

인공지능 단일법 제정의 방향성

- 유럽연합, 미국과의 비교를 중심으로 -

▶ 목 차 ◀

- | | |
|----------------------------|----------------------|
| I. 서론 | IV. AI 단일법 제정 방향 제언 |
| II. 국내 AI 관련 법제 현황 | 1. 22대 국회 인공지능 법안 검토 |
| 1. AI 개념 | 2. 제언 |
| 2. AI 관련 기존 법제 | V. 결론 |
| III. 유럽연합(EU)과 미국의 인공지능 법제 | |
| 1. 유럽연합(EU) | |
| 2. 미국 | |
| 3. 소결 | |

【국문초록】

인공지능(AI) 기술의 눈부신 발전에 따라 범세계적으로 법제도적 대응의 중요성이 대두되고 있다. 특히 화제가 되고 있는 인공지능의 잠재적 위험성과 그로 인한 사회적, 경제적 영향을 고려하였을 때, 국내에도 AI 단일법의 제정이 필요하다는 주장이 제기되고 있다. 본 연구는 이러한 문제의식을 바탕으로 유럽연합(EU)과 미국의 인공지능 관련 법제도를 비교·분석하고, 우리나라의 AI 단일법 제정 방향성을 제언하고자 한다. EU와 미국 법제의 비교를 통해 각 법제도의 장단점을 도출하고, 우리나라에 적합한 법적 대응 방안을 모색하는 것이 본 연구의 목적이다. 본 연구에서는 크게 ① 규제와 혁신의 균형, ② 유연한 법률 체계 구축, ③ 국제 표준과의 조화, ④ 산업계와의 협력 강화의 관점에서 AI 기술의 발전과 안전한 활용을 위한 구체적인 법적 조치를 제안한다. 본 연구는 AI 단일법 제정의 필요성을 강조하며, 법제도적 공백을 해소하고 사회적 신뢰를 구축하기 위한 실질적인 방향성을 제시하는 것에 그 목적이 있다.

주제어(Keyword) : 인공지능 단일법, 유럽연합(EU) AI ACT, 미국 AI 법제, 인공지능 법안, 규제와 혁신

I. 서론

인공지능(AI)의 급속한 발전은 다양한 산업과 일상에 변화를 일으키고 있으며, 이에 대응하여 세계 각국은 국가 차원에서 전략을 수립하고 과감한 투자와 규제 개선을 추진하는 한편, 기술 및 인재 교류와 공통된 표준 마련을 위해 국제적 협력을 이어가고 있다.¹⁾ 유럽연합(EU)은 2024년 3월 AI법을 제정했으며, 미국 역시 AI 관련 법 제정을 추진 중이다.

AI 법 제정에서의 핵심은 바로 규제에 대한 논의이다. 데이터 수집 단계에서의 개인정보 침해 문제와 생성형 AI의 저작권 침해 문제가 부각되면서 규제의 필요성은 더욱 중요해졌다. 그러나 지나치게 엄격한 규제는 기술 혁신을 저해할 수 있다는 우려도 제기되고 있다. 최근 메타(구 페이스북)와 애플이 강력한 규제가 시행되고 있는 유럽 시장에서 일부 서비스 철수를 고려하거나 실제로 사업을 축소한 사례는²⁾ 과도한 규제가 기업의 혁신을 억제하고 시장 진출에 어려움을 줄 수 있음을 보여준다.

이러한 배경 속에서 우리나라도 AI 관련 법 논의가 활발히 진행되고 있다. 21대 국회에서 7개의 법안이 발의된 데 이어, 22대 국회에서는 개원 이후 현재까지 6개의 인공지능 관련 법안이 발의되었다. 우리나라 AI 기술은 아직 발전 단계에 있기 때문에, 법 제정 초기 단계에서의 제정 방향성 설정은 AI 기술 발전의 초석을 다지는 중요한 사안이다. 따라서 법안 논의 초기 단계에서 규제의 방향성을 올바르게 설정하기 위한 관련 연구는 장기적인 관점에서 반드시 필요하다.

본 논문에서는 우리나라의 인공지능 관련 법령을 검토한 후, EU와 미국의 AI 법 제정 현황을 비교·분석할 것이다. 이를 바탕으로 22대 국회의 인공지능 법률안의 주요 내용을 검토하고, AI 기술 발전을 촉진하면서도 사회적 책임을 담보할 수 있는 완화된 방향성의 법안을 중심으로 제안해보고자

1) 김송옥, “AI 법제의 최신 동향과 과제 -유럽연합(EU) 법제와의 비교를 중심으로-”, 한국비교공법학회, 『공법학연구』 제22권 제4호, 2021, 1면.

2) 이상덕, “골치아픈 유럽, 차라리 안할래”...메타·애플, 새 모델 출시 포기, 매일경제, 2024. 7. 18. <https://www.mk.co.kr/news/it/11070826> (2024. 9. 1. 최종 접속)

한다.

II. 국내 AI 관련 법제 현황

1. AI 개념

AI는 인간의 지능인 학습, 추론, 인지, 논증 등을 가진 컴퓨터 시스템 또는 이를 연구하는 기술과 과학을 이르는 말로, 쉽게 말해 인간의 지능을 모방한 기술을 뜻한다. AI는 인간의 관여 없이 스스로 추론하며 알고리즘을 생성할 수 있고, 경험을 통해 학습한 내용을 바탕으로 스스로 문제를 해결할 수 있다.³⁾ 기존에는 지적 능력을 지닌 인간만이 수행할 수 있었던 작업들을 인간보다 훨씬 빠르고 정확하게 수행할 수 있기에, AI 기술의 발전은 산업의 가치를 높이는 혁신을 가져올 것이며, 우리의 일상에도 다양한 영향을 줄 것이다. 이러한 AI 기술의 예시로 2022년 OpenAI에서 출시한 초거대 생성형 AI 서비스인 ChatGPT가 있다. ChatGPT는 의미 있는 응답을 생성하는 대화형 인공지능으로, 두 달 만에 이용자 수 1억 명을 돌파하며 전 세계에 큰 파장을 불러일으켰다. ChatGPT의 출시 이후 AI는 더 급격한 속도로 발전하고 있다. 현재 다양한 빅테크 기업들이 대규모 언어 모델과 AI 서비스의 출시 계획을 발표하였으며, 해외뿐만 아니라 국내 기업에서도 한국어 기반의 초거대 언어 모델의 개발과 AI 서비스를 적극적으로 추진하고 있다.⁴⁾

2. AI 관련 기존 법제

AI 기술의 발전이 빠르게 이루어지는 현재 추세에 맞추어, 전 세계적으

3) 이대회, “인공지능에 의한 개인정보의 자동 처리에 대한 투명성 확보에 관한 연구”, 한국법학원, 「저스티스」 제200권 제200호, 2024, 299면.

