

미국 인공지능 행정명령의 정책 방향과 거버넌스 작동 구조

정책 흐름과 실행 메커니즘의 통합적 분석

서홍석 | 법제처 법제교류협력담당관, 변호사

1. 서론

최근 생성형 인공지능 기술의 급속한 발전에 대응하여 각국은 다양한 방식으로 법제와 정책을 정비하고 있다. 그 중에서도 인공지능 분야에서 압도적인 기술 패권을 보유한 미국이 어떠한 방식으로 이 문제에 접근하는지는 그 자체로 중요한 분석 대상이다. 유럽연합(EU)이 「AI Act」¹⁾를 통해 위험 기반의 포괄적 규율 체계를 구축하고, 한국이 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법」을 제정하여 단일법 중심의 제도 설계를 시도한 것과 달리, 미국은 연방 차원의 통합 입법 없이 이 문제에 대응해 왔다. 연방 차원에서 「국가 인공지능 이니셔티브법(National Artificial Intelligence Initiative Act of 2020)」²⁾을 중심으로, 「정부인공지능법(AI in Government Act of 2020)」³⁾, 「인공지능 교육법(AI

Training Act)」⁴⁾, 「미국 인공지능진흥법(Advancing American AI Act)」⁵⁾ 등 다수의 개별 법률이 제정되어 왔으나, 이들은 주로 연방정부의 인공지능 도입·활용, 연구개발 지원, 인력 양성 등에 관한 것으로서 인공지능 전반을 포괄하는 단일한 규율 체계를 이루고 있지는 않다.

이러한 가운데, 미국의 인공지능 정책은 대통령 행정명령(Executive Order)⁶⁾을 중심으로 정책 방향을 설정하고 이를 지속적으로 조정해 나가는 방식으로 전개되고 있다. 실제로 미국은 EO 13859(2019)⁷⁾를 시작으로 EO 13960(2020)⁸⁾, EO 14110(2023)⁹⁾, EO 14141(2025)¹⁰⁾, EO 14179(2025)¹¹⁾ 등을 거쳐, 2025년에는

4) Artificial Intelligence Training for the Acquisition Workforce Act, Pub. L. No. 117-207, 136 Stat. 2238 (2022). 이 법은 행정관리예산국(Office of Management and Budget, OMB) 국장에게 연방 행정부의 조달, 프로그램 관리 등 인공지능 관련 업무를 담당하는 공무원(covered workforce)을 대상으로 인공지능 교육 프로그램을 개발·제공하도록 요구하고, 해당 프로그램을 통해 인공지능의 작동 원리, 활용 가능성 및 위험요소에 대한 이해를 확보하도록 하는 것을 주요 내용으로 한다.

<https://www.govinfo.gov/app/details/COMPS-17079> (최종방문일: 2026. 4. 14.)

5) Advancing American AI Act, incorporated in James M. Inhofe National Defense Authorization Act for Fiscal Year 2023, Pub. L. No. 117-263, div. G, title LXXII, subtitle B (2022). 이 법은 연방정부의 인공지능 활용에 관한 원칙과 정책 수립을 위하여 행정관리예산국(Office of Management and Budget, OMB) 국장에게 지침 마련을 요구하고, 각 연방기관이 인공지능 활용 사례에 관한 인벤토리(AI use case inventory)를 작성·보고하도록 하는 등, 연방 차원의 인공지능 활용에 대한 거버넌스 체계를 구축하는 것을 주요 내용으로 한다.

<https://www.govinfo.gov/content/pkg/PLAW-117publ263/html/PLAW-117publ263.htm> (최종방문일: 2026. 4. 14.)

6) 행정명령(Executive Order)은 연방법 또는 헌법에 의해 대통령에게 부여된 권한을 근거로 발령되는 대통령 명령으로, 의회의 승인 없이 연방행정기관에 대해 법적 구속력을 가지는 행정지시이다. 미국 헌법 제2조제1항·제3항 및 관련 연방법률을 근거로 하며, 연방관보(Federal Register)에 공포됨으로써 효력이 발생한다.

7) 미국 인공지능 리더십 유지에 관한 행정명령(Maintaining American Leadership in Artificial Intelligence), Executive Order 13859, 84 Fed. Reg. 3967 (Feb. 14, 2019) [이하 EO 13859]. 이 행정명령은 이른바 ‘American AI Initiative’를 출범시키면서, 연방정부 차원의 인공지능 정책 기본 방향을 최초로 종합적으로 제시한 것으로 평가된다. 구체적으로 인공지능 연구개발 투자 확대, 기술 표준 개발 촉진, 인공지능 인력 양성, 공공 신뢰 확보 및 국제 협력 강화 등을 주요 정책 축으로 설정하고, 연방기관 간 협력을 통해 미국의 인공지능 분야 글로벌 리더십을 유지·강화하는 것을 목적으로 한다.

<https://www.federalregister.gov/documents/2019/02/14/2019-02544/maintaining-american-leadership-in-artificial-intelligence> (최종방문일: 2026. 4. 14.)

8) 신뢰할 수 있는 인공지능의 연방정부 활용 촉진에 관한 행정명령(Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government), Exec. Order No. 13960, 85 Fed. Reg. 78939 (Dec. 8, 2020) [이하 EO 13960]. 이 행정명령은 연방정부 내 인공지능 활용에 관한 기본 원칙을 제시하고, ‘신뢰할 수 있는 인공지능(trustworthy AI)’의 기준을 설정함으로써 공공부문 인공지능 거버넌스를 체계화한 것으로 평가된다. 구체적으로 각 연방기관에 대하여 인공지능 활용 사례에 관한 목록(AI use case inventory)의 작성·공개를 요구하고, 인공지능 시스템의 투명성·책임성 확보 및 위험 관리에 관한 지침을 마련하도록 하여, 연방정부 내부에서의 인공지능 활용에 대한 관리·감독 구조를 구축하는 것을 주요 내용으로 한다.

<https://www.federalregister.gov/documents/2020/12/08/2020-27065/promoting-the-use-of-trustworthy-artificial-intelligence-in-the-federal-government> (최종방문일: 2026. 4. 14.)

9) 안전하고 보안이 확보되며 신뢰할 수 있는 인공지능의 개발 및 이용에 관한 행정명령(Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence), Exec. Order No. 14110, 88 Fed. Reg. 75191 (Nov. 1, 2023) [이하 EO 14110]. 이 행정명령은 인공지능의 개발 및 활용 전반에 걸쳐 안전성·보안성·신뢰성을 확보하기 위한 포괄적 정책 및 규제 방향을 제시한 것으로서, 미국 인공지능 거버넌스가 기존의 정책적·권고적 접근을 넘어 실질적인 통제 메커니즘을 도입하는 전환점으로 평가된다. 특히 고위험 인공지능 시스템에 대한 안전성 검증 및 결과 보고, 생성형 인공지능 콘텐츠에 대한 표시 및 워터마킹 기술 개발, 개인정보 및 소비자 보호 강화, 연방정부의 인공지능 활용에 대한 감독 체계 정비 등을 주요 내용으로 포함하고 있으며, 연방기관뿐 아니라 민간 부문에 대해서도 간접적인 규범적 영향을 미치는 구조를 가진다.

<https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence> (최종방문일: 2026. 4. 14.)

10) 미국 인공지능 인프라 리더십 강화에 관한 행정명령(Advancing United States Leadership in Artificial Intelligence Infrastructure), Exec. Order No. 14141, 90 Fed. Reg. 5469 (Jan. 17, 2025) [이하 EO 14141]. 이 행정명령은 인공지능 기술 경쟁력 확보를 위하여 데이터센터, 전력망, 반도체 등 인공지능 인프라의 국내 구축을 촉진하는 것을 핵심 내용으로 하며, 미국의 인공지능 거버넌스가 규제 중심 접근을 넘어 산업·에너지·공급망 정책과 결합되는 단계로 진입하였음을 보여주는 것으로 평가된다. 특히 연방 토지에 대규모 인공지능 데이터센터를 구축하기 위한 부지 지정 및 민간 참여 구조를 설계하고, 청정에너지 기반 전력 공급, 송전 인프라 확충, 공급망 안정성 확보, 국가안보 차원의 보호조치 등을 포괄적으로 규정하고 있다. 다만, 이 행정명령은 이후 EO 14318에 의해 명시적으로 폐지되었다.

<https://www.federalregister.gov/documents/2025/01/17/2025-01395/advancing-united-states-leadership-in-artificial-intelligence-infrastructure> (최종방문일: 2026. 4. 14.)

AI 행동계획¹²⁾ 발표와 함께 다수의 후속 행정명령을 연달아 발령하고, 같은 해 12월에는 각 주(州)의 차원의 AI 규제에 대한 연방 차원의 대응을 담은 EO 14365¹³⁾를 발령하였다. 이러한 흐름은 미국의 인공지능 정책이 단일한 입법에 의해 안정적으로 정착되기보다는, 정치적·기술적 환경 변화에 따라 반복적으로 조정되는 특성을 가진다는 점을 보여준다. 아울러, 이 행정명령들은 단순한 정책 선언을 넘어 연방기관에 대한 구체적 임무 부여, 민간에 대한 간접적 규율, 재정 지원과 산업 정책의 동원 등을 통해 실질적인 집행 수단으로 기능한다는 점에 주목할 필요가 있다.

미국 인공지능 행정명령에 관한 기존 연구는 개별 행정명령의 내용 소개나 정책적 의의 분석에 상당부분 집중되어 있으며, 행정명령을 중심으로 정책의 형성·조정 과정과 실행 메커니즘을 통합적으로 분석한 법제적 연구는 상대적으로 제한적이다.¹⁴⁾ 한국의 인공지능기본법과의 비교를 통해 제도 설계상의 시사점을 도출하려는 시도 역시 아직은 충분히 이루어지고 있지 않다.

11) 미국 인공지능 리더십 저해요인의 제거에 관한 행정명령(Removing Barriers to American Leadership in Artificial Intelligence), Exec. Order No. 14179, 90 Fed. Reg. 8741 (Jan. 31, 2025) [이하 EO 14179]. 이 행정명령은 인공지능 기술 발전을 저해할 수 있는 규제적 장벽을 완화하고 혁신을 촉진하기 위한 정책 방향을 제시한 것으로서, 미국 인공지능 거버넌스가 안전성 및 통제 중심의 접근에서 기술 경쟁력과 산업 발전을 중시하는 방향으로 재조정되고 있음을 보여주는 규범으로 평가된다. 특히 연방정부의 인공지능 관련 규제 및 정책을 전반적으로 재검토하여 불필요한 제한을 완화하도록 하고, 연구개발 및 민간 혁신을 촉진하는 환경을 조성하는 것을 주요 내용으로 한다.
<https://www.federalregister.gov/documents/2025/01/31/2025-02172/removing-barriers-to-american-leadership-in-artificial-intelligence> (최종방문일: 2026. 4. 14.)

12) The White House, Winning the Race: America's AI Action Plan (July 2025) (EO 14179에 따라 수립된 미국 인공지능 정책의 종합 실행계획으로서, 인공지능 혁신 가속화, 인공지능 인프라 구축, 국제 인공지능 외교·안보 주도라는 세 가지 축을 중심으로 정책 방향을 제시하고, 다수의 행정명령을 통해 이를 구체화하는 구조를 취하고 있다). <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf> (최종방문일: 2026. 4. 14.)

