

ISSN 2734-0317

www.kita.net

KITA

통상 리포트

2024

VOL.07

인공지능(AI) 규제 주도권 확보를 위한 글로벌 경쟁 및 시사점

정해영 수석연구원

KITA 통상리포트 2024 VOL.07

인공지능(AI) 규제 주도권 확보를 위한 글로벌 경쟁 및 시사점

발행인 윤진식

편집인 조상현

발행처 한국무역협회 통상연구실

발행일 2024년 7월 26일

인공지능(AI) 규제 주도권 확보를 위한 글로벌 경쟁 및 시사점

Contents

요약

01. 인공지능(AI)의 긍정적 효과와 리스크	1
(1) AI 기술의 발전과 비즈니스 활용 확대	
(2) AI 기술 적용 확산에 따라 대두되는 리스크	
02. 주요국의 인공지능(AI) 규제 및 영향력 확대 전략	9
(1) 유럽연합(EU) : 위험 기반 포괄적 규제	
(2) 미국 : 가이드라인, 모범사례 등 기업의 자발적 준수에 의존	
(3) 중국 : 신속하고 반복적인 규제	
03. 인공지능(AI) 규제 협력 및 시사점	25

정해영 수석연구원 (02-6000-5601, hy.chung@kita.or.kr)

※ 본 리포트 인용 시,
반드시 출처를 명기하여 주시기 바랍니다.



요약

인공지능(AI)은 학습·추론·지각 등 인간의 지적 능력을 인공적으로 컴퓨터나 기계를 통해 구현하는 기술로, 글로벌 AI 시장 규모는 '23년 1,359억 달러에서 '30년 8,267억 달러로 연평균 29.4% 성장할 것으로 전망된다. AI 기술은 환경·건강·공공·금융·모빌리티·농업 등 다양한 분야에 긍정적인 경제적·사회적 효과를 가져다줄 것으로 기대되지만, AI의 긍정적 효과 이면에는 데이터 활용에 따른 위험 요소도 내포되어 있다. 정부는 혁신을 저해하지 않으면서 AI를 규제해야 하는 압박에 직면해왔다. AI 규제의 딜레마는 국가 통제와 혁신 사이의 균형으로, 지나친 규제는 AI의 창의적·발전적 잠재력을 저해할 수 있지만, 불충분한 규제는 사회 복지와 윤리 기준, 개인의 권리에 해로운 결과를 초래할 수 있다. 주요국과 기업의 기술 리더십 경쟁이 심화되며 AI 규제가 추진력을 얻고 있으나, 국가별로 여전히 AI 규제와 개발에 대해 상이한 견해와 목표가 대립하고 있다. 유럽연합(EU)과 미국, 중국은 AI에 대한 서로 다른 비전으로 기술의 미래에 대한 통제권을 놓고 경쟁하며 디지털 세계에서 영향력을 확대하고 있다.

데이터·개인정보 보호 등 시민의 권리에 초점을 맞춘 디지털 기술 규제를 선도적으로 도입해 온 유럽연합(EU)은 세계 최초의 포괄적 AI 규제인 AI법안이 EU 이사회의 최종 승인을 받으며 '26년 발효할 예정이다. AI법은 AI 시스템과 모델을 EU 시장에 출시·서비스하거나 AI 시스템에서 생성된 결과물을 EU에서 사용하려는 제3국 소재 기업에게도 적용된다. AI법은 AI 시스템의 위험 정도를 '허용될 수 없는 위험', 고위험, 제한적 위험, 최소한의 위험으로 분류하여 차등 규제하는 것이 핵심으로, 사람의 안전·생계·기본권에 허용될 수 없는 위험을 초래하는 AI 시스템은 엄격히 금지하고 있으며, 부정적인 영향을 미칠 수 있는 고위험 AI 시스템에 대해서는 엄격히 규제하고 있다. 챗봇, 딥페이크 등 콘텐츠를 생성·조작 등 제한적 위험을 초래하는 AI 시스템에는 투명성 의무가 적용되며, 그 외 위험이 낮은 AI 시스템은 현행 법률 이상의 추가 의무가 적용되지 않는다. AI법은 또한 범용 AI(GPAI) 모델과 시스템적 위험을 초래할 수 있는 범용 AI 모델을 구분하여 별도 규제하고 있다. EU는 개인정보 보호 규정인 GDPR과 마찬가지로 AI법을 통해 글로벌 표준을 형성하고 영향력을 확대하고자 한다.

다수의 빅테크 기업을 보유한 미국은 AI 규제에 대한 정부의 역할이 제한적이며, 가이드라인과 모범사례 등 기업의 자발적 준수를 중시해왔다. 미국은 언론의 자유, 자유로운 인터넷, 혁신에 대한 인센티브 보호를 중시해 왔으며, 미-중 기술패권 경쟁 격화에 따라 AI를 규제의 대상보다는 경제·지정학적 우위의 수단으로 간주해왔다. 빅테크 규제 부재에 대한 사회적 대가가 점차 부각되고 있으나, 의회의 포괄적인 AI 규제가 정치적 교착상태로 난항을 겪자, 바이든 행정부는 '23.7월과 8월 주요 AI 기업들로부터 안전하고 투명한 AI 기술 개발을 위한 자발적 조치를 약속받는데 이어, 10월 AI를 규제하는 첫 행정명령을 발표하며 국가 안보·건강·안전을 위협하는 AI 기술 개발과 이용에 대한 규제 의지를 표명했다. AI 행정명령은 AI 기술 개발 기업에게 일정 수준의 책임을 부여하고, 이중용도 프로젝트에 대한 정부 감독을 위한 기반을 마련했으나, 엄격한 준수요건과 법적 구속력을 갖춘 EU 방식과 달리, 규제보다 산업정책과 지침에 가깝다.

미국식 규제 모델은 기술 발전과 경제 성장을 촉진할 AI의 잠재력과 빠르게 진화하는 기술을 규제해야 하는 어려움 때문에 영국, 호주 등 일부 국가가 선호하고 있으나, 미-중 경쟁과 EU의 '브뤼셀 효과'를 의식한 미국이 EU와의 협력을 심화할 가능성도 있다.

중국은 AI 개발을 위한 포괄적인 계획의 일환으로 가장 먼저 AI를 규제한 국가로, 국가의 통제를 유지하면서 기술 혁신을 장려하는데 중점을 두고 있다. 중국은 EU처럼 AI를 포괄적으로 규제하기보다 새로운 AI 관련 이슈에 대해 단편적인 규칙을 신속하게 제정하고 지속적으로 수정해왔다. 중국 정부는 알고리즘 추천 서비스, 딥페이크, 생성형 AI에 관한 규범을 각각 '21년, '22년, '23년 통과시키며 새로운 AI 이슈별 규제를 신속히 도입해왔다. 생성형 AI에 대한 세계 최초 규제인 '생성형 AI 서비스 관리 잠정 방법'은 중국 주민을 대상으로 콘텐츠 생성 서비스를 제공하는 국내·외 생성형 AI 서비스 제공자에게 적용되며, 중국에서 개발되어도 국외 사용자만 대상으로 할 경우 적용되지 않는다. 잠정 방법의 특징은 생성형 AI가 만들어 내는 콘텐츠가 사회주의 핵심 가치를 건지하도록 규정하며, 국가권력, 사회주의 체제의 전복, 국가안보와 이익을 위태롭게 하는 내용을 포함하지 못하도록 금지하고 있다는 점이다. 또한 '23년 중국 국무원은 AI 법안의 입법 계획을 발표한 후 '24.3월 초안을 공개했는데, 중국의 AI 산업 발전에 초점을 두고 있다. AI 규제에 대한 중국의 접근법은 개인에 대한 피해를 완화하고 사회 안정과 국가 통제를 유지하는 동시에 AI 분야의 글로벌 리더십을 확보하여 규제 논의에 영향을 미치는데 초점을 두고 있으며, 이를 위해 국제표준 개발 노력을 지속하고, 디지털 실크로드를 바탕으로 AI 기반 감시 기술 및 디지털 인프라를 수출하고 있다.

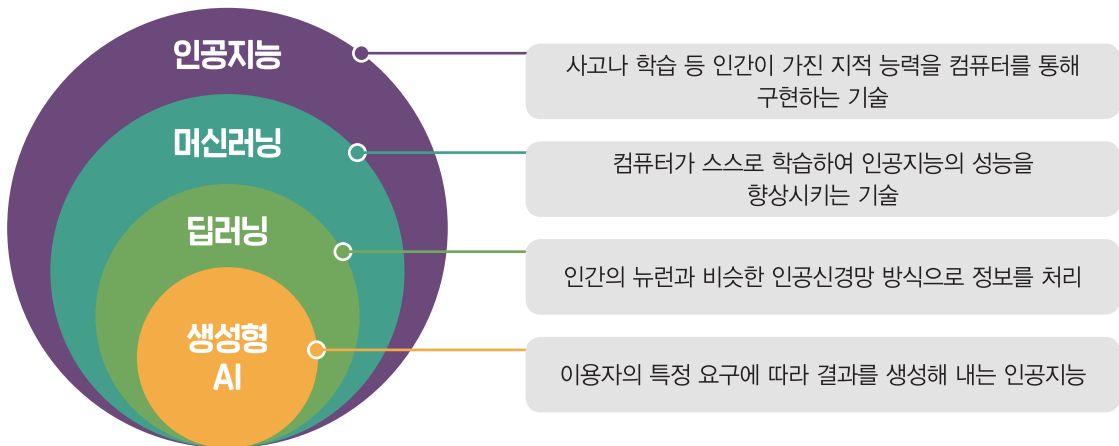
생성형 AI의 등장 이후 주요국들은 앞다퉀 AI 규제에 나섰으나, AI 기술은 국경을 넘어 전 세계적으로 영향을 미치기 때문에 국제 협력이 필요하다. AI 리스크 규제를 위한 국제 협력은 추진력을 얻고 있으나, 각국의 이해관계, 가치관, 역량의 차이로 의미 있는 합의 도달까지 난항이 예상된다. 비즈니스 효율화와 생산성 증대, 혁신을 통한 기업의 생존 차원에서 AI 도입은 더 이상 선택이 아닌 필수이나, AI 위험성을 심각하게 고려하는 전 세계 정부의 관심 증가로 AI 부작용과 사회적 비용 최소화를 위한 입법 노력이 빠르게 확산되고 있다. AI 활용 기업은 데이터 관련 법 위반 리스크 예방을 위한 규제 모니터링 등 선제적 노력이 필요하다.

1 인공지능(AI)의 긍정적 효과와 리스크

(1) AI 기술의 발전과 비즈니스 활용 확대

- 인공지능(AI)은 학습 · 추론 · 지각 등 인간의 지적 능력을 인공적으로 컴퓨터나 기계를 통해 구현하는 기술로, 다양한 형태로 존재하나 크게 머신러닝과 딥러닝으로 구분¹⁾

〈그림1. AI에 대한 분류 및 예시〉



자료: 최창열(2018) 참고하여 저자 정리

- 머신러닝은 가장 기본적이고 단순한 형태의 AI로, 명시적으로 규칙을 프로그래밍하지 않고 데이터로부터 의사결정을 위한 패턴을 기계가 스스로 학습하는 기술로서, 기본적인 채팅 어시스턴트, 앱, 음성 명령 및 기타 제한된 상호작용에서 사용되는 가장 일반적인 형태의 AI²⁾
- 딥러닝은 인공신경망 기반의 모델로 기계가 비정형 데이터로부터 특징 추출 · 판단까지 한 번에 수행할 수 있고, 인간과의 상호작용을 통해 사용자 정보를 통합 · 학습하는 과정에서 더욱 정교해지며, 대표적인 예시로는 OpenAI의 챗GPT와 구글의 제미니(Gemini) 등³⁾
 - '22년 말, 챗GPT가 출시되며 새로운 AI 트렌드로 급부상한 생성형 AI는 프롬프트에 대응하여 텍스트, 이미지, 음악 등을 생성할 수 있는 인공지능으로, 단순히 기존 데이터 분석이 아닌 새로운 콘텐츠 생성에 초점을 맞춘 인공지능

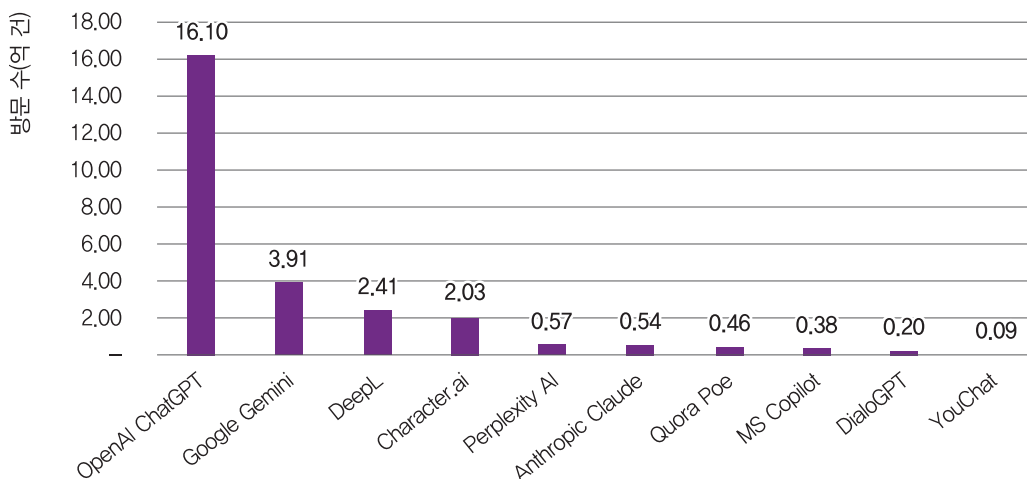
1) Amazon Web Services, "What is Artificial Intelligence (AI)?", AWS.

