

인공지능 관련 법 제도의 주요 논의 현황

임형주 법무법인(유) 율촌 변호사, IP&Technology 부문 신산업IP팀 팀장

1. 머리말

2016년 3월 구글 딥마인드가 개발한 인공지능(AI) 바둑 프로그램 '알파고(AlphaGo)'가 인간 최고수인 이세돌 9단에게 4:1로 완승을 거두었던 이른바 '알파고 쇼크'는 AI 발전의 잠재력과 현실에서의 활용 가능성에 대한 인식을 완전히 바꾸어 놓은 사건이었다. 이를 계기로 많은 기업들이 여러 최신 기술들을 바탕으로 한 AI 기술의 개발에 나서면서 관련 분야 산업 전반에 상당한 발전을 가져왔다. 동시에 소위 '4차 산업혁명'에 관한 담론이 확산되면서 AI 기술에 대한 법적·제도적 규제 필요성에 대한 논의가 본격적으로 촉발되었다.

알파고 쇼크 이후 비교적 잠잠한 것으로 보였던 AI 기술에 관한 관심은 최근 오픈AI가 개발한 대화형 AI 서비스 '챗지피티(ChatGPT)'가 공개되면서 다시 한번 뜨겁게 증폭되고 있다. 답변 내용의 진위 여부를 떠나 인간의 언어를 이해하고 인간과 대화하는 것과 같은 착각을 불러일으킬 정도로 정교화된 AI 모델의 등장은 당장 검색엔진 시장의 구도를 바꾸고 있는 것을 넘어 기존의 업무 방식, 나아가 우리의 삶 전반에 다양한 방식으로 활용되며 큰 변화를 가져다줄 것으로 예상된다.

'챗지피티' 외에도 최근 등장하는 여러 AI 관련 기술들은 우리 생활 전반에 상당한 변화를 가져올 것으로 기대된다. 이미지 생성형 AI의 하나인 '미드저니(Midjourney)'로 생성한 그림이 2022년 콜로라도 주립박람회 미술대회 디지털 아트부문에서 1위를 차지하는 파란을 일으키거나,¹⁾ AI가 작곡한 음악이 드라마 OST 등에 사용되기도 하는 등²⁾ 오로지 인간의 창작 활동의 전유물로 여겨져 온 '예술'의 영역에 대한 도전장을 내밀고 있다. 텍스트만으로 영상을 제작해주는 AI 모델들의 등장³⁾은 영화 등 영상물 제작과 관련된 산업을 크게 바꿀 것으로 전망된다. 이처럼 AI 기술의 다양한 활용이 그려갈 미래 모습에 대하여는 여러 업무를 경감하고 더욱 효율적이고 신속·정확한 처리를 바탕으로 삶을 더욱 윤택하게 해주리라는 긍정적 기대도 있지만, 한편으로는 AI 기술의 오·남용 또는 통제되지 않은 사용이 미칠 부정적 측면에 대한 우려도 함께 제기되고 있다.

챗지피티와 같은 모델을 사용하여 깃허브(GitHub)에서 자동 코드 완성 서비스를 제공하는 AI인 '코파일럿(Copilot)'은 최근 다른 개발자들의 코드를 무단으로 학습하였다는 이유로 집단 소송에 휘말렸고,⁴⁾ 미드저니(Midjourney) 또한 학습 과정에서 창작자들의 동의 없이 저작물을

1) <https://www.joongang.co.kr/article/25099483#home>

2) <https://www.donga.com/news/Economy/article/all/20220903/115282746/1>

3) <https://zdnet.co.kr/view/?no=20230321092530>

무단으로 사용하였다는 이유로 집단소송에 휘말리고 있는 등⁵⁾ AI를 개발하는 과정에서 제3자의 권리 침해 여부가 크게 논란이 되고 있다. 그 밖에도 AI 생성물로 인한 표절 등 제3자의 권리 침해 여부, AI 생성물에 대한 특허권, 저작권 등의 권리 인정의 필요 여부, 나아가 AI의 비윤리적 이용에 대한 우려 등 여러 문제들이 현실로 다가오고 있다. 그에 따라 법적·제도적 규율을 통한 통제의 필요성이 대두된 결과 최근 몇몇 국가들을 중심으로 AI 관련 기술의 통제를 위한 법률 제정 등에 대한 논의가 본격적으로 이뤄지고 있다.

본고에서는 AI 기술의 통제에 관한 법률, 이른바 ‘인공지능 기본법’의 주요 논의에 관한 현황을 살펴본다. 먼저 인공지능 기본법에 관한 개요 및 주요 쟁점들을 간략하게 다룬 뒤, 관련하여 주요 국가들의 동향과 우리나라의 현황을 소개한다. 인공지능 기본법에 관한 논의는 제정 필요성 등을 놓고 찬반 양론이 치열하게 대립하고 있으므로, 본고에서는 최대한 객관적인 입장에서 현황을 소개하는 수준으로 다루고자 한다.

2. 인공지능 기본법 개요 및 주요 쟁점

2.1 인공지능 기본법 개요

‘2001 스페이스 오딧세이’, ‘터미네이터’, ‘매트릭스’ 등 소위 AI 디스토피아를 다룬 여러 문학 작품이나 영화들을 통하여 오래 전부터 AI 기술 통제의 필요성에 대한 문제의식이 제기되어 왔다. 하지만 AI에 대한 법적·제도적 규제의 필요성이 본격적으로 논의된 계기는 전술한 알파고 쇼크이다.

알파고 쇼크 이후 널리 활용된 소위 딥러닝 기술에 바탕을 둔 AI 기술에 대하여 규제 및 통제의 필요성이 대두된 근본적인 원인은 딥러닝에 의하여 개발된 AI는 블랙박스라 한다는 점이다. AI에 일정한 데이터를 학습시키면 인간이 원하던 것과 동일하거나 거의 유사한 산출물을 얻을 수 있게 되지만, 어떠한 근거로 어떠한 과정을 통하여 그러한 결과가 나왔는지 정확하게 파악하기 어렵다는 문제가 있다.

