

유럽연합 인공지능법과 인공지능 규제

양 천 수 교수

영남대학교 법학전문대학원 교수

I Focus

I. 서론	131
II. 인공지능 규제 논의	133
III. 유럽연합 인공지능법의 배경	139
IV. 유럽연합 인공지능법 분석	148
V. 시사점 - 결론을 대신하여	164

제24-1호

요약문

ChatGPT 열풍이 보여준 것처럼 이제 인공지능은 우리 삶의 일부가 되고 있다. 비약적으로 발전한 인공지능의 성능으로 우리 삶이 점점 더 윤택해지는 것처럼 보인다. 그러나 사회의 많은 현상이 그렇듯이 인공지능 역시 빛과 어둠, 긍정적인 면과 부정적인 면을 동시에 가진다. 인공지능으로 인해 새로운 사회적 위험이 출현한다. 이에 따라 인공지능이 야기하는 위험 문제에 대처하기 위해 어떻게 인공지능을 규제할 것인지가 법학을 포함한 규범학 전반에 중요한 도전이 된다. 초기에는 인공지능 기술이 여전히 발전 과정에 있다는 점을 고려하여 법이 아닌 윤리로 인공지능을 규율할 필요가 있다는 목소리가 공감대를 형성하였다. 이에 따라 인공지능 윤리를 어떻게 설정할 것인지에 많은 논의가 진행되었다. 하지만 2022년 11월 ChatGPT가 출현하면서 상황이 변하게 되었다. ChatGPT로 대변되는 이른바 파운데이션 모델(foundation model)이 실현되면서 이제 윤리가 아닌 법으로 인공지능을 규제할 필요가 있다는 주장이 힘을 얻게 되었다.

이러한 상황에서 유럽연합은 2021년 4월 인공지능을 법으로 규제하는 규제법안을 선구적으로 내놓았고 여러 논의 끝에 2023년 12월 이를 제정하기로 합의하기에 이르렀다. 인공지능을 규제하는 포괄적인 규제법이 조만간 실정화될 예정이다. 최근 유럽연합에서 이루어지는 입법이 우리 법체계에 많은 영향을 미친다는 점을 고려할 때 유럽연합 인공지능법은 우리에게 중요한 의미가 있다. 이에 이 글은 유럽연합 인공지능법을 분석하고 이를 통해 몇 가지 시사점을 도출한다.

이에 따르면 유럽연합 인공지능법에서 다음과 같은 특징을 발견할 수 있다. 첫째, 위험 기반 접근법을 채택한다. 둘째, 위험에 따라 인공지능을 유형화한다. 셋째, 인공지능 유형에 맞추어 차별화된 규제를 적용한다. 넷째, 파운데이션 모델을 규제한다. 다섯째, 규제 수단으로서 영향평가와 투명성을 강조한다. 여섯째, 인공지능 거버넌스를 다원적으로 구축한다. 일곱째, 인공지능을 시장에 출시한 후에도 지속적으로 모니터링하고 감시한다. 유럽연합 인공지능법이 지닌 이러한 특징은 인공지능에 대한 규제를 설계할 때 유익한 참고가 된다.

I

서론

알파고 충격과 ChatGPT 열풍이 보여준 것처럼 이제 인공지능은 우리 삶의 일부가 되고 있다. ChatGPT가 예증하듯이 인공지능을 활용함으로써 우리 인간 존재의 역량은 비약적으로 향상되고 있다.¹⁾ 사회 전체적으로도 새로운 공리(utility)가 증진한다. 인공지능은 우리에게 더 나은 미래를 약속하는 것처럼 보인다.

그러나 사회의 많은 현상이 그렇듯이 인공지능 역시 빛과 어둠, 긍정적인 면과 부정적인 면을 동시에 가진다. 인공지능으로 인해 새로운 사회적 위험이 출현한다. 이에 따라 사회를 구성하는 우리의 권리나 이익이 침해될 위험이 증가한다. 예를 들어 인공지능이 인간이 수행하던 의사 결정 영역에 활용되면서 인공지능에 의한 차별 위험이 등장한다. 인공지능의 잘못된 투자로 주식 시장 전체에 장애가 발생하기도 한다.²⁾ 인공지능이 생성한 가짜뉴스가 정치적 공론장을 혼란에 빠트리기도 한다. 그리고 인공지능의 눈부신 역량으로 인간 존재가 담당하던 일자리를 인공지능이 대체한다. 인간을 위해 개발된 인공지능이 마치 ‘주인과 노예의 변증법’처럼 인간을 지배하고 소외시키는 문제가 출현하고 있는 것이다.

이러한 문제, 즉 인공지능이 야기하는 위험 문제에 대처하기 위해 어떻게 인공지능을 규제할 것인지가 법학을 포함한 규범학 전반에 중요한 도전이 된다.³⁾ 이에 관해서는 이미 알파고 충격 당시부터 논의가 진행되었다. 그렇지만 당시에는 인공지능 기술이 여전히 발전 과정에 있다는 점을 고려하여 경성 규범이 아닌 연성 규범, 달리 말해 법이 아닌 윤리로 인공지능을 규율할 필요가 있다는 목소리가 공감대를 형성하였다. 이에 따라 인공지능 윤리를 어떻게 설정할 것인지에 많은 논의가 진행되었다.⁴⁾ 하지만 2022년 11월 ChatGPT가 출현하면서 상황이 변하게 되었다. ChatGPT로 대변되는 이른바 파운데이션 모델(foundation model)이 실현되면서 이제 인공지능은 그들만의 것이 아닌 우리 모두의 것으로 자리매김하였고 이로 인해

1) ChatGPT 열풍에 관해서는 이시한, 『GPT 제너레이션: 챗GPT가 바꿀 우리 인류의 미래』(북로망스, 2023) 참고.

2) 이에 관해서는 크리스토퍼 스타이너, 박지유 (옮김), 『알고리즘으로 세상을 지배하라』(에이콘, 2016), 7쪽 아래 참고.

3) 이 문제에 관해서는 양천수, 『인공지능 혁명과 법』(박영사, 2021) 참고.

4) 예를 들어 양천수, “인공지능과 윤리: 법철학의 관점에서”, 『법학논총』(조선대) 제27집 제1호(2020. 4), 73-114쪽 참고.

인공지능의 위험 역시 직접적으로 우리를 위협하는 것으로 각인되었다. 이처럼 상황이 변하면서 이제 윤리가 아닌 법으로 인공지능을 규제할 필요가 있다는 주장이 힘을 얻게 되었다. 매년 인공지능 기술이 예상을 뛰어넘어 급격하게 발전하면서 인공지능을 향한 두려움이 막연한 것에서 지금 여기에 다가온 현실적인 것으로 우리를 위협하게 된 것이다.

이러한 상황에서 유럽연합은 2021년 4월 인공지능을 법으로 규제하는 규제법안을 선구적으로 내놓았고 여러 논의 끝에 2023년 12월 이를 제정하기로 합의하기에 이르렀다.⁵⁾ 인공지능을 규제하는 포괄적인 규제법이 조만간 실정화될 예정이다. 최근 유럽연합에서 이루어지는 입법이 우리 법체계에 많은 영향을 미친다는 점을 고려할 때 유럽연합 인공지능법은 우리에게 중요한 의미가 있다. 이에 이 글은 유럽연합 인공지능법을 분석하고 이를 통해 몇 가지 시사점을 도출하는 것을 목표로 한다.⁶⁾

5) 이에 관해서는 (<https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>) 참고(검색 일자: 2024. 1. 29.).

6) 유럽연합 인공지능법은 현재 공식적으로 제정된 것은 아니기에(2024. 2. 9. 기준) 이 글에서는 이를 '유럽연합 인공지능법안'과 '유럽연합 인공지능법'을 맥락에 따라 혼용하기로 한다.

II

인공지능 규제 논의

1. 인공지능에 관한 두 가지 시각

인공지능에 어떻게 대처할 것인가에는 서로 다른 두 가지 시각을 언급할 수 있다. 이는 인공지능이 지닌 양면성, 즉 효용과 위험 가운데 어느 쪽을 중시하는가와 관련을 맺는다. 지원(support)을 지향하는 시각과 규제(regulation)를 지향하는 시각이 그것이다.

전자의 시각은 인공지능이 우리에게 가져다줄 효용, 사회 전체의 공리를 중시한다. 인공지능이 우리 사회 전체의 플랫폼을 발전적인 방향으로 바꿀 것으로 파악한다.⁷⁾ 인공지능을 사회의 새로운 성장 동력으로 이해한다. 이러한 시각을 취하면 인공지능은 규제 대상이라기보다는 지원 및 진흥의 대상으로 자리매김한다. 따라서 이에 바탕을 둔 법은 지원법 또는 진흥법의 성질을 가진다.

이에 반해 후자의 시각은 인공지능이 지닌 위험에 초점을 맞춘다. 인공지능이 우리 개인과 사회 전체에 어떤 위험과 해악(harm)을 유발할 것인지를 중시한다. 이에 따라 인공지능은 규제 대상으로 각인된다. 인공지능이 지닌 위험을 어떻게 안전하게 관리할 것인가가 관심 대상이 된다. 이러한 맥락에서 인공지능을 규율하는 법은 규제법의 성격을 지닌다.

ChatGPT에 관한 사회적 관심과 열풍이 끓어오르던 출시 초반만 하더라도 인공지능에 관해 전자의 시각이 우세하였다. 이에 더해 인공지능 기술은 한참 개발 중에 있다는 점이 강조되었다. 기술 경쟁의 중요성이 강조되었다. 이러한 배경에서 인공지능에 관한 규제가 필요할 때도 경성 규범인 법보다는 연성 규범인 윤리가 규제 매체로 선호되었다. 한편으로는 인공지능의 효용을 극대화하면서도 다른 한편으로는 '신뢰할 수 있

7) 이러한 시각을 보여주는 예로는 이즈미다 료스케, 이수형 (옮김), 『구글은 왜 자동차를 만드는가: 구글 vs 도요타, 자동차의 미래를 선점하기 위한 전쟁의 시작』(미래의창, 2015) 참고.

는 인공지능’(trustworthy AI)을 구현하기 위해 전 세계적으로 다양한 인공지능 윤리가 선포되었다.⁸⁾

하지만 ChatGPT가 보여준 역량으로 전 세계가 이에 열광 및 충격을 받으면서 인공지능이 우리의 예상보다 훨씬 빠르게 발전하고 있음이 확인되었다. 이와 더불어 ChatGPT가 야기하는 이른바 환각(hallucination), 개인정보 및 지식재산권 침해 등이 사회적으로 문제가 되면서 이제 인공지능을 법으로 규제할 필요가 있다는 목소리가 전 세계적으로 높아졌다. 유럽연합 인공지능법은 바로 이러한 배경에서 제정에 이르게 된 것이다.

2. 인공지능의 위험

가. 위험 개념

인공지능을 법으로 규제하고자 하는 이유는 인공지능이 위험(risk)을 지니기 때문이다. 여기서 위험이란 무엇인지 간략하게 검토한다.⁹⁾ 위험은 그 무엇인가를 침해할 가능성을 뜻한다. 여기서 무엇을 침해에 대한 기준으로 삼을 것인지에 관해서는 여러 논의가 있지만 보통 이익(interest)이 기준이 된다. 이러한 이익은 ‘개인/사회’라는 구별을 활용하면 크게 두 가지로 나눌 수 있다. 개인에게 귀속되는 이익과 사회 전체에 귀속되는 이익이 그것이다. 이 가운데 전자의 경우로 대표적으로는 권리를 들 수 있다. 그리고 후자의 예로는 사회 전체의 공리(utility)나 공익(public interest)을 언급할 수 있다. 위험은 바로 이러한 개인의 권리나 공익 등을 침해할 가능성, 달리 말해 아직 현실화되지 않은 가능성을 뜻한다. 그 점에서 위험은 미래 지향적인 개념이다. 그렇기에 위험이 실제로 현실화될 것인지는 오직 가능성으로만, 더욱 구체적으로 말하면 ‘가능성/개연성/확실성’으로 표현되는 확률로만 판단할 수 있다. 그 점에서 위험을 관리한다는 것은 미래에 관한 불확실성을 관리한다는 것을 뜻한다.

나. 인공지능의 위험

인공지능에는 어떤 위험이 있는가? 이는 어떤 기준에 따라 위험을 관찰하는가에 따라 다양하게 나눌 수 있다.

가령 위험의 강도, 즉 개인과 사회의 이익 침해 강도에 따라 인공지능의 위험을 유형화하는 방법을 고려할 수 있다. 유럽연합 인공지능법안이 채택하는 ‘허용할 수 없는 위험/고위험/제한된 위험/저위험 또는 최

8) 신뢰할 수 있는 인공지능에 관해서는 박도현, “신뢰할 수 있는 인공지능의 이론적 고찰”, 『법철학연구』 제25권 제2호(2022. 8), 321-364쪽 참고.

9) 위험에 관한 논의로는 양천수, “위험·재난 및 안전 개념에 대한 법이론적 고찰”, 『공법학연구』 제16권 제2호(2015. 5), 187-216쪽 참고.

소 위험'이라는 구별이 그 예가 된다.

또는 인공지능을 구성하는 요소, 가령 하드웨어나 소프트웨어, 데이터 등에 따라 위험을 구별하는 방법도 생각할 수 있다. 이는 인공지능의 위험 가운데 무엇을 어떻게 규율할 것인가를 파악하는 데 도움이 된다. 더 불어 이러한 구별 방식은 위험의 강도를 기준으로 한 구별과 연결된다. 인공지능이 허용할 수 없는 위험이 나 고위험을 지니는가는 결국 인공지능을 구성하는 각 요소가 지닌 위험들이 상호 작용함으로써 판단되는 것이기 때문이다. 위험의 강도에 따라 인공지능의 위험을 구별하는 방식은 유럽연합 인공지능법을 분석할 때 살펴볼 것이기에 아래에서는 인공지능을 구성하는 요소를 기준으로 하여 인공지능이 어떤 위험을 지니는지 간략하게 살펴본다.¹⁰⁾

1) 인공지능의 구성 요소

인공지능이 구현되려면 크게 세 가지 요소가 필요하다. 하드웨어와 소프트웨어 그리고 데이터가 그것이다. 인공지능이 작동하려면 먼저 하드웨어, 즉 물리적 세계에서 작동하는 기계적 체계가 필요하다. 특히 인공지능이 로봇과 결합하여 활용될 때는 이러한 하드웨어적 측면이 중요한 부분이 된다. 이러한 하드웨어에 관해서는 무엇보다도 반도체가 핵심적인 역할을 한다. 오늘날 인공지능 기술이 급속하게 진보할 수 있었던 배경에는 기하급수적으로 발전하는 반도체 기술이 중요한 역할을 하였다. 다음 알고리즘을 핵심으로 하는 소프트웨어가 필요하다. 알고리즘이 갖추어져야 기계적 체계인 인공지능 하드웨어는 비로소 인공지능으로서 작동할 수 있다. 나아가 데이터가 필요하다. 데이터가 입력 및 활용되어야 인공지능 알고리즘은 학습을 할 수 있고 이를 통해 비로소 우리가 원하는 역량을 발휘할 수 있다.