2) LG CNS AI빅데이터연구소(2020.3.23.), "딥러닝, 데이터로 세상을 파악하다(1)", LG CNS 및 Thormundsson, Bergur(2024.6.19.), "Artificial intelligence (AI) worldwide - statistics & facts", Statista.

3) Ibid.

- 글로벌 생성형 AI 시장 규모는 '23년 448.9억 달러에서 연평균 24.4% 성장하여 '30년 2,070억 달러로 확대될 것으로 예상되며, 생성형 AI는 미디어·엔터테인먼트를 시작으로 자동차·운송, 금융, 헬스케어 분야에서 사용되며 점유율이 빠르게 확산될 것으로 예상⁴⁾
- '24.4월 기준, 가장 많이 사용되는 AI 플랫폼은 오픈AI의 '챗GPT'로 타 플랫폼 대비 압도적으로 많은 16억 1천만 건의 방문수를 기록했으며, 이어서 구글의 '제미나이', 번역 AI인 '딥엘' 순⁵⁾

〈그림2. 2024년 가장 많이 사용되는 AI 플랫폼〉



자료: Ahmed(2024) 참고하여 저자 정리

- 구글, 마이크로소프트, 아마존, X 등 빅테크들은 다양한 스타트업들 통해 가장 발전된 생성형 AI를 개발하기 위해 경쟁하고 있으나⁶⁾, AI 모델을 훈련시키는 데 저작권이 있는 텍스트와 이미지를 사용했다는 비난도 받음⁷⁾
- 현재 대부분의 AI가 상용화 단계에 이르지 못했으나, 생성형 AI는 향후 2년에서 5년 내 실질적 혁신 성과를 나타낼 것으로 예측⁸⁾

□ 글로벌 AI 시장 규모는 '23년 1,359억 달러에서 '30년 8,267억 달러로 연평균 29.4% 성장 전망⁹⁾

- 다양한 데이터가 폭발적으로 생성되며 알고리즘·머신러닝 기술 발전, 자율 AI에 대한 시장 수요 증가 및 비용 절감 등으로 글로벌 AI 시장은 빠르게 성장할 것으로 예상

4) Statista Research Department(2024.2.9.), "Generative artificial intelligence (AI) market size worldwide from 2020 to 2030", Statista.

5) Ahmed, Mashaid(2024.4.9.), "20 Most Used Artificial Intelligence (AI) Platforms of 2024", Insider Monkey.

6) Field, Hayden and Leswing, Kit(2024.3.30.), "Generative AI 'FOMO' is driving tech heavyweights to invest billions of dollars in startups", CNBC.

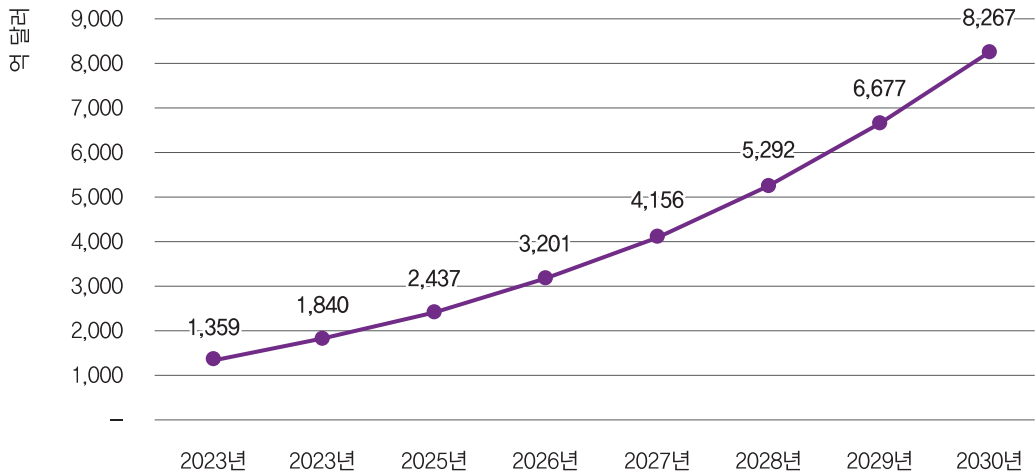
7) Robins-Early, Nick(2024.4.9.), "New bill would force AI companies to reveal use of copyrighted art", The Guardian.

8) Gartner(2023.8.17.), "What's New in Artificial Intelligence from the 2023 Gartner Hype Cycle", Gartner.

9) Thormundsson, Bergur(2024.6.20.), "AI market size worldwide from 2020-2030", Statista.

- － 머신러닝, 자연어 처리, 컴퓨터 비전, 생성형 AI 등 AI가 의료, 금융, 제조, 소매 등 다양한 분야에 빠르게 적용

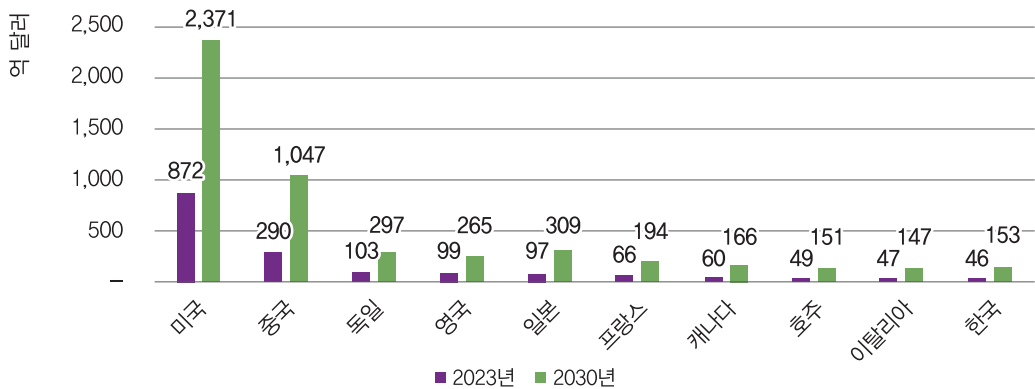
〈그림3. 글로벌 AI 시장 규모 전망〉



자료: Thormundsson(2024)

- 국가별 시장규모로는 미국이 가장 큰 AI 시장 규모를 갖고 있으며, 이어서 중국, 독일, 영국, 일본 순이며 한국은 10위¹⁰⁾
- － 국내 AI 시장은 '23년 46억 달러 규모를 형성하고 있으며, '30년까지 연평균 18.7% 성장하여 153억 달러 규모로 8위에 오를 전망

〈그림4. 국가별 AI 시장 규모 및 전망〉



자료: Statista Market Insights(2024) 참고하여 저자 정리

10) Statista Market Insights(2024.3.), "Artificial Intelligence", Market Insights, Statista.

□ AI 기술은 환경·건강·공공·금융·모빌리티·농업 등 다양한 분야에 긍정적인 경제적·사회적 효과를 가져다줄 것으로 기대

○ AI는 기존 산업과의 융합을 통한 디지털화 촉진 등 신성장 동력이 될 것으로 기대

－ AI의 가치는 데이터 기반의 예측력으로, 전략, 재무, 인사, 상품·서비스 개발, 생산 관리, 마케팅·영업 등 기업 업무 프로세스 전반의 디지털화를 촉진하고 의사결정에 영향

－ AI 기반 스타트업이 활성화되면서 AI는 점차 핵심 성장 기반으로서 응용 산업의 동반 성장과 기술 혁신을 주도할 것으로 기대

○ AI는 다양한 산업 분야의 비즈니스 전 단계에 적용되어 프로세스 단계별 비용 절감, 실행속도 단축, 복잡성 감축, 관계 전환, 혁신 촉진, 신뢰 강화 등 새로운 가치를 창출하는데 큰 역할을 수행¹¹⁾

〈표1. AI를 통한 새기업가치 창출 예시〉

새기업가치	설 명	예 시
비용 절감	AI와 지능형 자동화 솔루션 도입을 통해 상대적으로 저가치·반복적 업무를 자동화하여 효율성·품질 개선 및 비용 절감	자연어 처리를 활용한 데이터 입력 및 진료예약 관리 자동화
실행속도 단축	자연시간을 최소화해 운영·사업성과 도출에 걸리는 시간 단축	예측적 통찰력을 활용한 합성 의약품 임상시험으로 신약 승인절차 가속화
복잡성 감축	복잡한 데이터에서 패턴 파악능력 등이 향상된 애널리틱스 기술을 활용, 이해와 의사결정 개선	기계설비 유지보수 필요를 예측해 공장 다운타임 감축
관계 전환	사람과 기술 간 소통방식 전환을 통해 사람이 기계의 방식에 맞추는 대신 사람의 방식 그대로 소통	고객의 정서를 이해하고 대응할 수 있는 대화형 봇 활용, 고객 요구에 더욱 효과적으로 대응
혁신 촉진	AI를 활용해 혁신적인 신제품, 시장, 비즈니스 모델을 창출함으로써 진입할 시장과 성공전략 재정의	SNS에서 발굴한 고객요구와 취향 등에 기반한 새로운 제품과 컨셉 제언
신뢰 강화	사기 행위와 사이버 공격으로부터 사업체 보호, 일관적 서비스·품질 유지, 투명성 개선을 통한 브랜드 신뢰 강화	사이버 공격 발생 전 종류 식별 및 예측

자료: 딜로이트 AI 연구소(2023)

－ 한국과 미국 기업의 데이터를 분석한 결과에 따르면 AI 기술을 비즈니스에 활용하고 있는 기업에서 도입하지 않은 기업보다 더 높은 매출액 상승이 나타났으며, AI 도입이 다양한 유·무형의 내·외부자원과 결합되었을 때 AI의 긍정적 효과가 발생¹²⁾

11) 딜로이트 AI 연구소(2023), “인공지능(AI) 활용서: 6대 산업별 활용사례”, 2023년 8월 Deloitte Insights, Deloitte.

12) 김태균(2022), “기업 성과를 높이는 인공지능”, 월간 SW중심사회 2022년 5월호, 소프트웨어정책연구소.