AI 기술을 통제하기 위한 다양한 시도들, 특히 인공지능 기본법으로 대표되는 AI 기술에 대한 법적·제도적 규제는 블랙박스와 같은 AI의 특성으로 인하여 실제로 발생하였거나 발생할 우려가 있는 여러 권리 침해 문제를 최소화하기 위한 목적에서 도입 또는 검토되고 있다. 이에 관하여는 여러 다양한 쟁점들이 있으나, 아래에서는 주요 쟁점으로 논의되고 있는 사항들을 중심으로 소개한다.

2.2 인공지능 기본법의 주요 쟁점들

2.2.1 AI 기술의 금지 범위 및 규제의 적용 범위

머신러닝을 통한 AI 개발이 본격화되고 초기 모델들이 출시되는 과정에서 여러 AI 기술 오용 사례가 나타나면서, AI 규제에 대한 논의에서는 AI 윤리에 관한 쟁점이 중점적으로 다뤄졌다. 마이크로소프트사의 챗봇 ‘테이(Tay)’가 학습 과정에서 편견과 혐오를 가감 없이 드러내며 24시간만에 송출을 중단한 사건,⁴⁾ 네덜란드 세무당국이 AI를 활용하여 거짓 신고를 필터링하는

4) <https://zdnet.co.kr/view/?no=20221110105848>

5) <https://www.ekoreanews.co.kr/news/articleView.html?idxno=64876>

과정에서 저임금, 소수민족 가정들 상당수가 부당하게 조사 대상으로 선정된 사건⁷⁾ 등을 통하여 이미 존재하는 수많은 자료를 처리하여 알고리즘에 따라 결과물을 내놓는 AI 모델의 특성상 잘못된 정보, 의도치 않은 선입견이 포함된 정보 등이 사용될 가능성이 드러났다. 이에 따라 AI에 대한 규제는 AI 개발 및 사용 과정에서의 윤리 문제에 대한 가이드라인을 포함할 필요성이 제기되었다. 특히 챗지피티 등장 이후에는 인공지능 기술의 위험성 문제가 제기되면서 인공지능의 개발 자체를 일시 중단할 것을 요구하는 주장이 나오기도 하였다.⁸⁾

<표 1> 주요국 인공지능 법안 및 가이드라인 동향

미국	알고리즘책임법안 입법 추진 (2022년 3월)	• AI 도입 촉진을 위한 자율 규제
	통신정보관리청(NTIA) 인공지능 책임성에 대한 공개 여론 수렴(2023년 4월)	• 인공지능 책임성, 투명성, 책임주체 관련 여론 조사 • 합리적 규제 수준 도출을 위해 인공지능 책임성 규제 수준에 대한 대중과 기업의 인식 조사
영국	인공지능 규제백서 (2023년 3월)	• 혁신 도모 및 규제 불확실성 감소 등 성장과 번영 추진 • 공공 신뢰 제고 • 법 규제 최소화 및 규제혁신 추진 간 균형 확보 • 분야별 유연하고 합리적인 규제 추진
	AI 및 데이터 보호 가이드 일부 개정(2023년 3월)	• 직무성과 거버넌스 투명성 및 적법성 등 일부 내용 추가 • 인공지능 규제백서에 맞춰 변하는 환경에서 AI의 미래 지향적 활용 위한 개정
EU	인공지능법안 (2021년 4월)	• AI 위험 수준 4가지로 구분해 수준에 따른 규제 부과
	AI Act 관련 EU 이사회 수정의견(2022년 11월)	• 국가안보, 연구목적 및 비전문적 영역에 대한 적용 제외 • 고위험 AI의 적용범위 등 분류체계 일부 수정 • 고위험 인공지능 규제 요건 준수의 기술적 실현 가능성 향상 및 수범자 부담 경감 • 혁신 지원 대책 강화

(자료: NIA)

이러한 맥락에서 인공지능 기본법 제정에 관하여 논의되는 가장 핵심적인 쟁점은 인공지능 기본법에 따른 규제의 범위를 어디까지 적용할 것인지 여부, 나아가 일정한 AI 개발을 법적으로 금지할지 여부 및 법적으로 금지하는 범위를 어떻게 설정할 것인지 여부이다. 예컨대 딥페이크(Deepfake)와 같은 영상 또는 음성 합성 기술은 범죄에 악용될 우려가 있으니 그 자체로 금지되어야 한다는 의견과 엔터테인먼트 분야에서 가수 등 유명인이 직접 출연하지 않고도 영상 제작이 가능하며 현재도 널리 활용되고 있는 모션캡처 등을 활용한 CG 작업을 용이하게 하므로 긍정적으로 바라보는 의견이 모두 존재한다.⁹⁾ 이에 따라 어떤 기술을 구체적으로 금지하고 허용할지, 허용하는 경우 어떤 수준에서 이를 분류하고 통제할지 등에 관한 객관적인 기준의 수립이 요구되고 있다.

6) <https://www.dongascience.com/news.php?idx=11158>

7) <https://spectrum.ieee.org/artificial-intelligence-in-government>

8) https://www.chosun.com/economy/tech_it/2023/04/10/IAMSK27X6JGVJN66UIUE4WM3KQ/

9) <https://www.scourt.go.kr/portal/gongbo/PeoplePopupView.work?gubun=44&seqNum=2608>

2.2.2 AI 학습을 위한 데이터 수집 등에 대한 규제

알파고 쇼크 이후 딥러닝 기술이 AI 모델 개발의 핵심적인 기술로 자리 잡으면서 AI 개발은 점점 더 많은 수준의 데이터 학습을 필요로 하고 있다. 데이터를 확보하는 과정에서 기존에는 소위 크롤링을 비롯한 다양한 방법을 통하여 월드와이드웹에 존재하는 데이터를 수집 및 활용하는 방법이 많이 사용되어 왔다. 하지만 이 같은 데이터의 수집 및 이용과정에서 전술한 코파일럿, 미드저니 등의 저작권 침해 논란을 포함하여 지적재산권 및 개인정보 등 제3자의 권리에 대한 침해 가능성이 있어 이에 대한 규제의 필요성 또한 제기되고 있다.