13) 인공지능에 관한 국가 정책 프레임워크 보장에 관한 행정명령(Ensuring a National Policy Framework for Artificial Intelligence), Exec. Order No. 14365, 90 Fed. Reg. 58499 (Dec. 11, 2025) [이하 EO 14365]. 이 행정명령은 주(州) 차원에서 확산되고 있는 인공지능 규제에 대응하여 연방 차원의 정책 일관성을 확보하기 위한 것으로서, 미국 인공지능 거버넌스에서 연방과 주 간의 규제 충돌 문제를 본격적으로 다루는 것으로 평가된다. 특히 Section 1에서는 인공지능 기업이 과도한 규제 없이 혁신할 수 있는 환경의 필요성을 강조하는 한편, 일부 주법이 인공지능 시스템의 설계 및 운영에 특정한 제약을 가할 수 있다는 점을 지적하고 있다. 이러한 맥락에서 연방정부가 국가 차원의 최소한의 통일된 정책 기준을 설정하고, 주별 규제와의 관계를 재조정하려는 시도를 보여준다는 점에서 의의를 가진다.
<https://www.federalregister.gov/documents/2025/12/16/2025-23092/ensuring-a-national-policy-framework-for-artificial-intelligence> (최종방문일: 2026. 4. 14.)

14) 국내에서 미국 인공지능 행정명령을 다룬 주요 선행연구로는 신동민, 「인공지능(AI): 트럼프 1기 대통령실 주요 문서 분석 및 2기 정책방향 조망에의 함의」, 『한국과 세계』 제7권 제2호 (2025), 521-542면 (트럼프 1기 행정부 시기 대통령실 및 대통령실 문서를 중심으로 인공지능 정책의 형성과정을 시계열적으로 분석하고, 상위 정책 문서와 하위 실행 문서 간의 체계적 연계를 통해 정책 방향이 구체화되는 구조를 제시한 연구); 김지현, 「미국 연방정부의 인공지능 행동계획 및 행정명령」, 국회도서관, 『최신외국인법정보』 제288호 (2025.12.16) (트럼프 2기 행정부의 인공지능 행동계획(AI Action Plan)을 중심으로, 관련 행정명령들을 혁신·인프라·외교·안보 등 정책 축에 따라 체계적으로 정리하고, 인공지능 정책이 규제보다는 산업 진흥 중심으로 전환되고 있음을 개관한 입법정보 보고서); 서홍석, 「디지털 권리장전 구현을 위한 미래법제 연구」, 법제처 국외훈련보고서(2025) (미국의 인공지능 관련 연방법을 및 행정명령을 중심으로 주요 제도와 정책 흐름을 개관하고, 이를 바탕으로 한국 법제에 대한 시사점을 도출한 연구) 등이 있다.
국외 문헌으로는 Nicol Turner Lee, “How the White House Executive Order on AI Ensures an Effective Governance Regime,” Brookings Institution (Mar. 28, 2024) (미국 인공지능 행정명령을 AI Bill of Rights, NIST AI 위험관리 프레임워크, 민간 자율규제 등과 결합된 ‘정부 전체 접근(whole-of-government approach)’의 일환으로 분석하고, 인공지능 거버넌스가 단일 규제수단이 아니라 복합적 정책 수단의 결합을 통해 형성된다는 점을 강조한 정책 분석 자료) 등이 있다.
이 글은 위와 같은 선행연구를 바탕으로 하되, 미국 인공지능 정책을 행정명령을 중심으로 재구성하고 그 형성·조정 과정과 실행 메커니즘을 법학적으로 분석하는 데 목적이 있다. 기존 연구가 개별 행정명령의 내용이나 정책적 의의를 설명하거나 정책 문서 전반의 구조를 분석하는 데 중점을 둔 것과 달리, 이 글은 행정명령을 상호 연관된 규범으로 파악하여 정책이 형성되고 변화하며 구체적 실행으로 이어지는 과정을 살펴본다. 특히 행정명령이 연방기관에 대한 지침 설정, 민간에 대한 간접적 규율, 후속 정책 및 입법과의 연계를 통해 실질적인 정책 수단으로 기능한다는 점에 주목하고, 이를 바탕으로 우리나라 인공지능기본법 체계에 대한 시사점을 도출하고자 한다.

본 연구는 미국 인공지능 행정명령을 중심으로 두 가지 물음을 제기한다. 첫째는, 미국 인공지능 정책은 어떠한 흐름 속에서 형성되고 조정되어 왔는가(What)이고, 둘째는, 그 정책은 어떠한 거버넌스 구조와 정책 도구를 통해 실제로 작동하는가(How)이다. 이하에서는 이 두 물음을 통합적으로 분석함으로써 EO 13859(2019)부터 EO 14365(2025)까지의 미국 인공지능 행정명령이 정책 흐름 위에서 어떤 위치를 점하고 어떤 메커니즘을 작동시키는지를 입체적으로 규명하고, 이를 바탕으로 한국 인공지능기본법의 거버넌스 설계에 대한 비교법적 시사점을 도출하고자 한다.

II. 행정명령의 정책 방향: 진자 운동적 조정 구조

1. 혁신 촉진과 글로벌 리더십: 초기 정책의 출발점(EO 13859, EO 13960)

미국의 인공지능 정책은 초기 단계에서 기술 경쟁력 확보와 글로벌 리더십 유지에 중점을 두었다. 트럼프 1기 행정부가 2019년 2월 발령한 EO 13859(미국 인공지능 리더십 유지)¹⁵⁾는 인공지능을 국가 경쟁력의 핵심 요소로 규정하고, 기관 간 역할 조정을 위한 거버넌스 체계¹⁶⁾와 AI 규제 접근방식에 대한 지침 마련 의무¹⁷⁾를 규정함으로써 후속 행정명령들의 제도적 토대를 마련하였다. 특히 상무부 장관이 NIST 소장을 통해 AI 기술 표준 개발을 위한 연방 차원의 참여 계획을 수립하도록 요구¹⁸⁾한 것은, 이후 미국 AI 거버넌스의 핵심 축으로 기능하게 되는 NIST의 제도적 위상을 확립하는 출발점이 되었다.

2020년 12월 발령된 EO 13960(신뢰할 수 있는 인공지능의 연방정부 활용 촉진)¹⁹⁾은 연방정부 내 AI 활용을 위한 9가지 원칙을 제시하고, AI 활용 사례 인벤토리 구축²⁰⁾과 기관 간 협의체를 통한 조율 체계²¹⁾를 마련하였다. 규제 강화보다는 연방정부 내 AI 활용의 책임 있는 촉진에 초점을 두었다는 점에서, 두 행

15) EO 13859 Section 1 (“Policy and Principles”). 인공지능을 국가 경쟁력의 핵심 요소로 규정하고 (a) 기술 혁신 촉진, (b) 기술 표준 개발 및 장벽 제거, (c) 인재 양성, (d) 공공 신뢰 및 시민 자유 보호, (e) 국제 환경 조성 및 기술 보호의 다섯 가지 원칙을 제시하였다.

16) EO 13859 Section 3 (“Roles and Responsibilities”). 국가과학기술위원회(NSTC) 산하 인공지능 특별위원회(Select Committee on AI)를 통해 연방 기관 간 협력을 조율하도록 하였다.

17) EO 13859 Section 6 (“Guidance for Regulation of AI Applications”). 명령 발효 후 180일 이내에 OMB 국장이 모든 연방 기관장에게 AI 관련 규제 및 비규제 접근방식에 대한 지침을 담은 메모랜덤을 발행하도록 규정하였다.

18) EO 13859 Section 6(d). 상무부 장관이 NIST 소장을 통해 AI 기술 표준 및 관련 도구 개발을 위한 연방 차원의 참여 계획을 수립하도록 하였다. 이 계획은 연방 차원의 AI 표준화 우선순위 설정, 표준 개발 기구에서의 미국의 기술 리더십 역할 확보, 관련 기회와 도전 과제 분석을 포함하도록 하였으며, 민간·학계·비정부기구 등 이해관계자와의 협의를 거쳐 수립하도록 하였다.

19) EO 13960 Section 3 (“Principles for Use of AI in Government”). (a) 합법성, (b) 목적지향성, (c) 정확성·신뢰성·효율성, (d) 안전성·보안성·회복탄력성, (e) 이해가능성, (f) 책임성·추적가능성, (g) 정기적 모니터링, (h) 투명성, (i) 책임성의 9가지 원칙을 규정하고 있다.

20) EO 13960 Section 5 (“Inventory of AI Use Cases”). 각 기관은 비밀 및 민감 정보를 제외한 AI 활용 사례 인벤토리를 구축하고 이를 기관 간 공유 및 대중에 공개하도록 하였다.

21) EO 13960 Section 6 (“Interagency Coordination”). 연방최고정보책임자협의회(CIO Council)로 하여금 각 기관이 참여할 수 있는 기관 간 협의체 및 포럼의 목록을 공개하도록 하고, 각 기관이 자신의 권한과 임무에 따라 적절한 협의체에 참여하도록 하였다.

정명령은 기술 변화에 대한 빠른 정책 대응을 행정명령이라는 수단으로 구현하는 초기 형태를 보여준다.

2. 안전·형평성 가치 중심의 규제 강화: 정책 전환과 그 내용(EO 14110)

생성형 인공지능 기술의 급속한 확산에 따른 사회적 우려가 증대되면서, 바이든 행정부는 2023년 10월 EO 14110(안전하고 보안이 확보되며 신뢰할 수 있는 인공지능의 개발 및 이용)을 발령하여 정책 방향을 안전·형평성 중심으로 전환하였다. 이 행정명령은 듀얼유즈 기반 모델²²⁾ 개발자에 대한 안전 테스트 결과 보고 의무²³⁾를 포함한 8가지 정책 원칙²⁴⁾을 제시하고, NIST에 생성형 AI 및 듀얼유즈 기반 모델에 대한 안전 지침·벤치마크 개발을 지시하였다. FTC 등 기존 집행기관을 통한 소비자 보호 및 반차별 조치²⁵⁾도 강화함으로써, 미국의 인공지능 정책을 단순한 기술 촉진에서 위험 관리와 권리 보호를 포함하는 포괄적 규율 체계로 확장하려는 시도를 보여주었다. 특히 EO 14110은 정책 원칙 수립과 기술 표준 개발 과정에서 산업계·학계·시민사회의 참여를 광범위하게 요구하고(Section 2, Section 4.1), NIST로 하여금 생성형 AI 및 듀얼유즈 기반 모델에 대한 안전 지침을 개발하게 함으로써, 공동거버넌스(co-governance)²⁶⁾적 접근을 정책 명령의 형태로 구현한 사례이기도 하다.

아울러, EO 14110은 안전·보안 외에도 형평성(equity)과 다양성(diversity)을 정책의 핵심 가치로 명시하였다. Section 2(d)는 AI 정책이 형평성과 시민권 증진에 대한 행정부의 헌신과 일치해야 함을 선언하고 알고리즘 차별 심화에 대한 우려를 명시하였으며, Section 7은 ‘형평성과 시민권 증진’을 독립된 조로 편성하여 고용·주거·형사사법 등 전 영역에서 AI로 인한 차별 방지 조치를 규율하고 있다. 이는 AI 정책이 기술적 위험 관리를 넘어 특정 사회적 가치를 제도적으로 내재화하는 방향으로 전개되었음을 의미한다.

그러나, 이 행정명령은 발령 15개월 만에 트럼프 2기 행정부에 의해 폐지²⁷⁾됨으로써, 행정명령 중심 정

책 운영이 가져오는 구조적 불안정성을 드러냈다. EO 14110 하에서 FTC가 AI 관련 기만적 영업행위를 일제 단속한 ‘Operation AI Comply’²⁸⁾의 일환으로 이루어진 Rytr LLC에 대한 동의명령(2024)— 생성형 인공지능 기반 콘텐츠 생성 도구가 허위·기만적 콘텐츠 작성에 활용될 수 있다는 점을 문제 삼아 관련 기능의 제공을 금지하는 조치를 취한 사건 —이 그 전형적인 예이다. EO 14179 발령 이후 트럼프 2기 FTC는 이 동의명령을 2025년 12월 22일 취소하였으며, 이는 AI 행동계획의 FTC에 대한 직접 지시에 따른 것으로, 기관 집행 기조가 행정명령의 정책 방향 전환에 따라 실질적으로 재편될 수 있음을 보여준다²⁹⁾.