〈표2. 미국과 한국 기업의 AI 기술 도입 성과〉

미국 (348개 상장기업)	- 자동화 AI 도입 기업에서 단기적인 매출원가 절감 효과 - 중강 AI 도입 기업에서는 단기적 이익 증가나 비용 감소 효과는 없었으나 미래에 대한 기대로 더 높은 시장가치를 받음
한국 (300개 중소기업)	- AI 기술 도입 기업에서 매출액 상승 발생 - AI 기술 도입을 통해 제품·서비스 개발, 마케팅, 고객 응대의 효율성 제고 - AI 기술 도입 기업 중 클라우드 컴퓨팅과 데이터베이스 센터를 같이 활용한 기업에서 더 높은 매출액 상승 발생 - AI 기술 도입 기업 중 내부 연구개발을 진행해 온 기업에서 더 높은 매출액 상승 발생

자료: 김태균(2022)

- 딥러닝과 머신러닝의 발전이 세계 경제에 미치는 영향력으로 인해 모든 선진국에서 경쟁적으로 AI 발전에 나설 것으로 예상
- AI는 지정학적·군사적 힘의 원천이 되며, AI 개발을 주도하는 국가에 상당한 전략적 이점을 제공할 것으로 예상

(2) AI 기술 적용 확산에 따라 대두되는 리스크

〈표3. AI의 긍정적 혜택과 리스크 예시¹³⁾〉

AI의 긍정적 혜택		AI의 리스크	
의료 혁신	- 진단·치료: 의료 진단의 정확도 제고, 맞춤형 치료계획, 질병 조기 발견	프라이버시 침해	- 데이터 보안: AI 시스템이 수집하는 개인정보의 해킹·오용 위험
	- 가상 간호사: 환자 상태 모니터링, 치료 지침 제공, 병원내 업무 효율화		- 감시: AI 기술이 감시 카메라와 결합되어 개인의 프라이버시 침해
산업 자동화	- 제조업: 생산 공정의 자동화·효율성 제고, 품질관리 개선, 예측 유지보수	일자리 대체	- 자동화: AI 발전으로 일부 작업이 자동화되면서, 특정 일자리 도태 위기
	- 물류·공급망 관리: 공급망 최적화, 물류경로 계획, 재고 관리 효율화		- 재교육 필요: AI에 대체되지 않기 위해 새로운 기술 학습·재교육 압박 직면
금융 서비스	- 위험 관리: 위험 평가, 사기 탐지 강화, 투자 포트폴리오 최적화	편견과 차별	- 알고리즘의 편향성: 데이터에 내재된 편견이 결과에 반영될 수 있어 특정 그룹에 대한 차별 초래
	- 고객 서비스: 챗봇·가상 비서의 신속한 고객 대응, 맞춤형 금융서비스 제공		- 투명성 부족: AI의 결정과정 불투명
소비자 경험 향상	- 개인 맞춤형 마케팅: 고객 데이터 분석 통한 맞춤형 마케팅, 소비자 행동 예측	윤리적 문제	- 의사결정 책임: AI가 중요한 결정을 내리는 과정에서 책임소재 불분명
	- 엔터테인먼트: 맞춤형 콘텐츠 추천		- 안전성: 자율주행차, 의료 AI 등 시스템의 오작동으로 심각한 결과 초래

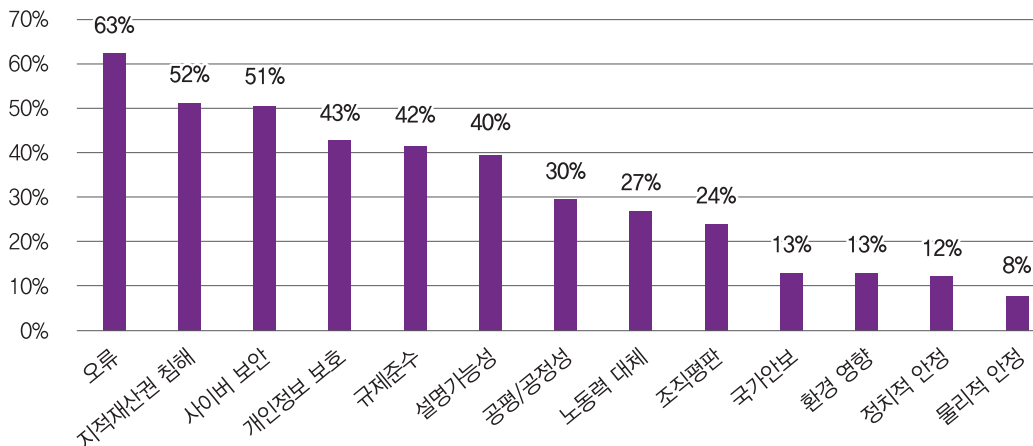
자료: Grand View Research(2024) 참고하여 저자 정리

13) Grand View Research(2024), AI Trust, Risk & Security Management Market Report, 2030, Grand View Research.

□ AI의 긍정적 효과 이면에는 데이터 활용에 따른 위험 요소 내포

- AI가 빠르게 발전하자 제품·서비스에 내장된 AI 기술이 사용자의 안전과 차별 금지, 표현의 자유, 인간의 존엄성, 개인정보 보호 등 기본권을 위태롭게 할 수 있다는 우려가 커짐에 따라 정부가 AI를 규제해야 한다는 주장이 확대
 - 오픈AI의 CEO 샘 알트먼과 애플의 공동창업자 스티브 워즈니악 등도 규제되지 않은 AI가 개인과 사회에 심각한 위협을 가하며 통제할 수 없는 피해를 초래할 수 있다고 경고¹⁴⁾
 - 챗봇은 보고서에 가짜 각주를 삽입하는 등 거짓말, 조작, ‘환각’을 일으킬 수 있고, 생성형 AI는 정치인과 연예인의 딥페이크 이미지와 동영상 제작할 수 있어 사기·기만·허위정보 확산에 악의적으로 활용될 우려¹⁵⁾
 - AI는 많은 노동자를 대체하며 기존 불평등 확대, 사회적 결속력 약화를 야기하고, 선거에 악용되는 경우 시민의 정치적 자율성 침해, 민주주의 훼손 및 강력한 감시 도구로서 개인의 기본권과 시민의 자유를 훼손할 가능성¹⁶⁾
- AI 리스크 관리 필요성이 커짐에 따라 기업들은 데이터의 효율적 활용뿐만 아니라 데이터 보안 및 위험 관리 등 노력을 기울이기 시작

〈그림5. 기업들이 고려하는 생성형 AI 관련 리스크〉



자료: Singla, Alex, et al.(2024) 참고하여 저자 정리

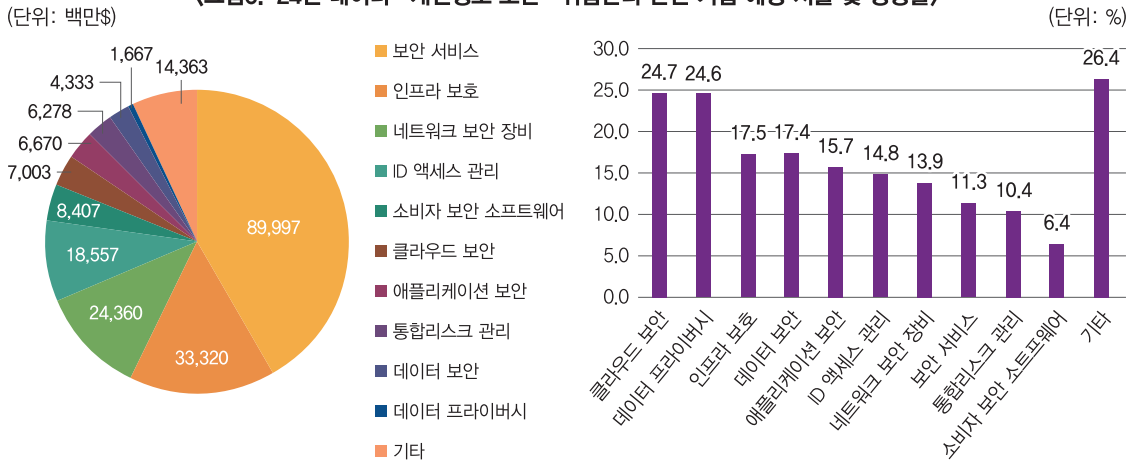
14) Bradford, Anu(2023.6.27.), "The Race to Regulate Artificial Intelligence", Foreign Affairs.

15) Newman, Matthew and Swift, Mike(2024.1.2.), "Comment: Global AI regulation taking diverse approaches, from rules-based to principles-based", MLex.

16) Acemoglu, Daron(2021), "Harms of AI", NBER Working Papers 29247, National Bureau of Economic Research.

- 기업들이 고려하는 생성형 AI 관련 리스크는 오류(inaccuracy), 지적재산권 침해, 사이버 보안, 개인정보 보호, 규제 준수 등¹⁷⁾
- 데이터 · 개인정보 보안 · 위험관리 관련 기업의 지출이 가장 많은 분야는 보안서비스로 '24년 900억 달러 규모의 기업지출이 예상되며, 이어서 인프라 보호 및 네트워크 보안장비 순이며 성장률이 가장 높은 분야는 클라우드 보안과 데이터 프라이버시¹⁸⁾

〈그림6. '24년 데이터 · 개인정보 보안 · 위험관리 관련 기업 예상 지출 및 성장률〉



자료: Gartner

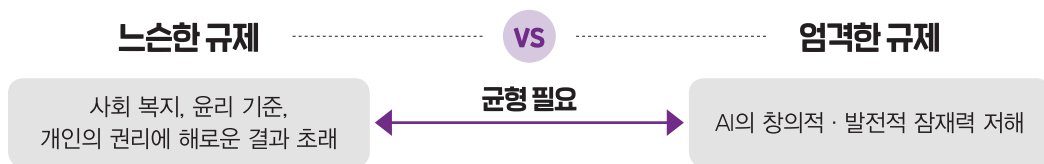
- AI는 긍정적 효과를 창출하는 동시에 리스크도 발생시켜 이를 최소화하기 위한 윤리적 고려와 규제가 필요
 - AI의 발전이 사회에 긍정적인 영향을 미치도록 하기 위해서는 기술발전과 함께 사회적 · 윤리적 논의가 지속적으로 이루어져야 함
- 빅테크는 비판과 감시 속에서도 AI 기술을 발전시키기 위해 경쟁해야 하고, 정부는 혁신을 저해하지 않으면서도 AI를 규제해야 하는 압박에 직면
 - AI 규제의 딜레마는 국가 통제와 혁신 사이의 균형¹⁹⁾
 - 지나친 규제는 AI의 창의적 · 발전적 잠재력을 저해할 수 있으나, 불충분한 규제는 사회 복지, 윤리 기준, 개인의 권리에 해로운 결과를 초래할 수 있음

17) '24년 2월 22일부터 3월 5일까지 AI를 도입한 기업의 전 직급 대상 1,052명에게 설문조사 진행. Singla, Alex, et al.(2024), *The state of AI in early 2024: Gen AI adoption spikes and starts to generate value*, McKinsey & Company.

18) Gartner(2023.9.28.), "Gartner Forecasts Global Security and Risk Management Spending to Grow 14% in 2024", *Gartner*.

19) Pernot-Leplay, Emmanuel(2024), "The AI Regulation Dilemma: A China, EU & U.S. Comparison".

〈그림7. AI 규제 딜레마〉



자료: 저자 정리

- 개인정보 보호, 보안, 경쟁력, 기술 리더십 등 각국의 국가목표, 사회·정치적 구조, 문화적 규범 등을 반영하여 국가별 AI 규제에 대한 다양한 접근법 채택
- AI 규제는 기술 혁신과 혜택을 촉진하는 동시에 잠재적인 위험으로부터 보호하고 국가의 목표와 가치에 부합하는 최적의 경로를 찾는 과정
- 주요국과 기업의 기술 리더십 경쟁이 심화되며 AI 규제가 추진력을 얻고 있으나, 국가별로 여전히 AI 규제 및 개발에 대해 상이한 견해와 목표가 대립²⁰⁾
 - 각기 다른 가치와 인센티브에 기반한 다양한 규제 패러다임이 등장하며 각국은 글로벌 디지털 경제에 대한 영향력을 확대하고 AI 우위를 차지하기 위한 경쟁적인 비전을 발전
- 美, 中, EU는 AI에 대한 서로 다른 비전으로 기술의 미래에 대한 통제권을 놓고 경쟁하며 디지털 세계에서 영향력을 확대 중
 - EU의 AI법은 일본, 인도, 브라질 등에 영향을 주었으나, 미국, 영국, 호주 등은 급성장하고 있는 산업에 대한 엄격한 규제가 투자와 혁신을 저해할 수 있어 규제보다 투자 유치 및 관련 연구·스타트업 허브를 지향²¹⁾
 - EU의 일반데이터보호규정(GDPR)이 전 세계 개인정보 보호법의 표준이 되었는데 AI법이 AI 규제의 글로벌 표준이 될 수 있을지는 아직 미지수²²⁾

20) Mok, Charles(2023.11.7.), "Global Competition for AI Regulation, or a Framework for AI Diplomacy?", *The Diplomat*.

21) Newman, Matthew and Swift, Mike(2024.1.2.), "Comment: Global AI regulation taking diverse approaches, from rules-based to principles-based", MLex.

22) Ibid.

