

2024년 4월 15일

## Preview

지난달 유럽의회가 세계 최초로 인공지능에 관한  
일반법인 AI ACT를 통과시켰습니다.

이 법률은 추후 각국에 도입될 인공지능 법률의 기준이 될 것이 예상되므로  
그 내용을 숙지할 필요가 있습니다.

이번 뉴스레터에서는 유럽 AI법의 주요 내용을 살펴봅니다.

## 리걸 이슈 Legal Issues

유럽연합의 인공지능 법률(EU AI ACT) 소개

## 민후 소식 Minwho News

방송프로그램 제작사의 저작권침해 혐의 형사고소 사건에서 승소

개인정보 보호법 시행령 개정에 따른 CPO 지정 관련 자문

## 리걸이슈

## 유럽연합의 인공지능 법률(EU AI ACT) 소개

2024. 3. 13. 유럽의회가 AI ACT(이하 유럽 AI법)를 통과시키면서 세계 최초로 인공지능에 관한 일반법이 통과되었다. 위 법률은 추후 각국에 도입될 인공지능에 관한 일반법의 기준이 될 것으로 예상 되므로, AI 비즈니스 관련 업종에서는 그 내용을 반드시 숙지할 필요가 있다.

유럽 AI법의 주된 내용은 AI의 위험성 정도에 따라 규제의 내용을 달리 하는 위험 기반 규제(risk-based approach)이므로 이를 살피고, 또한 챗GPT와 같은 범용AI에 대한 규제가 도입된 점도 주목할만한 부분이므로 그 내용을 살펴본다.

(※아래 내용은 필자의 임의적 번역에 따른 것으로, 정확한 전달을 위해 주요 내용은 가급적 원문을 병기하였다.)



원준성 파트너 변호사

T. 02-538-3427

E. wonjs@minwho.kr

## [위험 기반 규제]

유럽 AI법은 AI를 위험성의 정도에 따라 ▲ 수인불가 위험(Unacceptable Risk), ▲ 고위험(High-risk), ▲ 제한적 위험(Limited risk), ▲ 최소 위험(Low or minimal risk)으로 구별하고, 이에 상응하는 각기 다른 내용의 규제를 도입하였다.

▲ 수인불가 위험 AI 시스템은 8가지이다(예외로 허용되는 경우는 생략하였음). ① 잠재적, 조작적 또는 기만적 기술을 사용하여 행동을 왜곡하고 정보에 입각한 의사 결정을 저해하여 심각한 피해를 초래하는 시스템(deploying subliminal, manipulative, or deceptive techniques to distort behaviour and impair informed decision-making, causing significant harm.), ② 연령, 장애 또는 사회 경제적 상황과 관련된 취약성을 악용하여 행동을 왜곡하여 심각한 피해를 초래하는 시스템(exploiting vulnerabilities related to age, disability, or socio-economic circumstances to distort behaviour, causing significant harm.), ③ 생체정보의 분류로써 인종, 정치적 견해 성적 취향 등 민감한 속성 추론하는 생체정보 분류 시스템(biometric categorisation systems inferring sensitive attributes), ④ 사회적 행동 또는 개인적 특성에 따라 사람을 평가하거나 분류하여 불리한 대우를 유발하도록 사회점수를 매기는 시스템(social scoring, i.e., evaluating or classifying individuals or groups based on social behaviour or personal traits, causing detrimental or unfavourable treatment of those people.), ⑤ 프로파일링 또는 성격 특성만을 기반으로 개인의 범죄행위 가능성을 평가하는 시스템(assessing the risk of an individual committing criminal offenses solely based on profiling or personality traits.),

⑥ 인터넷이나 CCTV 영상에서 불특정 안면 이미지를 스크래핑하여 안면인식 데이터베이스를 컴파일하는 시스템(compiling facial recognition databases by untargeted scraping of facial images from the internet or CCTV footage.), ⑦ 사람의 감정을 추론하기 위한 시스템(inferring emotions in workplaces or educational institutions), ⑧ 법 집행을 위한 공공장소 실시간 원격 생체 인식(RBI) 시스템('real-time' remote biometric identification (RBI) in publicly accessible spaces for law enforcement)

이러한 수인불가 위험은 절대적으로 금지된다(Unacceptable risk is prohibited.)