3. 혁신 우선으로의 복귀: 이념 편향 배제와 연방-주 갈등(EO 14179, EO 14365)

트럼프 2기 행정부는 2025년 1월 23일 EO 14179(미국 인공지능 리더십 저해요인의 제거)를 발령하여 이념적 편향이나 사회적 의제로부터 자유로운 AI 개발을 정책 방향으로 선언(30)함과 동시에 AI 행동계획 수립을 지시³¹⁾하고 EO 14110(안전하고 보안이 확보되며 신뢰할 수 있는 인공지능의 개발 및 이용) 하에서 시행된 조치들의 재검토를 명하였다³²⁾.

이 기조는 2025년 7월 23일 동시에 발령된 행정명령들에서 더욱 구체화된다. 그 중에서도 EO 14319(연방정부 내 편향 AI 방지)³³⁾는 정책 방향의 관점에서 가장 주목할 만한 사례이다. 이 행정명령은 인공지능 모델에 이념적 편향(ideological bias)이 개입할 경우 그 결과가 왜곡될 수 있다는 문제의식을 전제로, ‘이념적 중립성(ideological neutrality)’과 ‘진실 추구(truth-seeking)’를 핵심 원칙으로 제시하고, 역사적 정확성(historical accuracy), 과학적 탐구(scientific inquiry), 객관성(objectivity) 등을 우선

28) FTC의 ‘Operation AI Comply’란 FTC가 2024년 9월 AI 관련 기만적 영업행위를 단속하기 위해 실시한 일련의 집행 조치를 묶어 부른 명칭으로, DoNotPay, Ascend Ecom, Ecommerce Empire Builders, Rytr, FBA Machine 등 여러 기업에 대한 조사 및 동의명령이 이에 포함된다. FTC, Press Release, “FTC Announces Crackdown on Deceptive AI Claims and Schemes” (Sept. 25, 2024), <https://www.ftc.gov/news-events/news/press-releases/2024/09/ftc-announces-crackdown-deceptive-ai-claims-schemes>. (최종방문일: 2026. 4. 14.)

29) 앞의 각주 12) 참조. AI 행동계획은 FTC에 대해 전임 행정부 시기에 개시된 조사를 재검토하고, AI 혁신에 과도한 부담을 주는 최종명령 동의명령 금지 명령을 수정하거나 취소하도록 지시하였다. FTC, Press Release, “FTC Reopens and Sets Aside Rytr Final Order in Response to the Trump Administration’s AI Action Plan” (Dec. 22, 2025). <https://www.ftc.gov/news-events/news/press-releases/2025/12/ftc-reopens-sets-aside-rytr-final-order-response-trump-administrations-ai-action-plan> (최종방문일: 2026. 4. 14.)

30) EO 14179 Section 1 (“Purpose”). 미국의 AI 리더십 유지를 위해 “이념적 편향(ideological bias) 또는 사회적 의제(engineered social agendas)로부터 자유로운 AI 시스템”의 개발이 필요함을 선언하고, 기존 AI 정책 중 혁신의 장애가 되는 것을 폐지한다고 밝혔다. Section 2 (“Policy”)는 “인류의 번영, 경제적 경쟁력, 국가안보 증진을 위해 미국의 글로벌 AI 패권을 유지·강화하는 것”을 국가 정책으로 선언하였다.

31) EO 14179 Section 4. 대통령실 과학기술보좌관(APST), AI·암호화폐 특별보좌관, 국가안보보좌관(APNSA)이 주도하고, 경제정책보좌관, OMB 국장 및 관련 부처장과 협력하여 180일 이내에 AI 행동계획을 수립·제출하도록 지시하였다.

32) EO 14179 Section 5(a). EO 14110에 따라 취해진 정책 지침 규정 명령 등을 즉시 검토하여, 본 행정명령 Section 2의 정책과 불일치하거나 장애가 되는 조치를 식별하고, 해당 조치를 중단·수정·폐지하도록 지시하였다. 즉시 폐지가 불가능한 경우 가능한 모든 면제 조치를 취하도록 하였다.

33) 연방정부에서의 편향된 인공지능 방지에 관한 행정명령(Preventing Woke AI in the Federal Government), Exec. Order No. 14319, 90 Fed. Reg. 35389 (July 28, 2025). 이 행정명령은 연방정부가 조달하는 대규모 언어모델(LLM)에 대하여 ‘진실성(truth-seeking)’과 ‘이념적 중립성(ideological neutrality)’을 요구하는 원칙을 제시하고, 이에 부합하는 인공지능 시스템만을 조달하도록 하는 기준을 설정하였다. <https://www.federalregister.gov/documents/2025/07/28/2025-14217/preventing-woke-ai-in-the-federal-government> (최종방문일: 2026. 4. 14.)

22) 듀얼유즈 기반 모델(dual-use foundation model)이란 광범위한 데이터로 학습되고, 수백억 개 이상의 매개변수를 포함하며, 다양한 맥락에 적용 가능한 AI 모델로서, 생물·화학·방사능 핵(CBRN) 무기 개발 지원, 사이버 공격 활용, 인간 통제 회피 등 안보·공중보건·경제안보에 심각한 위험을 초래할 수 있거나 그러한 용도로 쉽게 변형될 수 있는 모델을 말한다. EO 14110 Section 3(k).

23) EO 14110 Section 4.2 (“Ensuring Safe and Reliable AI”). 국방생산법(Defense Production Act, 50 U.S.C. 4501 et seq.) 상 권한을 근거로, 듀얼유즈 기반 모델 개발자에게 ① 훈련·개발 활동 현황, ② 모델 가치의 소유 및 보안 조치, ③ AI 레드팀 테스트 결과를 연방정부에 지속적으로 보고하도록 의무화하였다. 또한, 대규모 컴퓨팅 클러스터 보유자에게도 그 존재와 규모를 보고하도록 하였다.

24) EO 14110 Section 2 (“Policy and Principles”). 8가지 정책 원칙: ① 안전하고 보안이 보장되는 AI, ② 책임 있는 혁신·경쟁 협력, ③ 미국 노동자 보호, ④ 형평성과 시민권 보호, ⑤ 소비자 보호, ⑥ 프라이버시 및 시민 자유, ⑦ 연방정부의 AI 역량 강화, ⑧ 국제 AI 리더십.

25) EO 14110 Section 8 (“Protecting Consumers, Patients, Passengers, and Students”); Section 7 (“Advancing Equity and Civil Rights”). 법무부·FTC 등이 소관 분야에서 AI 관련 소비자 보호 및 반차별 조치를 취하도록 지시하였다.

26) 공동거버넌스(co-governance)란 정부·기업·시민사회 등 다양한 이해관계자가 정책 설계와 실행에 공동으로 참여하는 거버넌스 방식을 가리킨다. “Co-Governance and the Future of AI Regulation,” 138 Harv. L. Rev. 1609 (2025). <https://harvardlawreview.org/print/vol-138/co-governance-and-the-future-of-ai-regulation/> (최종방문일: 2026. 6. 17.)

27) Executive Order 14148 (“Initial Rescissions of Harmful Executive Orders and Actions”), 90 Fed. Reg. 8237 (Jan. 28, 2025), Section 2(ggg). EO 14110은 트럼프 대통령 취임 당일인 2025년 1월 20일 EO 14148 Section 2(ggg)에 의해 명시적으로 열거·폐지되었다. 이후 EO 14179 Section 5는 각 연방기관에 EO 14110에 따라 취해진 정책 지침 중 새로운 정책 방향과 불일치하는 것을 식별하여 중단·수정·폐지하도록 추가로 지시하였다. Executive Order 14179, 90 Fed. Reg. 6987 (Jan. 31, 2025), Section 5.

시하도록 요구한다. 아울러 이러한 기준을 충족하는 인공지능 시스템만을 연방 조달 대상으로 삼도록 함으로써, 연방정부의 구매력을 통해 해당 기준을 간접적으로 관철하려는 구조를 취하고 있다. 이는 인공지능이 특정 이념이나 정치적 목적에 종속되지 않고, 중립성과 진실성을 유지해야 한다는 정책적 방향을 제시한 것으로 볼 수 있다. 인공지능 모델이 이념적 편향을 내재할 경우, 그 편향은 알고리즘의 불투명성에 가려진 채 객관적 정보인 것처럼 이용자에게 전달된다는 점에서, 전통적 미디어의 편향과는 질적으로 다른 위험을 낳는다. 특히 연방정부가 AI를 정책 결정과 공공서비스 전달에 광범위하게 활용하는 맥락에서는, 편향된 AI가 국가의 행정 언어 자체를 특정 방향으로 구조화할 수 있다는 점에서 그 위험은 더욱 증폭된다. EO 14319는 이러한 문제의식에서 출발하여, 연방 조달을 매개로 AI 시스템의 중립성과 객관성을 제도적으로 담보하려는 시도로 이해할 수 있다.

나아가, EO 14319는 혁신 촉진의 관점에서도 의미를 가진다. 바이든 행정부의 EO 14110 체제 하에서는 형평성·다양성 심사, 사회적 영향 평가 등 이념적 가치 지향에 기반한 요건들이 AI 개발 및 연방 조달의 진입 조건으로 작용하였고, 이는 기업의 혁신 속도를 제약하는 요인으로 비판받아 왔다. EO 14319는 이러한 요건들을 ‘이념적 편향’으로 규정하고 배제함으로써, 사실상 AI 개발과 연방 조달에 대한 규제 장벽을 낮추는 효과를 함께 수행한다. 즉, 이념 편향 배제와 혁신 촉진은 이 행정명령에서 상호 보완적 목표로 기능하고 있다.

이러한 정책 방향의 전환은 2025년 12월 11일 발령된 EO 14365(인공지능에 관한 국가 정책 프레임워크 보장)³⁴⁾에서 AI 규제를 둘러싼 연방과 주 간의 갈등으로 확장된다. 이 행정명령은 2025년 한 해에만 미국 각 주에서 1,000건 이상의 AI 법안이 발의된 상황을 배경으로 발령되었다. 트럼프 대통령은 2025년 11월 18일 Truth Social에 “50개 주가 각각 규칙과 승인 절차에 개입한다면 오래 유지되지 못할 것”이라고 밝히며 단일 연방 기준의 필요성을 주장하였고, 12월 11일 EO 14365에 서명하면서 “하나의 규칙서(One Rulebook)”가 필요하다고 재확인하였다. EO 14365는 나아가 콜로라도주의 알고리즘 차별 방지법(SB 24-205)을 과도한 규제의 대표 사례로 직접 지목하였다³⁵⁾.

EO 14365의 발령에 이어 2026년 4월 9일 xAI는 콜로라도 연방지방법원에 SB 24-205³⁶⁾의 시행 금지

를 구하는 소송을 제기하였다. xAI는 소장에서 EO 14365의 문제의식을 원용하면서, 콜로라도법이 Grok의 출력이 특정한 주의 정책적 관점에 부합하도록 변경되도록 강제함으로써 자사의 표현의 자유를 침해한다고 주장하였다³⁷⁾. 이 소송은 행정명령 중심의 정책 전환이 민간 기업의 헌법 소송과 결합하여 주 AI 규제를 둘러싼 연방-주 갈등이 사법적 경로로 전개되고 있음을 보여주는 사례이다.

