업들은 AI를 책임있고 안전하게 사용할 것을 약속하는 8개 조항을 발표하였다.

143) 미국 테네시 주 홈페이지 <<https://aboutblaw.com/bdpV>> 최종 방문일 2024. 9. 11.>

144) 미국 오리건 주 홈페이지 <<https://legiscan.com/OR/drafts/SB1571/2024>> 최종 방문일 2024. 9. 11.>

145) 미국 미시시피 주 홈페이지 <<https://legiscan.com/MS/bill/SB2577/2024>> 최종 방문일 2024. 9. 11.>

146) 미국 백악관 홈페이지 <<https://www.whitehouse.gov/briefing-room/statements-releases/2023/09/12/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-eight-additional-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/>> 최종방문일 2024. 9. 10.>

서약서에는 ① 가짜 뉴스와 딥페이크를 방지하기 위해 텍스트나 이미지, 사운드, 동영상 등 AI 생성 파일에 워터마크를 추가, ② AI 오남용 모니터링 외부 팀을 구성, ③ 정부와 기업에 안전 정보를 공유, ④ 사이버 보안에 투자, ⑤ 보안 취약점, AI의 환각 및 편향성 문제를 보고, ⑥ 최첨단 AI 모델을 활용한 기후변화와 신약 개발 등 사회 문제를 우선 해결하는 안 등이 포함되었다.

(3) EU

(가) 인공지능법(The EU Artificial Intelligence Act)

2024년 5월에 EU 이사회(European Council)는 세계 최초로 AI를 포괄적으로 규제하는 「The EU Artificial Intelligence Act」를 최종 승인하였다. 인공지능법은 AI를 잠재적 위험에 따라 금지된 AI(unacceptable risk), 고위험 AI(high risk), 제한적 위험 AI(limited risk), 최소한의 위험 AI(minimal risk)의 네 가지 범주로 분류하였으며, 딥페이크는 제한된 위험 AI(limited risk)에 해당한다.

인공지능법은 딥페이크와 관련하여, AI 시스템 공급자(provider)와 사용자(deployer)에게 투명성 의무(Transparency Obligation)를 부과하고 있다.

인공지능법은 합성된 오디오·이미지·비디오 또는 텍스트 콘텐츠를 생성하는 AI 시스템의 제공자(provider)에게 해당 결과물이 기계가 판독할 수 있는 형식으로 표시하고, 인위적으로 생성 또는 조작된 것으로 감지할 수 있어야 한다는 표시 및 감지 의무를 부여한다(인공지능법 제50조 제2항).

그리고 인공지능법은 딥페이크에 해당하는 이미지·오디오 또는 비디오 콘텐츠를 생성하거나 조작하는 AI 시스템의 배포자(deployer)에게 결과물이 인위적으로 생성 또는 조작되었음을 공개할 의무를 부여한다(인공지능법 제50조 제4항).

(나) 디지털서비스법(Digital Service Act)

유럽연합(EU)에서는 유럽연합 기본권 헌장(Charter of Fundamental Rights of the European Union)의 기본권을 효과적으로 보호하기 위하여 「디지털서비스법」을 제정하였다.

「디지털서비스법」은 단순 전달, 캐싱, 호스팅으로 분류되는 중개서비스에 적용되며, 사업자별로 차등적인 주의의무 등을 규정하고 있는데, 초대형 온라인 플랫폼(VLOP)과 초대형 온라인 검색엔진(VLOSE)에 딥페이크를 사용된 영상이 업로드된 경우 이를 표시할 의무를 부여한다.

초대형 온라인 플랫폼(VLOP)과 초대형 온라인 검색엔진(VLOSE)의 서비스 제공자는 실존인물·물체·장소 등 진실인 것처럼 보이도록 생성되었거나 조작된 이미지·오디오 또는 비디오로 구성된 정보 항목이 온라인 플랫폼에 노출된 경우 눈에 잘 띄는 표시로 구별할 수 있어야 하고, 서비스 이용자도 쉽게 표시할 수 있는 기능을 제공해야 한다(디지털서비스법 제35조 제1항).

(4) 일본

일본은 딥페이크 성착취물 관련 법 조항이 마련되어 있지 않다. 대신 형법상 명예훼손, 저작권법 위반, 외설물 배포, 판매 목적 외설물 소지·보관죄 등의 혐의로 처벌하도록 하고 있다.

(5) 중국

(가) 정보 서비스의 심층합성 관리 규정

중국 국가인터넷정보판공실(国家互联网信息办公室)은 딥페이크 규제를 통해 안전한 인터넷 환경을 조성하기 위해 인터넷 2023년 1월 10일부터 「정보

서비스의 심층합성 관리 규정」(互联网信息服务深度合成管理规定)을 시행하고 있다.¹⁴⁷⁾

본 규정은 심층합성기술(딥페이크 기술) 이용은 허용하되 허위정보의 규제와 및 심층합성기술의 표시 등 서비스 제공자 플랫폼 인터넷 서비스 기술 지원자 등이 준수해야 할 사항을 규정하고 있다.¹⁴⁸⁾

(나) 생성형 인공지능 서비스관리를 위한 임시조치

중국은 2023년 네트워크정보국 등 7개 부서가 공동으로 ‘생성형 인공지능 서비스관리를 위한 임시조치’를 발표하였다. 생성형 인공지능 서비스제공자에게 ‘인터넷 정보서비스 심층 종합관리에 관한 규정’에 따라 사진, 비디오 및 기타 생성된 콘텐츠에 표시를 하도록 규정하고 있다.

(6) 영국

2024년 4월에 온라인 안전법(Online Safety Act)을 개정하여 성적으로 노골적인 딥페이크(sexually explicit deepfake) 영상을 제작할 경우, 공유할 의도가 없고 단지 피해자에게 불안감·굴욕감 또는 고통을 안겨주고 싶을 뿐이더라도 형사처벌하도록 하였다.

147) 중화인민공화국 중앙인민정부, https://www.gov.cn/zhengce/zhengceku/2022-12/12/content_5731431.htm 보충적으로 비공식 영문 자료 참고, <https://www.chinalawtranslate.com/en/deep-synthesis/>

148) 채은선, 「디지털 법제 Brief」, 한국지능정보사회진흥원, 2023. 1. 31.

(7) 프랑스

(가) SREN법¹⁴⁹⁾

프랑스는 2024년 5월 21일에 SREN법을 채택하여 동의 없이 이루어진 딥 페이크(non-consensual deep fakes)를 금지하고 있다. SREN법은 프랑스 형법 제226-8조를 보완하여 알고리즘 처리를 통해 생성되고 사람의 이미지나 발언을 나타내는 시각적 또는 청각적 콘텐츠를 해당 콘텐츠의 동의 없이 공유하는 행위를 명시적으로 금지하고 있다. 단, 콘텐츠가 알고리즘에 의해 생성되었다는 것이 명백하거나 명시적으로 언급된 경우는 예외적으로 적용되지 않는다.

한편, 프랑스는 정보조작대처법을 통해 선거전 3개월 동안 법원은 온라인 플랫폼에 허위정보를 게시하지 못하도록 명령할 수 있다.

(나) 지적재산권법 개정

저작인격권으로서 공표권을 규정하고 있는 조항에 저작물이 AI시스템에 의하여 생성된 경우 AI생성물이라는 것과 그러한 생성에 이르도록 한 주체의 성명을 표시하도록 하였다.

(8) G20 히로시마 정상회의 선언

히로시마 프로세스 국제 행동강령¹⁵⁰⁾은 AI 생성 콘텐츠 표시와 관련된 내용을 포함하고 있다. 이 행동강령은 워터마크 등에 의하여 이용자들이 AI 생성 콘텐츠를 확인할 수 있도록 하는 메커니즘 개발하고, 이용자가 AI 시스

149) 프랑스 의회 홈페이지 <<https://www.legifrance.gouv.fr/jorf/id/JORFTEXT000049563368>> 최종 방문일 2024. 9. 11.>

150) <https://digital-strategy.ec.europa.eu/en/library/hiroshima-process-international-code-conduct-advanced-ai-systems> 최종 방문일 2024. 9. 11.

템과 상호작용하고 있다는 것을 알 수 있도록 하기 위하여 표시하거나 부인하는 것과 같은 메커니즘을 실행하도록 하고 있다.

다. 입법 방향

(1) 허위 영상물 규제와 인공지능 생성물 표시 의무

허위 영상물 규제와 인공지능 생성물 표시 의무는 구별할 필요성이 있다. 인공지능 생성물 표시와 딥페이크 기술을 이용한 가짜뉴스, 이미지, 영상물에 대한 규제는 모두 인공지능을 활용하는 것과 관계되지만, 양자는 그 성격을 달리한다. 전자는 인공지능 생성물을 인간의 창작물로 오인시키는 것에 문제가 있지만, 후자는 이를 넘어서 생성형 인공지능을 악용하여 허위의 사실이나 존재를 진짜인 것처럼 오인시킴으로써 개인, 사회, 국가적으로 큰 피해를 발생시키는 문제점이 있다.

인공지능 생성물 표시는 인간 창작물과 인공지능 생성물을 구별하는 것에 초점을 맞추므로, 인공지능 생성물임을 표시하기만 한다면 공표하는 것 자체는 문제되지 않으며, 인공지능 생성물임을 표시하지 않는 것에 대한 규제만 하면 된다. 그러나 허위 영상물은 허위의 이미지나 영상물에 의하여 해당 인물이나 사실이 진정한 것처럼 속여서 유포하는 것을 목적으로 하므로, 허위의 이미지나 영상물을 공표하는 것 자체를 규제할 필요성이 있다는 점에서 차이가 존재한다.

(2) 허위 영상물 금지

현행 국내 법률상 딥페이크를 통해 제작된 허위 영상물에 대한 소지·시청 행위에 대하여 별도로 처벌하는 규정이 존재하고 있지 않다. 이 외에도 반

포할 목적이 없을 경우에는 허위 영상물의 제작행위에 대한 처벌을 할 수 없어 처벌 공백이 있다.

최근 학생들 사이에서도 딥페이크를 통한 성착취물 제작 또는 허위 영상물의 제작이 급증하고 있으며, AI의 특성상 한번 반포된 이후에는 회수가 쉽지 않다는 점을 통해 보다 강력한 제재가 필요할 것으로 보인다.

영국은 딥페이크를 통해 영상을 제작할 경우, 공유할 의도가 없고 단지 피해자에게 불안감, 굴욕감 또는 고통을 안겨주고 싶을 뿐이더라도 형사처벌을 하고 있으며, 이와 같이 반포할 목적이 없는 딥페이크 영상 제작에 대해서도 처벌하도록 하는 방안을 고려할 수 있다. 최근 정부는 딥페이크 등 허위 영상물 소지·구입·시청 행위를 처벌하는 규정을 신설하고, 제작·유통에 대한 처벌 기준을 상향하는 법률 개정을 추진하기로 하며, 국무조정실에서 ‘딥페이크 성범죄 대응을 위한 첫 번째 범정부 대책 회의’를 열고 입법 대책을 마련하고 있어서 그 귀추가 주목된다.

다만, 현재는 딥페이크를 통한 성착취물 제작 및 반포에 초점이 맞춰져 있으나, 다양하게 발생하는 허위 영상물 관련 문제점들을 효과적으로 해결하기 위하여 인공지능 기본법 또는 허위 영상물 관련 특별법에서 이를 예방·방지하는 시책을 추진하도록 하는 내용을 담는 방안도 고려할 수 있다.

(3) 인공지능 생성물 표시 의무

인공지능 기술 발전으로 가상 정보의 제작 비용은 낮아지고 생산량은 급격하게 증가하고 있는데, 이 과정에서 가상 정보의 구분 곤란성과 오인 가능성을 악용한 범죄와 사회적 혼란이 끊이지 않고 있다.

이에 따라 인공지능으로 만든 가상 정보에 대하여 표시를 하도록 한다면 이를 통해 사람들은 인공지능으로 만든 가상 정보와 실제 사실을 쉽고 효율적으로 구분할 수 있게 될 것이다.¹⁵¹⁾

EU, 미국 등 주요국에서도 입법이나 정책을 통해 인공지능 생성물 표시 의무화를 추진하고 있다는 점에서 우리나라 역시 이러한 입법 방향을 긍정적으로 고려할 필요가 있다.

더욱이, 인공지능 생성물 표시는 허위조작정보 금지뿐만 아니라 저작권 침해 예방에도 효과적일 것으로 생각된다. 이에 따라 특정 분야의 법률이 아닌 인공지능 전반을 규율하는 기본법에 규정하는 방안을 고려할 수 있다.

그러나 가상 정보 표시에 대한 기술적 회피·조작 가능성, 실제 사실에 대한 거짓 표시 가능성 등이 있으므로, 이를 해결하지 못한다면 표시 제도의 유용성·신뢰성이 저하될 우려가 있다.¹⁵²⁾

인공지능 생성물인지 여부를 판별하기 위한 기술을 개발하는 시도가 지속적으로 이루어지고 있지만, 현재 기술 수준은 완벽하게 신뢰할 수 있는 정도로 딥페이크 영상을 탐지하지 못하는 것으로 알려져 있다.¹⁵³⁾

이러한 점을 고려할 때, 현 시점에서는 적용 콘텐츠나 수범자를 좁히고 미준수 시 제재는 최소한으로 하되, 인공지능 생성물 탐지 기술이나 표시 의무 우회를 방지하는 기술 등의 발전 상황을 지켜보면서 규제를 넓혀가는 방안을 고려할 필요가 있다.

151) 정준화, 「정보통신망 이용촉진 및 정보보호 등에 관한 법률 일부개정법률안(의안번호 2126048) 입법영향분석- 인공지능으로 만든 가상 정보임을 표시하도록 하는 의무」, 국회입법조사처, NARS 입법영향분석 기획보고서 제7호, 2024. 8. 23. 11쪽

152) 정준화, 위 보고서, 14쪽

153) 김병필, 「생성형 AI 산출물 표시 의무화 연구」, 문화체육관광부·한국저작권위원회, AI-저작권 제도개선 워킹그룹 발제 자료집, 2023. 10. 49쪽

5. 인공지능 윤리 및 규제체계 입법 방향

가. 총론

인공지능 윤리 규제체계와 관련하여, EU와 같은 수평적·하향식 규제체계는 장점을 갖고 있다. 사람들은 인공지능 기술에 대하여 여전히 불확실성·불투명성으로 인한 불안감을 갖고 있으며, 최근의 생성형 인공지능의 환각 현상으로 인해 부정확성의 우려도 있어 인공지능 기술의 사회적 수용성을 높이기 위해 신뢰도 향상이 시급하고 중요한 상황으로서, 포괄적이고 전 분야에 걸쳐 촘촘하게 적용되는 규제체계는 인공지능 기술에 대한 신뢰를 높일 수 있다.

특히 생성형 인공지능 서비스의 기반이 되는 기반 모델(Foundation Model)은 인공지능 기술에 범용성을 부여하고 거의 전 분야에 적용되고 있고 그 성능이 비약적으로 발전하고 있어, 이에 대한 개발·생산·배포·유통·활용 전반에 걸쳐 일반적 규제요구사항을 담은 규제체계가 요구되고 있다.

인공지능 기술·시스템에 공통적으로 적용되는 의무를 도출할 수 있다. 예를 들어 인적 감독 의무, 학습데이터 요약 공개, 출시 후 모니터링, 로그 보관 등의 의무는 일반적으로 인공지능 기술·시스템을 더 잘 이해하고 신뢰를 높이는 데 도움이 된다.¹⁵⁴⁾ 이처럼 인공지능에 대한 수평적·하향식 규제체계의 여러 장점은 EU에서 인공지능법이 통과되는데 작용하고, 미국 일부 주에도 영향을 주었을 것으로 보인다.

반면, 인공지능 기술은 지속적으로 발전하고 있고 AI 시스템의 상당수는 아직 초기 단계이거나 제대로 사업화되지 않고 있어, 강력하고 경직된 규제를 적용하는 것이 적절한지에 대한 논란이 있다.

위험 기반 접근법을 취한 경우에도 어떤 제품이나 서비스에 적용된 인공

154) EU 홈페이지 <<https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>> 최종 방문일 2024. 9. 2.>

지능 기술이 금지되는 인공지능에 해당하는지, 고위험 인공지능에 해당하는지 불분명한 경우도 많을 것으로 예상된다.

나아가 실증적 근거가 부족하고, 되돌릴 수 없는 손해가 분명히 예견되는 상황도 아니어서 사전주의적 접근 방식이 적절한지 의문이 들 수 있고, 현실에 아직 존재한 적 없는 기술에 대한 규제는 연구개발 과정을 위축시키고 기업들의 경쟁력을 손상시킬 우려가 있다는 견해도 있다.¹⁵⁵⁾

우리나라는 EU와는 달리 자체 인터넷 포털과 모바일 인스턴트 메시지 서비스를 제공하고 있고, 해외 정보기술 시장에도 진입하고 있어 EU와는 입장이 다른 측면도 고려할 필요가 있다.

이런 점을 고려할 때, 규제의 대상이 되는 분야를 한정적으로 열거하고, 사업자에게 과도하지 않은 수준의 설명의무 등 최소한의 의무만을 부여하며, 위반 시 강한 행정벌을 부과하기 보다는 당국의 조사 및 공표 정도의 낮은 수준의 제재를 고려할 수 있다. 이후 인공지능 기술의 발전상을 따라가며, 인공지능으로 인하여 피해를 입은 구체적 사례를 분석하고 이를 축적하여 적절한 규제 범위와 제재 수준을 정해가는 것이 바람직하다고 판단된다.

나. 위험 기반 규제 방식 도입 여부

EU 인공지능법의 경우 인공지능을 위험성에 따라 차등 규율하는 위험 기반 규제 방식을 채택하였고, 금지된 인공지능은 원천적으로 개발 및 이용이 금지되며, 고위험 인공지능의 경우 개발자 및 배포자에게 상당한 수준의 의무를 부여하였다.

위험 기반 규율은 위험평가를 통해 문제의 위험성 및 발생가능성을 인지하고 대응순위를 설정함으로써 제재의 수위 및 역량의 분배를 통해 사안별 규

155) 고훈수·임용·박상철, 「유럽연합 인공지능법안의 개요 및 대응방안」, 서울대 인공지능정책 이니셔티브 「DAIG」 제2호, 2021

율에 이르는 규제방식을 말한다.¹⁵⁶⁾

인공지능은 거의 모든 분야에 이용될 수 있는데 사람의 권리·의무에 미치는 영향을 천차만별이므로, 인공지능 기술 및 제품·서비스 전반을 규율하는 기본법이 일률적으로 동일한 수준의 규제를 적용하는 것은 타당하지 않으며 그 위험성에 따라 차등 규율하는 위험 기반 규제 방식이 적절하다고 판단된다.

다만, 인공지능 기술을 이용한 사업이 아직 초기 단계인 경우가 대부분이며 그 위험성이 실증적으로 분석된 결과도 찾아보기 어려우므로, 금지된 인공지능을 규정하고 원칙적으로 개발과 이용을 차단하는 규제는 아직 시기상조라고 보인다.

현 단계에서는 공공분야를 중심으로 국민의 권리·의무에 상당한 영향을 줄 수 있는 분야를 엄선하여 ‘고위험 인공지능’ 분야로 정하고, 낮은 수준의 설명의무 등 개발자·배포자에게 큰 부담이 되지 않는 규제를 도입하는 것이 적절하다고 본다.

다. 생성형(범용) 인공지능 규제 여부

생성형 인공지능은 대화, 이야기, 이미지, 동영상, 음악 등 새로운 콘텐츠와 아이디어를 생성할 수 있는 인공지능인데, 허위조작정보(딥페이크), 여론조작, 환각, 개인정보·저작권 침해 등 위험 요소가 점점 구체화되고 있다. 더욱이 생성형 인공지능은 범용성을 갖고 다양한 제품·서비스의 기반 요소로 활용될 수 있어 그 파급력이 상당하다.

이에 따라 EU 인공지능법은 범용 인공지능에 대한 별도 규제를 도입하였으며, 범용 인공지능 모델 제공자에게 기술문서 작성 및 업데이트, 저작권법

156) 남구현, 「인공지능 규율에 있어 위험기반 접근 방식의 적합성 고찰」, 법학논총 제36권 제1호, 2023

준수, 학습에 사용된 콘텐츠에 관하여 충분히 상세한 요약서 작성·공개 등의 의무를 부과하였다.

이처럼 생성형 인공지능 모델은 개별 인공지능 제품·서비스의 기반 요소로서 범용성과 고성능을 특징으로 하므로 별도의 규제를 통해 위험을 저감할 필요가 있다.

라. 연산능력에 따른 규제 여부

EU 인공지능법은 범용 AI 모델의 학습 사용 누적 컴퓨팅이 10^{25} FLOPs¹⁵⁷⁾ 보다 크면 ‘고영향 성능’으로 추정하고 ‘구조적 위험(systemic risk)’을 갖고 있다고 보아 구조적 위험 파악 등 공급자의 의무를 가중하고 있다. 미국 AI 행정명령은 10^{26} FLOPs 이상의 성능, 또는 10^{23} FLOPs 이상의 성능을 보이면서 주로 생물학적 서열 데이터를 사용하는 경우 정부 보고 대상으로 규정한다. 캘리포니아주 SB-1047 법안은 클라우드 컴퓨팅 평균 시장가 기준으로 1억달러 이상의 개발비용이 들거나 10^{26} FLOPs 이상의 성능으로 학습된 AI 모델을 주요 규제 대상으로 삼고 있다.

이와 같이 높은 수준의 연산 능력으로 학습된 AI모델은 높은 성능을 보일 수 있고 그에 따른 위험도 증가하기에 추가 의무를 부담하는 것이 적절하다고 볼 수 있다.

그러나 AI 성능이 하루가 다르게 향상되고 있어 성능 기준을 경직된 법률 형식으로 규정하는 것이 효과적인지 의문이 있으며, FLOPs 성능 외에 양질의 훈련데이터, 미세조정, RAG(Retrieval-Augmented Generation, 검색 증강 생성)를 통해 성능이 향상될 수 있어, FLOPs 성능 기준을 넘어 포괄적인 접근법이 필요하다는 견해도 있다.¹⁵⁸⁾

157) 1초에 수행할 수 있는 부동(浮動) 소수점 연산의 횟수

이런 점을 고려할 때, 법률에서 연산능력에 따른 기준을 설정하기보다는 ‘고성능’ 등과 같이 불확정적인 개념을 도입하고 필요시 대통령령 등으로 기준으로 설정하고 의무를 부과하는 방안을 고려할 수 있다.

마. 제공자와 배포자 등 구분 여부

EU 인공지능법은 제공자(provider)와 배포자(deployer) 등을 구분하고 각기 다른 의무를 부과하고 있다. EU 인공지능법상 제공자는 AI 시스템 또는 범용 AI 모델을 개발한 주체, 배포자는 AI 시스템을 이용하는 주체로 구분된다.

인공지능 기술은 현재 인공지능 핵심 기술 또는 기반 모델을 제공하는 주체와 이를 이용하여 사업을 추진하는 주체가 대체로 구분되는데, 제공자는 인공지능 시스템을 훈련하는데 사용되는 데이터의 특징을 설명할 수 있는 반면, 배포자는 통상 인공지능 시스템이 어떻게 사용되고 있는지, 그 사용이 원래 목적과 일치하는지, 감독이 잘 이루어지는지 등 실제 운용과 관련된 사항을 잘 이해하는 차이가 있다.

이러한 점을 고려하여 제공자·배포자 등 인공지능 생태계에서 활동하는 각 주체별로 그에 적합한 의무를 부과하는 방안을 고려할 수 있다.

바. 거버넌스

인공지능 법제의 주요 규제 대상이 될 고위험 인공지능 영역의 제품 및 서비스는 우리나라의 경우 여러 부처에 걸쳐 규제가 중첩 적용될 가능성이 높고, 어느 부처·기관의 소관인지 불분명한 경우들이 발생할 것으로 예상된다.

158) Ingrid Stevens, “Regulating AI: The Limits of FLOPs as a Metric”, 2024. 5. 1. <<https://medium.com/@ingridwickstevens/regulating-ai-the-limits-of-flops-as-a-metric-41e3b12d5d0c> 최종 방문일 2024. 9. 2.>

만약 어느 한 부처가 인공지능 법안이나 정책을 내놓는 경우 다른 부처들도 이를 가만히 지켜볼 수만은 없어 경쟁적으로 법안·정책을 내놓는 ‘규제 레이스’가 촉발될 우려도 있다. 이에 따라, 개별 규제들의 중첩적인 집행으로 인한 사회적 비용이 발생하지 않도록 인공지능 관련 부처들간의 조정이 활성화되어야 한다.¹⁵⁹⁾

규제 집행의 실효성 확보와 인공지능 리스크 책임 분배를 해결하기 위해 컨트롤타워 역할을 할 수 있는 독립적인 조직의 필요성이 제기되고 있으며, 인공지능 관련 정책과 감독을 전담하는 독립적인 기구는 업무수행의 독립성, 공정성, 전문성을 확보해야 하고 각 부처별, 인공지능 관련 기구들과의 원활한 협력관계를 구축해야 할 것이다.¹⁶⁰⁾

한편, EU 인공지능법은 필수 조직으로 과학·기술 분야 지원을 위해 인공지능 분야에서 전문지식을 갖추고 AI 시스템 또는 범용 AI 모델의 제공자로부터 독립된 전문가로 구성된 과학패널을 구성하도록 하였다. 과학패널은 EU 인공지능법의 시행 및 집행 지원에 관한 조언, AI 사무국에 범용 AI 모델의 위험성 경고, 범용 AI 모델 및 시스템의 성능을 평가하는 도구 및 방법의 개발 지원 등 정책적으로 중요한 역할을 할 것으로 전망된다. 우리나라 역시 정책을 자문하는 전문가 집단의 구성을 고려할 필요가 있으며, 단순히 행정조직을 자문하는데 그치는 것이 아니라 실질적인 영향력과 구속력을 가질 수 있는 체계를 고려할 필요가 있다.

이와 관련하여, 정부는 2024년 5월에 「국가인공지능위원회의 설치 및 운영에 관한 규정 제정령안」 입법예고를 하였으며, 2024년 8월 6일부터 제정되고 시행되었다.

159) 고학수·임용·박상철, 위 논문

160) 이효진, 「우리 ‘인공지능법 제정안’의 인공지능 규율 방향에 관한 소고」, 법학논문집 제47집 제2호(중앙대학교 법학연구원), 2023. 8.