4) KISTEP 과학기술정책센터, “생성형 AI 관련 주요 이슈 및 정책적 시사점”, 2023, 1면.

로 인공지능을 규범적으로 규율하려는 움직임이 보이고 있다. 그 예로 2024년 3월 유럽연합(EU)이 제정한 AI법이 있다. 현재 우리나라에서도 AI 법 제정에 관한 논의가 진행되고 있으나, 아직 AI법과 같은 인공지능법이 마련되지는 않은 상태이다.

AI 법 제정의 필요성과 관련한 논의에 앞서, 우리나라의 AI 관련 기존 법제를 살펴볼 필요가 있다. 이와 관련하여 행정기본법, 개인정보 보호법, 신용정보의 이용 및 보호에 관한 법률 순으로 살펴보겠다.

현행 법령상 AI 육성에 관하여 직접적인 규범은 없지만, AI의 활용과 관련된 일반적인 규정은 행정기본법에 마련되어 있다. 우리나라는 행정기본법 제20조를 통해 인공지능 기술을 적용한 시스템을 포함한 완전히 자동화된 시스템으로 처분을 할 수 있다고 규정함으로써, AI의 자동적 처분을 인정하고 있다.⁵⁾ 이를 통해 자동적 처분에 관한 입법적 기준을 마련하여, AI에 의한 기본권 침해 소지를 사전에 예방하였다.

AI의 가장 핵심적인 요소는 데이터이기에, AI 산업 발전을 위해 대량의 데이터 수집은 필수적이다. AI가 개인정보가 포함된 대량의 데이터를 수집하고 처리하는 과정에서 정보 주체 개개인의 동의를 받는 것은 쉽지 않은데, 이와 관련하여 개인정보보호법은 정보 주체가 아닌 곳으로부터 수집한 개인정보에 대한 의무를 명시하여 정보 주체의 권리를 보장하고 있다. 먼저, 개인정보보호법 제20조를 통해 정보 주체가 아닌 곳으로부터 수집한 개인정보를 처리하는 경우, 정보 주체의 요구가 있으면 개인정보의 수집 출처와 처리 목적, 개인정보 처리의 정지를 요구하거나 동의를 철회할 권리가 있다는 사실을 정보 주체에게 알려야 함을 명시하고 있다.⁶⁾ 또한, 개인정보보호법 제37조의2에 의해 정보 주체는 AI 기술을 적용한 시스템을 포함한 자동화된 시스템으로 개인정보를 처리하여 이루어지는 결정이 자신의 권리 또는 의무에 중대한 영향을 미치는 경우 그 결정을 거부할 수 있는 권리를 가진다고 명시하고 있다.⁷⁾

5) 행정기본법 제20조.

6) 개인정보보호법 제20조.

7) 개인정보보호법 제37조의2.

신용정보의 이용 및 보호에 관한 법률 또한 개인정보의 이용에 관한 내용을 명시하고 있다. 개인정보보호법이 정보 주체의 권리를 보장하는 내용을 담고 있다면, 신용정보의 이용 및 보호에 관한 법률에는 신용정보 주체의 동의가 있었다고 객관적으로 인정되는 범위 내에서, 사회관계망 서비스 등에 직접 또는 제3자를 통해 공개한 정보에 해당하는 개인 신용정보의 경우 정보 주체 동의를 면제하도록 하는 내용을 담고 있다. 신용정보의 이용 및 보호에 관한 법률 제15조에 의하면 신용정보회사 등이 개인 신용정보를 수집할 때에는 신용정보 주체의 동의를 받아야 한다. 다만, 항상 신용정보 주체의 동의를 받아야 하는 것은 아니며, 신용정보 주체가 스스로 사회관계망 서비스 등에 직접 또는 제3자를 통하여 공개한 정보의 경우는 이미 해당 신용정보 주체의 동의가 있었다고 본다.⁸⁾

이처럼 현재 우리나라에는 행정기본법, 개인정보 보호법, 신용정보의 이용 및 보호에 관한 법률 등을 통해 AI와 관련한 내용을 규정하고 있다. 다만, AI 기술의 급진적 발전에 대응하기에는 기존의 법에 존재하는 허점이 많다는 우려의 목소리가 커지고 있다. 인공지능이 지닌 잠재적 위험을 사전에 예방하고 인공지능의 안전한 사용을 보장하기 위해서는 우리나라도 AI법 제정이 필요할 것이다.

Ⅲ. 유럽연합(EU)과 미국의 인공지능 법제 현황

1. 유럽연합(EU)

유럽의 인공지능 법제는 기본권 침해 상황을 우려하여, 법 조항을 통해 명확히 규정하고 있다.

유럽은 2018년 4월 25일, 「정책추진안: 유럽을 위한 AI(Communication:

8) 신용정보의 이용 및 보호에 관한 법률 제15조.

Artificial Intelligence for Europe)」에서 인공지능에 대한 규제 논의를 시작하였다. 신뢰할 수 있는 AI를 위한 윤리 가이드라인, AI 백서 등을 거쳐 2021년 4월에는 세계 최초의 인공지능 법안, AI Act(이하 AI 법)를 발의하여, 도입을 앞두고 있다.