4. 소결: 진자 운동적 조정 구조의 유연성과 불안정성

이상의 흐름을 종합하면, 미국의 인공지능 정책은 혁신 촉진(EO 13859·13960) → 안전·형평성 가치 중심의 규제 강화(EO 14110) → 혁신 우선 복귀 및 이념 편향 배제(EO 14179 이후)라는 진자 운동적 조정 구조를 나타내고 있다. 행정명령이라는 수단은 이러한 반복적 재조정을 가능하게 하는 유연성을 제공하는 동시에, 정권 교체에 따라 정책 방향이 급격히 전환될 수 있다는 구조적 불안정성을 내포한다.

주목할 점은, 이 진자 운동이 단순한 반복이 아니라는 것이다. 각 전환의 국면에서 새로운 정책 수단과 법적 갈등이 추가됨으로써 미국 AI 거버넌스의 전체 구조는 전환이 거듭될수록 더욱 복잡하고 다층적인 양상을 띠게 된다. EO 14365를 둘러싼 연방-주 갈등과 xAI v. Colorado 소송이 그 대표적인 사례이다. III장에서는 이러한 정책 전환을 실제로 작동시키는 거버넌스의 구체적 메커니즘을 분석한다.

III. 거버넌스 작동 구조: 행정명령이 정책으로 전환되는 방식

1. 연방기관의 기능별 역할 분산: 안정적 핵심과 유동적 외피

(1) NIST: 정치적 전환으로부터 독립적인 기술 표준의 허브

미국 인공지능 거버넌스의 첫 번째 층위는 연방행정기관 간의 기능별 역할 분산이다. 이 구조에서 NIST(National Institute of Standards and Technology, 국가표준기술연구원)는 행정부 교체와 무관하게 AI 기술 표준과 위험관리 프레임워크의 허브로 기능한다. NIST AI RMF(AI Risk Management Framework, 인공지능 위험관리 프레임워크)는 정권 교체에 따른 정책 전환에도 불구하고 연성법 기

34) EO 14365는 주(州) 차원의 인공지능 규제로 인한 분산적 규율을 문제 삼고, 연방 차원의 최소한의 통일된 정책 기준을 마련하기 위한 방향을 제시하는 한편, 주법에 대한 소송 대응(AI Litigation Task Force), 보조금 조건 설정, 연방 기준 도입 등을 통해 주 규제를 조정·제한하려는 정책 수단을 규정하고 있다.

35) Donald J. Trump, post on Truth Social (Nov. 18, 2025), 인용: Kanishka Singh, “Trump warns against AI ‘overregulation,’ says US needs to have one federal standard,” Reuters (Nov. 18, 2025), <https://www.reuters.com/world/trump-warns-against-ai-overregulation-says-us-needs-have-one-federal-standard-2025-11-18/>; Bryan Metzger et al., “Trump signs executive order restricting states’ ability to regulate AI,” Business Insider (Dec. 12, 2025), <https://www.businessinsider.com/trump-ai-executive-order-one-rule-state-regulations-2025-12>. 2025년 한 해에만 미국 50개 주에서 1,000건 이상의 AI 관련 법안이 발의되었다는 사실도 Business Insider 기사에 명시되어 있다. (이상 최종방문일: 2026. 4. 14.)

36) Colo. Rev. Stat. §§ 6-1-1701 to 6-1-1711 (Senate Bill 24-205, enacted 2024, effective June 30, 2026). 이 법은 고용, 주거, 교육, 금융, 의료 등 중대 결정에 사용되는 고위험 인공지능 시스템의 개발자 및 배포자에게 알고리즘 차별 방지, 위험관리 프로그램, 영향평가, 소비자 공시 등의 의무를 부과한다. <https://leg.colorado.gov/bills/sb24-205> (최종방문일: 2026. 4. 14.)

37) Complaint, X.AI LLC v. Weiser, No. 1:26-cv-01515 (D. Colo. filed Apr. 9, 2026). 이 사건에서 원고는 콜로라도 AI 법(SB 24-205)이 인공지능 모델의 출력 내용을 특정 방향으로 형성하도록 강제함으로써 수정헌법 제1조에 따른 표현의 자유를 침해하고, 주의 적용을 통해 주간 통상을 과도하게 제한하며, 주요 개념의 불명확성으로 인해 적법절차에 위반되고, 다양성 증진 목적 AI 활용에 대한 예외 조항이 수정헌법 제14조 평등보호 조항에 위반된다고 주장하였다.

<https://www.courthousenews.com/wp-content/uploads/2026/04/grok-ai-bill-colorado-complaint.pdf> (최종방문일: 2026. 4. 14.) 이후 법무부는 2026년 4월 24일 1964년 민권법(Civil Rights Act of 1964, 42 U.S.C. § 2000h-2)에 따라 해당 소송에 참가하였으며(intervention), 이는 연방정부가 주 AI법을 대상으로 한 소송에 직접 참가한 첫 사례로서 EO 14365가 실제 소송 전략으로 구체화된 것으로 평가된다. Kaitlyn E. Stone et al., “DOJ Intervenes in Lawsuit Challenging Colorado’s ‘Algorithmic Discrimination’ Law,” National Law Review (May 1, 2026), <https://natlawreview.com/article/doj-intervenes-lawsuit-challenging-colorados-algorithmic-discrimination-law/>; Norton Rose Fulbright, “X.AI sues, DOJ intervenes, enforcement of Colorado’s AI Act suspended” (May 2026), <https://www.nortonrosefulbright.com/en-us/knowledge/publications/de3ad9de/xai-sues-doj-intervenes-enforcement-of-colorado-ai-act-suspended> (최종방문일: 2026. 6. 16.)

반의 기술 표준이 어떻게 안정적 핵심으로 기능할 수 있는지를 보여주는 가장 전형적인 사례이다. EO 13859와 EO 14110은 NIST를 AI 표준 형성의 중심기관으로 직접 지목하였으며, EO 13960 역시 자발적 합의 표준 체계를 통해 같은 방향을 강화하였는바³⁸⁾, 이 사실 자체가 NIST의 역할이 앞서 본 진자 운동적 정책 전환으로부터 상대적으로 독립적임을 보여준다.

NIST의 AI RMF(AI 100-1)³⁹⁾는 이 구조의 핵심 산물이다. 통제(Govern)-파악(Map)-측정(Measure)-관리(Manage)의 4개 기능으로 구성된 이 프레임워크는 법적 구속력은 없으나 정부·기업·시민사회가 공동으로 개발한 사실상의 표준으로 기능하고 있다. 이는 일률적 규제 대신 위험 기반 접근(risk-based approach)을 강조하는 미국식 거버넌스 철학을 반영하며, AI 정책의 정치적 전환이 반복되는 상황에서도 기술 규범의 연속성을 유지하는 안정적 핵심으로 작동한다. AI RMF가 정부·기업·시민사회의 공동 참여로 개발되었다는 사실은, 공동거버넌스적 개발 방식이 연성 표준의 정치적 독립성에 기여하는 한 요인임을 시사한다.

(2) 집행기관의 정책 중속성

FTC는 소비자 보호 관점에서 AI의 오용이나 기만적 행위에 대한 집행을 담당하는바, FTC의 집행 기조는 NIST의 표준 역할과 달리, 행정명령의 정책 방향 전환에 상당 부분 중속되는 특성을 보인다. 이를 가장 선명하게 보여주는 사례가 Rytr 동의명령 취소이다.

앞서 II장에서 살펴본 바와 같이, 바이든 행정부 FTC는 2024년 9월 Rytr LLC에 대한 동의명령에서, AI 생성 허위 리뷰를 생성할 수 있는 서비스를 제공하는 것 자체가 FTC법 제5조⁴⁰⁾ 위반이라는 이론을 제시하였다. 그러나, 트럼프 2기 FTC는 2025년 12월 22일 2-0 표결로 이 동의명령을 취소하면서⁴¹⁾, “단순히 잠재적으로 문제 있는 방식으로 사용될 수 있다는 이유만으로 기술이나 서비스를 단죄하는 것은 법과 질서 있는 자유에 부합하지 않는다”고 밝혔다. 이러한 번복은 AI 행동계획의 FTC에 대한 직접 지시에 따른 것으로, 집행 기관의 법적 판단이 행정명령의 정책 기조에 의해 실질적으로 재편될 수 있음을 보여준다.

38) EO 13859 Section 6(d); EO 14110 Section 4(a)(i). EO 13859는 NIST 소장을 통해 AI 기술 표준 개발을 위한 연방 차원의 참여 계획 수립을 지시하였고, EO 14110은 NIST로 하여금 생성형 AI 및 듀얼유즈 기반 모델에 대한 안전 지침 벤치마크를 개발하도록 지시하였다. EO 13960 Section 4(c)는 NIST를 직접 지목하지는 않으나, 각 기관에 자발적 합의 표준을 활용하도록 지시함으로써 같은 방향을 강화하였다.

39) National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1 (2023). AI RMF는 통제(Govern)-파악(Map)-측정(Measure)-관리(Manage)의 4개 핵심 기능으로 구성되며, 법적 구속력은 없으나 정부·기업·시민사회가 공동으로 개발한 사실상의 표준으로 기능하고 있다.

40) Federal Trade Commission Act § 5, 15 U.S.C. § 45. 동 조항은 “불공정하거나 기만적인 행위 또는 관행(unfair or deceptive acts or practices)”을 금지하는 일반조항으로서, 미국 소비자 보호법 체계의 핵심 규정이다. 여기서 ‘기만(deception)’은 합리적인 소비자를 오도할 가능성이 있고 그 오도가 중요성(materiality)을 가지는 경우를 의미하며, FTC Policy Statement on Deception (1983)에서 그 판단 기준이 구체화되어 있다.

41) 앞의 각주 29) 참조. FTC는 취소 이유로 ① 원래 complaint가 FTC법 제5조 위반을 뒷받침하지 못한다는 법적 근거와, ② 해당 명령이 신생 AI 산업의 혁신에 과도한 부담을 준다는 정책적 근거를 함께 제시하였다.

이 밖에도 법무부(AI 소송 태스크포스 운영), 에너지부(AI 안전 평가 도구 개발 및 전력망 구축), 상무부(주법 평가 및 연방 보조금 조건 설정), 국방부(국가안보 AI 거버넌스) 등 각 부처가 자신의 법적 권한 범위 내에서 AI 관련 임무를 수행하는 분산 구조가 형성되어 있다⁴²⁾.

2. 정책 도구의 결합: 규제·재정·산업정책의 통합

미국의 인공지능 정책은 단순한 규제 수단에 국한되지 않고, 다양한 정책 도구를 결합하여 작동한다. EO 14179(미국 인공지능 리더십 저해요인의 제거)의 지시에 따라 수립된 AI 행동계획⁴³⁾은 이러한 통합적 접근의 가장 명확한 표현이다. 이 계획은 AI 혁신 가속화, 인프라 구축, 국제 외교·안보 주도의 세 축을 중심으로 규제 완화, 재정 지원, 인력 양성, 기술 수출을 통합하는 구조를 취하고 있다.

EO 14318(데이터센터 인프라 연방 인허가 절차의 가속화)⁴⁴⁾은 100메가와트를 초과하는 데이터센터 프로젝트를 중심으로, 연방 인허가 절차의 신속화와 규제 간소화, 재정 지원 및 연방 부지 활용을 결합하는 방식으로 인공지능 인프라 구축을 촉진한다. EO 14277(미국 청소년을 위한 인공지능 교육 증진)⁴⁵⁾·14278(미래 고임금 숙련직 일자리를 위한 미국 근로자 준비)⁴⁶⁾은 각각 K-12 인공지능 교육 확대

42) 법무부의 AI 소송 태스크포스는 EO 14365 Section 3에 근거한다. 에너지부의 AI 안전 평가 도구 개발은 EO 14110 Section 4.1(b)에 따라 개시되어 이미 완료된 조치이며, 전력망 구축 역할은 AI 행동계획 Pillar II에 근거한다. 상무부의 주법 평가 및 BEAD 보조금 조건 설정은 EO 14365 Section 4-5에 근거한다. BEAD(Broadband Equity, Access, and Deployment)란 연방정부가 광대역 인터넷 인프라 구축을 위해 각 주에 지원하는 보조금 프로그램으로, EO 14365 Section 5는 과도한 AI 규제를 시행하는 주를 이 프로그램의 비배분 자금 지원 대상에서 제외하도록 하였다. 국방부의 역할은 AI 행동계획 Pillar II에 근거하며, 고보안 데이터센터 구축 및 DOD AI 프레임워크 정비를 포함한다.

