

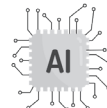
이슈페이퍼 2024-1

미국의 AI 개발 및 이용에 관한 행정명령: 주요내용과 의생명산업에 대한 시사점

유계환 / 한국지식재산연구원 연구위원

이원복 / 이화여자대학교 법학전문대학원 교수





이슈페이퍼 2024-1

미국의 AI 개발 및 이용에 관한 행정명령: 주요내용과 의생명산업에 대한 시사점

※ 이 번역의 저작권은 저희 생명의료법연구소에 있습니다. 이 글을 인용할 경우에는 다음과 같이 출처 표시를 해주시기 바랍니다: 이화여대 생명의료법연구소, "미국의 AI 개발 및 이용에 관한 행정명령: 주요내용과 의생명산업에 대한 시사점", 이슈페이퍼 (2024)

※ 이 글은 저자들의 개인적인 의견이며 생명의료법연구소의 의견과 일치하지 않을 수 있습니다.

I 검토 배경

■ (개요) 2023년 10월30일, 미국 바이든 대통령은 『안전성, 보안성 및 신뢰성을 갖는 AI의 개발과 활용에 관한 행정명령 (Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, Executive Order 14110』 (이하 ‘행정명령 14110’)¹에 서명함

• 행정명령 14110은 전 세계 국가가 인공지능의 가능성을 두고 경쟁하는 중요한 시기에 발령된 것으로, AI 모델이 일반에 공개되기 전에 그 안전성을 확보하여 국가 존망의 위협이나 정보 유출을 방지하고 소비자 및 노동자의 권리를 지키는 것을 목적으로 하고 있음²

- 행정명령 14110은 미국 연방차원에서 발령한 최초의 법적 구속력을 가진 AI 규제로, 50개 이상의 연방 정부기관에 대해 약 150개의 요건을 부과하는 광범위한 명령임

- 행정명령 14110은 바이든 정부의 책임 있는 혁신을 위한 포괄적 전략의 일환으로 안전하고 신뢰할 수 있는 AI 개발을 추진하기 위한 주요 AI 기업들의 자발적인 약속들을 기반으로 하고 있음

■ (추진 배경) AI기술 및 시스템의 지속적인 성장 및 활용 가속화와 함께 AI 시스템이 야기하는 위험성에 대한 인식이 증가하면서 전 세계적으로 AI 규제 논의가 진행 중이며,³ 그간 美정부도 안전하고 신뢰할 수 있는 AI 기술 개발을 추진하기 위해 노력해 왔음

• 2020년 12월 3일, 트럼프 대통령은 『연방정부에서 신뢰할 수 있는 인공지능(AI) 사용 촉진에 관한 행정명령 (Executive Order on Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government)』에 서명함⁴

• 2022년 10월, 백악관은 『AI 권리 장전 청사진 (Blueprint for an AI Bill of Rights)』을 공표하였으며,⁵ 동 청사진은 AI·자동화 시스템의 설계, 개발 및 보급에 관한 지침으로서, 미국 대중의 권리를 보호하기 위해 5가지 핵심 원칙을 제시하고 있음

1) 동 행정명령의 전문은 <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/>

2) 미국 대통령의 행정명령은 입법적인 기능을 갖고 있는데, 우리나라의 행정입법과 달리 법률에서 구체적으로 위임하지 않은 사항에 대하여도 권한을 행사할 수 있다는 점에서 차이가 있다. 조재현, “미국 대통령의 행정명령에 관한 연구 - 행정명령이 대통령의 입법권한 행사에 미치는 시사점을 중심으로 -”, 동아법학 제94호 (2022).

3) 2023년 6월 14일, 유럽의회가 EU AI Act 수정안 채택한 후 유럽의회, 유럽연합 집행위원회 및 유럽연합 이사회에서 최종 법률안을 협의, 지난 2023년 12월 8일 27개 EU 회원국 대표는 AI Act 법안에 잠정 합의했음. 동 AI 규제 법안은 EU 각료이사회와 유럽의회의 최종 승인을 거쳐 시행될 예정이며, 유럽의회는 내년 초 투표를 진행할 예정임. 중국은 지난 2023년 7월13일 ‘생성형 인공지능 서비스 관리를 위한 잠정 조치 (生成式人工智能服务管理暂行办法)’를 발표, 8월15일부터 시행하고 있음. 또한 지난 2023년 10월 31일에는 G7 국가가 AI 기업의 행동강령에 합의한 바 있음.

<AI 권리장전 청사진의 핵심원칙⁶⁾>

(1) 안전하고 효과적인 시스템(Safe and Effective Systems)

: 안전하지 않거나 비효율적인 시스템으로부터 보호받을 수 있을 것

(2) 알고리즘 차별로부터의 보호(Algorithmic Discrimination Protections)

: 자동화된 시스템의 알고리즘에 의한 차별에 직면하지 않도록 하며, 시스템은 공정한 방식으로 사용되고 설계되어야 함

(3) 데이터 프라이버시(Data Privacy)

: 수집된 데이터의 남용으로부터 보호받을 수 있어야 하며, 개인정보 활용에 대한 통제가 가능해야 함.

