

해외규제 모니터링 제5호

최신 글로벌 인공지능(AI) 규범 동향 리포트

“본 뉴스레터는 법무부 국제법무지원과 ‘규제 리포트’와 인공지능법 전문가이신 서울대 법학전문대학원 박상철 교수님의 칼럼으로 구성되어 있습니다.”

규제 리포트

최신 글로벌 인공지능(AI) 규범 동향 리포트 - 유럽·미국·중국·일본 등 주요 국가들의 제정 규범을 중심으로 -

법무부 국제법무지원과

1. ‘AI 시대’의 도래에 따른 범국가적 인공지능 규범제정

- “언어를 사용하고, 추상적인 개념을 형성하며, 현재는 오직 인간만이 해결할 수 있는 유형의 문제를 해결할 수 있는 기계.” 1955년 8월 31일, 미국의 컴퓨터공학자 존 매카시(John McCarthy) 교수는 다트머스 대학교에서 주관한 컴퓨터 정보학회에서 위 기계를 “인공지능(Artificial Intelligence, ‘AI’)”이라고 최초로 명명한 후, 2011년 별세하기 전까지 AI가 가져올 미래의 변화에 대해 끊임없이 연구했습니다.
- 그 후 2022년, 미국의 스타트업인 OpenAI社가 개발한 실시간 채팅로봇 모델인 ‘Chat GPT’가 출시되어 전 세계 이용자들이 활용할 수 있게 되면서 AI 기술의 확산이 가져올 현대사회의 변화와 이로 인해 새로이 대두될 사회 문제들에 대한 인식 또한 제고되었습니다.
 - ▶ “5년 동안 업무지원 부서 직원 2만6000명 중 30%가 AI와 자동화로 대체될 것”
- 미국 IBM CEO 아르빈드 크리슈나 인터뷰 中 (2023)
 - ▶ “AI가 스스로 생각해 사람을 해치는 ‘킬러 로봇’으로 변할 수 있다. 그동안의 AI 연구를 후회한다.”
- 캐나다 토론토대 교수 제프리 힌턴 인터뷰 中 (2023)
- 이에, 주요 국가들은 자국민이 AI 기술로부터 얻는 효용성을 사회적·윤리적 문제 없이 올바르게 누릴 수 있도록 범정부 차원의 정책토론회 등을 실시하면서 구체적인 대응 전략을 모색하고 있고 나아가 AI가 가져올 위험성을 미연에 방지하고자 각종 규범들을 제정하기 시작했습니다.
- 이번 리포트에서는 유럽·미국·동아시아 주요 국가들의 AI 규범제정 동향을 간략히 살펴보고 AI 시대를 마주하며 살아갈 우리 기업·국민들에게 미칠 영향들에 관해 알아보도록 하겠습니다.

2. “유럽”의 규제 동향 : ‘종합규범’으로서의 EU「AI법」

- 유럽연합은 '24. 8. 1. 세계 최초로 AI 기술의 정의와 적법 활용 요건 등을 적시하고 관내 기업 및 회원국 국민에게 법적 구속력을 갖춘 EU AI법(이하 'AI법')을 발효했습니다.
- 「AI법」의 가장 중요한 두 가지 특징은 첫째, AI 기술을 그 위험 수준에 따라 등급별로 그 허용의 범위·한계를 명확히 한 점(‘위험수준 등급체계’)이고, 둘째, 본 규범을 위반했을 경우 그 AI 기술의 공급자·수입업자·유통업자·배포업자에게 과징금 등 책임을 묻는 규정(‘규제조항’)을 제정한 점입니다.
 - 우선 「AI법」상 ‘위험수준 등급체계’란, AI 기술이 가져올 위험성을 4개 등급으로 분류한 후, 그 등급별로 규제를 차등 적용하는 체계를 뜻합니다. 특히 1등급에 해당되는 ‘금지’ AI의 경우, 유럽 전역에서 '24년 內 그 수입·수출 일체의 행위가 불가능하도록 조치했습니다.
 - 한편 2등급인 ‘높은 위험’ AI 기술을 유럽 관내에서 활용하기 위해서는, 공급자 등 책임자들이 ①사용자들에게 당해 시스템의 작동 메커니즘 등을 투명하게 설명해야 하고, ②시스템이 중요한 결정을 내릴 경우 인간이 직접 그 결정 과정에 의무적으로 개입하도록 설계했음을 입증해야 합니다.
 - 그 외 3등급인 ‘제한된 위험’ AI 모델의 공급자 등은 사용자에게 ‘본 소프트웨어는 AI를 활용한 것임’을 명시할 의무를 부담하는 반면, 4등급인 ‘최소 위험 AI’는 현행법상으로는 별다른 규제조항은 없는 것으로 분석됩니다.