43) 앞의 각주 12) 참조. 동 계획은 AI 혁신 가속화, AI 인프라 구축, 국제 AI 외교·안보 주도의 3개 축을 중심으로 규제 완화, 재정 지원, 인력 양성, 기술 수출을 통합하는 구조를 취하고 있다.

44) 데이터센터 인프라 연방 인허가 절차의 가속화에 관한 행정명령(Accelerating Federal Permitting of Data Center Infrastructure), Exec. Order No. 14318, 90 Fed. Reg. 35385 (July 28, 2025). 이 행정명령은 인공지능 데이터센터 및 관련 인프라를 국가안보 및 경제적 핵심 기반으로 규정하고, 연방 인허가 절차의 간소화, 환경심사(NEPA) 적용 완화, 연방 토지 활용 확대, 재정 지원 및 FAST-41 절차를 통한 신속 심사 등을 통해 데이터센터 인프라 구축을 가속화하려는 정책 방향을 제시한다. 이 행정명령은 기존의 인공지능 인프라 구축 전략을 제시하였던 EO 14141를 폐지하면서, 인프라 구축의 방향 제시에서 나아가 실제 인허가 절차와 규제 부담을 완화하는 실행 단계로 정책이 전환되었음을 보여준다. 즉, EO 14141이 데이터센터·전력망·반도체 등 인공지능 인프라의 국내 구축 필요성과 정책 방향을 제시하는 데 중점을 두었다면, EO 14318은 환경심사 완화 및 인허가 절차의 신속화를 통해 이러한 인프라 구축을 현실화하는 수단에 초점을 두고 있다는 점에서 차이를 보인다. <https://www.federalregister.gov/documents/2025/07/28/2025-14212/accelerating-federal-permitting-of-data-center-infrastructure> (최종방문일: 2026. 4. 14.)

45) 미국 청소년을 위한 인공지능 교육 증진에 관한 행정명령(Advancing Artificial Intelligence Education for American Youth), Exec. Order No. 14277, 90 Fed. Reg. 17519 (Apr. 28, 2025). 이 행정명령은 인공지능 시대에 대응하기 위한 국가 인재 양성을 목표로, K-12 교육에서의 인공지능 리터러시 확대, 교원 대상 인공지능 교육 및 연수 강화, 백악관 인공지능 교육 태스크포스 설치, 민간 협력을 통한 교육 프로그램 개발 및 ‘Presidential Artificial Intelligence Challenge’ 도입 등을 통해 인공지능 교육 기반을 구축하려는 정책 방향을 제시한다. <https://www.federalregister.gov/documents/2025/04/28/2025-07368/advancing-artificial-intelligence-education-for-american-youth> (최종방문일: 2026. 4. 14.)

46) 미래 고임금 숙련직 일자리를 위한 미국 근로자 준비에 관한 행정명령(Preparing Americans for High-Paying Skilled Trade Jobs of the Future), Exec. Order No. 14278, 90 Fed. Reg. 17525 (Apr. 28, 2025). 이 행정명령은 재산업화 및 신기술 확산에 대응하기 위한 노동력 양성을 목표로, 연방 직업훈련 프로그램의 통합·재편, 성과 중심 평가 강화, 숙련직 중심 인력 양성, 등록형 도제제도(Registered Apprenticeships)의 확대 및 신산업 분야로의 확장 등을 추진함으로써, 인공지능을 포함한 신기술 환경에 적합한 숙련 노동력 기반을 구축하려는 정책 방향을 제시한다. <https://www.federalregister.gov/documents/2025/04/28/2025-07369/preparing-americans-for-high-paying-skilled-trade-jobs-of-the-future> (최종방문일: 2026. 4. 14.)

및 직업훈련·도제제도 강화를 내용으로 하는 행정명령으로서, 인공지능 교육 및 직업훈련 정책을 혁신 전략의 핵심 요소로 통합하였으며, EO 14320(미국 인공지능 기술 스택의 수출 촉진)⁴⁷⁾은 하드웨어·데이터·모델·응용을 포함한 ‘full-stack’ 인공지능 기술 패키지의 해외 수출을 지원하는 행정명령으로서 산업 정책과 외교안보 정책을 인공지능 거버넌스 도구로 결합하는 시도를 보여준다. EO 14363(제네시스 미션)⁴⁸⁾은 연방 과학 데이터와 고성능 컴퓨팅 자원을 통합하여 인공지능 기반 과학기술 혁신을 추진하는 국가 연구개발 프로젝트로서, 이러한 정책 도구 결합의 가장 야심찬 사례에 해당한다. 이러한 다양한 정책 도구의 결합은 행정명령·재정·표준·수출·소송 등 복합적 수단이 병렬적으로 작동하는 구조를 보여주며, 미국의 인공지능 거버넌스가 단일 정책 수단에 의존하지 않는 특징을 갖고 있음을 시사한다.

3. 연방과 주의 관계: 간접 조정에서 직접 선점 전략으로

미국 인공지능 거버넌스에서 연방과 주의 관계는 전통적으로 연방 보조금이나 정책 지침을 통한 간접적 조정의 방식으로 운영되어 왔다. 그러나, 2025년 한 해에만 미국 50개 주에서 1,000건 이상의 AI 관련 법안이 발의되는 등⁴⁹⁾ 주 차원 규제가 폭발적으로 증가하면서, 연방 차원의 단일 규범 형성 필요성이 제기되었다⁵⁰⁾. 이에 연방 입법을 통해 주(州) 규제를 일원화하려는 시도가 이루어졌으나, 상원에서 99-1의 압도적 표결로 ‘빅 뷰티풀 법안’에서 10년간 주 AI 규제를 금지하는 모라토리엄 조항⁵¹⁾이 삭제되면서

좌절되었고⁵²⁾, 트럼프 행정부는 이에 대응하여 행정명령을 통한 접근으로 전환하였다.⁵³⁾ EO 14365(인공지능에 관한 국가 정책 프레임워크 보장)는 바로 이러한 맥락에서 발령된 것으로, 주 차원의 인공지능 규제를 단순히 조정의 대상으로 보는 것이 아니라, 연방 차원에서 통제하거나 배제할 수 있는 대상으로 설정하고 있다는 점에서 기존 접근과 구별되며, 연방-주 관계를 근본적으로 재구성하는 전환점으로 평가할 수 있다.

여기서 주목할 것은 EO 14365가 ‘선점(preemption)’이라는 개념을 전면에 내세우고 있다는 점이다. 선점은 미국 헌법 제6조의 최고법규조항(Supremacy Clause)⁵⁴⁾에 근거하여 연방법이 주법에 우선하고, 양자가 충돌하는 경우 주법의 효력이 배제되는 전통적 법리이다. 그러나, EO 14365에서의 선점은 이러한 법리가 확정된 형태로 적용되는 것이 아니라, 행정명령을 통해 그와 유사한 효과를 정책적으로 구현하려는 시도라는 점에서 기존의 선점 법리 적용과 구별된다. 즉, 이 행정명령은 선점 법리를 직접 적용하기보다는 이를 정책적 수단으로 차용하여, 주 규제에 대한 연방 차원의 통제 가능성을 확대하려는 접근을 취하고 있다.

EO 14365의 구조는 이러한 점을 분명히 보여준다. 이 행정명령은 주 규제에 대응하기 위하여 단일한 수단이 아니라, 소송, 재정, 입법을 결합한 다층적 구조를 설계하고 있다. 첫째, 소송 단계에서 법무부는 AI 소송 태스크포스(AI Litigation Task Force)를 설치하여 주 인공지능 법률이 기존 연방법에 의해 선점되는지 여부를 근거로 직접적인 법적 이익을 제기하도록 하고 있다(Section 3)⁵⁵⁾. 이는 주 법률의 효력을 사후적으로 통제하는 사법적 경로를 정책적으로 적극 활용하는 방식이다. 둘째, 재정 단계에서 연방 정부는 BEAD 프로그램 등 연방 보조금의 배분을 주 규제와 연계하여, 일정한 인공지능 법률을 유지하

EO 14277은 유치원생부터 고등학생까지의 AI 교육을, EO 14278은 ‘인공지능 주도 경제’에서의 근로자 재교육 및 지원을 각각 대상으로 한다는 점에서 차이가 있다.

47) 미국 인공지능 기술 스택의 수출 촉진에 관한 행정명령(Promoting the Export of the American AI Technology Stack), Exec. Order No. 14320, 90 Fed. Reg. 35393 (July 28, 2025). 이 행정명령은 미국 인공지능 산업의 글로벌 확산과 기술 패권 유지를 목표로, 하드웨어·데이터·모델·응용을 포괄하는 ‘full-stack’ 인공지능 패키지의 해외 수출을 지원하는 ‘미국 인공지능 수출 프로그램(American AI Exports Program)’을 신설하고, 외교·금융 수단을 결합하여 우선 수출 프로젝트를 선정·지원함으로써 미국 중심의 인공지능 기술 및 규범의 국제적 확산을 도모하는 정책 방향을 제시한다. <https://www.federalregister.gov/documents/2025/07/28/2025-14218/promoting-the-export-of-the-american-ai-technology-stack> (최종방문일: 2026. 4. 14.)

48) 제네시스 미션 출범에 관한 행정명령(Launching the Genesis Mission), Exec. Order No. 14363, 90 Fed. Reg. 55035 (Nov. 28, 2025). 이 행정명령은 인공지능을 활용한 과학기술 혁신 가속화를 목표로, 연방 과학 데이터와 고성능 컴퓨팅 자원을 통합한 ‘제네시스 미션’을 출범시키고, 에너지부를 중심으로 국가 연구개발 역량을 결집하여 과학 분야별 기초모델 개발, AI 기반 연구 자동화 및 국가 핵심 과제 해결을 추진하는 한편, 연방기관·민간·학계 간 협력을 통해 인공지능 기반 과학 연구 생태계를 구축하려는 정책 방향을 제시한다. <https://www.federalregister.gov/documents/2025/11/28/2025-21665/launching-the-genesis-mission> (최종방문일: 2026. 4. 14.)

49) Baker Botts LLP, “U.S. Artificial Intelligence Law Update: Navigating the Evolving State and Federal Regulatory Landscape,” 2026. 1. (미국의 주별 AI 입법 동향과 연방-주 규제 갈등을 종합적으로 정리), <https://www.bakerbotts.com/thought-leadership/publications/2026/january/us-ai-law-update>; Kevin Frazier & Adam Thierer, “1,000 AI Bills: Time for Congress to Get Serious About Preemption,” Lawfare, May 9, 2025 (2025년 미국에서 약 1,000건의 AI 관련 법안이 발의되었음을 지적하며 연방 선점 필요성을 주장), <https://www.lawfaremedia.org/article/1-000-ai-bills-time-for-congress-to-get-serious-about-preemption>(이상 최종방문일: 2026. 4. 14.)

