

글로벌 인공지능(AI) 규제 of 확산: 우리 기업의 새로운 비즈니스 전략이 필요한 결정적 시기

March 27, 2024

오늘날 빠르게 진화하는 기술 환경 속에서 “인공지능(Artificial Intelligence, “AI”)”을 둘러싼 규제 환경도 빠르게 진화하고 있습니다. 지금까지 AI는 개인 정보의 수집/사용/공개와 관련된 범위 내에서 각국의 개인정보 보호 법령의 적용을 받아왔습니다. 그러나 AI 기술 발전 과정에서 예측하지 못한 사회적, 경제적, 윤리적 리스크 발생 가능성이 크게 우려됨에 따라 **AI 신뢰성과 안전성을 제고하기 위한 각국의 입법 시도가 활발해지고** 있습니다.

해외 각국은 각자 상이한 방식으로 AI에 대한 규제를 시도하고 있으나, 데이터 보호 등 법적 의무를 기업에만 한정하기보다 **AI 산업 생태계 참여자들의 전반적인 책임을 강조하고, AI 위험성을 주기적으로 모니터링하도록 함으로써** 식별된 위험에 대한 완화 조치를 시의적절하게 취할 수 있도록 한다는 점에서 유사한 특징을 보입니다.

실제 EU, 미국 등 글로벌 AI 규범을 정립하려는 해외 주요국의 움직임이 크게 증가한 가운데, **AI 원칙 설정 및 표준화를 통해 AI 거버넌스를 구축하기 위한 국제협력도 강조되고** 있습니다. 이러한 상황에서 우리 기업도 이 같은 새로운 글로벌 규제 환경에서 자사의 위험관리 전략을 평가하고 새로운 전략을 마련할 필요성이 대두되었습니다. 이에, 아래에서는 특히 최근의 **EU와 미국 발 규제 동향**에 대해 집중적으로 살펴보고 우리 기업의 대응 전략을 모색해보고자 합니다.

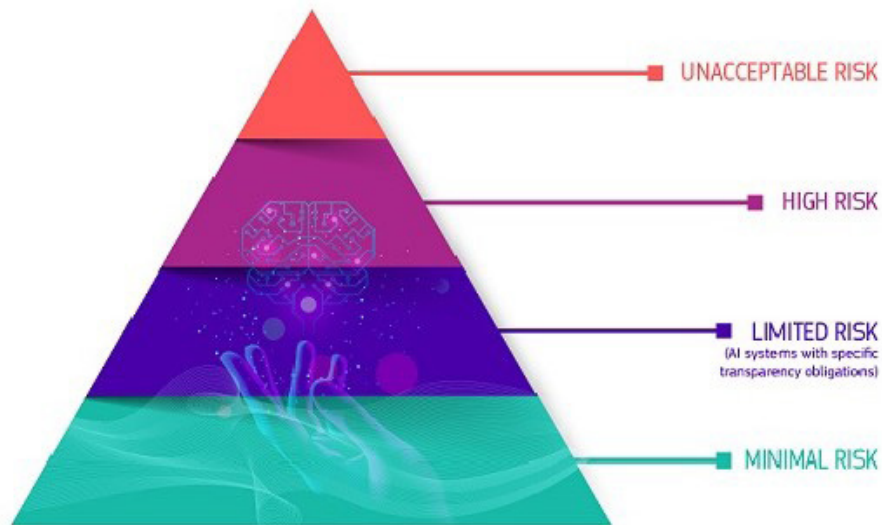
1. EU의 AI 관련 규제 동향: 글로벌 AI 규제의 선두

EU 의회는 2024. 3. 13. 인공지능을 규제하기 위하여 전세계 최초로 AI 법안(Regulation Laying Down Harmonised Rules on Artificial Intelligence, 이하 ‘**EU AI 법**’)을 통과시켰습니다. EU AI법은 현재까지 AI에 관하여 이루어진 가장 포괄적인 규제 프레임워크이며, 2024. 5, 6월경에 발효되어 2024. 하순부터 시행될 것으로 예상됩니다(구체적인 시행 시점은 각 조항별로 차이가 있습니다).

EU AI 법의 주 목적은 강력한 **AI 시스템이 초래할 수 있는 인권, 안전, 윤리 문제를 감소**시키는 데에 있으며, AI 시스템의 배포 및 사용에 대한 위험(Risk) 기반 접근 방식을 채택하여 AI 시스템의 잠재적인 위험과 영향 수준에 따라 **위험도를 4 단계로 나누어 차등 규제**합니다.