이에 대하여는 사전 허용되지 않은 크롤링 등의 행위 자체를 위법한 것으로 판단하여 데이터의 권리자들에 대한 보호가 필요하다는 견해와 AI 기술 발전 필요성을 근거로 소위 텍스트·데이터 마이닝(Text Data Mining, TDM) 등 AI 학습을 위한 목적의 데이터 수집 및 이용에 대하여는 일정한 요건을 갖추는 경우 면책 규정을 두어야 한다는 견해가 대립하고 있다. 또한 이에 관한 규정을 두는 경우 인공지능 기본법에 두어야 할지, 관련 지적재산권 법률 및 개인정보보호법 등 개별 법률에 두는 것이 바람직한지에 관하여도 의견이 나뉘고 있다.

2.2.3 알고리즘 투명성과 설명요구권 도입 여부

AI 규제에 대한 논의의 초기에는 예컨대 자율주행 자동차가 등장하는 경우, AI 오류로 인한 사고 발생을 어떻게 통제할 것인지, 사고가 발생한 경우 그것이 AI의 오류에 인한 것인지 어떻게 검증할지 등 주로 결과물에 대한 통제와 책임의 필요성에 대한 관점에서 접근이 이루어졌다. 관련하여 알파고 쇼크 이후 각국 정부, 다국적 기업, 국제기구, 전문가 협회 등은 AI 시스템이 준수하여야 하는 여러 원칙들에 대하여 발표하였는데, 해당 원칙들에서 공통적으로 규정된 원칙이 AI 시스템의 투명성(Transparency) 확보에 관한 사항이다.

유럽연합 집행위원회의 AI 고위급 전문가그룹(AI HLEG, High-Level Expert Group on Artificial Intelligence)이 2019년 4월 발표한 '신뢰가능한 AI를 위한 윤리 가이드라인(이하 EC AI 윤리 가이드라인)'에 의하면, '투명성'이란 ① AI 시스템에 의한 의사결정은 사람에게 의하여 이해 및 추적될 수 있어야 하고, ② 개인이 AI 시스템의 의사결정으로부터 상당한 영향을 받은 경우 그 결정을 내린 이유에 대한 설명을 요구할 수 있어야 하며, ③ 소통 상대방이 AI 시스템이라는 점을 통지받을 수 있어야 하고, ④ 기본권 보장을 위하여 필요한 경우 AI 시스템과의 소통을 거부할 수 있어야 함을 주요 내용으로 한다.¹⁰⁾ 위 각원칙들로부터 AI 시스템 적용 사실에 대한 고지의무, 적용거부권, 이의제기권 등 다양한 권리보장 및 의무 부과가 논의되고 있으나, 가장 중요한 이슈는 설명요구권이다.

설명요구권(right to explanation)이란 AI 서비스 이용자가 AI에 의한 의사결정의 결과의 타당성에 대한 의문이 있을 경우 그 AI의 알고리즘의 배열, 주요 변수 등을 비롯한 작동 원리에 대한 설명을 요구할 수 있는 권리를 의미한다. AI 시스템으로부터 영향을 받는 개인이 AI 시스템의 결정에 대하여 이의를 제기하거나, 그로 인한 피해를 구제받기 위한 선행 조건으로서 중요한 의미를 갖는 권리로 여겨지고 있다.¹¹⁾

¹⁰⁾ European Commission AI HLEG (High-Level Expert Group on Artificial Intelligence), "Ethics guidelines for trustworthy AI", Brussels: European Commission, 2019, 18면 이하.

그러나 AI 시스템에 대한 설명요구권, 특히 알고리즘에 대한 설명요구권은 자칫 그 설명에 필요한 정도가 지나치게 구체적인 경우에는 AI 시스템에 관한 해당 기업의 영업비밀을 공개하는 것과 다르지 않을 수 있다는 우려가 있고, 특히 경쟁자의 영업비밀 취득을 위한 수단으로 악용될 우려가 있다는 이유로 도입을 반대하는 입장도 있다.

2.2.4 AI 산출물에 대한 책임의 귀속 및 AI에 대한 법인격 부여 필요 여부

2.2.3절에서 언급된 AI 산출물에 대한 통제와 책임의 필요성에 관하여, 예컨대 현행 자동차손해배상보장법 제29조의2에서는 자율주행자동차의 사고 발생시 기존의 운전자 책임을 유지하되 차체 결함으로 인하여 사고가 발생한 경우에는 그 제조자에게 책임을 구상할 수 있도록 하는 등 현재로서는 해당 AI의 개발자 또는 운용자에게 책임이 귀속되는 것을 전제로 하고 있다. 그러나 향후 AI 관련 기술이 더욱 고도화되고 AI 시스템이 더욱 자율화되는 경우에는 AI에 별도의 법인격 등 주체로서의 권리의무를 부과할 필요성이 있는지 여부가 논의되고 있다. 인공지능이 특허법상 발명자가 될 수 있는지에 관하여 최근 진행 중인 소송 또한 같은 맥락이다.¹²⁾

AI에 대한 법인격 등 권리의무 주체로서의 성격을 부여할 필요가 있다는 입장은 AI의 개발자에게 책임을 부과하는 것은 자칫 자기책임의 원칙에서 벗어난 과도한 규제가 될 수 있다는 입장이다. AI의 특성상 개발자 또한 예상하지 못했던 다양한 결과가 발생할 수 있기 때문이다. 반대로 현재 AI의 수준을 고려할 때 별도의 권리의무 주체성을 인정할 필요가 있다고 보기 어려울 뿐 아니라, 이는 AI의 개발자 또는 운용자의 책임 회피수단으로 악용될 여지가 있다는 점에서 반대하는 견해도 상당하다.