구체적으로, 인공지능 산업의 진흥 및 인공지능 신뢰 기반 조성을 위한 주요 정책 등에 관한 사항을 효율적으로 심의·조정하기 위하여 대통령 소속으로 국가인공지능위원회를 설치하고 기획재정부 등 주요 부처, 디지털플랫폼정부위원회, 대통령실 과학기술보좌관 및 인공지능 관련 전문지식과 경험이 풍부한 사람으로 구성하도록 하였다.

위 위원회를 대통령 소속으로 설치함으로써 해당 위원회의 실질적 권한을 강화하고 부처별 조율을 주도할 수 있는 기반을 마련했다고 볼 수 있다. 다만 해당 규정은 한시적인 대통령령으로서 지속성에 의문이 있으므로, 법률 규정을 통해 그 지위와 권한을 안정적으로 규정할 필요가 있다. 과거 4차산업혁명위원회가 대통령령에 근거하여 설치되었고 인공지능을 비롯한 여러 정책을 추진하였지만 한시적인 기구로서 폐지된 사례를 반면교사로 삼을 필요가 있다. 향후 해당 위원회는 각 부처에 자료 제출을 요구하거나 법안을 발의할 수 있도록 권한을 강화하는 방안을 검토할 수 있다.

또한 위 위원회에 위원으로 각계 전문가들이 대거 위촉될 수 있도록 하고 분야별 분과위원회·특별위원회 및 자문단의 설치하도록 하여, 인공지능 기술 및 정책·법제 전문가들이 참여할 기회를 만들었다고 평가된다. 민간전문가들의 목소리가 실질적으로 반영될 수 있도록, 민간위원이나 분과위원회 등의 의견에 반하는 결정을 하는 경우 이에 대한 근거를 남기도록 하는 방안 등을 고려할 수 있다.

III. 인공지능 저작권 규제체계

1. 인공지능과 저작권

가. 저작물 등의 정의

‘저작물’의 정의는 ‘인간의 사상 또는 감정’을 표현한 창작물로서, 저작권은 인간의 사상 또는 감정이 표현된 저작물에 부여되는 권리이다. ‘표현’을 보호하므로, 어떤 화가의 화풍 등 스타일은 저작권으로 보호되지 않는다. 최소한의 창작성만 있으면 보호가 된다.

이에 따라, 인터넷상 공개된 글, 이미지, 영상 등은 상당수 저작물로 인정될 가능성이 높다. 통상 인공지능은 웹 크롤링 등을 통해 인터넷상 데이터를 수집하여 학습하므로 저작물을 저작권자 동의 없이 이용하는 경우가 다수 발생하게 된다.

한편, ‘저작자’의 정의는 ‘저작물을 창작한 자’로서 통상 사람을 말하는 것으로 이해된다. 따라서 원숭이 등 동물은 저작자가 될 수 없으며, 인공지능 또한 창작의 도구로 활용될 뿐 그 자체를 저작자로 인정하기는 어렵다.

나. 우리나라 정부 및 공공기관 정책 현황

문화체육관광부와 한국저작권위원회는 2023년 12월 인공지능 기술 상용화로 인한 시장 혼란을 최소화하기 위해, 생성형 인공지능 사용 시 유의사항, 저작권 등록 등 주요 사항을 정리한 ‘생성형 인공지능(AI) 저작권 안내서’를 발표하였다.¹⁶¹⁾ 이후 문화체육관광부와 한국저작권위원회는 ‘2024 인공지능-저작

161) 문화체육관광부 보도자료, “인공지능(AI)-저작권 안내서 발표로 시장의 불확실성 해소하고, 안무·건축 등 ‘저작권 사각지대’ 없앤다”, 2023. 12. 26.

권 제도개선 워킹그룹'을 발족하고 ▲인공지능 학습에 저작물을 어떤 방식으로 이용할 수 있는지, ▲학습데이터의 공개 여부, ▲인공지능 산출물의 법적 성격과 저작권 침해 여부 등 인공지능 학습과 산출 단계에서의 쟁점에 대한 정책 방안을 모색하고 있다.¹⁶²⁾

정부출연연구소, 국책연구기관들도 인공지능 관련 저작권 정책을 다각도로 연구하고 있다. 정보통신정책연구원은 2023년 「생성형 AI와 저작권 현안」을 통해 생성형 인공지능의 저작권 침해와 공정이용 이슈를 다루었다.¹⁶³⁾ 한국법제연구원은 2024년 「생성형 AI를 둘러싼 최근 저작권 분쟁 동향과 시사점」을 통해 외국에서 진행되는 생성형 인공지능 관련 저작권 소송에 대하여 분석하였다.¹⁶⁴⁾

다. 인공지능 저작권 관련 분쟁 현황

최근 생성형 인공지능 도입이 증가함에 따라 법적 분쟁이 상당수 발생하고 있다. OpenAI, Stability AI 등 생성형 인공지능을 개발하여 서비스를 제공하는 기업들을 상대로 소송이 제기되고 있으며, 대부분 기반 모델의 학습 방법, 학습에 사용된 데이터, 이용자 프롬프트를 통한 결과물 등에서 저작권 침해가 성립한다는 이유로 소송이 진행되고 있다. 주목할 만한 소송을 정리하면 아래와 같다.

162) 문화체육관광부 보도자료, “인공지능과 저작권 쟁점별 구체적 정책 방안 마련한다”, 2024. 2. 16.

163) 김윤명, 「생성형 AI와 저작권 현안」, KISDI AI Outlook, 정보통신정책연구원, 2023. 6.

164) 김창화, 「생성형 AI를 둘러싼 최근 저작권 분쟁 동향과 시사점」, 규제법제리뷰 제24-1호, 2024. 2.

[표 11] 인공지능 저작권 관련 주요 소송

※ 표의 순서는 시간에 역순

일자·지역	소송 당사자	경과
2024. 8. 19. 미국 캘리포니아 북부지법	Andrea Bartz, et al. v. Anthropic PBC	• 안드레아 바츠, 찰스 그레이버, 커크 월러스 존슨은 앤트로픽의 AI 제품인 클로드가 수천 권의 책을 허가나 보상 없이 사용함으로써 저작권을 침해했다고 주장하며 앤트로픽을 상대로 소 제기 • 저작권 침해에 대한 집단 소송 승인, 금전적 손해배상, 금지명령 구제를 요청 • 저자들은 Anthropic이 불법 복제 콘텐츠를 기반으로 비즈니스를 구축하여 AI가 오리지널 저작물 시장과 경쟁하거나 시장을 희석시키는 콘텐츠를 생산할 수 있게 되었다고 주장
2024. 6. 27. 미국 뉴욕 남부지법	Center for Investigative Reporting, Inc. v. OpenAI, Inc., et al.	• 탐사보도센터(CIR)는 OpenAI와 Microsoft를 상대로 소 제기 • 자신의 콘텐츠에 대한 저작권 저보를 제거하고 무단으로 사용하여 ChatGPT 및 Copilot과 같은 AI 제품을 학습시키는 데 사용함으로써 저작권을 침해하고 DMCA를 위반했다고 주장
2024. 6. 24. 미국 메사추세츠 지방법원	UMG Recordings, Inc., et al. v. Suno, Inc., et al.,	• UMG Recordings(UMG)는 AI를 이용한 음악 생성 서비스를 제공하고 있는 Suno Inc. 등 및 Uncharted Labs, Inc. 등을 상대로 저작권 침해 소 제기
2024. 5. 16. 미국 뉴욕 남부지법	Lehrman, et al. v. LOVO, Inc.	• 성우 폴 레어만과 리네아 세이지가 자신의 목소리를 무단으로 사용했다며 텍스트 음성 변환 서비스 업체인 LOVO를 상대로 소 제기 • LOVO가 “모든 목소리를 복제할 수 있다”고 허위 광고하고 동의 없이 성우의 신원을 사용해 더빙 작품을 제작함으로써 각 주 개인정보 보호법 및 연방 랜햄법(Lanham Act)을 위반했다고 주장

일자·지역	소송 당사자	경과
2024. 4. 30. 미국 뉴욕 남부지법	Daily News, LP, et al. v. Microsoft Corp., et al.	• 시카고 트리뷴과 올랜도 센티널을 포함한 8개 뉴스 출판사가 Microsoft와 OpenAI가 수백만 건의 기사를 허가나 대가 없이 ChatGPT 및 Copilot과 같은 AI 제품을 학습하는 데 사용했다며 저작권 침해 및 상표권 희석 혐의로 소 제기
2024. 4. 26. 미국 캘리포니아 북부지법	Zhang et al. v. Google LLC and Alphabet Inc.	• 시각 예술가 그룹은 구글과 그 모회사인 알파벳이 구글의 인공지능 이미지 생성기인 Imagen 훈련에 자신의 저작권이 있는 작품을 무단으로 사용했다며 구글과 알파벳을 상대로 소 제기 • 원고들은 Imagen의 학습 과정에서 자신의 작품을 복제하여 AI 모델이 저작권을 침해하는 2차적저작물이 되었다고 주장
2024. 5. 8. 미국 캘리포니아 북부지법	Nazemian, et al. v. NVIDIA Corp.	• 작가인 아브디 나제미안, 브라이언 킨, 스투어트 오난은 NVIDIA를 상대로 저작권 침해 소 제기 • 원고측은 NVIDIA가 허가 없이 자신들의 책인 Books3 데이터 세트 일부를 NeMo Megatron 모델 훈련에 사용했다고 하면서 NVIDIA가 Books3가 포함된 더 파일 데이터 세트의 사용을 인정했다고 주장
2024. 2. 28. 미국 뉴욕 남부지법	Raw Story Media, et al. v. OpenAI, Inc., et al.	• 뉴스 매체인 Raw Story Media와 AlterNet Media는 ChatGPT가 저작자 표시나 저작권 정보 없이 자신들의 저작권이 있는 저널리즘을 재가공한다고 주장하며 OpenAI를 상대로 소 제기 • 원고 측은 OpenAI가 자신들의 저작물에 대한 모델을 훈련시키면서 저작권 관리 세부 정보를 고의로 삭제하여 ChatGPT가 적절한 승인 없이 저널리즘의 상당 부분을 모방하거나 복제하는 결과물을 생산하도록 유도했다고 주장

일자·지역	소송 당사자	경과
2024. 2. 28. 미국 뉴욕 남부지법	The Intercept Media, Inc. v. OpenAI, Inc.	• 인터셉트 미디어는 OpenAI, Microsoft 및 계열사가 저자 정보 없이 자사의 기사를 사용하여 ChatGPT를 훈련시킴으로써 DMCA를 위반했다고 주장하며 소 제기
2024. 1. 25. 미국 캘리포니아 중앙지법	Main Sequence, et al. v. Dudesy LLC, et al.	• 고인이 된 코미디언 조지 칼린의 유족은 Dudesy 등이 허가 없이 AI를 사용하여 칼린의 목소리와 외모를 모방하였다는 이유로 소 제기 • (경과) 2024. 4. 2. 양측은 합의에 이르렀고, Dudesy는 칼린의 이미지, 음성 또는 초상을 승인 없이 사용하는 것을 금지하는 영구적인 금지 명령에 동의
2024. 1. 5. 미국 뉴욕 남부지법	Basbanes v. Microsoft Corp. and OpenAI, et al.	• 언론인인 니콜라스 바스벤즈와 니콜라스 게이지는 영리를 목적으로 대규모 언어 모델(LLM)을 훈련하기 위해 자신의 저작물을 무단으로 사용했다며 Microsoft와 OpenAI를 상대로 저작권 침해 소 제기
2023. 12. 27. 미국 뉴욕 남부지법	New York Times Company v. Microsoft Corp. and OpenAI, et al.	• 뉴욕타임스는 저작권 침해, DMCA 위반, 불공정 경쟁, 상표권 희석 등의 혐의로 OpenAI와 Microsoft를 상대로 소 제기 • 뉴욕타임스는 피고들이 자사의 저작권이 있는 기사를 무단으로 사용하여 AI 모델을 학습시킴으로써 자사의 비즈니스와 신뢰도를 훼손했다고 주장
2023. 11. 21. 미국 뉴욕 남부지법	Sancton v. OpenAI Inc., Microsoft Corporation, et al.	• 언론인인 줄리안 생턴은 자신의 저서인 '지구 끝의 미친 집'과 기타 저작권이 있는 저작물을 허가 없이 AI 모델 훈련에 사용했다며 OpenAI와 Microsoft를 상대로 저작권 침해 혐의로 소 제기
2023. 10. 18. 미국 테네시 중부지법	Concord Music Group, Inc. v. Anthropic PBC	• 유니버설 뮤직과 콩코드 뮤직 그룹을 비롯한 여러 음악 퍼블리셔는 AI 스타트업인 Anthropic이 저작권이 있는 노래 가사를 무단으로 사용하여 AI 모델을 학습시켰다는 혐의로 소 제기

일자·지역	소송 당사자	경과
		• 미국 테네시주에서 제기된 이 소송은 Anthropic 이 롤링 스톤즈, 가스 브룩스, 케이티 페리 등의 아티스트의 가사를 무단으로 사용했다고 주장
2023. 9. 19. 미국 뉴욕 남부지법	Authors Guild, et al. v. OpenAI, Inc.	• 작가 조합과 존 그리섬과 조지 R.R. 마틴을 비롯한 저명한 작가들은 OpenAI가 언어 모델을 학습시키기 위해 허가 없이 자신들의 저작물을 무단으로 복제했다며 저작권 침해로 소 제기 • 이들은 OpenAI의 모델이 자신의 책을 기반으로 파생 저작물을 생성하여 사용자가 저렴하게 또는 무료로 텍스트를 제작할 수 있게 함으로써 소설 작가 시장을 위협한다고 주장
2023. 9. 8. 미국 캘리포니아 북부지법	Chabon v. OpenAI, Inc.	• 작가 마이클 세이본, 데이비드 헨리 황 등은 OpenAI가 자신들의 저작물을 동의 없이 GPT 모델 훈련에 사용했다며 소 제기 • 작가들은 ChatGPT가 자신의 저작물에 대한 상세한 분석을 생성하며, 이는 모델이 자신의 자료로 학습되었음을 시사한다고 주장 • 2023. 11. 9. 이 사건은 Tremblay v. OpenAI, Inc. 사건과 병합되어 진행 중
2023. 7. 11. 미국 캘리포니아 북부지법	J.L., C.B., K.S., et al., v. Google LLC	• 원고들은 구글이 이용자 동의 없이 웹 스크랩 데이터와 개인정보를 이용하여 바드 챗봇과 같은 AI 제품을 개발했다는 이유로 소 제기 • 캘리포니아 CCPA 및 지식재산권 위반이 주장됨 • 2024. 6. 6. 원고 청구 기각 판결
2023. 7. 7. 미국 캘리포니아 북부지법	Silverman, et al. v. OpenAI, Inc.	• 사라 실버맨, 크리스토퍼 골든, 리처드 카드리는 저작권 침해, DMCA 위반, 부당이득을 이유로 OpenAI를 상대로 소 제기 • 이들은 자신들의 저작권이 있는 책이 동의 없이 ChatGPT 훈련에 사용되었다고 주장 • 2023. 11. 8. 이 사건은 Tremblay v. OpenAI, Inc. 사건과 병합 • 2024. 2. 12. 법원은 원고의 주장 부분 기각

일자·지역	소송 당사자	경과
2023. 7. 7. 미국 캘리포니아 북부지법	Kadrey, et al. v. Meta Platforms, Inc.	• 위 소송과 동일한 원고들은 Meta를 상대로 유사한 내용의 소 제기 • 원고들의 저작권이 있는 책 중 상당수가 EleutherAI라는 연구기관에서 수집한 데이터셋에 포함되어 있었으며, 이는 LLaMA 훈련의 일환으로 복사 및 수집된 것이라고 주장 • 2023. 11. 9. 법원은 LLaMA의 결과물이 원고들의 저작물과 실질적으로 유사하지 않다는 이유로 원고 청구 대부분 기각
2023. 6. 28. 미국 캘리포니아 북부지법	Tremblay v. OpenAI, Inc.	• 작가들은 OpenAI가 ChatGPT 개발 과정에서 작가들이 저술한 책의 텍스트를 포함한 대량의 데이터를 무단으로 사용함으로써 저작권 침해, DMCA 위반 및 부정경쟁행위를 했다고 주장 • OpenAI가 고의로 ChatGPT가 저작물의 일부 또는 요약을 저작자 표시 없이 출력하도록 하여 부당이득 취했다고 주장 • 2023. 11. 8. Silverman, et al. v. OpenAI, Inc. 및 Chabon v. OpenAI, Inc. 사건과 병합
2023. 6. 28. 미국 캘리포니아 북부지법	P.M. et al v. OpenAI LP et al	• 12명의 미성년자가 OpenAI와 Microsoft를 상대로, 이용자 동의 없이 미성년자를 포함한 개인정보를 불법 수집하여 ChatGPT 및 Dall-E와 같은 AI 제품을 개발했다고 주장하며 소 제기 • 2023. 9. 15. 원고들은 소 취하를 하였으며, 이는 법정 외 합의하였음을 시사
2023. 7. 14. 미국 조지아 북부지법	Walters v. OpenAI LLC	• 라디오 진행자인 원고는 ChatGPT가 한 기자에게 원고가 사기 혐의로 기소된 피고인이라고 잘못 설명하였다는 이유로, OpenAI를 상대로 손해배상 청구
2023. 4. 3. 미국 캘리포니아 중앙지법	Young v. NeoCortext, Inc.	• NeoCortext는 딥페이크 앱인 Reface의 개발사이며, 유명인인 원고는 NeoCortext가 자신의 퍼블리시티권을 침해했다고 소 제기

일자·지역	소송 당사자	경과
2023. 2. 15. 미국 캘리포니아 북부지법	Flora, et al., v. Prisma Labs, Inc.	• 이미지 생성앱 Lensa A.I.,의 개발사 Prisma Labs에 대하여 원고는 얼굴 형상 스캔을 수집·이용하였음에도 이용자에게 제대로 알리지 않았다고 주장 • 2023. 8. .8. 법원은 피고의 손을 들어주며 중재를 강제하는 신청을 승인
2023. 2. 3. 미국 델라웨어 지방법원	Getty Images (US), Inc. v. Stability AI, Inc.	• 원고는 피고 회사의 Stable Diffusion AI 가 수백만 장의 원고 소유 사진을 무단으로 사용하여 AI 모델을 학습시켰다며 저작권 침해로 미국에서 소 제기 • 원고는 위 AI 결과물에 워터마크가 수정된 경우가 많아 원고 이미지와 혼동을 야기하여 상표권 침해도 주장
2023. 1. 13. 미국 캘리포니아 북부지법	Andersen, et al. v. Stability AI LTD., et al.	• 3명의 아티스트들은 스테이블 디퓨전을 포함한 AI 이미지 생성 프로그램이 자신의 작품을 무단으로 사용하여 저작권을 침해하는 파생 저작물을 생성하여 막대한 침해를 발생시켰다고 소 제기 • 2023. 10. 30. 법원은 사라 앤더슨의 Stability AI에 대한 저작권 침해 주장을 제외하고는 피고들의 소 각하 청구 모두 인용 • 미국에서는 저작권 침해 민사소송 제기 위해 저작물 등록해야 하나 원고는 소 제기 전 저작물을 등록하지 않음 • 다만, 법원은 원고들에게 소장 수정 기회 제공(실질적 유사성 등 입증 필요 시사)
2023. 1. 12. 영국 런던 고등법원	Getty Images (UK), Inc. v. Stability AI Ltd.	• 위 미국 소송과 유사한 내용으로 원고가 Stability AI를 상대로 영국에서 소 제기
2022. 11. 10. 미국 캘리포니아	J. Doe v. Github, Inc., et al.	• 원고들은 Github, Microsoft, OpenAI를 상대로 피고들의 AI 도구인 Copilot이 DMCA 위반, 계약 위반, 부당이득 위반을 이유로 소 제기

일자·지역	소송 당사자	경과
북부지법		• 원고는 Copilot이 저작자 표시 없이 소스코드를 사용하여 학습되었으므로 오픈소스 라이선스를 위반했다고 주장 • 2024. 7. 9. 법원은 DMCA 위반이라는 원고 주장 기각 • 계약 위반 및 오픈소스 라이선스 위반 여부는 심리 중
2020. 5. 6. 미국 델라웨어 지방법원	Thomson Reuters Enterprise Centre GmbH et al v. ROSS Intelligence Inc.	• 원고는 피고가 생성 AI 기반 법률 연구 플랫폼의 학습데이터로 사용하기 위해 (라이선스가 거부된 후) 자사의 Westlaw 데이터베이스 전체를 복사했다고 주장

자료 : TFL Media, Inc., “From ChatGPT to Getty v. Stability AI: A Running List of Key AI-Lawsuits”; CourtListener.com(<https://www.courtlistener.com/>), Justia Connect(<https://connect.justia.com/>) 재구성

2. 학습데이터로서의 저작물 이용

가. 저작권법상 보호대상인 학습데이터

학습데이터는 ① 재산권으로 보호되는 것과 ② 인격권 등 재산권 이외의 권리로 보호되는 것으로 구분될 수 있다. 이 중에서 저작권 등 지식재산권은 ①에 해당한다.

지식재산권으로 저작권법, 콘텐츠산업진흥법, 부정경쟁방지 및 영업비밀보호에 관한 법률 등을 들 수 있다.

나. 인공지능의 저작물 학습 관련 특징

최근 연구에 따르면 생성형 인공지능이 학습데이터의 이미지를 기억했다가 그대로 복제하는 현상이 발생할 수 있는 것으로 파악되며,¹⁶⁵⁾ AI 학습에 대한 공정이용 인정 여부에 판단할 때 고려할 필요가 있다

구글, 딥마인드, UC버클리, 프린스턴, ETH취리히 대학의 AI 연구원들은 Stable Diffusion 등 생성형 인공지능이 일부 이미지를 기억(memory)하고 있어 개인정보 보호와 저작권에 영향을 줄 수 있음을 보여주었다.

연구원들은 Stable Diffusion에서 사용된 1억 6천만 개의 이미지 데이터셋에서 35만 개의 기억 가능성이 높은 이미지를 테스트했으며, 이 중 직접 일치하는 이미지 94개와 지각적으로 근접 일치하는 이미지 109개를 추출하여 약 0.03%의 복제율을 보였으며, 이미지 데이터셋 크기에 비해 AI 모델 용량이 작아 많은 데이터를 기억할 수 없기 때문에 통계적으로 적게 나타났다고 설명하였다.

구글 Imagen AI 모델에서 가장 많이 중복된 상위 1천 개의 이미지를 대상으로 실험한 결과 2.3%의 복제율이 나타남을 보였으며, 이 경우에도 확률이 높지는 않지만 무시하기는 곤란하며, AI가 생성한 결과물의 저작권 침해 가능성이 존재할 수 있다는 점을 나타낸다.

165) Nicholas Carlini et al., 「Extracting Training Data from Diffusion Models」, arXiv:2301.13188, 2023.

[그림 4] 생성형 인공지능에 의해 복제된 이미지 사례

(윗줄의 이미지는 원본, 아래줄의 이미지는 추출 이미지)



자료: Nicholas Carlini et al.

생성형 인공지능이 자신이 학습한 저작물과 동일하거나 유사한 결과물을 내놓는 확률이 무시할 수 없는 수준이라면, AI 학습에 대하여 공정이용을 인정하면 원저작권자의 권리를 침해할 위험성이 있어 이를 고려할 필요가 있다.

다. 학습데이터 이용에 따른 저작권 침해

인공지능 모델이 저작물을 학습하는 경우 저작물의 복제·전송이 이루어진다. 정당한 권원에 의하지 않는 저작물 이용은 저작권 침해에 해당하는데, 저작권자의 이용허락이나 저작권 침해에 대한 제한(공정이용 등)이 정당한 권원에 해당한다.

그런데 대규모 데이터를 기계적으로 분석학습하는 경우 수많은 저작권자들로부터 개별적으로 동의를 받는 것이 사실상 불가능하다. 즉, 공정이용에 해당하지 않는다면 저작권 침해가 이루어지게 된다.

라. 학습데이터 이용의 면책 여부 검토

(1) 개요

저작권 면책과 관련한 입법 방식은 통상 포괄적 공정이용과 TDM(Text-Data

Mining) 면책 조항으로 구분된다. 한국과 미국은 포괄적 공정이용을 규정하고 있으며, 주요국의 저작권 면책 관련 입법 현황은 다음과 같다.¹⁶⁶⁾

[표 12] 주요국의 저작권 면책 관련 입법 현황

구분	국가	관련 법령	주요 내용
공정이용 일반조항	한국	「저작권법」 제35조의5	저작물의 통상적인 이용 방법과 충돌하지 않고 저작자의 정당한 이익을 부당하게 해치지 않는 경우 저작물 이용 가능
	미국	「연방저작권법」 제107조	공정이용의 요건을 충족하는 경우 비평·논평· 시사보도·교수·학문·연구 등 목적으로 저작물 이용 가능
TDM	EU	「EU 디지털 단일시장 저작권 지침」	① 연구기관 및 문화유산기관의 학술목적 TDM(배제 불가), ② 적법하게 접근할 수 있는 저작물 등에 대한 TDM(저작권자가 opt-out 행 사 가능) 수행 목적으로 저작물 및 데이터베이 스의 복제·추출 허용
	영국	「저작권법」 제29A조	비상업적 목적에 한하여 TDM 허용
	일본	「저작권법」 제30조의4, 제47조의5	· ‘정보해석 용도 제공’ 등 저작물에 표현된 사상이나 감정을 자신 또는 타인이 향수 할 목적이 아닌 경우 면책 · ‘정보해석의 실시 및 결과 제공’ 등 컴퓨터를 이용한 정보처리를 통해 새로운 지식·정보 를 창출함으로써 저작물의 이용촉진에 기여 하는 행위를 하는 자가 일정한 행위에 부수 하여 경미하게 이용하는 것은 면책
공정이용 일반조항 + TDM	싱가 포르	「저작권법」 제243조, 제244조 등	· 목적 내 이용 및 저작물 등에 대한 적법 한 접근 권한 등의 요건 하에, 컴퓨터 데 이터 분석 및 그 준비작업 목적으로 저작 물 이용 허용 ※ 공정이용 규정이 있음에도 목적 제한 등이 없는 TDM 규정 신설

166) 애플경제, “챗GPT 계기, 데이터 개방 위한 ‘TDM’ 부각”, 2023. 3. 27. <<https://www.apple-economy.com/news/articleView.html?idxno=71167>> 최종방문일 2024. 9. 10.>

EU, 영국, 일본, 싱가포르 등은 TDM 활동에 대한 면책 규정을 두고 있다. 미국, 한국은 별도로 TDM 활동에 대한 명문의 면책조항을 두고 있는 것은 아니나, 공정이용 원칙을 통해 TDM의 허용 여부를 사례별로 판단하고 있다.