EU AI 법은 AI 시스템이 초래하는 위험을 “허용할 수 없는 위험(Unacceptable Risk)”, “높은 위험(High Risk)”,

“제한된 위험(Limited Risk)”, “최소의 위험(Minimal Risk)”의 4등급으로 나누어 각 위험 등급별로 관여자들의 의무를 설정하고 있습니다.



출처: 유럽연합 집행위원회 보고서

(1) 허용할 수 없는 위험 등급 (Unacceptable Risk)

사람의 안전에 “**허용할 수 없는 위험**”을 초래하는 AI 시스템은 금지됩니다. ▲사람이나 특정 취약 집단에 대한 인지 행동 조작, ▲사회적 행동, 사회 경제적 지위 또는 개인적 특성으로 사람들을 분류하는 사회적 점수 평가(Social Scoring), ▲안전인식 등 실시간 원격 생체인식 시스템이 “허용할 수 없는” 위험을 초래하는 것으로 간주되며, EU AI 법은 이러한 AI 시스템은 원칙적으로 금지하고 있습니다.

<금지되는 AI 시스템>

- 연령, 장애 또는 사회경제적 상황과 관련된 취약점을 이용하여 행동을 왜곡 조작
- 민감정보(예: 정치적, 종교적, 철학적 신념, 성적 취향, 인종)를 이용한 생체 인식 분류 시스템
- 사회적 행동 또는 개인적 특성에 기반한 사회적 점수화(Social Scoring)
- 프로파일링 또는 인격적 특성만을 근거로 범죄를 범하는 개인의 위험을 평가
- 인터넷 또는 CCTV 영상에서 얼굴 이미지의 표적 스크랩을 통해 안전 데이터베이스 구성
- 직장 및 교육 기관에서의 감정 추론
- 법 집행을 위한 목적으로 공개적으로 접근 가능한 공간에서 ‘실시간’ 원격 생체 인식 시스템 사용

이러한 AI 시스템은 강간 테러 등 중대범죄 용의자 수색 등 극히 예외적인 경우에 법원의 사전 허가를 받는다는 조건으로 허용됩니다.

금지된 AI 시스템을 사용한 기업에게는 **최대 3,500만 유로 또는 회사의 전 세계 연간 매출액의 최대 7%**에 해당하는 벌금이 부과될 수 있습니다.

(2) 높은 위험 등급 (High Risk)

사람의 안전이나 기본권에 “높은 위험”을 초래하는 AI 시스템은 완전히 금지되지는 않지만, **시장에 출시되기 전에 엄격한 위험 평가 및 관리 의무가 부과됩니다.** 의료 교육 고용 금융 등 필수적인 공공 민간 서비스와 법 집행, 이주 및 국경관리 등 국가의 주요 시스템과 관련된 AI 시스템이 높은 위험을 초래하는 것으로 분류됩니다.

<높은 위험 AI 시스템>

- 시민의 생명과 건강을 위험에 빠뜨릴 수 있는 중요한 인프라(예: 교통)
- 교육에 대한 접근 및 전문적인 교육 또는 직업 훈련(예: 시험 채점)
- 제품의 안전 구성 요소(예: 로봇 보조 수술의 AI 적용);
- 고용, 근로자 관리 및 자영업에 대한 접근(예: 채용 절차를 위한 이력서 분류 소프트웨어)
- 필수 민간 및 공공 서비스(예: 신용 점수, 시민 대출 기회 거부)
- 국민의 기본권을 침해할 수 있는 법 집행(예: 증거의 신뢰성 평가)
- 이민, 망명 및 국경 통제 관리(예: 비자 신청 자동 심사)
- 사법 행정 및 민주적 절차(예: 법원 판결을 검색하기 위한 AI 솔루션)

높은 위험을 초래하는 AI 시스템을 사용하려는 자는 위험 평가 및 관리 시스템을 도입하여 명확하고 적절한 정보 제공, 위험을 최소화하기 위한 적절한 인적 감독 조치, 높은 수준의 견고성, 보안 및 정확성 등을 준수하는 **적절한 보호 장치를 마련**하여야 합니다. 즉, 높은 위험을 초래하는 AI 시스템의 제공자는 해당 AI 시스템이 기본권, 공공 안전 또는 공중 보건에 심각한 위험을 초래하지 않는다는 것을 입증할 수 있어야 합니다.

EU AI법은 범용 AI(General Purpose AI, GPAI) 모델 그 자체를 AI 시스템으로 평가하여 규율하고 있지는 않지만, 그 모델이 AI 시스템과 통합되는 경우에는 아래와 같은 의무를 부담하도록 규정하고 있습니다.