이러한 세 가지 요소를 갖추면 기계적 체계로서 인공지능은 구현된다. 그러나 인공지능은 자기 목적적인 존재는 아니다. 이는 인간을 위한 수단적인 존재이다. 인간을 위한 인공지능이어야 한다. 이렇게 인공지능이 인간을 위한 유용한 수단으로 작동할 수 있으려면 사회 각 영역에서 사용될 수 있어야 한다. 사회에서 사용되어야 인공지능은 우리 인간의 측면에서 볼 때 비로소 완전하게 구현될 수 있다. 그 점에서 앞에서 언급한 세 가지 요소에 더하여 사회적 사용이라는 요소가 인공지능의 작동 요소에 추가될 수 있다.

인공지능의 위험은 이러한 네 가지 요소에 각각 존재한다. 각 요소에 존재하는 위험이 상호 작용함으로써 인공지능의 위험이 종합적으로 결정된다. 그중에서도 데이터와 알고리즘에 위험이 집중된다.

2) 데이터의 위험

인공지능이 제대로 작동하려면 다양하면서 많은 양의 데이터가 필요하다. 이때 인공지능이 학습하는 데

10) 이에 관해서는 양천수, “지능정보기술의 위험과 법적 대응 방안: 알고리즘에 대한 대응을 중심으로 하여”, 『법학연구』(충북대) 제32권 제1호(2021. 6), 351-384쪽 참고.

활용하는 데이터 가운데는 개인정보와 저작물이 포함된다. 이에 따라 다음과 같은 위험이 출현한다. 개인 정보 침해 위험과 저작권 침해 위험이 그것이다. 가령 인공지능이 대량의 데이터를 학습하는 과정에서 정보 주체의 개인정보 자기결정권을 침해할 수 있다. 또는 저작자의 동의 없이 저작물을 인공지능 학습에 이용함으로써 저작권을 침해할 수 있다. 최근 ChatGPT와 관련하여 특히 이 문제가 불거졌다.

3) 알고리즘의 위험

인공지능의 위험에 관해 가장 주목의 대상이 된 것은 알고리즘이 지닌 위험이다. 알고리즘은 인공지능이 어떻게 작동하는지, 가령 어떤 판단을 내리는지에 영향을 미치기에 알고리즘이 가진 위험은 인공지능이 작동하는 방식 및 산출 결과에 직결된다. 이러한 위험으로 크게 두 가지를 거론할 수 있다. 첫째는 알고리즘이 정확하지 않게 작동할 위험, 즉 오작동의 위험이다. ChatGPT의 환각 이슈가 이를 잘 예증한다. 특히 가짜 뉴스가 정교하게 구현되고 있기에 알고리즘의 오작동 위험은 사회 전체에 큰 영향을 미칠 수 있다. 둘째는 알고리즘의 편향 위험이다. 알고리즘에 따라 진행되는 판단과 결정이 편향을 지녀 특정 개인이나 젠더, 집단 등을 차별하는 위험이 그것이다.

4) 하드웨어 및 사회적 사용의 위험

인공지능을 구성하는 물리적 하드웨어 및 인공지능이 사회적으로 사용되는 것에 관해서도 위험이 존재한다. 하드웨어, 특히 인공지능 반도체에 관해서는 공급망 보안 침해 위험이 존재한다. 인공지능을 구성하는 핵심 반도체가 사전에 해킹 또는 조작되어 인공지능 시스템 전체의 보안이 침해될 수 있다.¹¹⁾ 나아가 인공지능이 사회, 특히 실제 공간에서 작동하면서 개별 권리주체의 권리나 이익 등을 침해할 위험을 언급할 수 있다. 자율주행자동차가 교통사고를 일으켜 다른 자동차의 운전자나 보행자에게 피해를 입히는 경우를 들 수 있다.¹²⁾ 이외에 앞으로 더욱 크게 다가올 위험은 인공지능의 성능이 비약적으로 발전하여 인간 중심의 노동시장을 인공지능이 대체하는 것을 언급할 수 있을 것이다.¹³⁾

11) 이 문제에 관해서는 양천수·하민경, 『새로운 보안체계 구축에 관한 글로벌 규범 현황 및 법제 대응 방안 연구』(한국법제연구원, 2023) 참고.

12) 이에 관해서는 양천수, “자율주행자동차에 관한 법적책의 방향: 의의·위험·규제”, 『모빌리티연구』 제3권 제2호(2023. 9), 53-72 쪽 참고.

13) 이 문제를 다루는 제리 카플란, 신동숙 (옮김), 『인간은 필요없다』(한즈미디어, 2016) 참고.

3. 인공지능 규제 유형

인공지능은 어떻게 규제할 수 있을까? 달리 말해 인공지능의 위험은 어떻게 관리하는 게 적절한가? 이에 대해 우리가 선택할 수 있는 규제 유형을 검토할 필요가 있다. 규제 유형은 크게 다음과 같이 구별할 수 있다.

가. 자율규제와 타율규제

첫째, 자율규제와 타율규제를 구별할 수 있다. 이는 투입되는 규제의 강도와도 관련을 맺는다. 자율규제는 인공지능 관련 수법자가 자율적으로 인공지능의 위험을 규제하는 유형을 말한다. 윤리와 같은 연성 규범으로 인공지능을 규제하는 방안이 그 예다.¹⁴⁾ 반면 타율규제는 법과 같은 경성 규범으로 인공지능 관련 수법자를 강제적으로 규제하는 유형이다. 강제성이라는 타율성을 사용하는 것이다. 인공지능에 관해 ‘지원’이라는 시각을 지향하면 전자의 규제 유형을 선호한다. 이에 반해 인공지능에 대한 ‘규제’를 강조하는 시각에서 볼 때는 후자의 규제 유형을 선택하게 된다.

주의해야 할 점은 양자가 일도 양단적으로 선택되는 것은 아니라는 점이다. 두 유형은 규제 대상이 지닌 특성이나 그 대상이 처한 상황을 고려하여 적절하게 혼합될 수 있다. 자율규제와 타율규제를 동시에 사용하는 것이다. 유럽연합 인공지능법안이 바로 이러한 방식을 취한다. 단순화해서 말하면 고위험 인공지능에는 타율규제를, 저위험 인공지능에는 자율규제를 적용한다.

나. 행위 중심적 규제와 결과 중심적 규제

둘째, 행위 중심적 규제와 결과 중심적 규제로 유형화할 수 있다. 이는 독일 형법학에서 논의되는 행위반가치(행위불법)/결과반가치(결과불법) 구별을 원용한 것이다.¹⁵⁾ 여기서 행위 중심적 규제는 인공지능이 수행하는 행위나 결정 그 자체의 위험을 문제 삼는다. 예를 들어 알고리즘이 편향적인 결정을 하는 경우 이러한 결정 자체를 차별 금지 위반으로 규제하는 것이다. 또한 알고리즘이 투명성이나 설명 가능성이라는 기준을 이행하지 않는 경우 이를 규제하는 것도 이러한 예로 볼 수 있다.¹⁶⁾ 이에 반해 결과 중심적 규제는 인공지능이 특정한 권리나 이익을 침해하는 결과를 산출하는 경우 이러한 결과를 규제하는 유형을 말한다. 이는 법에서, 이른바 책임법에서 흔히 볼 수 있는 규제 유형에 해당한다.¹⁷⁾

14) 인공지능 윤리에 관해서는 인공지능과 가치 연구회, 『인공지능윤리: 다원적 접근』(박영사, 2021) 참고.

15) 이에 관해서는 심재우, “형법에 있어서 결과불법과 행위불법”, 『법학논집』 제20집(1982. 12), 127~170쪽 참고.

16) 투명성이나 설명 가능성은 인공지능 윤리에서 중요한 원칙으로 자리매김한다. 이에 관해서는 심우민, “알고리즘 투명성에 대한 규범적 접근방식”, 『디지털 윤리』 제2권 제1호(2018. 6), 1~12쪽 참고.

17) 책임법에 관해서는 양천수, 『책임과 법』(박영사, 2022) 참고.

다. 사전 규제와 사후 규제

셋째, 시간성을 기준으로 하여 사전 규제와 사후 규제를 구별할 수 있다. 사전 규제는 인공지능이 지닌 위험이 현실화되기 이전에 이를 규제하는 것을 말한다. 이러한 예로 사전 영향 평가 제도를 언급할 수 있다. 사전 영향 평가를 실시하여 규제 대상에 취약점이 발견된 경우 이에 적절한 조치를 취하게 하는 것이다. 인공지능에 설계주의를 적용하는 것도 사전 규제의 예로 볼 수 있다.¹⁸⁾ 이에 대해 사후 규제는 인공지능이 위법한 결과를 산출한 시점 이후에 비로소 이를 규제하는 방식을 말한다. 이는 전통적인 책임법이 즐겨 활용하는 방식이다. 그 점에서 사후 규제는 결과 중심적 규제와 연결된다.

라. 결과 비난 규제와 반성적 규제

넷째, 결과 비난 규제와 반성적 규제를 언급할 수 있다. 결과 비난 규제는 인공지능이 산출한 결과 그 자체를 비난하는 데 초점을 맞추는 규제를 말한다. 이는 형벌 이론에서 말하는 응보 이론과 연결된다.¹⁹⁾ 이에 반해 반성적 규제는 인공지능 규제 수범자가 규제를 통해 인공지능을 반성적으로 개선 또는 보완하게끔 하는 데 초점을 맞추는 규제를 말한다. 그 점에서 반성적 규제는 형벌 이론에서 말하는 특별 예방 이론과 연결된다.²⁰⁾ 이러한 예로 투명성 규제를 들 수 있을 것이다. 왜냐하면 인공지능 알고리즘에 투명성을 요구함으로써 알고리즘이 좀 더 투명해지고 설명 가능하도록 개선할 수 있기 때문이다.

18) 설계주의에 관해서는 成原慧, “ア-キテクヤの自由の再構築”, 松尾陽 (編), 『ア-キテクヤと法』(弘文堂, 2016) 참고.

19) 응보 이론을 포함하는 형벌 이론에 관해서는 빈프리트 하세머, 배종대 · 윤재왕 (옮김), 『범죄와 형벌: 올바른 형법을 위한 변론』(나남, 2011) 참고.

20) 특별 예방 이론에 관해서는 프란츠 폰 리스트, 심재우 · 윤재왕 (옮김), 『마르부르크 강령: 형법의 목적사상』(강, 2012) 참고.

III

유럽연합 인공지능법의 배경

1. 인공지능에 관한 유럽연합 법정책의 배경

인공지능에 관한 유럽연합의 법적 규제를 살펴보기 전에 최근 유럽연합이 어떤 방향의 정책 혹은 법정책을 지향하는지 살펴본다. 그 이유는 최근 유럽연합은 단순한 국가연합의 단계를 넘어 준연방국가의 모습을 보이고 있고 이에 따라 정책 역시 일관된 방향과 체계 아래 추진하고 있기 때문이다.²¹⁾ 요컨대 오늘날 유럽연합은 로마법이 적용되었던 과거 로마제국이나 공통법(ius commune)이 규율되었던 중세의 신성로마제국이 그랬던 것처럼 공통된 규범적 가치와 기준을 설정하여 이를 정책 및 법정책에 담아낸다. 인공지능에 관한 유럽연합의 법정책 역시 이러한 배경 아래에서 파악할 필요가 있다.

가. 지속가능성

먼저 지속가능성(sustainability)을 언급할 수 있다. 유럽연합은 지속가능성을 중요한 정책 목표로 설정한다. 시간성의 측면에서 바꾸어 말하면 단기간이 아닌 장기간의 견지에서 미래 지향적인 정책을 펼친다는 것이다. 이는 유럽 그린딜(The European Green Deal)에서 잘 확인할 수 있다.²²⁾ 기후 변화에 대응하기 위해 탄소 중립, 즉 탈탄소 경제를 구현하고자 하는 유럽 그린딜은 구체적인 로드맵으로 크게 두 가지를 제시한다. 첫째는 자원의 효율적 사용을 확대함으로써 현재의 탄소 경제를 청정·순환 경제로 전환하는 것이다. 둘째는 기후 변화를 원래 상태로 되돌림으로써 생물 다양성을 복원하고 환경 오염을 줄이는 것이다.

21) 이에 관해서는 타마르 헤르초그, 이영록 (옮김), 『유럽법약사』(민속원, 2023) 참고.

22) 유럽 그린딜에 관해서는 (https://www.eeas.europa.eu/delegations/south-korea/%EC%9C%A0%EB%9F%BD%EA%B7%B8%EB%A6%B0%EB%94%9C_ko?s=179) 참고(검색 일자: 2024. 2. 1.).