3. 인공지능 기본법 제정 관련 주요국 동향

3.1 유럽연합(EU, European Union)

EU는 2018년 5월 EU 일반데이터보호규정(GDPR, General Data Protection Regulation)을 시행하며 자동화된 의사결정 알고리즘에 대한 개인정보 주체의 권리를 명시했다. 2019년 4월에는 인공지능 윤리 가이드라인(Ethics Guidelines for trustworthy AI)¹³⁾을 통하여 AI 시스템 전주기에 걸쳐 반드시 지켜야 할 3가지 구성 요소 및 7가지 원칙을 발표하였고, 2020년 2월에는 EU AI 백서(White Paper on Artificial Intelligence)를 통하여 유럽 내 단일 AI 생태계 조성을 위한 정책 및 규제 프레임워크를 제시하는 등 알파고 쇼크 이후 초창기부터 AI 관련 법적·제도적 규제를 선도적으로 마련해왔다.

위와 같이 여러 차례 거듭해온 논의를 바탕으로 EU 집행위원회는 2021년 4월 유럽의회에 「인공지능에 관한 통일규범(인공지능법)의 제정 및 일부 연합제정법들의 개정을 위한 법안 (Proposal for a Regulation laying down harmonized rules on artificial intelligence(Artificial Intelligence Act) and amending certain Union Legislative Acts)」(이하 'EU AI법안')을 발의¹⁴⁾하며

11) Edwards, Lilian & Veale, Michael, "Slave to the Algorithm? Why a 'Right to an Explanation' Is Probably Not the Remedy You Are Looking For", 16 Duke Law & Technology Review 18, 2017, 21면; Kaminski, Margot E., "The right to explanation, explained". 34 Berkeley Tech. L.J. 189, 2019, 204면 등.

12) <https://www.aitimes.kr/news/articleView.html?idxno=27054>

13) <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

14) European Commission, COM/2021/206 final, Brussels, 21.4.2021

인간의 안전과 생계에 위협이 될수 있는 AI를 전면 금지하되, 그렇지 않은 AI의 경우에도 기본권을 침해할 가능성이 있는 경우에는 사전평가 등 일정한 의무사항을 준수하여야 비로소 서비스를 출시할 수 있도록 하는 등 강력하고도 포괄적인 규제를 예고하였다. EU AI법안은 일부 수정을 거쳐 2023년 5월 11일 유럽의회 상임위원회를 통과하였으며, 다음달 본회의 표결을 앞두고 있다.

EU AI법안이 유럽의회 본회의를 통과할 경우 주요국들 가운데 최초로 인공지능에 관한 포괄적인 규제가 될 예정이다. EU AI법안의 주요 내용들을 간략히 소개하면 아래와 같다.

3.1.1 규제의 적용 범위 – 포괄적으로 적용하되 사적 범위의 사용에 대한 적용 제외

EU AI법안에 따른 규제의 적용 범위는 크게 인적 범위와 물적 범위로 나누어볼 수 있다.

먼저 인적 범위에 관하여 살펴보면, EU AI법안은 EU에서 AI 시스템을 시장에 출시하거나 서비스하는 모든 공급자, EU 내에 소재하는 인공지능 시스템의 사용자에게 적용되고, 인공지능의 산출물이 EU 내에서 사용되는 경우에 한하여는 제3국에 소재한 인공지능 시스템의 공급자 및 사용자에게 적용된다. 다만, 제3국의 공공당국 또는 국제기구가 EU 또는 EU 회원국과 체결한 국제협약에 따라 법 집행 및 사법공조를 위하여 인공지능 시스템을 사용하는 경우와 사적인 목적의 인공지능 사용의 경우에는 EU AI법안의 적용을 받지 않는다.

물적 범위에 관하여 살펴보면, EU AI법안이 적용되는 AI 시스템에는 하나 이상의 기술과 기법(기계학습, 논리·지식 기반 접근법, 통계적 접근법, 베이지 추정법, 검색 및 최적화법 등)을 이용하여 개발되고, 인간이 정의한 목적에 따라 상호작용하는 환경에 영향을 미치는 콘텐츠, 예측, 추천, 의사결정 등의 결과물을 생성할 수 있는 소프트웨어를 모두 포함하되, 군사 목적으로만 개발 및 사용되는 AI 시스템은 적용에서 제외되었다.

3.1.2 규제의 접근방법 – 위험수준별(Risk-based) 규제

EU AI법안은 법안의 규제 대상이 되는 AI 시스템을 인간의 건강과 안전, 기본권에 미치는 위험 정도에 따라 ① 허용불가(unacceptable risk), ② 고위험(high risk), ③ 제한된 위험(limited risk), ④ 최소위험(minimal risk)으로 구분하고, 각 분류에 따라 법적 의무를 달리하여 부과하는 위험수준별 규제를 채택하였다.

EU의 가치를 침해하는 매우 유해한 AI 사용으로, 인간의 안전, 생계 및 기본권에 대한 명백한 위협이 되는 AI 시스템으로 정의된 ‘허용불가’ AI에는 인간의 자유의지를 회피하여 행동을 조작하는 AI 시스템, 공공기관이 AI 기술을 활용하여 개인의 사회적 신용 등을 평가하는 AI 시스템 등이 예시로 제시되었다. ‘허용불가’ 등급에 해당하는 AI 시스템은 그 개발 및 사용을 엄격하게 금지하였다.

인간의 안전, 생계 및 기본권에 부정적인 영향을 끼치는 AI 시스템인 ‘고위험’ AI 시스템에는 생체인식시스템을 비롯하여 중요 기반시설의 관리 및 운용에 사용되거나 교육 및 직업훈련, 근로자 관리, 필수 공공/민간 서비스에 사용되는 AI 시스템 등이 예시로 제시되었다. ‘고위험’ 등급에 해당하는 AI 시스템은 출시 전 ① 적절한 위험관 리시스템을 운용하고 있는지, ② 고품질의 데이터셋(Dataset)을 사용하였는지, ③ 당국에 상세한 문서를 제공하였는지, ④ 결과 추적을 위

한 로그를 기록 및 관리하고 있는지, ⑤ 사용자에게 투명하게 정보를 제공하는지, ⑥ 인간의 감독이 이루어지고 있는지, ⑦ 정확성, 신뢰성 및 보안이 유지되고 있는지 등의 요건을 준수하고 있는지 여부에 대한 적합성 평가를 거쳐 이를 통과한 경우에만 서비스될 수 있도록 하였다.