<일반 목적 AI 모델>

- GPAI 모델을 훈련시키는데 사용되는 콘텐츠의 주요 내용 공지 등 투명성 확보를 위한 조치
- GPAI 모델을 훈련시키고 평가하는 프로세스와 그 결과 등에 관한 기술 서류를 AI 시스템 제공자에게 공개할 것
- 지적재산권법 준수를 위한 정책을 수용할 것

2. 미국의 AI 관련 규제 동향: AI 규제 논의의 중심

(1) 연방 정부 차원

미국 법무부 Lisa Monaco 차관은 “**전반적인 컴플라이언스 노력의 일환으로 기업의 AI 관련 위험 관리 능력을 평가하기 시작할 것**”이며, 이를 통해 미국 법무부의 신규 “**Justice AI**” 이니셔티브도 **공개**할 것이라고 발표하면서 “기업의 준법감시인들이 이에 주목해야 할 것”이라고 발언하였습니다.

이에 따라 한국 기업 중 미국에 진출한 기업은 미국 내 자회사뿐만 아니라 본사 차원에서도 기존 컴플라이언스 기준에 따라 AI 관련 위험을 제대로 평가하고 자사에 취약점은 없는지 신속히 검토를 시작할 필요가 있습니다.

한편, 지난 2024. 2. 1.에는 ‘2024 인공지능 환경영향법(Artificial Intelligence Environmental Impacts Act of 2024)(이하 ‘**US AI 환경영향법**’)' 이 발의되었습니다. OpenAI 연구리서치 내용에 따르면 2012년 이후 AI 모델 훈련에 필요한 컴퓨팅 성능이 약 3.4개월마다 2배로 증가함에 따라, 2040년까지 정보통신기술(ICT) 산업 전반에서 발생하는 온실가스 배출량이 전세계 배출량의 14%에 달할 것으로 예상되고 있습니다. 이러한 상황에서, AI가 환경에 미치는 영향을 연구할 필요성 또한 대두되기 시작한 것입니다.

AI가 환경에 미치는 영향을 정확히 파악하는 것을 목적으로 하는 이 법안은 구체적으로 다음과 같은 내용을 포함하고 있습니다.

- AI의 환경영향에 대한 연구: 환경보호국(EPA, Environmental Protection Agency)은 2년 이내 AI의 환경 영향에 대한 연구 수행. 에너지와 물 소비량 연구, 전자폐기물 발생 관련 하드웨어의 수명주기 조사 및 그 외 AI 애플리케이션이 환경에 미칠 수 있는 모든 긍정적, 부정적 영향을 평가
- AI 환경영향에 관한 컨소시엄 소집: AI가 환경에 미치는 영향을 측정하기 위한 지표 개발을 목적으로 하는 이해관계자 컨소시엄 소집
- AI 환경영향의 자발적 보고 시스템 구축: AI 개발, 운영기업이 AI가 환경에 미치는 영향을 자발적으로 보고할 수 있는 시스템 구축
- 의회 보고: 제정일로부터 4년 이내에, EPA, 에너지부, 국립기술표준원(NIST)은 컨소시엄의 조사 결과, 자발적 보고시스템에 대한 설명 및 추가 입법/행정조치에 대한 권고사항을 포함한 보고서를 공동으로 제출

미국에서 AI 환경영향법이 제정 및 시행되면 이에 따라 AI 기술이 환경에 미치는 영향들에 대한 상세한 연구 및 조사가 수행되고 이를 바탕으로 환경에 미치는 부정적 영향을 최소화하기 위한 각종 규제가 도입될 것으로 예상됩니다. 결국 기업들은 이러한 새로운 규제 도입 가능성을 염두에 두면서 AI 기술 개발과정에서 에너지 효율성 제고 등 환경에 미치는 부정적 영향을 줄여 나감과 동시에 관련 보고 시스템을 구축하는데 관심을 기울여야 할 것입니다.

(2) 주 정부 차원

미국 캘리포니아주 의회에서는 현재 약 30개 이상의 AI 관련 법안이 심사 중에 있으며, 담당부처 차원의 규제안도 제안되는 등, 미국 내 그 어느 주정부보다도 활발하게 AI 관련 규제에 대해 논의가 진행되고 있습니다.

특히 캘리포니아주 개인정보보호국(California Privacy Protection Agency, “CPPA”)은 “자동화된 의사결정 기술(Automated Decision-making Technology, “**ADMT**”)"에 관한 규정(이하 ADMT 규정”) 초안을 발표했습니다. CPPA는 ADMT 규정 초안을 바탕으로 올해 내내 최종안의 작성 및 공식적인 채택 절차를 진행해 나갈 예정입니다.