지속가능성을 지향하는 유럽연합의 정책에서 크게 두 가지 시사점을 찾을 수 있다. 우선 유럽연합은 단기적인 성과에 집착하지 않는다는 점이다. 나아가 사회, 더 나아가 세계 전체는 유기적으로 연결되어 있다는 생태주의적 관점을 바탕으로 삼는다는 점이다.²³⁾ 이는 근대법의 기반이 되었던 분할적 관점, 즉 주체와 주체 및 객체를 명확하게 분리하는 사고방식이 오늘날 더 이상 유효하지 않다는 시각을 함의한다.²⁴⁾

나. 인권 보호

다음으로 인권 보호를 언급할 수 있다. 유럽인권협약과 유럽인권법원이 예증하듯이 인권은 유럽연합이 지향해야 하는 중요한 규범적 가치이자 기준이다. 유럽연합에서 인권은 단순한 도덕적 권리에 머무는 것이 아니라 법적 권리로 자리매김한다. 이와 더불어 유럽연합은 인권의 개념적 외연을 지속적으로 확장한다. 이러한 인권은 제4차 산업혁명이 진행되면서 이슈가 되는 디지털 전환, 빅데이터, 플랫폼 경제, 인공지능 등에서 중요한 규범적 기준이 된다. 최근 관심의 초점이 되는 공급망(supply chain) 문제에서도 인권은 환경과 더불어 중요한 보호 대상, 달리 말해 실사(due diligence) 대상으로 자리 잡는다. 유럽연합의 공급망 실사 법안(Supply chain due diligence directive)이 이를 잘 예증한다.²⁵⁾

다. 역내 산업 및 시장 보호

나아가 최근 들어 유럽연합이 역내 산업 및 시장 보호를 강화한다는 점을 언급할 필요가 있다. 요컨대 유럽연합이 블록 경제와 보호무역주의를 강화하는 것이다. 냉전과 이념 대결이 끝나고 세계화(globalization) 흐름이 전 세계를 휩쓸면서 자유무역주의가 지배적인 경제 기조가 되었다. 중국이 세계의 공장으로 떠오르고 신자유주의의 물결 아래 이른바 세계 시장 및 세계 경제 체계가 형성되었다. 이에 따라 WTO의 영향력이 강화되었다. 초국가적 사회 그리고 초국가적 법도 실현 가능한 것으로 희망적으로 예측되었다. 더불어 세계 각국은 FTA를 체결함으로써 자유 시장 영역을 국가 상호적으로 확장하였다. 이에 따라 심지어 과연 국가가 오늘날에도 여전히 필요한지에 의문이 제기되기도 하였다. 한 국가의 권력을 넘어서는 초국적 기업이 세계 경제 체계에서 새로운 권력 주체로 부상하였다.

그러나 미국과 중국의 패권 경쟁 및 갈등이 심화하고 러시아-우크라이나 전쟁, 이스라엘-하마스 전쟁 등이 발발하면서 세계화의 환상은 깨지고 자유무역주의는 한계를 맞고 있다. 블록 경제를 중심으로 한 새로운 냉전이 세계를 엄습한다. 시장에 빼앗겼던 권력을 국가가 되찾고 있는 것이다. 이에 따라 그 존재 필요성이

23) 생태주의적 관점에 관해서는 다치바나 다카시, 김경원 (옮김), 『생태학적 사고법』(바다출판사, 2021) 참고.

24) 이를 보여주는 프리초프 카프라 · 우고 마테이, 박태현 · 김영준 (공역), 『최후의 전환: 지속 가능한 미래를 위한 커먼즈와 생태법』(경희대학교출판문화원, 2019) 참고.

25) 공급망 실사 법안의 공식 명칭은 Proposal for a Directive of the European Parliament and the Council on Corporate Sustainability Due Diligence and amending Directive (EU) 2019/1937이다. 공식 약칭은 Corporate Sustainability Due Diligence Directive (CSDDD)이다.

의심되었던 국가의 의미가 새롭게 복원된다. 이에에는 크게 두 가지 이유를 제시할 수 있다. 첫째, 제4차 산업혁명이 진행되면서 과학기술이 경제 성장에 결정적인 요소가 되고 이에 과학기술을 둘러싼 선진국 간의 패권 경쟁이 격화되고 있다는 것이다. 데이터나 인공지능, 반도체, 배터리, 양자컴퓨터 기술 등을 둘러싼 미국, 중국, 유럽연합의 경쟁이 이를 예증한다. 둘째, 러시아-우크라이나 전쟁이 보여준 것처럼 국가 경제의 기반이 되는 핵심 물자의 공급망이 불안정해졌다는 것이다.

이 같은 상황에서 유럽연합은 무엇보다도 미국이나 중국, 러시아 등으로부터 역내 산업 및 시장을 보호하기 위해 다양한 법규범을 마련해 시행한다. 대표적인 예로 DSA(Digital Service Act)와 DMA(Digital Market Act)를 들 수 있다.²⁶⁾ 이들 법규범은 구글이나 아마존, 메타, 넷플릭스와 같은 플랫폼 기업을 겨냥한다. 플랫폼 경제가 새로운 경제 체제로 안착하면서 플랫폼 기업의 시장 지배력이 강화된다. 플랫폼 기업은 빅데이터 분석 및 고객 맞춤형 서비스로 무장하여 시장 지배를 확장한다. DSA와 DMA는 디지털 서비스와 디지털 시장 규제로 이들 플랫폼 기업을 통제함으로써 유럽연합의 역내 산업 및 시장을 보호하고 유럽연합 시민들의 기본권을 보호하고자 한다. 이를테면 개인 데이터를 활용한 고객 맞춤형 광고를 규제하거나 경쟁 제한 등을 근거로 하여 플랫폼 기업의 시장 지배를 통제한다.

이 글이 분석 대상으로 삼는 유럽연합의 인공지능법안도 이러한 기능을 담당한다.²⁷⁾ 위험 기반 접근법에 따라 인공지능을 ‘허용할 수 없는 위험(unacceptable risk)의 인공지능/고위험 인공지능/저위험 또는 최소 위험 인공지능’으로 유형화함으로써 미국과 특히 안면 인식 기술을 강화하는 중국에 대해 유럽연합의 인공지능 산업 및 시장을 보호하고자 한다.

유럽연합의 반도체법(EU Chips Act)이나 배터리법(EU Battery Regulation)도 마찬가지로 맥락에서 파악할 수 있다.²⁸⁾ 오늘날 인공지능 기술을 포함한 디지털 기술이 하드웨어의 측면에서는 반도체를 필수 요소

26) DAS의 공식 명칭은 Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC이다. 이에 관한 소개는 (<https://www.eu-digital-services-act.com/>) 참고(검색 일자: 2024. 2. 1.). 이에 관한 연구로는 이병준, “유럽연합 디지털 서비스법을 통한 플랫폼 규제: 디지털 서비스법 초안의 주요내용과 입법방향을 중심으로”, 『소비자법연구』 제7권 제2호(2021. 5), 181~210쪽 참고. DMA의 공식 명칭은 Regulation (EU) 2022/1925 of the European Parliament and of the Council of 14 September 2022 on contestable and fair markets in the digital sector and amending Directives (EU) 2019/1937 and (EU) 2020/1828이다. 이에 관한 소개로는 https://digital-markets-act.ec.europa.eu/legislation_en 참고(검색 일자: 2024. 2. 1.).

27) 유럽연합 인공지능법안의 공식 명칭은 Proposal for a REGULATION OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL LAYING DOWN HARMONISED RULES ON ARTIFICIAL INTELLIGENCE (ARTIFICIAL INTELLIGENCE ACT) AND AMENDING CERTAIN UNION LEGISLATIVE ACTS이다. 이는 (<https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=CELEX:52021PC0206>)에서 확인할 수 있다(검색 일자: 2024. 2. 1.).

28) 반도체법의 공식 명칭은 Regulation (EU) 2023/1781 of the European Parliament and of the Council of 13 September 2023 establishing a framework of measures for strengthening Europe’s semiconductor ecosystem and amending Regulation (EU) 2021/694이다. 이에 관한 소개로는 심소연, “유럽연합(EU)의 「반도체법(Chips Act)」”, 『최신 외국입법정보』 제234호(2023. 10. 31) 참고. 배터리법의 공식 명칭은 Regulation (EU) 2023/1542 of the European Parliament and of the Council of 12 July 2023 concerning batteries and waste batteries, amending Directive 2008/98/EC and Regulation (EU) 2019/1020 and repealing Directive 2006/66/EC이다.

로 한다는 점에서 반도체 설계 및 생산 전 과정에 관한 주권을 누가 보유하는지는 디지털 기술 전체에 관한 주권과 직결된다. 이에 유럽연합은 반도체 설계 및 생산에 관한 주권을 확보하고자 한다. 달리 말해 유럽연합 역내에서 반도체에 관한 전 과정을 실행할 수 있도록 한다. 반도체법이 이를 규율한다. 나아가 탄소 중립 시대에 대응할 수 있는 모빌리티(mobility)는 전기자동차이고 이에 가장 중요한 요소가 배터리라는 점에서 유럽연합은 배터리에 관한 주권도 확보하고자 한다. 배터리법이 이에 기여한다.

라. 유럽연합 역내 공급망 안전 강화

유럽연합의 역내 산업 및 시장을 보호하기 위해 필요한 정책 가운데 하나가 역내 공급망의 안전을 강화하는 것이다. 러시아-우크라이나 전쟁으로 러시아가 유럽연합의 실제적인 위협으로 다시 부상하면서 이는 매우 현실적인 요청이 되었다. 당장 석유나 천연가스와 같은 에너지원의 안정적인 공급이 의문에 처하게 되었다. 이때 공급망 안전은 두 가지 측면으로 유형화할 수 있다. 물리적 측면의 공급망 안보와 디지털 측면의 공급망 보안이 그것이다.

물리적 측면에서는 핵심 물자를 공급하는 공급망을 물리적 공간에서 안전하게 보장하는 게 중요하다. 예를 들어 전쟁이나 테러 등으로 식량이나 천연가스, 석유 등과 같은 핵심 물자의 공급망이 물리적인 실제 세계에서 단절되는 것을 막을 필요가 있다. 이는 한 국가의 경제 안보와도 직결되기에 공급망 안보(supply chain defense)로 지칭되기도 한다.

디지털의 측면에서는 공급망을 통해 발생하는 정보 보안 공격, 가령 공급망의 디지털 네트워크를 대상으로 하는 해킹이나 랜섬웨어 공격 등에 대처함으로써 디지털 영역에서 공급망 보안(supply chain security)을 구현해야 한다. 후자의 경우에는 특히 2021년을 전후로 하여 미국에서 문제가 되었다.

앞에서 언급한 반도체법이나 배터리법은 유럽연합의 공급망을 안전하게 보장하는 데 이바지한다. 공급망 실사 법안도 이에 간접적으로 도움을 준다. 기후 변화 및 인권 보호에 적절하게 대처할 수 있도록 기업의 공급망을 통제함으로써 결과적으로 이에 친화적인 유럽연합 역내 기업들의 공급망이 강화되도록 하는 것이다. 더불어 NIS, NIS2 지침 등은 공급망의 디지털 보안 공격에 대처할 수 있도록 한다.²⁹⁾

29) NIS 지침(Directive on security of network and information systems)의 공식 명칭은 Directive (EU) 2016/1148 of the European Parliament and of the Council of 6 July 2016 concerning measures for a high common level of security of network and information systems across the Union이다. NIS2 지침은 이를 개정한 것으로 공식 명칭은 Directive (EU) 2022/2555 of the European Parliament and of the Council of 14 December 2022 on measures for a high common level of cybersecurity across the Union, amending Regulation (EU) No 910/2014 and Directive (EU) 2018/1972, and repealing Directive (EU) 2016/1148이다.

2. 유럽연합 규제체계의 경향

최근 유럽연합이 마련하는 규제체계에서는 다음과 같은 경향을 발견할 수 있다.

가. 위험 기반 접근법

첫째, 규제를 설계할 때 위험 기반 접근법(risk based approach)을 취한다는 점이다. 이는 크게 세 가지 의미를 담는다. 우선 인공지능법안이 보여주듯이 규제 대상이 지닌 위험을 평가하여 이를 허용할 수 없는 위험, 고위험, 저위험 등과 같이 유형화한다는 것이다. 다음 각 위험 유형에 적합한 강도의 규제를 설계 및 투입한다는 것이다. 재생의료 산업에 관한 규제에서 독일이 주로 사용하는 ‘이익-위험-관계’(Nutzen-Risiko-Verhältnis) 원칙이 이를 예증한다. 나아가 사후 지향적(ex post) 규제가 아닌 사전 지향적(ex ante) 규제에 중점을 둔다는 것이다. 그 이유는 위험 개념 자체가 미래 지향적 개념이기 때문이다.³⁰⁾

나. 투명성과 설명 가능성

둘째, 투명성과 설명 가능성을 중시한다는 것이다. 이 점도 인공지능법안에서 발견된다. 규제 대상이 되는 행위 주체, 특히 사회에서 이른바 강자에 속하는 행위 주체가 특정한 행위나 의사 결정을 할 때는 이를 투명하게 그리고 설명 가능하게 해야 한다는 것이다. 이는 사회 전반에 ‘절차적 정당화’를 강화하는 규제를 적용하는 것이라 말할 수 있다.

다. 사전 영향 평가

셋째, 영향 평가(impact assessment), 그중에서도 사전 영향 평가를 강조한다는 것이다. 이는 위험 기반 접근법과 관련이 있다. 위험 기반 접근법을 취하려면 규제 대상이 어떤 위험을 지니는지 평가할 수 있어야 하는데 이때 적용하는 방법이 영향 평가인 셈이다. 이러한 영향 평가는 최근 유럽연합이 마련하는 여러 법적 규제에서 즐겨 사용된다. 한편 영향 평가, 특히 사전 영향 평가를 강조하는 것은 아래에서 언급하는 절차주의적 규제와도 관련을 맺는다.

라. 절차주의적 규제

넷째, 절차주의적 규제를 즐겨 사용한다는 것이다. 여기서 절차주의적 규제란 타율규제와 자율규제를 혼합한 규제의 일종에 해당한다.³¹⁾ 이를테면 피규제자에 특정한 규제를 해야 한다는 점은 법규범으로 정하되

30) 이에 관해서는 양천수, “위험·재난 및 안전 개념에 대한 법이론적 고찰”, 『공법학연구』 제16권 제2호(2015. 5), 187-216쪽 참고.

31) 절차주의적 규제에 관해서는 Graf-Peter Calliess, Prozedurales Recht (Baden-Baden, 1999) 참고.

이를 어떻게 구체화하여 적용할 것인지는 피규제자의 자율에 맡기는 방식을 들 수 있다. 이를 잘 보여주는 예가 공급망 실사 법안이다. 이에 따르면 유럽연합은 규제 대상이 되는 기업이 공급망 실사를 해야 한다는 점은 유럽연합 지침(directive)으로 규율한다. 다만 이를 어떻게 구체화하여 실시할 것인지는 규제 대상이 되는 기업의 자율적인 선택에 맡긴다.