(4) 공지 및 설명(Notice and Explanation)

: 자동화된 시스템이 언제 어떻게 이용되었으며, 그 결과가 어떠한 영향을 미치는지에 관하여 알 수 있어야 함

(5) 인간의 대안, 고려사항과 대안(Human Alternatives, Consideration, and Fallback) 상황에 따라 시스템 대신 사람을 선택할 수 있도록 하여, 직면한 문제를 신속히 고려하고 해결할 수 있는 사람으로 전환 가능해야 함

- 다만 『AI 권리 장전 청사진』은 연방 정부, 이해관계자, 대중 간의 개인정보 보호에 관한 지속적인 논의를 촉진하기 위한 것이지만, EU의 AI법과 달리 AI 배포 및 시행을 위한 세부 정보 또는 메커니즘에 대한 금지가 없으므로, 민간 부문에 미치는 영향은 제한적일 것으로 평가됨

4) 동 명령에 대한 개요는 “미국 트럼프 대통령, 신뢰할 수 있는 인공지능 사용 촉진에 관한 행정 명령 서명”, 한국지식재산연구원, IP News 제2020-50 권호(2020.12.15.) (https://www.kiip.re.kr/board/trend/view.do?bd_gb=trend&bd_cd=1&bd_item=0&field=searchTC&query=%EC%8B%A0%EB%A2%B0%ED%95%A0%20%EC%88%98%20%EC%9E%88%EB%8A%94&po_item_gb=¤tPage=3&po_no=20094) 참조.

5) The White House, OSTP, The Blueprint for an AI Bill of Rights: Making Automated Systems Work for the American People, 보고서 전문은 <https://www.whitehouse.gov/wp-content/uploads/2022/10/Blueprint-for-an-AI-Bill-of-Rights.pdf> 에서 다운로드 가능

• 2023년 7월과 9월, 총 15개의 선도적인 AI 기업의 자발적인 AI 안전 서약(Voluntary Commitments)을 받는 등 AI 규제와 관련된 정책을 준비해 옴

- 2023년 7월 21일, 백악관은 7개의 주요 AI 기업(Amazon, Anthropic, Google, Inflection, Meta, Microsoft, OpenAI)을 소집하여 안전하고 투명한 AI 기술 개발을 위해 이들 기업으로부터 ‘자발적인 AI 안전 서약’을 확보했다고 발표함⁷

- 2023년 9월 12일, 백악관은 AI기술의 안전하고 신뢰할 수 있는 개발을 추진하고자 8개의 기업(Adobe, Cohere, IBM, Nvidia, Palantir, Salesforce, Scale AI, Stability)과 추가로 ‘자발적인 AI 안전 서약’을 체결함⁸

■ (특징) 이번의 행정명령 14110은 연방정부기관에 대하여 법적 구속력을 가지는 것으로 민간사업자의 권리의무나 벌칙을 직접 정하고 있는 것은 아님

• 주(州)정부 차원의 규제가 도입·시도된 적은 있으나, 연방차원에서는 구속력이 없는 가이드라인이나 기업에 의한 자율규제 중심으로 논의되어 왔음

- 행정명령 14110은 기존의 논의와 달리 향후 기존 법령 또는 새로운 입법 등을 통해 미국에 구속력이 있는 AI 규제가 도입됨을 의미하고 있어 크게 주목받고 있음

- 부루스 리드(Bruce Reed) 백악관 부비서실장(Deputy Chief of Staff)은 성명을 통해 행정명령 14110은 “AI의 안전, 보안, 신뢰성과 관련해 전 세계 정부가 취한 조치 중 가장 강력한 조치”라고 밝힘⁹

- 행정명령 14110은 기존 법령에 따라 혹은 새로운 입법 등을 통해 일부 AI 개발자에 대한 보고의무와 AI 이용자에 대한 규율 검토를 포함한 일정한 조치를 소정의 기간 내(항목에 따라 30일부터 540일 이내)에 강구하도록 관계 당국의 장에게 지시하고 있음

6) The White House, OSTP, Blueprint for an AI Bill of Rights, <https://www.whitehouse.gov/ostp/ai-bill-of-rights/> 에서 확인 가능.

7) The White House, FACT SHEET: Biden-~~Joe~~Harris Administration Secures Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI(2023.7.21.) <https://www.whitehouse.gov/briefing-room/statements-releases/2023/07/21/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-leading-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/> 에서 확인 가능.

8) The White House, FACT SHEET: Biden-~~Joe~~Harris Administration Secures Voluntary Commitments from Eight Additional Artificial Intelligence Companies to Manage the Risks Posed by AI(2023.9.12) <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/> 에서 확인 가능.

II 주요내용

■ (지도 원칙) 행정명령 14110은 Sec.2에서 안전하고 책임 있는 AI의 개발과 이용을 촉진할 목적으로 다음의 8가지의 지도 원칙(guiding principles)과 우선순위를 밝히고 있음

- AI의 안전성을 높이기 위한 테스트와 기준 확립
- 미국 시민의 프라이버시 보호
- 공정성 확보를 위한 시책의 추진
- 소비자·환자·학생 보호
- 근로자 보호
- AI를 활용한 이노베이션 촉진
- 미국의 국제적 리더십 발휘
- AI의 효과적이고 안전한 활용을 지시

• 본 이슈페이퍼에서는 행정명령 14110의 주요 구성요소를 백악관이 공개한 팩트시트의 주요 원칙 순서에 따라 서술함¹⁰

<행정명령 14110의 주요 원칙 및 해당 Section>

주요원칙 섹션	섹션
New Standards for AI Safety and Security	Sec.4.
Protecting Americans' Privacy	Sec.9.
Advancing Equity and Civil Rights	Sec.7.

