

November 14, 2023

인공지능(AI) 규제와 통상분쟁·마찰 가능성

전세계에서 참여한 100여명의 각국 지도자, 기업인, 학자 등은 2023년 11월 1일, 2일 양일동안 영국에서 개최된 제1회 Global AI Safety Summit에서 빠르게 변하고 발전하는 인공지능(Artificial Intelligence, “AI”)이 가진 잠재적 위험과 그 경감 방안을 논의하였습니다. 참가국들은 AI로 인하여 발생할 수 있는 위험을 과학적·객관적으로 평가하고, AI 기술을 안전하게 활용할 수 있는 방안을 모색하기로 하는 선언(‘Bletchley Declaration on AI Safety’)을 채택하고, 영국은 AI의 안전성을 연구하고 정보를 교환하기 위하여 AI Safety Institute를 설립하기로 하였습니다.² 앞으로 AI가 인간 삶의 다양한 측면에 막대한 영향을 미칠 것이 예상되는 상황에서 국제사회와 각국 정부, 기업 등의 지도자가 AI의 잠재적 위험성과 협력방안을 논의하고, 앞으로도 논의를 이어 가기로 하였다는 점에서 큰 의미가 있습니다.

그럼에도 불구하고 각국 정부는 AI를 어떻게 규제할지에 관하여 기본적으로 다른 태도를 취하고 있습니다. 유럽연합은 세계 최초로 AI에 관한 포괄적인 규제법안(EU AI Act)을 논의하고 있고, 한국을 비롯하여 미국, 중국 등 다른 국가들도 AI의 규제방안을 모색하고 있습니다. 그러나 AI 관련 규제의 국제표준을 선점하고 AI 개발경쟁에서 뒤처지지 않으려는 각국 정부의 기본적인 입장 차이로 인하여 AI 규제에 지역별, 국가별로 차이가 발생하고, 통상 분쟁과 마찰이 발생할 가능성을 배제할 수 없습니다. 이하에서 EU와 미국 등이 논의하는 AI 규제방안을 간단히 살펴보고, 향후 통상 분쟁이 나 마찰 가능성을 살펴보겠습니다.

I. 각국의 최근 동향

1. 유럽연합(EU)의 인공지능법안 논의

유럽집행위원회(European Commission)는 2021년 4월 인공지능법안(Proposal for a Regulation on a European Approach for Artificial Intelligence)을 제안하였고, 이후 유럽이사회(European Council)는 2022. 12., 유럽의회(European Parliament)는 2023. 6. 각각 법안에 관한 입장을 채택하였습니다. 현재 유럽집행위원회, 유럽이사회 및 유럽의회 간 3자 회합(trilogue negotiations procedure)이 진행 중입니다. 최근 유럽의회가 채택한 법안(의회안)의 주요 내용은 아래와 같습니다.

1 영국, 미국, 중국, 독일, 프랑스, 일본, 호주, 캐나다, 이탈리아, 싱가포르, 국제연합(UN), 유럽연합(EU) 등 각국 지도자가 참여하였고, Tesla, DeepMind, Meta, OpenAI 등 AI 관련 기업들과 학자들도 참여하였습니다. 우리나라 역시 윤석열 대통령(화상), 이종호 과학기술정보통신부 장관 등이 참석하였습니다.

2 논의를 이어 가기 위하여, 대한민국 정부는 2024년 중순경 화상 미니 정상회의(mini virtual summit)를, 프랑스 정부는 2024년 하반기 차회 대면 정상회의를 각각 개최하기로 하였습니다.

- (위험 구분에 따른 규제) 인공지능의 위험을 그 정도에 따라 ① 허용될 수 없는 위험(an unacceptable risk), ② 고위험(a high risk), ③ 저위험(low or minimal risk)으로 구분.
- (허용될 수 없는 위험) 공공장소에서의 실시간 원격 생체인식 시스템, 민감정보를 활용한 생체인식 분류 시스템, 스크래핑을 통해 안면인식 DB를 구축하는 시스템 등의 사용을 엄격하게 금지.
- (고위험) 자연인의 건강, 안전 또는 기본권에 중대한 위해를 가할 위험이 있는 경우 등에는 고위험 AI로 간주하고, 고위험 AI 시스템 서비스 제공자, 배포자 및 기타 당사자는 △ 위험관리 시스템의 구축 및 유지, △ 품질 기준 충족, △ 기술문서 작성 및 유지, △ 기록관리 △ 투명성 등 높은 수준의 의무 부담.
- (저위험) 고위험 AI에 해당하지 않는 경우에도, 허위임에도 사실로 오인되는 텍스트, 오디오 또는 시각 콘텐츠를 생성하거나 조작하고, 사람들의 동의 없이 말하거나 하지 않은 것을 말하거나 한 것처럼 묘사할 수 있는 AI 시스템(일명 '딥 페이크') 등의 경우에는 고위험 AI와 마찬가지로의 투명성 의무 부여.