AI 법은 유럽뿐만 아니라 유럽에 영향을 미치는 모든 대상에 적용되는 법률이다. AI 법의 목적은 다음과 같다. 첫째, 통일된 법체제로 역내 시장을 개선한다. 둘째, 신뢰할 수 있는 인공지능 활용을 촉진하고, 인공지능 시스템의 유해한 영향으로부터 기본권을 높은 수준으로 보호한다.⁹⁾ 인공지능이 빠르게 진화하는 기술 분야라는 점을 고려하여, 기존의 GDPR과 별개로 인공지능을 전담하여 적용하는 법이다.

AI 법은 리스크 규제 방식을 따른다. 리스크 기반 규제 방식이란, 위험 정도에 따라 규제 정도를 달리하는 방식으로, 인공지능 분야에서는 크게 고위험 인공지능과 일반 인공지능으로 분류한다. AI 법의 경우, ① 용인할 수 없는 위험(unacceptable risk), ② 고위험(high risk), 그리고 ③ 낮은 위험 혹은 최소 위험(low or minimal risk)으로 세분하여 인공지능 규제를 달리한다.¹⁰⁾

각 위험도에 따른 인공지능 시스템을 요약하면 다음과 같다.

<표 1> AI 법의 위험도에 따른 인공지능 시스템 분류

위험도	내용	조항
용인할 수 없는 위험 (unacceptable risk)	그 자체로 금지된 인공지능을 뜻하며, 인간에게 상당한 피해를 유발할 수 있거나 편향을 생성할 수 있는 인공지능 시스템	제5조

9) Artificial Intelligence Act, 제1조.

10) European Commission, “Explanatory Memorandum, Proposal for a Regulation of the European Parliament and of the Council: Laying Down Harmonised Rules on Artificial Intelligence and Amending Certain Union Legislative Acts”, COM(2021) 206 final, p. 12.

	1. 사회적 점수 매기는 인공지능 2. 자연인 범죄 행위 가능성 예측 인공지능 3. 사람의 의사결정 왜곡 등으로 중대한 해를 미치는 것이 목적인 인공지능 4. 피해를 유발하는 취약성 탐지가 목적인 인공지능 5. CCTV 등으로 얼굴 인식 데이터베이스를 만드는 인공지능 위의 예시처럼 기본권을 침해하는 인공지능은 금지됨.	
고위험 (high risk)	고위험 단계의 인공지능은 제한적으로 허용된 인공지능 시스템을 뜻함. 인간의 의사결정 결과에 실질적 영향이 없거나, 기본권 위험이 중대하지 않은 경우는 고위험 인공지능으로 분류하지 않음. 1. 생체 인식 시스템(목적에 따라, 사용 방식에 따라 달라질 수 있음) 2. 시민의 생명에 직결된 중요 인프라를 관리하는 인공지능 시스템 3. 학습 평가, 채용 평가 등 평가에 사용되는 인공지능 시스템 4. 필수 민간 서비스, 공공 서비스 등 혜택 부여 적격성 판단 인공지능 시스템 5. 법 집행 목적 인공지능 시스템 6. 사법 행정, 민주적 절차 침해 할	제6조, 제7조

	수 있는 인공지능 시스템 7. 제품의 안전 구성요소로 사용되는 인공지능 시스템	
	필수 요건은 ① 위험 관리 시스템 구축, ② 데이터 거버넌스, ③ 기술 문서 작성 및 업데이트, ④ 로그 기록 자동 보관, ⑤ 투명성 및 정보 제공, ⑥ 인간의 관리 감독, ⑦ 정확성/견고성/사이버 보안	제8조~ 제15조
낮은 위험 (low or minimal risk)	제한적 위험 단계의 인공지능은, 딥페이크나 챗봇과 같이 기본권을 침해할 우려는 있으나 중대하지 않은 인공지능 시스템이다. 이러한 인공지능에는 투명성 의무만 부과된다.	제69조

유럽은 다음과 같이 인공지능을 위험도에 따라 명확히 분류하고, 일정 수준 이상의 위험이 되는 인공지능에 대해서는 엄격한 규제를 취하고 있다. 그러나, AI 법 제6장 혁신 지원 조치에서는 인공지능의 발전을 위해 규제 정도를 경우에 따라 낮추는 조항이 있다. 가령, 특정 조건을 만족한 경우, 규제 샌드박스를 활용하여 비교적 자유롭게 고위험 인공지능을 실험할 수 있다.¹¹⁾

한편, AI 법에서는 기존의 기관을 활용하여 인공지능을 관리하는 것이 아니라, 인공지능 전담 기관을 이용하여 관리하게 규정하고 있다.¹²⁾

AI 법의 벌칙은 대부분 벌금으로 규정되어 있다. 기업의 경우, 최대 3,500만 유로 또는 전 세계 연간 매출의 7% 중 더 큰 금액의 벌금을 부과받을 수 있다.¹³⁾ 다만, 위반의 유형과 심각성에 따라 벌금의 크기는 달라질 수 있다.¹⁴⁾ 뿐만 아니라, 유럽연합 단체 등에 대한 과징금, 범용 인공지능 제공

11) 자세한 사항은 Artificial Intelligence Act, 제57조-제61조 참고.

12) 자세한 사항은 Artificial Intelligence Act, 제7장 참고.

13) Artificial Intelligence Act, 제99조 3.

자에 대한 과징금도 규정되어 있다.¹⁵⁾ 벌금을 제외하고, 운영 중단이나 데이터 삭제 등이 언급되어 있다.¹⁶⁾

EU는 세계적인 추세와 비교하더라도 굉장히 명확하게 구체적으로 규제를 하고 있기 때문에 규제의 강도에 대한 비판이 있다. 스탠퍼드 대학의 ‘인간중심 인공지능 연구소(HAI)의 보고서 ‘파운데이션 모델 공급자들은 EU AI 법 초안을 준수하는가?’에서는 ‘주요 파운데이션 모델 기업 대부분은 AI 법을 준수하지 못할 것’이라고 언급하였다. 특히, EU의 투명성 부분을 언급하며, ‘파운데이션 모델 제공 기업 대부분은 데이터 및 연산 기능의 주요 특성에 대한 정보를 제공하지 않고 있다.’라고 밝히기도 했다.¹⁷⁾ 또한, IBM CEO 크리슈나는 2024 밀컨 컨퍼런스에서 AI 법의 규제가 지나치다는 취지의 입장을 밝히며, 인공지능 기술 규제가 아니라, 인공지능 기술의 악용 규제에 초점을 맞춰야 한다고 비판했다.¹⁸⁾