중국과 호주는 TDM 관련 규정과 포괄적 공정이용을 규정하지 않고 있다. 중국은 2023년 8월에 생성형 인공지능 서비스에 대한 잠정 행정 조치가 발효되었는데, AI 서비스제공자들에게 학습데이터가 다른 사람의 IP를 침해하지 않도록 요구하는 내용이 포함되어 있다.¹⁶⁷⁾ 호주에서도 최근 활발한 논의가 진행되고 있다.

한편, TDM 규정을 도입한 대부분의 국가들도 생성형 인공지능의 출현 이전에 입법이 이루어져 생성형 인공지능에 대한 적용 여부에 관해서는 논란이 있는 상황이다.

(2) 포괄적 공정이용

(가) 현황

포괄적 공정이용은 한국, 미국, 싱가포르 등에 입법되어 있는 저작권 제한 사유이다.

미국은 「연방저작권법」 제107조에 규정하고 있으며, 한국은 2011년 한미 FTA협정 당시 위 조항을 한국 저작권법에 거의 그대로 도입하였다.

한국 저작권법

제35조의5(저작물의 공정한 이용) ① 제23조부터 제35조의4까지, 제101조의3부터 제101조의5까지의 경우 외에 저작물의 일반적인 이용 방법과 충돌하지 아

167) Yi Wu, “Interim Administrative Measures for Generative Artificial Intelligence Services”, China Briefing, 2023. 7. 27. <<https://www.china-briefing.com/news/how-to-interpret-chinas-first-effort-to-regulate-generative-ai-measures/>> 최종방문일 2024. 9. 11.>

니하고 저작자의 정당한 이익을 부당하게 해치지 아니하는 경우에는 저작물을 이용할 수 있다.

② 저작물 이용 행위가 제1항에 해당하는지를 판단할 때에는 다음 각 호의 사항등을 고려하여야 한다.

1. 이용의 목적 및 성격
2. 저작물의 종류 및 용도
3. 이용된 부분이 저작물 전체에서 차지하는 비중과 그 중요성
4. 저작물의 이용이 그 저작물의 현재 시장 또는 가치나 잠재적인 시장 또는 가치에 미치는 영향

미국 저작권법(United States Code Title 17–Copyrights)

제107조 배타적 권리에 대한 제한 : 공정이용(Fair Use)

제106조 및 제106조의 A의 규정에도 불구하고 비평, 논평, 시사보도, 교수(학습용으로 다수 복제하는 경우를 포함한다), 학문 또는 연구 등과 같은 목적을 위하여 저작권으로 보호되는 저작물을 복제하거나 음반으로 제작하거나 또는 기타 제106조 및 제106조의 A에서 규정한 방법으로 사용하는 경우를 포함하여 공정 사용하는 행위는 저작권 침해가 되지 아니한다. 개별적인 사건에 있어서 저작물의 사용이 공정이용인가의 여부를 결정함에 있어서 다음을 참작하여야 한다.

- (1) 당해 사용이 상업적 성질의 것인지 또는 비영리적 교육목적을 위한 것인지 등 그 사용의 목적 및 성격
- (2) 저작권으로 보호되는 저작물의 성격
- (3) 사용된 부분이 저작권으로 보호되는 저작물 전체에서 차지하는 양과 상당성, 그리고
- (4) 이러한 사용이 저작권으로 보호되는 저작물의 잠재적 시장이나 가치에 미치는 영향

위의 모든 사항을 참작하여 내려지는 결정인 경우에, 저작물이 미발행되었다는 사실 자체는 공정사용의 결정을 방해하지 못한다.

미국은 TDM과 관련된 개별 면책 조항을 마련하지 않고, 공정이용 조항을 통해 해결을 도모하고 있다. TDM 목적의 저작물 복제 등이 공정이용으로 판단될 가능성에 대해 긍정적이라는 견해가 있다. 주요 근거로는 TDM 과정에서 저작물 이용이 새로운 정보나 가치 또는 의미를 창출하는 변형적 이용(transformative use)에 해당한다는 점과, TDM 결과가 원저작물의 수요를 대체하기 어렵다는 점이 있다.

일부에서는 TDM을 ‘비표현적 이용(non-expressive use)’으로 규정하여, 공정이용 분석의 제1요소인 ‘이용의 목적·성격’과 관련된 변형적 이용의 범주로 묘사한다. ‘비표현적 이용’은 ‘비소비적 이용(non-consumptive use)’으로도 불리며, 이는 변형적 이용의 한 속성을 생산적 이용으로 표현하는 것과 일맥상통한다. 그러나 이러한 기준은 컴퓨터 분석을 위한 데이터베이스화 및 저장과 같은 경우를 설명하는 데 한계가 있으며, 생성형 인공지능과 같은 ‘표현적’ 이용 사례를 고려하지 못한다는 비판도 있다.

(나) 미국의 공정이용 관련 판례¹⁶⁸⁾

해당 조항은 추상적인 요건들로 구성되어 있어, 어떤 사안이 공정이용에 해당하는지 케이스별로 살펴보아야 한다. 인공지능 저작물 학습이 공정이용에 해당하는지 판단하는데 앞서 일반적인 공정이용 법리를 알아볼 필요가 있으며, 이에 따라 미국의 주요 판례를 살펴본다. 다음의 판례들은 소위 ‘비표현적 이용(non-expressive use)’에 관한 것으로서, TDM이나 AI 학습에서도 비표현적 이용이 이루어진다는 점에서 시사점을 도출할 수 있다.

168) 이대회, 「Andy Warhol 케이스의 ‘변형적 이용(Transformative Use)’의 해석과 AI 학습데이터 이용에 대한 영향」, 한국저작권위원회, 이슈리포트 2023-8, 2023. 11.

1) A.V. ex rel. Vanderhye v. iParadigms, LLC 케이스(2009)

이 케이스는 표절방지를 위한 온라인 시스템에 관한 것인데, 법원은 이 시스템이 표절방지를 위한 데이터베이스로 활용하기 위하여 원고(학생)가 제출하였던 논문을 계속 저장하는 것이 공정이용에 해당한다고 판시하였다. 법원은 그 근거로서 iParadigms가 저작물(논문)을 완전히 다른 목적(표절방지)을 위하여 이용(저장)하고 있고, 교육기관에게 상당한 공익을 제공하고 있고, 상업적 성격(원고가 저작물을 영리를 위하여 이용)은 저작물 이용의 변형적 성격을 고려하면 중요한 의미를 가지지 않으며, 원저작물을 변경하거나 원저작물에 추가하지 않더라도 변형적일 수 있으며, 저작물의 이용이 표현적 콘텐츠와 전혀 관계되지 않고 표절을 찾아내고 이를 억제하기 위한 것이라는 것 등을 제시하였다.

2) Authors Guild, Inc. v. HathiTrust 케이스(2014)

Authors Guild 케이스는 HathiTrust가 Google이 디지털 스캔한 복제물에 의하여 제공하였던 검색서비스에 관한 것인데, AI 학습데이터와 연관되는 것은 일반공중이 모든 디지털 복제물에 대하여 특정 키워드를 이용하여 검색할 수 있도록 하는 서비스(full-text search)이다.¹⁶⁹⁾ 검색결과는 키워드가 포함된 페이지 숫자와 해당 키워드가 각 페이지에 나타난 횟수에 불과하고 키워드를 포함하고 있는 저작물은 제공하지 않는다. 따라서 이용자들은 해당 키워드가 있는 페이지나 책의 일부분을 볼 수가 없다.

169) Google은 2004년 여러 대학의 소장 장서를 스캐닝하였는데, 2008년 13개 대학은 HathiTrust Digital Library(HDL)를 운영하여 1,000만 이상의 디지털 복제본에 대하여 3가지 서비스를 제공하기로 하였다. HDL은 ① 일반공중이 검색할 수 있는 서비스 외에, ② 시각장애자가 저작물 전체에 접속할 수 있는 서비스와 ③ 회원 도서관이 가지고 있던 장서의 멸실 등에 대하여 디지털 복제본으로 보유할 수 있도록 하는 서비스를 제공하였다. 법원은 위 3개 서비스가 모두 공정이용에 해당한다고 판시하였다.

연방 항소법원은 검색결과와 목적, 성격, 표현, 의미 및 메시지가 원저작물의 페이지와 다른 것으로서, 원저작물과 HDL 검색결과는 유사하지 않으며, 저작자들이 검색하게 할 목적으로 서적을 작성하였다는 증거가 없으므로, 원창작의 대상이나 목적을 대체하지 않으며, HDL은 원저작물을 단순히 재포장·재공표하거나 새로운 방식으로 제시(presentation)하는 것이 아니라, 원저작물에 다른 목적 및 성격이 있는 새로운 무엇인가를 추가하는 것으로서, 변형적 이용(transformative use)에 해당한다고 판시하였다.

3) Andy Warhol Foundation for the Visual Arts, Inc. v. Goldsmith 케이스(2023)

[그림 5] Goldsmith의 사진(원저작물)과 Warhol의 작품 ‘Orange Prince’(새로운 저작물)



자료 : 미국 연방대법원, ANDY WARHOL FOUNDATION FOR THE VISUAL ARTS, INC. v. GOLDSMITH ET AL. 판결문

2016년 잡지사는 AWF로부터 Orange Prince에 대한 이용허락을 받아 잡지에 이를 사용하였는데, Goldsmith를 작가로 표시하지도 않았다. 이 케이스의 쟁점은 2016년 AWF가 잡지사에게 이용허락한 행위와 관련하여, 이용허락의 대상이 되었던 Orange Prince를 이용허락한 행위, 즉 AWF가 2016년 이용허락에 의하여 저작물을 잡지에 이용하도록 하는 행위가 변형적 이용에 해당하는가 여부이다.

변형적 이용 여부를 판단함에 있어서 ‘원저작물의 대상(또는 창작)을 대체(supersede)’하는 것에 불과하거나 ‘추가적인 목적(further purpose)’이나 ‘다른 성격(different character)’이 있는지 여부가 쟁점이다.

연방대법원은 ① Goldsmith의 사진과 AWF의 2016년 Orange Prince 이용허락행위는 Prince에 관한 기사에 사용하도록 허락한 것으로서 실질적으로 목적이 동일하고, ② AWF가 잡지사에게 Orange Prince를 이용허락하는 행위는 Goldsmith 사진의 대상을 대체하는 것으로 볼 수 있으며, ③ AWF가 Goldsmith의 사진을 이용하는 행위가 상업적이라는 것은 AWF의 이용행위가 공정이용이 되지 않는 방향으로 작용한다고 판단하였다.

(다) 최근 사건에서의 공정이용 관련 당사자 주장¹⁷⁰⁾

1) 사건의 경위

Concord Music Group 등(이하 ‘Concord’)은 2023년 7월에 Anthropic PBC(이하 ‘Anthropic’)에 대하여 저작권 침해 소송을 제기하였다. 이는 음악저작물의 저작권자가 가사(lyric)에 대하여 제기한 AI 저작권 소송 케이스로 저작물 이용의 증명과 공정이용 여부에 대한 것이다.

위 사건 당사자들의 주장은 각 저작권자와 인공지능 개발자의 입장을 잘 정리하고 있어 다음에서 정리하였다.

170) 이대회, 「음악 생성 AI 모델에 대한 저작권 소송과 학습데이터 이용의 공정이용 여부 주장 분석」, 한국저작권위원회, 이슈리포트 2024-24, 2024. 8.

2) 저작물 이용의 증명

Concord는 이용자들이 Claude를 통하여 광범위한 범위의 가사를 단어 그대로의(verbatim), 또는 단어 그대로인 것과 동일한 가사를 생성할 수 있다고 주장하였다. 예컨대 Claude에 ‘What are the lyrics to Friends in Low Places by Garth Brooks’를 프롬프트로 입력하였을 경우, 가수 Garth Brooks가 부른 ‘Friends in Low Places’(1990년 출시)라는 가사와 사실상 동일한 가사가 제공됨을 주장하였다.

3) 저작권자의 공정이용 부인 주장

가) 저작물 이용의 목적과 성격

Concord는 Anthropic이 Claude API에 접근하고 산출물(가사)을 생성할 때마다 수익을 얻고, 이용자(소비자)에게도 구독 서비스 비용을 부과함으로써, Anthropic의 저작물 이용이 상업적이라며 아래와 같이 주장하였다.

첫째로, Concord는 Anthropic의 이용에 대해 이용자의 요구에 따라 저작물의 창의적 표현을 제공한다는 점에서 동일하기 때문에 Anthropic의 이용이 변형적 이용이 되지 않는다고 주장하였다.

둘째로, Concord는 Anthropic이 원고의 저작물을 학습데이터로 사용하는 것은 원저작물을 비평·비판하기 위한 것이 아니고, 원저작물에 관하여 얻을 수 없는 정보를 제공하기 위한 것도 아니기 때문에, 역시 변형적 이용이 아니라고 주장하였다.

셋째, Concord는 Anthropic이 원고의 저작물과 동일하거나 거의 동일한 AI 산출물을 제공하는 것은 변형적 이용이 되지 않는다고 주장하였다.

넷째, Concord는 Anthropic이 새로운 기술을 이용하여 저작물을 모두 복제하는 것은 창의적 변형과 전혀 유사하지 않다고 주장하였다.

마지막으로 Concord는 Anthropic의 Claude가 원고 저작물의 형태를 변형

하거나 다른 텍스트와 결합하는 경우 2차적 저작물 작성권을 침해하는 것으로서 변형적이지 않다고 주장하였다.

나) 저작물의 성격

Concord는 음악저작물(가사 포함)이 저작권이 보호하고자 하는 목적의 핵심에 속하는 것으로서 일반 공중에게 배포하기 위한 창의적 표현이 포함되어 있는 유형이라거나 저작권 보호의 핵심에 해당하는 것이라고 주장하였다.

다) 이용된 부분의 양과 상당성

Concord는 Anthropic의 전체 저작물을 AI 모델의 학습데이터로 복제하고, 저작물 전체 또는 거의 전체에 가까운 정도의 산출물을 생성한다고 주장하였다. 또한 Claude가 가사의 일부분을 포함하는 산출물을 생성하는 경우에도 그 자체 또는 다른 가사와 함께 혼합되어, 저작물의 가장 뚜렷하고 인식가능한, 즉 특징적인 요소를 복제한다고 주장하였다.

라) 저작물 시장과 가치에 미치는 영향

Concord는 Anthropic이 아무런 제약을 받지 않고 저작물을 이용하는 경우 저작물 시장을 손상할 수 있다고 다음과 같이 주장하였다.

첫째, Anthropic 서비스를 광범위하게 사용하게 된다면 가사에 대한 이용허락 시장이 수축되어 저작물의 가치를 감소시키고, 이용허락에 대한 수요를 약화시키며, 이용허락을 받은 주체가 사업유지하는 것을 위태롭게 한다.

둘째, Anthropic은 저작권자의 2차적 저작물 작성권을 침해하고 이에 따라 2차적 저작물 시장을 손상한다.

셋째, Anthropic은 현재 존재하고 있는 시장을 약화시킬 뿐만 아니라 AI 학습데이터로 이용되기 위하여 지금 막 성장하려는 시장까지 억누를 위험성이 있다.

4) 개발자의 반론

가) 중간과정의 복제와 공정이용

Anthropic은 AI 모델의 산출물이라는 완제품에서는 저작물이 복제되지 않고 학습과정에서 저작물을 사용하는 것이라고 하여 중간과정의 복제물이라고 주장하였다. 다만 이는 프로그램 역분석에서 사용되는 주장으로서 음악 가사 저작물에도 적용될지 여부는 법원의 판단을 지켜봐야 한다.

나) 저작물 이용의 목적과 성격

첫째, Anthropic은 생성형 인공지능 모델을 훈련시키기 위하여 토큰 데이터셋의 일부로서 Concord 저작물의 가사를 이용하는 것은 변형적이라고 주장하였다. 즉, 가사는 복제물 그대로 저장하는 것이 아니라 모델 내의 통계적인 가중치(weight)를 얻기 위하여 작은 토큰으로 분리하는 것을 근거로 삼고 있다.

둘째, Anthropic은 가사를 이용하는 것이, 사람이 가사를 창작하는 것과 동일한 목적이 아니라 새로운 목적, 곧 신경망(neural network)으로 하여금 사람의 언어가 어떻게 작동하는지 교육시키기 위한 것이므로, 변형적이라고 주장하였다.

셋째, Anthropic은 변형적 이용에 대한 원고의 주장이 사실이 아니거나 법에 어긋나는 것이라고 지적하고 있다. 즉 Claude의 학습 목적은 이용자의 요구에 따라 원고 저작물의 창작성 있는 표현을 전달하기 위한 것이거나, 비판이나 비평을 추가하지 않고 저작물을 전체적으로 복제하는 경우 변형적 이용이 되지 않는다는 원고의 주장을 배척하고 있다. 또한 Anthropic은 영리기업이라고 하여 결과가 달라지지 않는다고 주장하고 있다.

다) 저작물의 성격 및 이용된 부분의 양과 상당성

Anthropic은 얼마나 많이 있는 그대로 복제를 하였느냐가 아니라, 학습 목적을 달성하기 위하여 필요한 부분 이상이 사용되지 않았다는 것이 중요하다고 주장하였다. 따라서 원저작물과 다른 목적을 위하여 사용하는 경우, 저작물을 있는 그대로 전체를 복제하는 것도 정당화될 수 있는데, 바로 이 사건이 여기에 해당한다는 것이다.

라) 저작물 시장 및 가치에 대한 영향

첫째, Anthropic은 학습데이터 이용허락 시장이 존재하지 않는다고 주장하였다. 원고는 Claude를 학습시키기 위하여 이용허락을 받지 않고 저작물을 이용함으로써 현재 막 자라나고 있는 AI 학습데이터 시장을 억누를 위험성이 있다고 주장하였는데, Anthropic은 ‘가사’에 대한 범용 AI 학습데이터로서의 시장이 존재한다거나 미래에 개발될 것으로 법원이 판단할 근거가 없다고 주장하고 있다.

둘째, Anthropic은 Concord가 학습데이터 이용에 대한 사용료 징수를 ‘희망’한다는 것이 법적으로 무의미한 것이라고 주장하였다. 원고는 AI 기업이 다른 유형의 저작권자와 거래한 것(뉴스, 이미지에 대한 이용허락)을 지적하고 있는데, 이러한 거래들이 AI 모델을 학습시키기 위한 이용허락에 해당하는지 여부에 대한 아무런 증거가 존재하지 않는다는 것이다. 또한 이용허락이 몇 개 존재한다고 하더라도, 다른 것을 사유로 하여 이루어진 개별적인 이용 허락이 아니라, 진정한 시장이 존재할 것을 요구하는 공정이용에 따른 결과를 변경시키지 않는다고 주장하고 있다.

셋째, Anthropic은 Claude가 Concord의 저작물을 학습자료로 이용한다고 하여 현재 존재하는 이용허락 시장을 대체할 위험은 존재하지 않는다고 주장하였다. 원고의 주장은 Anthropic이 학습데이터 이용에 대하여 이용허락을 받지 않았다는 것이고, 법원은 개별적인 공정이용 케이스에서 피고가 이용허락

을 받지 않았다는 것에 대해서만 결정하여야 한다고 주장하고 있다. Anthropic의 이러한 주장은 Concord의 가사를 학습데이터로 이용한다고 하더라도 가사 서비스를 제공하는 시장이 영향받지는 않을 것이고, 기존의 가사 서비스 이용 허락 시장에 대한 영향을 고려할 것이 아니라 학습데이터로 사용하는 것에 대하여 이용허락을 받지 않았다는 것에 한정하여 판단하여야 한다는 것이다.

(라) 시사점

AI를 학습(훈련)시키는 행위 자체는 비표현적 이용에 해당하고 상업적 목적을 위한 것이 아니라고 볼 여지가 있다. 저작권이있는 자료를 이용하여 인공지능 프로그램과 제품을 ‘훈련’시키는 것은 ‘혁신성’이 인정될 수 있으며, 인공지능 프로그램을 훈련하기 위한 ‘입력 자료’로 이용하는 것은 공정이용으로 인정될 가능성이 높다는 견해가 있다.¹⁷¹⁾ 미국의 법원은 비표현적, 변형적 이용 시 공정이용을 비교적 넓게 인정하는 경향이 있었다.

그러나 생성형 인공지능 모델이 출시된 이후에는 상업적 성격을 가지고 표현적인 AI 결과물을 생성한다. 생성형 인공지능 모델이 무료로 이용된다고 하더라도 API 등을 통하여 이익을 얻는 것이 일반적이므로 상업적 성격이 충분히 인정될 수 있다. AI 생성물이 학습데이터에 포함되어 있는 저작물에 대한 수요를 단순히 대체하거나 대신할 수 있다고 볼 여지가 있다. 생성형 인공지능이 학습데이터인 저작물에 상응하는 산출물을 제공하므로 변형적 이용이라고 판단된 기존 케이스(표절 확인이나 검색결과 제공 등)와는 다르다고 인정될 가능성이 높다.

Andy Warhol 케이스에서 법원은 원저작물과 변형된 저작물 간 목적으로 실질적으로 동일하다면 공정이용 인정 가능성이 낮아진다고 판단하였다. 생성

171) Gary Myers, 「Artificial Intelligence and Transformative Use After Warhol」, Washington and Lee Law Review Online, Volume 81, Issue 1, 2023. 12.

형 인공지능의 결과물의 목적이 원저작물과 동일하다면 공정이용 가능성이 낮아질 것이다. 반대로, 인공지능 개발자는 인공지능 결과물과 원저작물 간 목적이 실질적으로 다르다는 점을 입증하면 공정이용이 인정될 여지가 있을 것이다.¹⁷²⁾

(3) TDM 면책 조항

(가) 개요

인공지능의 학습데이터 사용을 허용하기 위해 국내에 TDM 면책 규정을 도입하거나 적용하자는 견해도 있다. TDM은 텍스트·데이터 마이닝(Text and Data Mining)의 약자로서, 컴퓨팅 자원을 활용해 대량의 데이터를 분석하고 그로부터 일정한 패턴이나 구조를 추출해 의미 있는 추론을 이끌어 내는 기술을 말한다. 우리나라에서는 2021년 TDM 면책 규정을 저작권법에 도입하는 내용의 저작권법 전부개정법률안이 발의되었고, 이후 다수의 법안이 발의되었지만, 저작권자들의 반대 등으로 통과되지 못하고 임기만으로 폐기되었으며, 현재 제22대 국회에서는 아직 관련 법안이 발의되지 않은 상황이다.

(나) 우리나라 제21대 국회 TDM 면책 규정 법안 폐기 경위

1) 제21대 국회 TDM 면책 조항 발의 법안 개요

우리나라 제21대 국회에서는 2021년 1월 15일에 도종완의원이 대표발의한 저작권법 전부개정법률안 제43조를 시작으로, 총 4개의 TDM 면책 조항에 관한 저작권법 개정안(이하 ‘TDM 법안’)이 발의되었으나 결국 제21대 국회 만료로 최종 폐기되었다.

172) Edward D. Lanquist and Dominic Rota, 「The Impact of the Supreme Court's 'Goldsmith' Decision on Copyright Enforcement Against AI Tools」, Baker Donelson, 2023. 9. 26.

[표 13] 제21대 국회 TDM 면책 조항 발의안

제안일자	의안번호	의안명	약칭	의결결과
2021.01.15	제2107440호	저작권법 전부개정법률안 (도종환의원 등 13인)	도종환의원안	임기만료 폐기
2022.10.31	제2117990호	저작권법 일부개정법률안 (이용호의원 등 14인)	이용호의원안	임기만료 폐기
2023.06.08	제2122537호	저작권법 일부개정법률안 (황보승희의원 등 10인)	황보승희의원안	임기만료 폐기
2023.09.25	제2124685호	저작권법 일부개정법률안 (이인영의원 등 16인)	이인영의원안	임기만료 폐기

이하 각 TDM 법안이 폐기된 경위와 법안 도입에 관한 쟁점 및 이해관계자의 입장이 어떠하였는지 정리하고, 이를 토대로 향후 입법 과정에서 고려할 점을 정리한다.

2) TDM 법안의 내용

도종환의원안은 저작물을 복제·전송할 수 있는 요건으로 저작물에 표현된 사상이나 감정의 향유를 금지하여 인간이 분석과정에 참여하는 것을 허용하지 않으며(이하 ‘비향유적 이용 요건’), 저작물에 적법하게 접근할 것을 규정하여 해킹, 불법 다운로드 등을 통한 복제는 배제(이하 ‘적법한 접근 요건’)하고 있다는 점에서 저작권 제한 범위를 설정하였다.¹⁷³⁾ 또한, TDM을 통해 생성된 복제물은 정보분석을 위하여 필요한 한도에서 보관할 수 있다고 규정하고 있다.

173) 국회 문화체육관광위원회, 「저작권법 전부개정법률안(도종환의원 대표발의, 의안번호 제2107440호) 검토보고서」, 2021. 2.

도중환의원안 제43조(정보분석을 위한 복제·전송)

- ① 컴퓨터를 이용한 자동화 분석기술을 통해 다수의 저작물을 포함한 대량의 정보를 분석(규칙, 구조, 경향, 상관관계 등의 정보를 추출하는 것)하여 추가적인 정보 또는 가치를 생성하기 위한 것으로 저작물에 표현된 사상이나 감정을 향유하지 않는 경우에는 필요한 한도 안에서 저작물을 복제·전송할 수 있다. 다만, 해당 저작물에 적법하게 접근할 수 있는 경우에 한정한다.
- ② 제1항에 따라 만들어진 복제물은 정보분석을 위하여 필요한 한도에서 보관할 수 있다.

이용호의원안은 도중환의원안과 마찬가지로 비향유적 이용 요건 및 적법한 접근 요건을 포함하고 있으며, 복제·전송 및 정보분석을 필요한 한도 안에서 복제물을 보관할 수 있음을 정하고 있다.

여기에서 더 나아가 이용호의원안은 ① 복제물 보관을 위하여 복제방지조치 등 대통령령으로 정하는 필요한 조치를 해야 할 의무를 규정하고, ② 정보분석의 결과물에 대하여 정당한 권리자로부터 저작물의 복제·전송에 대한 이용의 허락을 받지 않은 경우에도 교육·조사·연구 등 비상업적 목적 및 저작물의 창작 목적으로 접근하는 경우에는 부정경쟁방지법 등 다른 데이터 보호법률의 제한으로부터 예외가 됨을 추가로 정하고 있다.