구체적으로, ADMT 규규정은 기업의 ADMT 활용으로부터 소비자(직원 및 입사지원자까지 포함)를 광범위하게 보호하기 위한 것으로 최종 확정되면 기업에게는 아래와 같은 의무가 부과됩니다.

- 소비자에게 기술 사용, 의사결정 과정, 소비자의 거부권(및 정보 접근권)에 대해 사전 고지
- 직원의 보상, 업무 할당 또는 배정, 채용 결정 등 소비자의 삶에 중대한 영향을 미치는 결정에 대해 소비자가 ADMT 사용을 거부할 수 있도록 허용

- 소비자에게 ADMT 사용 정보에 대한 액세스 권한을 제공

그런데 현재 업계에서는 ADMT 규정 초안의 일부 사항이 CPPA에 부여된 직무 권한을 넘어서는 우려를 표명하고 있기 때문에 CPPA로서는 최종안 작성 및 채택 과정에서 우려가 제기된 사항을 수정해 나갈 가능성이 높습니다.

현재로서는 ADMT 규정이 공식적으로 채택된 것은 아니지만 캘리포니아주에 진출하거나 진출하려는 우리 기업/개인 등은 언제라도 위와 같은 규제가 도입될 수 있다는 점을 염두에 두어야 할 것이고 이러한 규제가 다른 주 차원으로 확대될지, 그리고 궁극적으로 연방 차원으로 확대될지에 대해서도 지속적으로 관심을 기울일 필요가 있습니다.

3. 기타 AI를 둘러싼 다양한 이슈

세계 최초의 포괄적인 AI 규제인 EU AI 법은 전 세계적으로 운영되는 미국 기반 비즈니스뿐만 아니라 해외 진출을 계획 중인 우리 기업에게도 상당한 영향을 미칠 것입니다. 기업들은 EU AI 법의 적용 범위를 검토하고 자사의 AI 시스템 사용에 대해 EU AI 법이 적용되는지 판단해야 합니다. 특히 EU의 일반 개인정보보호법(General Data Protection Regulation, “**GDPR**”)이 세계적으로 개인정보 규제를 선도하면서 우리나라 규제에도 많은 영향을 미쳤다는 점을 고려할 때 EU AI 법상 규제 또한 우리나라 규제 형성에 상당한 영향을 미칠 가능성이 높다고 할 것입니다.

특히 국가 및 산업 간 경계를 초월한 AI 도입 · 운영이 활성화됨에 따라 해외 각국은 각국의 개인정보 보호 법령 개정을 통해 글로벌 스탠다드 및 주요국의 법제도와와의 상호운용성을 확보하고자 노력하고 있으며, AI 신뢰성에 대한 상호 검증 및 인증을 위한 국제표준을 마련하려는 시도도 지속하고 있습니다.

이 외에도 AI 기술은 미-중 기술패권 경쟁의 핵심에 해당합니다. 특히 승자독식의 특성이 큰 AI 시장은 초기 시장 선점이 무엇보다도 중요합니다. 이에 따라 그간 미국은 AI 기술개발과 특히 관련성이 높은 반도체에 대한 對 중국 수출통제를 진행하고 있습니다.

더욱이 환경보호와 관련된 규제 또한 중요한 논의의 대상이 되고 있습니다. AI 수요 증가로 에너지 소비가 크게 증가하고 있음에도 전세계적인 탈탄소 정책 흐름은 신규 데이터센터 설립의 걸림돌이 되고 있습니다. AI 혁신을 위한 에너지 집약적인 AI 시스템 개발과, 환경적인 낭비를 막을 수 있는 효율적 AI 사용에 관한 고민도 필요한 시점입니다.

4. 우리 기업의 대응 방안

AI와 관련하여 완전히 새로운 전략 보다는 우선 기존에 마련한 각종 전략에서 AI와 관련된 리스크를 충분히 보완하는 방안을 강구해볼 수 있습니다. 즉, 우리 기업은 기존에 구축한 규제 컴플라이언스 프로그램을 AI와 관련된 리스크에 적용해볼 수 있으며, 기존 프로그램으로 해결되지 않는 리스크는 신속히 식별하여 기존 프로그램을 보완하거나 새로운 프로그램을 구축할 필요가 있습니다. AI 위험과 규제가 모두 증가하는 지금 이야말로 경영자가 리더십을 발휘할 수 있는 절호의 기회입니다.