2. 미국

앞서 논의된 유럽의 AI 법제와 비교해보았을 때, 미국은 이제껏 AI 관련 법에 있어 기술의 발전과 혁신을 지향하며 보다 유연한 방향의 접근 방식을 채택하였고 이는 미국이 세계적인 AI 선도국가로 자리매김하는 것의 중요한 발판이 되었다. 미국의 AI 법제화 과정은 2020년부터 본격적으로 시작된 여러 법안, 법률과 행정명령의 발표를 통해 구체화되었으며, 2023년도 중후반에 들어서면서 AI의 사회적 책임성과 윤리적 사용에 대한 국제적 요구가 증가함에 따라 미국 또한 유럽과 마찬가지로 강력한 법적 규제와 구체적

14) 자세한 사항은 Artificial Intelligence Act, 제99조 참고.

15) 자세한 사항은 Artificial Intelligence Act, 제100조, 제101조 참고.

16) AI 법 제12장 벌칙에는 언급되어 있지 않으나, 고위험 인공지능을 다루는 부분 등에서 언급되고 있다.

17) 홍윤지, “스탠퍼드大 “초거대 AI시스템 모델, EU규제 충족 못할 것”, 법률신문, 2023. 07. 06. <https://www.lawtimes.co.kr/news/188993> (2024. 9. 1. 최종 접속)

18) 박신영, “AI혁신 꺾을라…과도한 규제보다 부족한 게 낫다”, 한경신문, 2024. 05. 07. <https://www.hankyung.com/article/2024050718931> (2024. 9. 1. 최종 접속)

인 표준을 마련하고 있는 추세다.¹⁹⁾ 본 절에서는 미국의 주요 AI 관련 법안들과 정책들을 살펴보며 2020년부터 2024년 현재까지의 방향성을 분석해 보고자 한다.

AI 법제화의 출발점이라고 볼 수 있는 「미국 국가 인공지능 구상법 2020(National Artificial Intelligence Initiative Act of 2020)」이 포함하는 대다수의 조항은 미국의 AI 연구 및 개발을 지원하고 촉진시키는 것이 주 목적이었다.²⁰⁾ 예컨대, 해당 법률의 제5101조 b항에서는 계획을 수행할 때 계획사무국, 기관연합위원회 및 대통령이 적절하다고 판단하는 기관장을 통해 보조금, 협력계약, 시험대 및 데이터와 컴퓨팅 자원에 대한 접근을 포함한 인공지능 연구 개발에 대한 지속적이고 일관적인 지원을 수행하여야 한다고 명시되어 있다.²¹⁾ 또한 제5103조 d항 2호에 명시되어 있듯, 연방정부는 AI 연구개발 및 시연 분야에서 지도력과 투자가 필요하다고 판단되는 영역을 결정하고 이에 우선순위를 부여하여 자금이나 데이터 시설 등을 지원하여야 한다.²²⁾ 이렇듯 위 법률은 AI가 공익을 위해 활용될 수 있도록 하는 제도적 장치를 마련하는 것에 그 초점을 두었다.

그러나 이를 근거로 AI 시대 초기의 미국이 규제의 필요성을 완전히 배제하였다고 보기는 어렵다.²³⁾ 예컨대, 「생성적 적대 신경망 출력물 확인법(Identifying Outputs of General Adversarial Networks Act, 2020)」은 AI 기술, 개중에서도 생성형 AI 기술을 구체적으로 정의하며 이에 대한 규

19) 법제처 미래법제혁신기획단, “인공지능(AI) 관련 국내외 법제 동향”, 2024, 3면.

20) 법무법인 화우, “생성형 AI의 국제 규제 동향 -미국, EU, 중국을 중심으로 한 규제 동향 훑어보기”, 2024, 2면.

21) National Artificial Intelligence Initiative Act, SEC. 5101. National Artificial Intelligence Initiative. (b).

22) National Artificial Intelligence Initiative Act, SEC. 5103. Coordination by Interagency Committee. (d). (2).

23) EU 집행위원회는 2021년 발표한 AI Act의 초안을 통해 고위험 AI 프로그램에 대한 명확한 규제 프레임워크를 설정하였는데, 이러한 접근은 미국에도 영향을 미쳐 비슷한 형태의 규제 논의 진행으로 이어졌다. 예컨대 미국의 여러 AI 관련 법안에서 AI 시스템의 투명성, 소비자 보호, AI 기술의 안전한 사용 등 AI 기술의 적절한 규제와 관련된 항목에서 EU의 AI Act 초안과 유사한 조항들을 포함하였음을 확인할 수 있다 (구체적인 사항은 “The EU and U.S. are starting to align on AI regulation” 참고).

제를 중점적으로 다룬다. 해당 법률에서는 생성적 적대 신경망의 기능과 결과물 및 콘텐츠를 합성, 조작하는 기술에 대한 자발적 표준의 개발을 위해 공공 및 민간 부문의 참여 가능성을 고려하여야 한다는 점을 언급한다.²⁴⁾ 또한 조작·합성된 콘텐츠 탐지를 위한 보고서를 제출하도록 규정하고 있다.²⁵⁾ 해당되는 보고서에는 디지털 미디어 기업과 같은 민간부문과 협력하여 생성적 대립 신경망의 기능 및 출력물을 탐지하고 규제할 수 있는 연구 기회의 가능성을 평가하는 내용이 포함된다. 언급된 조항의 공통점은 모두 생성형 AI 기술이 초래할 수 있는 사회적 위험을 인지하고, 이를 규제하기 위한 법적 틀을 마련하고 있다는 점이다. 해당 법률은 생성형 AI의 악용으로 인한 허위 정보 및 조작된 생성물의 확산을 염두에 두고 이를 방지하기 위한 법적 기반을 제공하며, AI 기술로 인해 야기될 수 있는 윤리적 문제들을 해결하기 위한 초기 대응으로 평가받는다.