50) Sam Altman, CEO, OpenAI, Testimony before Senate Commerce Committee, “Winning the AI Race: Strengthening U.S. Capabilities in Computing and Innovation” (May 8, 2025). 알트만은 청문회에서 “50개의 서로 다른 규제 체계를 어떻게 준수할지 상상하기 매우 어렵다”고 발언하였다. transcript available at Tech Policy Press (May 9, 2025), <https://www.techpolicy.press/transcript-sam-altman-testifies-at-us-senate-hearing-on-ai-competitiveness> (최종방문일: 2026. 4. 14.)

51) Kevin Frazier & Adam Thierer, “Understanding the Proposed AI Moratorium: Answers to Key Questions,” R Street Institute, 2025. 6. 3. (미국 의회에서 논의된 주(州) AI 규제 모라토리엄과 연방 선점 문제를 설명), <https://www.rstreet.org/ai-moratorium-questions/> (최종방문일: 2026. 4. 14.)

52) One Big Beautiful Bill Act(H.R. 1, 119th Cong.)는 트럼프 행정부의 주요 재정·정책 패키지 법안으로, 하원 통과안에는 주(州)의 인공지능 규제를 10년간 금지하는 모라토리엄 조항이 포함되어 있었으나, 2025년 7월 1일 상원에서 99대 1의 표결로 해당 조항이 삭제되었다. Omer Tene et al., Federal AI Moratorium Dies on the Vine as Senate Passes the One Big Beautiful Bill Act, Goodwin Law (July 3, 2025), <https://www.goodwinlaw.com/en/insights/publications/2025/07/alerts-practices-aiml-federal-ai-moratorium-dies-on-the-vine> (최종방문일: 2026. 4. 14.)

53) 트럼프 행정부는 입법을 통한 주 AI 규제 선점을 반복적으로 추진하였으나 모두 좌절되었고, EO 14365는 이러한 입법 실패 이후 행정명령을 통해 같은 목표를 추구한 것으로 평가된다. 구체적으로 2025년 7월 ‘빅 뷰티풀 법안(One Big Beautiful Bill Act)’에 삽입된 10년간 주 AI 규제 금지 모라토리엄 조항이 상원에서 99-1로 삭제되었고, 이후 국방수권법(NDAA)에 유사 조항을 삽입하려는 시도도 실패하였다. Benton Institute for Broadband & Society, “Trump Executive Orders Shape Federal AI Regulation and Override State Actions” (Dec. 12, 2025) <https://www.benton.org/blog/trump-executive-orders-shape-federal-ai-regulation-and-override-state-actions>; Alston & Bird, “President Trump Signs Executive Order Aiming to Curb State AI Regulation” (Dec. 19, 2025) <https://www.alstonconsumerfinance.com/executive-order-state-ai-regulation/>; DLA Piper, “White House Releases the National Policy Framework for Artificial Intelligence: Key Points” (Mar. 20, 2026), <https://www.dlapiper.com/en-kr/insights/publications/2026/03/white-house-releases-the-national-policy-framework-for-ai-key-points> (이상 최종방문일: 2026. 4. 14.).

54) U.S. Const. art. VI, cl. 2.

55) EO 14365 Section 3. 법무부 장관에게 명령 발효 후 30일 이내에 AI 소송 태스크포스를 설치하도록 지시하였다. 태스크포스의 임무는 연방 AI 정책에 반하는 주법을 식별하고, 통상 조항(Commerce Clause) 위반, 기존 연방 규정에 의한 선점(preemption), 기타 위법성을 근거로 법적 이익을 제기하는 것이다.



는 주에 대해서는 재정적 불이익을 부과할 수 있는 구조를 마련하고 있다(Section 5)⁵⁶⁾. 이는 연방 보조금을 정책 유인으로 활용하여 주의 입법 선택에 영향을 미치는 방식이다. 셋째, 입법 단계에서는 향후 연방법을 통해 주 인공지능 법률을 명시적으로 선점하는 통합적 정책 프레임워크를 마련하도록 요구하고 있다(Section 8)⁵⁷⁾. 이러한 구조는 사법적 판단, 재정적 유인, 입법적 조치를 결합함으로써, 전통적인 선점 법리의 효과를 정책적으로 재현하려는 시도로 이해할 수 있다.

아울러, 연방거래위원회(FTC)는 정책 성명을 통해 일정한 주 법률이 연방거래위원회법에 의해 선점될 수 있는 기준을 제시함으로써, 이러한 소송 및 입법 전략을 법리적으로 뒷받침하는 역할을 수행한다(Section 7)⁵⁸⁾. 이는 선점 여부에 관한 해석 기준을 행정적으로 제시함으로써, 소송 단계에서의 법적 주장과 입법 단계에서의 규범 형성을 연결하는 기능을 수행한다.

이와 같은 정책적 접근은 민간 기업의 헌법 소송으로도 이어지고 있다. 앞서 II 장에서 살펴본 xAI v. Colorado 사건(2026. 4. 9.)은 EO 14365가 형성한 정책 방향과 궤를 같이하는 사례로서, 원고는 EO 14365의 문제의식을 인용하면서 콜로라도 주의 인공지능 규제가 헌법에 위반된다고 주장하였다. 이는 행정명령이 민간 주체의 소송 전략과 결합하여 주 법률의 효력을 직접적으로 다투는 경로를 형성하고 있음을 보여준다.

그러나, 이러한 방식은 헌법적 한계를 수반한다. 연방 보조금과 주 정책을 연계하는 방식은 연방 보조금에 부가된 조건이 주의 실질적 선택 가능성을 박탈하는 수준에 이르는 경우 과도한 강제에 해당할 가능성이 있으며⁵⁹⁾, 주 법률에 대한 선점 주장 역시 연방법의 범위와 해석에 따라 그 효력이 제한될 수 있다. 따라서, EO 14365를 중심으로 한 이러한 선점 전략이 연방-주 관계의 재구성으로 이어질 수 있는지 여부는 향후 사법적 판단과 입법 과정에서 중요한 쟁점으로 전개될 것으로 보인다.

56) EO 14365 Section 5. 상무부 장관이 명령 발효 후 90일 이내에 광대역 형평성 접근 배치(BEAD) 프로그램의 '비배치(non-deployment) 자금' 지원 조건에 관한 정책 공지를 발행하도록 지시하였다. 해당 공지는 상무부가 '과중한 AI 법률'을 보유한 것으로 식별한 주가 이 자금에 대해 자격이 없음을 명시해야 한다.

57) EO 14365 Section 8. AI 암호화폐 특별보좌관과 대통령실 과학기술보좌관이 공동으로 주 AI 법률을 선점하는 연방 통일 정책 프레임워크 수립을 위한 입법 권고안을 준비하도록 지시하였다. 다만, 아동 안전 보호, AI 컴퓨팅 및 데이터센터 인프라, 주정부의 AI 조달 활용에 관한 주법은 선점 대상에서 제외하도록 하였다. 이는 EO 14365가 행정명령 차원의 즉각적 조치와 함께 의회 입법을 통한 항구적 연방 통일 기준 마련을 병행 추진하고 있음을 보여준다.

58) EO 14365 Section 7. FTC 의장이 AI 모델에 대한 불공정·기만적 행위 금지(15 U.S.C. § 45) 적용 방식에 관한 정책 성명을 발행하도록 지시하였다. 이 성명은 AI 모델의 진실된 출력을 변경하도록 요구하는 주법이 FTC법상 기만 행위 금지 조항에 의해 선점(preempt)되는 상황을 설명해야 한다.

59) NFIB v. Sebelius, 567 U.S. 519 (2012). 연방대법원은 이 사건에서 Medicaid 확대와 관련하여, 새로운 조건을 수용하지 않는 주에 대해 기존 Medicaid 재원까지 박탈하도록 한 구조가 주에 대한 “gun to the head”에 해당하는 강제(coercion)라고 판단하였다. 이는 연방 보조금에 부가된 조건이 단순한 재정적 유인을 넘어, 주의 실질적 선택 가능성을 박탈하는 경우 헌법상 한계를 초과할 수 있음을 보여준다.
<https://supreme.justia.com/cases/federal/us/567/519/> (최종판문일: 2026. 4. 14.)

4. 소결: 안정적 핵심과 유동적 외피의 이중 구조

기관 역할 분산, 정책 도구의 결합, 연방-주 조정이라는 세 층위는 상호 보완적으로 결합하여 미국 AI 거버넌스의 실질적 작동 구조를 구성한다. 이 구조의 핵심적 특징은 안정과 유동의 공존이다. NIST를 중심으로 한 연성법 기반의 기술 표준은 정치적 전환에도 불구하고 기술 규범의 연속성을 유지하는 안정적 핵심으로 기능하는 반면, FTC의 집행 기조, 조건부 재정 지원, 소송 전략은 행정명령의 정책 방향에 따라 실질적으로 재편되는 유동적 외피를 형성한다. 행정명령 하나가 발령될 때마다 수십 개 연방기관이 각자의 법적 권한 범위 안에서 움직이고, 재정·표준·수출·소송이라는 복합적 수단이 병렬적으로 동원된다. 이것이 미국 AI 거버넌스가 단일 입법 없이도 실질적 규범 질서를 형성할 수 있는 이유이며, 동시에 정권 교체에 따른 구조적 불안정성을 내포하는 이유이기도 하다.

〈미국 인공지능 행정명령 현황 및 정책 방향〉

행정명령	발령일 ⁶⁰⁾	제목(약칭)	정책방향	주요내용
EO 13859	2019. 2. 11.	미국 AI 리더십 유지	혁신 촉진	American AI Initiative 출범, NIST 표준화 역할 부여
EO 13960	2020. 12. 3.	신뢰할 수 있는 AI의 연방정부 활용 촉진	혁신 촉진	9대 원칙 제시, AI 활용 사례 인벤토리 구축
EO 14110	2023. 10. 30.	안전하고 신뢰할 수 있는 AI 개발 및 이용	규제 강화	듀얼유즈 모델 안전 테스트 보고 의무, 형평성·시민권 보호
EO 14141	2025. 1. 14.	미국 AI 인프라 리더십 강화	혁신 복귀	데이터센터 전력망 구축 촉진(EO 14318에 의해 폐지)
EO 14148	2025. 1. 20.	유해 행정명령 초기 폐지	혁신 복귀	EO 14110 폐지
EO 14179	2025. 1. 23.	미국 AI 리더십 저해요인 제거	혁신 복귀	규제 장벽 완화, AI 행동계획 수립 지시
EO 14277	2025. 4. 23.	미국 청소년 AI 교육 증진	혁신 복귀	K-12 AI 리터러시 확대
EO 14278	2025. 4. 23.	미래 고임금 숙련직 일자리를 위한 근로자 준비	혁신 복귀	직업훈련·도제제도 강화
EO 14318	2025. 7. 23.	데이터센터 인프라 연방 인허가 절차 가속화	혁신 복귀	인허가 간소화, EO 14141 폐지
EO 14319	2025. 7. 23.	연방정부 내 편향 AI 방지	혁신 복귀	이념적 중립성, 진실추구 원칙, 연방 조달 기준 설정
EO 14320	2025. 7. 23.	미국 AI 기술 스택 수출 촉진	혁신 복귀	full-stack AI 패키지 해외 수출 지원
EO 14363	2025. 11. 24.	제네시스 미션 출범	혁신 복귀	연방 과학 데이터와 고성능 컴퓨팅 통합, AI 기반 R&D
EO 14365	2025. 12. 11.	AI에 관한 국가 정책 프레임워크 보장	혁신 복귀	주(州) AI 규제 통제, AI 소송 태스크포스 설치

60) 발령일은 대통령 서명일 기준임. 연방관보(Federal Register) 게재일과 상이할 수 있음.