9) CNBC, "Biden issues U.S.' first AI executive order, requiring safety assessments, civil rights guidance, research on labor market impact"(2023.10.30.) <https://www.cnbc.com/2023/10/30/biden-unveils-us-governments-first-ever-ai-executive-order.html> 에서 확인 가능.

10) The White House, FACT SHEET: President Biden Issues Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence, (2023.10.30) <https://www.whitehouse.gov/briefing-room/statements-releases/2023/10/30/fact-sheet-president-biden-issues-executive-order-on-safe-secure-and-trustworthy-artificial-intelligence/> 에서 확인 가능.

<행정명령 14110의 주요 원칙 및 해당 Section>

주요원칙 섹션	섹션
Standing Up for Consumers, Patients, and Students	Sec.8.
Supporting Workers	Sec.6
Promoting Innovation and Competition	Sec.5.
Advancing American Leadership Abroad	Sec.11.
Ensuring Responsible and Effective Government Use of AI	Sec.10.

1. AI 안전 및 보안에 대한 새로운 표준 확립 (New standards for AI safety and security)

■ AI 안전 및 보안을 위한 지침, 표준 및 모범 사례 등의 개발

• AI 레드팀 테스트 절차를 포함한 가이드라인과 생성 AI의 리스크 매니지먼트 프레임워크 등에 관한 가이드라인 등의 마련(section 4.1)

• 상무부 장관은 미국 국립표준기술연구소(NIST) 소장을 통해 에너지부 장관, 국토안보부 장관 및 상무부 장관이 적절하다고 판단하는 기타 관련 기관의 장과 협력하여 안전하고 신뢰할 수 있는 AI 시스템 개발을 보장하기 위한 조치를 취해야 함

- NIST는 AI 위험관리 프레임워크(AI Risk Management Framework, NIST AI 100-1)의 생성형 AI 용 보조자료를 개발하고, AI 개발자가 AI 시스템을 대중에게 공개하기 전에 테스트하기 위한 AI 레드팀 구성* 절차 지침 등을 수립할 것

* AI 레드팀 구성(AI red-teaming)이란, 통제된 환경에서 인공지능 개발자와 협력하여 인공지능 시스템의 결함 및 취약점을 찾기 위한 구조화된 테스트 노력을 의미함. AI 레드팀은 대부분 적대적인 방법을 채택하여 인공지능 시스템의 유해하거나 차별적인 결과, 예측할 수 없거나 바람직하지 않은 시스템 동작, 제한 사항 또는 시스템 오용과 관련된 잠재적 위험과 같은 결함 및 취약성을 식별하는 전담 '레드팀'에 의해 수행됨(section3(d))

- 국가 및 경제 안보, 공중 보건 및 안전에 심각한 위험을 초래하는 기반 모델(foundation model)을 개발하는 회사는 모델을 훈련할 때 정부에 알리고 테스트 결과를 정부와 공유해야 함

- 특히, ‘듀얼 유스 기반 모델’(dual use foundation models)* 개발자들은 국방생산법(Defense Production Act)에 따라 AI 시스템 공개 전 안전성 테스트(AI 레드팀)의 결과나 듀얼 유스 기반 모델의 트레이닝, 개발, 제조에 관한 정보 등을 연방 정부에 보고할 의무를 부과할 것을 요구(section 4.2(a)(i)).

* 「듀얼 유스 기반 모델」이란, 악용되면 국가의 안보, 경제, 공중위생이나 안전에 대한 심각한 위험을 초래할 수 있는 AI 모델을 의미. 구체적으로는 ①CBRN(화학, 생물, 방사선, 핵) 무기의 설계 등을 용이하게 하는 것, ②사이버상의 취약성을 발견하기 쉽게 하여 강력한 사이버 공격을 가능하게 하는 것, ③ 인간의 제어나 감시를 회피하는 것을 가능하게 하는 것 등임(section 3(k)).

- 기타 대규모 계산 클러스터(large-scale computing cluster)의 소유자 등에 대한 보고의무(section 4.2(a)(ii)), 서비스형 인프라스트럭처(Infrastructure as a Service(IaaS)) 제공자에게 외국인과의 거래에 관한 보고의무를 부과할 것(section 4.2(c))

• 생성형 AI의 위험성을 고려하여, AI가 생성한 콘텐츠를 식별하는 능력을 높이기 위한 디지털 콘텐츠 인증(digital content authentication) 및 생성 콘텐츠의 라벨링(watermarking 등) 등에 관한 지침 수립(section 4.5)

• 국토안보부는 이러한 표준을 주요 인프라 영역에 적용하고, AI 안보위원회(AI Safety and Security Board)를 설립

2. 미국인의 개인정보 보호(Protecting Americans' Privacy)

■ 의회가 모든 미국인, 특히 아동의 개인정보 보호를 강화를 위한 초당적인 데이터 프라이버시 법안을 통과시킬 것을 촉구

• 미국 국립과학재단(National Science Foundation)이 지원하는 연구 조정 네트워크(Research Coordination Network)에 자금을 지원해 암호화 도구와 같은 개인 프라이버시를 보호하는 연구 및 기술 강화

- AI가 개인에 관한 정보의 수집이나 이용을 촉진하는 것 및 개인에 관한 추론을 실시하는 것을 AI에 의한 악용가능성이 있는 개인정보 위험으로 지적(section 9(a))
- 연방 기관이 AI 시스템에 사용되는 기술을 포함하여 개인정보 보호 기술의 효과를 평가할 수 있는 지침을 개발하도록 함(section 9(b))
- 연방 정부가 프라이버시 확장기술(Privacy Enhancing Technology)의 개발을 지원할 것 등(section 9(c))