2021년에 들어서면서 미국의 AI 법제화는 AI 기술의 발전을 지원하거나 규제한다는 명목에 기인한 일차원적 접근을 넘어, 그 사회적 책임과 윤리적 사용을 보장하기 위한 보다 명확한 규제적 조치들을 포함하기 시작했다. 그중에서도 「미국 인공지능 진흥법(Advancing American AI Act, 2021)」은 AI 시스템이 공정하고 투명하게 운영되도록 하는 것을 주요 목표로 하였다.²⁶⁾ 예컨대, 위 법률에서는 감사관이 인공지능 시스템의 사용을 감독하고, 해당 시스템이 사생활 보호와 시민권에 미치는 영향을 포함한 여러 위험 요소를 평가할 수 있는 교육과 투자를 강화할 것을 규정하고 있다.²⁷⁾ 이러한 조치는 AI 시스템이 예상치 못한 방식으로 작동하거나 사회적 편향을 초래할 위험을 규제하기 위한 역할을 마련하기 위함이다. 또한 제7225조 a

24) Identifying Outputs of General Adversarial Networks Act, SEC. 4. NIST SUPPORT FOR RESEARCH AND STANDARDS ON GENERATIVE ADVERSARIAL NETWORKS. (b). OUTREACH. (2).

25) Identifying Outputs of General Adversarial Networks Act, SEC. 5. REPORT ON FEASIBILITY OF PUBLIC-PRIVATE PARTNERSHIP TO DETECT MANIPULATED OR SYNTHESIZED CONTENT.

26) 국회도서관, “미국의 인공지능 입법 현황과 바이든 행정부의 행정명령”, 2024, 3면.

27) Advancing American AI Act, SEC. 7224. PRINCIPLES AND POLICIES FOR USE OF ARTIFICIAL INTELLIGENCE IN GOVERNMENT. (c).

항의 2목, 3목에서는 각각 기관이 AI 시스템의 사용 사례를 투명하게 보고 하며 보고된 내용이 윤리적 기준을 충족할 것, AI 시스템의 사용 현황을 투명하게 공개하며 공개된 정보가 사생활 보호 및 민감한 정보와 관련된 법적 요구사항을 준수할 것을 요구하고 있다.²⁸⁾ 주로 AI 기술이 사회적 책임을 다하면서 운영되도록 보장하고, 그 과정에서 발생할 수 있는 윤리적 문제를 예방하는 것에 중점을 둔 조항이다. 언급되었던 조항들에서도 드러나듯이, 해당 법률은 AI 시스템의 운영에서 발생할 수 있는 모든 문제들에 관하여 명확한 책임주체를 지정하고 AI 기술의 안전한 사용을 보장하는 제도적 조치임이 부각된다. 이는 AI 기술의 발전을 지향하는 동시에 그러한 기술이 사회적, 윤리적 기준을 충족하도록 하는 데 중요한 역할을 한다.

이러한 흐름 속, 2023년 바이든 정부가 발표한 「행정명령 제14110호 (Executive Order 14110 on Safe, Secure, and Trustworthy AI, 2023)」는 AI 기술의 개발과 사용을 발전시키기 위함 포괄적인 지침과 더불어 사회적 및 경제적 위험을 최소화하기 위한 규제를 모두 제시하고 있는데, 이는 2020년 트럼프 정부가 발표했던 「행정명령 제14110호(Executive Order 13960 on Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government)」와 대조적으로 AI의 위험성을 규제하는 것에 초점을 맞추었다는 평가를 받는다.²⁹⁾ 「행정명령 제14110호」 제4조의 4.1(a) 항목에서는 상무부장관이 국립표준기술원(NIST)장을 통해 생성형 AI를 포함한 AI 시스템의 안정성과 신뢰성을 보장하기 위한 지침과 모범사례 등을 개발할 것을 요구한다.³⁰⁾ 위 조항에서는 특히 AI 시스템이 사용되기 전에 발생가능한 위험을 예상하고 완화하여 대비할 수 있는 방법을 마련하는 것이 중요함을 강조한다. 또한 AI 개발자들이 안전하고 보안이 보장된 AI 시스템을 배포할 수 있도록 하기 위하여 일종의 보안 검증 절차인 레드 팀ing(Red Teaming)³¹⁾ 시험을 수행하도록 하는 지침을 마련하도록 하는 한

28) Advancing American AI Act, SEC. 7225. AGENCY INVENTORIES AND ARTIFICIAL INTELLIGENCE USE CASES. (a).

29) 국가안보전략연구원, “바이든 행정부의 첫 인공지능(AI) 행정명령과 시사점”, 2023.

30) Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, SEC. 4. 1. (a).

편,³²⁾ 중소기업청(SBA) 장관이 AI 관련 목적으로 자본 접근 프로그램에 대한 인식을 제고하고 AI 관련 전문성을 가진 투자 자금이 중소기업 투자회사(SBIC) 라이선스를 신청할 수 있도록 지원할 것을 규정하고 있다.³³⁾ 이와 같은 조항들은 AI 기술이 사회적으로 책임 있게 사용될 수 있도록 보장하는 동시에 AI 기술의 경제적 적용과 그 범위의 확대를 지원하는 규제적 조치들로 분류된다.

미국의 AI 관련 법제화는 연도를 거치며 보다 구체적이면서도 포괄적인 방향으로 발전해 왔다. 초기에는 AI 기술의 채택 및 연구 개발 육성에 중점을 두었으나, 시간이 지날수록 국제적인 트렌드에 발맞추어 AI 기술의 사회적 책임과 윤리적 사용, 안전성, 투명성에 관한 규제적 요구를 강화시키는 방향으로 나아가고 있다. 현재 미국의 법제는 AI 기술의 사회적, 경제적 영향의 체계적 관리를 추구하면서도 AI 기술의 발전과 혁신을 위한 적절한 틀을 마련 중에 있다. 다만 일면에서는 미국의 AI 관련 법제가 AI의 위험을 평가, 규제하는 과정에서 충분히 엄격한 규제가 적용되지 않고 있다는 지적도 존재한다.³⁴⁾ 미국 역시 EU가 개척한 길을 따라 AI에 대한 규제를 강화하는 추세이나, 기존 EU 인공지능법의 구체적임에는 미치지 못한다는 비판은 피해갈 수 없는 형국이다.