절차주의적 규제 아래에서 피규제자는 규제를 어떻게 스스로 구체화해 자신에게 적용할지를 판단하기 위해 위험 기반 접근법과 사전 영향 평가를 활용하게 된다. 사전 영향 평가를 실시하여 자신에게 어떤 위험이 있는지 규명한 후 이를 적절하게 규율할 수 있는 규제를 자율적으로 구체화한다. 이 점에서 위험 기반 접근법과 사전 영향평가는 절차주의적 규제 아래에서 기능적으로 결합한다.

마. 사후 추적 조사

다섯째, 사후 추적 조사를 강화한다는 것이다. 이는 이른바 '사려 깊은 경계'(prudent vigilance) 모델에 바탕을 둔다.³²⁾ 이에 따라 규제가 적용되는 시간적 범위를 확장한다. 과거/현재/미래, 즉 사후 · 현재 · 사전 모두에 지속적으로 규제를 투입하는 것이다. 이를 통해 사후 추적 조사와 같은 규제가 수용된다. 예를 들어 재생 의료와 관련된 특정한 행위나 제품을 허가한 이후에도 사후적으로 그 위험 등을 추적 조사하여 이를 환류하는 것이다.

3. 유럽연합의 규제 형식

유럽연합이 정립할 수 있는 규제 형식은 크게 구속력이 있는 규제, 즉 법적 규제와 구속력이 없는 규제로 나눌 수 있다. 전자의 예로 규정(regulation), 지침(directive), 결정(decision)이 있다. 후자의 예로 권고(recommendation)를 들 수 있다.³³⁾

이 가운데 규정과 지침은 추상적 · 일반적인 형식을 취한다. 그 점에서 일반적인 법규범에 해당한다. 이에 반해 결정은 행정부가 발령하는 행정 처분이나 법원이 내리는 판결처럼 개별적 · 구체적인 형식을 취한다. 규범 창설(Normbegründung)과 규범 적용(Normanwendung) 또는 법규범과 판결을 형식적으로 구별하는 이분법에 따르면 결정은 법규범보다는 판결에 가깝다고 평가할 수 있다.

32) 이에 관해서는 Brian Patrick Green, "Six Approaches to Making Ethical Decisions in Cases of Uncertainty and Risk", Markkula Center for Applied Ethics (2019) 참고. 이는 (<https://www.scu.edu/ethics/focus-areas/technology-ethics/resources/six-approaches-to-making-ethical-decisions-in-cases-of-uncertainty-and-risk/>)에서 확인할 수 있다(검색 일자: 2024. 2. 1.).

33) 이에 관해서는 庄司克宏, 『はじめてのEU法』(有斐閣, 2020) 참고. 한편 'Regulation'은 '명령'으로 번역되기도 한다. 실제로 Regulation은 독일어로는 'Verordnung'으로 공식 표현된다.

규정과 지침은 다음과 같은 점에서 구별된다. ‘직접 적용/간접 적용’이라는 도식을 활용하면 규정은 유럽 연합 회원국에 직접 적용되는 반면, 지침은 간접 적용된다는 것이다. 요컨대 지침은 회원국에 의해 국내법으로 전환되어야 비로소 효력이 있다. ‘원칙(principle)/규칙(rule) 구별’을 활용하면 규정이 규칙으로서 회원국을 직접 구속한다면, 지침은 원칙으로서 회원국이 정립해야 하는 국내법의 방향을 제공한다. 그 점에서 상세하고 직접적인 내용을 담는 규정보다 지침은 상대적으로 간소한 내용을 규율한다.

4. 유럽연합 인공지능법의 입법과정

2021년 4월 21일에 시작된 유럽연합 인공지능법 입법과정은 최근 드디어 막바지에 접어들었다. 지난 2023년 12월 8일 유럽의회와 이사회의 협상가들이 인공지능법안에 관해 협상하고 이에 합의하였기 때문이다. 아직은 인공지능법의 최종 조항이 공식적으로 제시되지는 않았다.³⁴⁾ 하지만 2021년부터 시작된 인공지능법 입법과정이 이제 끝맺음으로 나아가고 있다는 점은 확인할 수 있다. 아래에서는 이러한 입법과정을 간략하게 정리한다.

가. 인공지능법 제안 이전

인공지능법안은 2021년 4월 21일 유럽연합 집행위원회가 공식적으로 제안하였다. 여러 연구가 지적하듯이 인공지능법안은 갑자기 도출된 것이 아니라 이전부터 진행된 인공지능 규제 논의, 즉 ‘신뢰할 수 있는 인공지능’(trustworthy AI)을 구현하기 위한 논의의 연장선상 혹은 최종 과정에서 제안된 것이다.³⁵⁾

인공지능 기술이 급속하게 발전하면서 한편으로는 인공지능 기술 발전을 지원하면서도 다른 한편으로는 인공지능의 위험을 어떻게 관리할 것인가가 유럽연합에도 중요한 과제가 되었다. 이에 인공지능을 어떻게 규제할 것인가가 논의되기 시작하였다. 유럽연합에서도 처음에는 윤리라는 연성 규범으로 인공지능을 규제하는 방안이 논의되었다. 이를 위해 2018년 6월 “인공지능에 관한 고위 전문가 그룹”(High-Level Expert Group on AI, 이하 ‘고위 전문가 그룹’으로 약칭)을 발족하였다. 이러한 고위 전문가 그룹이 2018년 12월에 초안을 발표하고 그 후 이해관계자들의 의견을 수렴하여 2019년 4월에 최종안으로 발표한 것이 바로 「신뢰할 수 있는 인공지능을 위한 윤리 가이드라인」(The Ethics Guideline for Trustworthy Artificial Intelligence)이다.³⁶⁾ 이에 연계하여 고위 전문가 그룹은 2020년 7월 「신뢰할 수 있는 인공지능을 위한 평

34) 2024년 2월 9일 현재 시점.

35) 이에 관해서는 김송욱, “AI 법제의 최신 동향과 과제: 유럽연합(EU) 법제와의 비교를 중심으로”, 『공법학연구』 제22권 제4호(2021. 11), 117-121쪽 참고.

36) 이에 관해서는 이보연, “유럽연합의 인공지능 관련 입법 동향을 통해 본 시사점”, 『법학논문집』(중앙대) 제43집 제2호(2019. 8), 13-14쪽 참고.

가 항목」(Assessment List for Trustworthy Artificial Intelligence)을 발표하였다.³⁷⁾

나. 인공지능 입법과정

1) 유럽연합 집행위원회 제안

유럽연합 집행위원회가 2021년 4월 21일에 제안한 인공지능법안은 그 이전에 정립된 인공지능 윤리 체계의 다음 단계로 마련된 것이다. 인공지능법안이 제안된 배경에는 인공지능 윤리로는 인공지능이 지닌 위험을 만족스럽게 규제하기 어렵다는 문제의식이 작용한 것으로 보인다.³⁸⁾ 그리고 현실적으로는 유럽연합에 비해 인공지능 기술이 급격하게 발전하던 미국과 중국에 대처하기 위한 조치로 해석할 수 있다. 한편 유럽 연합 집행위원회가 제안한 인공지능법안은 인공지능을 ‘규정’(regulation), 즉 법이라는 경성 규범으로 규율하는 본격적인 포괄적 법안이기에 전 세계적으로 많은 관심의 대상이 되었다. 우리나라에서도 법안이 제치된 이후 많은 연구가 축적되었다.³⁹⁾ 동시에 과연 이러한 법안이 실제로 입법될 수 있는지에 관해서도 흥미로운 관찰 대상이 되었다.

2) 유럽의회 수정 통과

유럽연합 집행위원회가 인공지능법안을 제안한 후 2년 이상이 지난 시점인 2023년 6월 14일 인공지능법안은 유럽의회를 통과한다. 이 과정에서 최초 법안은 상당 부분 수정 및 보완된다. 가장 중요한 이유는 그 사이 인공지능에 관해 매우 큰 이슈가 발생하였기 때문이다. 바로 Open AI의 ChatGPT 이슈가 그것이다. 2022년 11월 ChatGPT가 출시되면서 이는 곧바로 전 세계적인 열풍을 일으켰다. 더불어 생성형 인공지능(generative AI)에 관한 관심과 우려도 증폭시켰다. 이어 2023년 3월 GPT 4가 출시되면서 ‘파운데이션 인공지능’(foundation AI)이라는 새로운 인공지능 모델이 등장하였다. 문제는 ChatGPT와 같은 언어적 생성형 인공지능과 파운데이션 인공지능은 2021년에 제안된 인공지능법안이 충분히 고려하지 못한 규율 대상이었다는 점이다. 이에 유럽의회는 이를 반영하는 수정안을 통과시킨 것이다.

3) 최종 합의

그리고 2023년 12월 8일 유럽의회를 통과한 인공지능법안은 드디어 유럽의회 및 이사회 간 최종 합의

37) 이에 관해서는 유재홍·추형석·강송희, 『유럽(EU)의 인공지능 윤리 정책 현황과 시사점: 원칙에서 실천으로』, Issue Report IS-114(소프트웨어정책연구소, 2021. 3. 25.), 12-14쪽 참고.

38) Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative Acts, p. 3.

39) 여러 연구 가운데 몇 가지를 언급하면 이경선, “EU 인공지능 규제안의 주요 내용과 시사점”(정보통신정책연구원, 2021); 한국지능정보사회진흥원 정책본부 지능화법제도팀, “EU 인공지능법(안)의 주요 내용과 시사점”(한국지능정보사회진흥원, 2021); 김준권, “EU 인공지능명령안의 주요 내용과 그 시사점”, 『헌법재판연구』제8권 제2호(2021. 12), 65-100쪽 등 참고.

에 도달한다. 인공지능법안이 최종 관문을 넘어서는 것이다. 물론 인공지능법안이 실제로 제정 및 시행되기까지는 여전히 시간이 필요하다는 관측도 제시된다.⁴⁰⁾ 그렇지만 세계 최초로 인공지능을 포괄적으로 규율하는 법이 눈앞에 있다는 점은 의미가 결코 적지 않다.

40) 정병일, “인공지능 규제 법안 유럽의회 통과”, 『AI타임스』(2023. 6. 15)(<https://www.aitimes.com/news/articleView.html?idxno=151738>)(검색 일자: 2024. 1. 29.).

IV

유럽연합 인공지능법 분석

아래에서는 유럽연합 인공지능법이 어떤 내용을 담고 있는지 분석한다. 다만 이에 관해서는 그동안 상당한 연구가 축적되었기에 아래에서는 쟁점을 중심으로 하여 필요한 한도에서 논의를 전개하겠다. 더불어 인공지능법은 유럽연합 집행위원회가 제시한 초안과 유럽의회에서 통과된 안 그리고 최종 합의된 안 사이에 주목할 만한 차이가 있으므로 아래에서는 이를 구별하여 서술하도록 한다.

1. 유럽연합 집행위원회 법안

가. 전체 구조

유럽연합 집행위원회가 제안한 인공지능법안은 다음과 같이 구성된다.⁴¹⁾ 먼저 법안 설명(Explanatory Memorandum)과 총 89개 항으로 구성된 전문(Recital)이 제시된다. 이어 12개의 편(Title)과 85개의 조(Article)로 구성된 본문이 제시된다. 그중 가장 중요한 부분을 규율하는 제3편, 즉 고위험 인공지능 시스템을 규율하는 제3편은 다시 5개의 장(Chapter)으로 구성된다. 이외에 거버넌스를 규율하는 제6편은 2개의 장, 시장 출시 후 모니터링 등을 규율하는 제8편은 3개의 장으로 구성된다. 이는 다음 <표 1>과 같다.

41) 이에 관해서는 Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative Acts, p. 12-16(이하 인공지능법안은 'Proposal for AI Act'로 약칭하여 인용한다.); 홍석한, "유럽연합 '인공지능법안'의 주요 내용과 시사점", 『유럽헌법연구』 제38호(2022. 4), 250-251쪽 참고.

표 1 집행위원회 인공지능법안의 전체 구조

Title		내용
Explanatory Memorandum		법안 설명
Recital 1-89		법안 제안 이유
Title I		적용 범위와 개념 정의(Scope and Definitions)
Title II		금지된 인공지능 실행(Prohibited Artificial Intelligence Practices)
Title III		고위험 인공지능 시스템(High-Risk AI Systems)
Title III	Chapter 1	고위험 인공지능 시스템의 유형
	Chapter 2	고위험 인공지능 시스템이 갖추어야 하는 요건
	Chapter 3	고위험 인공지능 시스템 공급자와 사용자 및 기타 당사자의 의무
	Chapter 4	인증 당국과 피인증 기구(Notifying Authorities and Notified Bodies)
	Chapter 5	기준, 적합성 평가, 인증서, 등록
Title IV		특정한 AI 시스템에 관한 투명성 의무(Transparency Obligations for Certain AI Systems)
Title V		혁신 지원에 관한 조치(Measures in Support of Innovation)
Title VI		거버넌스(Governance)
Title VI	Chapter 1	유럽 인공지능 위원회(European Artificial Intelligence Board)
	Chapter 2	국가 관할 당국(National Competent Authorities)
Title VII		EU 자립형 고위험 인공지능 시스템 데이터베이스(EU Database for Stand-Alone High-Risk AI Systems)
Title VIII		시장 출시 후 모니터링, 정보 공유, 시장 감시(Post-Market Monitoring, Information Sharing, Market Surveillance)
Title VIII	Chapter 1	시장 출시 후 모니터링
	Chapter 2	사고 및 악의적 작동에 관한 정보 공유(Sharing of Information on Incidents and Malfunctions)
	Chapter 3	집행(Enforcement)
Title IX		행동 강령(Codes of Conduct)
Title X		비밀 유지와 벌칙(Confidentiality and Penalties)
Title XI		권한 위임과 위원회 절차(Delegation of Power and Committee Procedure)
Title XII		최종 조항(Final Provisions)

* 출처: 법안을 참고하여 직접 구성

나. 규범 목적과 핵심 규율 내용

1) 규범 목적

집행위원회 인공지능법안은 규범 목적으로 다음을 제시한다.⁴²⁾ 첫째는 유럽연합에서 출시 및 사용되는 인공지능 시스템이 안전하고 기본권에 관한 법과 유럽연합의 가치를 존중하도록 하는 것이다. 둘째는 인공지능에 대한 투자와 혁신을 촉진할 수 있도록 법적 안정성(legal certainty)을 보장하는 것이다. 셋째는 인공지능 시스템에 적용될 수 있는 기본권 및 안전 요건을 규율하는 기존 법에 관한 거버넌스와 효과적인 집행을 강화하는 것이다. 넷째는 합법적이고 안전하며 신뢰할 수 있는 인공지능 애플리케이션을 위한 단일 시장이 발전하는 것을 촉진하고 시장이 분절화되는 것을 막는 것이다.