2. 미국의 자발적 AI 약속과 행정명령 발표

미국 연방정부의 주도 하에 Amazon, Google, Meta, Microsoft, OpenAI 등 미국의 AI 개발 관련 7개 기업은 2023년 7월경 AI로 인한 위험을 적절하게 관리하겠다는 내용의 '자발적 AI 약속(Voluntary AI Commitments)'에 참여하였습니다. 이는 기업이 AI 상품을 출시하기 전에 적절하게 내·외부적으로 안전성을 점검하고, 공공의 신뢰를 확보하기 위한 다양한 조치를 취하는 것 등을 그 내용으로 합니다.

그리고 미국 바이든 대통령은 2023년 10월 30일 '안심할 수 있고, 안전하며 신뢰할 수 있는 인공지능 개발과 사용에 관한 행정명령(Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence)'에 서명했습니다. 이는 AI를 규제하는 미국 대통령의 첫 행정명령으로, 기술 혁신을 위한 AI의 가능성과 이로 인한 위험의 균형을 적절하게 맞추기 위하여 연방기관들에게 AI를 안전하게 사용하도록 지시하고 민간 부문에서 안전한 AI 사용을 촉진하도록 하고 있습니다. 행정명령은 국가 안보, 경제 안보, 공중 보건 안전에 중대한 위해를 가할 수 있는 AI 모델을 개발하는 기업은 이를 연방정부에 보고하고 안정성 검증 결과를 공유하도록 하고 있습니다. 그리고 행정명령은 모든 연방 정부기관으로 하여금 일정한 기한 내에 AI에 관한 각종 기준, 안전 조치, 보고 요건 등을 시행하고, 연방 정부기관의 AI 관련 활동을 관리 및 조율하기 위하여 백악관에 AI 위원회(White House AI Council)를 설치하도록 하고 있습니다.

행정명령의 8개 목적은 ① AI의 안정성과 보안 확보, ② 사생활 보호, ③ 평등과 인권의 증진, ④ 소비자 보호, ⑤ 근로자 보호, ⑥ 혁신과 경쟁 강화, ⑦ 국제적 리더십 발휘, ⑧ 연방정부의 책임감 있고 효과적인 AI 사용이고, 이를 달성하기 위한 다양한 내용을 포함하고 있습니다.

3. 중국의 생성형 AI 서비스 관리 방안 마련

중국은 2023년 7월 13일 '생성형 AI 서비스 관리 방법'을 채택하였고, 이는 2023년 8월 15일 발효되었습니다. 이 관리 방안은 사회주의 핵심가치관을 지킬 것을 규정하여 중국이 중시하는 가치에 반하거나 사생활, 개인정보 등

개인의 권리를 침해하는 콘텐츠를 규제하는 등 생성형 AI로 생산되는 콘텐츠의 규제를 그 주된 내용으로 하고 있습니다. 또한, 중국은 2023년 10월 18일 Global AI Governance Initiative를 발표하여, AI 규제에 관한 자신의 입장을 제시하였는데 대부분 미국 등 서구의 그것과 상이하지 않지만, AI 기술의 공정한 활용을 강조하여 상대적으로 개발도상국의 이해관계를 중시하는 모습을 보이고 있습니다.

II. 협력과 경쟁, 그리고 통상마찰 가능성

Global AI Safety Summit 등 다양한 국제적 논의에서 AI의 장점을 활용하고 기술 개발을 촉진하면서도 그 잠재적 위험을 적절하게 규제하기 위하여 국제협력이 중요하다는 점에 대해 국제사회와 여러 국가의 지도자들이 동의하는 것으로 보입니다. 실제로도 다양한 방식의 협력방안이 논의되고 있습니다. 예컨대, 미국과 EU는 2023년 10월 30일 무역기술위원회(Trade and Technology Council)에서 공동으로 초안한 '위험 기반 AI 접근 방식을 이해하는 데 필수적인 주요 AI 용어 65개와 이에 대한 EU 및 미국의 해석, EU·미국 공동 정의의 초기 목록 (65 key AI terms essential to understanding risk-based approaches to AI, along with their EU and U.S. interpretations and shared EU-U.S definitions)'에 대한 이해관계자의 의견을 요청하는 등 AI 관련 규제에 있어 협력하는 모습을 보이고 있습니다.