3. 소결

인공지능 기술의 발전에 따라, 유럽과 미국 모두 인공지능 규제 방안을

31) 레드티밍(Red Teaming)은 특정 조직의 보안 수준을 평가하고 개선하기 위해 자체적으로 실제 공격을 시도하는 정책을, 레드팀(Red Team)은 이러한 작업을 전담하는 조직 내의 부서를 뜻한다.

32) Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, SEC. 4. 1. (a). (ii).

33) Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, SEC. 5. 3. (d).

34) SAFE Innovation Framework를 포함한 일부 프레임워크의 규제사항이 다소 엄격하지 못하다는 비판의 시각이 존재한다(구체적인 사항은 “Comparing the EU AI Act to Proposed AI-Related Legislation in the US”참고).

마련하고 있다. 두 지역 모두 리스크 기반 규제 방식 따라, 위험 정도에 따라 인공지능을 구분하여 대응하고 있다. 유럽의 경우, 인공지능을 4단계를 구분하여, 법 조문에 명시하고 이에 따른 대응 정도 역시 정해진 규제 방식을 취한다. 반면, 미국의 경우는 보다 유연한 접근 방식을 따른다. 가령, 각 주의 연방기관이 인공지능의 위험 정도를 평가해 규제하는 것이 있다.

유럽의 규제 방식의 경우, 명확하게 법으로 규제를 하기 때문에 기업이 규제 정도를 예상 할 수 있고, 고위험 인공지능에 대한 선제적으로 대처가 용이하다.³⁵⁾ 그러나, 인공지능 개발이 늦춰질 수 있고, 인공지능 개발에 맞춰 법을 바꾸기에 불리할 수 있다.

미국의 규제 방식의 경우, 2023년 하반기 이후 강력한 규제를 마련하고 있기는 하나, 그 전까지 자율적인 규제 방식을 취했다. 따라서, 기업의 기술 발전 저해에 대한 우려가 유럽에 비해 적다. 또한, 상황에 맞춰 다르게 규제할 수 있다는 장점이 있다. 하지만, 기존의 규제 방식을 따를 경우, 인공지능 규제 정도가 약해, 고위험 인공지능에 대한 대처에 불리하다는 단점 역시 존재한다.

두 접근 방식 모두 AI 법제에 있어 중요한 시사점을 제공하며, 한국은 이들 사례를 참고해 자국의 상황에 맞는 균형 잡힌 AI 단일법을 제정할 필요가 있다.

IV. AI 단일법 제정 방향 제언

1. 22대 국회 인공지능 법안 검토

35) McKinsey london 수석 이사 예툰데 다다는, 유럽의 규제 방식의 명확한 프레임 워크가 세계 기업이 생성 AI를 채택하는 데에 도움이 된다고 주장했다.

McKinsey & Company, Europe's Regulatory Environment Can Accelerate Adoption of Generative AI, McKinsey & Company, 2024. 10. 17. <https://www.mckinsey.com/featured-insights/lifting-europes-ambition/videos-and-podcasts/europes-regulatory-environment-can-accelerate-adoption-of-generative-ai> (최종접속일: 2024. 9. 1.)

22대 국회 개원 후 인공지능 관련 법안은 여당에서 3개, 야당에서 3개 총 6개의 법안이 발의되었다. 발의 순서에 따라 「인공지능 산업 육성 및 신뢰 확보에 관한 법률안 (2024.5.31.)」, 「인공지능 발전과 신뢰 기반 조성 등에 관한 법률안」 (2024.6.17.), 「인공지능산업 육성 및 신뢰 확보에 관한 법률안」 (2024.6.19.), 「인공지능산업 육성 및 신뢰 확보에 관한 법률안」 (2024.6.19.), 「인공지능기술 기본법안」 (2024.6.28.), 「인공지능 개발 및 이용 등에 관한 법률안」 (2024.7.4.)이다. 법안들은 대체로 정의, 위원회, 관계기관 설립, 산업 진흥, 윤리·신뢰, 벌칙, 기타 내용을 다룬 각 장으로 이루어졌고, 세부적인 내용은 법안에 따라 다르다. 각 법안에서 인공지능 기술 발전과 직접적인 관련이 있는 정의 관련 조항들과 윤리·신뢰 관련 조항들이 어떻게 다른지 비교 분석하고자 한다.

법안들은 고위험 인공지능 시스템의 정의와 금지된 인공지능 기술의 사용에 대한 규정을 포함하고 있다. 예를 들어, 안철수 의원의 「인공지능 산업 육성 및 신뢰 확보에 관한 법률안」 (2024.5.31.)은 고위험 인공지능 시스템을 사람의 생명이나 안전에 직접적인 영향을 미칠 수 있는 기술로 정의하며, 이에 대한 엄격한 규제와 모니터링이 필요함을 강조한다. 반면, 정점식 의원의 「인공지능 발전과 신뢰 기반 조성 등에 관한 법률안」 (2024.6.17.)은 보다 넓은 범위에서 고위험 인공지능 시스템을 정의하며, 사회적, 경제적 위험도 함께 고려한다. 이와 달리, 권철승 의원의 「인공지능 개발 및 이용 등에 관한 법률안」 (2024.7.4.)은 고위험 인공지능 시스템의 정의를 구체적으로 나열하기보다는, 이를 판단할 수 있는 기준과 절차를 제시하는 데 초점을 맞추고 있다.

각 법안은 인공지능 개발사업자와 인공지능 이용사업자를 다르게 정의하고 있다. 김성원 의원의 「인공지능산업 육성 및 신뢰 확보에 관한 법률안」 (2024.6.19.)은 개발사업자를 인공지능 기술을 연구, 개발, 상업화하는 자로 명확히 규정하고 있으며, 이용사업자는 인공지능 기술을 직접 사용하거나 응용하는 자로 정의하고 있다. 민형배 의원의 「인공지능기술 기본법안」 (2024.6.28.)은 이에 대해 좀 더 포괄적인 정의를 제공하며, 인공지능 관련 다양한 사업 유형을 모두 포괄하는 것을 목표로 하고 있다.