【EU 「AI법」상 위험수준(단계)별 정리】

위험	단계 명칭	주요 내용
4등급	최소 위험 AI (Minimal Risk)	▶ 정의 : 시민의 권리나 안전에 영향이 미미한 시스템 ▶ 예시 : AI 활용 게임, 스팸메일 필터링 시스템 등
3등급	제한된 위험 AI (Limited Risk)	▶ 정의 : 인간과 상호작용하거나, 콘텐츠를 생성·조작할 시스템 ▶ 예시 : 챗봇, AI 이미지·동영상 제작 프로그램 등
2등급	높은 위험 AI (High Risk)	▶ 정의 : 건강·안전·기본권·환경권에 상당한 위험을 가할 시스템 ▶ 예시 : 판결예측 시스템, 대출 신용평가 시스템, 자율주행차량 등
1등급	금지 AI (Unacceptable)	▶ 정의 : 인간의 기본권·핵심가치에 명백한 위험을 가할 시스템 ▶ 예시 : 실시간 원격 생체인식, 전투로봇, 성별·인종 분류시스템

- 다음으로, 「AI법」상 규제조항들은 크게 ▲과징금 처분, ▲운영중단 조치, ▲데이터 삭제 조치 등으로 정리되어 있습니다.
 - 우선 「AI법」을 위반한 기업의 경우, 최대 3,500만 유로 또는 쉐 세계 연간 매출액의 7% 중 더 큰 액수를 과징금으로 부담하도록 규정했고, 위반의 정도가 심각할 경우 그 회원국 관내 국가의 감독기관 등은 당해 기업에게 ‘AI 시스템 운영 중단’을 명할 수도 있게 규정했습니다.
 - 또한, AI 시스템 운영 과정에서 사용자들의 개인정보 등을 수집·처리할 경우 그 이용자들의 명시적인 동의를 얻어야 하며(이 경우 EU 개인정보법인 ‘GDPR’ 준수 의무 부담), 만일 그 처리과정이 부적법했음이 밝혀질 경우

관내 규제당국은 AI 시스템 운영자에게 '정보삭제명령'을 부과할 수 있도록 규정했습니다.

- 「AI법」은 '25. 2.부터 총칙과 인공지능 활용 금지 행위에 대한 규정들이, '26. 8.부터는 고위험 AI 시스템 규제를 비롯한 대부분의 규정이 효력을 발휘하며, 36개월의 유예기간을 부여받은 ChatGPT, Gemini, Copilot 등 일부 범용 AI 서비스에 대해서도 '27. 8.부터 효력을 발휘할 예정입니다.

3.“미국”의 규제 동향 : 2023년‘백악관 행정명령’ 등

- 미국 백악관은 '23. 10. 30. AI 기술의 투명성 향상과 새로운 기준 마련을 주요 목표로 하는 “AI 행정명령(Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence)”을 발표하였습니다.
- 미국 ‘AI 행정명령’은 ▲안전성 평가 의무, ▲개인정보 보호, ▲형평성·시민권 제고, ▲국제적 AI 리더십 확보 등을 주요 골자로 제정되었으나, EU「AI법」과는 달리 위반 시 제재가 별도로 규정되어 있지는 않습니다.
- 한편, 미국 캘리포니아 의회는 AI 기술이 활용되는 구체적인 상황에서의 개별적 규제(AI 워터마크 표시법, 디지털 복제물 생산 금지법, 정치적 목적 딥페이크 활용 금지법 등)를 담은 주(州) 법안들을 입안하였고 최근 주지사(Governor)의 서명을 거쳐 '24. 9.부로 효력을 발휘하게 되었습니다.
- '20년대 초반까지만 하더라도 미국은 AI 기술 개발의 선두국가로서 규제보다는 양성에 집중했습니다. 그러나 글로벌 플랫폼 서비스들의 무분별한 AI 기술 오남용을 방지하고 국민의 정보 주권을 보호할 필요성이 대두되면서 현재는 주(州) 단위의 보완 입법들이 다수 발의되는 추세입니다.