이용호의원안 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 정하

는 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

황보승희의원안도 도종환의원안과 마찬가지로 저작물에 표현된 사상이나 감정을 향유하지 아니하고, 정보분석의 대상이 되는 저작물에 적법하게 접근해야 함을 요건으로 두고 있으며, 정보분석을 필요한 한도 안에서 복제물을 보관할 수 있음을 정하고 있다. 다만, 황보승희의원안은 이용방법으로서 복제·전송 외에 ‘2차적저작물의 작성’을 추가하여 정하고 있다.

황보승희의원안 제35조의5(정보분석을 위한 복제·전송 등)

① 컴퓨터를 이용한 자동화 분석기술을 통하여 다수의 저작물을 포함한 대량의 정보를 해석(패턴, 트렌드 및 상관관계 등의 정보를 추출·비교·분류·분석하는 경우를 말한다. 이하 이 조에서 “정보분석”이라 한다)함으로써 추가적인 정보 또는 가치를 생성하기 위하여 다음 각 호의 요건을 모두 갖춘 경우에는 필요한 범위 안에서 저작물을 복제·전송하거나 2차적저작물을 작성할 수 있다.

1. 해당 저작물에 대하여 적법하게 접근할 것
2. 해당 저작물에 표현된 사상이나 감정을 향유하는 것을 목적으로 하지 아니할 것

② 제1항에 따라 만들어진 복제물은 정보분석을 위하여 필요한 범위 안에서 보관할 수 있다.

이인영의원안은 이전의 법안들과 달리 일본 저작권법의 요소인 비향유적 이용 요건을 제외하였고, ① (주체의 제한) 대학·연구기관 또는 그 밖에 대통

령령으로 정하는 기관이나 단체에서 수행할 것, ② (목적) 교육·조사·연구 등 비상업적 목적으로 이용할 것, ③ (적법한 접근 요건) 정보분석의 대상이 되는 해당 저작물에 대하여 적법하게 접근할 것을 요건으로 정하였다. 복제물의 보관은 기존 법안들과 마찬가지로 정보분석을 위해 필요한 한도 내에서 허용하며, 복제 방지 및 보안 등 기술보안조치를 하도록 정하고 있다.

이인영의원안 제35조의5(교육·조사·연구 목적의 정보분석을 위한 복제·전송 등)

① 컴퓨터를 이용한 자동화 분석 기술을 통하여 추가적인 정보 또는 가치를 생성하기 위하여 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. 대학·연구기관 또는 그 밖에 대통령령으로 정하는 기관이나 단체에서 수행할 것
2. 교육·조사·연구 등 비상업적 목적으로 이용할 것
3. 정보분석의 대상이 되는 해당 저작물에 적법하게 접근할 것

② 제1항에 따라 저작물을 복제하는 자는 정보분석을 위하여 필요한 한도 내에서 복제물을 보관할 수 있다. 이 경우 저작권 및 그 밖에 이 법에 따라 보호되는 권리의 침해를 방지하기 위하여 복제방지 및 보안 등 대통령령으로 정하는 필요한 조치를 하여야 한다.

3) 법안 발의의 경위 및 내용 분석

TDM 법안들은 공통적으로 정보분석을 위한 일련의 저작물 이용 행위에 대해 현행법상 포괄적인 저작재산권 제한 규정인 공정이용 조항이 적용될 여지는 있으나 동 규정은 저작자의 이익을 부당하게 해치지 않는 범위 내에서 저작물의 공정한 이용이 가능하다고 명시하고 있을 뿐이어서 데이터마이닝이 해당 규정에 따라 면책될 수 있는지 여부가 불확실하므로, 개정안과 같이 정보분석을 위한 저작물의 복제·전송을 명시적으로 허용하는 경우 관련 사업자의 법적 안정성을 도모하고자 하는 취지를 담고 있다.¹⁷⁴⁾

도종환의원안은 의원입법의 형식으로 발의되었지만, 실질적으로는 저작권법의 소관 정부부처인 문화체육관광부가 주도하여 2020년 1월 ‘저작권법 전부개정 연구반’을 구성하고 동년 2월 ‘저작권 비전 2030’을 발표한 이후 한국저작권위원회와 선행 연구 및 해외사례 검토를 진행하고, 업계 이해관계자들과 전문가의 의견을 수렴하였으며, 도종환의원실과 함께 온라인 공청회를 거쳐 일반인 의견을 청취하여 법안을 작성하였다.¹⁷⁵⁾ 도종환의원안이 오랜 준비기간을 거쳐 만들어진 것임은 분명하나, 그 이후 발의된 이용호의원안, 황보승희의원안, 이인영의원안을 비롯한 TDM 규정의 경우에는 우리나라의 기술 및 경제 현실에 대한 실증사실에 기반한 것이라기보다는 외국 법제들의 각 요소를 발췌해 만든 법이라는 평가가 있다.¹⁷⁶⁾

먼저, 각 법안의 TDM 정의를 살펴보면, 다음과 같다. 이는 유럽연합의 DSM 지침 및 이를 참고한 독일 등 유럽연합 회원국의 정의 내용과 유사하다.

[표 14] 제21대 국회 발의 법안의 TDM 정의 비교

TDM 정의	
도종환의원안	컴퓨터를 이용한 자동화 분석기술을 통해 다수의 저작물을 포함한 대량의 정보를 분석(규칙, 구조, 경향, 상관관계 등의 정보를 추출하는 것)하는 것

174) 국회 문화체육관광위원회, 「저작권법 전부개정법률안(도종환의원 대표발의, 의안번호 제2107440호) 검토보고서」, 2021. 2.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이용호의원 대표발의, 의안번호 제2117990호) 검토보고서」, 2022. 12.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의, 의안번호 제2122537호) 검토보고서」, 2023. 9.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이인영의원 대표발의, 의안번호 제2124685호) 검토보고서」, 2023. 11.

175) 류시원, 「인공지능 시대 저작권 정책 형성절차에 관한 제언- 텍스트·데이터 마이닝 예외규정 입법 논의를 중심으로 -」, 법제처, 2024. 3., 116쪽

176) 류시원, 위 논문, 119쪽

TDM 정의	
이용호의원안	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)
황보승희의원안	컴퓨터를 이용한 자동화 분석기술을 통하여 다수의 저작물을 포함한 대량의 정보를 해석(패턴, 트렌드 및 상관관계 등의 정보를 추출·비교·분류·분석하는 경우를 말한다. 이하 이 조에서 “정보분석”이라 한다)
이인영의원안	컴퓨터를 이용한 자동화 분석 기술을 통하여 추가적인 정보 또는 가치를 생성하기 위하여 다수의 저작물을 포함한 대량의 정보를 분석(규칙, 구조, 경향 및 상관관계 등의 정보를 추출하는 경우를 말한다. 이하 이 조에서 “정보분석”이라 한다)하는 것
EU DSM	패턴, 트렌드, 상관관계 등의 정보 생산을 위해 디지털 형태 텍스트와 데이터를 분석하는 것을 목적으로 하는 자동화 기술

TDM을 정의하는 방식은 다양할 수 있으며, 특히 TDM을 통해 정보를 분석하는 과정에서 그 학습의 데이터가 되는 정보가 여러 법제에서 규율하고 있는 다양한 형태의 정보에 해당할 가능성이 열려있으므로, 법제의 체계적 정합성을 위한 고민도 수반되어야 할 것이다.

[표 15] 제21대 국회 발의 법안별 TDM 규정 주요 내용 비교 분석

	도종환의원안	이용호의원안	황보승희의원안	이인영의원안
상업적 이용가능성	O	O	O	X
주체 제한	X	X	X	대학·연구기관 등
적법한 접근 요건	O	O	O	O
비향유적 이용 요건	O	O	O	X
이용방법	복제·전송	복제·전송	복제·전송·2차적 저작물 작성	복제·전송
결과물 보관	O	기술적 조치의무	O	기술적 조치의무

	도종환의원안	이용호의원안	황보승희의원안	이인영의원안
옵트아웃제도	X	X	X	X
강행규정성	-	-	-	-

한편, 각 법안에 규정된 주요 요소를 살펴보면, 도종환·이용호·황보승희의원안은 상업적 이용가능성을 인정하고, 주체를 제한하지 않으며, 적법한 접근 요건과 비항유적 이용요건을 정하고 있고, 복제·전송의 이용을 허용하며(황보승희의원안의 경우 2차적저작물 작성도 포함함), 결과물인 복제물 보관을 허용하고 있으나, 옵트아웃제도 및 명시적인 강행규정성을 규정하고 있지 않다는 점에서 대동소이하다. 이 법안들은 특히 비항유적 이용요건, 주체·목적 제한 및 상업적 이용의 가능성 등을 고려할 때 일본 저작권법의 요소를 일부 가져온 것으로 보이며, TDM 정의조항, 이용방법, 적법한 접근 요건, 결과물 보관 여부 및 그 범위에 관하여는 EU의 DSM 지침을 반영한 것으로 보인다.¹⁷⁷⁾

이인영의원안의 경우 상업적 이용을 불허하고, 대학·연구기관 등 주체를 한정하며, 적법한 접근 요건을 도입하고, 기술적 조치의무를 도입한 결과물 보관을 허용하며, 옵트아웃제도를 도입하고 있어 이러한 면에서 EU의 DSM 3조의 규율 형태를 닮아있다.

4) 제21대 국회 발의 TDM 법안의 쟁점 및 이해관계자의 입장 정리

각 법안 발의에 따라 국회 소관위원회¹⁷⁸⁾에서 검토된 주요 쟁점과 그 의견을 정리하면 다음과 같다.

177) 류시원, 위 논문, 120쪽

178) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(도종환의원 대표발의) 검토보고」, 2021.2.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이용호의원 대표발의) 검토보고」, 2022.12.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이인영의원 대표발의) 검토보고」, 2023. 11.의 내용을 종합하여 정리하였다.

가) “영리적 목적”의 이용¹⁷⁹⁾

영리적 목적의 이용도 허용하는 경우에도 보상체계가 보장되지 않으면 권리자들에 대한 과도한 권리 침해라는 저작권자들의 반대의견이 있다.

비록 이인영의원안과 같이 교육·조사·연구 등 비상업적 목적의 데이터마이닝에 한하여 저작재산권을 제한하는 경우에도, 상업적 목적과 비상업적 목적을 구분하는 것이 쉽지 않고 비상업적 데이터마이닝의 결과물이 상업적으로 이용되는 것까지 배제된다고 보기 어려우므로, 데이터마이닝을 위해 아무런 대가 없이 저작물을 복제·전송할 수 있도록 하는 것은 헌법상 보장된 재산권의 과도한 제한이 된다는 반대의견이 있다.

나) “적법한 접근”

‘적법한 접근’의 요구와 관련하여 그 의미가 분명하지 않아 해석에 관한 분쟁이 초래될 수 있다는 점에서 신중한 검토가 필요하다는 견해가 있다.¹⁸⁰⁾

다) “저작물에 표현된 사상이나 감정을 향유하지 아니할 것”

정보분석을 위한 저작물 이용을 허락하는 요건 중 하나인 ‘저작물에 표현된 사상이나 감정을 향유하는 것을 목적으로 하지 아니할 것’의 의미가 추상적이어서 해석·적용이 모호하므로 충분한 검토가 필요하다는 의견이 있다.¹⁸¹⁾

179) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9.; 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이인영의원 대표발의) 검토보고」, 2023. 11.

180) 차상욱, 「저작권법상 인공지능 학습용 데이터셋의 보호와 쟁점」, 경영법률 제32권제1호, 2021

181) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9. 8쪽

라) 이용 방법 : 정보분석의 결과물(생성물)에 관한 문제

황보승희의원안은 ‘2차적저작물의 작성’을 이용방법으로 추가하여 정하고 있다. 여기서 이용방법 중 하나로 2차적저작물 작성을 추가한 취지가 데이터셋 작성을 위한 변형을 고려한 것으로 파악됨에 따라 기능적 작업에 불과한 데이터셋 작성 과정에서 새로운 창작성이 부가될 개연성이 있는지 단언할 수 없다는 점, 이 규정은 자칫 정보분석의 결과물로 원저작물과 유사한 생성물을 만들어내는 행위까지도 허용하는 것으로 오해될 여지가 있다는 점을 이유로 바람직하지 않다는 의견이 있다.¹⁸²⁾

한편, 문화체육관광위원회 검토보고서에서는 2차적저작물 작성 추가의 취지를 정보분석 결과물로서 2차적저작물을 만들어내는 행위를 포괄하려는 것으로 파악하고, 그러한 결과물이 2차적저작물로 성립할 수 있는지에 대한 의문 및 저작재산권자의 이익에 대한 과도한 제한 우려를 표시하고 있다.¹⁸³⁾

이용호의원안은 권리자의 허락 없이 정보분석한 결과물을 이용할 수 있는 예외 규정을 두고 있다.¹⁸⁴⁾ 이에 대해 정보분석의 결과물은 이용허락 없이 사용한 저작물과 경제적 대가를 지불한 저작물, 저작물 아닌 데이터 등을 종합적으로 분석한 결과일 수 있다는 점에서, 그 범위에 이용허락 없이 사용한 저작물이 포함된다는 사유만으로 「저작권법」에서 타법상 데이터에 관한 권리를 배제하는 것이 타당한지 여부에 대해서는 추가적인 검토가 필요해 보인다는 의견이 있다. 특허청은 저작물의 창작 목적을 이유로 「부정경쟁 방지 및 영업

182) 류시원, 위 논문 121쪽 각주33

183) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9. 8쪽

184) 저작재산권자의 권리를 제한하여 분석된 정보분석 결과물에 대해서는 비상업적 목적 또는 저작물의 창작 목적으로 적법하게 접근하는 경우 「부정경쟁방지 및 영업비밀보호에 관한 법률(‘21.12.7. 개정)」 등 데이터 보호에 관한 다른 법률의 규정에도 불구하고 이용이 가능하도록 규정함으로써 데이터 보호 관련 법률의 적용을 배제하고 있음.

비밀보호에 관한 법률」 관련 데이터 일반을 사용할 수 있게 하는 것은 해당 법률의 보호 법익과 상충된다는 의견을 제시했다.¹⁸⁵⁾

마) 정보분석 과정을 통해 생성된 복제물의 보관

황보승희의원안은 정보분석 과정을 통해 생성된 복제물을 필요한 범위 안에서 보관할 수 있도록 규정하고 있으나 복제방지조치 등은 의무화하고 있지 않다는 점에서, 정보분석에 활용되는 저작물 보관 과정에서의 복제물 유출 등 저작권 침해 문제를 해소하기 위한 방안을 검토할 필요성이 있다는 의견이 있다.¹⁸⁶⁾

복제물의 보관기간에 대해서도 독일의 사례를 참조하여 데이터마이닝 수행 후 해당 복제물을 삭제할 의무를 규정하는 것이 바람직하다는 의견도 있다.¹⁸⁷⁾

바) TDM 법안에 대한 관련 저작권자(협회) 의견¹⁸⁸⁾

TDM 면책 조항의 도입은, 저작물 이용 행위에 대해 현행법상 포괄적인 저작재산권 제한 규정인 공정이용 조항이 아니라 정보분석을 위한 저작물의 복제·전송을 명시적으로 허용하여 AI 산업을 진흥하고 법적안정성을 확보하고자 하는 취지가 있다.

반면, 저작권이 제한되는 입장에 선 저작권자들은 이에 관하여 대체로 반

185) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이용호의원 대표발의) 검토보고」, 2022.12. 13쪽

186) 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9.

187) 최상필, 「빅데이터의 분석과 활용을 위한 TDM의 저작권적 쟁점과 입법론」, 계간 저작권, 2022.3.

188) 황보승희의원안에 대한 2023. 6. 제출 의견 : 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의) 검토보고」, 2023. 9. 16~18쪽; 이인영의원안에 대한 2023. 11. 제출 의견 : 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이인영의원 대표발의) 검토보고」, 2023. 11. 19쪽을 반영하여 TDM 법안에 관한 저작권자(협회) 의견 [표 39] 작성하였다.

대의견을 제시하고 있으며, 영리적 이용(상업적 이용)과 비영리적 이용을 구분하고 영리적 이용을 하는 경우에는 적절한 보상체계가 보장되어야 한다는 입장이다(일각에서는 영리적 사용과 비영리적 사용을 구분하는 것 자체가 어렵다는 의견도 있다). 상업적 이용에 관하여 현재 업계에서는 저작물을 이용하고 사용료를 지급하는 계약을 체결하여 보상이 이루어지고 있는데, 오픈아웃 제도나 보상규정이 없는 TDM 규정이 도입될 시 저작권자들에 경제적 손해가 발생한다는 우려가 큰 것으로 보인다.

[표 16] TDM 법안에 관한 저작권자(협회) 의견

협회명 (의견제출일자)	의견 주요 내용
한국언론 진흥재단 (‘23.6.20.)	- 인공지능 학습 및 빅데이터 분석을 위하여 다량의 저작물을 사용하는 경우에 대해 저작권 침해 위법성조각사유를 규정하는 것은 정보 기업인 언론사의 저작재산권을 중대하게 침해할 수 있는 소지가 있음. 또한, 상업적 목적으로 2차적 가공물을 생산·판매하는 경우 위 개정 조항으로는 이를 구별하기가 어려울 것으로 판단 - 대부분의 해외 언론사는 뉴스 콘텐츠 유료화 정책을 통해 기사를 불법 다운로드받지 못하도록 하고 있으나, <u>국내의 경우 인터넷 포털에서 인쇄·텍스트복사·SNS·이메일 공유 등을 통해 얼마든지 기사를 복제할 수 있어 무분별한 뉴스 복제·전송이 난무할 수 있음</u>
한국문학예술저 작권협회 (‘23.6.21.)	해외 주요국가에서도 TDM 관련하여 서로 상이한 입법을 하는 등 많은 의견이 있는 바, 우리나라에서도 <u>이용목적이나 출처표시 등 관련한 다양한 쟁점에 대하여 이해관계자들간에 충분한 논의가 있는 연 후에 협의를 기초로 법안이 발의되기를 요청함</u>
한국방송 작가협회 (‘23.6.22.)	- 작가협회는 AI 데이터마ining 등 정보분석을 위해 협회 관리저작물을 이용하고자 하는 이용자와 저작물이용허락계약을 체결하고 사용료를 정산한 바 있으나, 그러한 이용허락절차 및 사용료 지급 없이 이용 가능케 하는 것은 저작권자의 권리를 지나치게 제한함(보상체계가 필요함) - 이용자가 이용하고자 하는 저작물과 이용목적 및 범위를 특정하고 저작권자와 충분히 협의하여 이용할 수 있음
한국음악 실연자연협회 (‘23.6.22.)	찬성

협회명 (의견제출일자)	의견 주요 내용
한국음악 저작권협회 (‘23.6.22.)	• 개정안은 헌법이 보호하는 창작자의 재산권을 조건 없이 희생하게 하는 것으로, 이 법이 시행되는 경우 창작산업계에 돌이킬 수 없는 피해가 발생할 수 있음 • 개정안이 허용하고 있는 저작권 제한의 범위가 지나치게 넓으며, 이와 같은 제한으로 발생할 수 있는 창작자의 피해를 보전할 수 있는 안전장치도 마련되어 있지 않음 - Opt-out 제도를 규정하고 있지도 않고, 우리나라에는 사적복제보상금제도 등도 존재하지 않고, 개정안은 이용목적에 제한이 없어, <u>인공지능 사업자가 인공지능의 상업적인 이용을 위해 저작물을 학습에 이용하는 경우에도 창작자는 적절한 보상을 받을 수 없음</u> - <u>2차적 저작물을 작성할 권한까지는 인공지능 학습과 관련해 반드시 허용되어야 하는 권한이라고 볼 수도 없음</u> - <u>적법한 접근 요건은 “적법한 이용허락을 받았을 것”을 전제로 해야 함</u> • 인공지능 학습데이터의 출처표시는 창작자의 권리행사와 향후 인공지능 관련 법제 마련의 필요성 등 연구를 진행하기 위해서 반드시 필요 - 현재까지 TDM 규정 등 도입 필요성에 대한 경제적 분석은 명확히 이루어진 바 없는데, 이는 인공지능 학습데이터의 출처 표시의무가 법제화되어 있지 않기 때문임 - 출처표시 의무화를 통해 향후 인공지능 산업계를 위해 어떠한 법안이 필요할지 구체적으로 검토할 수 있을 것이며, 창작자들 역시 인공지능 사업자가 사용한 저작물의 출처를 표시해줘야 저작물 이용 허락을 할 수 있음 • 종합하면, 본 개정안은 창작산업계에 심각한 피해를 가져올 수 있고, 저작권법의 입법 취지에도 부합하지 않으며, 지나치게 넓은 저작권 제한 범위를 갖고 있음에도, 창작자 보호를 위한 충분한 안전장치를 갖고 있지 않아 수정 필요 - 본 개정안의 적용 범위를 비영리적 목적 또는 연구 목적의 복제·전송으로 제한하고, 창작자들이 원할 경우 권리를 유보할 수 있도록 요청 - 상업적 목적의 이용과 관련하여서는 추후 공정이용에 관한 법리의 해석과 사적복제보상금제도의 도입, 인공지능 학습 자료의 출처 표시에 대한 충분한 논의가 이루어진 후에 관련 논의를 시작하여도 충분 - 또한, 인공지능 학습에 사용된 저작물의 출처 표시가 의무적으로 이루어질 수 있도록 관련 개정안 수정 필요

협회명 (의견제출일자)	의견 주요 내용
한국방송 실연자권리협회 (‘23.6.23.)	인공지능 기술 발전을 위하여 저작물의 이용이 허락되어야 할 필요성이 매우 크기는 하나, 아무런 대가 없이 저작물 이용이 허락될 경우 ‘인공지능 창작 알고리즘’의 생성물로 인하여 저작재산권자의 권리가 심각하게 침해될 것이므로 보상금 법정허락 제도가 보완되어야 할 것임
대한출판문화협회 (‘23.11.1)	저작물로 인정되는 텍스트 및 데이터 이용(또는 해당 개정안에 의해 산출된 결과물을 외부에 공개할 시)에는 반드시 저작권자의 사전동의를 받아야 한다는 조항의 삽입이 필요함
언론진흥재단 (‘23.11.1)	뉴스저작권자의 권리를 심각하게 침해할 수 있음
한국문학예술 저작권협회 (‘23.11.1)	- TDM 관련하여 연구나 논의가 진행중이므로 결과를 지켜보는 것이 필요함. 이미 제안 중인 다른 입법(안)(도종환 의원, 이용호 의원, 황보승희 의원 등 대표발의)과의 차이를 구분하기 어려움 - 해외 주요국가에서도 TDM 관련하여 서로 상이한 입법을 하는 등 많은 의견이 있고 논의가 진전되면서 기존 권리자 보호에 무게가 실리고 있음. 우리나라에서도 출처표시, 권리자 보상방안 등 권리자 보호에 대한 의견이 반영되기를 요청함
한국출판인회의 (‘23.11.1)	- 교육·조사·연구 등의 비상업적 목적으로 저작물을 마음껏 복제·전송할 수 있도록 하는 걸 무분별하게 법제화하려는 것은 저작권자와 출판권자에 대한 <u>최소한의 권리 보호나 사적 재산에 대한 보호</u>를 일절 고려하지 않은, 개정안의 시행에 따른 파급효과를 구체적으로 검토하지 않은 것으로 판단됨 - 비상업적 용도일지라도 저작물을 이용하는 데에는 <u>정당한 대가</u>가 필요하다는 기본 인식을 견지하고 그에 따른 출판권자와 저작권자에 대한 <u>정당한 보상 체계</u>를 마련해야 함. 이에 본 개정안은 정부 예산 지원 등이 뒷받침되어야 할 것임

5) 소결

우리나라 저작권법의 TDM 조항을 도입하기 위하여 일본 저작권법이 규정하고 있는 특유의 요건인 비향유적 이용 요건 및 유럽의 국가들이 채택하고 있는 적법한 접근 요건을 포함하여, 주체, 목적에 관한 제한 여부, 상업적 이용을 허용할 것인지 여부, 이용 방법을 제한할 것인지 여부, 결과물 보관을 허용할 것인지, 허용한다면 어느 범위까지 허용할 것인지, 옵트아웃 제도를 도입

할 것인지(관련하여 보상제도를 도입할 필요는 없는지) 등 외국 법제들의 체계를 참고하여 TDM 조항의 도입에 부수하는 여러 쟁점과 논의를 살피는 과정은 바람직할 것이다.

TDM 면책 조항의 도입은 AI 산업의 국제적 경쟁력을 제고하는 법적 근거가 될 것이나, 동시에 업계의 현실적 이해관계와 저작권 등 데이터를 보호하며 그 균형을 맞춰야 할 필요가 있다. 이를 위해서는 우리나라의 상황에 맞는 독자적인 규율 체계를 마련하는 방안도 고려되어야 한다.

(다) 주요국의 TDM 면책 조항

1) EU

2019년에 제정된 유럽연합 DSM 저작권 지침(이하 ‘DSM 지침’)은 TDM에 관한 규정을 담고 있다. DSM 지침 제2조 제2항은 TDM을 ‘패턴, 트렌드, 상관관계 등의 정보를 생성하기 위해 디지털 형태의 텍스트와 데이터를 분석하는 모든 자동화된 분석 기술’로 정의하고 있으며, 이와 관련하여 제3조 및 제4조의 이원적 규율 체계를 갖고 있다.

- 제3조: 연구기관 및 문화유산기관이 학술연구 목적의 TDM을 수행할 때 적법하게 접근 가능한 저작물에 대해 복제 및 추출을 허용함. 이 규정은 강행규정으로, 계약이나 기술적 보호수단으로 배제할 수 없음.
- 제4조: 모든 주체와 목적에 대해 TDM을 허용하되, 권리자가 옵트아웃할 수 있는 조건을 둠. 이는 제3조와 달리 강행규정이 아니며, 권리자는 이를 통해 TDM 수행을 제한할 수 있음.

유럽연합 DSM 저작권 지침 (2019년)¹⁸⁹⁾

제3조 과학적 연구 목적의 텍스트 및 데이터 마이닝

(1) 회원국은 연구기관과 문화유산기관이 과학적 연구 목적으로 그들이 합법적인 접근 권한을 가지는 저작물이나 그 밖의 보호대상을 텍스트 및 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, 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 및 데이터 마이닝을 위한 예외와 제한

(1) 회원국은 텍스트 및 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 및 데이터 마이닝 목적으로 필요한 한 보관될 수 있다.

(3) 제1항에 규정된 예외와 제한은 제1항에 언급된 저작물과 그 밖의 보호

189) 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.

대상의 이용이 권리자에 의해, 콘텐츠가 온라인으로 공중에게 이용 제공되는 경우에 기계가독형 수단 등, 적절한 방법으로 명시적으로 유보되지 않았다는 것을 조건으로 적용되어야 한다.