서구 국가와 중국 사이에 안보, 경제 등 여러 분야에 존재하는 긴장관계에도 불구하고 AI를 적절하게 규제하기 위해 중국의 참여와 협력이 필요하다는 인식도 존재합니다. 설사 대부분 국가들이 안전한 AI 기술 개발을 위하여 적절한 규제를 하더라도, 일부 국가가 위험한 AI 기술 개발을 허용한다면, 이러한 규제는 결국 효용이 없게 될 것입니다. 중국 역시 AI 규제에 관한 논의에서 다른 국가와 협력할 의향을 보이고 있었습니다. 그렇기 때문에 최근 영국에서 개최된 Global AI Safety Summit에 중국 역시 초대되었고, 최종적으로 채택된 선언에도 참여하였다고 이해할 수 있습니다. 다가오는 2023년 11월 15일 개최되는 미국 바이든 대통령과 중국 시진핑 주석 간의 정상회담에서도 군사 분야의 AI 사용 규제에 관한 논의가 있을 것으로 예상되고 있습니다.

다른 한편, AI 기술의 개발이나 규제를 둘러싼 국제표준을 선점하고 기술경쟁에서 앞서가기 위한 경쟁도 치열하게 이루어지고 있습니다. 앞서 살펴본 각국의 규제방안 마련이 그러한 경쟁 중 하나라고 할 수 있습니다. 중국의 AI 기술 개발을 막기 위한 미국의 수출통제 역시 대표적인 사례라 할 수 있습니다. 중국이 기술 개발에 있어 군과 민간 간의 융합('military-civil fusion') 전략을 취하자 중국군을 대상으로 하는 제재와 수출통제의 효과를 높이기 위해, 미국 정부는 2022년 10월 고성능 기술이 포함된 그래픽카드 등 AI 반도체를 중국 기업에 수출하지 못하도록 하는 등의 수출규제를 발표하였습니다. 최근에도 미국 정부는 2023년 10월 18일 규제대상인 AI 반도체의 범위를 넓히는 등 수출규제를 강화하는 조치를 발표하였습니다. 이처럼 미국과 중국은 AI에 대한 규제 방안을 마련하기 위하여 협력이 필요하다는 점에 인식을 같이 하면서도, AI 기술 개발에 있어서는 치열하게 경쟁하고 긴장관계를 유지하고 있습니다.

또한, 각국이 취하는 AI 규제의 내용이 상이하면 통상마찰이 발생할 가능성이 있다는 점도 지적되고 있습니다. 예컨대, EU는 AI 기술에 의해 발생할 수 있는 위험을 구분하고 위험별로 규제의 내용을 정하는 일종의 top-down 방식을 채택한 반면, 미국은 연방 정부기관에 대해서는 AI 기술의 혁신과 안전성을 조화하기 위한 방안을 마련하

도록 지침을 마련하면서 동시에 민간 기업의 자발적인 참여를 중시하는 방식을 채택한 것으로 보입니다. 영국은 새로이 광범위한 규제를 도입하기 보다는 기존의 규제 체계 내에서 AI 안정성 확보를 목표로 합니다. 또한, 중국은 AI 기술에 따라 자국이 중시하는 사회주의적 가치에 반하거나 개인의 권리를 침해하는 내용의 콘텐츠 등이 생성되는 것을 막는 등 AI 기술이 생성하는 콘텐츠를 규제하는 방식의 규제 방안을 채택하였습니다. 이처럼 각 국가가 상이한 내용의 규제를 도입한다면, AI 서비스를 제공하거나 AI 관련 제품을 판매하고자 하는 기업 입장에서는 어떤 법적 규제를 따라야 하는지에 혼란이 있을 수 있습니다. 나아가 그러한 규제를 준수하기 위한 자원이 충분하지 않은 기업들은 AI 기술 개발 자체에 어려움을 겪을 가능성도 있습니다. 그러므로 AI 규제에 있어 각국이 공통분모를 찾지 못한다면 이를 둘러싸고 다양한 통상마찰이 발생할 가능성을 배제할 수 없습니다.