2 주요국의 인공지능(AI) 규제 및 영향력 확대 전략

〈표4. EU, 미국, 중국의 AI 규제 및 영향력 확대 전략 비교〉

	EU	미국	중국
디지털 규제	기본권 중심 - 정부 개입 통한 기본권 보호 및 공정한 디지털 경제 조성	시장 중심 - 자유로운 언론·인터넷, 혁신에 대한 인센티브 보호에 중점	국가 중심 - 디지털 기술을 검열, 감시, 선전의 도구로 활용
AI 규제	위험 기반 포괄적 규제 - 거의 모든 부문을 포괄하여 광범위하게 적용	가이드라인, 모범사례 등 기업의 자발적 준수에 의존 - AI 기술의 성장·발전 여지 필요	신속하고 반복적인 AI 규제 - 새로운 AI 이슈에 대한 해결책 신속 제언, 민첩성 확보
약점	AI 대처 위한 민첩성 부족 - '21년 AI법 초안은 생성형 AI에 부적합, 법안 수정 필요	명확한 규제 모델 부재 - 연방 개인정보 보호법 부재	AI 규제의 안정성 부족 - 기업의 개발·투자 저해 가능성
영향력 확대 전략	규제 선도 - 디지털 기술에 대한 글로벌 규제 표준 설정(브뤼셀 효과)	기술 선도 - 기술 우위 및 리더십 유지	기술 및 규제 선도(개도국 위주) - 국제 표준 개발, 디지털 실크로드

자료: 저자 정리

(1) 유럽연합(EU) : 위험 기반 포괄적 규제

□ EU는 데이터·개인정보 보호 등 시민의 권리에 초점을 맞춘 디지털 기술 규제를 선도적으로 도입

- EU는 AI를 파괴적 잠재력을 지닌 디지털 혁신으로 간주하여 빅테크에게만 맡기기보다, 정부가 개입하여 개인의 기본권을 보호하고, 사회의 민주적 구조를 보존하며, 디지털 경제의 혜택을 공정하게 분배하고자 함²³⁾

- EU는 ▲개인정보를 보호하는 일반데이터보호규정(GDPR), ▲빅테크 등 디지털 게이트키퍼의 지배력을 축소하고 경쟁을 보호할 의무를 부과하는 디지털시장법(Digital Markets Act), ▲온라인 플랫폼이 호스팅하는 콘텐츠에 대한 책임을 지는 규정을 마련한 디지털서비스법(Digital Services Act) 등을 통해 권리 중심의 디지털 규제 시행

□ EU 집행위가 '21년 발의한 세계 최초의 포괄적 AI 규제인 AI법안(AI Act)은 '24.5월 EU 이사회의 최종 승인을 받으며 '26년 발효 예정

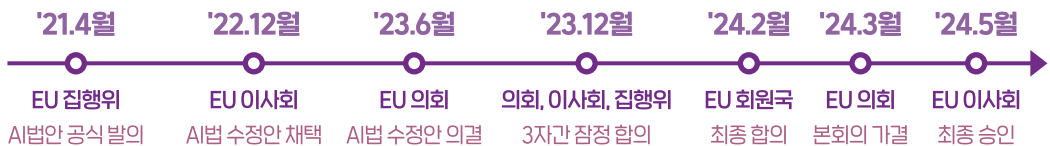
- AI법안 발의 후 EU 의회·이사회·집행위는 안면인식 기술 사용, 챗GPT 등 파운데이션 모델 규제 방법 등에 대해 격렬한 토론을 거쳐 결국 합의안을 도출²⁴⁾

23) Bradford, Anu(2023.6.27.), "The Race to Regulate Artificial Intelligence", Foreign Affairs.

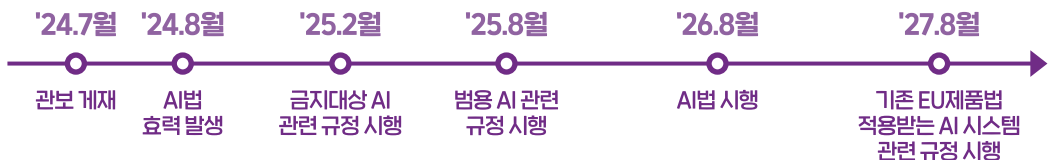
24) 김익현(2024.5.21.), "EU, 세계 첫 'AI법' 확정...이사회 최종 승인", ZDNET Korea.

- 생성형 AI 조항은 '21년 AI법안 초안에는 포함되지 않았으나, 이듬 해 챗GPT가 각광을 받으며 오남용 우려가 확산되자 관련 조항을 추가
- 프랑스, 독일, 이탈리아는 과도한 규제가 자국 AI 스타트업에 대한 투자를 위축시킬 수 있다고 우려하며 파운데이션 모델 규제에 반대했으나, 파운데이션 모델²⁵⁾과 동 모델을 기반으로 구축된 범용 시스템²⁶⁾을 구분하는 2단계 접근 방식에 합의

〈그림8. EU AI법 진행 경과〉



- AI법은 EU 관보에 게재된 날로부터 20일 후 효력이 발생되며, 효력 발생 후 24개월 경과 후 시행되므로 '26년 중 시행 예상
- 예외적으로 AI법의 금지대상 AI 관련 규정은 효력 발생 후 6개월, 범용 AI(GPAI) 관련 규정은 12개월, 기존 EU 제품법의 적용을 받는 AI 시스템의 경우 36개월 시점부터 시행²⁷⁾

〈그림9. EU AI법 시행 예상 시점 구분²⁸⁾〉

자료: Maples Group(2024.7.2.) 참고하여 저자 정리

- AI법은 AI 시스템과 범용 AI 모델을 EU 시장에 출시·서비스하려는 EU 소재 제공·배포자 및 AI 시스템에서 생성된 결과물을 EU에서 사용하려는 제3국 소재 배포·제공자에게도 적용
 - 공공·민간기관에서 군사·국방·국가안보 목적으로 출시·서비스·사용되는 AI 시스템, 과학적 연구·개발만을 목적으로 특별히 개발·서비스된 AI 시스템·모델 등에는 미적용
- AI법은 AI 시스템의 위험 정도를 ▲허용될 수 없는 위험, ▲고위험, ▲제한된 위험, ▲최소한의 위험으로 분류하여 차등 규제²⁹⁾

25) 대표적으로 Open AI의 GPT-4.

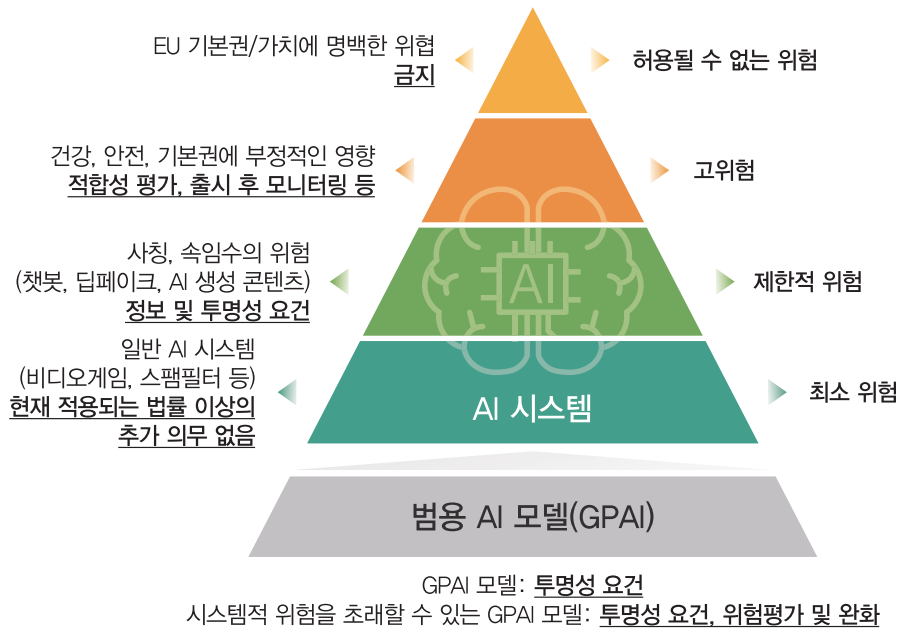
26) 오디오, 코드, 이미지, 텍스트, 비디오 등 새로운 콘텐츠를 생성할 수 있는 AI 앱인 챗GPT 등.

27) Madiega, Tambiama(2024.3.11.), *Artificial intelligence act*, Briefing, European Parliamentary Research Service (EPRS).

28) Maples Group(2024.7.2.), "AI Act Set to Come into Force on 1 August 2024", Maples Group.

29) Madiega, Tambiama(2024.3.11.), *Artificial intelligence act*, Briefing, European Parliamentary Research Service (EPRS).

〈그림10. EU AI법상 AI 시스템의 위험 구분〉



자료: Madiega(2024)

- 사람의 안전 · 생계 · 권리에 ‘허용될 수 없는 위험’을 초래하는 AI 시스템은 엄격히 금지

〈참고1. EU AI법상 금지된 AI 시스템〉

- ▶ 잠재적 · 조작적 · 기만적 기법 사용, 행동을 왜곡하고 의사결정을 저해하여 심각한 피해를 초래하는 시스템
- ▶ 연령, 장애 또는 사회 · 경제적 상황과 관련된 취약점을 악용하여 심각한 피해를 초래하는 시스템
- ▶ 인종, 정치적 견해, 노동조합 가입 여부, 종교 · 철학적 신념, 성적 취향 등을 유추하는 생체인식 분류 시스템 (법 집행 목적의 합법적인 라벨링 · 필터링은 제외)
- ▶ 사회적 행동 · 개인적 특성에 따라 사람을 평가 · 분류하여 해롭거나 불균형적인 대우를 초래하는 시스템
- ▶ 법 집행을 위한 공공장소 ‘실시간’ 원격 생체인식 시스템(납치 · 성착취 · 실종자 수색 · 실질적이고 임박한 특정 위협 방지 또는 심각한 범죄 용의자 식별 등의 특정 목적은 제외)
- ▶ 프로파일링 · 성격 특성만을 기반으로 개인의 범죄 위험성을 평가하는 시스템(범죄 행위와 관련된 객관적이고 검증가능한 사실에 근거한 평가는 예외)
- ▶ 인터넷 · CCTV 영상에서 불특정 스크래핑을 통해 안면인식 데이터베이스를 생성 · 확장하는 시스템
- ▶ 직장 · 교육기관에서 감정을 추론하는 AI 시스템(의료 · 안전상의 이유 제외)

— 금지 대상 AI 위반시 최대 3,500만 유로 또는 직전 회계연도의 전세계 연간 매출 7% 이하 중 더 높은 금액의 과태료 부과

- 사람의 건강, 안전, 기본권에 부정적인 영향을 미칠 수 있는 고위험 AI 시스템에 대해서는 엄격히 규제
 - 고위험 AI 시스템은 ▲의료기기·교통수단·장난감 등 부문별 EU 법률이 적용되는 제품의 안전 구성요소로서 제3자의 적합성 평가 대상인 경우, ▲생체인식, 핵심 기간산업, 교육·직업훈련, 고용·인사관리, 필수 공공·민간 서비스 및 그 혜택에 대한 접근·수혜, 법 집행, 이민·망명·출입국 관리, 사법행정과 민주적 절차 등 부속서 III에 기재된 8개 분야 등 두 범주로 구성
 - AI 시스템이 자연인의 건강, 안전 또는 기본권에 중대한 위해를 가할 위험이 없는 경우 고위험으로 간주되지 않으나(‘필터 조항’), 개인정보를 자동으로 처리하여 사람의 다양한 측면을 평가하는 개인 프로파일링의 경우 항상 고위험으로 간주
 - 고위험 AI 시스템 제공자는 역내 제품 판매·사용 전 내부통제(자체 평가) 또는 공인기관을 통해 ‘적합성 평가’(conformity assessment)를 받아야 하며, 추후 개발될 유럽 통일규격(harmonised standards)을 준수하는 경우에는 적합한 것으로 인정
 - 경우에 따라 AI 시스템이 EU 법률을 준수하는지 ‘기본권 영향평가’(fundamental rights impact assessment)를 받고, 제품 출시 후에는 모니터링을 실시하고 필요시 시정 조치를 취해야 함

〈표5. EU AI법상 고위험 AI의 주요 준수 의무〉

시장 출시 전	시장 출시 후
기본권 영향평가(Fundamental Rights Impact Assessment) – 개발된 AI 시스템을 공급받아 활용하고자 하는 사용자(Deployer)가 개별 서비스 활용 맥락에서 AI 시스템이 기본권에 미칠 수 있는 영향 평가	– 생성로고 최소 6개월 이상 유지 – 법 미준수사항 발생시 즉시 시정조치 및 타 관계자들에게 통지 – AI Office ³⁰⁾ 와 협력하여 준수를 입증할 수 있는 모든 정보와 문서 공유
적합성 평가(Conformity Assessment) – 제공자(Provider)의 고위험 AI의 의무규정 준수여부 평가	– 전 수명주기 전반에 걸친 위험관리 시스템 구축 – 기본권 침해 유발 중대사건 및 시스템 장애 발생 시 규제당국 보고 – 중대한 변경사항 발생시 적합성 평가 재 실시

- 챗봇 등 사람과 상호 작용하는 시스템, 감정 인식 시스템, 생체인식 분류 시스템, 딥페이크 등 이미지·오디오·비디오 콘텐츠를 생성·조작하는 AI 시스템 등 제한된 위험을 초래하는 AI 시스템에는 투명성 의무 적용
 - 사용자는 챗봇과 상호작용한다는 사실을 인지해야 하며, 이미지·오디오·비디오 콘텐츠를 생성·조작하는 AI 시스템의 배포자는 범죄 예방에 사용되는 매우 제한된 경우를 제외하고 해당 콘텐츠가 인위적으로 생성·조작되었다는 사실을 공개해야 함

30) EU집행위 내 설치되는 AI법 총괄기구.