(4) 이 조항은 이 지침 제3조의 적용에 영향을 미쳐서는 안 된다.

DSM 지침의 이원적 구조는 다양한 이해관계를 조정하기 위한 것으로 평가되지만, 제3조의 적용 범위가 지나치게 협소하며, 제4조는 옵트아웃으로 인해 실효성이 제한된다는 비판이 제기되고 있다. 한편, 제3조와 제4조가 일정 범위 내에서 TDM 분석 결과물의 보유를 허용하기는 하나, TDM을 위한 제3자의 데이터셋 제공 등을 면책 대상으로 명시하지 않은 것은 TDM 면책 규정으로서의 의의를 다소 약화시키는 요소라고 볼 수 있다.¹⁹⁰⁾

2) 일본

일본은 2018년 저작권법 제30조의4를 도입하여, ① 저작물에 표현된 사상 또는 감정의 향수를 목적으로 하지 않으면서, ② 저작권자의 이익을 부당하게 침해하지 않는 경우에 한하여, ③ 필요하다고 인정되는 한도에서 해당 저작물을 이용할 수 있도록 허용하였다.

저작권법(2023. 6. 14. 개정)¹⁹¹⁾

제30조의4 (저작물에 표현된 사상 또는 감정의 향유를 목적으로 하지 아니하는 이용)

저작물은 다음의 경우, 그 밖의 해당 저작물에 표현된 사상 또는 감정을 직접 즐기거나 타인이 즐기게 하는 것을 목적으로 하지 아니하는 경우에는 필요하

190) 류시원, 「저작권법상 텍스트·데이터 마이닝(TDM) 면책규정 도입 방향의 검토」, 선진상사법률연구 제101호, 2023, 357~360쪽

191) 세계법령정보센터, 일본 저작권법(2023.05.26.공포, 2024.01.01.시행) 번역본

다고 인정되는 한도에서 어떠한 방법으로든 이용할 수 있다. 다만, 해당 저작물의 종류 및 용도, 해당 이용의 상황에 비추어 저작권자의 이익을 부당하게 침해하게 되는 경우에는 그러하지 아니하다.

1. 저작물의 녹음, 녹화, 그 밖의 이용과 관련된 기술의 개발 또는 실용화를 위한 시험에 제공하는 경우

2. 정보해석(다수의 저작물, 그 밖의 대량의 정보로부터 해당 정보를 구성하는 언어, 음, 영상, 그 밖의 요소와 관련된 정보를 추출, 비교, 분류, 그 밖의 해석을 하는 것을 말한다. 제47조의5제1항제2호에서 같다)에 제공하는 경우

3. 제1호 및 제2호의 경우 외에 저작물의 표현에 대한 사람의 지각에 의한 인식을 수반하지 아니하고 해당 저작물을 컴퓨터를 통한 정보처리과정에서의 이용, 그 밖의 이용(프로그램저작물의 경우에는 해당저작물을 컴퓨터에서 실행하는 것을 제외한다)에 제공하는 경우

제47조의5 (컴퓨터에 의한 정보처리 및 그 결과의 제공에 부수하는 경미한 이용 등)

① 컴퓨터를 이용한 정보처리를 통해 새로운 지식 또는 정보를 창출함으로써 저작물의 이용 촉진에 기여하는 다음 각 호의 행위를 하는 자(해당 행위의 일부를 하는 자를 포함하며, 해당 행위를 정령으로 정하는 기준에 따라 수행하는 자로 한정한다)는 공중에 대한 제공 등(공중에 대한 제공 또는 제시를 말하며, 송신가능화를 포함한다. 이하 같다)이 된 저작물(이하 이 조 및 제47조의6제2항제2호에서 “공중제공 등 저작물”이라 한다. 공표된 저작물 또는 송신가능화된 저작물로 한정한다)에 대하여 해당 각 호의 행위의 목적상 필요하다고 인정되는 한도에서 해당 행위에 부수하여 어떠한 방법으로든 이용(해당 공중제공 등 저작물 중 그 이용에 제공되는 부분이 차지하는 비율, 그 이용에 제공되는 부분의 양, 그 이용에 제공될 때의 표시의 정확도, 그 밖의 요소에 비추어 경미한 것으로 한정한다. 이하 이 조에서 “경미한 이용”이라 한다)을 할 수 있다. 다만, 해당 공중제공 등 저작물과 관련된 공중에 대한 제공 등이 저작권을 침해하는 것임(국외에서의 공중에 대한 제공 등의 경우에는 국내에서 했다면 저작권 침해가 되는 것임)을 알면서 해당 경미한 이용을 하는 경우, 그 밖에 해

당 공중제공 등 저작물의 종류 및 용도, 해당 경미한 이용상황에 비추어 저작권자의 이익을 부당하게 침해하게 되는 경우에는 그러하지 아니하다.

1. 컴퓨터를 이용하여 검색을 통해 찾는 정보(이하 이 호에서 “검색정보” 라 한다)가 기록된 저작물의 제호 또는 저작자명, 송신가능화된 검색정보와 관련된 송신처 식별부호(자동공중송신의 송신처를 식별하기 위한 문자, 번호, 기호, 그 밖의 부호를 말한다. 제113조제2항 및 제4항에서 같다), 그 밖의 검색정보의 특정 또는 소재에 관한 정보를 검색하고, 그 결과를 제공하는 행위

2. 컴퓨터를 통한 정보해석을 실시하고, 그 결과를 제공하는 행위

3. 제1호 및 제2호의 행위 외에 컴퓨터에 의한 정보처리를 통하여 새로운 지식 또는 정보를 창출하고, 그 결과를 제공하는 행위로서 국민생활의 편리성 향상에 기여하는 것으로 정령으로 정하는 행위

② 제1항 각 호의 행위의 준비를 하는 자(해당 행위의 준비를 위한 정보의 수집, 정리 및 제공을 정령으로 정하는 기준에 따라 수행하는 자로 한정한다)는 공중제공 등 저작물에 대하여 같은 항 규정에 따른 경미한 이용의 준비를 위하여 필요하다고 인정되는 한도에서 복제나 공중송신(자동공중송신의 경우에는 송신가능화를 포함한다. 이하 이 항 및 제47조의6제2항제2호에서 같다)을 하거나 그 복제물을 배포할 수 있다. 다만, 해당 공중송신 등 저작물의 종류 및 용도, 해당 복제 또는 배포의 부수 및 해당 복제, 공중송신 또는 배포 상황에 비추어 저작권자의 이익을 부당하게 침해하게 되는 경우에는 그러하지 아니하다.

일본의 TDM 면책 규정은 주요 규정인 제30조의4 외 제47조의5를 두어 이원적으로 규정되어 있다. 일본 저작권법 제30조의4와 제47조의5는 목적이나 주체에 따라 면책 범위를 구분하는 DSM 지침과 달리, 정보해석을 위한 비향수적 이용과 그에 부수하는 경미한 이용을 각각 허용하여 면책의 공백을 최소화하는 이원화를 설정하였다.

일본 저작권법은 형식상 공정이용과 같은 일반 면책조항은 아니지만, 제30조의4에서 매우 넓은 면책 범위를 제공하며, ‘적법한 접근’ 요구나 이용 방법에 제한이 없다. 이처럼 일본의 TDM 면책 규정은 주체, 목적, 이용 방법에

제한이 없고 저작권자에 대한 보상의무도 없어, TDM에 가장 우호적인 입법례로 평가된다.¹⁹²⁾¹⁹³⁾

3) 싱가포르 : ‘공정이용 + TDM’ 규정

싱가포르 정부는 디지털 시대의 권리 보호와 창작물 이용 환경 개선을 위해 2016년경부터 논의를 시작하였으며, 약 3년간의 집중적인 검토와 공개 의견 수렴을 거쳐 2019년 1월 17일에 TDM 개별 면책 규정 도입을 포함한 저작권법 개정 계획을 발표하였고, 이후, 2021년 9월에 싱가포르 의회를 통과한 개정 저작권법은 2021년 11월 21일부터 시행되었다.

싱가포르는 상업적, 비상업적 목적을 구분하지 않고, 옵트아웃도 허용하지 않는 강력한 규정을 이미 시행 중이다. 이 규정은 유럽연합 DSM 지침보다 더 넓은 면책 범위를 제공하고 있다. 싱가포르와 영국의 입법 사례는 ‘적법한 접근’ 요건을 유지하여 권리자 이익을 보호한다는 점에서 유럽연합 DSM 지침과 공통점을 가진다.

싱가포르 저작권법 제244조는 데이터 분석 과정에서의 복제뿐만 아니라, 복제물의 저장, 보유, 그리고 특정 목적이나 범위 내에서 제3자에게 제공하는 행위도 허용하고 있다.

또한, 복제물의 저장과 보유가 허용되고, TDM을 위한 복제 및 제공이 ‘발행’으로 간주되지 않으며, 이를 배제하거나 제한하는 계약 조건은 무효로 간주된다. 싱가포르 저작권법은 공정 이용(fair use) 규정도 포함하고 있지만, 상업적·비상업적 목적이나 주체의 제한이 없는 포괄적인 TDM 면책 조항을 도입한 것이 특징이다.¹⁹⁴⁾¹⁹⁵⁾

192) 박정훈·박다효, 「[EU, 일본] 유럽연합과 일본의 TDM 규정과 그 적용 범위 — AI 학습을 위한 저작물 이용과 관련하여」, 한국저작권위원회, 이슈리포트, 2024. 7.

193) 류시원, 위 논문, 368~369쪽

194) 이철남, 「싱가포르」 인공지능과 지식재산법의 교차에 관한 이슈 보고서 발표,

2021 저작권법(Copyright Act 2021)¹⁹⁶⁾

제2절 공정한 이용

제190조 허가된 이용과 같은 공정한 이용

- (1) 저작물의 공정한 이용은 저작물의 허가된 이용과 같다.
- (2) 다음 각 호의 어느 하나를 공정하게 이용하는 행위는 보호받는 실연의 허가된 이용과 같다.
- (a) 실연
 - (b) 실연의 녹화물

제191조 공정한 이용 여부 판단 시 고려 사항

제192조, 제193조 및 제194조를 전제로, 저작물이나 보호받는 실연(실연의 녹화물을 포함한다)이 공정하게 이용되었는지 판단할 때 다음 각 호를 포함하여 관련사항을 모두 고려하여야 한다.

- (a) 상업적 이용인지 비영리 교육 목적인지 여부를 포함한 이용 목적과 성격
- (b) 저작물 또는 실연의 성격
- (c) 저작물이나 실연 전체 대비 이용된 부분의 분량과 중요도
- (d) 해당 이용이 저작물이나 실연의 잠재 시장성 또는 가치에 미치는 영향

제192조 특정 목적 이용 시 적절한 출처 표시 추가 요건

(생략)

제193조 공정하게 이용된 저작물 내 저작물 또는 녹화물의 공정 이용 간주

(생략)

제194조 연구 및 조사를 위한 합리적 분량 복제의 공정 이용 간주

(생략)

한국저작권위원회, 저작권 동향, 2024.

195) 류시원, 위 논문, 365~366쪽

196) 세계법령정보센터, 싱가포르 2021 저작권법(2022.10.25.개정) 번역본

제8절 컴퓨터 데이터 분석

제243조 해석: 컴퓨터 데이터 분석

이 절에서 저작물이나 보호받는 실연의 녹화물과 관련하여 "컴퓨터 데이터 분석"이란 다음 각 호를 포함한다.

(a) 컴퓨터프로그램을 이용하여 저작물이나 녹화물로부터 정보 또는 데이터를 식별하고 추출하여 분석하는 행위

(b) 저작물이나 녹화물을 정보 또는 데이터 유형의 예시로 삼아 해당 유형의 정보 또는 데이터 관련 컴퓨터프로그램의 기능을 향상시키는 행위

예시

제(b)호에 따른 컴퓨터 데이터 분석의 예로 화상을 이용하여 컴퓨터프로그램을 학습시켜 화상을 인식할 수 있도록 하는 경우를 들 수 있다.

제244조 컴퓨터 데이터 분석을 위한 복제 또는 공중송신

(1) 제(2)항의 조건이 충족되면 이용자("갑")가 다음 각 호 중 어느 하나를 복제하는 행위는 허가된 이용에 해당한다.

(a) 저작물

(b) 보호받는 실연의 녹화물

(2) 조건은 다음 각 호와 같다.

(a) 복제본 제작 목적이 다음 각 목의 어느 하나일 것

(i) 컴퓨터 데이터 분석

(ii) 컴퓨터 데이터 분석을 위한 저작물 또는 녹화물 준비

(b) 갑은 복제물을 그 밖에 다른 목적으로 사용하지 아니할 것

(c) 갑은 다음 각 목을 위한 경우를 제외하고 복제물을 다른 자에게 유출(공중송신 등 방법을 불문한다)하지 아니할 것

(i) 갑이 수행한 컴퓨터 데이터 분석 결과 검증

(ii) 갑이 수행한 컴퓨터 데이터 분석 목적과 관련한 합동 연구 또는 조사

(d) 이용자는 복제 대상 자료(이 조에서"1차 복제물"로 부른다)에 적법한 접근 자격이 있을 것

예시

㉔ 갑이 유료화 벽을 우회하여 1차 복제물에 접근한 경우 갑은 1차 복제물에 적법한 접근 자격이 없다.

㉕ 갑이 데이터베이스 사용 조건을 위반 (제187조에 따라 무효가 되는 조건은 제외한다)하여 1차 복제물에 접근한 경우, 갑은 1차 복제물에 적법한 접근 자격이 없다.

(e) 다음 각 목의 어느 하나를 충족할 것

(i) 1차 복제물이 저작권 불법 복제물이 아닐 것

(ii) 1차 복제물이 저작권 불법 복제물인 경우, 다음에 해당할 것

(A) 갑이 저작권 침해 사실을 알지 못하는 경우

(B) 1차 복제물을 악성 불법 온라인 사이트(사이트가 제325조에 따른 접근 차단 명령 대상인지 여부를 불문한다)에서 얻은 경우, 갑이 그 사실을 알지 못하였거나 합리적으로 알 수 없었던 경우

(iii) 1차 복제물이 저작권 불법 복제물인 경우, 다음에 해당할 것

(A) 규정된 목적을 위하여 저작권 불법복제물 사용이 불가피한 경우

(B) 갑이 그 밖에 다른 목적의 컴퓨터 데이터 분석 수행을 위하여 해당 복제물을 사용하지 아니한 경우

(3) 명확하게 정리하면, 제(1)항의 복제물 제작은 복제물 저장 또는 보유를 포함한다.

(4) 다음 각 호의 조건이 충족되면 갑이 저작물이나 보호받는 실연의 녹화물을 공중송신하는 행위는 허가된 이용에 해당한다.

(a) 제(1)항이 적용되는 상황에서 제작된 복제물을 공중송신 할 것

(b) 갑은 다음 각 목의 어느 하나를 위한 경우를 제외하고 다른 자에게 복제물을 유출(공중송신 등 방법을 불문한다)하지 아니할 것

(i) 갑이 수행한 컴퓨터 데이터 분석 결과 검증

(ii) 갑이 수행한 컴퓨터 데이터 분석 목적과 관련한 합동 연구 또는 조사

(5) 이 법의 목적상, 이 조가 적용되는 상황에서 자료의 복제물을 제공하는 행위는 다음 각 호에 해당한다.

(a) 자료(또는 자료에 포함된 저작물 또는 녹화물) 공표에 해당하지 아니함

(b) 자료(또는 자료에 포함된 저작물)의 저작권 존속 기간 결정 시 고려하지 아니함

4) 영국 : ‘공정거래(Fair Dealing) + TDM’

영국은 2014년 6월 1일에 TDM과 관련한 개별 면책조항인 CDPA 제29A조를 도입하였다.

CDPA 제29A조는 TDM과 관련하여 비상업적 연구 목적으로 컴퓨터 분석을 위해 저작물을 복제하는 행위를 저작권 침해로 보지 않도록 규정하고 있다. 이 조항에 따라, 저작물에 대해 적법한 접근 권한을 가진 자가 비상업적 연구를 목적으로 복제를 수행할 때, 복제물에 충분한 출처를 명시하는 한도 내에서 저작권 침해로 간주되지 않는다. 그러나 이 복제물을 저작권자의 허락 없이 타인에게 양도하거나 다른 용도로 사용하는 경우에는 면책되지 않는다는 한계가 있다. 또한, 이 조항에 의해 허용된 복제를 금지하거나 제한하는 계약 조항은 무효로 간주하는 조항을 두고 있는데 이를 통해 해당 조항이 강행규정임을 분명히 하고 있다.

영국 저작권법(Copyright, Design and Patent Act 1988 Section 29A) ¹⁹⁷⁾
제29A조(비영리적인 연구를 위하여 본문과 데이터를 분석하기 위한 복제물) (1) 저작물에 합법적으로 접근할 수 있는 사람이 다음과 같은 저작물의 복제물을 만드는 경우, 이는 해당 저작물의 저작권을 침해하는 것으로 간주하지 아니한다. (a) 저작물에 합법적으로 접근할 수 있는 사람이 비영리적인 목적의 연구를 위하여 컴퓨터를 사용하여 저작물에 기록된 내용을 분석할 수 있도록 복제물을 제작한다. (b) 실행 가능성이나 다른 이유로 인하여 복제물을 확인할 수 있는 경우, 복제물을 충분히 확인한다. (2) 이 조에 따라 저작물의 복제물을 만들었고 저작물의 복제물이 다음 중 어느 하나에 해당하는 경우, 이는 저작물의 저작권을 침해하는 것이다. (a) 복제물을 다른 사람에게 양도한다. 저작권자가 해당 양도를 승인하는 경우는 제외한다.

197) 세계법령정보센터, 영국 저작권, 의장 및 특허에 관한 법률(1988) 번역본

- (b) 제(1)항제(a)호에서 언급한 목적 이외의 목적을 위하여 복제물을 사용한
다. 저작권자가 해당 사용을 승인하는 경우는 제외한다.
- (3) 이 조에 따라 만든 복제물을 추후에 처리하는 경우
- (a) 해당 취급과 관련하여 해당 복제물을 침해 복제물로 취급한다.
- (b) 해당 취급이 저작권을 침해하는 경우, 추후의 모든 목적과 관련하여 해당
복제물을 침해 복제물로 취급한다.
- (4) 제(3)항에서 "처리"는 업무상 판매 또는 임대하거나 판매나 업무를 위하여
제안 또는 노출시키는 것을 말한다.
- (5) 저작권을 침해하지 아니하면서 이 조를 통하여 복제물을 만드는 것을 계약
조건이 방지하거나 제한하는 경우, 해당 계약 조건은 집행할 수 없다.

이러한 CDPA 제29A조는 영국에서 TDM 면책 규정을 명시적으로 도입한
적극적인 입법 사례로 평가된다. 최근에는 비영리적 목적의 연구뿐만 아니라
영리적 목적을 포함하여 목적을 불문한 TDM 면책 규정을 도입하려는 논의가
계속되고 있다.¹⁹⁸⁾¹⁹⁹⁾

5) 독일

독일은 2017년 제한된 범위의 TDM 면책 규정을 자국 저작권법제에 도입
하였으며, 유럽연합의 DSM 지침이 제정되면서 지침에 따라 국내 법제를 이에
맞춰 개정하였다.

기존의 ‘비상업적 목적’이라는 요구 사항이 주체 및 목적의 제한으로 대
체되었으며 복제 행위 자체에 대해서는 더 이상 비상업적 목적을 요구하지 않
게 되었다. 단, 특정 목적을 위해 복제물을 공공에게 제공할 경우에는 여전히
‘비상업적 목적’을 요구하는데, 이 경우 그 목적이 완료되면 지체 없이 종료되

198) 류시원, 위 논문, 363~365쪽

199) 유혜정, 「영국 지식재산청, 텍스트와 데이터마이닝(TDM) 관련 저작권법 개정
추진」, 한국저작권위원회, 저작권 속보, 2022. 7. 6.

어야 한다는 것이 제60d조 4항에 명시되어 있다.

또한 개정법에서는 학술연구 목적으로 이루어지는 TDM 과정에서의 복제에 대해서는 보상금 지급 의무가 면제되었다. 이는 제60h조 2항 3호에서 다루고 있으며, 학술연구 활동을 촉진하고 연구자들에게 부담을 덜어주는 방향으로 규정이 변경된 것으로 보인다. 그러나 이 경우에도 권리자의 옴트아웃이 가능하므로 라이선스를 통한 경제적 보상 수취 기회는 잔존한다.²⁰⁰⁾

독일 저작권법(Urheberrechtsgesetz)²⁰¹⁾

제44b조(텍스트 및 데이터 마이닝)

- (1) 텍스트 및 데이터 마이닝은 특히 패턴, 추세 및 상관관계에 대한 정보를 획득하기 위해 하나 이상의 디지털 또는 디지털화된 저작물을 자동으로 분석하는 것이다.
- (2) 텍스트 및 데이터 마이닝을 위해 합법적으로 접근할 수 있는 저작물의 복제가 허용된다. 복제품은 텍스트 및 데이터 마이닝에 더 이상 필요하지 않으면 삭제된다.
- (3) 제2항제1문에 따른 사용은 권리보유자가 유보하지 않은 경우에만 허용된다. 온라인으로 접근할 수 있는 저작물에 대한 사용 유보는 기계 판독이 가능한 형태인 경우에만 유효하다.

제60d조(학술 연구 목적의 텍스트 및 데이터마이닝)

- (1) 텍스트 및 데이터 마이닝을 위한 복제(제44b조제1항과 제2항제1문)는 다음 조항에 따라 학술 연구 목적으로 허용된다.
- (2) 연구기관은 복제할 권리가 있다. 다음 중 어느 하나에 해당하는 경우 학술 연구를 수행하는 대학, 연구소 또는 그 밖의 기관은 연구기관에 해당한다.
 1. 비상업적인 목적을 추구하는 경우
 2. 모든 수익을 학술 연구에 재투자하는 경우

200) 류시원, 위 논문, 361쪽

201) 세계법령정보센터, 독일_저작권법_번역본(국회도서관)

3. 국가 공인 계약의 범위 내에서 공익을 위해 활동하는 경우 연구기관에 특정한 영향을 미치고 학술 연구결과에 우선적으로 접근할 수 있는 민간 기업과 협력하는 연구기관은 제1문에 따른 권리가 없다.
- (3) 또한 다음의 경우 복제할 권리가 있다.
 1. 일반 대중의 접근이 허용된 도서관과 박물관, 기록보존소 및 시청각유산 부문의 시설(문화유산기관)
 2. 상업적 목적을 추구하지 않는 개별 연구원
- (4) 비상업적 목적을 추구하는 제2항과 제3항에 따라 권리를 부여받은 사람은 다음의 사람에게 제1항에 따른 복제에 공개적으로 접근할 수 있도록 할 수 있다.
 1. 공동 학술 연구를 위해 특정적으로 분류된 사람들의 집단
 2. 학술 연구의 질을 검증하기 위한 개별 제3자

공동 학술 연구 또는 학술 연구 질 검증이 완료되는 즉시 일반 대중에게 접근이 가능하도록 제공하는 것은 종료되어야 한다.
- (5) 제2항과 제3항제1호에 따라 권한을 부여받은 사람은 학술적 연구 또는 학술적 발견의 검증을 위해 필요한 경우 무단 사용을 방지하는 적절한 보안 예방 조치와 함께 제1항에 따라 복제물을 보관할 수 있다.
- (6) 권리보유자는 제1항에 따라 복제물에 의해 네트워크와 데이터베이스의 보안 및 무결성이 저해되는 것을 방지하도록 필요한 조치를 취할 권한이 있다.

(라) TDM 관련 주요 해외 입법 및 정책 동향

1) 영국

영국 지식재산청(The Intellectual Property Office)은 2022년 6월 28일에 텍스트 및 데이터마이닝(TDM)을 통한 저작물 이용 시, 영리·비영리 목적을 불문하고 이를 허용하는 저작권법 개정 의사를 저작권법 조항은 개정하겠다고 밝혔으나, 영국 음악 및 관련 산업계의 반대에 부딪혔다.

2023년 2월 1일에 ‘인공지능이 창작자의 지식재산권에 미치는 잠재적 영

향'이라는 주제로 이루어진 토론에서, 조지 프리먼(George Freeman) 부장관은 영국 지식재산청이 제안한 '모든 목적'의 텍스트 및 데이터마이닝을 허용하는 저작권법 개정을 추진하지 않겠다는 의사를 밝혔으며, 현재는 자율규제방식이 주도적이다. 다만, 모든 목적의 TDM을 허용하는 저작권법 개정 계획안의 취소 결정은 중국적인 것이 아니며, 권리와 책임 그리고 보상에 관한 균형을 맞추기 위한 잠정 중단의 상태로 보인다.²⁰²⁾²⁰³⁾

2) 호주

호주 법무부는 저작권과 인공지능을 주제로 라운드테이블 회의를 개최하였고, 2023년 8월 28일부터 29일까지 열린 제3차 회의에서 처음으로 인공지능과 생성형 인공지능에 대한 저작권 문제를 논의하였다. 이 회의에는 저작권자와 이용자 대표 47개 단체가 참석하였으며, 인공지능 기술과 저작권 문제에 대한 의견을 공유하였다. 또한, 호주 법무부는 '인공지능 및 저작권 전문가 자문 그룹(AI and Copyright Expert Reference Group)'을 운영하여 향후 인공지능 문제를 지속적으로 논의할 계획을 밝혔다.

호주 산업과학자원부는 인공지능 규제에 대한 논의를 주도하고 있으며, 2023년에 '호주의 안전하고 책임감 있는 인공지능(Safe and Responsible AI in Australia)'을 위한 의견 수렴을 실시하였다. 다수의 의견이 저작권 문제에 집중되었고, 현재 산업과학자원부는 EU 인공지능법과 캐나다의 AIDA에서 제안된 위험 기반 접근 방식을 고려하여 인공지능 기술의 책임 있는 활용을 위한 조치를 마련하고 있다.

또한, 뉴 사우스 웨일즈(NSW) 의회는 2023년 6월 27일에 '인공지능에 대

202) 유혜정, 위 보고서

203) 구문모, 「[영국] 영국 정부, '모든 목적'의 텍스트 및 데이터마이닝 허용 저작권법 개정 계획 철회 결정」, 한국저작권위원회, 저작권 동향, 2023. 2.