한편 챗봇과 같이 인간과 상호작용하는 시스템, 딥페이크 등과 같이 콘텐츠를 생성하거나 조작하는 시스템 등이 포함된 '제한된 위험' AI의 경우에는 강도 높은 투명성 원칙을 적용하여 ① 인간에게 AI 시스템과 상호작용하고 있다는 사실을 고지받을 수 있도록 시스템을 설계 및 운영하여야 하고, ② 딥페이크 등과 같은 AI 산출물의 경우에는 그 산출물이 인위적으로 생성되거나 조작되었다는 사실을 공개할 의무를 부과하였다.

마지막으로 인간의 권리나 안전에 영향이 거의 없는 기타 모든 최소위험 AI 시스템의 경우에는 추가적인 법적 의무 없이 기존 법규에 따라 개발 및 사용을 허가하였으나, 해당 시스템을 개발하고 운영하는 업계에서 자율적인 행동강령을 수립하고 준수할 것을 권장하였다.

3.1.3 규제의 방법 및 내용 – 유럽 인공지능 위원회(EAIB, European Artificial Intelligence Board)의 설치, 국가별 규제 및 벌금의 부과

EU AI법안은 고위험 AI 시스템에 대한 적합성 평가를 비롯한 규제를 원칙적으로는 각 회원국이 지정 또는 설치하는 감독당국에 의하여 시행되도록 규정하면서도, 각 당국이 효율적으로 협력하고, 새로운 문제에 대한 지침과 분석을 조정하여 법안의 일관성 있는 적용을 보장하기 위한 기구로 유럽 인공지능 위원회의 설치를 예정하고 있다.

적합성 평가 절차가 진행 중이라고 하더라도 감독당국에서 필요한 일정한 사유에 따라 해당 회원국 내에서의 시장 출시와 서비스 제공을 승인할 수 있도록 하여 규제의 유연성을 확보하되, 적합성 평가 결과 이후에도 감독당국의 필요에 따라 시장철수, 리콜 등 위험에 비례하는 조치를 사후적으로 요구할 수 있도록 하였다. 규정을 준수하지 않을 경우 최고 3000만 유로 또는 글로벌 연간 총 매출액의 6% 이하의 벌금을 부과할 수 있도록 하였다.

3.2 미국

법률과 같은 형식에 의한 엄격한 규제를 목표로 하여 논의가 오래 전부터 진행되어 온 EU와 달리, 미국은 2016년 10월 '국가 AI 연구개발 전략 계획(National Artificial Intelligence Research and Development Strategic Plan, 2019년 업데이트됨)',¹⁵⁾ 'AI의 미래를 위한 준비(Preparing for the Future of Artificial Intelligence)',¹⁶⁾ 등 백악관에서 발표한 보고서를 통하여 AI에 대한 공적 연구개발을 위한 전략적 계획을 제시하고 있다. 또 2019년 2월 트럼프 행정부의 행정명령 13859¹⁷⁾를 통하여 AI 개발 관련 미국의 리더십을 확보하고 유지하기 위한 전략을 제시하는 등 초기에는 AI에 대한 규제보다는 AI 관련 기술개발에 행정부의 관심이 집중되었다. 미국은 현재에도 비교적 규제의 유연한 운용을 지향하는 것으로 알려져 있다.¹⁸⁾

¹⁵⁾ <https://www.nitrd.gov/pubs/National-AI-RD-Strategy-2019.pdf>

¹⁶⁾ https://obamawhitehouse.archives.gov/sites/default/files/whitehouse_files/microsites/ostp/NSTC/preparing_for_the_future_of_ai.pdf

¹⁷⁾ <https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence>

¹⁸⁾ <https://news.mt.co.kr/mtview.php?no=2023050118270880859>

한편 미국 하원은 2019년 4월 AI 시스템에 사용되는 알고리즘의 편향성과 차별적 결과를 방지하기 위하여 미국 연방거래위원회(FTC, Federal Trade Committee)에게 알고리즘 관련 권한을 부여하고, 일정한 영향평가를 수행하도록 요구하는 알고리즘 책임법(Algorithmic Accountability Act of 2019)¹⁹⁾을 발의하기도 하였으나, 연방 차원에서 법률 형태의 입법을 통한 AI 규제 움직임이 구체적으로 나타나고 있지는 않다.

그에 따라 미국의 AI 관련 현행 규제는 주로 AI 유해성으로부터 국민의 권리를 보호하고 위험을 관리하기 위한 기준이 제시된 연방 차원의 이니셔티브들이 존재하는 가운데, 개별 주정부 등에서 AI 관련 규제를 개별적으로 마련하는 방식으로 구성되고 있다. 비록 법적 강제성은 없으나, 몇몇 주요 연방의 이니셔티브를 간략히 소개하면 아래와 같다.

3.2.1 AI 권리장전을 위한 청사진(Blueprint for AI Bill of Right)²⁰⁾

백악관 과학기술정책실이 미국 대중의 권리 보호를 위하여 AI 시스템의 설계·개발·보급 등에 적용될 기준을 권고한 지침이다. 해당 지침에서는 AI 시스템에 관하여 적용되어야 할 5대 원칙으로 ① 안전하고 동시에 효율적인 시스템을 구축할 것, ② 알고리즘에 의한 차별이 발생하지 않도록 공정한 방식으로 설계되고 사용될 것, ③ 데이터 주체들이 무분별한 데이터 수집으로부터 보호될 수 있어야 하며, 활용된 데이터에 관한 정보를 정보주체에게 제공할 것, ④ AI 시스템이 언제 어떻게 이용되었으며, 그 결과가 어떠한 영향을 미칠 수 있는지에 관한 충분한 설명이 이루어질 것, ⑤ 개별적인 상황에 따라 AI 대신 인간이 상대방으로 선택되어야 할 수 있어야 하며, 문제 발생시 인간의 개입에 의한 신속한 해결 조치가 취하여질 수 있을 것 등이 제시되었다.