▲ 고위험(High-risk) AI는 ① 의료기기, 교통수단, 장난감 등과 같은 제품 안전성과 관련, 즉 안전 구성요소나 제품으로 사용되는 것으로서 Annex II 기재 EU 법률이 적용되거나 혹은 그 법률에 따라 제3자의 적합성 평가 대상이 되는 경우와(used as a safety component or a product covered by EU laws in Annex II AND required to undergo a third-party conformity assessment under those Annex II laws), ② 생체인식(Biometrics), 핵심 기간 산업(Critical infrastructure), 교육 및 직업 훈련(Education and vocational training), 고용 및 인사관리(Employment, workers management and access to self-employment), 필수 공공·사적 서비스 및 그 혜택에 대한 접근 및 이용(Access to and enjoyment of essential private services and essential public services and benefits), 법 집행(Law enforcement), 이민, 망명과 출입국관리(Migration, asylum and border control management), 사법행정과 민주적 절차(Administration of justice and democratic processes) 등 ANNEX III 기재의 8가지 분야의 두 범주로 나뉘어 규정된다. ③ 그리고 경제 상황, 건강, 선호도, 관심사, 위치 등의 개인정보를 자동으로 처리하여 사람의 삶의 다양한 측면을 평가하는 개인 프로파일링의 경우(if it profiles individuals) 항상 고위험으로 간주된다.



고위험 AI의 제공자(Provider)는, AI 라이프사이클 전반에 걸친 위험관리 시스템 구축(risk management system), 오류 및 목적 확인을 위한 데이터 거버넌스 수행(data governance), 규정 준수 확인을 위한 기술문서의 작성 및 제출(technical documentation), 규정 준수 확인을 위한

기술문서의 작성 및 제출(technical documentation), 이벤트 자동 기록보관(record-keeping), 사용지침 제공(instructions for use), 배포자의 인적 감독(human oversight), 적절한 수준의 정확성, 견고성, 보안성(appropriate levels of accuracy, robustness, and cybersecurity), 품질관리 시스템의 구축(quality management system) 등의 의무를 준수해야 한다.

▲ 제한적 위험(Limited risk) AI는 투명성 의무(Transparency obligation)를 준수하면 되므로, 고위험 군에 비하여 상대적으로 적은 규제를 받는다.

합성된 콘텐츠를 생성하는 제공자(generating synthetic audio, image, video or text content), 감정인식 혹은 생체분류 시스템의 배포자(Deployers of an emotion recognition system or a biometric categorisation system), 딥페이크 콘텐츠 배포자(Deployers of an AI system that generates or manipulates image, audio or video content constituting a deep fake), 공익문제에 관한 정보의 공공제공을 위한 텍스트 생성 시스템의 배포자(Deployers of an AI system that generates or manipulates text which is published with the purpose of informing the public on matters of public interest)는 그것이 AI 시스템에 의한 것임을 표시, 고지, 공개하여야 합니다. 한편 범용AI(GPAI)도 투명성 의무의 대상이다.

▲ 위에 속하지 않는 최소 위험(Low or minimal risk) AI에는 특별한 규제를 두지 않았다.

### [범용AI(GPAI)에 대한 규제]

유럽 AI법은 범용AI(General purpose AI - GPAI)에 대한 별도의 규제도 마련하였다.

범용AI란 대량의 데이터를 자기지도학습으로 훈련되고(trained with a large amount of data using self-supervision at scale), 상당한 범용성을 가지며(displays significant generality), 시장에 어떻게 배치되었는지와 관계없이 다양한 작업들을 능숙하게 수행할 있으며(capable to competently perform a wide range of distinct tasks regardless of the way the model is placed on the market), 하위 앱이나 시스템에 통합될 수 있는(can be integrated into a variety of downstream systems or applications) AI 모델로 정의되는데, 잘 알려진 챗GPT가 대표적인 예라 할 수 있다.



범용 AI 시스템(GPAI System)은 범용 AI모델을 기반으로 하는 AI 시스템을 뜻하며, 직접 사용되는 경우는 물론 다른 AI 시스템과 통합되어 사용되는 경우도 포함된다.

한편 범용AI 모델에 시스템적 위험(systemic risk)이 있는 경우에는 강화된 규제가 적용된다. GPAI 모델의 학습훈련에 사용된 누적 컴퓨팅의 양이  $10^{25}$  FLOPS(위원회에 의해서 수정 가능) 이상인 경우 시스템적이라 간주된다. 또한 위원회는 GPAI 모델이 높은 영향력(high impact capabilities)을 가지고 있다고 판단되면 시스템적이라 결정할 수 있다.

시스템적 위험이 있는 GPAI 모델의 제공자는, 적대적 테스트를 포함하는 모델 평가를 수행해야 하고(Perform model evaluations, including conducting and documenting adversarial testing to identify and mitigate systemic risk), 원인을 포함하여 가능한 시스템 위험을 평가하고 완화해야 하며(Assess and mitigate possible systemic risks, including their sources), 심각한 사건과 이에 대한 조치를 기록하여 관련 국가 기관에 보고해야 하며(Track, document and report serious incidents and possible corrective measures to the AI Office and relevant national competent authorities without undue delay), 적절한 수준의 사이버 보안을 확보해야 한다(Ensure an adequate level of cybersecurity protection).