- 대량의 합성 콘텐츠를 생성하는 AI 시스템 제공자는 워터마크 등 신뢰할 수 있고 상호 운용가능하며 효과적인 기술·방법을 구현하여 해당 결과물이 사람이 아닌 AI 시스템에 의해 생성·조작되었음을 표시해야 함
- 직장에 AI 시스템을 배치하는 고용주는 근로자에게 이를 통지해야 함
- 스팸 필터, 비디오 게임 등 위험이 낮거나 최소한의 위험만 있는 다른 모든 AI 시스템은 GDPR 등 현재 적용되는 법률 이상의 추가 의무 미적용

□ AI법은 범용 AI 모델과 시스템적 위험을 초래할 수 있는 범용 AI 모델을 구분하여 별도 규제³¹⁾

〈표6. 범용 AI 모델 및 시스템적 위험이 있는 범용 AI 모델 제공자 의무 구분〉

시스템적 위험이 있는 범용 AI 모델
범용 AI 모델
- 최신 기술문서 작성·유지³²⁾ - AI 시스템의 하위 제공자에게 관련 정보와 문서 제공 - 워터마킹 등 저작권 지침을 존중할 수 있는 정책 마련 - 학습에 사용된 콘텐츠에 대한 상세 요약본 공개 - 역외 사업자는 역내 대리인 지정
- 적대적 테스트 등 모델 평가 수행 - 시스템 위험 평가·완화 및 사이버 보안 확보 - 기본권 침해 등 심각한 사건 추적·문서화·보고 및 시정조치 이행

자료: 저자 정리

- AI법은 ‘범용 AI 모델’을 대량의 데이터를 자기지도학습으로 훈련되고, 상당한 범용성을 보이며, 다양한 작업을 능숙하게 수행할 수 있고, 다양한 하위 시스템·앱에 통합될 수 있는 모델로 정의³³⁾
- ‘범용 AI 시스템’은 범용 AI 모델을 기반으로 하는 AI 시스템으로, 직접 사용되는 경우와 다른 AI 시스템과 통합되어 사용되는 경우 모두 포함
- 범용 AI 모델 제공자는 제한된 위험을 초래하는 AI 시스템에 적용되는 투명성 의무를 준수해야 함
 - 추가적으로 ▲기술문서 작성·유지, ▲범용 AI 모델을 자신의 AI 시스템에 통합하려고 하는 AI 제공자에게 관련 정보와 문서 제공, ▲저작권 지침을 존중할 수 있는 정책 마련 및 학습 데이터에 대한 상세요약 게시 의무 부과

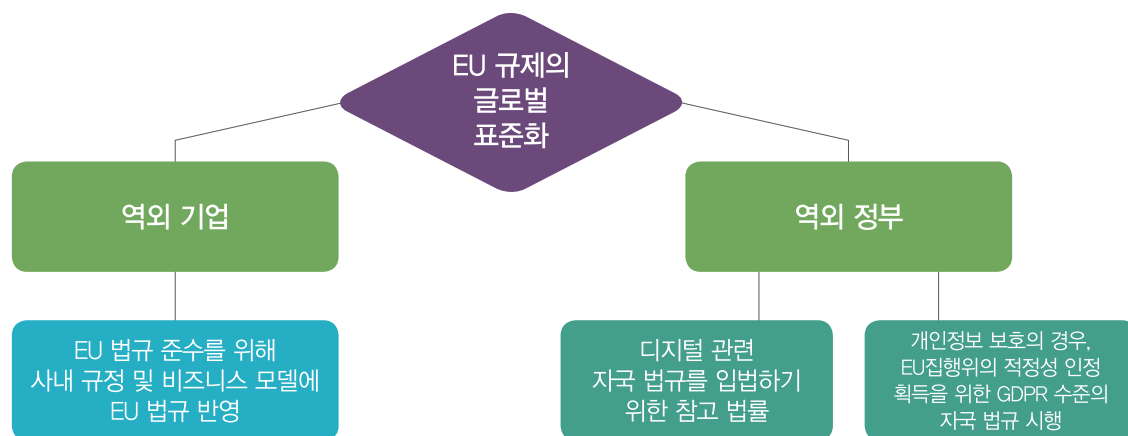
31) Madiega, Tambiama(2024.3.11.), *Artificial intelligence act*, Briefing, European Parliamentary Research Service (EPRS).

32) 무료 오픈소스를 통해 접근가능한 AI 모델은 기술문서 공개 등 일부 의무 면제.

33) 오픈AI의 챗GPT가 대표적인 예.

- 범용 AI 모델 제공자는 EU 이행규정을 통해 제정될 실행강령을 통해 AI법 준수여부를 입증할 수 있을 것으로 예상
 - 시스템적 위험이 있고 시장에 중대한 영향을 미칠 수 있는 범용 AI 모델에는 강화된 규제를 적용
 - 범용 AI 모델의 학습훈련에 사용된 누적 컴퓨팅의 양이 10^{25} FLOPs³⁴⁾ 이상인 경우 시스템적 위험이 있는 모델로 간주되며, 높은 영향력을 갖고 있다고 EU 집행위가 판단하는 모델은 추가로 지정될 수 있음³⁵⁾
 - 시스템적 위험이 있는 범용 AI 모델에는 ▲적대적 테스트 등 모델 평가 수행, ▲시스템 위험 평가·완화, ▲심각한 사건 및 대응조치 기록 및 관련 국가기관 보고, ▲적정 수준의 사이버 보안 확보 등의 의무 부과
- EU는 GDPR과 마찬가지로 AI법을 통해 글로벌 표준을 형성하고 영향력을 확대하려 함

〈그림11. EU 규제의 브뤼셀 효과〉



자료: 저자 정리

- 빅테크는 전 세계적으로 제품과 서비스를 표준화하기 위해 EU의 엄격한 규제를 글로벌 비즈니스 운영 전반에 확대 적용하는 경우가 많으며, 이를 ‘브뤼셀 효과’라고 함³⁶⁾

34) 현재 오픈AI의 GPT-4, Meta의 Llama 2 및 구글의 Gemini 등 극소수 모델이 해당 위험임계값을 초과하며 매우 높은 계산능력을 갖고 있어, 이로 인한 잠재적 위험성을 주의 깊게 관리해야 한다는 것을 시사.

35) 컴퓨터의 처리 속도를 측정하는 FLOPs(초당 부동 소수점 연산)는 시간이 지남에 따라 기술 및 산업 변화를 반영하여 조정될 것으로 보이며, EU 집행위가 학습 데이터의 규모, 최종 사용자 수, 자율성 및 확장성 등에 대한 전반적인 평가를 바탕으로 해당 모델이 FLOP 임계값과 동등한 기능·영향을 미치는 것으로 판단하는 경우, 시스템적 위험이 있는 모델로 지정할 수 있음.

36) Bradford, Anu(2023.6.27.), “The Race to Regulate Artificial Intelligence”, Foreign Affairs.

- 역내에서 취득한 데이터를 사용하여 알고리즘을 학습시키는 AI 개발자는 역외에서도 EU의 AI법에 구속되며, EU의 규제에서 벗어나려면 역내 데이터 없이 완전히 새로운 알고리즘을 개발해야 함
- EU는 역내뿐만 아니라 역외에서 AI 규제를 형성하고, EU의 권리 중심 규범을 세계화하며, 그 과정에서 EU의 디지털 영향력 확대에 주도적인 역할을 행사하고자 함
- 미국 빅테크에 대한 반발이 거세지는 가운데 일본, 인도, 브라질 등 다수 국가가 시장 중심의 틀에서 벗어나 디지털 경제에 대한 통제권을 되찾으려 하고 있어, 점차 EU의 디지털 규제를 많이 참고하는 추세³⁷⁾
- EU는 GDPR과 마찬가지로 AI법이 전세계적으로 모범적인 법적 프레임워크로 조기에 활용되어 글로벌 관행을 형성하길 희망

37) Newman, Matthew and Swift, Mike(2024.1.2.), "Comment: Global AI regulation taking diverse approaches, from rules-based to principles-based", MLex.

(2) 미국 : 가이드라인, 모범사례 등 기업의 자발적 준수에 의존

□ 다수의 빅테크 기업을 보유한 미국은 AI 규제에 대한 정부의 역할이 제한적이며, 가이드라인과 모범사례 등 기업의 자발적 준수를 중시

- 미국은 언론의 자유, 자유로운 인터넷, 혁신에 대한 인센티브를 보호하는데 중점을 두고 있으며, 디지털 기술을 경제적 번영과 정치적 자유의 원천이자 사회 변화와 진보를 위한 수단으로 간주하며 AI 규제에 미온적³⁸⁾
 - AI 기술이 규제보다 성장하고 발전할 여지가 필요하다는 논리와 데이터 · 개인정보 보호, 플랫폼의 책임 등 빅테크 규제 법률이 부족하다는 비판이 충돌해왔음
- 미-중 기술패권 경쟁 격화, 지정학적 긴장 고조에 따라 미국은 AI를 경제성장 촉진 및 기술 · 군사적 우위 확보의 기회로 간주하여 규제의 대상보다는 경제 · 지정학적 우위의 수단으로 봄³⁹⁾
 - AI의 안전 · 보안 · 신뢰성이라는 동일한 목표를 달성하기 위해 EU가 입법과 규제에 중점을 두고 있는 반면, 미국은 기술 우위 및 리더십 유지에 초점
- 미국의 시장 중심 정책은 부를 창출하고 기술발전을 촉진했으나, 빅테크 규제 부재의 대가 또한 부각되고 있음⁴⁰⁾
 - 디지털 광고 기술 독점으로 경쟁업체에 피해를 주는 등 시장지배력의 반복적 남용, 사용자 개인정보 착취에 관한 허위정보 및 폭로 확산으로 빅테크에 대한 불신 확산
 - 자국 인터넷 사용자에게 대한 미국 빅테크의 막강한 영향력을 줄이기 위해 각국 정부는 디지털 시장에 대한 통제권을 확보하고 주요 빅테크를 견제하기 위해 노력하고 있으며, 미국에서도 빅테크에 대한 정부의 감독 강화를 요구
- 의회 내 포괄적인 AI 규제가 난항을 겪자 미국에서의 AI 규제는 행정부 주도로 진행⁴¹⁾
 - 혁신 저해와 기술 리더십 약화를 우려하는 목소리와 의회의 정치적 교착상태로 AI 기술을 규제하기 위한 법안의 합의가 지연되는 사이, 미국은 행정부 주도의 AI 정책으로 글로벌 AI 규제 경쟁에서 주도권을 유지

□ 바이든 행정부는 '23.7월 7개 주요 AI 기업들로부터 안전하고 투명한 AI 기술 개발을 위한 자발적 조치를 약속받는데 이어⁴²⁾, 8월 8개 기업으로부터 자발적 약속 추가 확보⁴³⁾

38) Bradford, Anu(2023.6.27.), "The Race to Regulate Artificial Intelligence", Foreign Affairs.

39) Ibid.

40) Ibid.

41) Mok, Charles(2023.11.7.), "Global Competition for AI Regulation, or a Framework for AI Diplomacy?", The Diplomat.

〈표7. AI 기업이 약속한 AI 개발원칙 주요 내용〉

구 분	주요 내용
안 전	- AI 시스템 출시 전 내·외부 보안 테스트 실시 - 정부·시민사회·학계와 AI 위험 관리에 관한 정보(모범사례, 우회시도, 기술협력 등) 공유
보 안	- 사이버 보안 및 내부 위협 보호 장치에 투자하여 독점·미공개 모델 가중치 보호 - 제3자가 AI 시스템의 취약점을 발견·보고할 수 있도록 지원 노력
신뢰	- 사용자가 AI로 생성된 콘텐츠임을 알 수 있도록 하는 기술 개발(워터마크 등) - AI 시스템의 능력, 한계, 적절·부적절한 사용영역 등을 공개 - 편견·차별 방지, 개인정보 보호 등 AI 시스템이 초래할 수 있는 사회적 위험 연구 - 암 예방, 기후변화 대응 등 사회적 과제에 대응하는 첨단 AI 시스템 개발·배포

자료: The White House(2023.7.21.) 참고하여 저자 정리

〈표8. 자발적 조치를 약속한 15개 AI 기업〉

1차('23.7월, 7개사)	2차('23.9월, 8개사)
      	       

□ 바이든 행정부는 '23.10월 AI를 규제하는 첫 행정명령을 발표하며, 국가안보·건강·안전을 위협하는 AI 기술 개발과 이용에 대한 규제 의지를 표명⁴²⁾

○ 총 8개 부분으로 구성되어 있는 행정명령은 (1) AI 안전·보안 기준, (2) 개인정보 보호, (3) 평등과 시민권 향상, (4) 소비자·환자·학생 보호, (5) 노동자 지원, (6) 혁신·경쟁 촉진, (7) 미국의 리더십 강화, (8) 연방정부의 AI 사용·조달을 위한 지침 개발 등에 관한 내용 포함

— 행정명령의 핵심 부분은 AI 안전·보안 기준으로, 가이드라인과 표준에 대한 규칙을 제정하고 잠재적 이중용도 파운데이션 모델 개발자가 모델 훈련 및 레드팀 안전 테스트 결과에 대한 정보를 연방정부에 보고하도록 요구⁴³⁾

42) The White House(2023.7.21.), "FACT SHEET: Biden-Harris Administration Secures Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI".