한 의회 조사위원회’를 구성하여 관련 논의를 진행 중이다. 호주 정부기관은 인공지능 저작권 문제를 논의하고 있지만, 저작권법의 큰 변화보다는 투명성 의무 강화와 같은 규정을 도입할 가능성이 높다.²⁰⁴⁾

3) 프랑스

2023년 9월 12일에 프랑스 하원의 Guillaume VUILLET 의원 등 8명의 의원이 ‘저작권을 통한 인공지능 규제’를 목표로 지식재산권법(Code de la propriété intellectuelle)을 개정하는 법률안을 제출하였다. 법안의 주요 내용은 AI가 학습데이터를 이용하는 행위는 저작권자의 권리에 해당하고, AI 산출물에 대한 권리자를 정하고 이러한 권리는 집중관리단체에 의하여 관리되며, AI 산출물임을 표시하도록 하고, 출처불명인 저작물을 학습데이터로 사용하여 AI 산출물을 만들어 내는 AI 시스템 운영자에 대하여 세금을 부과하고 이러한 세금이 창작을 장려하고 집중관리단체의 이익이 되도록 하는 것이다.²⁰⁵⁾

4) 캐나다

캐나다 연방정부는 2023년 10월 12일에 혁신과학경제개발부(ISED)가 주도하여 'AI와 사물 인터넷(IoT)을 위한 최신 저작권 프레임워크(A Consultation on a Modern Copyright Framework for Artificial Intelligence and the Internet of Things)'에 관한 의견수렴을 시작하여 2024년 1월 25일에 종료되었으며, ISED는 올해 안에 이와 관련된 최종보고서를 공개할 예정이다.²⁰⁶⁾

204) 송선미, 「[호주] AI 저작권 관련 호주 정부의 대응 현황」, 한국저작권위원회, 저작권 동향, 2024. 2.

205) 이대희, 「프랑스, ‘AI와 저작권 관련 지식재산권법’ 개정안 발의」, 한국저작권위원회, 저작권 동향, 2023. 12.

206) 캐나다 정부 홈페이지, “Government of Canada launches consultation on the implications of generative artificial intelligence for copyright” (2024. 6. 19.)
<https://www.canada.ca/en/innovation-science-economic-development/news/2023/10/g>

현재까지는 별도의 결과 보고서가 발표되지 않았지만, 이해관계자들이 제출한 일부 내용이 공개되었다. AI와 관련된 저작권 문제, 특히 텍스트 및 데이터 마이닝(TDM), AI 생성물의 저작권 소유, 그리고 AI 시스템의 저작권 침해 문제 등이 대표적이다.²⁰⁷⁾

5) 미국

가) 정책 논의

미국 저작권청은 2023년 3월에 인공지능 생성물의 저작권 등록 가이드라인을 발간하고, 2023년 4월과 5월에 생성형 인공지능의 저작권 이슈에 대한 의견 수집을 위한 공청회를 열었다. 또한, 2023년 6월과 7월에는 약 2천 명이 참석한 생성형 인공지능 관련 웨비나를 진행하여, 인공지능 시대의 저작권 정책 수립을 위한 기초자료를 축적하고 있다. 이어서, 2023년 8월 30일에는 생성형 인공지능과 관련된 현황 조사 및 이해관계자 의견 청취를 위한 공개 의견 수렴 절차를 개시하였다.

이 공개 의견 수렴 절차에서는 저작권청이 34개의 주요 질문과 하위 질문들을 포함한 광범위한 질의 목록을 제시하였다. 주요 질문들은 인공지능 생성물의 저작권 보호 여부 및 방식, 퍼블리시티권 보호 문제, 그리고 특정 업계나 영역에서 발생하는 고유한 쟁점들이다. 이 과정에서 제출된 의견은 두 차례 연장되어 2023년 12월 6일까지 진행되었으며, 총 1만 건 이상의 의견이 연방 규정 포털에 공개되었다.²⁰⁸⁾ Google, Meta Platforms, Stability AI와 같은 AI

overnment-of-canada-launches-consultation-on-the-implications-of-generative-artificial-intelligence-for-copyright.html

207) 차상욱, 「캐나다, 생성형 AI와 저작권에 대한 협의 관련 의견수렴 내용 공개」, 한국저작권위원회, 저작권 동향, 2024. 8.

208) 미국 저작권 사무소(U.S. Copyright Office) 홈페이지
<https://www.copyright.gov/policy/artificial-intelligence/>

기술기업과 Authors Guild, creative 등의 단체가 상반된 의견을 제출하였고, 이를 뒷받침하는 연구 자료와 판례들도 함께 제공되었다.²⁰⁹⁾

나) 법안 발의²¹⁰⁾

2024년 7월 11일에 Cantwell 의원 등 3명의 미국 연방상원의원은 ‘COPIED Act’라는 명칭의 법안을 제출하였다. 이 법안은 콘텐츠의 이력(provenance, 출처) 및 인공지능(AI) 산출물의 파악, AI 시스템 이용자가 콘텐츠 이력을 포함시킬 수 있도록 하는 것 등을 내용으로 하고 있다.

첫째, 연방기관으로 하여금 AI 생성 콘텐츠 파악, 워터마크, 콘텐츠 이력 정보 등에 대하여 연구·개발하거나 표준을 제정할 의무를 부과하는 것을 내용으로 하고 있다(§5).

둘째, 생성형 인공지능 툴(tool, AI 시스템이나 앱) 제공자에게 사용자가 AI 산출물에 대한(연방기관에 의하여 개발되거나 표준이 제정된) ‘콘텐츠 이력 정보’ 표시 의무를 부여하고 있다. 이러한 의무는 AI 산출물에 대해서뿐만 아니라 AI 툴에 의하여 생성되거나 상당히 수정된 디지털 형태의 콘텐츠에 대해서도 적용된다(§6). 콘텐츠 이력정보(content provenance information)는 디지털 콘텐츠의 출처나 이력(history)을 기록한 정보로, ‘최신의 기술에 의한, 기계가 읽을 수 있는 정보’로 정의된다(§2).

셋째, 누구든지 불공정하거나 기망행위를 위하여 고의로 콘텐츠 이력정보를 제거, 수정, 변경, 무력화하는 것은 금지된다(§6(b)(1)). 플랫폼도 이력 정보를 제거 등을 함으로써 이용자가 이러한 정보에 접근할 수 없도록 하는 것도 금지되며, 다만 보안연구를 위하여서만 허용된다(§6(b)(2)). 저작권 관리정보의

209) 류시원, 「미국」美 저작권청, 인공지능 산출 이미지 ‘SURYAST’ 등록거절, 한국저작권위원회, 이슈리포트, 2024. 3.

210) 이대희, 「AI 산출물 식별 및 이력정보에 관한 미국의 법안」, 한국저작권위원회, 저작권 동향 2024. 9.

제거 등이 금지되는 것과 같이 이력정보의 제거 등도 금지된다.

넷째, 콘텐츠 이력 정보를 가진 저작물이나 이력 정보가 제거·분리되었다는 것을 알거나 알 수 있었던 저작물을, AI 또는 알고리즘을 사용하는 시스템을 훈련시키거나 AI 산출물을 생성·수정하기 위하여, 사용하는 것은 금지된다. 이에 해당하지 않기 위하여서는 해당 저작물을 소유하는 자의 동의를 얻거나 보수 지급 등 해당 저작물 이용에 관한 조건을 준수해야 한다(§6(c)).

다섯째, COPIED Act의 집행은 불공정하거나 기망적 행위를 규율하는 연방거래위원회(Federal Trade Commission, FTC)나 주 검찰총장에 의하여 이루어진다(§7).

위 법안은 콘텐츠의 출처나 이력 정보를 기록하도록 의무화하여 저작물이 무단으로 이용되는 것을 예방하는 효과가 있을 것으로 평가되는데, 이러한 법안이 도입될 경우 허위조작정보를 색출하는 효과도 부가적으로 있을 것으로 예상된다.

6) 일본

일본 정부는 인공지능이 제기하는 문제들을 검토하기 위해 ‘AI 전략회의(이노베이션 정책 강화 추진을 위한 지식인 회의)’를 조직하고, 2023년 5월부터 11월까지 6차례 회의를 진행하였다. 특히, 2023년 9월 8일 5차 회의에서 생성형 인공지능의 지적재산권 리스크를 검토할 별도 위원회인 ‘AI 시대의 지적재산권 검토회’가 출범하였다. 이 검토회는 2023년 10월 4일에 첫 회의를 개최하고, 2024년 1월까지 총 5회의 회의를 진행했으며, 2024년 4~5월까지 7차례 회의를 예정하고 있다. 검토회는 생성형 인공지능과 관련된 지적재산권 문제를 다루며, 법적 과제와 대응책을 종합적으로 검토하고 있다.

회의 내용은 공개되며,²¹¹⁾²¹²⁾ 정부 각 부처와 민간 전문가들이 참여한다.

211) AI전략회의(AI戰略)

<https://www8.cao.go.jp/cstp/ai/index.html>

검토회는 저작권과 AI 학습, 생성물 이용에 관한 저작권 침해 사례, AI 관련 기술 대응책, 수익 환원책 등 다양한 주제를 포괄한다. 2024년 들어 생성형 인공지능과 저작권 논점이 집중적으로 논의되고 있으며, 회의 자료와 의견 수렴 결과가 모두 공개되어 일반 공중이 참여할 수 있는 기회를 제공한다. 이러한 작업은 인공지능 관련 지적재산권 정책이 구체적 사실에 기반해 논의되도록 하며, 정책의 일관성을 높이는 데 기여할 것으로 기대된다.

(4) 주요국 법제 비교²¹³⁾²¹⁴⁾

지금까지 살펴본 유럽연합, 독일, 싱가포르, 일본, 영국, 미국의 저작권 관련 법제를 종합적으로 비교해 보면 다음과 같다.

212) AI검토회(AI時代の知的財産権検討会)

https://www.kantei.go.jp/jp/singi/titeki2/ai_kentoukai/kaisai/index.html

213) 류시원, 위 논문, 370쪽

214) Trina Ha et al, 「When Code Creates: A Landscape Report on Issues at the Intersection of Artificial Intelligence and Intellectual Property Law」, Singapore Management University, 2024. 2. 28.

[표 17] 주요국 저작권 법제 현황

구분	유럽연합		독일		싱가포르	일본	영국	미국
	제3조	제4조	제60d조	제44b조	제244조	제30조의4	제29A조	제107조
시행일	2021. 6. 7.		2021. 6. 7.		2021. 11. 21.	2019. 1. 1.	2014. 6. 1.	1978. 1. 1.
상업적 이용	불가	가능	불가	가능	가능	가능	불가	가능
주체	연구기관 문화예술 기관	기타 (영리 기업 등)	연구기관 문화예술 기관	기타 (영리 기업 등)	-	-	-	-
특유 요건	적법한 접근		적법한 접근		적법한 접근	저작물에 표현된 사상이나 감정의 향유(향수)를 하지 않을 목적	적법한 접근, 충분한 사사표시 (acknowledge ment)	-
이용 방법	복제, 전송 ²¹⁵⁾	복제, 전송 ²¹⁶⁾	복제, 전송 ²¹⁷⁾	복제, 전송 ²¹⁸⁾	복제·준비 작업, 결과 검증·합동연 구 목적 송신	모든 이용방법 (공중에 양도 가능)	복제	모든 이용방법
결과물 보관 (삭제의무)	보안을 갖춰 기간 제한 없이 저장·유지 가능	TDM 목적의 필요 한도 내에서 유지 가능	보안을 갖춰 기간 제한 없이 저장· 유지 가능	TDM에 더 이상 필요하지 않은 경우 삭제	-	-	-	-
옵트 아웃제도	X	O	X	X ²¹⁹⁾	X	X	X	X
계약에 의한 배제 가능성	X	O	-	-	X	-	X	-

- 215) 저작권자인접권자편집저작물(DB) 저작자의 복제권, 데이터베이스제작자-언론 간행물 발행자의 복제권, 전송권 제한
- 216) 저작권자인접권자편집저작물(DB) 저작자-컴퓨터프로그램 저작자의 복제권. 데이터베이스제작자-언론간행물 발행자의 복제권, 전송권 제한
- 217) 저작권자의 복제권, 데이터베이스제작자의 복제권 제한
- 218) 저작권자의 복제권, 컴퓨터프로그램저작자의 2차적저작물작성권, 데이터베이스 제작자의 복제권 제한
- 219) 독일 저작권법상 명시적인 옵트아웃제도를 도입하고 있지 않으나 AI 학습에 이용된 오픈데이터셋을 둘러싼 사건인 ‘Robert Kneschke v. LAION e.V.’ 사건에서 DSM 제5조 제1항과 이를 이행한 독일 저작권법 제44b조에 의해 옵트아웃

TDM 면책 관련 법제를 마련한 주요국들은 TDM 생태계 내 다양한 이해관계를 조정하기 위해 각기 다른 접근 방식을 취하고 있다. EU의 DSM 지침은 학술연구 목적의 TDM과 일반적인 TDM을 구분하여 이원적으로 규율하고 있으며, 학술연구 목적의 TDM에 대해서는 계약이나 기술적 수단으로 면책을 배제하지 못하도록 강제하는 반면, 일반 TDM에 대해서는 오픈아웃을 허용함으로써 균형을 맞추고 있다. 이러한 이원적 접근은 TDM 이용 환경이 다른 회원국들의 상반된 이해관계를 반영한 결과로, 일정한 제한이 있음에도 불구하고 TDM 활동을 면책 대상으로 포함하려는 시도를 반영한 것으로 볼 수 있다.

[표 18] 주요국 TDM 정의

국가	TDM 정의
유럽 연합	패턴, 트렌드, 상관관계 등의 정보 생산을 위해 디지털 형태 텍스트와 데이터를 분석하는 것을 목적으로 하는 자동화 기술
독일	특정 패턴, 추세 및 상관관계에 대한 정보를 얻기 위해 하나 이상의 디지털 또는 디지털화된 작업으로 구성된 자동화된 분석
싱가포르	컴퓨터프로그램을 이용하여 저작물이나 녹화물로부터 정보 또는 데이터를 식별하고 추출하여 분석하는 행위, 저작물이나 녹화물을 정보 또는 데이터 유형의 예시로 삼아 해당 유형의 정보 또는 데이터 관련 컴퓨터프로그램의 기능을 향상시키는 행위
일본	다수 저작물 등 대량의 정보로부터 해당 정보를 구성하는 언어, 소리, 영상 등의 요소에 관한 정보를 추출, 비교, 분류 등 분석하는 것(저작물에 표현된 사상 또는 감정을 스스로 향유하거나 다른 사람에게 향유하게 하는 것을 목적으로 하지 않는 경우)

제가 적용됨을 전제하여 재판이 진행되고 있다[이일호, 「[독일] AI 학습에 이용된 오픈 데이터셋을 둘러싼 사건」, 한국저작권위원회, 저작권 동향(2024) 참조].

영국과 싱가포르는 상업적/비상업적 목적의 구분 없이 옵트아웃을 허용하지 않는 강력한 규정을 도입함으로써 EU보다 더 넓은 범위의 면책을 제공하고 있다. 특히, 싱가포르의 입법례는 ‘적법한 접근’ 요건을 유지하면서도 강행 규정을 적용함으로써, TDM을 위한 복제물의 보관 및 제3자 제공을 허용하는 등 포괄적인 면책 규정을 마련하였다.

일본은 DSM 지침과 유사하게 TDM 면책 규정을 이원화하여 정보해석 그 자체를 위한 비향수적 이용과, 정보해석에 부수하는 경미한 이용을 구분하고 있다. 일본 저작권법 제30조의4는 비향수적 이용을, 제47조의5는 경미한 이용을 각각 면책 대상으로 규정하며, TDM 활동 전반에 걸쳐 실효성을 높이기 위한 규정이라고 볼 수 있다.

반면, 미국은 개별 면책 규정을 마련하지 않고 공정이용(fair use) 조항을 통해 유연하고 탄력적으로 TDM 관련 저작물 이용을 허용한다. 이용 방법이나 주체, 목적에 관한 명시적 제한이 없고, 적법한 접근 요건이나 보상 의무도 부과되지 않는다. 다만, TDM의 변형적 이용 가능성을 인정할 경우 상업적 목적의 유무가 공정이용 판단에 크게 영향을 미치지 않을 것으로 보인다.

마. 학습데이터 이용 시 저작권자 보상

(1) 개요

저작권법은 저작재산권을 제한하지만, 저작재산권자에게 보상금을 지급하도록 하는 규정들을 두고 있다. 저작권법 제25조(학교교육 목적 등의 이용), 제31조(도서관등에서의 복제 등), 제35조의4(문화시설에 의한 복제 등)가 그에 해당한다.

생성형 인공지능의 저작물 학습에 대하여 저작재산권이 제한되더라도, 위 조항들과 유사하게 저작재산권자에게 보상금을 지급하도록 규정하는 방안을

고려할 수 있다. 다만, 이에 대하여는 이해관계자들 간 깊이 있는 논의가 필요하다.

(2) 보상 관련 해외사례

유럽연합의 경우, DSM 지침의 도입에 따라 별도의 보상의무 대신 옵트아웃제도를 도입하고 있으며, 사적 영역에서 보상을 수취할 수 있는 가능성을 열어두고 있다.

독일의 경우, 제60h조 제1항에서 ‘본 절에 따른 이용’에 대해 공정한 보상금을 수취할 권리를 규정하고 있다. 그러나 DSM 지침의 도입으로 TDM 조항을 개정하면서 비-학술연구 목적의 이용에 관한 DSM 지침 제4조에 대응하는 제44b조를 제60h조와 다른 절에 규정하였다. 그 결과, 비-학술연구 목적의 TDM 이용도 보상금지급의무의 대상에 포함되지 않는다. 따라서 권리자는 옵트아웃제도를 이용하여 경제적 보상을 수취할 수 있을 뿐이다.²²⁰⁾

일본의 경우 별도의 보상의무 규정을 두고 있지 않다. TDM 조항을 도입하고 있으면서 보상의무 규정과 옵트아웃제도를 모두 도입하지 않고 있는 일본의 저작권법 조항은 사실상 TDM 면책 대상이 되는 경우 권리자에 대한 보상의무가 발생하지 않는다는 점에서 자유로운 데이터의 이용이 보장된다.

미국의 경우, 별도의 TDM 규정을 두고 있지는 않다. 다만, 최근 저작권청(USCO)이 최근 공시한 의견수렴절차 개시통지서(Notice of Inquiry)에는 총 34개의 질문과 그 하위 질문들을 제시하고 있는데, 데이터 이용을 강제허락 받는 제도의 도입을 논의하면서 강제허락에 따른 보상금 요율이나 이용조건은 어떻게 정할 것인지, 보상금은 어떻게 할당·보고·분배해야 할 것인지, 확대된 집중관리제도가 적합하거나 바람직한 접근방법인가를 묻고 있으며, 생성형 인

220) 류시원, 위 논문 361쪽

공지능에 관하여 저작권자/창작자에게 보상을 제공하는 라이선스 제도의 작동 방법, 옵트아웃제도의 도입 여부에 관한 의견 또한 수렴 중에 있다.²²¹⁾

(3) 시사점

이처럼, 현재 주요 외국 법제는 대체로 TDM 조항에 권리자에 대한 보상금 지급 의무 규정을 두지 않고 있으며, 경우에 따라 옵트아웃제도를 운영하여 사적 영역에서의 보상을 열어두고 있다.

우리나라의 경우에도 권리자의 재산상 권리 침해에 대한 보호수단을 마련해 이해관계를 조정해야 한다는 취지의 의견이 제출되고 있다. 우리나라의 문화산업 현실에 맞추어 보상체계를 마련하는 방안에 대해 고민해 볼 필요가 있을 것이다.

바. 학습데이터 공개 관련

EU 인공지능법은 범용 AI에 대하여 규제하고 있다. GPAI-M²²²⁾ 공급자는 일반적으로 (i) 학습 및 테스트 절차, 학습 콘텐츠 및 성능 평가결과 등을 포함하는 기술 문서의 작성, 유지 및 이를 감독 당국 및 해당 모델을 사용하는 공급자에게 제공할 의무(투명성 의무) (ii) EU 디지털 시장 저작권 지침(DSM)을 존중(respect)하는 정책을 수립, 시행할 의무를 진다.

시스템적 위험²²³⁾이 있는 GPAI-M 공급자는 위 일반적 의무에 더하여 (i)

221) 이대회, 「미국 Copyright Office의 AI 저작권 쟁점 의견 청취」, 한국저작권위원회, 이슈리포트, 제2023-11호, 2023. 12.

222) 적절한 기술적 도구와 방법론을 기반으로 평가된 ‘높은 영향력’이 인정되면 시스템적 위험이 있는 GPAI-M으로 분류될 수 있으며, EU 인공지능 위원회가 직권으로 해당 여부를 결정할 수도 있다. 한편 10^{25} FLOPS 컴퓨팅 파워 이상을 가진 컴퓨터로 훈련된 GPAI-M은 시스템적 위험이 있는 것으로 간주된다

최신 기술이 반영된 표준화된 프로토콜 및 수단을 사용한 모델 평가(특히, 적대적 테스트(adversarial test))를 수행하고 그 결과를 문서화할 의무, (ii) 시스템적 위험을 평가하고 이를 완화할 의무, (iii) 중대 사고 발생시 이를 감독 당국에 보고할 의무, (iv) 적절한 수준의 사이버 보안을 보장할 의무 등을 부담한다.

저작권에서 보호하는 텍스트 및 데이터를 포함한 범용인공지능모델의 사전학습 및 학습에 사용되는 데이터에 관한 투명성을 높이기 위하여 해당 모델의 제공자는 범용 모델의 학습에 사용된 콘텐츠에 관한 ‘충분히 상세한 요약’을 작성하고 공개하는 것이 적절하다. 영업비밀 및 기업비밀정보를 보호할 필요성을 고려하면서 이 요약은 기술적으로 상세하기 보다는 저작권자를 포함한 정당한 이익이 있는 당사자가 예를 들어, 대규모 민간 또는 공공 데이터베이스 또는 데이터 아카이브 등 해당 모델의 학습에 들어간 주요 데이터 컬렉션 또는 세트를 열거하고 사용된 그 밖의 데이터 소스를 서술식으로 설명하여 유럽연합 법령에 따른 권리를 행사하거나 집행하는 것을 촉진하기 위하여 그 범위에서 대체로 포괄적이어야 한다. 인공지능사무국은 간단하고 효과적이며 제공자가 서술식으로 요구하는 요약을 제공할 수 있는 요약 서식을 제공하는 것이 적절하다.

한편, DSM 지침의 경우 TDM이 아무런 제한없이 허용되는 것은 아니다. 권리보유자가 자신의 저작물이나 보호대상에 대한 TDM을 원하지 않는다면 필요한 조치를 할 수 있다. 즉 권리보유자는 온라인에 게시된 콘텐츠의 경우 기계가 판독 가능한 수단 등 적절한 방식으로 자신의 저작물이나 기타 보호대상을 사용할 권리를 명시적으로 유보할 수 있다. 기계가 판독 가능한 수단으로는 “robots.txt” 파일이 대표적이다. TDM에 대한 권리보유자의 유보는 옵트

- 223) ‘시스템적 위험’이란 GPAI-M로 인하여 EU시장에 중대한 영향을 미치며 공공 보건, 안전, 안보, 기본권 또는 EU 사회 전체에 실질적으로 또는 합리적으로 예측가능한 부정적인 영향을 미칠 위험을 말한다.

아웃(Opt-out) 권리이다. 이러한 권리가 보다 효과적으로 행사되기 위해서 권리보유자는 자신의 어떤 저작물이 TDM 목적으로 복제되어 AI의 학습에 이용되었는지 알아야 한다. 따라서 AI 학습데이터의 요약 공개의무는 DSM 지침 제4조 제3항이 권리보유자에게 부여한 옵트아웃 권리를 보완한 것이라고 판단된다. AI 학습 목적으로 채굴 및 추출된 모든 저작물이 포함된 보고서가 없는 경우, 권리자가 자신의 저작물이 소프트웨어에 주입되었다는 사실을 발견하는 것은 거의 불가능하기 때문이다.

3. 인공지능 산출물의 보호

가. 인공지능 산출물 보호 관련 국내외 동향

(1) 미국²²⁴⁾

미국 저작권청 심사위원회는 2022년 2월 인공지능 창작 미술작품의 저작권 등록 거절을 최종 결정하였다. 저작권법은 오직 인간 정신의 창조적 힘에 의해 만들어진 지적 노동의 결실만을 보호한다고 판단하였다. 저작권청 실무 지침은 인간의 창작이 결합된 작품 예로 ‘원숭이가 찍은 사진’, ‘성령을 작자로 명명하는 노래’를 제시하였다.

미국 연방대법원을 비롯한 법원은 지금까지 저작권 보호 대상을 ‘인간에 의한 창작물’로 일관되게 제한하고 있다.

224) 이해영, 「미국저작권청(USCO), 인공지능이 만든 저작물 등록 거절」, 한국저작권보호원, 해외 저작권 보호 동향, 2022. 7. 14.

[그림 6] 스테판 탈러의 ‘A Recent Entrance to Paradise’



자료 : 이해영, 「미국저작권청(USCO), 인공지능이 만든 저작물 등록 거절」, 한국저작권보호원, 해외 저작권 보호 동향, 2022. 7. 14.

한편 AI가 저작자에 해당하는지 여부를 검토한 미국 법원 판례는 아직 없다. AI DABUS의 발명품에 대한 특허출원을 특허청이 등록거절 결정한 사안에서 버지니아 동부지방법원은 특허법에 따라 ‘발명자’는 자연인이라고 판단한 바 있다.

(2) 중국

원고는 2023. 2. 24. 생성형 인공지능인 ‘스테이블 디퓨전(Stable Diffusion)’에 프롬프트를 입력하는 방식으로 이미지를 생성하여 SNS에 게재하였으며, 피고는 그 이미지를 무단으로 사용하였다.

1심 법원인 베이징인터넷법원은 2023. 12. 원고의 이미지에 대한 저작권을 인정하였다. 원고의 지적인 노력을 통해 이미지가 최종적으로 생성되었으므로, 해당 이미지는 지적 성과에 해당하며, 원고는 인물과 그 표현방식 등의 요소에 대하여 프롬프트를 설정하고 화면 내 요소의 배치와 구도에 있어서 매개변수를 조절함으로써 원고의 개인화된 표현으로 인정할 수 있으므로 ‘독창

성’ 요건도 인정되고, 본질적으로 인공지능을 도구로 이용하여 창작한 것이므로 인공지능이 아닌 사람이 창작한 것으로 보아 원고의 저작권을 인정하고, 피고의 행위가 저작권 침해에 해당한다고 판결하였다((2023)경0497민초11279 호).²²⁵⁾

위 베이징 인터넷 법원의 판결은, AI를 사용하여 생성한 이미지라고 하더라도 최종 결과물 생성에 이르기까지 인간의 창작적인 기여와 노력이 투입되었다면 저작물로 인정될 수 있다는 입장을 취한 것으로 볼 수 있다.