3. 평등 및 인권 증진(Advancing equity and civil rights)

■ 주로 형사 사법 및 공적 급부에 있어서의 AI의 공평한 활용에 대해서 언급

• 바이든-해리스 행정부는 이미 AI 권리 장전 청사진을 발표하고 각 기관에 알고리즘 차별을 방지하도록 행정명령을 내린 바 있으며, AI가 평등과 인권을 증진할 수 있도록 추가 조치를 지시

- AI 알고리즘이 형사 사법, 의료 및 주거 등에서 차별을 악화시키는 방향으로 사용되지 않도록 임대인, 연방 지원 프로그램 및 연방 계약자에게 명확한 지침을 제공
- AI와 관련된 인권 침해를 조사하고 기소하는 모범 사례에 대한 교육, 기술 지원, 법무부와 연방 인권 사무소(Office for Civil Rights) 간의 조정을 통해 알고리즘에 의한 차별 문제 해결
- 선고, 가석방 및 보호 관찰, 재판 전 석방 및 구금, 위험 평가, 감시, 범죄 예측 및 치안 예측, 포렌식 분석에 AI를 사용하는 모범 사례를 개발하여 형사사법시스템 전반의 공정성 보장
- 채용에 있어서의 AI의 이용에 의해 차별이 생기지 않게 하기 위한 연방 계약업자용의 지침을 공표하도록 함(section 7.3).
- 주택거래와 관련하여 관계 당국의 장관에게 AI의 이용에 의해 초래되는 차별 경감을 위한 추가 지침 마련을 장려하고 있음

4. 소비자, 환자 및 학생의 권리 보호(Standing up for consumers, patients and students)

- AI가 미국인의 삶을 개선할 수 있도록 소비자, 의료, 운수 및 교육에 있어서의 AI의 이용에 관한 규제의 책정을 관계 당국에 검토하도록 요구하고 있으며, 향후 구체적인 규제가 이루어질 가능성이 높음

- (소비자) 사기, 차별, 프라이버시에 대한 위협으로부터 소비자 보호 및 금융안정성에 대한 위협을 포함하여 AI의 이용으로부터 발생할 수 있는 위험에 대처하기 위한 규제 방안 마련할 것을 장려(section 8(a))

- 기존의 규제나 지침이 AI에 적용되는 부분을 강조하고, 명확히 규정할 것을 지시하고 있음.

- 이 명확화에는, 규제 대상 주체가 이용하는 제3자의 AI 서비스에 대해서 실사를 실시해 감시하는 책임을 명확화 하는 것이나, AI의 투명성이나 규제 대상 주체의 AI 이용에 관한 설명 능력에 대한 요구사항의 명확화가 포함된다고 하고 있어, 책임 있는 AI 이용을 위한 구체적인 규제의 도입이 예상됨.

- (의료, 공중 위생 및 복지 분야에서의 AI의 이용) 의료, 공중 보건, 보건 분야 복지와 관련한 다양한 목표 설정(section 8(b))

- HHS AI 태스크포스 설립: 보건부 장관은 국방부 장관 및 재향군인회와 협의하여 AI 태스크포스를 구성한다. 이 태스크포스는 1년 이내에 의료 서비스에 책임 있는 AI 배치에 초점을 맞춘 전략 계획을 개발한다. 여기에는 연구, 안전, 의료 서비스 제공 및 공중 보건이 포함된다. 구체적으로 의료 분야의 예측 AI 개발 및 사용; AI 안전 및 성능 모니터링; AI 기술의 평등성; 안전, 개인 정보 보호 및 보안 표준; 사용자 안내를 위한 AI 애플리케이션 문서화; 지역 보건 기관과 협력하여 AI 모범 사례 개발; AI를 사용한 업무 효율화를 다룬다.

- AI 기술에 대한 품질 보증 전략: 의료 서비스의 AI 기술이 품질 표준을 유지하도록 하는 전략을 개발. 여기에는 시판 전 평가 및 리얼월드 데이터에 기반한 시판 후 감독이 포함됨.

- 연방 차별 금지법 준수: 의료 서비스 제공자가 AI와 관련된 연방 차별 금지법을 준수하도록 하기 위한 조치를 모색. 여기에는 기술 지원, 규정 미준수 민원에 대한 대응이 포함됨.

- AI 안전 프로그램 수립: 보건부 장관은 국방부 장관 및 재향군인회와 협력하여 AI 안전 프로그램을 수립. 이 프로그램은 AI 관련 임상 오류 식별; 모범 사례 개발; 이해 관계자 교육 등을 포함.

- 신약 개발에서 AI 규제: 신약 개발에서의 AI 사용을 규제하기 위한 전략을 개발. 여기에는 규제 목표 설정, 필요한 규칙 제정, 권한 부여, 리소스 요구 사항 평가 등이 포함됨.