43) The White House(2023.9.12.), "FACT SHEET: Biden-Harris Administration Secures Voluntary Commitments from Eight Additional Artificial Intelligence Companies to Manage the Risks Posed by AI".

44) The White House(2023.10.30.), "FACT SHEET: President Biden Issues Executives Order on Safe, Secure, and Trustworthy Artificial Intelligence".

— 행정명령의 성공 여부는 ‘자발적 약속’에 근거한 AI 모델 개발자들의 협력이 좌우⁴⁵⁾

〈표9. AI 행정명령 주요 내용〉

조 치	세부 내용
AI 안전·보안 기준	— 국가안보, 경제안보, 공중보건·안전에 심각한 위험을 초래하는 파운데이션 모델 개발 기업은 모델 훈련시 연방정부에 통보하고 모든 레드팀 안전 테스트 결과 공유 — AI 시스템의 안전·보안·신뢰성을 위한 표준·도구·테스트 개발 — AI가 생성한 콘텐츠 탐지 및 공식 콘텐츠 인증을 위한 표준·모범사례를 수립하여 AI를 이용한 사기·속임수로부터 보호 등
개인정보 보호	— 개인정보 보호 연구 및 기술 개발·사용을 가속화하기 위한 연방정부 지원 — 상업적 가용정보 수집·사용 방법 평가 및 연방기관의 개인정보 보호 지침 강화 — 연방기관이 개인정보 보호 기술의 효과를 평가할 수 있는 가이드라인 개발 등
평등과 시민권 향상	— AI 알고리즘이 차별을 악화하지 않도록 임대인·연방 계약업체 등에게 명확한 지침 제공 — AI 관련 시민권 침해 조사·기소 모범사례 교육·기술지원 등을 통해 알고리즘 차별 해결 — 선고, 가석방·보호관찰, 구금, 위험평가, 감시, 범죄 예측·예측 치안, 법의학 분석에 AI를 사용하는 모범사례를 개발하여 형사사법 시스템 전반의 공정성 보장 등
소비자·환자 보호	— 의료 분야 및 약품 개발에서 AI의 책임감 있는 사용 촉진 등
노동자 지원	— 일자리 대체, 노동기준, 형평성, 건강, 안전, 데이터 수집에 관한 원칙·모범사례 개발 — 노동시장에 대한 AI의 잠재적 영향에 대한 보고서를 작성하고, 노동종단에 직면한 근로자를 위한 연방지원 강화 방안 연구 등
혁신·경쟁 촉진	— 의료·기후변화 등 중요 분야 AI 연구에 대한 지원금 확대 통해 AI 연구 촉진 — 소규모 개발자에게 기술지원 제공, 중소기업의 AI 혁신 상용화 지원 — 비자 기준, 인터뷰·심사 현대화·간소화를 통해 중요 분야에 전문성을 갖춘 고숙련 이민자·비이민자의 미국 유학·체류·취업 기회 확대 등
국제무대에서 미국의 리더십 강화	— AI 협력을 위한 양자, 다자, 다중 이해관계자 참여 확대 — 지속가능 발전과 중요 인프라에 대한 위험 완화 등 글로벌 과제를 해결하기 위해 안전하고 책임감 있는 해외 AI 개발·사용 촉진 등
연방정부의 책임감 있고 효과적인 AI 사용	— 권리·안전 보호, AI 조달 개선 등 명확한 기준을 포함한 기관의 AI 사용지침 발행 — 신속·효율적인 계약을 통해 AI 제품·서비스의 신속·효과적 확보 지원 — AI 전문가의 채용 가속화 등

자료: The White House(2023.10.30.) 참고하여 저자 정리

— AI 행정명령은 AI 기술 개발 기업에게 일정 수준의 책임을 부여하고, 이중용도 프로젝트에 대한 정부 감독을 위한 기반을 마련했으나, 엄격한 준수요건과 법적 구속력을 갖춘 EU 방식과 달리, 규제보다 산업개발 정책과 지침에 가깝음

45) Mok, Charles(2023.11.7.), “Global Competition for AI Regulation, or a Framework for AI Diplomacy?”, *The Diplomat*.

46) Ibid.

〈표10. EU 인공지능법과 미국의 AI 행정명령 비교⁴⁷⁾〉

	EU 인공지능법안(AI Act)	美 AI 행정명령(AI EO)
공 통 점		
파운데이션 모델	AI 파운데이션 모델 개발에 엄격한 의무사항을 부여함으로써 AI 파운데이션 모델을 통제하고자 함	
테스트 및 모니터링	생애주기 전반에 걸쳐 AI 시스템을 테스트 및 모니터링할 필요성 부각	
	- 개발자에게 출시 전 테스트와 출시 후 모니터링 절차를 준수할 것을 요구	- 오남용 및 비윤리적 문제를 방지하기 위해 테스트 및 모니터링의 중요성 언급
개인정보 보호	AI 시스템을 훈련시키기 위한 예외를 허용하지 않는 등 개인정보 보호의 중요성 강조	
사이버보안	사이버보안 표준 준수 요구	
차 이 점		
집 행	EU 단일시장 및 모든 산업에 바로 적용가능한 통일된 규제 제정	민간산업을 직접 규제하지 않으며 강제력 없음 - AI 개발 관련 산업표준, 가이드라인, 관행, 규제를 만들 권한을 정부부처에 부여
주요 초점	위험수준에 따라 AI 시스템을 분류하는 위험 기반 규제 프레임워크	표준 및 가이드라인에 중점, AI 관련 위험을 완화하기 위한 국제협력 중요성 강조
기 타	기업의 준수 측면에 중점	이민, 교육, 노동 관련 이슈 포함

자료: Fauchoux, Vincent and Perissinotto, Juliana(2023).

- AI 행정명령은 향후 연방정부가 채택할 표준·가이드라인과 함께 美 국내 AI 정책이 다양한 채널을 통해 전 세계에 영향을 미칠 것으로 예상⁴⁸⁾

- 방대한 美 조달시장은 기업이 미국의 AI 표준과 가이드라인을 준수하도록 동기를 부여하며, 미국의 AI 표준은 국제표준화기구(ISO)·국제전기표준회의(IEC) 등 국제기구의 공식 절차를 통해 더욱 발전하고 세계화가 될 것으로 전망

- 기술 발전과 경제 성장을 촉진할 AI의 잠재력과 빠르게 진화하는 기술을 규제해야 하는 어려움 때문에 일부 국가는 미국식 규제 모델을 선호할 가능성이 있으나, 미-중 경쟁과 ‘브뤼셀 효과’를 의식한 미국이 EU와의 협력을 심화할 가능성도 존재

- 영국, 호주 등은 과도한 규제가 투자와 혁신을 저해할 것이라고 우려하며, EU의 경직된 규칙 기반 프레임워크보다 가볍고 유연한 원칙 기반 시스템을 선호⁴⁹⁾

47) Fauchoux, Vincent and Perissinotto, Juliana(2023.9.11.), “Artificial intelligence: Joe Biden’s executive order vs. the European AI Act – Common Goals?”, DDG.

48) Meltzer, Joshua P.(2023.11.1.), “Toward international cooperation on AI governance—the US executive order on AI”, Brookings.

49) Newman, Matthew and Swift, Mike(2024.1.2.), “Comment: Global AI regulation taking diverse approaches, from rules-based to principles-based”, MLex.

- 빅테크는 미국 요건을 대부분 포괄하는 EU의 AI법을 준수하기 위해 관행을 구축할 것으로 예상되어 미국은 EU가 ‘브뤼셀 효과’를 통해 미국 시장을 일방적으로 규제하기보다 공동으로 규제를 설정하는 것을 선호할 것으로 예상⁵⁰⁾
- 중국의 글로벌 영향력 확대에 대해 미국과 EU 모두 우려하고 있어 미-EU 양자간 협력 확대를 모색할 전망⁵¹⁾
 - 바이든 행정부는 기술 우위를 차지하기 위한 경쟁에서 민주적 가치에 부합하는 규범을 장려하고 디지털 권위주의에 맞서기 위해 EU 및 기타 동맹국들과 협력하고자 하는 의지를 거듭 강조
 - 미국과 EU는 AI 규제에 대한 상호 입장 차이에도 불구하고 혁신 촉진, 기본권 보호, 민주주의 수호 등을 위해 공동 AI 표준을 개발하는 등 협력을 심화할 가능성

50) Bradford, Anu(2023.6.27.), “The Race to Regulate Artificial Intelligence”, Foreign Affairs.

51) Ibid.

(3) 중국 : 신속하고 반복적인 규제

□ 중국은 AI 개발을 위한 포괄적인 계획의 일환으로 가장 먼저 AI를 규제한 국가로, 국가의 통제를 유지하면서 기술혁신을 장려하는데 중점

- 중국은 세계 최고의 기술 강국이 되기 위해 국가 주도의 디지털 규제를 채택⁵²⁾
 - 중국은 AI가 사회 안정과 공산당의 정치적 통제를 훼손하지 않도록 통제하며 기술발전을 장려
 - 초기 자국 기술 산업의 성장을 촉진했으나, 최근 몇 년 동안 ‘공동 번영’이라는 명분으로 기술 산업에 대해 적극적으로 단속하며, 기술 산업이 국가를 압도하지 못하도록 노력
 - 사회 안정을 유지한다는 명목으로 자국민에 대한 대규모 검열·감시·선전 능력을 향상시킬 AI 기술에 막대한 보조금을 지급
- 중국은 EU처럼 AI를 포괄적으로 규제하기보다 새로운 AI 관련 이슈에 대해 단편적인 규칙을 신속하게 제정하고 지속적으로 수정⁵³⁾

〈표11. EU와 중국의 AI 규제 방식 비교〉

E U	중 국
– 위험 기반 포괄적 규제 AI법 – 금지 대상 AI, 고위험 AI – 제한적 위험을 지닌 AI(챗봇, 딥페이크, AI 생성 콘텐츠) – 범용 AI(GPAI)	– 이슈별 단편적 규제 신속 도입 인터넷 정보서비스 알고리즘 추천 관리 규정 인터넷 정보서비스 심층 종합 관리 규정(딥페이크) 생성형 AI 서비스 관리 잠정방법

자료: 저자 정리

- 중국의 개인정보 보호법(个人信息保护法)은 EU GDPR의 원칙과 개념에서 크게 영향을 받았으나, AI 규제에 대해서는 독자적인 방향을 설정
- 중국 정부는 알고리즘 추천 서비스⁵⁴⁾, 딥페이크⁵⁵⁾, 생성형 AI⁵⁶⁾에 관한 규범을 각각 '21년, '22년, '23년 통과시키며 새로운 AI 이슈별 규제를 신속히 도입
- 정부와 사용자는 기술 발전에 따른 위험에 신속하게 대응할 수 있으나, 장기적·거시적 발전을 도모하기 어렵고, 기업의 개발과 투자도 저해할 가능성

52) Bradford, Anu(2023.6.27.), “The Race to Regulate Artificial Intelligence”, Foreign Affairs.

53) Pernot-Leplay, Emmanuel(2024), “The AI Regulation Dilemma: A China, EU & U.S. Comparison”.

54) 인터넷 정보 서비스 알고리즘 추천 관리규정(互联网信息服务算法推荐管理规定).

55) 인터넷 정보 서비스 심층 종합 관리규정(互联网信息服务深度合成管理规定).

56) 생성형 인공지능 서비스 관리 잠정방법(生成式人工智能服务管理暂行办法).