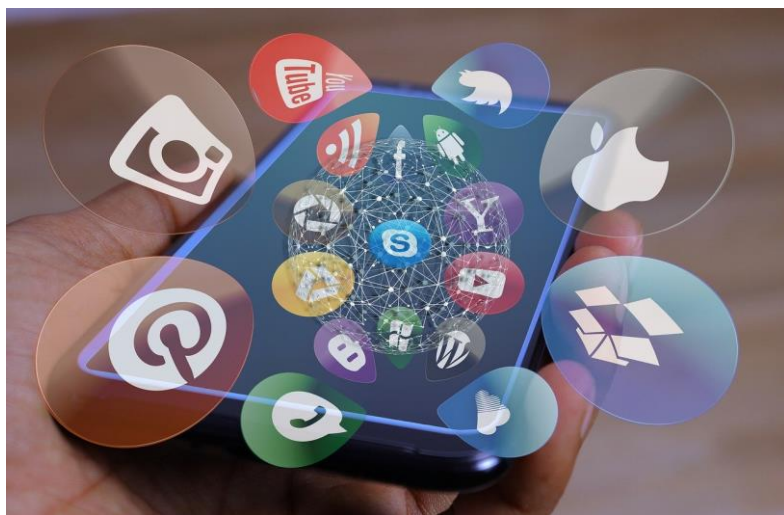
AI를 통한 비즈니스를 계획하는 사업자는 자신의 사업모델이 어떤 종류에 속하여 어떤 규제를 받는 대상인지를 꼼꼼히 확인하여 불측의 피해를 입는 일이 없도록 해야 할 것이다.

## 민후 소식

## 방송프로그램 제작사의 저작권침해 혐의 형사고소 사건에서 승소

법무법인 민후는 저작권침해 행위로 피해를 입은 피해자를 대리한 저작권침해 형사고소 사건에서 유죄판결을 이끌어 승소하였습니다.

의뢰인(피해자)은 가상캐릭터를 활용한 시청각 콘텐츠를 제작하고 온라인 플랫폼에 방송하고, 굿즈 판매 등 관련 사업을 영위 하는 회사입니다.



의뢰인은 타인에 의하여 의뢰인의 가상캐릭터 등 저작물이 무단으로 SNS에 게시되는 일이 발생하자 본 법인에 대응을 요청하였습니다.

본 법인은 가상캐릭터와 이를 활용한 시청각 콘텐츠 역시 저작권법상 저작물에 해당할 수 있으며,

이를 타인이 무단으로 게시하는 행위는 저작권법위반에 해당한다는 점을 이 분야에 대한 전문성을 바탕으로 상세히 설명하며 피고인을 고소하였습니다.

의뢰인은 타인에 의하여 의뢰인의 가상캐릭터 등 저작물이 무단으로 SNS에 게시되는 일이 발생하자 본 법인에 대응을 요청하였습니다.

본 법인은 가상캐릭터와 이를 활용한 시청각 콘텐츠 역시 저작권법상 저작물에 해당할 수 있으며, 이를 타인이 무단으로 게시하는 행위는 저작권법위반에 해당한다는 점을 이 분야에 대한 전문성을 바탕으로 상세히 설명하며 피고인을 고소하였습니다.

이에 수사기관은 본 법인의 주장을 받아들여 피고인을 기소하였고, 법원 역시 본 법인의 주장을 받아들여 피고인의 저작권법위반죄를 인정하는 유죄 판결을 내렸습니다.

## 민후 소식

## 개인정보 보호법 시행령 개정에 따른 CPO 지정 관련 자문

2024. 3. 12.경 개인정보 보호법 시행령 중 개인정보보호책임자(CPO)의 자격에 관한 개정사항이 있었고, 의뢰인은 개정법령에 따른 CPO의 자격기준에 대하여 문의하였습니다.

본 법인은 개인정보 보호법 시행령 [별표1]의 개인정보보호경력, 정보보호 경력 등의 의미를 검토 하고, 이에 해당될 수 있는 것과 그렇지 않은 것을 분석해 구별하여 자문하였습니다.

또한 본 법인은 개인정보 보호법의 개인정보 보호책임자(CPO)와 정보통신망법의 정보보호 최고책임자(CISO)에 관한 관련법령을 분석하여 각 취지와 자격요건의 공통점과 차이점을 안내하여 CPO 및 CISO 지정에 관한 기준을 자문하였습니다.

최신 개정 내용을 바탕으로 정확한 기준을 제시한 본 법인의 자문으로 의뢰인은 CPO 및 CISO 지정에 관한 법적 리스크를 줄일 수 있었습니다.

본 뉴스레터의 내용 또는 기타 법률 문의가 필요하신 경우,  
법무법인 민후로 연락주시면 담당 변호사님의 답변을 받으실 수 있습니다.