- (운수 관련 분야) 운수 분야에서 AI의 안전하고 책임있는 개발 및 사용 촉진을 위해 운수 분야 AI 사용 관련 정보, 기술 지원, 지침의 필요성 평가 등을 지시(section 8(c))

- 정부에 의한 자율주행 기술에 관한 현황 조사 등을 포함하는 취지

- (교육분야 AI 이용) 교육에 있어서 차별이 없도록 책임 있는 AI의 개발 및 배포 강조(section 8(d))

- 교육 분야에서는 개인 맞춤형 튜터링과 같은 AI 기반 교육 도구를 채택할 수 있도록 교육자를 지원하는 리소스를 만들어 교육을 혁신할 수 있는 AI의 잠재력을 구체화함

- 소비자 보호나 의료, 공중위생 및 복지 분야에서의 AI의 이용에서는 구체적인 규제 내용을 담고 있음에 반해, 운수와 교육 분야에 대해서는 구체적인 규제 내용은 언급하고 있지 않음

- (통신분야) (section 8(e)) 연방통신위원회(FCC)로 하여금 연방 규제가 이루어지지 않는 스펙트럼 사용의 효율성 개선, 연방 및 민간 간의 스펙트럼 공유 촉진 등 스펙트럼 관리를 강화하는데 AI를 어떻게 활용할 수 있는지 검토하도록 함. 또한 AI를 이용한 스팸 통화나 문자 및 문제를 해결하기 위한 법령 제정도 검토하도록 함.

5. 근로자 역량 강화(Supporting workers)

- AI는 미국의 일자리와 직장을 변화시키고 생산성 향상의 기대와 직장 내 감시·편견·실직의 위험성을 동시에 안겨주고 있음. 이러한 위험을 완화하고, 근로자의 단체교섭 능력의 지원과 누구나 이용 가능한 인재 교육 및 개발에 투자하기 위해 대통령은 다음과 같은 조치를 지시

- AI의 잠재적 노동시장 영향에 대한 보고서를 작성하고, AI로 인해 노동 중단에 직면한 근로자에 대한 연방 정부의 지원 강화 방안 연구 등

- 일자리 대체, 노동 기준, 직장 내 평등, 건강 및 안전, 데이터 수집에 관한 원칙과 모범사례를 개발하여, AI에 의한 근로자의 피해를 완화하고 혜택을 극대화하기 위한 원칙과 모범 사례를 개발(section 6)- 채용에 있어서의 AI의 이용에 의해 차별이 생기지 않게 하기 위한 연방 계약업자용의 지침을 공표하도록 함(section 7.3).

- 주택거래와 관련하여 관계 당국의 장관에게 AI의 이용에 의해 초래되는 차별 경감을 위한 추가 지침 마련을 장려하고 있음

- 이러한 원칙과 모범 사례를 통해 고용주가 근로자에게 과소 보상하거나, 입사 지원서를 불공정하게 평가하거나, 근로자의 조직화 능력을 저해하지 않도록 지침을 제공함으로써 근로자에게 도움이 될 것임

6. 혁신과 경쟁 촉진(Promoting innovation and competition)

- (AI 인재 유치) 기존 권한을 활용하여, 비자 기준, 인터뷰 및 심사를 현대화하고 간소화하여 고도로 숙련된 이민자와 핵심 분야의 전문성을 갖춘 비이민자의 미국 유학, 체류 및 취업 기회 확대(section 5.1)

- (혁신 촉진) 연구자와 학생이 AI 데이터에 액세스할 수 있도록 하는 도구인 National AI Research Resource 시범 운영을 통해 미국 전역에서 연구를 촉진하고, 이를 위해 의료 및 기후 변화와 같은 주요 분야의 보조금 확대(section 5.2.)

- 특허심사관 및 출원인에 대한 발명자 책임과 발명 프로세스에서의 AI의 이용에 관한 지침의 공표, AI 관련 IP 리스크 경감을 위한 연수, 분석, 평가 프로그램의 개발이나 사회적·세계적인 과제에 관한 연구에 있어서의 AI의 잠재적 역할에 관한 보고서의 공표, 미국에 있어서의 AI 규제의 사고방식을 나타내는 중요 문서 공개 등을 요구(section 5.2(c),(d),(h))

- (경쟁 촉진) 소규모 개발자와 기업가에게 기술 지원과 리소스에 대한 액세스 제공, 중소기업의 AI 혁신 기술 상용화 지원, 연방거래위원회(FTC)의 권한 행사를 장려함으로써 공정하고 개방적이며 경쟁력 있는 AI 생태계 촉진(section 5.3)

- AI 시장에서의 공정한 경쟁을 확보하고 AI에 의한 손해로부터 소비자와 노동자를 보호하기 위해 연방 거래위원회의 권한 행사가 장려되는 취지로, 향후의 권한 행사의 방침을 언급하고 있음

- 미국은 경쟁법을 적극적으로 역외 적용하고 있으므로, 동 규정은 향후 우리 기업에도 적용될 가능성이 있음

7. 해외에서 미국의 리더십 증진(Advancing American leadership abroad)

- 국무부는 상무부와 협력하여 AI의 이점을 활용하고 위험을 관리하며 안전을 보장하기 위한 강력한 국제 프레임워크를 구축하기 위한 노력을 주도

- AI 협력을 위한 양자, 다자, 다중 이해관계자 참여를 확대
- 국제 파트너 및 표준 기관과 함께 중요한 AI 표준의 개발 및 구현을 가속화하여 기술의 안전, 보안, 신뢰성, 상호 운용성을 보장
- 지속 가능한 개발을 촉진하고 중요 인프라에 대한 위험을 완화하는 등 글로벌 과제를 해결하기 위해 해외에서 안전하고 책임감 있고 권리를 보장하는 AI 표준 개발 및 배포 추진