생성형 인공지능은 사람이 프롬프트를 입력하면 그에 따른 결과물을 출력하는데, 이러한 프롬프트를 아이디어나 소재의 제공 수준에 불과한 것으로 평가할지, 아니면 이용자의 창작적인 표현이 나타난 것으로 평가할 수 있는지 문제될 수 있다. 최종 산출물에 이르기까지 어느 정도로 프롬프트나 매개변수를 입력, 조정해야 창작적 기여로 평가될 수 있을 것인지가 문제인데, 구체적인 사례가 풍부하게 축적되어야 이에 대한 정확한 평가가 가능할 것이다.²²⁶⁾

(3) 한국

우리나라 국회는 2020년 12월 ‘인공지능 저작물’ 개념 명시한 개정안 발의(임기만료 폐기). ‘인공지능 저작물의 저작자’를 ‘인공지능 서비스를 이용하여 저작물을 창작한 자 또는 인공지능 저작물의 제작에 창작적 기여를 한 인공지능 제작자·서비스 제공자 등’으로 정의하고 저작물을 공표한 때로부터 5년간 보호, 해당 저작물 등록 시 인공지능이 제작한 작품 표기 의무를 규정한 바 있다.²²⁷⁾

225) 백지연, 「[중국] 베이징인터넷법원, 생성형 AI 산출물의 저작권을 인정」, 한국저작권위원회, 저작권 동향 2023 제11호, 2023. 12.

226) 김우균 외, “생성형 AI 저작권 관련 중국 판례 동향”, 법률신문, 2024. 6. 17. <<https://www.lawtimes.co.kr/LawFirm-NewsLetter/199153> 최종 방문일 2024. 9. 10>

나. 인공지능 산출물의 저작자 결정

인간과 AI를 공동저작자로 인정한 사례로 인도와 캐나다 사례를 들 수 있다.

(1) 인도

인도 저작권청은 2020년 11월 세계 최초로 인공지능에게 저작자의 지위를 인정하였는데, ‘RAGHAV 인공지능 페인팅 앱(RAGHAV Artificial Intelligence Painting App)’이라고 불리는 인공지능 앱이 생성한 미술 작품을 이 앱의 소유자와 공동저작자로 하는 저작권 등록 신청을 승인하였다. 신청인은 인공지능 앱을 단독 저작자로 하는 저작권 등록이 거절되자 이 앱의 소유자인 신청인을 공동저작자로 하여 다시 저작권 등록을 신청하였다. 신청인은 ‘일몰(Suryast)’이라는 제목으로 인공지능 앱이 생성한 미술 작품을 등록하였다. 인공지능 페인팅 앱은 다양한 예술 스타일을 훈련받았는데 고흐의 ‘별이 빛나는 밤’과 인공지능 앱의 소유자가 촬영한 사진으로 미술 작품을 생성하는 데이터셋으로 이용하였다. 그러나 2021년 11월 25일 인도 저작권청은 이 작품에 대한 저작권 등록 철회를 통지하였다. 인도 저작권청은 철회통지를 통해 신청인에게 인공지능 앱의 법적 상태(legal status)에 대한 정보를 제공하도록 요청하였고, 저작권법 제2조(d)(iii) 및 제2조(d)(vi)를 검토하도록 하였다. 저작권 등록신청인이자 공동저작자는 저작권법과 저작권법 규칙은 저작권법 제2조(d)(vi)에 규정된 ‘사람(person)’의 의미를 정의하고 있지 않고 일반법(General Clauses Act)에서는 이 의미를 인공적인 사람과 법적인 사람을 포함하는 개념으로 정의하고 있다고 주장하면서, 저작권법 제2조(d)(vi)는 컴퓨터를 이용한 저작물의 경우 이용자를 저작자로 규정하고 있어서 등록 전체를 철회할 수 없다는 점 등을

227) 저작권법 일부개정법률안, 주호영의원 대표발의, 의안번호 2106785, 2020. 12. 21.

주장하며 이의신청을 하였다. 현재까지 인도저작권청이 최종 결정을 내리지 않고 있어서 ‘RAGHAV’ 인공지능 앱은 공동저작자의 지위를 유지하고 있다.²²⁸⁾

(2) 캐나다

캐나다 지식재산청(Canada Intellectual Property Office)은 2021년 12월 1일 미술작품의 공동저작자 중 한 명을 인공지능 앱으로 하는 저작권 등록 신청을 승인하였다. 인도 저작권청에 등록된 것과 동일한 작품이 캐나다 지식재산권 청에서도 저작권으로 등록되었다. 이 인공지능 앱은 캐나다에서도 인공지능이 저작자로 등록된 첫 번째 사례이다. 캐나다 지식재산청의 저작권 등록은 실질적인 심사의 대상은 아니지만 신청서에 기재된 세부사항에 대해서는 공식적인 확인 절차를 진행하고 있다. 캐나다 저작권법이 저작자의 의미를 규정하지 않아 인공지능이 저작자가 될 수 있는지 여부가 명확하지는 않았지만 이번 저작권 등록을 통해 새로운 해석이 가능할 것으로 보인다.²²⁹⁾

4. 인공지능 산출물에 의한 저작권 침해

가. 일반 법리

(1) 의거 관계

학습완료 모델에는 학습데이터가 남아 있지 않아 의거성 입증에 어려움이 발생한다.

228) 송선미, 「AI 창작물의 저작권 보호에 관한 해외 동향」, 한국저작권위원회, 이슈리포트 2022-02. 2022. 3.

229) 송선미, 위 보고서

그런데 현재 기술수준으로는 심층학습 AI의 출력값이 어느 입력값, 노드, 가중치에 의해 산출되었는지 경로 파악이 불가능하며, 학습데이터에 원저작물이 포함되었다는 점을 입증할 수 있더라도 AI 산출물이 그 중 어느 원저작물에 의거하여 작성되었는지 특정이 곤란하다.²³⁰⁾

기계학습은 데이터들의 패턴, 특징을 추출하여 인식하는 것이고 특정 데이터의 표현에 의존하는 것이 아니라 학습데이터에 저작물이 있었다고 하더라도 의거성이 부정될 여지가 존재한다.

반면, AI 생성물이 기존 저작물과 실질적으로 유사하다면 그 의거를 인정할 수 있다는 견해도 존재한다.

(2) 실질적 유사성

실질적 유사성은 구체적 사실관계에 달려 있으며, 검토자마다 의견이 상이할 수 있다.

Andersen v. Stability 사건에서 법원은 실질적 유사성 입증 부족하다고 판단한 바 있어, 인공지능 저작권 소송에서의 쟁점 중에 하나로 논의되고 있다.²³¹⁾ 한편, 실질적 유사성을 가진 저작물이 전체에서의 비중도 고려할 필요가 있다.

(3) 책임 귀속

AI 산출물을 지배·관리하여 권리나 경제적 이익을 취득할 수 있는 자는 AI에 의한 저작권 침해에 대해서 법적 책임을 부담해야 한다는 견해가 존재한다.²³²⁾ 다만, 개발자와 이용자의 공동 침해자 여부는 명확하지 않다.

230) 전웅준, 「AI 관련 저작권, 개인정보, 규제 쟁점 및 입법적 과제」, 국회입법조사처, 2024 국가비전 입법정책 컨퍼런스, 2024. 4.

231) 이대희, 「AI 저작권 관련 Sarah Andersen v. Stability AI 소송 분석」, 한국저작권위원회 이슈리포트 2023-15, 2023. 12.

232) 전웅준, 위 발표자료

나. 침해 인정 사례²³³⁾

2024년 2월 8일 중국 광저우 인터넷 법원은 생성형 인공지능 서비스로 만들어진 이미지에 대해 저작권 침해를 인정한 판결을 하였다.

츠부라야 프로덕션은 일본의 대표적인 캐릭터인 ‘울트라맨’ 시리즈의 저작권을 보유한 회사이다. 원고인 상해신창화문화발전유한공사는 츠부라야 프로덕션으로부터 중국 내에서 울트라맨 저작물에 대한 독점적인 라이선스를 부여받았다.

피고는 텍스트를 이미지로 변환해 주는 생성형 인공지능 서비스를 제공하는 웹사이트를 운영하고 있었는데, 2023년 말 원고는 해당 웹사이트에서 ‘울트라맨’과 관련된 프롬프트를 입력하면 울트라맨과 유사한 이미지를 생성할 수 있다는 것을 발견했다. 이에 원고는 피고가 저작권 침해를 했다고 주장하며 소송을 제기했다.

원고는 복제권, 각색권(저작물 변형을 통해 새롭게 창작할 수 있는 권리), 정보네트워크중계권(대중이 저작물을 특정 시간과 장소에서 접근할 수 있도록 제공할 권리)을 침해했다고 주장하며, 피고가 울트라맨과 유사한 이미지 생성을 중단할 것, 울트라맨과 관련된 학습 데이터를 삭제할 것, 그리고 300,000위안의 손해배상을 청구했다.

이에, 광저우 인터넷 법원은 피고의 AI 서비스가 울트라맨과 실질적으로 유사한 이미지를 생성했으며, 복제권과 각색권을 침해했다고 판단했다. 복제권과 각색권 침해가 인정된 이상 정보네트워크중계권 침해 여부는 별도로 판단하지 않았다. 법원은 피고가 유사한 이미지 생성을 방지하기 위한 조치를 취하고, 10,000위안을 배상할 것을 명령했다.

하지만 법원은 피고가 직접 AI 모델을 개발하여 학습시킨 것이 아니며 타사의 모델을 사용한 것이기 때문에, 원고가 요구한 울트라맨 관련 학습 데

233) 법률신문, 위 기사

이터 삭제 요구는 기각했다.

피고가 손해배상 책임을 지기 위해서는 과실이 인정되어야 하는데, 이 사건에서 법원은 중국의 ‘생성형 인공지능 잠정관리 방법’²³⁴⁾을 적용하여 피고의 과실 여부를 판단했다.

피고는 자신이 제공하는 생성형 인공지능 모델이 자사가 개발한 것이 아니기 때문에 해당 법령의 적용을 받지 않는다고 주장했다. 그러나 법원은 ‘생성형 인공지능 잠정관리 방법’에 따라 프로그래밍 가능한 인터페이스를 통해 서비스를 제공하는 자도 생성형 인공지능 서비스 제공자에 해당하므로, 피고도 이 법령의 적용을 받는다고 판결했다.

법원은 피고가 생성형 인공지능 서비스를 운영하면서 법령에 규정된 사용자 규정과 고지 의무를 이행하지 않았고, 민원 제기 및 신고 체제를 구축하지 않았으며, AI 산출물에 대한 표시 의무도 다하지 않았다고 판단했다. 따라서 피고의 손해배상 책임을 인정하고, 필요한 조치를 취하지 않은 점을 들어 피고가 법령을 위반했다고 판결했다.

위 판결은 생성형 인공지능의 결과물이 저작권을 침해했다는 최초의 판결로서 의미가 있다. 다만, 실제로 저작권 침해 이미지를 생성한 것은 피고가 아니라 그 인공지능 서비스를 이용한 이용자가 프롬프트를 입력해 이미지를 생성한 것이므로, 생성형 인공지능 서비스 제공자에 대해서만 저작권 침해 책임을 지도록 한 점은 한계로 평가된다.

234) 중국의 ‘생성형 인공지능 잠정관리 방법’은 생성형 인공지능의 건전한 발전과 규범의 적용을 추진하고 국가안보와 사회 공익을 수호하며 공민, 법인 그 밖의 조직의 정당한 권리를 보호하기 위하여 제정된 법령으로서(제1조), 지식재산권 등 다른 사람의 권리 존중, 데이터 학습 시의 준수 사항, 생성형 인공지능으로 만들어진 콘텐츠에 대한 표시 의무, 민원 제기 및 신고 체제 구축 의무 등을 정하고 있다.

5. 인공지능 저작권 규제체계 입법 방향

가. 총론

인공지능의 활용이 확대되면서 저작권을 둘러싼 논쟁이 심화되고 분쟁이 발생하고 있다. 특히, 생성형 인공지능 기반 모델이 저작물을 학습할 때 저작권 침해가 면책될 수 있는지, 즉 저작재산권이 제한될 수 있는지가 주요 쟁점이다.

인공지능의 저작물 학습 시 저작권 면책을 인정할지 여부에 관하여, 저작권 보호를 통한 문화 창달과 인공지능 기술의 편익 증진을 모두 고려한 조화로운 방안을 모색할 필요가 있다.

한편, 저작권법상 포괄적 공정이용을 비롯한 저작재산권의 제한 규정은 저작인격권에 영향을 미치지 않는다(저작권법 제38조). 이에 따라, 본 보고서에서는 저작재산권에 한정하여 논의한다.

나. 포괄적 공정이용 조항 외에 TDM 면책 조항 필요 여부

우리나라는 미국 등과 같이 포괄적 공정이용 조항을 도입하였으므로 그 해석에 따라 인공지능 저작물 학습을 면책시키면 된다는 주장이 제기될 수 있다. 인공지능 기술이 나날이 발전하고 있어 일반조항을 유연하게 적용할 필요성이 인정될 수 있다.

그러나 우리나라는 미국과는 달리 대륙법 체계를 채택하고 있어 포괄적 조항의 해석만으로 한계가 있으며 도입된지 상당한 기간이 지났음에도 판례가 많이 축적되지 못한 상황이므로 수범자들의 예측가능성이 저해될 수 있다. 포괄적 공정이용 조항과 TDM 면책 조항은 그 입법 취지와 요건도 다르므로, TDM 면책 조항을 별도로 도입을 하는 것이 보다 직접적인 문제 해결 방식이라고 볼 수 있다.

다. 영리 목적 인정 여부

생성형 인공지능 개발자(제공자)는 통상 영리 목적으로 개발하며, 배포자 역시 영리 목적으로 이용하는 경우가 다수이다. 인공지능 개발에 상당한 자원과 비용이 드는 만큼 일부 공공 분야 외에는 비영리 목적의 인공지능 개발은 찾기 어려울 수 있다.

나아가, 영리 목적과 비영리 목적의 구분이 불분명한 경우도 발생할 수 있어 TDM 면책을 비영리 목적인 경우에만 적용하는 방식의 실효성이 높지 않을 수 있다.

이에 따라, 영리 목적인 경우에도 TDM 면책을 인정하되, 옵트아웃 권리 인정이나 보상 인정 등을 통해 저작권자의 권리 인정을 통해 이해관계자 간 균형을 도모하는 방안을 고려할 수 있다.

라. 저작권자의 옵트아웃 권리 인정 여부

저작권자가 자신의 저작물을 인공지능 학습에 사용하지 않도록 명시적인 표시를 한다면 TDM 면책을 인정하지 않는 방안을 고려할 수 있다. 저작권자에게 옵트아웃을 부여하지 않는 상황에서 제공·실행자를 무조건 면책시키는 것은 바람직하지 않다는 견해가 있다.²³⁵⁾

다만, 저작권자의 옵트아웃 권리를 넓게 인정해 줄 경우, 인공지능이 학습할 수 있는 저작물의 범위가 상당히 제한될 수 있는 우려가 있다. 이에 따라, 저작권자의 옵트아웃 표시 방법을 제한하는 방안, 저작권자가 옵트아웃을 하더라도 정당한 보상을 하도록 하는 방안 등을 통해 이해관계자 간 균형을 찾아볼 수 있을 것이다.

235) 박상철, 「인공지능의 법적 규율 I : 범용모델과 생성모델」, 서울대학교 법학 제65권 제1호, 2024. 3. 105쪽

마. 대상 권리 범위

통상 TDM 면책은 복제·전송에 한정하여 인정하도록 되어 있는데, 2차적 저작물작성권까지 인정되도록 할지 문제가 된다.

TDM 면책은 저작권자를 제한하는 규정으로서 저작권자에게 불측의 손해가 발생할 우려가 있는 점을 고려하면 그 범위를 확대하는 것은 바람직하지 않을 수 있다. 복제·전송에 한하여 인정하고, 추후 생성형 인공지능 학습 및 이용 상황을 보면서 면책 범위를 확대할지 여부를 결정할 필요가 있다.

바. 적법한 접근 관련

해킹과 같은 정보통신망 침해 행위를 통하여 저작물에 접근한다면 이는 적법한 접근이라고 보기 어려울 것이다.

그런데 로봇배제표준(robotx.txt)을 이용하여 크롤링을 차단하는 의사표시를 하였음에도 이를 우회하여 접근하였을 때, 적법한 접근이라고 볼 수 있는지 문제가 된다.

이는 앞의 ‘라. 저작권자의 옵트아웃 권리 인정 여부’ 항의 옵트아웃 표시 방법과 관련되는데, 제도적으로 저작권자의 권리를 확대한다면 그에 따라 적법한 접근의 범위를 넓게 인정하고, 반대로 저작권자의 권리가 상대적으로 축소된다면 적법한 접근의 범위를 좁게 인정하도록, 유연하게 제도를 설계할 필요가 있다.

사. 학습데이터 출처 표시 여부

학습데이터 출처를 엄격하게 표시하도록 한다면, 상당한 양의 데이터·저작물을 학습시키는 인공지능 개발자에게는 매우 큰 부담이 될 것이다.

반면, 저작권자가 자신의 저작물이 무단으로 인공지능 학습에 사용되었는

지 여부를 확인하기 위해서는, 인공지능 학습 시 학습데이터 출처 표시가 되어야 할 필요가 있다.

이에 따라, EU 인공지능법과 같이 인공지능이 학습한 데이터의 출처에 대한 요약을 표시하도록 하여, 저작권자가 추후 분쟁을 제기하고자 할 때 근거로 삼을 수 있도록 하는 방안을 고려할 수 있다. 다만, 어느 정도로 상세하게 표시해야 할지에 대하여 결정이 쉽지 않으며, 저작권 분쟁 과정에서 저작권자가 자신의 권리가 침해됐음을 입증하는 방법 및 법원의 인정 사례, 인공지능 개발자의 부담 등을 종합적으로 고려할 필요가 있다.

아. 복제물의 보관기간 관련

복제물의 보관기간에 대해서는 독일의 사례를 참조하여 데이터마이닝 수행 후 해당 복제물을 삭제할 의무를 규정하는 것을 고려할 수 있다.

자. 보상 여부

TDM 면책 시 보상 의무를 입법화 할 것인지 여부에 대하여, 주요국 TDM 면책 조항이나 포괄적 공정이용 조항에서 보상 의무를 규정하고 있지 않다는 점에서 우리나라의 독자적인 입법화에 대하여는 심도있는 논의가 필요해 보인다. 보상 의무 여부만을 두고 판단할 것이 아니라, 옵트아웃 권리 인정 여부, 영리 목적 이용 시 면책 여부 등 여러 쟁점들을 종합적으로 고려하여 제도를 설계할 필요가 있다.

IV. 결론

1. 연구 요약

가. 인공지능 윤리 규제체계 현황

인공지능 윤리 원칙의 규범력과 강제력 확보를 위해 정책이나 입법이 필요하며, 본 연구는 인공지능 윤리와 관련하여 법으로써 규율할 수 있는 규제 체계에 중점을 두었다.

해외 입법 사례 중에 특히 EU와 미국을 주목하였다. 최근 발효된 EU 인공지능법은 위험 기반 접근법을 채택하였으며, 인공지능 기술 개발 및 활용과 관련한 이해관계자를 제공자, 배포자, 수입업자 및 유통업자로 구분하고 역할 별로 의무사항을 상세히 규정하였다. 범용 인공지능 기반 모델 규제, 강한 행정규제, EU와 회원국 간 다층적 거버넌스, 혁신지원 제도 등도 특징이다. 세계 최초의 통합적이고 옴니버스 형태를 가진 인공지능 규제로서 다른 국가의 입법에도 영향을 미치고 있다.

EU의 통합적인 접근방식은 전 분야에 걸쳐 인공지능의 위험성을 사전에 예방한다는 장점이 있지만, 여러 분야를 한데 묶어 규율하다보니 법률이 복잡하고 아직 명확하지 않은 규제 영역이 많아 향후 법 시행 상황을 지켜볼 필요가 있다. EU 인공지능법을 무조건적으로 수용하기보다는 그 특징과 장점을 잘 살리되 미국 등 다른 국가의 사례와 우리나라의 특수성 등을 고려하여 우리나라의 고유한 제도로 발전시키는 방안을 모색할 필요가 있다.

미국은 연방정부 차원의 포괄적인 규제 법률은 제정되어 있지 않지만, 행정명령을 통해 공공분야에서부터 인공지능 정책을 추진하고 있으며, 캘리포니아주(州), 콜로라도주 등 개별 주에서 상당한 강한 규제가 담긴 법률이 제정되고 있다.

EU와 캐나다는 상대적으로 중앙집권적인 집행이 가능한 수평적 규제, 즉 대부분의 분야와 애플리케이션에 걸쳐 적용되도록 설계된 규제를 도입하는 접근 방식을 취한다. 반면, 미국, 영국, 중국은 보다 분산된 접근 방식을 취하고 있으며, 여러 규제 기관에서 시행하는 수직적 규제, 즉 인공지능의 특정 애플리케이션이나 특정 부문에만 적용되는 규제에 대한 의존도가 더 높다. 미국의 접근 방식은 구속력이 없는 원칙, 위험 관리에 대한 자발적 지침, 연방 차원에서 새로운 인공지능 관련 법률을 개발하기보다는 기존의 부문별 법률을 적용하는 것이 특징이다. EU와 미국 방식 중 어느 한 방식을 규범적으로 선택하여 적용하기 보다는, 국민의 권리를 보호하면서도 인공지능 기술·산업 발전을 촉진할 수 있는 방안을 유연하게 혼합하여 적용하는 방안을 고려할 필요가 있다.

허위조작정보(딥페이크)는 최근 사회적 문제로 대두되고 있으며, 국내외에서 허위 영상물을 금지하거나 인공지능이 생성한 영상물·이미지라는 점을 표시하도록 하는 규제가 입법화되거나 추진 중에 있다.

딥페이크를 통한 허위 영상물의 제작 및 반포를 금지하는 규제들이 분야별로 입법화되고 있다. 이에 더해, 다양하게 발생하는 허위 영상물 관련 문제점들을 효과적으로 해결하기 위하여 인공지능 기본법 또는 허위 영상물 관련 특별법에서 이를 예방·방지하는 시책을 추진하도록 하는 내용을 담은 방안을 고려할 수 있다.

인공지능으로 만든 가상 정보에 대하여 표시를 하도록 한다면 이를 통해 사람들은 인공지능으로 만든 가상 정보와 실제 사실을 쉽고 효율적으로 구분할 수 있으며, 허위조작정보 금지뿐만 아니라 저작권 침해 예방에도 효과적일 것으로 전망된다. 이에 따라 특정 분야의 법률이 아닌 인공지능 전반을 규율하는 기본법에 규정하는 방안을 고려할 수 있다.

그러나 가상 정보 표시에 대한 기술적 회피·조작 가능성, 실제 사실에 대한 거짓 표시 가능성 등이 있으므로, 이를 해결하지 못한다면 표시 제도의 유

용성·신뢰성이 저하될 우려가 있다.

이러한 점을 고려할 때, 현 시점에서는 적용 콘텐츠나 수범자의 범위를 좁히고 미준수 시 제재는 최소한으로 하되, 인공지능 생성물 탐지 기술이나 표시 의무 우회를 방지하는 기술 등의 발전 상황을 지켜보면서 규제를 넓혀가는 방안을 고려할 필요가 있다.

나. 인공지능 저작권 규제체계 현황

최근 생성형 인공지능 도입이 증가함에 따라 저작권 침해가 문제되고 있으며, 해외에서 법적 분쟁이 다수 발생하고 있다.

인공지능 저작권과 관련된 쟁점으로 인공지능이 생성한 저작물에 저작권이 부여될 수 있는지 여부가 쟁점이 될 수 있는데, 인공지능에 권리능력이나 행위능력 등을 인정하기 어려워 순수하게 인공지능 자체에 저작권을 부여하기 어려우며, 사람인 저작자가 인공지능을 도구로서 사용한 경우에는 기존의 법리에 따라 창작성이 인정된다면 저작권이 인정될 수 있다. 그리고 인공지능이 생성한 결과물이 다른 저작권을 침해하는지 여부는 기존의 저작권 침해 법리를 따르면 될 것으로 보인다. 이에 따라 본 연구에서는 인공지능이 저작물을 학습하는 경우 저작권이 면책되는지 여부에 관하여 중점적으로 살펴보았다.

저작권 면책과 관련한 입법 방식은 통상 포괄적 공정이용과 TDM 면책조항으로 구분된다. 미국은 포괄적 공정이용 조항을 두고 있으며, 인공지능 학습은 비표현적이고 상업적 목적이 없다고 해석될 수 있으며, 저작권이 있는 자료를 인공지능 훈련에 사용하는 것은 혁신성을 인정받아 공정이용으로 판단될 가능성이 있다. 특히, 미국 법원은 비표현적이고 변형적인 이용에 대해 공정이용을 비교적 넓게 인정하는 경향이 있었다. 그러나 생성형 인공지능 모델은 상업적 성격이 충분히 인정될 수 있으며, 인공지능 생성물이 학습데이터인 저작물에 대한 수요를 대체할 수 있다고 볼 여지가 있다. 최근 미국 연방대법

원 판결에 따르면, 생성형 인공지능의 결과물의 목적이 원저작물과 동일하다면 공정이용 가능성이 낮아질 것이다. 반대로, 인공지능 개발자는 인공지능 결과물과 원저작물 간 목적이 실질적으로 다르다는 점을 입증하면 공정이용이 인정될 여지가 있을 것이다.

EU, 영국, 일본, 독일, 싱가포르 등은 TDM 면책 규정을 도입하였다. TDM 면책 관련 법제를 마련한 주요국들은 TDM 생태계 내 다양한 이해관계를 조정하기 위해 각기 다른 접근 방식을 취하고 있다. EU의 DSM 지침은 학술연구 목적의 TDM과 일반적인 TDM을 구분하여 이원적으로 규율하고 있으며, 학술연구 목적의 TDM에 대해서는 계약이나 기술적 수단으로 면책을 배제하지 못하도록 강제하는 반면, 일반 TDM에 대해서는 옵트아웃을 허용함으로써 균형을 맞추고 있다.

영국과 싱가포르는 상업적/비상업적 목적의 구분 없이 옵트아웃을 허용하지 않는 강력한 규정을 도입함으로써 EU보다 더 넓은 범위의 면책을 제공하고 있다. 특히, 싱가포르의 입법례는 ‘적법한 접근’ 요건을 유지하면서도 강행 규정을 적용함으로써, TDM을 위한 복제물의 보관 및 제3자 제공을 허용하는 등 포괄적인 면책 규정을 마련하였다.