○ 중국은 모든 파운데이션 모델을 출시 전 정부에 등록하도록 요구⁵⁷⁾

– 중국 정부는 외국 AI 기업의 제품을 승인하지 않으며 사실상 자국 기업을 보호하고 있으며, 이를 통해 경쟁을 억제하고 온라인 발언에 대한 통제를 강화⁵⁸⁾

□ 중국은 '23.7월 7개 부처⁵⁹⁾ 공동으로 생성형 AI에 대한 세계 최초 규제인 '생성형 AI 서비스 관리 잠정 방법(生成式人工智能服务管理暂行办法, 이하 '잠정 방법)'을 발표하고 한 달 뒤 전격 시행

○ 잠정 방법은 ▲생성형 AI의 건전한 발전, ▲표준화된 적용 추진, ▲국가안보 및 사회 공공이익 수호, ▲권익 보호를 목표로 기술개발 촉진, 준수 의무 등을 명시하고 있으며, '23.4월 발표된 초안보다 다소 완화된 수준의 규제를 도입⁶⁰⁾

〈표12. 중국 생성형 AI 잠정관리 방법 주요 내용〉

일반 조항	– 적용 대상 정의 – 주요 원칙 및 준수·금지사항 규정
기술개발 및 관리	– 혁신적 응용 장려 – 다양한 분야간 협력 장려 – 자발적 혁신과 국제협력 장려 – 자원 공유 장려
서비스 규범화	– 생성형 AI로 인한 사회문제에 적극 대응 – 생성형 AI 서비스를 제공·사용하는 법적 요건 규정 – 세밀한 책임 구분, 생성형 AI 관리 체계 구축 규정
감독 및 법적책임	– 보안 평가 및 알고리즘 등록의 요구사항 명시 – 정보공개 요구사항 규정

○ 적용 대상은 생성형 AI 기술을 이용하여 중국 주민을 대상으로 텍스트·이미지·음성·동영상 등 콘텐츠 생성 서비스를 제공하는 자로, 중국 주민을 대상으로 하지 않는 산업단체·기업·연구 기관 등은 제외

– 국외 생성형 AI 서비스 제공자도 중국 주민을 대상으로 하는 경우 적용되며, 생성형 AI 서비스가 중국에서 개발되어도 국외 사용자만 대상으로 할 경우 적용받지 않음

○ 잠정 방법의 특징은 생성형 AI가 만들어 내는 콘텐츠가 사회주의 핵심 가치를 견지하도록 규정하며, 국가권력, 사회주의 체제의 전복, 국가안보와 이익을 위태롭게 하는 내용을 포함하지 못하도록 함

57) '23년말 기준, 22개 기업이 AI 모델을 등록.

58) Ryan-Mosley, Tate, et al.(2024.1.5.), *What's next for AI regulation in 2024?*, What's Next in Tech, MIT Technology Review.

59) 국가인터넷정보판공실(国家互联网信息办公室), 국가발전개혁위원회, 교육부, 과학기술부, 공업정보화부,公安部, 국가광파전시총국(国家广播电视总局) 등.

60) 초안에 있던 규정 위반시 벌금(최대 10만 위안), 3개월 내 시정명령 이행 등 일부 엄격한 조항이 최종안에서 제외되고, 초안에 없었던 생성형 AI 응용 장려, 기술개발 및 교류 지원 내용이 추가.

- 여론 형성의 속성을 가지거나 사회 선동력을 가진 생성형 AI 서비스 제공 시에는 '안전성 평가'와 '알고리즘 등록'을 실시하도록 규정
- 안전성 평가와 알고리즘 등록 제도는 EU의 AI법이나 미국의 행정명령에도 비슷하게 규정되어 있으나, 중국에서 실시하는 안전성 평가는 기준이 명확하지 않고 안전성 여부를 정부가 판단하여 정부의 자의성이 비교적 큼⁶¹⁾
- 잠정 방법의 적용 대상인 AI 서비스 제공자는 ▲연구·개발 과정에 데이터 라벨링 규칙 제정 및 품질 평가, ▲생성형 콘텐츠 정확성·신뢰성을 높이면서 차별, 프라이버시 침해 등 문제 정보 제거 등의 의무를 준수해야 함
 - 또한 이미지·동영상 등 AI로 생성된 콘텐츠에 표기를 해야 하며, 사용자가 입력한 정보와 사용기록을 보호하고 개인정보 보호를 위한 개인의 요청을 적시 처리해야 하며, 정상적인 사용을 보장하기 위해 안정적·지속적 서비스를 제공하고 합법적인 사용을 유도해야 함
- 제공자가 잠정 방법을 위반하는 경우, 주무부서는 관련 법규⁶²⁾에 따라 처벌하거나 경고·시정명령을 할 수 있으며, 이를 거부하거나 관련 상황이 심각한 경우 서비스를 중지하거나 치안관리 또는 형사 책임을 물 수도 있음
 - 제공자는 불법적인 내용이나 사용자 불법 행위 발견시 내용 삭제, 기능 제한, 서비스 중단 등 조치 후 주무부서에 보고해야 함
- 잠정 방법은 전반적으로 기술 개발부터 운영·관리 과정에 의무를 부과하는 성격이 크지만, 기술 개발을 촉진하고 산업 발전을 지원하는 방안도 포함
 - 고품질 콘텐츠 생성 장려, 기술 혁신·활용과 위험 예방을 위한 기업·기관 협력 지원, 생성형 AI 알고리즘부터 반도체 등 기본기술 자립 장려, 훈련 데이터 개방과 효율적 이용 등 조항을 포함

□ '23년 국무원은 AI 법안의 입법 계획을 발표한 후 '24.3월 초안 공개

- 공식적으로 '전문가 제안'⁶³⁾이라고 불리는 이번 초안은 미-중 AI 경쟁에 대한 대응책의 일환으로 AI 시스템 사용자 보호보다는 AI 산업 발전에 초점⁶⁴⁾
 - '23.8월 사회과학원이 공개한 전문가 제안은 중국의 AI 개발을 감독할 '국가 AI 사무국' 신설, 파운데이션 모델에 대한 독립적이고 연례적인 '사회적 책임보고서' 제출, 위험성이 높은 AI 분야의 '네거티브 리스트'를 설정하여 정부 승인 없이 연구를 할 수 없도록 제안

61) 윤성혜(2023.11.10.), "법을 제정속도 느린 중국, AI만은 예외다", 프레시안.

62) 인터넷보안법, 데이터안전법, 개인정보보호법, 과학기술진보법 등.

63) 지난 8월 중국 사회과학원이 공개한 전문가 제안에 따른 것.

64) Wang, Juan(2024.4.17.), "Chinese experts suggest negative list system be included in proposed AI law", MLex.

- 이번 초안은 세제 혜택과 대학 인재 양성 등 정책적 지원을 포함해 국가가 해야 할 기여를 강조하며 전반적으로 AI 산업의 신속한 발전을 위한 법적 구조 마련에 중점을 두고 있으나, AI 제공자가 중국의 데이터 및 개인정보 보호법을 준수하도록 요구하고, 개인의 행동 · 습관, 관심사 · 취미, 경제 · 건강 · 신용정보 분석 · 평가 등 특정 데이터 사용을 제한
- 동 초안은 알고리즘 차별 방지, 개인정보 보호, AI가 생성한 콘텐츠에 대한 라벨링 등 미국의 AI 행정 명령 내용 일부를 반영했는데, AI 규제와 개발은 본질적으로 초국가적이기 때문에 AI 관련 우려사항과 해결법에는 미-중간 관점이 일치⁶⁵⁾
- 초안은 AI 제공자가 학습 목적으로 저작권 자료의 사용 방법을 포함하고 있으나, 저작물 보호보다는 산업 발전에 유리한 방향으로 규정⁶⁶⁾
 - AI 프로그램은 AI가 생성한 콘텐츠를 수정하거나 저작권 보호를 신청할 수 없지만, 텍스트 또는 이미지 생성기를 사용한 사람은 결과물을 충분히 변경하면 저작권 보호 신청 가능⁶⁷⁾
- AI 규제에 대한 중국의 접근법은 개인에 대한 피해를 완화하고 사회 안정과 국가 통제를 유지하는 동시에 AI 분야의 글로벌 리더십을 확보하여 규제 논의에 영향을 미치는데 초점⁶⁸⁾
 - 중국은 AI 분야 경쟁우위를 확보하기 위한 국제표준 개발 노력을 지속하고 있으며, 디지털 실크로드를 바탕으로 AI 기반 감시 기술 및 기타 디지털 인프라를 수출 중⁶⁹⁾
 - 중국은 AI 기반 감시 기술 분야에서 세계를 선도하고 있음에도 불구하고, 생성형 AI 시스템 개발에서는 미국에 계속 뒤처지고 있는데, 파운데이션 모델을 훈련하기 위해 사용할 수 있는 데이터를 제한하는 중국의 검열 규정이 원인 중 하나
 - 권위주의 국가들은 혁신과 정치적 통제를 결합한 중국식 접근법을 선호하고 있으며, 중국은 글로벌 AI 이니셔티브와 외교적 노력을 통해 중국을 개발도상국의 AI 리더로서 입지를 굳히고자 함

65) Costigan, Johanna(2024.3.22.), “China’s New Draft AI Law Prioritizes Industry Development”, *Forbes*.

66) AI법 초안 제24조에 따르면 AI 제공자는 “사용 행위가 데이터의 원래 목적이거나 기능과 다르지 않고, 데이터 권리자의 정당한 권익을 부당하게 해치지 않는 경우” 저작권이 있는 데이터를 학습에 사용할 수 있음.

67) '23.11월 베이징 인터넷 법원이 AI 이미지 생성기를 사용하여 만든 저작물에 대해 사용자가 충분히 결과물을 변경했으므로 저작권 보호를 받을 수 있다고 판결한 것과 맥락을 같이 함.

68) Pernot-Leplay, Emmanuel(2024), “The AI Regulation Dilemma: A China, EU & U.S. Comparison”.

69) Bradford, Anu(2023.6.27.), “The Race to Regulate Artificial Intelligence”, *Foreign Affairs*.

3 인공지능(AI) 규제 협력 및 시사점

□ 생성형 AI의 등장 이후 주요국들은 앞다퉀 AI 규제에 나섰으나, AI 기술은 국경을 넘어 전 세계적으로 영향을 미치기 때문에 국제 협력이 필요

- AI 관련 연구 · 혁신 · 표준화를 위해 국가 간 긴밀하게 협력하고 있으며, 대부분의 AI 국가전략은 AI 개발과 국제 협력의 연관성을 높이 평가⁷⁰⁾
 - 국가마다 전략적 우선순위, 법 · 경제 · 인구 · 지리적 특성 등이 다르기 때문에 국제 협력을 통해 동일한 규범을 도출하기 보다는, 공통 원칙과 가치를 바탕으로 AI 솔루션 개발을 위해 경쟁하면서도 공동성과를 추구할 수 있는 '유익한 경쟁'(co-opetition)의 장을 형성하는 것이 국제 협력의 주된 목적

〈참고2. AI 규제 관련 국제협력의 중요성〉

- ▶ AI의 연구 · 개발은 규모의 경제가 중요한 자원집약적 분야
- ▶ 공통 민주적 원칙에 기반한 국제 협력은 책임 있는 AI 개발 및 신뢰 구축에 도움
- ▶ 서로 다른 AI 규제 방식은 AI 기술의 혁신과 확산을 저해
- ▶ 핵심 요소에 대한 AI 규제의 조율은 전문화된 AI 시스템 개발을 촉진
- ▶ 상품 · 데이터 이동에 대한 부당한 제한은 AI의 잠재적 혜택을 크게 감소
- ▶ 글로벌 도전과제를 해결하기 위해 AI 솔루션의 잠재력을 활용하려면 국제 협력 강화 필요
- ▶ 개방성과 민주주의, 표현의 자유와 인권을 보호하기 위해 비슷한 입장을 보유한 국가 간의 협력 중요

자료: Kerry, Cameron F, et al.(2021.10.)

□ AI 리스크 규제를 위한 국제 협력은 추진력을 얻고 있으나, 국가별로 상이한 견해와 목표가 규범 간 상호운용성을 저해

- OECD, G20, G7, UN 등 국제기구에서 글로벌 표준을 개발하려는 노력이 진행 중이나, 각국의 이해관계, 가치관, 역량의 차이로 의미 있는 합의 도출까지 난항이 예상
- '23.11월과 '24.5월 런던과 서울에서 개최된 AI 정상회의에서는 고위급 원칙을 지지하고 대화를 계속 하겠다는 원론적 성명 위주의 선언 발표

70) Kerry, Cameron F, et al.(2021.10.), *Strengthening International Cooperation on AI*, Progress Report, Brookings.