- 미국이 AI의 글로벌 스탠다드를 만들어 나가기 위한 각종 대책이 열거되고 있는 것 외에, 안전하고 책임감 있고 지속가능한 AI 발전을 위한 동맹국 및 파트너와의 제휴의 중요성이 강조되고 있음

8. 정부의 책임감 있고 효과적인 AI 사용 보장(Ensuring responsible and effective government use of AI)

- (美연방정부에 대한 지침) AI는 규제, 관리, 혜택 지급을 위한 기관의 역량 강화, 비용 절감, 정부 시스템 보안 강화 등 장점이 있으나, AI의 사용은 차별이나 안전하지 않은 결정과 같은 위험을 초래할 수 있으므로, 정부 전반에 걸쳐 AI 전문가의 신속한 채택을 가속화하고, 권리와 안전을 보호하기 위한 명확한 표준과 부처 및 기관의 AI 사용에 대한 명확한 지침을 제공하도록 함

• AI 이용에 관하여 각 기관에 대하여 다음과 같은 사항을 권고하고 있음(section 10.1(b)(viii))

- 생성형 AI를 위한 AI 레드 팀을 포함한 AI의 외부 테스트
- 생성형 AI에 대해 차별적이거나, 오해를 초래하거나, 선동적이거나, 안전하지 않거나 기만적인 생성물에 대한 테스트 및 보호 조치, 아동 성적 학대 자료의 작성에 대한 테스트 및 보호 조치 및 개인의 동의 없이 생성적인 이미지(식별가능한 개인의 신체 또는 신체의 일부의 성적인 디지털 묘사 포함)를 생성하는 것에 대한 테스트 및 보호 조치
- 생성형 AI의 생성물에 워터마크(watermark) 및 기타 라벨을 지정하기 위한 합리적인 조치
- 필요한 최소한의 위험 관리 관행. 여기에는 필요에 따라 AI 권리 장전 청사진과 NIST의 AI 위험 관리 프레임워크에서 파생된 사례 포함(section 10.1(b)(iv))
- AI 제품의 효율성과 위험 완화에 대한 공급업체의 주장에 대한 독립적인 평가
- 조달한 AI의 문서화 및 감독
- 기관에 대한 가치를 극대화할 수 있도록 조달된 AI의 지속적인 개선을 위해 계약자에 인센티브 제공
- AI와 관련된 행정명령 및 기타 문헌에 명시된 원칙에 따라 AI에 대한 교육 실시

• 인사관리처(Office of Personnel Management), 미국 디지털 서비스(U.S. Digital Service), 미국 디지털 부대(U.S. Digital Corps), 대통령 혁신 펠로우십(Presidential Innovation Fellowship) 이 주도하는 정부 차원의 AI 인재 확충의 일환으로 AI 전문가의 신속한 채용을 시작하여야 함(10.2.(c))



시사점

1. 전반적인 평가

- 향후 미국 내에서는 행정명령 14110으로 지시된 내용에 근거해 각 정부기관이 구체적인 조치를 강구하게 되며, 선도적인 AI 기술시장인 미국의 AI 규제 법률은 미국에 진출한 전 세계 기업에 영향을 미칠 것이므로, 우리 기업에 대한 영향을 고려하며 미국의 규제 동향을 지속적으로 주시할 필요가 있음.

- AI의 개발이나 이용에 관한 법규제를 도입·검토하는 움직임이 각국에서 가속화될 가능성이 존재함.

* 27개 EU 회원국 대표는 2023년 12월 8일, AI Act 법안에 잠정 합의한 바 있으며, 중국은 ‘생성형 인공지능 서비스 관리를 위한 잠정 조치’를 2023년 8월15일부터 시행 중인 상황으로, 일부 국가에서 규제가 도입되는 움직임을 보이고 있음.

- 각국의 정책방향은 다양하나, 이번 바이든 대통령의 행정명령 14110은 앞으로 미국도 연방정부 차원에서 개입할 것을 시사한 바, 우리나라를 비롯해 다른 주요 국가의 정책에도 영향을 미칠 수 있을 것임.

- 행정명령 14110은 (1) 합법적인 용도로 사용하기 위하여 개발되는 기반모델이 생화학무기를 만든다거나 사이버 공격에 악용되는 등의 위험을 구체적으로 명시하는 등 AI가 야기하는 국가 안보 문제를 상당한 강조하고 있는 점, (2) 일부 산업 분야에만 적용되는 것이 아니라 산업 전반에 적용되는 것을 예상하고 있지만, 의료·교육·운수·통신 등 몇 개 산업 분야에는 상당히 구체적인 계획을 제시하는 점 등은 EU의 AI Act 법안과의 눈에 띄는 차이이며, 미국 바이든 행정부가 AI 규제에 대하여 갖고 있는 시각을 시사함.

- 또한 직접적인 규제를 하는 대신 표준(standard) 확립에 적극적으로 개입할 의사를 보인 것은 의미가 있는 부분. GDPR을 통하여 EU가 개인정보 보호 규범에서 국제적인 주도권을 확보한 현상이¹¹ 되풀이되는 것을 방지하지는 않으려는 의도로 읽힘.

- 한편, 지도 원칙 중 채용에 있어서의 AI의 이용에 의해 차별이 생기지 않게 하기 위한 연방 계약업자용의 지침(section 7.3)이나, 각 부처에 대한 AI 이용 지침(section 10.1(b)(viii))은 각 부처에 대한 규제 또는 권고에 해당하지만, AI 이용의 리스크 관리의 관점에서 우리 기업들도 참고할만한 내용이라고 사료됨.