3.2.2. 인공지능과 알고리즘 공정성 이니셔티브(Artificial Intelligence and Algorithmic Fairness Initiative)²¹⁾

미국 평등고용기회위원회(EEOC, Equal Employment Opportunity Commission)가 법무부와 함께 2022년 5월 고용주의 AI 사용에 관하여 발표한 지침이다. 이 지침은 고용주가 AI를 사용하여 채용, 성과 모니터링, 급여 또는 승진을 결정하는 경우 미국 장애인법(American Disabilities Act)을 위반하여 장애인에 대한 불법적인 차별을 초래할 수 있음을 경고하는 한편, 이를 해소하기 위한 방안으로 ① 법적으로 요구되는 합리적인 편의를 제공하여 AI 시스템 적용에 따른 편향성을 해소하기 위한 노력을 수행하고, ② 의도의 유무를 떠나 장애로 인한 불이익 부과의 가능성을 최소화하기 위하여 다양한 관점에서 충분한 변수가 고려되고 설정되도록 AI 시스템을 개발 및 운영하며, ③ AI 시스템은 직무에 필요한 능력이나 자격 등을 객관적으로 측정하는 단계만을 수행하고, 이를 점수화하지 않고 직접적인 수치로만 측정하도록 하며, 등급 등의 평가행위는 해당 측정 결과를 바탕으로 인간의 개입에 의하여 직접 평가하도록 하는 등의 방안을 권고하였다.

19) <https://www.congress.gov/bill/116th-congress/house-bill/2231/all-info>

20) <https://www.whitehouse.gov/wp-content/uploads/2022/10/Blueprint-for-an-AI-Bill-of-Rights.pdf>

21) <https://www.eeoc.gov/laws/guidance/americans-disabilities-act-and-use-software-algorithms-and-artificial-intelligence>

3.3 일본

AI 규제에 틀에 관한 일본 정부의 공식 입장은 2019년 3월 발표된 '인간 중심의 AI 사회 원칙 안'²²⁾에 제시된 AI 사회 실현을 위한 7가지 원칙으로 정리되어 있다. 위 7개의 원칙은 ① AI의 이용이 인간의 기본적 인권을 침해해서는 안되고, 인간 스스로가 판단의 주체가 되어야 한다는 '인간 중심의 원칙', ② 양극화가 발생하지 않도록 다양한 계층에서 폭넓은 교육이 이뤄지고 상호작용할 수 있는 환경이 조성되어야 한다는 '교육, 리터러시 원칙', ③ 개인정보가 이용될 경우 정보주체의 자기결정권이 존중되고 프라이버시가 확보되어야 한다는 '프라이버시 확보 원칙', ④ AI 이용에 따른 위험 평가와 위험 저감을 위한 연구개발이 추진되고 리스크 관리가 되어야 하며, 특정 AI에 대한 의존이 이루어져서는 안된다는 '보안확보 원칙', ⑤ 지배적인 지위를 이용한 부당한 데이터 수집 및 정보주체의 권리 침해가 발생하여서는 안되며, AI에 의한 부와 사회적 영향력의 부당한 편중이 발생해서는 안된다는 '공정경쟁 확보의 원칙', ⑥ AI는 인종, 성별, 국적, 정치적 신념, 종교 등에 따라 부당하게 차별해서는 안되며, AI를 이용하고 있다는 사실 및 AI에 이용되는 데이터의 취득 방법 및 사용 방법이 충분히 설명될 수 있어야 하고 AI 산출물의 적정성을 담보할 수 있어야 한다는 '공평성, 설명책임 및 투명성의 원칙', ⑦ 데이터가 독점되지 않고 효율적으로 이용될 수 있는 환경이 마련되어야 하며, AI 개발이 가속화될 수 있도록 다양한 측면에서의 제도적 뒷받침이 이루어져야 한다는 '혁신의 원칙' 등이다.²³⁾ 일본 또한 미국과 마찬가지로 법률 등의 방법을 통한 AI에 대한 구체적 규제에 대하여는 EU에 비하여 상대적으로 소극적인 입장으로 알려져 있으나,²⁴⁾ 향후 AI 관련 규제가 마련될 경우 위 7가지 원칙을 보다 구체화하는 형태로 마련될 것으로 예상된다.

4. 인공지능 기본법 제정 관련 국내 동향

AI 규제에 관한 우리나라의 논의는 주로 AI 규제 및 산업진흥 일반에 관한 부분과 AI 개발 등의 과정에서 발생하는 권리침해에 대한 부분으로 나뉘어 진행되고 있다.

전자의 경우 과학기술정보통신부의 주도로 AI 산업에 대한 법제도적 뒷받침과 관련 생태계 구축을 위한 근거법 마련을 위하여 기존 국가정보화 기본법을 2020년 6월 9일 지능정보화 기본법으로 개정하고, 지능정보 기술 기준 및 정보보호시스템의 성능·신뢰도 관련 기준 등을 정비한 시행령을 법 개정과 동시에 시행하였다. 2020년 4월에는 「지능정보사회 윤리 가이드라인」을 발표하여 ① 공공성, ② 책무성, ③ 통제성, ④ 투명성의 4대 원칙 및 각 원칙별 개발자·공급자·이용자 등의 주체별 준수사항들이 규정된 세부지침을 마련하였다.²⁵⁾ 같은해 12월에는 대통령 직속 4차산업혁명위원회 전체회의에서 「인공지능(AI) 윤리기준」을 마련하여 AI 전체 주기에 걸쳐 충족되어야 하는 10가지 요건으로 ① 인권보장, ② 프라이버시 보호, ③ 다양성 존중, ④ 침해금지, ⑤ 공공성, ⑥ 연대성, ⑦ 데이터 관리, ⑧ 책임성, ⑨ 안전성, ⑩ 투명성 등이 제시되었다.²⁶⁾

22) <https://ai.bsa.org/wp-content/uploads/2019/09/humancentricai.pdf>

23) 유재홍, "일본의 인공지능 전략 동향: AI 전략 2019", 월간 SW 중심사회 60, 2019.

https://www.spri.kr/posts/view/22689?code=data_all&study_type=&board_type=industry_trend

24) 위 주18 참조.