검인증제도는 인공지능 시스템이 윤리적이고 신뢰할 수 있는지 검증하는 절차로, 법안들에서 중요한 비중을 차지하고 있다. 안철수 의원의 법안은 인공지능 시스템의 검인을 위해 독립된 기관을 설립하고, 이를 통해 윤리적 기준과 기술적 신뢰성을 평가받도록 규정하고 있다. 조인철 의원의 「인공지능산업 육성 및 신뢰 확보에 관한 법률안」(2024.6.19.)은 검인증제도를 강화하며, 인공지능 시스템이 일정 기준을 충족하지 못할 경우 시장에서의 사용을 제한할 수 있는 강력한 규제 조치를 제안하고 있다. 반면, 민형배 의원의 법안은 검인증제도의 절차를 간소화하고, 기술 혁신을 저해하지 않도록 유연한 접근 방식을 취하고 있다.

우선 허용 후 사후 규제는 인공지능 기술의 혁신을 촉진하기 위한 접근 방식으로, 각 법안에서 다르게 적용되고 있다. 정점식 의원의 법안은 인공지능 기술의 신속한 상용화를 위해 우선 허용 후 사후 규제의 원칙을 적극적으로 반영하고 있으며, 이에 대한 사후 규제 절차를 상세히 설명하고 있다. 권칠승 의원의 법안 역시 이 원칙을 받아들이고 있으나, 보다 엄격한 사후 규제 메커니즘을 도입하여 위험성이 높은 인공지능 기술에 대한 신속한 대응을 강조하고 있다. 반면, 김성원 의원의 법안은 우선 허용보다는 사전 규제에 무게를 두고, 보다 신중하게 인공지능 기술의 도입을 추진하는 방향을 제시하고 있다.

이처럼 현재 국회에 발의된 법안들의 방향성과 규범 체계도 다양하여 AI 업계 전반에 혼란을 야기하는 것을 막기 위해서라도 AI 단일법 제정이 조속하게 이루어져야한다. EU와 미국의 인공지능 관련 법제를 참고한 후, 규제 정도에 대한 방향성과 척도를 설정하여 관련 법을 제정할 필요성이 있다.

2. 제언

(1) 규제와 혁신의 균형

사회적 가치를 보장하면서도 AI 기술의 혁신을 촉진할 수 있는 방향으

로 AI 단일법을 제정해야 한다. 한국과학기술평가원의 ‘EU 인공지능(AI) 규제 현황과 시사점’에 따르면, 허용과 제한의 이분법적인 접근보다는 AI 기술혁신·산업 진흥과 리스크 완화 측면을 모두 고려한 법제가 필요하다.³⁶⁾ 특히, 국내 인공지능 산업이 아직 발전 초기 단계이며, 글로벌 시장 선점이 중요한 상황임을 고려할 때 해외 사례에 따른 규제보다는 글로벌 경쟁력 확보와 국내 AI 산업의 상황에 맞는 규제를 마련할 필요가 있다.³⁷⁾ EU는 역내 기업 보호와 기술 주도권 확보, 미국은 자율적인 규제를 통한 혁신이라는 자국의 상황에 맞는 법규를 제정하고 지원책을 마련하고 있는 만큼, 우리나라도 이익이 실현 가능한 규제의 방향으로 균형있는 규범 체계를 마련해야 한다.

(2) 유연한 법률 체계 구축

OECD AI 원칙 2.3은 “정부는 신뢰할 수 있는 AI 시스템의 연구·개발 단계에서 배포 및 운영 단계로의 전환을 지원하는 민첩한 정책 환경을 조성해야 하며, 이를 위해 AI 시스템을 테스트하고, 적절한 경우 이를 확장할 수 있는 통제된 환경을 제공하는 실험적 접근을 고려해야 한다”라고 명시하고 있다. AI 기술은 발전 속도가 빠르기 때문에, AI 단일법은 기술 발전에 유연하게 대응할 수 있는 구조로 설계되어야 한다. 이를 위해 유럽의 인공지능 법의 규제 샌드박스 제도를 도입해 새로운 AI 기술이 적은 규제 속에서 테스트될 수 있도록 해야 하며, 권칠승 의원의 「인공지능 개발 및 이용 등에 관한 법률안」(2024.7.4.)에서 제안한 것처럼 AI의 정의에 대한 내용이나 구체적인 사안에 대해서는 대통령령으로 위임할 수 있도록 제정하는 방향이 유연한 법률 체계를 구축할 수 있을 것이다.

(3) 국제 표준과의 조화

AI 기술은 국제적으로 경쟁이 치열한 분야이기 때문에, 각국이 서로 다

36) 한국과학기술평가원, “EU 인공지능(AI) 규제 현황과 시사점”, 2024, 7면.

37) 과학기술정보방송통신위원회, 「인공지능 산업 육성 및 신뢰 확보에 관한 법률안」 검토보고서, 2024, 13면.

른 규제와 표준을 적용할 경우 기술 개발 및 상용화에 혼란이 초래될 수 있다. 한국과학기술평가원의 ‘EU 인공지능(AI) 규제 현황과 시사점’에 따르면, 각국이 취하는 AI 규제 차이로 통상마찰이 발생할 가능성을 언급하며, 이에 대한 대비가 필요하다.³⁸⁾ 따라서 한국은 앞서 살펴본 EU와 미국의 법제 등 국제 표준을 고려한 법제를 제정하여, 이를 통해 AI 기술의 글로벌 호환성을 보장하고, 한국 기업들이 해외 시장에서 경쟁력을 가질 수 있도록 해야 한다. 또한, 국제 협력을 강화해 AI 기술과 관련된 국제적 법적 문제 해결에 기여해야 한다.