특히, 중국은 AI 기술이 공정하게 사용될 수 있어야 한다는 입장을 취하고 있는데, 이는 AI 기술을 개발할 수 있는 자원이 부족한 개발도상국의 이익을 대변하는 것이라 할 수 있습니다. 향후 AI 관련 국제표준 정립 과정에서 AI 기술 개발에 적극 참여하기 어려운 개발도상국의 입장과 미국, 영국 등 AI 기술 개발을 적극 추진하는 국가의 입장이 충돌할 가능성을 배제할 수 없습니다.

III. TBT 협정 등 통상규범 위반 가능성

AI 기술의 향후 발전과 응용에 따라 현재 각국이 취하는 규제로 인하여 WTO 협정 등 통상규범 위반 상황이 발생할 가능성도 있습니다. 예컨대, AI에 규제가 AI가 탑재된 상품에 적용될 경우, WTO 무역에 대한 기술장벽에 관한 협정(Agreement on Technical Barriers to Trade, "TBT 협정")의 준수 여부가 문제될 수 있습니다. TBT 협정은 기술규정에 대하여 다른 WTO 회원국의 상품이 자기나라의 동종상품 및 그 밖의 국가를 원산지로서 하는 동종상품보다 불리한 취급을 받지 아니하도록 보장하고(제2.1조), 기술규정이 국제무역에 불필요한 장애를 초래할 목적으로 또는 그러한 효과를 갖도록 해서는 안 된다는 최소무역제한 원칙을 두고 있는 등(제2.2조) 공정한 무역을 위한 여러가지 규범을 정하고 있습니다. 따라서 앞으로 미국, EU, 중국 등 WTO 회원국의 AI 규제로 인하여 TBT 협정상 의무 위반 상황이 발생할 가능성도 배제할 수 없습니다.

또한, TBT 협정은 관련 국제표준이 존재하거나 그 완성이 임박한 경우, WTO 회원국이 원칙적으로 이러한 국제표준을 자기나라 기술규정의 기초로 사용하도록 하고 있습니다(제2.4조). 따라서 향후 AI에 대한 국제표준이 마련된다면, WTO 회원국은 국내법이 해당 국제표준과 조화될 수 있도록 노력할 필요가 있고, 만약 이를 위반하면 TBT 협정 등 통상규범에 위배될 가능성이 있습니다.

IV. 시사점

스마트폰, 스마트 가전, 자율주행차량 등 제품을 제조하는 국내 기업들, 미국, EU, 중국 등 소비자들이 사용하는 글로벌 어플리케이션을 개발하는 국내 기업들이 자사의 제품 또는 서비스에 AI 시스템을 활용할 가능성이 높다는 점을 고려할 때, 우리 기업 또한 앞서 살펴본 미국, EU, 중국 등 각국 규제의 영향을 받을 수밖에 없습니다. 그러므로 우리나라 정부와 기업들은 향후 다른 국가나 지역에서 도입될 가능성이 있는 규제에 대비하고, 국제 논의에 적극

참여하여 필요하다면 통상규범적 시각에서 시정을 요청하는 등 국제표준이나 규제 논의에서 우리 기업과 산업의 이익이 반영될 수 있도록 할 필요가 있습니다.

우리나라 국회에서 2020. 7. 13. 「인공지능 연구개발 및 산업진흥, 윤리적 책임 등에 관한 법률안」이 발의된 이래 지금까지 인공지능 관련 법안이 최소 13건이 발의되었습니다. 또한, 과학기술방송정보통신위원회는 2023. 2. 14. 「인공지능 산업 육성 및 신뢰 기반 조성에 관한 법률안」 처리를 위원회 대안으로 합의하는 등 AI 관련 입법적 노력이 계속되고 있습니다. 우리나라의 AI 관련 입법과 규제 마련 과정에서는 다른 국가와 국제사회의 논의 동향을 지속적으로 주시하면서 우리나라 실정에 부합하는 입법을 하여야 우리 기업과 산업에 발생할 수 있는 어려움을 낮출 수 있을 것입니다.

관련 구성원

한창완

변호사

T 02.3404.1076

E changwan.han@bkl.co.kr

김지은

변호사

T 02.3404.1265

E jieun.kim@bkl.co.kr