25) 과학기술정보통신부 보도자료, 「AI 윤리 기준 정리 및 실천방안 마련」, 2020. 4. 20.

26) 과학기술정보통신부 보도자료, 「과기정통부, 사람이 중심이 되는 인공지능(AI) 윤리 기준 마련」, 2020. 12. 22.

이를 바탕으로 기존 인공지능 규제에 관한 기본법의 기능 수행을 위하여 제안된 여러 법률안들에 대한 국회과학기술정보방송통신위원회의 심사 끝에 2023년 2월 14일 「인공지능산업 육성 및 신뢰 기반 조성에 관한 법률안(대안)」이 법안심사소위원회를 통과하여 전체회의 의결 및 법제사법위원회의 체계·자구심사를 거쳐 본회의에 제출될 예정이다.

한편 위와 같은 AI 대상 일반 규제 및 윤리원칙과 더불어, AI 개발 등의 과정에서 발생하는 권리침해에 대하여는 개인정보보호위원회, 특허청, 문화체육관광부 등 관계 부처의 주도로 관련 법률의 개정이 이미 이루어졌거나 개정안이 국회에 계류 중인 상황이다.

아래에서는 이미 개정된 법률 또는 국회에 계류 중인 법률안을 간략하게 소개한다.

4.1 개인정보보호법(2023. 9. 15. 법률 제19234호로 시행될 예정인 것)

2023년 2월 27일 국회 본회의를 통과하여 2023년 9월 15일 시행될 예정인 개인정보보호법은 AI 등 자동화된 결정에 대한 거부 및 설명요구권을 규정함으로써, AI에 의한 개인정보 수집 및 이용에 관한 규제를 마련하였다.

보다 구체적으로, 개정 개인정보보호법 제4조, 제37조의2는 “완전히 자동화된 개인정보 처리에 따른 결정을 거부하거나 그에 대한 설명 등을 요구할 권리”를 정보 주체의 권리로 규정하여 AI에 의한 개인정보 처리에 대한 적용거부권과 설명요구권을 명문화하는 한편, 위 자동화된 결정의 기준, 절차, 개인정보가 처리되는 방식 등을 정보주체가 쉽게 확인할 수 있도록 공개하도록 함으로써 결과적으로 개인정보를 처리하는 AI 시스템에 관하여는 상당히 높은 수준의 투명성 원칙을 관철시켰다. 다만 위 조항에 따른 구체적인 절차 및 공개 등에 관한 구체적인 사항은 시행령에 위임함으로써, 투명성 원칙의 운용과정에서의 유연성을 확보할 수 있도록 하였는데, 추후 시행령의 구체적 내용에 따라 개인정보를 처리하는 AI 시스템에서 부담하는 공개의 수준 등이 크게 달라질 것으로 예상된다.

4.2 부정경쟁방지 및 영업비밀 보호에 관한 법률(2022.4. 20. 법률 제18548호로 시행된 것)

2021년 11월 11일 국회 본회의를 통과하고 2022년 4월 20일부터 시행 중인 부정경쟁방지 및 영업비밀 보호에 관한 법률(이하 ‘부정경쟁방지법’)은 제2조 제1호 (가)목에 소위 ‘데이터 부정사용’ 행위를 신설하여 일정한 요건을 갖춘 데이터의 경우 ① 접근권한이 없는 자가 그 데이터를 부정한 수단으로 취득하거나 그 취득한 데이터를 사용·공개하는 행위, ② 계약관계 등에 따라 접근권한이 있는 자가 부정한 이익을 얻거나 데이터 보유자에게 손해를 입힐 목적으로 데이터를 사용·공개·제3자제공하는 행위, ③ 위 행위들이 개입된 사실을 알면서도 그 데이터를 취득·사용·공개하는 행위, ④ 데이터의 기술적 보호조치를 무력화하는 행위 등을 부정경쟁행위의 한 유형으로 규정하였다.

이는 AI 학습 등을 위한 데이터의 수집 및 이용이 이뤄지고 있음에도 불구하고 위법한 데이터 수집 행위에 대한 규제를 개정 부정경쟁방지법 제2조 제1호 (파)목[구 (가)목]의 ‘기타 성과도용 행위’ 등 일반조항에 의존할 수밖에 없던 문제를 해결하기 위한 것이다. 위 규정의 도입으로 데이터 수집·이용의 위법성에 관한 일응의 기준이 마련되었으나, 여전히 일반 대중에 공개되어 있는 데이터의 수집·이용에 관하여는 그 위법성의 경계가 불분명하다는 점은 한계로 지적된다.

4.3 저작권법 전부개정법률안(의안번호 2107440, 도종환 의원 대표발의)²⁷⁾

2021년 1월 15일 도종환 의원 등 13인에 의하여 발의된 저작권법 전부개정안은 2006년 이후 약 15년 만에 추진되는 저작권법 전부개정안으로, AI 학습을 위한 저작물 이용의 면책 규정을 신설하는 내용을 포함하고 있다(개정법률안 제43조).

앞서 논의된 것과 같이 AI 학습을 위한 과정에서 다양한 형태의 데이터 수집 및 이용이 이뤄지고 있는 상황에서 학습의 대상이 되는 데이터가 저작권법상 보호되는 저작물에 해당하는 경우 저작권에 대한 침해 여부가 문제되어 왔으며, 이에 대한 공정이용 조항(저작권법 제35조의5)의 적용 여부가 불분명하다는 지적이 있어 왔다. 이에 저작권법 전부개정법률안은 AI 학습을 위한 이용에 관한 면책 근거를 마련하는 동시에 허용되는 이용행위의 요건을 명확히 하여 AI 개발자들의 필요와 저작권자의 권익 사이에 균형을 이룰 수 있도록 의도하고 있다.