- '23.11월 'AI 안전성 정상회의'(AI Safety Summit)에서는 고성능 범용 AI 모델의 실존적 위험에 대응하기 위해 국제협력 방침을 포함한 '블레츨리 선언'(Bletchley Declaration)을 발표했으며, 美·中 등 주요국이 처음으로 AI가 심각한 피해를 초래할 수 있다는 점에 공감하고 공식적으로 협력을 다짐한 공동선언이란 점에서 상징성이 큼
- '24.5월 'AI 서울 정상회의'에서는 G7·EU·싱가포르·호주, OECD·UN 및 글로벌 기업이 참석하여 각국 AI 안전연구소 간 네트워크를 조성하고 AI 안전성 제고를 위한 상호 협력할 것을 다짐⁷¹⁾
- 미국과 EU는 무역기술위원회(Trade and Technology Council, TTC)를 통해 AI 포함 양국의 기술과 무역정책을 조율하고 있으며, 현재 우리나라 포함 다수 국가가 체결하고 있는 디지털통상 협정에도 AI 관련 협력 조항이 포함
 - 우리나라가 체결한 '한-싱가포르 디지털동반자협정'(Digital Partnership Agreement, DPA)와 '디지털 경제동반자협정'(Digital Economy Partnership Agreement, DEPA) 등 다수 디지털통상 협정에서 AI 기술의 안전하고 윤리적 활용을 위한 협력 규정을 포함
- 비즈니스 효율화 및 생산성 증대, 혁신을 통한 기업의 생존 차원에서 AI 내재화는 선택이 아닌 필수이나, AI 위험성을 심각하게 고려하는 전 세계 정부의 관심 증가로 AI 부작용·사회적 비용 최소화를 위한 입법 노력이 빠르게 확산되고 있어, AI 활용 기업은 데이터 관련 법 위반 리스크 예방을 위한 규제 모니터링 등 선제적 노력 필요
 - AI를 비즈니스에 내재화하기 전에 AI·데이터 관련 규제 요건을 모니터링하고 준수 여부를 판단하기 위해 선제적으로 기업 내부에 자체 AI·데이터 거버넌스를 구축하거나 독립적 감사를 수행해야 할 필요성
 - AI 활용을 통한 비즈니스 가치 극대화 차원에서 데이터 흐름을 최적화하려는 노력을 기울이는 동시에, 데이터 및 개인정보 보호 관련 법률 위반 리스크를 최소화하기 위해 내부 정책과 관행을 확립하는 노력도 중요
 - AI는 데이터 활용을 전제로 하기 때문에 데이터 보안 사고에 근본적으로 취약하며, 이로 인해 비즈니스 전반에 리스크를 내포하고 있어 전사적 차원에서 사전 리스크 관리체계 구축 필요
 - 데이터 활용시 기업 내·외부 데이터 결합·이전을 시도할 경우 법 위반 등 규제 리스크도 AI 활용시 주의 필요
 - 법률이나 행정명령 등 AI에 대한 규율을 시도하고 있는 주요국의 추세를 고려할 때 국내에서도 AI에 대한 법률을 제정하려는 시도가 지속될 것으로 예상⁷²⁾

71) 대한민국 대통령실 보도자료(2024.5.21.), “尹 대통령, AI 서울 정상회의 공동 주재하며 AI 기술 선도국 간 안전·혁신·포용의 3대 글로벌 AI 거버넌스 목표 합의 이끌어”.

72) 이종현(2024.4.23.), “[AI웨이브 2024]⑥ AI 규제 법제화 속도내는 미국·유럽…한국은 지지부진?”, 디지털데일리.

-
- 한국은 기본적으로 AI 활성화 기조를 유지하고 있으나, AI의 바람직한 발전을 위해서는 일정한 규범도 필요하다는 입장
 - AI 윤리, 신뢰성, 투명성, 데이터 보호 등 사회적으로 부정적 영향을 미칠 가능성이 있는 이슈에 대해서 각 부처별로 가이드라인, 해설서 등을 통해 규율할 것으로 예상
 - AI 활성화를 저해하는 수준에서 AI 관련 법령이 마련되지 않도록 업계 의견 적극 개진 필요

[표]

1. AI를 통한 新기업가치 창출 예시 (p.4)
2. 미국과 한국 기업의 AI 기술 도입 성과 (p.5)
3. AI의 긍정적 혜택과 리스크 예시 (p.5)
4. EU, 미국, 중국의 AI 규제 및 영향력 확대 전략 비교 (p.9)
5. EU AI법상 고위험 AI의 주요 준수 의무 (p.12)
6. 범용 AI 모델 및 시스템적 위험이 있는 범용 AI 모델 제공자 의무 구분 (p.13)
7. AI 기업이 약속한 AI 개발원칙 주요 내용 (p.17)
8. 자발적 조치를 약속한 15개 AI 기업 (p.17)
9. AI 행정명령 주요 내용 (p.18)
10. EU 인공지능법과 미국의 AI 행정명령 비교 (p.19)
11. EU와 중국의 AI 규제 방식 비교 (p.21)
12. 중국 생성형 AI 잠정관리 방법 주요 내용 (p.22)

[그림]

1. AI에 대한 분류 및 예시 (p.1)
2. 2024년 가장 많이 사용되는 AI 플랫폼 (p.2)
3. 글로벌 AI 시장 규모 전망 (p.3)
4. 국가별 AI 시장 규모 및 전망 (p.3)
5. 기업들이 고려하는 생성형 AI 관련 리스크 (p.6)
6. '24년 데이터 · 개인정보 보안 · 위험관리 관련 기업 예상 지출 및 성장률 (p.7)
7. AI 규제의 딜레마 (p.8)
8. EU AI법 진행 경과 (p.10)
9. EU AI법 시행 예상 시점 구분 (p.10)
10. EU AI법상 AI 시스템의 위험 구분 (p.11)
11. EU 규제의 브뤼셀 효과 (p.14)

[참고]

1. EU AI법상 금지된 AI 시스템 (p.11)
2. AI 규제 관련 국제협력의 필요성 (p.25)

참고자료

[국내 문헌 · 자료]

- 김익현(2024.5.21.), “EU, 세계 첫 ‘AI법’ 확정…이사회 최종 승인”, ZDNET Korea.
- 김태균(2022), “기업 성과를 높이는 인공지능”, 월간 SW중심사회 2022년 5월호, 소프트웨어정책연구소.
- 대한민국 대통령실 보도자료(2024.5.21), “尹 대통령, AI 서울 정상회의 공동 주재하며 AI 기술 선도국 간 안전 · 혁신 · 포용의 3대 글로벌 AI 거버넌스 목표 합의 이끌어”.
- 딜로이트 AI 연구소(2023), “인공지능(AI) 활용서: 6대 산업별 활용사례”, 2023년 8월 Deloitte Insights, Deloitte.
- 윤성혜(2023.11.10.), “법률 제정속도 느린 중국, AI만은 예외다”, 프레시안.
- 이종현(2024.4.23.), “[AI웨이브 2024]⑥ AI 규제 법제화 속도내는 미국 · 유럽…한국은 지지부진?”, 디지털데일리.
- 최창열(2018.12.6.), “자바스크립트 개발자는 머신러닝을 어떻게 바라봐야 하는가?”, (주)투비소프트, <http://tobetong.com/?p=9393>.
- LG CNS AI빅데이터연구소(2020.3.23.), “딥러닝, 데이터로 세상을 파악하다(1)”, LG CNS, <https://www.lgcns.com/blog/cns-tech/ai-data/8864/>.

[해외 문헌 · 자료]

- Acemoglu, Daron(2021), “Harms of AI”, NBER Working Papers 29247, National Bureau of Economic Research.
- Ahmed, Mashaid(2024.4.9.), “20 Most Used Artificial Intelligence (AI) Platforms of 2024”, Insider Monkey, <https://www.insidermonkey.com/blog/20-most-used-artificial-intelligence-ai-platforms-of-2024-1287025/>.
- Amazon Web Services, “What is Artificial Intelligence (AI)?”, AWS, <https://aws.amazon.com/what-is/artificial-intelligence/>.
- Bradford, Anu(2023.6.27.), “The Race to Regulate Artificial Intelligence”, Foreign Affairs, <https://www.foreignaffairs.com/united-states/race-regulate-artificial-intelligence>.
- Costigan, Johanna(2024.3.22.), “China’s New Draft AI Law Prioritizes Industry Development”, *Forbes*.

- Fauchoux, Vincent and Perissinotto, Juliana(2023.9.11.), "Artificial Intelligence: Joe Biden's executive order vs. the European AI Act – Common Goals?", DDG, <https://www.ddg.fr/actualite/ai-joe-bidens-executive-order-vs-the-european-ai-act-common-goals>.
- Field, Hayden and Leswing, Kif(2024.3.30.), "Generative AI 'FOMO' is driving tech heavyweights to invest billions of dollars in startups", CNBC.
- Gartner(2023.8.17.), "What's New in Artificial Intelligence from the 2023 Gartner Hype Cycle", *Gartner*.
- Gartner(2023.9.28.), "Gartner Forecasts Global Security and Risk Management Spending to Grow 14% in 2024", *Gartner*.
- Grand View Research(2024.), *AI Trust, Risk & Security Management Market Report, 2030*, Grand View Research.
- Kerry, Cameron F, et al.(2021.10.), *Strengthening International Cooperation on AI*, Progress Report, Brookings.
- Madiaga, Tambiana(2024.3.11.), *Artificial intelligence act*, Briefing, European Parliamentary Research Service (EPRS).
- Maples Group(2024.7.2.), "AI Act Set to Come into Force on 1 August 2024", Maples Group.
- Meltzer, Joshua P.(2023.11.1.), "Toward international cooperation on AI governance—the US executive order on AI", Brookings.
- Mok, Charles(2023.11.7.), "Global Competition for AI Regulation, or a Framework for AI Diplomacy?", *The Diplomat*, <https://thediplomat.com/2023/11/global-competition-for-ai-regulation-or-a-framework-for-ai-diplomacy/>.
- Newman, Matthew and Swift, Mike(2024.1.2.), "Comment: Global AI regulation taking diverse approaches, from rules-based to principles-based", MLex.
- Pernot-Leplay, Emmanuel(2024), "The AI Regulation Dilemma: A China, EU & U.S. Comparison". <https://pernot-leplay.com/ai-regulation-china-eu-us-comparison/>.
- Robins-Early, Nick(2024.4.9.), "New bill would force AI companies to reveal use of copyrighted art", *The Guardian*.
- Ryan-Mosley, Tate, et al.(2024.1.5.), *What's next for AI regulation in 2024?*, What's Next in Tech, MIT Technology Review.
- Singla, Alex, et al.(2024), *The state of AI in early 2024: Gen AI adoption spikes and starts to generate value*, McKinsey & Company.

Statista Market Insights(2024.3.), “Artificial Intelligence”, Market Insights, Statista, <https://www.statista.com/outlook/tmo/artificial-intelligence/worldwide>.

Statista Research Department(2024.2.9.), “Generative artificial intelligence (AI) market size worldwide from 2020 to 2030”, Statista.

The White House(2023.7.21.), “FACT SHEET: Biden–Harris Administration Secures Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI”.

The White House(2023.9.12.), “FACT SHEET: Biden–Harris Administration Secures Voluntary Commitments from Eight Additional Artificial Intelligence Companies to Manage the Risks Posed by AI”.

The White House(2023.10.30.), “FACT SHEET: President Biden Issues Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence”.

Thormundsson, Bergur(2024.6.19.), “Artificial intelligence (AI) worldwide – statistics & facts”, Statista.

Thormundsson, Bergur(2024.6.20.), “AI market size worldwide from 2020–2030”, Statista.

Wang, Juan(2024.4.17.), “Chinese experts suggest negative list system be included in proposed AI law”, MLex.

〈2024년 KITA 통상리포트 발간 목록〉

No.	제목	작성자	발간일자
7	인공지능(AI) 규제 주도권 확보를 위한 글로벌 경쟁 및 시사점	정해영	24.07.26
6	미-중 전략경쟁, 레거시 반도체로 전이되나? - 주요국의 레거시 반도체 정책 현황 및 시사점	이유진	24.07.12
5	美 의회 대중국 견제 입법 동향 및 시사점	한아름	24.06.10
4	한국의 FTA 체결현황 점검	통상연구실	24.05.24
3	미국의 경제안보 · 핵심기술 통제 전략 강화 및 시사점	정해영/이정아/한주희/고성은	24.02.27
2	디지털세 주요 내용 및 입법 동향	강금윤	24.02.20
1	글로벌 공급망에 켜진 또 다른 경고등 - 강제노동 규제 동향과 우리 기업 대응방안	한아름	24.02.08



TRADE REPORT

2024



한국무역협회
국제무역통상연구원

서울시 강남구 무역센터 트레이드타워 47층
2024. 07. 26