이는 다음과 같이 분석할 수 있다. 첫 번째 목적은 인공지능의 위험과 관련을 맺는다. 인공지능의 위험을 안전하게 관리하여 유럽연합이 법규범으로 보장하는 기본권 및 규범적 가치를 보장하겠다는 것이다. 그 점에서 첫 번째 목적은 인공지능에 대한 규제를 지향한다. 이에 반해 두 번째 목적은 인공지능법으로 법적 안정성과 예측 가능성을 보장함으로써 인공지능에 대한 투자와 혁신을 구현하고자 하는 데 있다. 그 점에서 두 번째 목적은 인공지능에 대한 지원을 지향한다. 이렇게 보면 첫 번째 목적과 두 번째 목적은 상반된다. 인공지능법안은 이렇게 서로 부딪히는 규범 목적을 조화롭게 구현하고자 한다. 요컨대 인공지능법안은 목적 간의 조화(harmonised)를 매우 강조한다. 세 번째 목적은 이러한 두 목적을 구현하는 데 필요한 거버넌스와 집행체계를 구축하는 것과 관련된다. 그 점에서 세 번째 목적은 첫 번째 목적과 두 번째 목적을 실현하는 데 필요한 수단의 지위에 해당한다. 네 번째 목적은 유럽연합의 인공지능 시장이 견고하게 유지되는 것과 관련된다. 이러한 네 가지 목적을 키워드로 요약하면 ‘위험 관리/투자 및 혁신/거버넌스/단일한 시장’이라 말할 수 있다.

2) 핵심 규율 내용

이러한 규범 목적을 구현하기 위해 인공지능법안은 다음과 같은 내용을 핵심 규율 내용(subject matter)으로 설정한다(제1조).

첫째는 유럽연합 역내에서 시장에 출시되고 서비스에 투입되며 이용되는 인공지능 시스템을 규율하는 조화로운 규칙을 마련하고 특정한 인공지능 실행을 금지하는 것이다(제1조 (a)).

둘째는 고위험 인공지능 시스템에 관한 특별한 요건과 이러한 인공지능 시스템의 운영자가 부담해야 하는 의무를 설정하는 것이다(제1조 (b)).

42) Proposal for AI Act, p. 3.

셋째는 자연인과 상호 작용을 의도하는 인공지능, 감정 인식 시스템, 생체 분류 시스템 그리고 이미지, 소리, 영상 콘텐츠를 생성 또는 조작하는 인공지능 시스템에 적용되는 조화로운 투명성 규칙을 설정하는 것이다(제1조 (c)).

넷째는 시장 모니터링과 감시에 관한 규칙을 마련하는 것이다(제1조 (d)).

다. 적용 대상

인공지능법안은 다음과 같은 주체를 적용 대상으로 예정한다(제2조).

첫째는 유럽연합 역내에서 인공지능 시스템을 시장에 출시하거나 서비스에 적용하는 공급자이다. 이때 공급자가 유럽연합 역내 또는 제3국에서 설립되었는지는 고려되지 않는다(제2조 제1항 (a)).

둘째는 유럽연합 역내에 있는 인공지능 시스템 사용자이다(제2조 제1항 (b)).

셋째는 제3국에 자리한 인공지능 시스템의 공급자와 사용자이기는 하지만 이러한 인공지능 시스템에 의해 생산되는 산출물이 유럽연합 역내에서 사용되는 경우의 공급자와 사용자이다(제2조 제1항 (c)).

여기서 다음과 같은 점을 주목할 수 있다.

우선 인공지능 법안은 적용 대상이 되는 주체를 기능적인 측면에서 확장한다는 것이다. 원칙적으로는 유럽연합 역내에서 활동하는 인공지능 시스템의 공급자와 이용자를 적용 대상으로 한다. 하지만 기능적인 측면에서 볼 때 인공지능이 산출한 결과물이 유럽연합 역내에서 사용되는 때에는 역내에 자리하지 않는 공급자와 이용자도 규율 대상에 포섭하는 것이다.

나아가 인공지능 시스템 그 자체가 아니라 인공지능 시스템을 공급하거나 사용하는 공급자와 사용자 등과 같은 인적 주체를 적용 대상으로 삼는다는 것이다. 인공지능 시스템 자체는 배제된다. 이는 인공지능의 법인격성을 둘러싼 논의가 인공지능법안의 적용 대상에 관해서는 불필요하다는 점을 보여준다.⁴³⁾ 달리 말하면 인공지능이 지닌 위험을 관리해야 할 책임은 인공지능 시스템 그 자체가 아닌 이를 운용하는 인적 주체가 부담해야 함을 명확히 하는 것으로 이해할 수 있다.

43) 인공지능의 법인격성 논의에 관해서는 양천수, “현대 지능정보사회와 인격성의 확장”, 『동북아법연구』제12권 제1호(2018. 5), 1-26쪽 참고.

라. 인공지능의 유형화

1) 위험 기반 접근법

인공지능법안에서 흥미로운 부분은 잘 알려진 것처럼 인공지능을 유형화한다는 것이다. 이때 유형화의 기준이 되는 것은 바로 인공지능이 지닌 위험이다. 그 점에서 인공지능법안은 위험 기반 규제 접근법(risk based regulatory approach)을 사용한다. 여기서 위험의 기준이 되는 것은 바로 자연인(natural person), 즉 인간 존재의 건강, 안전, 기본권에 대한 침해 가능성이다.

앞에서도 살펴본 것처럼 인공지능법안이 위험 기반 접근법을 취한다는 점은 인공지능에 대한 규제를 사전적으로(ex ante), 즉 미래 지향적으로 앞당긴다는 점을 의미한다. 법적 규제로 인공지능의 위험을 사전에 관리하겠다는 것이다. 그 점에서 인공지능법안은 인공지능에 대한 규제의 적용 시기를 앞당긴다. 그만큼 강한 규제임을 뜻한다.

2) 위험 및 인공지능의 유형화

이러한 위험 기반 접근법에 따라 인공지능법안은 위험을 크게 세 가지로 구별한다. 허용할 수 없는 위험(unacceptable risk), 고위험(high risk), 저위험(low risk) 또는 최소 위험(minimal risk)이 그것이다.⁴⁴⁾ 이 외에 '제한된 위험'(limited risk)이 독자적인 유형으로 제시되기도 한다.⁴⁵⁾ 이러한 구별은 인공지능이 창출하는 위험을 크게 '고위험/저위험'으로 구별하는 데서 출발한다. 그중 고위험을 '허용할 수 없는 위험'과 '고위험'으로 다시 나누고 인공지능의 작동 방식에 따라 저위험 가운데 제한된 위험을 구별하는 것이다.

인공지능법안은 이러한 위험에 대응하여 인공지능을 크게 다음과 같이 유형화한다. 금지되는 인공지능, 고위험 인공지능, 제한된 위험의 인공지능, 저위험 또는 최소 위험 인공지능이 그것이다. 여기서 제한된 위험의 인공지능은 특정한 방식으로 제한된 위험을 창출하는 그 때문에 투명성 규제가 적용되는 인공지능을 뜻한다.

이때 주의해야 할 점은 인공지능을 유형화하는 방식은 인공지능 시스템 그 자체를 기준으로 하는 게 아니라는 점이다. 제안 설명(Explanatory Memorandum)에도 나와 있듯이 인공지능의 유형화는 이른바 기능

44) Proposal for AI Act, p. 12.

45) 김한균, "고위험 인공지능에 대한 가치지향적 · 위험평가기반 형사정책", 『형사정책』 제34권 제1호(2022. 4), 17쪽 참고. '제한된 위험'은 유럽연합 집행위원회 법안에서는 정면에서 규정되지 않은 것으로 보인다. 이는 박해성 · 김법연 · 권현영, "인공지능 통제를 위한 규제의 동향과 시사점", 『정보법학』 제25권 제2호(2021. 8), 15-17쪽 참고. 그렇지만 고위험 인공지능이 아니면서도 투명성 요건이 적용되는 '특정한 인공지능 시스템'(certain AI systems)(제52조)과 관련하여 이러한 제한된 위험이 전제된 것으로 보인다. 말하자면 이러한 특정한 인공지능 시스템이 제한된 위험의 인공지능이라는 것이다. 이후 제한된 위험은 특히 유럽의회 차원에서 인공지능법안이 검토될 때 인공지능이 지닌 위험의 유형으로 자리매김한다.

주의에 바탕을 둔다. 인공지능이 어떻게 사용되고 실행되는지에 따라 인공지능을 유형화하는 것이다.⁴⁶⁾ 따라서 인공지능을 유형화한 후 이에 맞추어 규제를 투입하는 방식은 인공지능이 기능적 관점에서 어떻게 사용되는지에 따라 규제하는 것으로 이해할 수 있다.

3) 금지되는 인공지능 실행

이러한 인공지능 유형화에 따를 때 우선 허용할 수 없는 위험을 창출하는 인공지능 실행은 금지된다. 여기서 허용할 수 없는 위험이란 자연인인 인간 존재의 건강, 안전, 권리를 명백하게 위협하는 위험을 말한다.⁴⁷⁾ 특히 유럽연합이 강조하는 규범적 가치인 인간의 존엄을 명백하게 침해할 우려가 있는 위험이 이에 해당한다. 인공지능법안은 이러한 허용할 수 없는 위험을 인공지능이 창출하는 경우로 네 가지를 규정한다.⁴⁸⁾

첫째는 인공지능으로 잠재의식에 영향을 미침으로써 사람(person)의 행동(behaviour)을 중대하게 왜곡하는 것이다(제5조 제1항 (a)).

둘째는 인공지능으로 사람이 속한 집단의 취약점을 활용함으로써 사람의 행동을 중대하게 왜곡하는 것이다(제5조 제1항 (b)).

셋째는 공공기관 등이 인공지능을 활용하여 사람의 사회적 행동이나 성격 등을 사회적으로 평가함으로써 그 사람이나 그가 속한 집단에 불리한 처우를 하는 것이다(제5조 제1항 (c)).

넷째는 법을 집행하기 위한 목적으로 공적으로 접근할 수 있는 공간에서 '실시간' 원격 생체정보 활용 신원확인 시스템('real-time' remote biometric identification systems), 가령 안면 인식 시스템을 활용하는 것이다(제5조 제1항 (d)). 다만 여기에는 세 가지 예외가 있다. 첫째는 아동을 포함한 잠재적 범죄 피해를 추적하기 위해 활용하는 것이다. 둘째는 중대하고 임박한 자연인의 생명 또는 안전 침해 위협이나 테러 위협을 예방하기 위해 활용하는 것이다. 셋째는 유럽연합 법규범이 규정하는 중범죄자를 수사 및 기소하기 위해 활용하는 것이다(제5조 제1항 (d) (i)-(iii)).⁴⁹⁾

여기서 확인할 수 있듯이 금지되는 인공지능 실행은 인간의 존엄성을 침해할 수 있는 그래서 허용할 수

46) Proposal for AI Act, p. 12.

47) 김한균, 앞의 논문, 17쪽.

48) 이에 관한 소개는 김송옥, 앞의 논문, 122-123쪽 참고.

49) 그러나 이러한 예외를 인정할지에 논란이 전개되었고 결국 유럽의회는 이를 부정하는 쪽으로 인공지능법안을 수정하였다. 아래 IV.2. 참고. 이를 비판하는 문헌으로는 Marietje Schaake, "The European Commission's Artificial Intelligence Act", Human Centered Artificial Intelligence (Stanford University, 2021. 6), p. 3 참고. 이 문헌은 (https://hai.stanford.edu/sites/default/files/2021-06/HAI_Issue-Brief_The-European-Commissions-Artificial-Intelligence-Act.pdf)에서 확인할 수 있다(검색 일자: 2024. 2. 7.).

없는 위험을 창출한다. 인간의 존엄성을 구성하는 본질적인 징표가 무엇인지는 법적으로 판단하기 어려운 문제이지만 이에 관해 크게 두 가지가 제시되는 편이다. 인간 존재의 자율성을 침해하는 것과 인간 존재를 철저하게 도구화하는 것이 그것이다. 이는 서로 연결된다. 이에 더해 개인정보와 개인의 내밀한 영역의 불가침성이 강조되는 오늘날에는 개인정보를 남김없이 수집하여 해당 개인을 포괄적으로 프로파일링하는 것도 인간의 존엄성을 침해하는 예로 포섭할 수 있을 것이다.⁵⁰⁾ 금지되는 인공지능 실행으로 인공지능법안이 규정하는 것은 바로 이렇게 개인의 자율성을 중대하게 침해하거나 그를 철저하게 수단화하는 것 또는 개인을 철저하게 프로파일링하는 것에 해당한다는 점이 이를 시사한다.

한편 흥미로운 점은 법 집행 기관이 법을 집행하기 위해 '실시간' 원격 생체 정보 활용 신원 확인 시스템, 가령 안면 인식 기능을 수행하는 인공지능을 활용하는 것도 금지되는 인공지능 실행으로 규정한다는 것이다. 이는 법 집행이라는 이름 아래 형사 사법 기관이 이러한 인공지능을 활용해 수사나 기소 등을 하는 것도 인간 존재를 수단화할 수 있다는 문제의식에 바탕을 두는 것으로 보인다. 하지만 앞에서 언급하였듯이 여기에는 세 가지 예외가 있다. 첫 번째 예외가 범죄 피해자 보호를 위한 것이라면, 두 번째 예외는 경찰 활동 그리고 세 번째 예외는 형사 사법 활동과 관련을 맺는다. 이러한 예외는 유럽의회에서 문제가 되었다.

4) 고위험 인공지능

고위험 인공지능은 인공지능법안이 가장 중점적으로 규율하는 인공지능 유형이다. 고위험 인공지능이 무엇인지에 관해 일단 인공지능법안은 자연인의 건강, 안전 또는 기본권에 고위험을 창출하는 인공지능을 고위험 인공지능으로 규정한다(제7조). 하지만 이 같은 고위험 인공지능을 법규범에서 정면으로 상세하게 규정하는 것은 어렵다. 이는 법규범의 현실 적응력을 떨어뜨린다. 이에 인공지능법안은 두 가지 입법 기술을 활용한다. 부속서(Annex II, III) 연결, 부속서 업데이트 권한이 그것이다.