(4) 산업계와의 협력 강화

AI 단일법 제정 과정에서 산업계와의 협력은 필수적이다. 규제 기관과 기업 간의 소통을 강화하여, 기업들이 규제에 대한 명확성을 가질 수 있도록 하고, 규제가 혁신을 저해하지 않도록 해야 한다. 맥킨지 런던 수석이사 예툰데 다다는 규제를 통한 명확한 프레임워크 제공으로 기업이 AI도입을 용이하게 할 수 있다고 밝혔다.³⁹⁾ ITI의 부사장 존 밀러는 전 세계 정책 입안자들이 산업, 학계, 지역 이해 관계자와 협력하여 이 기술이 모든 곳에서 안전하고 신뢰할 수 있고 투명한 방식으로 계속 개발되고 사용되도록 하는 것이 중요하다고 언급하며, 산업계와 정책 입안자들의 소통을 강조했다.⁴⁰⁾ 따라서, 정부는 미국의 「생성적 적대 신경망 출력물 확인법」과 같이 기업의 의견을 반영한 현실적인 규제 방안을 마련하고, 규제 시행 후에도 산업계와 지

38) 한국과학기술평가원, “EU 인공지능(AI) 규제 현황과 시사점”, 2024, 7면.

39) McKinsey & Company, Europe’s Regulatory Environment Can Accelerate Adoption of Generative AI, McKinsey & Company, 2024. 10. 17. <https://www.mckinsey.com/featured-insights/lifting-europes-ambition/videos-and-podcasts/europes-regulatory-environment-can-accelerate-adoption-of-generative-ai> (최종접속일: 2024. 9. 1.)

40) Information Technology Industry Council (ITI), New ITI Global AI Policy Recommendations Promote Government, Industry, and Stakeholder Collaboration on AI, Information Technology Industry Council, 2021. 3. 21. <https://www.itic.org/news-events/news-releases/new-iti-global-ai-policy-recommendations-promote-government-industry-and-stakeholder-collaboration-on-ai> (최종접속일: 2024. 9. 1.)

속적으로 협력해 규제의 효과를 평가 및 개선하는 시스템을 도입해야 한다.

V. 결론

인공지능은 21세기 가장 중요한 기술 혁신 중 하나로, 그 영향력은 산업 전반과 사회 전체에 걸쳐 광범위하게 확산되고 있다. 한국이 이러한 흐름에 효과적으로 대응하려면 AI 단일법 제정을 통해 명확한 법적 틀을 마련해야 한다. 그러나 AI 법제는 단순한 규제 도입에 그치는 것이 아니라, AI 기술 발전을 촉진하면서도 사회적 가치를 보장할 수 있는 균형 잡힌 법률 체계가 되어야 한다.

유럽연합은 강력한 규제를 통해 AI 기술의 윤리적 측면을 강조하고 있으며, 미국은 기술 혁신을 촉진하는 자율 규제 방식을 채택하고 있다. 한국은 이 두 접근 방식을 종합하여, 기술 혁신과 규제를 조화시키는 법체계를 마련해야 한다. 특히 글로벌 표준과의 조화를 이루는 법안을 통해 한국의 AI 기술이 국제적 경쟁력을 가질 수 있도록 해야 한다.

AI 법 제정 과정에서 산업계와의 협력을 강화하고, 유연한 법률 체계를 구축하여 AI 산업 발전에 저해가 되지 않는 방향으로 법제를 발전시켜야 한다. 아울러 한국이 AI 시대를 선도적으로 맞이할 수 있는 기반을 다져야 한다.

궁극적으로 AI 법제는 기술 발전과 사회적 책임이 공존할 수 있는 환경을 조성하는 것을 목표로 해야 하며, 이를 통해 한국은 글로벌 AI 선도 국가로 자리매김할 수 있다. AI 단일법은 미래 기술의 방향성을 결정짓는 중요한 이정표가 될 것이며, 이를 통해 한국 사회는 AI 시대를 선도적으로 맞이할 준비를 갖추 수 있을 것이다.

[참고문헌]

[국내 논문]

- 김나래, “인공지능 규제에 관한 소고”, 동아대학교 법학연구소, 『東亞法學』 제97호, 2022.
- 김송옥, “AI 법제의 최신 동향과 과제 -유럽연합(EU) 법제와의 비교를 중심으로-”, 한국비교공법학회, 『공법학연구』 제22권 제4호, 2021.
- 박노형·김효권, “EU AI법 개인정보보호 관련 규정 검토”, 사법발전재단, 『사법』 제1권 제68호, 2024.
- 이대회, “인공지능에 의한 개인정보의 자동 처리에 대한 투명성 확보에 관한 연구”, 한국법학원, 『저스티스』 제200권 제200호, 2024.
- 이숙연, “인공지능 관련 규범 수립의 국내외 현황과 과제”, 법조협회, 『法曹』 제72권 제1호, 2023.
- 이효진, “인공지능(AI) 단일법 제정 필요성과 행정법학의 과제”, 국민대학교 법학연구소, 『법학논총』 제35권 제3호, 2024.

[기타 문헌]

- 국가안보전략연구원, “바이든 행정부의 첫 인공지능(AI) 행정명령과 시사점”, 이슈브리프 제480호, 2023.
- 국회도서관, “미국의 인공지능 입법 현황과 바이든 행정부의 행정명령”, 2024-1호(통권 제73호), 2024.
- 법무법인 화우, “생성형 AI의 국제 규제 동향 -미국, EU, 중국을 중심으로 한 규제 동향 훑어보기”, 2024.
- 법제처 미래법제혁신기획단, “인공지능(AI) 관련 국내외 법제 동향”, 2024.
- 한국과학기술평가원, “EU 인공지능(AI) 규제 현황과 시사점”, 2024.
- 한국정보산업연합회, “인공지능기본법 입법 추진현황 및 산업진흥 측면에서 본 이슈”, 2024.
- KISTEP 과학기술정책센터, “생성형 AI 관련 주요 이슈 및 정책적 시사점”,

2023.

Alex Engler, “The EU and U.S. are starting to align on AI regulation”,
Brookings, 2022.

Sofia Gracias, “Comparing the EU AI Act to Proposed AI-Related
Legislation in the US”, The University of Chicago Business Law
Review, 2024.