(1) 잠재적 리스크에 대비하기

투자자/창업자/기업 등은 AI 분야에 대한 투자 및 도입을 위해서는 법률/규제 전문가의 지식 및 경험을 충분히 활용할

필요가 있습니다. 이를 통해 사전에 AI 관련 잠재적 리스크를 파악하고 이에 대한 대비책을 세움으로써 AI 관련 규제 위험을 효과적으로 완화할 수 있습니다.

(2) 기존 컴플라이언스 프로그램을 최대한 활용하기

기업은 사용 중인 AI 시스템을 모니터링하고, 수집된 정보의 투명성을 확인하며, AI 시스템이 편향적으로 작동하거나 금지되는 행위를 하지 않는지 지속적으로 모니터링하여야 합니다.

또한 AI 시스템 운용을 위한 데이터 수집 과정에서 민감정보 등 금지되는 데이터를 수집하여 처리한다면 이는 추후 심각한 민형사상의 문제를 초래할 수 있습니다. 따라서 자체적인 모니터링은 필수적이며 이를 위해서는 이미 내부적으로 마련하여 운용 중인 컴플라이언스 프로그램을 적극 활용하는 것을 고려할 수 있습니다.

(3) AI 거버넌스 구축하기

AI 리스크는 그동안 기업이 경험해 보지 못한 식별되지 않은 위험을 초래할 뿐만 아니라 비즈니스 프로세스 전체에 분산되어 통제하기 어렵다는 특징을 보입니다. 이에 기업들은 AI 기본 원칙 및 내부 정책 가이드라인을 마련하는 것에 더해 비즈니스 프로세스 전반 및 AI 전주기에 걸쳐 AI의 잠재적 위험을 분석·식별하고, 평가·개선하며 위험을 체계적으로 관리할 수 있는 AI 거버넌스 및 데이터 매니지먼트 체계를 구축할 필요가 있습니다.

이를 통해 기업은 데이터 유출 등으로 인한 법 위반 및 비즈니스 리스크 발생 가능성을 최소화함으로써 기업 평판 손상, 고객 및 매출 손실 등 기업이 통제하기 어려운 외부 위험까지 사전에 관리할 수 있습니다.

AI 기술의 보편화는 피할 수 없는 미래인 만큼, 대다수 기업들이 AI 관련 리스크의 영향권에 들 것으로 보이므로 현재 EU 및 미국에서 시도되고 있는 법규정 제정 및 시행 과정을 눈 여겨 보면서 이에 상응한 국내의 AI 관련 규제 논의나 동향도 잘 지켜보고 대응해 나가야 할 것입니다.

법무법인(유) 세종 AI센터, 해외규제팀 및 환경팀은 최근 글로벌 경제 · 무역 · 안보 환경 변화의 핵심을 가져오고 있는 AI를 둘러싼 주요 국가의 규제/정책 및 컴플라이언스 이슈 변화에 신속히 대응할 수 있도록 우리 기업들에게 다양한 관련 정보 및 종합적 대응전략을 선제적으로 제공해드리고 있습니다. 담당 변호사 및 전문가는 아래와 같습니다.

법무법인(유) 세종 AI센터, 해외규제팀 및 환경팀

• 이용우 변호사	T. 02-316-4007	E. ywlee@shinkim.com
• 황성익 변호사	T. 02-316-4417	E. sihwang@shinkim.com
• 장준영 변호사	T. 02-316-4985	E. jyojang@shinkim.com
• 박효민 변호사	T. 02-316-4252	E. hmipark@shinkim.com
• 정홍규 변호사	T. 02-316-4266	E. hkchung@shinkim.com
• 박창준 변호사	T. 02-316-1660	E. cjpark@shinkim.com
• 윤호상 변호사	T. 02-316-2584	E. hsyoon@shinkim.com
• 고현정 변호사	T. 02-316-2811	E. hjko@shinkim.com
• 남수진 변호사	T. 02-316-4449	E. sjnam@shinkim.com
• 이상욱 변호사	T. 02-316-4538	E. swlee@shinkim.com
• 이지은 선임연구원	T. 02-316-1720	E. jeunlee@shinkim.com

SHIN & KIM

법무법인(유) 세종

법무법인(유) 세종 뉴스레터의 게재된 내용 및 의견은 일반적인 정보 제공 목적으로 발행된 것이며, 이에 수록된 내용은 법무법인(유) 세종의 공식적인 견해나 구체적인 사안에 관한 법률의견이 아님을 알려드립니다.

The content and opinions expressed within Shin & Kim LLC's newsletter are provided for general informational purposes only and should not be considered as rendering of legal advice for any specific matter.