가) 부속서와 연결하여 판단

첫째, 부속서 II(Annex II)와 부속서 III(Annex III)에 기재된 인공지능 시스템은 고위험 인공지능으로 고려된다(제6조 제1항 및 제2항).⁵¹⁾

부속서 II는 “유럽연합 조화 입법 리스트”(List of Union Harmonisation Legislation)를 규정하는 것으로 크게 Section A와 Section B로 구성된다. “새로운 입법 프레임워크에 바탕을 둔 유럽연합 조화 입법 리스트”(List of Union harmonisation legislation based on the New Legislative Framework)에 관한 Section A는 “기계, 완구, 레저용 선박 및 개인용 선박, 승강기 및 승강기용 안전 요소, 폭발 가능성

50) 이 문제에 관해서는 양천수, 『빅데이터와 인권: 빅데이터와 인권의 실제적 조화를 위한 법정정책 방안』(영남대학교출판부, 2016) 참고.
51) 부속서 II와 III은 (https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC_2&format=PDF)에서 확인할 수 있다(검색 일자: 2024. 2. 7.).

환경에 대비한 장비 및 보호 시스템, 무선 장비, 압력 장치, 케이블카, 개인 보호 장비, 가스 연소 기기, 의료 기기, 체외 진단 의료 기기” 등을 규정한다.⁵²⁾ “기타 유럽연합 조화 입법 리스트”(List of other Union harmonisation legislation)에 관한 Section B는 “민간 항공 보안, 2륜 및 4륜 차량, 농업 또는 임업용 차량, 철도 시스템, 원동기 차량과 트레일러 및 관련 시스템과 구성 요소, 무인 항공기 및 그 엔진과 프로펠러 및 원격제어 부품 및 장비” 등을 규정한다.⁵³⁾

여기서 다음과 같은 점을 확인할 수 있다. Section A는 주로 물리적인 현실 세계(real world)에서 위험을 야기할 수 있는 물건과 관련을 맺는다는 것이다. 더불어 Section B는 모빌리티(mobility)를 규정하는데 이 역시 물리적인 현실 세계에서 인간 존재에 위험을 야기할 수 있다는 것이다. 그 점에서 부속서 II가 규정하는 물건들은 고위험과 연결된다. 따라서 이러한 물건들에 설치되어 사용되는 인공지능 시스템 역시 고위험과 연결된다고 말할 수 있다.

“인공지능법안 제6조 제2항이 언급하는 고위험 인공지능 시스템”(HIGH-RISK AI SYSTEMS REFERRED TO IN ARTICLE 6(2))에 관한 부속서 III(Annex III)은 고위험 인공지능 시스템에 해당하는 자립형 인공지능의 영역을 다음과 같이 규정한다. “1. 자연인의 생체 인식 및 분류, 2. 주요 기반 시설의 관리 및 운영(도로교통의 관리 및 운영, 물, 가스, 난방, 전기 공급의 안전 요소), 3. 교육 및 직업 훈련, 4. 고용, 근로자 관리, 자영업에 대한 접근, 5. 필수적인 민간 및 공공 서비스에 대한 접근 및 향유, 6. 법 집행, 7. 이주, 망명, 국경 통제 관리, 8. 사법 및 민주주의 절차 관리” 등이 그것이다.⁵⁴⁾

이때 다음과 같은 특징을 발견할 수 있다. 부속서 II는 주로 그 자체 위험성을 가진 물건들을 규정하는 데 반해, 부속서 III은 자립형 인공지능이 수행하는 기능 영역이 지닌 위험성을 고려하여 규정한다는 점이다. 달리 말해 부속서 II는 설치형 인공지능, 부속서 III은 자립형 인공지능을 규정한다는 것이다.

나) 부속서 업데이트 권한

둘째, 인공지능법안은 유럽연합 집행위원회에 부속서 III을 업데이트할 수 있는 권한을 위임한다. 인공지능 기술은 우리의 예상을 매년 뛰어넘을 정도로 급속하게 발전하고 있다는 점, 이를 고려하지 않고 고위험 인공지능의 목록을 법이나 부속서에 고정적으로 규정하면 인공지능의 발전 속도에 법규범이 적절하게 대처하지 못한다는 점을 감안할 때 이는 적절한 입법 기술로 판단된다.

이때 집행위원회는 다음과 같은 단계를 거쳐 부속서 III을 업데이트할 수 있다. 우선 평가 대상이 되는 인공지능 시스템이 부속서 III 제1호에서 제8호까지 규정하는 영역에 사용되는지 검토한다(제7조 제1항 (a)).

52) 홍석한, 앞의 논문, 253쪽 각주(11)에서 인용하면서 약간의 수정을 가하였다.

53) 홍석한, 앞의 논문, 253쪽 각주(11)에서 인용.

54) 홍석한, 앞의 논문, 254쪽 각주(12)에서 인용.

다음 평가 대상 인공지능이 부속서 III이 기재하는 고위험 인공지능 시스템 이상으로 건강과 안전 또는 기본 권에 위험을 초래하는지 판단한다(제7조 제1항 (b)). 나아가 이렇게 위험을 평가할 때 인공지능법안 제7조 제2항 (a)~(h)호에서 규정하는 여러 요소, 즉 인공지능 시스템이 의도하는 목적이나 사용되는 범위, 인공지능 시스템이 지닌 위험이 미칠 잠재적 범위 등을 함께 고려한다(제7조 제2항).⁵⁵⁾ 요컨대 평가 대상이 되는 인공지능 시스템이 고위험을 창출하는지를 판단할 때 여러 요소를 종합적으로 검토하는 것이다.

5) 제한된 위험의 인공지능

인공지능법안 제52조는 “특정한 인공지능 시스템에 적용되는 투명성 의무”(Transparency obligations for certain AI systems)를 규정한다. 여기서 투명성 의무가 적용되는 인공지능이 제한된 위험의 인공지능에 해당한다.⁵⁶⁾ 인공지능법안 제52조는 이러한 인공지능으로 자연인과 상호작용을 의도하는 인공지능(이른바 챗봇), 감정 인식 시스템(emotion recognition system), 생체 정보 인식 시스템(biometric categorisation system), 이른바 ‘딥페이크’를 생성하는 생성형 인공지능을 규정한다.⁵⁷⁾ 이러한 인공지능은 고위험 인공지능으로 분류되지는 않지만, 인공지능이 작동하는 과정에 자연인이 개입하고 이로 인해 자연인에게 제한된 위험을 창출할 수 있기에 인공지능법안이 별도로 규정하는 것이다. 이러한 인공지능에는 고위험 인공지능처럼 투명성 규제가 적용된다.

6) 저위험 또는 최소 위험 인공지능

금지되는 인공지능, 고위험 인공지능, 제한된 위험의 인공지능을 제외한 나머지 인공지능은 저위험 또는 최소 위험 인공지능에 해당한다. 이러한 인공지능에는 기본적으로 행동 강령(codes of conduct)을 활용한 자율규제가 적용된다.

마. 인공지능 규제 방식

1) 규제 모델

인공지능은 어떻게 규제할 것인가? 이에 관해 인공지능법안 설명(Explanatory Memorandum)은 규제 개입의 강도에 따라 다섯 가지 선택지(options)를 고려한다.⁵⁸⁾ 이는 다음 <표 2>와 같다.

55) 이에 관한 상세한 내용은 김광수, “인공지능 알고리즘 규율을 위한 법적 동향: 미국과 EU 인공지능법의 비교를 중심으로”, 『행정법연구』 제70호(2023. 3), 189쪽 참고.

56) 김한균, 앞의 논문, 19쪽.

57) 이때 말하는 생성형 인공지능은 GAN(Generative Adversarial Networks) 기술이 적용된 인공지능을 말한다. 언어를 생성하는 GPT 계열의 인공지능은 여기서 간과되고 있다.

58) Proposal for AI Act, p. 9. 이를 소개하는 문헌으로는 김광수, 앞의 논문, 188쪽.

표 2 집행위원회 인공지능법안이 고려한 규제 모델

선택지	규제 모델
선택지 1	자발적 등급 표시 방식(voluntary labelling scheme)에 바탕을 둔 EU 입법 조치
선택지 2	영역 중심적, 특정 목적 중심적 접근
선택지 3	비례적인 위험 기반 접근법에 바탕을 둔 포괄적인(horizontal) EU 입법 조치
선택지 3+	비례적인 위험 기반 접근법에 바탕을 둔 포괄적인 EU 입법 조치 + 비고위험 인공지능 시스템(non-high-risk AI systems)에 관한 자율 규범(codes of conduct)
선택지 4	인공지능이 산출하는 위험에 상관없이 모든 인공지능 시스템에 적용되는 강제적 요건을 신설하는 포괄적인 EU 입법 조치

* 출처: 인공지능법안, p. 9 및 김광수, 188쪽을 참고하여 재구성

위 선택지는 규제 대상과 개입 강도를 조합한 것이다. 이에 따르면 선택지 1이 자율규제 방식에 가장 가깝다면, 선택지 4는 가장 포괄적이면서 개입 강도가 강한 타율규제 방식에 해당한다. 인공지능법안은 이 가운데 선택지 3+를 선택하였다. 모든 유형의 인공지능을 포괄하면서 자율규제와 타율규제를 종합한 규제법안을 만든 것이다.

인공지능법안은 ‘규정’(regulation)이라는 구속력이 가장 강한 형식을 이용하여 선택지 3+를 채택한다. 그 점에서 원칙적으로 포괄적이면서 강력한 타율규제 방식을 취한다고 말할 수 있다. 하지만 그렇다고 해서 인공지능법안은 민법상 불법 행위나 형법상 범죄 구성 요건처럼 금지 방식의 규제만을 선택하지는 않는다. 인공지능의 유형에 맞게 차별화된 강도의 규제를 선택해 적용하기 때문이다. 그 때문에 저위험 또는 최소 위험 인공지능에는 자율 규범인 행동 강령을 활용한 자율규제를 적용한다. 그 점에서 인공지능법안은 경성 규범과 연성 규범을 혼합한 규제 모델이라 할 수 있다.

2) 포괄적 규제

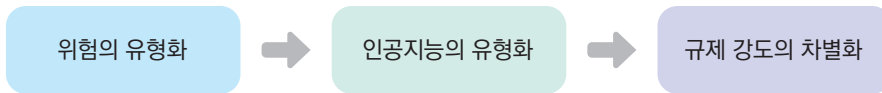
위에서 검토한 것처럼 인공지능법안은 포괄적 규제(horizontal regulation)를 채택한다. 이때 포괄적 규제란 인공지능 영역 전반에 걸쳐 투입되는 규제를 말한다. 그 점에서 포괄적 규제는 특정 영역에 한정되어 적용되는 영역적 규제(vertical regulation)와 구별된다.⁵⁹⁾ 인공지능법안은 이러한 포괄적 규제를 활용함으로써 유럽연합 역내에서 사용되는 인공지능 시스템 전반을 규제 및 관리하는 포괄적인 프레임워크(framework)를 제공한다.

59) 이에 관해서는 (<https://www.holistica.com/blog/regulating-ai-the-horizontal-vs-vertical-approach>) 참고(검색 일자: 2024. 2. 1.).

3) 차별화된 규제

인공지능법안은 적용 영역 면에서는 포괄적이지만 그렇다고 해서 개입 강도가 획일적인 규제를 취하지는 않는다. 인공지능법안은 규제 대상의 특성을 고려한 차별화된 규제 방식을 선택한다. 이때 규제 대상의 특성이란 인공지능이 창출하는 위험을 말한다. 요컨대 인공지능의 유형에 따라 규제 강도가 차별화되는 것이다. 이 점에서 인공지능 유형에 따라 차별화된 규제 강도를 도식으로 나타내면 [그림 1]과 같다.

[그림 1] 차별화된 규제의 구조



* 출처: 직접 작성

바. 규제 수단

인공지능법안은 각 유형의 인공지능에 다음과 같은 규제 수단을 투입한다.

1) 금지된 인공지능 실행

인공지능법안이 금지하는 인공지능 실행을 하는 것은 마치 민법상 불법행위를 저지르는 것과 같다. 이는 법질서가 허용하지 않는 위법한 작동이기에 이에겐 행정벌(administrative fines)이 부과된다(제71조 제3항 (a)).

2) 고위험 인공지능

인공지능법안이 가장 중점을 두는 고위험 인공지능에는 다음과 같은 규제 수단이 적용된다. 고위험 인공지능 시스템이 필수적으로 갖추어야 하는 요건을 설정하고(제3편 제2장), 고위험 인공지능 시스템 제공자의 의무(제3편 제3장), 수입업자의 의무(제3편 제3장), 배포업자(distributor)의 의무(제3편 제3장), 사용자 및 제3자의 의무(제3편 제3장)를 규정한다.

이는 다음과 같이 분석할 수 있다. 고위험 인공지능 시스템에 투입되는 규제 수단은 크게 두 가지로 구별된다. 인공지능 시스템 자체에 적용되는 규제 수단과 인공지능 시스템과 관련을 맺는 이해관계자에 적용되는 규제 수단이 그것이다. 그중 후자에서는 인공지능 시스템의 가치 사슬(value chain)이 고려된다. 이에 따라 인공지능 시스템 공급자 → 수입업자 → 시장 배포자 → 사용자 및 제3자라는 당사자가 도출된다. 인공지능법안은 이러한 당사자를 모두 규제 대상으로 포섭하여 각각 규제 수단을 마련한다.

이 가운데 고위험 인공지능 시스템 자체에 투입되는 규제 수단, 즉 반드시 갖추어야 하는 요건으로는 위험 관리 시스템(risk management system)(제9조), 데이터 및 데이터 거버넌스(제10조), 기술 문서

(technical documentation)(제11조), 기록 유지(record-keeping)(제12조), 사용자를 위한 투명성 및 정보 규정(제13조), 인간에 의한 감독(human oversight)(제14조), 정확함, 견고함 및 사이버 보안(accuracy, robustness and cybersecurity)(제15조)이 제시된다.