IV. 비교법적 분석: 두 모델의 구조적 대비와 한국 인공지능기본법의 설계 과제

1. 두 모델의 구조적 대비: 적응성과 안정성

2026년 1월 22일 시행에 들어간 한국의 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법」(이하 “인공지능기본법”)은 포괄적 단일법 체계를 취하고 있다는 점에서 미국의 분산형 행정명령 모델과 구조적으로 대조된다. 이 대조는 단순한 규범 형식의 차이가 아니라, 각 모델이 추구하는 거버넌스 가치의 차이를 반영하는 것으로 볼 수 있다.

미국 모델의 강점은 적응성(adaptability)에 있다. 행정명령은 의회 입법 없이 즉각 발령될 수 있고, 기술 환경이나 정치적 기조가 바뀌면 단기간에 전면 재편될 수 있다(EO 14110 → EO 14179). 이러한 유연성은 기술 변화의 속도에 대응하는 데 유리할 수 있으나, 그 반대급부로 예측가능성이 약화될 수 있다. 앞서 살펴본 FTC의 Rytr 동의명령 취소가 보여주듯, 기업과 시민 입장에서는 오늘의 규범이 내일도 유효할지 확신하기 어려운 상황이 발생할 수 있다.

한국 모델의 강점은 안정성(stability)에 있다. 법률로 제정된 인공지능기본법은 행정명령에 비해 개정·폐지의 절차적 요건이 엄격하여 상대적으로 지속성이 높으며, 3년 주기 기본계획(제6조)은 정책의 연속성을 제도적으로 뒷받침한다. 그러나, 이 안정성은 적응성의 상대적 결여라는 문제를 수반할 수 있다. 특히 인공지능과 같이 기술 변화가 급격히 이루어지는 영역의 경우, 법률 개정의 경직성으로 인해 규율 공백이 발생할 수 있다는 문제가 있다.

이 두 모델은 상호 배타적이라기보다는 서로의 약점을 비추는 거울로 이해하는 것이 적절하다. 미국 모델의 유연성—기술 변화와 정책 기조의 전환에 신속하게 대응할 수 있는 능력—은 한국 단일법 체계가 상대적으로 결여한 가치를 드러내며, 반대로 미국의 반복적 정책 전환이 초래하는 예측가능성의 약화는 한국 모델의 제도적 안정성이 갖는 의미를 역설적으로 보여준다. 문제는 두 가치—적응성과 안정성—를 어떻게 균형 있게 확보할 것인가이다.

2. NIST AI RMF의 시사점: 기술적 연속성과 제도적 유연성의 확보

앞서 Ⅲ장에서 살펴본 바와 같이, 미국 AI 거버넌스에서 NIST AI RMF는 행정명령의 반복적 전환 속에서도 기술 규범의 연속성을 유지하는 안정적 핵심으로 기능한다. 이러한 기능은 단일 규범 체계를 취하는 한국의 맥락에서는 다른 의미를 갖는다. 법률과 시행령 중심의 구조가 갖는 상대적 경직성을 고려할 때, RMF와 같은 연성 표준은 규범의 연속성을 유지하는 것을 넘어, 법률 체계를 보완하면서 기술 변화에 유연하게 대응할 수 있는 수단으로 작동할 수 있기 때문이다.

(1) 민간 참여를 통한 기술적 연속성 확보

NIST AI RMF가 정치적 전환으로부터 독립성을 유지할 수 있는 핵심 이유는, 법적 강제력이 아니라 산업계·학계·시민사회가 참여하는 자발적 형성 과정과 그에 대한 민간의 수용에 기반하기 때문이다. 이러한 참여 구조 속에서 형성된 기술 표준은 특정 정책 기조에 종속되지 않고, 실무 영역에서 지속적으로 활용되는 경향을 보인다. 이 점에서 NIST AI RMF는 행정명령의 변화와 무관하게 기술 규범의 연속성을 유지하는 기반으로 기능한다.

한국 인공지능기본법 제14조⁶¹⁾는 정부가 표준을 제정·개정·폐지하고 이를 관련 사업자에게 권고하는 구조를 취하면서도(제1항·제2항), 민간 표준화 사업에 대한 지원(제3항)과 국제 표준기구와의 협력(제4항)에 관한 근거를 함께 두고 있다. 이 구조를 적극적으로 활용한다면 NIST AI RMF의 강점을 상당 부분 구현할 수 있다. 결국 한국에 필요한 것은 현행 인공지능기본법 제14조의 틀 안에서 민간 참여 구조를 제도적으로 조직하고 그 결과물에 연속성을 부여하는 운용의 전환이다. 표준의 형성 과정 자체에 민간이 실질적으로 관여하는 구조가 마련될 때, 한국의 AI 기술 표준 역시 정치적 전환에도 일정한 연속성을 확보할 수 있을 것이다.

(2) 단일 규범 체계 하에서의 유연성 확보

한국의 인공지능기본법은 법률과 시행령으로 이어지는 단일 규범 체계를 취하고 있다. 이러한 구조는 예측가능성과 안정성을 제도적으로 보장하는 장점을 가지지만, 인공지능과 같이 기술 변화의 속도가 빠른 영역에서는 규범이 현실을 충분히 포섭하지 못하는 한계를 내포한다. 따라서, 문제는 단일 규범 체계 내부에서 기술 변화에 대응할 수 있는 유연성을 어떻게 확보할 것인가에 있다.

이 맥락에서 NIST AI RMF의 운용 방식은 중요한 시사점을 제공한다. NIST AI RMF는 법적 구속력을 갖는 규범이 아님에도 불구하고, 산업계·학계·시민사회가 참여하는 지속적 갱신 과정을 통해 변화하는 기술 환경을 반영한다. 즉, 법령 개정 절차와 별도로 내용이 지속적으로 갱신될 수 있는 구조를 형성하고 있다는 점에서, 기술 변화에 대한 대응력을 확보하고 있다.

한국의 경우 이러한 기능은 법률이나 시행령과 분리된 외부 규범이 아니라, 그 체계 내부에서 생성·운

61) 인공지능기본법

제14조(인공지능기술의 표준화) ① 정부는 인공지능기술, 학습용데이터, 인공지능의 안전성·신뢰성 등과 관련된 표준화를 위하여 다음 각 호의 사업을 추진할 수 있다.

1. 인공지능기술 관련 표준의 제정·개정 및 폐지와 그 보급
2. 인공지능기술 관련 국내외 표준의 조사·연구개발
3. 그 밖에 인공지능기술 관련 표준화 사업
- ② 정부는 제1항제1호에 따라 제정된 표준을 고시하여 관련 사업자에게 그 준수를 권고할 수 있다.
- ③ 정부는 민간 부문에서 추진하는 인공지능기술 관련 표준화 사업에 필요한 지원을 할 수 있다.
- ④ 정부는 인공지능기술 표준과 관련된 국제표준기구 또는 국제표준기관과 협력체계를 유지·강화하여야 한다.
- ⑤ 그 밖에 제1항 및 제3항에 따른 표준화 사업의 추진 및 지원 등과 관련하여 필요한 사항은 대통령령으로 정한다.

용되는 비구속적 규범을 통해 구현될 수 있다. 인공지능기본법은 인공지능 안전성·신뢰성 확보를 위하여 인공지능의 개발에 관한 가이드라인의 보급을 명시적으로 예정하고(제30조제1항제1호)⁶²⁾, 고영향 인공지능의 기준과 예시 등에 관하여 가이드라인을 수립·보급할 수 있도록 규정하고 있다(제33조제3항)⁶³⁾. 이러한 규정은 기술 변화에 따라 유동적으로 조정되어야 하는 사항을 법률 수준에서 직접 고정하기보다는, 가이드라인이라는 형식을 통해 지속적으로 보완하도록 설계되어 있음을 보여준다.

이러한 법령 구조를 전제로 할 때, AI 위험관리 가이드라인은 법령 외부에서 작동하는 규범이 아니라, 인공지능기본법 체계 내부에서 생성·운용되는 실무적 지침으로 이해하는 것이 타당하다. 인공지능기본법의 틀 안에서 이러한 가이드라인을 적극 활용하고, 기술 변화에 민감한 사항—고영향 인공지능의 범위, 위험 평가 기준, 투명성 의무의 구체적 내용 등—을 지속적으로 갱신하는 체계를 구축하는 것이 그 현실적인 경로이다. 이러한 가이드라인은 법적 구속력을 갖지 않으면서도, 법률에 근거하여 형성되고 산업계의 자발적 수용을 통해 사실상의 규범으로 기능한다는 점에서, 단일 규범 체계 내부에서 유연성을 확보하는 수단으로 작동한다.

다만, 이러한 방식이 실효성을 갖기 위해서는 가이드라인의 위상을 명확히 하는 것이 중요하다. 가이드라인은 법률·시행령의 해석과 적용을 보조하는 범위 내에서 운용되어야 하며, 법적 의무와의 경계를 명확히 유지하는 방식으로 설계될 필요가 있다. 이와 같은 조건이 충족될 때, 한국의 단일 규범 체계도 안정성을 유지하면서 기술 변화에 대한 실질적 대응력을 확보할 수 있을 것이다.

3. AI 인프라의 국가 전략적 배분: EO 14318·14363의 교훈

미국의 EO 14318(데이터센터 허가 가속화)과 EO 14363(제네시스 미션)은 AI 인프라—데이터센터, 전력망, 연방 과학 데이터셋—를 단순한 민간 투자 영역이 아니라 국가 전략 자원으로 재정의한다. 특히 EO 14318이 대규모 데이터센터 프로젝트에 대한 연방 허가를 신속히 처리하도록 하고 연방 부지 활용을 명시한 것은, AI 인프라 구축이 에너지·토지 정책과 불가분하게 연결된 국가 차원의 과제임을 보여준다.

이러한 인프라 중심의 경쟁 구도는 최근 국제 연구에서도 확인된다. 스탠퍼드대학교 인간중심 인공지

능 연구소(HAI)가 발표한 「2026 AI Index Report」는 주요 국가 간 모델 성능 격차가 빠르게 축소되는 상황에서, 인공지능 경쟁의 핵심이 전력과 냉각 인프라 확보로 이동하고 있음을 지적한다. 보고서에 따르면 글로벌 AI 데이터센터의 전력 수요는 급격히 증가하여 미국 뉴욕주 전체의 피크 전력 수요와 맞먹는 수준에 이르고 있으며, 대규모 모델 운영 과정에서 상당한 수준의 에너지와 수자원이 소모되는 것으로 나타난다⁶⁴⁾. 이러한 분석은 AI 경쟁력이 더 이상 알고리즘이나 데이터만으로 결정되지 않고, 이를 뒷받침하는 물리적 인프라에 의해 구조적으로 제약될 수 있음을 시사한다.

한국에서도 이러한 문제는 이미 현실화되고 있다. 한국은 인공지능 특허와 반도체 기술 등에서 높은 경쟁력을 보유하고 있음에도 불구하고, 데이터센터 전력 수요를 안정적으로 공급할 수 있는 전력망과 적절한 입지 확보 측면에서 제약을 받고 있다. 특히 전력 공급 구조와 부지 규제, 부처 간 정책 조정의 어려움이 결합되면서, 인공지능 인프라 구축이 개별 산업 정책의 문제가 아니라 국가적 조정이 필요한 영역으로 전환되고 있다. 그러나, 현행 인공지능기본법은 AI 인프라의 지리적 배분이나 전력망 연계에 관한 명시적 규정을 두고 있지 않다.

이 공백을 채우기 위해서는 인공지능기본법 제6조의 기본계획 법정 사항에 ‘AI 인프라의 입지·에너지 연계에 관한 관계 부처 공동 계획 사항’을 명시적으로 포함시키는 방향을 검토할 필요가 있다.⁶⁵⁾ 구체적으로는 데이터센터 입지 선정 시 전력 수급 안정성, 냉각수 가용성, 기존 산업단지와의 연계 가능성을 종합적으로 고려하는 평가 체계를 기본계획에 담고, 이를 관계 부처가 공동으로 수립하도록 하는 부처 간 협력 의무를 명시하는 것이다. 미국에서 EO 14318이 에너지부, 국방부, 환경보호청 등 여러 부처의 협력을 동시에 요구한 것과 같이, AI 인프라는 단일 부처가 아닌 범부처 조정의 대상으로 명시적으로 제도화되어야 한다.