【미국 2023년 백악관 “AI 행정명령” 개요】

주요 테마	세부 내용
안전성 평가 의무화	▶ AI 시스템의 안정성과 신뢰성을 확인할 수 있는 새로운 표준 개발 ▶ AI 생성 콘텐츠 표시 관련 표준 확립
개인정보 보호	▶ AI 활용 개인정보 보호기술 개발 우선 지원 ▶ 개인정보 상업적 이용 지침 강화
형평성 및 시민권 제고	▶ AI를 활용한 차별 심화 방지 ▶ 양형 등 형사사법체계에서 AI 시스템 도입
국제적 AI 리더십 확보	▶ 국가 데이터 제공 등 AI 학습용 데이터 접근성 향상 ▶ 국제 현안 해결 위한 AI 개발 협력 및 국제 AI 프레임워크 주도

4.“중국·일본”의 규제 동향 : 개별법과 가이드라인

- 중국 정부는 기본법 제정을 통한 종합적·포괄적인 AI 규제 대신, 새로운 AI 현안에 즉각 대응하여 관련 규범을 제정하는 방식을 택하고 있습니다.

【중국 AI 현안별 개별법 개요】

주요 AI 현안	중국의 개별법 제정 현황
윤리규범 ('21)	▶ 「차세대 인공지능 윤리규범」 * 新一代人工智能伦理规范
인터넷 서비스 알고리즘 ('21)	▶ 「인터넷 정보 서비스 알고리즘 추천 관리규정」 * 互联网信息服务算法推荐管理规定
딥 페이크 ('22)	▶ 「인터넷 정보 서비스 심층 종합 관리규정」 * 互联网信息服务深度合成管理规定
생성형 AI ('23)	▶ 「생성형 인공지능 서비스 관리 잠정방법」 * 生成式人工智能服务管理暂行办法 ▶ 「인터넷 안전기준 실천지침-생성형 인공지능 서비스내용 표식방법」 * 网络安全标准实践指南—生成式人工智能服务内容标识方法

- 특히 「차세대 인공지능 윤리규범」은 인공지능의 수 생애주기 관련 윤리 기준을 제시하고, ▲형평성·공정성 증진, ▲개인정보 보호, ▲통제 가능성 및 신뢰 보장, ▲윤리 제고, ▲책임 강화, ▲인간복지 증진의 6가지 기본 윤리 규범을 확립하였습니다.
- 한편 「생성형 인공지능 서비스 관리 잠정방법」의 경우 그 적용 대상과 관련, 국외 생성형 AI 서비스 제공자라 하더라도 중국 국민을 대상으로 서비스를 제공하는 경우 수범 대상에 포함한다는 점은 특히 유의할 부분입니다.
- 한편 일본 정부는 G7 히로시마 정상회의에서 마련한 ‘히로시마 AI 프로세스 종합 정책 프레임워크(Hiroshima AI Process Comprehensive Policy Framework)’의 연장선상에서 ‘기업을 위한 AI 가이드라인 초안’을 발표하는 등 AI가 가져올 새로운 법적 문제를 대비할 각종 지침들을 '24. 1. 제시했습니다.
- ‘기업을 위한 AI 가이드라인 초안’의 경우 ①인간의 존엄성, ②다양성과 포용성, ③지속 가능한 사회의 3가지 기본 가치의 보장을 목표로 하고, 이를 실현하기 위해 ▲인간 중심, ▲안전, ▲공정성, ▲개인정보 보호, ▲보안, ▲투명성, ▲책임, ▲교육·리터러시, ▲공정 경쟁, ▲혁신 등 10가지 원칙을 제언하고 있습니다.
- 한편, 일본은 기존 법률인 개인정보보호법 등을 지속 활용하면서도 비즈니스 분야에서 발생할 신규 법률 분쟁의 경우 중국과 유사하게 AI 현안에 따라 신규 지침을 제작·활용하는 동향을 보이고 있습니다.

5. 맺음말

- 우수한 AI 서비스는 언제든지 웹·앱 등을 통해 세계 각지에서 활용될 수 있는바, AI 서비스 또는 AI 기술을 접목한 제품을 개발·판매 중인 우리 기업들은 AI 규제 동향을 조속히 분석함으로써 글로벌 산업구조의 새로운 패러다임에 빠르게 적응할 수 있을 것입니다.
- 이와 관련하여 ①한국무역협회 편찬 통상리포트 7호 ‘AI 규제 주도권 확보를 위한 글로벌 경쟁 및 시사점’,

②법제연구원 편찬 유럽연합 시법 번역본, ③미국 백악관 ‘AI 행정명령’ 원본을 본 게시물과 함께 별도로 첨부했습니다.

- 법무부는 앞으로도 다양한 각국의 인공지능 규제 등을 적극 모니터링하고 외부 전문가들의 의견을 수렴하여, 해외 시장 진출을 희망하는 우리 기업들이 보다 안정된 환경에서 경영활동을 하는데 도움이 되도록 노력하겠습니다. 🇰🇷

담당자_ 국제법무지원과 사무관 **황현준**, 법무관 **이동건**, 인턴 **정정아**