고위험 인공지능 시스템 관련자, 그중에서도 공급자가 부담해야 하는 의무 가운데 주목할 만한 것으로는 품질 관리 체계(quality management system) 마련(제17조), 기술 문서 작성(제18조), 적합성 평가 수행(제19조), 자동 생성 로그 기록 보관(제20조), 시정 조치 수행(제21조), 정보 제공 의무(제22조), 관할 당국과 협력(cooperation with competent authorities), CE 인증 마크 부착(제16조 (i)) 등을 언급할 수 있다.⁶⁰⁾

3) 제한된 위험의 인공지능

제한된 위험의 인공지능, 즉 자연인과 상호 작용을 의도하는 인공지능이나 감정 인식 시스템 등에는 고위험 인공지능에 사용되는 것과 유사한 투명성 규제가 적용된다.

4) 저위험 또는 최저 위험 인공지능

저위험 또는 최저 위험 인공지능은 타율규제 대상은 아니다. 그렇지만 앞에서 본 것처럼 인공지능법안이 이를 규제에서 완전히 배제하는 것은 아니다. 이들 인공지능에는 자발적인 행동 강령 수립을 활용한 자율규제가 적용된다. 이때 행동 강령에는 고위험 인공지능에 적용되는 규제 방법이나 강도가 수용된다.⁶¹⁾ 엄격한 내용의 규제이지만 이를 자율적으로 준수하도록 하는 것이다.

사. 혁신 지원을 위한 수단

인공지능법안은 인공지능에 대한 규제, 즉 자유 제한을 뜻하는 좁은 의미의 규제만을 규율하지는 않는다. “인공지능에 관한 조화로운 규칙”(Harmonised Rules on Artificial Intelligence)라는 명칭이 시사하듯이 인공지능법안은 규제뿐만 아니라 인공지능 기술 혁신을 지원하는 수단도 담는다. 이를 예증하는 것이 법안 제7편이다. 예를 들어 제7편 제53조는 인공지능에 대한 규제 샌드박스(regulatory sandboxes)를 규정한 다. 인공지능 개발에 관해 규제 샌드박스를 적용할 수 있도록 한 것이다.

아. 거버넌스

인공지능법안은 제6편에서 거버넌스를 규정한다.⁶²⁾ 이에 따르면 유럽 인공지능위원회(제1장)와 국가 관할 당국(제2장)이 거버넌스로 제시된다. 그러나 인공지능법안 전체를 관찰하면 인공지능 규제와 관련된 거

60) 인공지능법안은 당사자 가운데 고위험 인공지능 시스템 공급자의 의무를 강조한다. 이에 따라 제16조에서 공급자가 준수해야 하는 의무를 열거하면서 제17조에서 제23조에 걸쳐 각 의무를 상세하게 규정한다.

61) 김한균, 앞의 논문, 19쪽.

62) 이에 관해서는 홍석한, 앞의 논문, 259쪽 아래; 정소영, “유럽연합 인공지능법안의 거버넌스 분석: 유럽인공지능위원회와 회원국 감독기관의 역할과 기능을 중심으로”, 『연세법학』 제39권(2022. 7), 57쪽 아래 참고.

버넌스가 이에 한정되지 않는다는 점을 확인할 수 있다. 이외에도 다양한 거버넌스를 언급할 수 있다. 인공지능법이 규율하는 거버넌스는 다음과 같이 구별할 수 있다. 인증 관련 거버넌스, 회원국 거버넌스 및 유럽연합 거버넌스가 그것이다.

1) 인증 관련 거버넌스

인증(notification) 관련 거버넌스로는 크게 인증 당국(notifying authority), 적합성 평가 기구(conformity assessment body), 피인증 기구(notified body)를 언급할 수 있다. 여기서 인증 당국은 적합성 평가와 관련된 절차 전반을 관할하는 거버넌스로 회원국에 설치된 당국을 말한다(제3조 제19호). 다음 적합성 평가 기구는 고위험 인공지능의 적합성 평가를 수행하는 기구를 말한다(제3조 제21호). 나아가 피인증 기구는 인증 당국에 의해 인증을 받고 인공지능법에 따라 적합성 평가를 수행하는 기구를 말한다(제3조 제22호).

2) 회원국 거버넌스

유럽연합 회원국 안에 설치되는 거버넌스로 국가 관할 당국(national competent authority)을 언급할 수 있다. 국가 관할 당국은 국가 감독 당국(national supervisory authority), 인증 당국 및 시장 감시 당국(market surveillance authority)으로 구성된다(제3조 제43호). 여기서 국가 감독 당국은 회원국 내에서 인공지능법의 집행과 적용을 책임지는 당국을 말한다(제3조 제42호). 시장 감시 당국은 인공지능 시스템을 시장에 출시한 후 이를 모니터링하고 시장을 감시하는 당국을 말한다(제61조-제68조). 인공지능법에 따르면 원칙적으로 국가 감독 당국이 인증 당국과 시장 감시 당국의 역할을 수행한다(제59조 제2항).

3) 유럽연합 거버넌스

유럽연합 차원의 거버넌스로 유럽연합 집행위원회와 유럽 인공지능위원회를 거론할 수 있다. 집행위원회는 유럽연합 차원에서 인공지능법이 집행 및 적용되는 것을 총괄한다. 이때 유럽 인공지능위원회는 인공지능법 집행 및 적용에 관해 집행위원회와 국가 감독 당국 사이의 협력과 조정에 기여하기 위해 집행위원회의 자문에 응하고 이를 보좌하는 역할을 한다(제56조 제2항). 유럽 인공지능위원회는 국가 감독 당국과 유럽 데이터 보호 감독관(European Data Protection Supervisor)으로 구성된다(제57조 제1항). 위원장은 집행위원회가 맡는다(제57조 제3항).

자. 시장 출시 후 감시

인공지능법이 지닌 특징적인 면 가운데 하나는 시장 출시 후 감시 등을 규정한다는 점이다. 이를테면 인공지능 시스템을 시장에 출시한 이후에도 이에 관한 모니터링, 정보 공유 및 감시를 강화한다. 이를 위해 인공지능법은 위에서 언급한 시장 감시 당국이라는 거버넌스도 설치한다. 이는 의약품 규제에서 강조되는 ‘사려 깊은 경계’(prudent vigilance) 모델을 받아들인 것이다. 인공지능이 시장에 출시된 이후 규제가 종료되지는 않는다는 점이다.

2. 유럽의회 수정안

가. 배경

유럽연합 집행위원회가 제안한 인공지능법안은 2023년 6월 14일 유럽의회를 수정 통과하였다.⁶³⁾ 찬성 499표, 반대 28표, 기권 93표로 가결되었다. 이때 주시해야 할 부분은 집행위원회 원안이 유럽의회를 통과 하면서 상당 부분 수정 및 보완이 이루어졌다는 점이다. 이의 가장 중요한 계기로 두 가지를 들 수 있다.

우선 2022년 11월 Open AI에 의해 ChatGPT라는 혁신적인 생성형 인공지능이 출시되었다는 점이다. ChatGPT가 전 세계적인 열풍이 일으키면서 이에 관해 크게 두 가지 문제가 제기되었다. 첫째, ChatGPT가 환각(hallucination)을 유발한다는 점이다. 둘째, ChatGPT 등 생성형 인공지능이 학습을 위해 활용하는 데이터의 저작권 및 개인정보 침해 이슈가 제기되었다는 점이다. 집행위원회 원안은 이를 충분히 고려하지 않았기에 유럽의회 차원에서 이를 보완해야 할 필요가 있었다.

나아가 집행위원회 원안은 금지되는 인공지능 실행에 세 가지 예외를 인정하는데 이에 문제가 제기되었다. 무엇보다도 법 집행 기구가 중대한 위험 예방 등을 이유로 실시간 안면 인식과 같은 ‘실시간 원격 생체 정보 활용 신원 확인 시스템’을 활용할 수 있게 한 점이 문제가 되었다. 이는 시민을 법 집행을 위한 수단으로 전락시킬 수 있다는 것이다.⁶⁴⁾

유럽의회는 이러한 점에 주목하여 집행위원회가 제안한 인공지능법안을 검토 및 수정하여 통과시켰다.

나. 주요 수정 내용

유럽의회가 수정한 내용 가운데 중요한 점 몇 가지를 언급하면 다음과 같다.⁶⁵⁾

먼저 인공지능법안의 목적을 다소 수정하였다. 집행위원회 원안에 비해 인간중심적이면서 신뢰할 수 있는 인공지능을 더욱 강조하였다. 이에 연결하여 인공지능이 유발하는 위험으로부터 건강과 안전, 기본 권, 민주주의 및 법의 지배 그리고 환경을 보호할 것을 강조하였다. 물론 이와 더불어 유럽연합 역내의 인

63) 정병일, “인공지능 규제 법안 유럽의회 통과”, 『AI타임스』(2023. 6. 15).

64) 형사사법 영역에서 인공지능을 활용하는 것에 관한 유럽의회 태도에 관해서는 European Parliament resolution of 6 October 2021 on artificial intelligence in criminal law and its use by the police and judicial authorities in criminal matters 참고(https://www.europarl.europa.eu/doceo/document/TA-9-2021-0405_EN.html)(검색 일자: 2024. 2. 3.). 이에 관해서는 정소영, “형사사법에서의 인공지능 사용에 대한 유럽의회 결의안: ‘인공지능에 의한 결정’ 금지에 관하여”, 『형사정책』 제34권 제2호(2022. 7), 137-172쪽 참고.

65) 이에 관해서는(<https://www.europarl.europa.eu/news/en/headlines/society/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>) 참고(검색 일자: 2024. 2. 3.). 또한 Tambiama Madiega/Samy Chahri, “Artificial intelligence act”, European Parliamentary Research Service (2023) 참고([https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI\(2021\)698792_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI(2021)698792_EN.pdf))(검색 일자: 2024. 2. 3.).

공기능 시장 혁신을 지원하고 시장 기능을 향상시키는 것도 목적으로 설정하기는 한다.⁶⁶⁾ 하지만 규제와 혁신을 조화롭게 하고자 한 집행위원회 원안보다 규제 방향으로 무게 중심이 더욱 이동했다는 점은 부정할 수 없다.

다음 ChatGPT 출현에 대응하여 인공지능 유형을 새롭게 추가하였다. 물론 집행위원회 원안에서 바탕이 된 구별, 즉 허용할 수 없는 위험/고위험/제한된 위험/저위험 또는 최소 위험이라는 위험 구별과 이에 연결된 금지되는 인공지능 실행/고위험 인공지능/제한된 위험의 인공지능/저위험 또는 최소 위험 인공지능이라는 구별은 유지된다.⁶⁷⁾ 유럽의회는 이에 추가하여 범용 목적의 생성형 인공지능(general purpose and generative AI), 즉 '파운데이션 모델'(foundation model)을 새롭게 설정하였다.⁶⁸⁾ 이에 따라 파운데이션 모델 제공자에 새롭게 여러 의무를 부과하였다. 예를 들어 파운데이션 모델 제공자는 고위험 인공지능이나 제한된 위험의 인공지능에 요구되는 투명성 요건을 충족해야 한다. 나아가 파운데이션 모델이 생성하는 내용이 무엇인지 밝혀야 하고 불법적인 내용을 생성하지 못하도록 막아야 한다. 또한 가장 논란이 되는 의무로서 파운데이션 모델이 저작권이 있는 데이터를 학습용으로 사용하는 때에는 이를 요약해서 공개적으로 제공해야 한다.

나아가 법을 집행하기 위한 목적으로 실시간 안전인식 기술과 같은 '실시간 원격 생체 정보 활용 신원 확인 시스템'을 활용할 때 이를 예외적으로 인정하던 경우를 삭제하였다.⁶⁹⁾ 오직 예외적인 경우, 가령 중범죄를 기소하기 위해 그리고 법원의 승인 아래에서만 그것도 '실시간'(real-time)이 아닌 '사후적인'(post) 원격 생체 정보 활용 신원 확인 시스템만을 사용할 수 있을 뿐이다.

같은 맥락에서 범행 전력이나 위치 등을 포함해 개인의 성격이나 특성 등을 활용하여 범죄 행위나 재범 위험성 등을 평가하여 사전에 범죄를 예방하고자 하는 예측적 치안(predictive policing)도 금지되는 인공지능 실행에 새롭게 도포되었다.⁷⁰⁾ 요컨대 경찰 및 형사사법 영역과 관련해 금지되는 인공지능 실행의 범위를 확대한 것이다.⁷¹⁾

66) Amendment 3: "The purpose of this Regulation is to promote the uptake of human centric and trustworthy artificial intelligence and to ensure a high level of protection of health, safety, fundamental rights, democracy and rule of law and the environment from harmful effects of artificial intelligence systems in the Union while supporting innovation and improving the functioning of the internal market." 수정안 원문은 (https://www.europarl.europa.eu/doceo/document/TA-9-2023-0236_EN.html) 참고(검색 일자: 2024. 2. 3.).

67) Tambiama Madiaga/Samy Chahri, "Artificial intelligence act", European Parliamentary Research Service (2023), p. 4.

68) Amendment 399: Article 28 b Obligations of the provider of a foundation mode 참고.

69) Amendment 220.

70) Amendment 224: "(d a) the placing on the market, putting into service or use of an AI system for making risk assessments of natural persons or groups thereof in order to assess the risk of a natural person for offending or reoffending or for predicting the occurrence or reoccurrence of an actual or potential criminal or administrative offence based on profiling of a natural person or on assessing personality traits and characteristics, including the person's location, or past criminal behaviour of natural persons or groups of natural persons;" 예측적 치안 문제에 관해서는 정채연, "예측적 치안과 형사정책: 쟁점과 과제", 『영남법학』 제57호(2023. 12), 67-92쪽 참고.

71) 형사사법에서 인공지능을 활용하는 문제에 관해서는 주현경, "형사절차에서 인공지능 위험평가 프로그램을 이용한 양형 판단", 『영남법학』 제57호(2023. 12), 35-66쪽 참고.