2. 의료 및 의생명과학 분야

- 이번 행정명령 14110에서 의료는 교육, 운수, 통신과 더불어 구체적으로 언급되는 소수의 업역 가운데 하나이고, 그 가운데에서도 의료에 대한 조항인 Sec. 8(b)는 다른 업역에 대한 조항에 비하여 분량이 상대적으로 많은 것으로 보아 의료 분야의 AI 개발에 대하여 미국 정부가 주목을 하고 있는 것으로 보임.
- 그렇다고 의료 분야에 대한 행정명령의 성격을 전적으로 “규제”라고 규정하기는 어려움. 진흥에 관한 많은 내용을 담고 있을 뿐만 아니라, 행정명령 14110의 모든 내용이 그러하듯 私人的 행동을 직접적으로 구속하는 것이 아니라 행정기관의 행동을 명령하는 내용이기 때문임. 다만, 의료는 대표적인 규제 분야이고, 전국민 의료보험이 없는 미국에서도 Medicare와 같은 공적 의료가 의미 있는 비중을 차지하므로, 행정기관에 대한 명령이 간접적으로 민간 의료에도 영향을 미치게 될 것임. 예컨대 보건부 장관으로 하여금 보건부가 제공하는 의료사회복지 프로그램에서 자동 의사결정이나 알고리즘으로 인한 처분에 대하여 수급자가 충분히 이해하고 불복할 수 있는 절차를 마련하라고 한 경우가 그러함.¹²
- 의료와 관련하여 행정명령 14110에서 강조되고 있는 가치인 차별 금지 및 평등, 프라이버시, 안전 등은 AI가 의료에 본격적으로 등장하기 이전부터 오랫동안 미국 의료에서 지적되어 온 문제들로, AI에 의하여 문제가 심화되는 것을 방지하고 역으로 AI를 통하여 해결책을 찾으려는 취지가 읽힘. 건강 불평등 (health disparity)은 미국 의료에서 가장 심각한 문제 중 하나이고 인종 사이의 경계를 두고 발생한다는 맥락을 배경에 두고 이를 이해하여야 함.
- 범용 기반 모델이 국가 안보를 위협하는데 악용되지 않도록 상당히 자세한 주문을 하고 있는데, 여기에는 오믹스(omics) 연구, 합성 생물학(synthetic biology) 연구 등 매우 구체적인 의생명과학 분야의 연구가 언급되고 있으므로,¹³ 국가 안보와 직접적인 관련이 없는 의생명과학 분야의 AI 개발에 대하여 미국 정부의 개입이 발생할 것으로 예상됨.

11) 컬럼비아 대학의 Anu Bradford 교수는 이를 “브루셀 효과”라고 칭했다. Bradford, A., “The Brussels Effect,” Nw. UL Rev., Vol. 107 (2012).

12) Sec. 7.2 (b).

13) Sec. 4.4.

부록

행정명령 14110의 구조

Section	주요내용
1. Purpose.	• 인공지능의 안전하고 안심되며 신뢰할 수 있는 개발과 이용에 관한 행정명령 14110의 목적 기술
2. Policy and Principles	• 정책 방침으로서 다음과 같은 8가지 지도원칙 및 우선순위 제시 - (a) AI 안전 및 보안에 대한 새로운 표준 확립 - (b) 혁신과 경쟁 촉진 - (c) 노동자 지원 - (d) 평등 및 인권 증진 - (e) (의료, 법률 등 주요 분야에서의) 미국인의 이익보호 - (f) 미국인의 개인정보 및 자유 보호 - (g) 정부의 책임감 있고 효과적인 AI 사용 보장 - (h) 해외에서 미국의 리더십 증진
3. Definitions. For purposes of this order	• 행정명령 14110에서 사용하고 있는 다양한 용어의 정의
4. Ensuring the Safety and Security of AI Technology	• AI의 안전성과 보안에 관한 가이드라인, 표준 및 모범사례 개발 • 안전하고 신뢰할만한 AI 보장 • 주요 인프라 및 사이버 보안에서의 AI 관리 • AI와 CBRN(Cheical, Biological, Radiological and Nuclear) 위협의 교차 부분에 있어서의 리스크 경감 • 합성 콘텐츠로 인한 위협 경감 • 널리 이용 가능한 모델 가중치를 이용한 이중 용도 기초 모델(Dual-Use Foundation Models)에 관한 의견 요청 • AI 훈련을 위한 연방 데이터의 안전한 공개 촉진 및 악의적 이용의 방지 • 국가안보각서(National Security Memorandum)개발의 지휘
5. Promoting Innovation and Competition	• 미국으로의 AI 인재 유치 • 혁신 촉진 • 경쟁의 촉진

Section	주요내용
6. Supporting Workers	• 노동자 지원을 위한 조치 강구
7. Advancing Equity and Civil Rights.	• 형사사법제도에서의 AI와 인권 강화 • 정부 혜택 및 프로그램과 관련된 인권 보호 • 광범위한 경제에 있어서의 AI와 인권 강화
8. Protecting Consumers, Patients, Passengers, and Students.	• 소비자, 의료(환자), 운수(승객) 및 학생의 보호
9. Protecting Privacy.	• AI에 의해 잠재적으로 악화되는 개인정보 보호와 관련된 위험 경감 조치
10. Advancing Federal Government Use of AI.	• AI 관리에 대한 지침 제공 • 정부의 AI 역량강화를 위한 조치
11. Strengthening American Leadership Abroad.	• 해외에서의 미국의 리더십 강화 방안
12. Implementation.	• 효율적인 실행을 위해 ‘백악관 AI 위원회(White House AI Council)’ 설치
13. General Provisions.	• 일반 조항으로서, 기존 법령에 따른 각 기관의 권한 및 기능에 대한 영향 배제