일본은 EU DSM 지침과 유사하게 TDM 면책 규정을 이원화하여 정보해석 그 자체를 위한 비향수적 이용과, 정보해석에 부수하는 경미한 이용을 구분하고 있다. 일본 저작권법 제30조의4는 비향수적 이용을, 제47조의5는 경미한 이용을 각각 면책 대상으로 규정하고 있으며, TDM 활동 전반에 걸쳐 실효성을 높이기 위한 규정이라고 평가될 수 있다.

우리나라에서는 2021년 TDM 면책 규정을 저작권법에 도입하는 내용의 저작권법 전부개정법률안이 발의되었고, 이후 여러 법안이 발의되었지만, 저작권자들의 반대 등으로 통과되지 못하고 임기만료로 폐기되었으며, 현재 제22대 국회에서는 아직 관련 법안이 발의가 되지 않은 상황이다.

2. 입법 방향

가. 인공지능 윤리 입법 방향

인공지능 윤리 규제체제와 관련하여, EU와 같은 수평적·하향식 규제체제는 포괄적이고 전 분야에 걸쳐 촘촘하게 적용될 수 있어 인공지능 기술에 대한 신뢰를 높일 수 있다. 특히 생성형 인공지능 기반 모델은 범용성을 갖고 있으며 거의 전 분야에 적용되고 있고 그 성능이 비약적으로 발전하고 있어, 이에 대한 개발·생산·배포·유통·활용 전반에 걸쳐 일반적 규제요구사항을 담은 규제체제가 요구되고 있다.

한편, 인공지능 기술은 나날이 발전하고 있고 인공지능 시스템의 상당수는 아직 초기단계이거나 제대로 사업화되지 않고 있어, 강력하고 경직된 규제를 적용하는 것이 적절한지에 대한 논란이 있다. 위험 기반 접근법을 적용할 경우 금지되는 인공지능이나 고위험 인공지능에 해당하는지 불분명한 경우가 발생할 것으로 예상되며, 실증적 근거가 부족하고, 되돌릴 수 없는 손해가 분명히 예견되는 상황도 아니어서 사전 접근 방식이 적절하지 않다는 견해도 있다.

아울러, 우리나라는 자체 인터넷 포털과 모바일 인스턴트 메시지 서비스를 제공하고 있고, 해외 정보기술 시장에도 진입하고 있어 EU와 같은 강한 규제를 적용하는 것이 적절한지 고민이 필요하다.

이런 점을 고려할 때, 규제의 대상이 되는 분야를 한정적으로 열거하고, 사업자에게 설명의무 등 최소한의 의무만을 부여하며, 위반 시 당국의 조사 및 공표 정도의 낮은 수준의 제재를 도입하는 것이 바람직하다고 판단된다. 이후 인공지능 기술의 발전상을 따라가며, 인공지능으로 인하여 피해를 입은 구체적 사례를 분석하고 이를 축적하여 적절한 규제 범위와 제재 수준을 정하는 것을 고려할 수 있다.

위험 기반 규제 방식과 관련하여 EU 인공지능법은 인공지능 시스템의 위험성을 기준으로 차등 규제를 도입하고 있다. 고위험 인공지능에 대해선 개발자와 배포자에게 엄격한 규제를 부과하는 반면, 그 외의 경우는 덜 엄격한 규제가 적용된다. 인공지능이 사람의 권리·의무에 미치는 영향은 천차만별이므로 일률적인 규제보다는 위험도에 따른 차별화된 규제 접근이 필요하다. 그러나 현재 인공지능 기술이 초기 단계인 만큼, 금지된 인공지능을 규정하고 전면적으로 차단하는 것은 시기상조이며, 공공분야를 중심으로 국민의 권리·의무에 상당한 영향을 줄 수 있는 분야를 엄선하여 ‘고위험 인공지능’ 분야로 정하고, 낮은 수준의 설명의무 등 개발자·배포자에게 큰 부담이 되지 않는 규제를 도입하는 것이 적절하다고 본다.

생성형 인공지능을 별도로 규제할지 여부에 관하여 살펴본다. 생성형 인공지능은 새로운 콘텐츠를 창출할 수 있으며, 그로 인해 발생하는 위험 요소로서 딥페이크나 개인정보 침해 등이 구체화되고 있다. 생성형 인공지능은 그 범용성 때문에 많은 제품과 서비스의 기반 요소로 사용되며, 그로 인한 파급력이 매우 크다. 이에 EU 인공지능법은 생성형 인공지능 모델 제공자에게 특정 의무를 부과하는 별도의 규제를 도입하였다. 생성형 인공지능의 특성상 이를 별도로 규제하여 위험을 최소화할 필요가 있다.

연산능력에 따른 규제를 도입할지 여부와 관련하여, EU 인공지능법에서는 특정 수준 이상의 연산 능력을 가진 인공지능 모델을 고위험 성능으로 간주하고 추가적인 규제를 적용하고 있다. 이는 인공지능 모델이 성능이 높아질수록 그에 따른 위험성도 커지기 때문에 규제가 필요하다는 논리다. 그러나 인공지능의 성능이 하루가 다르게 향상되고 있어 성능 기준을 경직된 법률 형식으로 규정하는 것이 적절하지 않을 수 있으며, 연산 능력 외에 다른 요소들을 고려한 포괄적인 접근이 필요할 수 있다. 이에 따라 법률에서 연산능력에 따른 기준을 설정하기보다는 ‘고성능’ 등과 같이 불확정적인 개념을 도입하고 필요시

대통령령 등으로 기준으로 설정하고 의무를 부과하는 방안을 고려할 수 있다.

제공자와 배포자의 의무를 구분할 필요가 있는지 여부에 대하여, 인공지능 생태계에서 제공자와 배포자의 역할이 서로 다르기 때문에 이들을 구분하여 각각의 의무를 부과하는 것이 적절하다고 본다. 제공자는 인공지능 시스템의 개발을 담당하며, 시스템의 훈련 데이터에 대한 이해도가 높은 반면, 배포자는 시스템이 실제로 어떻게 사용되는지에 대한 정보를 더 잘 알고 있으므로, 각 역할에 적합한 의무를 도입하는 것이 필요하다.

마지막으로 거버넌스에 관한 논의가 있다. 인공지능의 다양한 적용 분야에서 중복되는 규제가 발생할 수 있으며, 이를 방지하기 위해 부처 간의 조정이 중요하다. 또한, 규제 집행의 실효성 확보와 인공지능 리스크 책임 분배를 해결하기 위해 컨트롤타워 역할을 할 수 있는 독립적인 조직의 필요성이 제기된다. EU 인공지능법에서는 인공지능 관련 정책 자문을 담당하는 과학 패널을 구성하고 있으며, 우리나라도 이와 같은 독립적이고 전문적인 기구를 설립하여 실질적인 영향력을 갖도록 해야 한다. 2024년 설립된 국가인공지능위원회는 인공지능 관련 정책을 조율하고 추진할 수 있는 권한을 가지고 있으나, 한시적인 대통령령에 근거하고 있어 지속적인 법적 지위와 권한을 확보할 필요가 있다고 평가된다.

나. 인공지능 저작권 입법 방향

인공지능 저작권 규제체계 입법 방향과 관련하여, 우리나라에 포괄적 공정이용 조항이 규정되어 있지만, 우리나라는 미국과는 달리 대륙법 체계를 채택하고 있어 포괄적 조항의 해석만으로 한계가 있으며, 도입된지 상당한 기간이 지났음에도 판례가 많이 축적되지 못한 상황이므로 이 조항만으로는 수범자들의 예측가능성이 저해될 수 있다. 아울러 TDM 면책 조항과는 그 입법 취지와 요건도 다르므로, 포괄적 공정이용 조항 외에 TDM 면책 조항을 별도

로 도입을 하는 것이 보다 직접적인 문제 해결 방식이라고 볼 수 있다.

영리 목적의 이용 시 TDM 면책을 허용할 것인지 여부에 관하여, 생성형 인공지능 개발자(제공자)는 통상 영리 목적으로 개발하며, 배포자 역시 영리 목적으로 이용하는 경우가 다수이며, 영리 목적과 비영리 목적의 구분이 불분명할 수 있으므로, 영리 목적인 경우에도 TDM 면책을 인정하되 옵트아웃 권리 인정이나 보상 인정 등을 통해 저작권자의 권리 인정을 통해 이해관계자 간 균형을 도모하는 방안을 고려할 수 있다.

저작권자가 자신의 저작물을 인공지능 학습에 사용하지 않도록 명시적인 표시(옵트아웃)를 한다면 TDM 면책을 인정하지 않는 방안을 고려할 수 있다. 다만, 이 경우 인공지능이 학습할 수 있는 저작물의 범위가 상당히 제한될 수 있으므로, 옵트아웃 표시 방법을 제한하는 방안, 저작권자가 옵트아웃을 하더라도 정당한 보상을 하면 면책되도록 하는 방안 등을 통해 이해관계자 간 균형을 찾아볼 수 있을 것이다.

통상 TDM 면책은 복제·전송에 한정하여 인정하도록 되어 있는데, TDM 면책은 저작재산권을 제한하는 규정으로서 저작권자에게 불측의 손해가 발생할 우려가 있는 점을 고려하면 그 범위를 확대하는 것은 바람직하지 않을 수 있다. 따라서 복제·전송에 한하여 인정하고, 추후 생성형 인공지능 학습 및 이용 상황을 보면서 면책 범위를 확대할지 여부를 결정할 필요가 있다.

로봇배제표준(robotx.txt)을 이용하여 크롤링을 차단하는 의사표시를 하였음에도 이를 우회하여 접근하였을 때 저작물에 대한 적법한 접근이라고 볼 수 있는지 문제된다. 저작권자의 권리를 확대한다면 그에 따라 적법한 접근의 범위를 넓게 인정하고, 반대로 저작권자의 권리가 상대적으로 축소된다면 적법한 접근의 범위를 좁게 인정하도록 하는 등, 저작권자의 권리를 함께 고려하여 유연하게 제도를 설계할 필요가 있다.

학습데이터 출처를 엄격하게 표시하도록 한다면, 인공지능 개발자의 부담

이 커지는 반면, 저작권자가 자신의 저작물이 무단으로 인공지능 학습에 사용되었는지 여부를 확인하기 위해 인공지능 학습 시 학습데이터 출처 표시가 되어야 할 필요가 있다. 이에 대하여, EU 인공지능법과 같이 인공지능이 학습한 데이터의 출처에 대한 요약 표시하도록 하는 방안을 고려할 수 있지만, 어느 정도로 상세하게 표시해야 할지에 대하여, 각 이해관계자의 사정을 종합적으로 고려할 필요가 있다.

복제물의 보관기간에 대해서는 독일의 사례를 참조하여 데이터마이닝 수행 후 해당 복제물을 삭제할 의무를 규정하는 것을 고려할 수 있다.

TDM 면책 시 보상 의무를 입법화 할 것인지 여부에 대하여, 주요국 TDM 면책 조항이나 포괄적 공정이용 조항에서 보상 의무를 규정하고 있지 않은 점을 고려할 필요가 있으며, 옵트아웃 권리 인정 여부, 영리 목적 이용 시 면책 여부 등 여러 쟁점들을 종합적으로 고려하여 제도를 설계할 필요가 있다.

3. 시사점

앞 절에서 제시한 입법 방향은 다음의 특징을 갖고 있다. 첫째, 이해관계자 간 이익형량을 도모하고 있다. 인공지능 개발을 저해하지 않으면서 이용자 권리를 보호하도록 하고 있으며, 인공지능 개발자·이용자와 저작재산권 간 균형을 도모하는 방안을 모색하였다. 둘째, 국제적·법체계적 정합성을 추구하고 있다. 국가별 접근법에 차이가 있음을 확인하고 수평적·수직적 규제의 장단점을 두루 검토하였다. 셋째, 인공지능 기술 및 산업 발전의 수준을 고려하고 있다. 인공지능 기술은 하루가 다르게 진화하고 있고 산업은 아직 초기 단계이므로 사전적으로 강한 규제보다는 적절하고 가벼운 규제를 적용한 후 기술·산업 발전 상황에 따라 유연하게 적용할 것을 제안하였다.

본 연구는 국내외 인공지능 규제체계 논의를 폭넓게 검토하였다. 이를 참고하여 신중하게 설계된 규제체계는 인공지능이 갖고 있는 내재적인 불확실성·오류와 인공지능의 강력한 성능에 대한 오·남용을 최소화함으로써 인공지능 기술에 대한 신뢰성을 제고하고 사회적 수용을 높일 수 있을 것이다. 아울러 이해관계자들 간 갈등을 줄여 기술·산업 발전을 저해하지 않으면서도 기존 산업 관계자와 국민의 권리를 보호할 수 있다. 국제적 규제체계 정합성을 통해 우리나라 기업의 글로벌 진출도 지원할 수 있다.

특히 인공지능 분야는 강한 규제가 산업의 혁신을 저해할 우려가 크며, 이는 규제 중심의 EU 법제와 자율 중심의 미국 법제에서 어느 정도 확인할 수 있었다. 다만 EU도 규제샌드박스과 같은 혁신 지향적 조항을 두고 있다는 점, 미국의 연방 법률안과 각 주의 법률에서도 EU 방식의 강한 규제가 일부 반영되어 있다는 점이 나타나서 규제와 혁신 사이의 균형이 인공지능 입법의 궁극적인 지향점이 될 수 있음을 알 수 있다. 우리나라의 경우 아직 성숙하지 않은 인공지능 산업에 대하여 과도한 규제를 지양함으로써 궁극적으로 기술과 산업의 발전을 견인할 수 있을 것이다.

참고문헌

국내문헌

- 고학수·임용·박상철, 「유럽연합 인공지능법안의 개요 및 대응방안」, 서울대 인공지능정책 이니셔티브 DAIG 제2호, 2021.
- 구문모, 「[영국] 영국 정부, ‘모든 목적’의 텍스트 및 데이터마이닝 허용 저작권법 개정 계획 철회 결정」, 한국저작권위원회, 저작권 동향, 2023. 2.
- 김건우, 「인공지능 규제 거버넌스의 현재와 미래」, 파이돈, 2023.
- 김윤명, 「생성형 AI와 저작권 현안」, KISDI AI Outlook, 정보통신정책연구원, 2023. 6.
- 김창화, 「생성형 AI를 둘러싼 최근 저작권 분쟁 동향과 시사점」, 규제법제리뷰 제24-1호, 2024. 2.
- 남구현, 「인공지능 규율에 있어 위험기반 접근 방식의 적합성 고찰」, 법학논총 제36권 제1호, 2023
- 라기원, 「유럽연합 인공지능법(EU AI ACT)의 특성과 쟁점-우리나라 인공지능 입법에 대한 시사점」, 한국법제연구원, 2024.
- 류시원, 「저작권법상 텍스트·데이터 마이닝(TDM) 면책규정 도입 방향의 검토」, 선진상사법률연구 제101호, 2023.
- _____, 「인공지능 시대 저작권 정책 형성절차에 관한 제언- 텍스트·데이터 마이닝 예외규정 입법 논의를 중심으로 -」, 법제처, 2024.3.
- 박상철, 「인공지능의 법적 규율 I : 범용모델과 생성모델」, 서울대학교 법학 제65권 제1호, 2024. 3. 105쪽
- 박정훈 외, 「[EU, 일본] 유럽연합과 일본의 TDM 규정과 그 적용 범위 — AI

- 학습을 위한 저작물 이용과 관련하여」, 한국저작권위원회, 이슈리포트, 2024. 7.
- 백지연, 「[중국] 베이징인터넷법원, 생성형 AI 산출물의 저작권을 인정」, 한국저작권위원회, 저작권 동향 2023 제11호, 2023. 12.
- 법제처 미래법제혁신기획단, 「인공지능(AI) 관련 국내외 법제 동향」, 법제소식, 2024. 7.
- 송선미, 「AI 창작물의 저작권 보호에 관한 해외 동향」, 한국저작권위원회, 이슈리포트 2022-02. 2022. 3.
- _____, 「[호주] AI 저작권 관련 호주 정부의 대응 현황」, 한국저작권위원회, 저작권 동향, 2024. 2
- 심소연, 「규제 중심의 유럽연합 인공지능법(EU AI ACT)」, 국회도서관 최신 외국입법정보, 2024.
- 오유빈, 「미국의 인공지능 입법 현황과 바이든 행정부의 행정명령」, 국회도서관, 2024. 1.
- 유재홍 외, 「글로벌 AI 신뢰성 정책 동향 연구」, 연구보고서 RE-150, 소프트웨어정책연구소, 2023.
- 유혜정, 「영국 지식재산청, 텍스트와 데이터마이닝(TDM) 관련 저작권법 개정 추진」, 한국저작권위원회, 저작권 속보, 2022. 7. 6.
- 이근우, 「해외 주요국의 AI 규제 거버넌스 현황 및 시사점 - 영국의 AI 규제 거버넌스 현황」, 지능정보사회 법제도 이슈리포트(2024-01), 2024. 8.
- _____, 「Andy Warhol 케이스의 ‘변형적 이용(Transformative Use)’의 해석과 AI 학습데이터 이용에 대한 영향」, 한국저작권위원회, 이슈리포트 2023-8, 2023. 11.
- _____, 「미국 Copyright Office의 AI 저작권 쟁점 의견 청취」, 한국저작권위원회, 이슈리포트, 제2023-11호, 2023. 12.

- _____, 「프랑스, ‘AI와 저작권 관련 지식재산권법’ 개정안 발의」, 한국저작권위원회, 저작권 동향, 2023. 12.
- _____, 「AI 저작권 관련 Sarah Andersen v. Stability AI 소송 분석」, 한국저작권위원회 이슈리포트 2023-15, 2023. 12.
- _____, 「음악 생성 AI 모델에 대한 저작권 소송과 학습데이터 이용의 공정이용 여부 주장 분석」, 한국저작권위원회, 이슈리포트 2024-24, 2024. 8.
- 이철남, 「[싱가포르] 인공지능과 지식재산법의 교차에 관한 이슈 보고서 발표」, 한국저작권위원회, 저작권 동향, 2024.
- 이효진, 「우리 ‘인공지능법 제정안’의 인공지능 규율 방향에 관한 소고」, 법학논문집 제47집 제2호(중앙대학교 법학연구원), 2023. 8.
- 이혜영, 「미국저작권청(USCO), 인공지능이 만든 저작물 등록 거절」, 한국저작권보호원, 해외 저작권 보호 동향, 2022. 7. 14.
- 전응준, 「AI 관련 저작권, 개인정보, 규제 쟁점 및 입법적 과제」, 국회입법조사처, 2024 국가비전 입법정책 컨퍼런스, 2024. 4.
- 정준화, 「정보통신망 이용촉진 및 정보보호 등에 관한 법률 일부개정법률안(의안번호 2126048) 입법영향분석- 인공지능으로 만든 가상 정보임을 표시하도록 하는 의무」, 국회입법조사처, NARS 입법영향분석 기획보고서 제7호, 2024. 8. 23.
- 정지은, 「해외 주요국의 AI 규제 거버넌스 현황 및 시사점 - 일본의 AI 규제 거버넌스 현황」, 지능정보사회 법제도 이슈리포트(2024-01), 2024. 8.; KoTRA, 일본의 AI 정책과 실제 사례, 2024. 4. 4.
- 조원영 외, 「인공지능 신뢰체계 정립방안 연구」, 연구보고서 RE-123, 소프트웨어정책연구소, 2023.
- 차상욱, 「저작권법상 인공지능 학습용 데이터셋의 보호와 쟁점」, 경영법률 제32권 제1호, 2021.

- _____, 「캐나다, 생성형 AI와 저작권에 대한 협의 관련 의견수렴 내용 공개」, 한국저작권위원회, 저작권 동향, 2024. 8.
- 채은선, 「디지털 법제 Brief」, 한국지능정보사회진흥원, 2023. 1. 31.
- _____, 「美, AI 행정명령(2023.10.30.)의 주요 내용 및 이행 현황」, 디지털 법제 Brief, 2024. 7.
- 최경석, 「생명윤리에서의 법, 도덕 및 윤리의 역할과 한계」, 이화여자대학교 법학논집 제15권 제4호, 2011. 6.
- 최경진·이기평, 「AI 윤리 관련 법제화 방안 연구」, 글로벌법제전략 연구, 2021. 10.
- 최상필, 「빅데이터의 분석과 활용을 위한 TDM의 저작권적 쟁점과 입법론」, 계간 저작권, 2022.3.
- 한국정보산업연합회, 「인공지능기본법 입법 추진현황 및 산업진흥 측면에서 본 이슈」, 2024. 7.

외국문헌

- Alex Engler, 「The EU AI Act will have global impact, but a limited Brussels Effect」, Brookings Research」, 2022. 6. 8.
- _____, 「The EU and U.S. diverge on AI regulation: A transatlantic comparison and steps to alignment」, Brookings Research, 2023. 4. 25.
- Edward D. Lanquist and Dominic Rota, 「The Impact of the Supreme Court's 'Goldsmith' Decision on Copyright Enforcement Against AI Tools」, Baker Donelson, 2023. 9. 26.
- Gary Myers, 「Artificial Intelligence and Transformative Use After Warhol」, Washington and Lee Law Review Online, Volume 81, Issue 1, 2023. 12.

- Huw Roberts et al, 「A comparative framework for ai regulatory policy」, 2023.
- Nicholas Carlini et al., 「Extracting Training Data from Diffusion Models」, arXiv: 2301.13188, 2023.
- Suzan Slijpen, Mauritz Kop & I. Glenn Cohen, 「EU and US Regulatory Challenges Facing AI Health Care Innovator Firms」, 2024. 4. 6.
- Trina Ha et al, 「When Code Creates: A Landscape Report on Issues at the Intersection of Artificial Intelligence and Intellectual Property Law」, Singapore Management University, 2024. 2. 28.
- Ursula von der Leyen, 「“A Union that strives for more, My agenda for Europe By Candidate for President of the European Commission”, Political Guidelines for the Next European commission 2019-2024」, 2018, p. 13,

홈페이지

- 미국 백악관 홈페이지 <<https://www.whitehouse.gov/> 최종방문일 2024. 9. 11.>
- 미국 연방관보 홈페이지 <<https://www.federalregister.gov/> 최종방문일 2024. 9. 11.>
- 미국 의회 홈페이지 <<https://www.congress.gov/> 최종 방문일 2024. 9. 11.>
- 미국 저작권 사무소(U.S. Copyright Office) 홈페이지<<https://www.copyright.gov/#> 최종 방문일 2024. 9. 11.>
- 영국 정부 홈페이지 <<https://www.gov.uk/> 최종 방문일 2024. 9. 11.>
- 유럽연합 집행위원회 홈페이지 <https://commission.europa.eu/index_en 최종 방문일 2024. 9. 11.>
- 일본 AI전략회의(AI戦略) <<https://www8.cao.go.jp/cstp/ai/index.html> 최종방문일 2024. 9. 11.>
- 일본 AI검토회(AI時代の知的財産権検討会)<https://www.kantei.go.jp/jp/singi/titeki2/ai_kentoukai/kaisai/index.html 최종방문일 2024. 9. 11.>

중화인민공화국 중앙인민정부 홈페이지<<https://www.gov.cn/> 최종방문일 2024. 9. 11.>
캐나다 정부 홈페이지 <<https://www.canada.ca/en.html> 최종 방문일 2024. 9. 11.>
캘리포니아 법령 정보 홈페이지 <<https://leginfo.legislature.ca.gov/faces/home.xhtml>
최종 방문일 2024. 9. 11.>
코트리스너 홈페이지 <<https://www.courtlistener.com/> 최종 방문일 2024. 9. 11.>
저스티아커넥트 홈페이지 <<https://connect.justia.com/> 최종방문일 2024. 9. 11.>

기타

과학기술정보통신부 보도자료, “법정부 역량을 결집하여 AI 시대 미래 비전과 전략을 담은‘AI 국가전략’발표”, 2019. 12. 17.
과학기술정보통신부 보도자료, “인공지능 시대를 준비하는 법·제도·규제 정비 로드맵 마련”, 2020. 12. 23.
과학기술정보통신부 보도자료, “대한민국이 새로운 디지털 규범질서를 전 세계에 제시합니다!”, 2023. 9. 25.
문화체육관광부 보도자료, “인공지능(AI)-저작권 안내서 발표로 시장의 불확실성 해소하고, 안무·건축 등 ‘저작권 사각지대’ 없앤다”, 2023. 12. 26.
문화체육관광부 보도자료, “인공지능과 저작권 쟁점별 구체적 정책 방안 마련한다”, 2024. 2. 16.

국회 문화체육관광위원회, 콘텐츠산업 진흥법 일부개정법률안(강유정의원 대표발의, 의안번호 제2200048호) 검토보고서, 2024. 8.
국회 문화체육관광위원회, 「저작권법 전부개정법률안(도종환의원 대표발의, 의안번호 제2107440호) 검토보고서」, 2021. 2.

- 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이용호의원 대표발의, 의안번호 제2117990호) 검토보고서」, 2022. 12.
- 국회 문화체육관광위원회, 「저작권법 일부개정법률안(황보승희의원 대표발의, 의안번호 제2122537호) 검토보고서」, 2023. 9.
- 국회 문화체육관광위원회, 「저작권법 일부개정법률안(이인영의원 대표발의, 의안번호 제2124685호) 검토보고서」, 2023. 11.
- 애플경제, “챗GPT 계기, 데이터 개방 위한 ‘TDM’ 부각”, 2023. 3. 27. <<https://www.apple-economy.com/news/articleView.html?idxno=71167> 최종방문일 2024. 9. 10.>
- 김우균 외, “생성형 AI 저작권 관련 중국 판례 동향”, 법률신문, 2024. 6. 17. <<https://www.lawtimes.co.kr/LawFirm-NewsLetter/199153> 최종 방문일 2024. 9. 10>|
- Ingrid Stevens, “Regulating AI: The Limits of FLOPs as a Metric”, 2024. 5. 1. <<https://medium.com/@ingridwickstevens/regulating-ai-the-limits-of-flops-as-a-metric-41e3b12d5d0c> 최종 방문일 2024. 9. 2.>
- Stuart D. Levi et al, “Colorado’s Landmark AI Act: What Companies Need To Know”, Skadden Publication, 2024. 6. 24.; Colorado General Assembly <<https://leg.colorado.gov/bills/sb24-205> 최종 방문일 2024. 9. 1.>
- Yi Wu, “Interim Administrative Measures for Generative Artificial Intelligence Services”, China Briefing, 2023. 7. 27. <<https://www.china-briefing.com/news/how-to-interpret-chinas-first-effort-to-regulate-generative-ai-measures/> 최종방문일 2024. 9. 11.>