64) Stanford Institute for Human-Centered Artificial Intelligence, AI Index Report 2026, Stanford University, 2026, at 30 (Figure 1.2.4), <https://aiindex.stanford.edu/report/> (최종방문일: 2026. 4. 14.)

65) 인공지능기본법

제6조(인공지능 기본계획의 수립) ① 과학기술정보통신부장관은 관계 중앙행정기관의 장 및 지방자치단체의 장의 의견을 들어 3년마다 인공지능기술 및 인공지능산업의 진흥과 국가경쟁력 강화를 위하여 인공지능 기본계획(이하 “기본계획”이라 한다)을 제7조에 따른 국가인공지능전략위원회의 심의·의결을 거쳐 수립·변경 및 시행하여야 한다. 다만, 기본계획 중 대통령령으로 정하는 경미한 사항을 변경하는 경우에는 그러하지 아니하다.

② 기본계획에는 다음 각 호의 사항이 포함되어야 한다.

1. 인공지능등에 관한 정책의 기본 방향과 전략에 관한 사항
2. 인공지능산업의 체계적 육성을 위한 전문인력의 양성 및 인공지능 개발·활용 촉진 기반 조성 등에 관한 사항
3. 인공지능윤리의 확산 등 건전한 인공지능사회 구현을 위한 법·제도 및 문화에 관한 사항
4. 인공지능기술 개발 및 인공지능산업 진흥을 위한 재원의 확보와 투자의 방향 등에 관한 사항
- 4의2. 「공공데이터의 제공 및 이용 활성화에 관한 법률」에 따른 공공데이터를 이용한 학습용데이터 생성, 공공데이터 제공 등의 범위와 기준 및 그 활성화에 관한 사항
5. 인공지능의 공정성·투명성·책임성·안전성·접근성 확보 등 신뢰 기반 조성에 관한 사항
6. 인공지능기술의 발전 방향 및 그에 따른 교육·노동·경제·문화 등 사회 각 영역의 변화와 대응에 관한 사항
- 6의2. 인공지능기술의 이해 및 활용을 위한 교육의 지원 및 홍보에 관한 사항
7. 그 밖에 인공지능기술 및 인공지능산업의 진흥과 국제협력 등 국가경쟁력 강화를 위하여 과학기술정보통신부장관이 필요하다고 인정하는 사항(이하 생략)

62) 인공지능기본법

제30조(인공지능 안전성·신뢰성 감·인증등 지원) ① 과학기술정보통신부장관은 단체등이 인공지능의 안전성·신뢰성 확보를 위하여 자율적으로 추진하는 검증·인증 활동(이하 “검·인증등”이라 한다)을 지원하기 위하여 다음 각 호의 사업을 추진할 수 있다.

1. 인공지능의 개발에 관한 가이드라인 보급(이하 생략)

63) 인공지능기본법

제33조(고영향 인공지능의 확인) ① 인공지능사업자는 인공지능 또는 이를 이용한 제품·서비스를 제공하는 경우 그 인공지능이 고영향 인공지능에 해당하는지에 대하여 사전에 검토하여야 하며, 필요한 경우 과학기술정보통신부장관에게 고영향 인공지능에 해당하는지 여부의 확인을 요청할 수 있다.

- ② 과학기술정보통신부장관은 제1항에 따른 요청이 있는 경우 고영향 인공지능 해당 여부를 확인하여야 하며, 필요한 경우 전문위원회를 설치하여 관련 자문을 받을 수 있다.

- ③ 과학기술정보통신부장관은 고영향 인공지능의 기준과 예시 등에 관한 가이드라인을 수립하여 보급할 수 있다.

- ④ 그 밖에 제1항에 따른 확인 절차 등에 관하여 필요한 사항은 대통령령으로 정한다.

4. 국제 표준화 전략: 규범 수용에서 규범 형성으로

인공지능기본법 제14조제4항은 정부가 국제표준기구와 협력체계를 유지·강화하도록 규정하고 있다⁶⁶⁾. 다만, 현재 한국의 국제 AI 표준화 참여는 국내에서 형성된 규범을 국제적으로 확산하기보다는, 기존 국제표준에 대한 정합성을 확보하는 데 상대적으로 초점이 맞추어진 것으로 보인다.

미국의 경우, NIST가 개발한 AI 위험관리 프레임워크 등 기술 표준이 산업과 정책 영역에서 널리 활용되면서, 이러한 기준이 국제표준화 논의에도 일정한 영향을 미치는 구조가 형성되고 있다. 또한, 기술·인프라·서비스를 포괄하는 방식의 대외 확산 과정에서, 해당 기술에 내재된 관리 방식과 기준이 함께 확산되는 경향이 나타난다. 이는 국제표준이 반드시 조약이나 공식 규범의 형태로만 형성되는 것이 아니라, 기술의 보급과 산업적 채택을 통해 사실상 표준으로 자리 잡을 수 있음을 보여준다.

한국이 AI 강국으로 도약하기 위해서는 규범 수용자에서 규범 형성자로의 전환이 필요하다. 이를 위해서는 두 가지 조건이 선행되어야 한다. 첫째, 앞서 제안한 한국형 AI 위험관리 프레임워크가 국내에서 실질적으로 작동하고, 국제적으로 공유 가능한 형태로 정립되어야 한다. 자국 내에서 충분한 운용 경험이 축적되지 않은 규범은 국제 표준으로 제안되기 어렵기 때문이다. 둘째, 인공지능기본법의 시행 경험—특히 고영향 인공지능의 기준 설정과 검·인증 체계의 구축 및 지원 등—을 국제 표준화 논의에 적극적으로 연결하는 정책적·외교적 노력이 필요하다. EU에 이어 포괄적 단일 AI 기본법을 시행한 국가로서의 경험은, 국제 논의에서 단순한 규범 수용을 넘어 규범 형성 과정에 참여할 수 있는 중요한 기반이 될 수 있다.

V. 결론

이 글에서는 EO 13859(2019)부터 EO 14365(2025)에 이르는 미국의 인공지능 행정명령을 분석하여, ‘What(정책 방향)’과 ‘How(실행 구조)’라는 두 물음에 답하고, 이를 한국 인공지능기본법의 설계 과제와 비교법적으로 연결하고자 하였다.

첫 번째 물음에 대하여는, 미국 인공지능 정책이 혁신 촉진(EO 13859·13960) → 안전·형평성 가치 중심의 규제 강화(EO 14110) → 혁신 우선 복귀 및 이념 편향 배제(EO 14179 이후)라는 진자 운동적 조정 구조를 나타냄을 살펴보았다. 이 구조는 기술 변화에 대한 신속한 대응을 가능하게 하는 반면, 정권 교체에 따라 정책 방향이 급격히 전환될 수 있다는 구조적 불안정성을 내포한다. 나아가, 이 진자 운동은 단순한 반복이 아니어서, 전환이 거듭될수록 새로운 정책 수단과 법적 갈등이 누적되면서 미국 AI 거버넌스의 전체 구조는 점점 더 복잡하고 다층적인 양상을 띠게 된다. EO 14319는 이념적 편향을 내재한 AI가 국가 행정

의 언어 자체를 특정 방향으로 구조화할 수 있다는 문제의식에서 출발하여, 이념적 요건에 따른 규제 부담을 견어냄으로써 혁신의 공간을 확보하는 동시에, 연방 조달을 매개로 AI 시스템의 중립성과 객관성을 제도적으로 담보하려는 새로운 국면을 열었으며, EO 14365를 계기로 제기된 xAI v. Colorado 소송(2026)은 행정명령 중심의 정책 전환이 연방-주 갈등을 법적 분쟁의 형태로 현재화하는 양상을 보여준다.

두 번째 물음에 대하여는, 미국 AI 거버넌스가 ‘연성법 중심의 안정적 핵심’과 ‘행정명령 중심의 유동적 외피’의 이중 구조로 작동함을 분석하였다. NIST AI RMF는 정부·기업·시민사회가 공동으로 개발한 사실상의 표준으로서 행정부 교체와 무관하게 기술 규범의 연속성을 유지하는 안정적 핵심을 형성하며, 그 위에 규제 완화, 재정 지원, 기술 수출, 인력 양성을 통합하는 복합적 정책 도구들이 유동적 외피로 결합된다. FTC의 Rytr 동의명령 취소는 집행기관의 법적 판단이 행정명령의 정책 기조에 따라 실질적으로 재편될 수 있음을 보여주는 전형적 사례이다. 이 이중 구조는 단일 입법 없이도 실질적 규범 질서를 형성할 수 있는 미국식 거버넌스의 적응력을 설명하는 동시에, 정권 교체에 따른 구조적 불안정성이라는 내재적 취약성을 함께 내포한다.

비교법적 분석에서는 미국 모델(적응성 우선)과 한국 모델(안정성 우선)의 구조적 대비를 살펴보고, 네 가지 설계 과제를 도출하였다. 첫째, 한국은 정치적 전환으로부터 독립적인 AI 위험관리 프레임워크를 정부·산업·시민사회의 공동 참여 방식으로 구축해야 한다. 인공지능기본법 제14조의 틀 안에서 민간의 자발적 참여 공간을 최대한 확보하고 그 결과물에 제도적 연속성을 부여하는 운용의 전환이 그 현실적인 경로이다. 둘째, 법률과 시행령의 경직성을 보완하는 수단으로서 고영향 인공지능 가이드라인을 적극 활용하여 기술 변화에 유연하게 대응하는 체계를 구축해야 한다. 가이드라인은 법적 구속력을 갖지 않음면서도 법률에 근거하여 형성되고 법령 개정 절차와 별도로 지속적으로 갱신될 수 있다는 점에서, 단일 규범 체계 내부에서 유연성을 확보하는 현실적 수단이다. 셋째, AI 데이터센터·전력망 등 인프라의 입지 및 에너지 연계에 관한 사항을 기본계획의 법정 사항으로 명시하고 관계 부처의 공동 수립 의무를 부과함으로써, AI 인프라를 단일 부처 소관이 아닌 범부처 조정의 대상으로 명시적으로 제도화해야 한다. 넷째, 한국은 국제 AI 표준화 과정에서 규범 수용자에서 규범 형성자로의 전환을 모색해야 한다. 이를 위해서는 한국형 AI 위험관리 프레임워크의 국내 운용 경험을 축적하고, 인공지능기본법 시행 경험—특히 고영향 인공지능 기준 설정과 검·인증 체계—을 국제 표준화 논의에 적극적으로 연결하는 정책적·외교적 노력이 병행되어야 한다.

미국 AI 행정명령 거버넌스는 연성법 중심의 안정적 핵심과 행정명령 중심의 유동적 외피의 이중 구조로 작동하며, 이 구조가 보여주는 유연성은 한국 인공지능기본법이 보완해야 할 지점을, 불안정성은 경계해야 할 지점을 각각 드러낸다. 한국이 단일법의 안정성을 유지하면서도 기술 변화에 실질적으로 대응하려면, 그 안에 유연성의 공간을 의도적으로 설계하는 전략—‘구조 내 유연성(flexibility within structure)’—이 필요하다. 이것이 미국 행정명령 거버넌스의 경험이 한국에 주는 가장 중요한 시사점이다.

66) 인공지능기본법

제14조(인공지능기술의 표준화) ④ 정부는 인공지능기술 표준과 관련된 국제표준기구 또는 국제표준기관과 협력체계를 유지·강화하여야 한다.