문화체육관광부 또한 위 저작권법 개정안을 지지하는 입장으로 알려져 있으므로,²⁸⁾ 국회의 회기 종료로 위 전부개정법률안이 폐기되더라도 AI 학습 관련 면책 조항이 포함된 방향으로의 저작권법 개정은 향후에도 지속적으로 시도될 가능성이 있다.

4.4 인공지능 육성 및 신뢰 기반 조성에 관한 법률안(의안번호 2111261, 정필모 의원 대표발의, 위원회 대안반영예정)²⁹⁾

2023년 2월 14일 국회 과학기술정보방송통신위원회 법안2소위원회를 통과한 「인공지능 육성 및 신뢰 기반 조성에 관한 법률안(대안)」³⁰⁾은 통과될 경우 인공지능 규제 관련 기본법의 기능을 할 것으로 예상되는 법률안이다. 이 법률안은 AI 사업자가 준수하여야 할 윤리기준으로 ① 인간의 기본권을 침해하지 않을 것, ② 공동체 전체의 선(善)을 침해하지 않을 것, ③ 생태계의 지속가능성을 침해하지 않을 것, ④ 정보통신기술 및 생명과학기술 윤리에 관한 사항을 준수할 것, ⑤ 개발되는 기술 및 서비스의 목적과 기능을 설정하고 이를 준수할 것, ⑥ 사용연한을 정하고 폐기에 대한 지침을 마련할 것을 요구(제7조)하는 한편, 의료, 전기·가스·상수도 공급, 범죄 수사, 원자력시설의 관리·운영 등을 포함하여 사람의 생명·신체에 위험을 줄 수 있거나 인간의 존엄성을 해할 위험이 있는 특수한 영역의 인공지능에 대하여는 이용자에 대한 고지의무(제20조), 개발·제조·유통·서비스 제공되는 인공지능에 대한 신고의무(제21조)를 부과하고 위반하면 폐업을 포함한 여러 행정조치를 취할 수 있도록 하고 있다.

우리 국회에서 심사 중인 AI 법률안은 EU AI 법안과 비교할 때 AI의 위험성 제거와 규제의 내용을 포함하고 있으나, AI에 대한 규제보다는 산업 육성과 지원에 상대적으로 초점이 맞추어진 것으로 평가된다. 다만 위 법률안의 구체적인 내용은 관련 위원회 전체회의 및 법제사법위원회 등을 거치는 과정에서 달라질 여지가 상당하므로 향후 지속적인 추적관찰이 필요할 것이다.

²⁷⁾ https://likms.assembly.go.kr/bill/billDetail.do?billId=PRC_Q2T1M0X1D0M4W1T4M3O0R3Y4C7O3D2

²⁸⁾ https://www.mcst.go.kr/kor/s_notice/press/pressView.jsp?pSeq=18417

²⁹⁾ https://likms.assembly.go.kr/bill/billDetail.do?billId=PRC_Y2B1M0R6G2I2P1B0V2X9H4Z0X3M3J2

³⁰⁾ <https://www.assembly.go.kr/portal/bbs/B0000051/view.do?nttlId=2095056&menuNo=600101&sdate=&edate=&pageUnit=10&pageIndex=1>

5. 맺음말

본고에서는 인공지능 기본법의 주요 쟁점들 및 관련 주요 국가들의 동향과 우리나라의 현황을 간략히 소개하였다.

구체적 법률안이 확정되어 표결 등 구체적인 제정 단계를 앞두고 있는 EU를 제외하면, 우리나라를 포함한 많은 국가들에서는 AI 규제의 필요성이 다양한 영역에 걸쳐서 논의되면서도 한편으로는 AI에 대한 규제로 관련 산업의 발전이 저해될 것을 우려하여 법적·제도적 장치들을 통한 일반적인 규제 틀을 마련하는 것에는 아직 조심스러운 단계로 파악된다. 다만 AI로 인하여 발생하는 구체적 권리침해에 관하여는 개인정보보호법, 부정경쟁방지법, 자동차 손해배상보장법 등 개별 국면을 규율하는 법률들의 개정을 통하여 조심스럽게 접근이 시도되고 있다.

다만 아직 법률안 형태의 제정 단계에 있지 않은 국가들에서도 AI 개발 및 이용에 관하여 ‘안전성’, ‘비차별성’, ‘투명성’ 등과 같은 공통적 원칙이 일관되게 제시되고 있음을 고려하면, 현재 원칙, 이니셔티브, (구속력 없는) 권고안 등 여하한 형태를 불문하고 정부에서 공식적으로 발표되고 있는 문서에 포함된 위 내용들은 향후 AI 규제의 제도화가 본격적으로 진행될 경우 규제의 중요한 골격을 이룰 것으로 예상된다.

특히 2018년 EU의 GDPR이 제정된 뒤 이에 영향을 받아 한국의 개인정보보호법 등 여러 국가의 관련 법률의 개정을 통하여 개인정보 주체의 권리가 대폭 강화되고 그 과정에서 개인정보 처리에 관한 설명요구권 등이 본격적으로 도입되었던 점을 고려하면, EU의 AI법안이 통과될 경우 다른 국가들의 입법 또한 EU AI법안과 유사한 내용으로 급물살을 탈 가능성도 있을 것으로 예상된다.

규제 불확실성을 해소함으로써 관련 기술의 발전을 통한 산업 진흥과 육성을 도모하면서도 동시에 AI 학습 데이터의 권리자, AI 기술의 이용자를 포함한 이해관계자들의 권리가 충분히 보호될 수 있도록 균형 잡힌 인공지능 기본법이 조속히 입법되고 시행될 수 있기를 기대한다.

※ 출처: TTA 저널 제207호