규정

행정명령 14110 가운데 의료 AI 관련 규정들

제8조. 소비자, 환자, 승객 및 학생 보호

- 8.b.i. 본 행정명령의 발표로부터 90일 이내에, 미국 보건부(HHS) 장관은 국방부 장관 및 보건부 장관과 협의하여 보건부 인공지능(AI) 태스크포스를 설립하여야 한다. 설립 후 1년 이내에, 태스크포스는 보건 분야(연구 및 발견, 의약품 및 기기 안전, 의료 서비스 제공 및 비용 지급, 공중 보건 포함)에서 인공지능 및 인공지능 기반 기술의 책임 있는 배치 및 사용에 관한 정책과 프레임워크(적절한 경우 규제 조치를 포함)를 포함한 전략 계획을 수립해야 한다. 이 계획은 다음과 같은 영역에서 인공지능 배치를 촉진하기 위한 적절한 지침과 자원을 식별해야 한다:
 - 의료 서비스 제공 및 재정에서 예측 및 생성 인공지능 기반 기술을 개발, 유지 및 사용하는 것(품질 측정, 성능 개선, 프로그램 무결성, 혜택 관리 및 환자 경험을 포함)에는 인공지능이 생성한 산출물의 적용에 대한 적절한 인간 감독 고려
 - 규제기관, 개발자 및 사용자에게 제품 업데이트를 전달할 수 있는 수단을 통해 보건 분야에서 인공지능 기반 기술의 장기 안전성 및 실사용 성능 모니터링(임상적으로 관련성이 있거나 중요한 수정 및 모집단 간의 성능을 포함)
 - 보건 분야에서 사용되는 인공지능 기반 기술에 형평의 원칙을 통합하고, 새로운 모델을 개발할 때 영향을 받는 모집단에 대한 세분화된 데이터를 사용하거나 대표 모집단 데이터 세트를 채택하고, 기존 모델의 차별 및 편향에 대한 알고리즘 성능을 모니터링하며, 현재 시스템의 차별 및 편향을 식별하고 완화할 수 있도록 지원
 - 보건 분야에서 인공지능에 의해 강화된 사이버보안 위협에 대응하기 위한 조치를 포함하여 개인 식별 정보 보호를 위한 소프트웨어 개발 생명주기에 안전성, 개인정보 보호 및 보안 표준을 통합
 - 보건 분야의 현지 환경에서 사용자를 돕기 위한 문서를 개발, 유지 및 제공함으로써 인공지능의 적절하고 안전한 사용 결정

- 주, 지방, 부족, 지역 보건 기관과 협력하여 지역 환경에서 인공지능 활용을 위한 긍정적 사례 및 모범 사례를 발전

- 행정 부담을 줄이는 것을 포함하여 보건 분야 업무 효율성 및 만족도를 제고를 위한 인공지능 활용 방안 파악

■ 8.b.ii. 본 행정명령의 공표일로부터 180일 이내에, 보건부 장관은 관련 기관과 협의하여, 보건 분야에서 인공지능 기반 기술이 적절한 수준의 품질을 유지하는지 여부를 결정하는 전략을 개발해야 한다. 이 작업에는 인공지능 기반 의료 도구의 성능을 평가하기 위한 인공지능 품질 보증 정책 개발과, 실사용 데이터에 대한 인공지능 기반 의료기술 알고리즘 시스템의 성능에 대한 시장 진입 전 평가 및 시장 진입 후 감독을 가능하게 하기 위한 인프라 요구가 포함된다.

■ 8.b.iii. 본 행정명령의 공표일로부터 180일 이내에, 보건부 장관은 자신이 적절하다고 판단하는 관련 기관과 협의하여, 연방 재정 지원을 받는 의료 제공자들이 연방 차별금지법을 이해하고 준수하는 데 있어 적절한 조치를 고려해야 한다. 또한, 이러한 법률이 인공지능과 어떻게 관련되는지에 대해서도 고려해야 한다. 이러한 조치에는 다음이 포함될 수 있다:

- 인공지능과 관련된 연방 차별금지법 및 개인정보 보호법에 따른 의무와 이를 준수하지 않을 시 발생할 수 있는 결과에 대해 의료 제공자와 지불자들을 모아 기술적 지원을 제공
 - 인공지능과 관련하여 연방 차별금지법 및 개인정보 보호법의 불이행에 대한 불만이나 기타 보고에 대응하여 지침 발행 및 적절한 기타 조치
-

저자	유계환 / 한국지식재산연구원 연구위원 이원복 / 이화여자대학교 법학전문대학원 교수
발행일	2024.3
호수	2024-1
발행처	이화여자대학교 생명의료법연구소
주소	03760 서울특별시 서대문구 이화여대길 52, 이화여자대학교 법학관 418호
연락처	TEL) 02-3277-4227 FAX) 02-3277-4221
홈페이지	http://www.eible.org
이메일	admin@eible.org
디자인	www.studio7kg.com



이화여자대학교
생명의료법연구소
Ewha Institute for Biomedical Law & Ethics