3. 인공지능법안 합의

가. 배경

2023년 12월 8일 유럽의회와 이사회는 인공지능법안에 최종 합의한다.⁷²⁾ 이로써 인공지능법안은 공식 제정을 눈앞에 두고 있다. 이러한 합의가 의미가 있는 이유는 집행위원회가 제안한 원안에 관해 그동안 유럽의회와 이사회가 각기 독자적인 수정안을 마련하였기 때문이다.⁷³⁾ 이는 그만큼 인공지능법안에 관해 팽팽한 견해 대립이 존재한다는 점을 시사한다. 무엇보다도 법을 집행하기 위한 목적으로 안면인식과 같은 실시간 원격 생체 정보 활용 신원 확인 시스템을 예외적으로 활용할 수 있는지에 관해 유럽의회와 이사회 및 집행위원회는 견해를 달리하였다. 그런데 이에 관해 합의가 이루어졌다는 점은 일종의 절충안이 성공적으로 마련되었다는 점을 의미한다.

나. 주요 내용

합의가 이루어진 내용 중에서 크게 두 가지를 언급할 필요가 있다.⁷⁴⁾

첫째, 유럽의회가 삭제하였던 실시간 원격 생체 정보 활용 신원 확인 시스템의 예외적 허용이 다시 그러나 집행위원회 원안보다 제한적으로 인정되었다는 점이다.

둘째, 유럽의회가 금지한 범죄 예측 기술 가운데 사람의 개인적 특성에 기반을 둔 범죄 예측 기술은 금지하되 지리적 데이터에 기반을 둔 범죄 예측 기술은 허용하기로 했다는 점이다.⁷⁵⁾

72) (<https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>) 참고(검색 일자: 2024. 1. 29.).

73) 2023년 6월 14일 유럽의회 수정안이 마련되기 전에 이미 이사회는 집행위원회 원안에 관해 2022년 12월 6일 이사회 수정안을 채택하였다. 이에 관해서는 한국지능정보사회진흥원, “EU 인공지능법 입법현황과 시사점”, 『지능정보사회 법제도 이슈리포트』(2023-03), 3쪽 참고.

74) (<https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>)(검색 일자: 2024. 1. 29.).

75) 정채연, 앞의 논문, 85쪽.

V

시사점 - 결론을 대신하여

최종 합의에 이른 유럽연합 인공지능법안이 실제로 제정되면 이는 전 세계 최초의 포괄적인 인공지능 규제법이 된다. 따라서 유럽연합 인공지능법은 이후 제정될 전 세계의 인공지능법에 중요한 규범적 기준이 될 것이다.⁷⁶⁾ 예를 들어 유럽연합 인공지능법안이 채택한 위험 기반 접근법, 위험에 따른 인공지능의 유형화, 유형에 맞춘 차별화된 규제, 파운데이션 모델 규제, 규제 수단으로서의 영향 평가와 투명성의 강조, 다원적으로 마련된 인공지능 거버넌스, 인공지능의 시장 출시 후 지속적 감시 등은 인공지능에 대한 규제를 설계할 때 유익한 참고가 될 것이다.

실제로 우리나라에서도 유럽연합의 인공지능법안을 모델로 한 「인공지능 산업 육성 및 신뢰 기반 조성에 관한 법률안」(이하 ‘인공지능 법률안’으로 약칭함)이 2023년 2월 14일 국회 과학기술방송정보통신위원회에서 위원회 대안으로 통과한 바 있다. 우리 인공지능 법률안은 유럽연합의 인공지능법안처럼 위반 기반 접근법을 활용해 인공지능을 유형화한다. 특히 고위험 인공지능에 관해 다음과 같은 규제 방안을 마련한다. 사전 고지 의무 부과, 신뢰성 확보 조치, 위험 관리 방안 수립, 최종 결과 도출 과정 설명, 이용자 보호 방안이 그것이다. 다만 유럽연합 인공지능법안이 도입한 금지되는 인공지능 실행은 우리 인공지능 법률안에서는 고려하지 않는다. 그 점에서 우리 인공지능 법률안은 유럽연합 인공지능법안에 비해 인공지능 혁신 지원을 좀 더 지향한다고 평가할 수 있다.

이처럼 유럽연합 인공지능법안이 인공지능을 규율하기 위해 도입한 규제 모델이나 원칙, 수단, 거버넌스 등은 앞으로 인공지능에 대한 규제법을 신설할 때, 더 나아가 현대 과학기술과 관련해 새롭게 발생하는 규범적 문제를 규율할 때 유익한 참고가 될 것이다.

76) 유럽연합 인공지능법의 의미에 관해서는 홍석한, 앞의 논문, 263쪽 아래 참고.

그렇지만 유럽연합 인공지능법안에 관해 집행위원회, 유럽의회, 이사회 사이에 여러 견해 대립이 있었다는 점을 고려하면 인공지능법안이 과연 장점만을 지니는지에는 의문을 제기할 수 있을 것이다. 예를 들어 인공지능법안이 과연 인공지능에 대한 규제와 혁신을 조화롭게 추구하는지에 의문을 제기할 수 있다. 무엇보다도 위험 기반 접근법에 따라 인공지능을 유형화하면서 금지되는 인공지능 실행을 설정한 것, 그것도 예외를 좀처럼 인정하지 않는 방식을 선택한 것에는 논란이 있을 수 있다. 이에는 다음과 같은 방식을 선택하는 게 더 적절하지 않았을까 진단해 본다. 첫째는 금지되는 인공지능 실행에도 '원칙/예외' 모델을 좀 더 폭넓게 인정하는 것이다. 둘째는 인공지능 유형화를 법에 명시하지 말고 이에 관한 판단 권한을 규제 당국에 위임하는 것이다. 이에 근거하여 규제 당국이 합리적인 사전 위험 평가를 활용해 개별 인공지능이 어떤 유형의 인공지능에 해당하는지를 그때그때 판단하게끔 하는 방법이 인공지능 기술의 급속한 발전에 현실 적합적으로 대처하는 방안이지 않을까 고민해 본다.

참고문헌

[국내 문헌]

김광수, “인공지능 알고리즘 규율을 위한 법제 동향: 미국과 EU 인공지능법의 비교를 중심으로”, 『행정법연구』제70호(2023. 3).

김송옥, “AI 법제의 최신 동향과 과제: 유럽연합(EU) 법제와의 비교를 중심으로”, 『공법학연구』제22권 제4호(2021. 11).

김중권, “EU 인공지능명령안의 주요 내용과 그 시사점”, 『헌법재판연구』제8권 제2호(2021. 12).

김한균, “고위험 인공지능에 대한 가치지향적 · 위험평가기반 형사정책”, 『형사정책』제34권 제1호(2022. 4).

박도현, “신뢰할 수 있는 인공지능의 이론적 고찰”, 『법철학연구』제25권 제2호(2022. 8).

박해성 · 김범연 · 권현영, “인공지능 통제를 위한 규제의 동향과 시사점”, 『정보법학』제25권 제2호(2021. 8).

심소연, “유럽연합(EU)의 「반도체법(Chips Act)」”, 『최신 외국입법정보』제234호(2023. 10. 31).

심우민, “알고리즘 투명성에 대한 규범적 접근방식”, 『디지털 윤리』제2권 제1호(2018. 6).

심재우, “형법에 있어서 결과불법과 행위불법”, 『법학논집』제20집(1982. 12).

양천수, 『빅데이터와 인권: 빅데이터와 인권의 실제적 조화를 위한 법정정책 방안』(영남대학교출판부, 2016).

양천수, 『인공지능 혁명과 법』(박영사, 2021).

양천수, 『책임과 법』(박영사, 2022).

양천수, “위험 · 재난 및 안전 개념에 대한 법이론적 고찰”, 『공법학연구』제16권 제2호(2015. 5).

- 양천수, “현대 지능정보사회와 인격성의 확장”, 『동북아법연구』제12권 제1호(2018. 5).
- 양천수, “인공지능과 윤리: 법철학의 관점에서”, 『법학논총』(조선대) 제27집 제1호(2020. 4).
- 양천수, “지능정보기술의 위험과 법적 대응 방안: 알고리즘에 대한 대응을 중심으로 하여”, 『법학연구』(충북대) 제32권 제1호(2021. 6).
- 양천수, “자율주행자동차에 관한 법정책의 방향: 의의 · 위험 · 규제”, 『모빌리티연구』제3권 제2호(2023. 9).
- 양천수 · 하민경, 『새로운 보안체계 구축에 관한 글로벌 규범 현황 및 법제 대응 방안 연구』(한국법제연구원, 2023).
- 유재홍 · 추형석 · 강송희, 『유럽(EU)의 인공지능 윤리 정책 현황과 시사점: 원칙에서 실천으로』, Issue Report IS-114(소프트웨어정책연구소, 2021. 3. 25).
- 이경선, “EU 인공지능 규제안의 주요 내용과 시사점”(정보통신정책연구원, 2021).
- 이병준, “유럽연합 디지털 서비스법을 통한 플랫폼 규제: 디지털 서비스법 초안의 주요내용과 입법방향을 중심으로”, 『소비자법연구』제7권 제2호(2021. 5).
- 이보연, “유럽연합의 인공지능 관련 입법 동향을 통해 본 시사점”, 『법학논문집』(중앙대) 제43집 제2호(2019. 8).
- 이시한, 『GPT 제너레이션: 챗GPT가 바꿀 우리 인류의 미래』(북로망스, 2023).
- 인공지능과 가치 연구회, 『인공지능윤리: 다원적 접근』(박영사, 2021).
- 정병일, “인공지능 규제 법안 유럽의회 통과”, 『시타임스』(2023. 6. 15).
- 정소영, “유럽연합 인공지능법안의 거버넌스 분석: 유럽인공지능위원회와 회원국 감독기관의 역할과 기능을 중심으로”, 『연세법학』제39권(2022. 7).
- 정소영, “형사사법에서의 인공지능 사용에 대한 유럽의회 결의안: ‘인공지능에 의한 결정’ 금지에 관하여”, 『형사정책』제34권 제2호(2022. 7).
- 정채연, “예측적 치안과 형사정책: 쟁점과 과제”, 『영남법학』제57호(2023. 12).
- 주현경, “형사절차에서 인공지능 위험평가 프로그램을 이용한 양형 판단”, 『영남법학』제57호(2023. 12).
- 한국지능정보사회진흥원 정책본부 지능화법제도팀, “EU 인공지능법(안)의 주요 내용과 시사점”(한국지능정보사회진흥원, 2021).
- 한국지능정보사회진흥원, “EU 인공지능법 입법현황과 시사점”, 『지능정보사회 법제도 이슈리포트』(2023-03).

- 홍석한, “유럽연합 ‘인공지능법안’의 주요 내용과 시사점”, 『유럽헌법연구』제38호(2022. 4).
- 다치바나 다카시, 김경원 (옮김), 『생태학적 사고법』(바다출판사, 2021).
- 빈프리트 하세머, 배종대 · 윤재왕 (옮김), 『범죄와 형벌: 올바른 형법을 위한 변론』(나남, 2011).
- 이즈미다 료스케, 이수형 (옮김), 『구글은 왜 자동차를 만드는가: 구글 vs 도요타, 자동차의 미래를 선점하기 위한 전쟁의 시작』(미래의창, 2015).
- 제리 카플란, 신동숙 (옮김), 『인간은 필요없다』(한즈미디어, 2016).
- 크리스토퍼 스타이너, 박지유 (옮김), 『알고리즘으로 세상을 지배하라』(에이콘, 2016).
- 타마르 헤르초그, 이영록 (옮김), 『유럽법약사』(민속원, 2023).
- 프란츠 폰 리스트, 심재우 · 윤재왕 (옮김), 『마르부르크 강령: 형법의 목적사상』(강, 2012).
- 프리츠포 카프라 · 우고 마테이, 박태현 · 김영준 (공역), 『최후의 전환: 지속 가능한 미래를 위한 커먼즈와 생태법』(경희대학교출판문화원, 2019).

[기타 문헌]

- 成原慧, “アーキテクチャの自由の再構築”, 松尾陽 (編), 『アーキテクチャと法』(弘文堂, 2016).
- 庄司克宏, 『はじめてのEU法』(有斐閣, 2020).
- Graf-Peter Calliess, Prozedurales Recht (Baden-Baden, 1999).
- Brian Patrick Green, “Six Approaches to Making Ethical Decisions in Cases of Uncertainty and Risk”, Markkula Center for Applied Ethics (2019).
- Madiega/Samy Chahri, “Artificial intelligence act”, European Parliamentary Research Service (2023).
- Marietje Schaake, “The European Commission’s Artificial Intelligence Act”, Human Centered Artificial Intelligence (Stanford University, 2021. 6).
- Proposal for a Directive of the European Parliament and the Council on Corporate Sustainability Due Diligence and amending Directive (EU) 2019/1937.
- Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative Acts.
- Regulation (EU) 2022/1925 of the European Parliament and of the Council of 14 September 2022 on contestable and fair markets in the digital sector and amending Directives (EU) 2019/1937 and (EU) 2020/1828.

Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC.

Regulation (EU) 2023/1542 of the European Parliament and of the Council of 12 July 2023 concerning batteries and waste batteries, amending Directive 2008/98/EC and Regulation (EU) 2019/1020 and repealing Directive 2006/66/EC.

Regulation (EU) 2023/1781 of the European Parliament and of the Council of 13 September 2023 establishing a framework of measures for strengthening Europe's semiconductor ecosystem and amending Regulation (EU) 2021/694.

Directive (EU) 2016/1148 of the European Parliament and of the Council of 6 July 2016 concerning measures for a high common level of security of network and information systems across the Union.

Directive (EU) 2022/2555 of the European Parliament and of the Council of 14 December 2022 on measures for a high common level of cybersecurity across the Union, amending Regulation (EU) No 910/2014 and Directive (EU) 2018/1972, and repealing Directive (EU) 2016/1148.

European Parliament resolution of 6 October 2021 on artificial intelligence in criminal law and its use by the police and judicial authorities in criminal matters.

<https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>(검색 일자: 2024. 1. 29.).

<https://www.aitimes.com/news/articleView.html?idxno=151738>(검색 일자: 2024. 1. 29.).

<https://www.europarl.europa.eu/news/en/press-room/20231206IPR15699/artificial-intelligence-act-deal-on-comprehensive-rules-for-trustworthy-ai>(검색 일자: 2024. 1. 29.).

https://www.eeas.europa.eu/delegations/south-korea/%EC%9C%A0%EB%9F%BD%EA%B7%B8%EB%A6%B0%EB%94%9C_ko?s=179(검색 일자: 2024. 2. 1.).

<https://www.eu-digital-services-act.com/>(검색 일자: 2024. 2. 1.)

<https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=CELEX:52021PC0206>(검색 일자: 2024. 2. 1.)

<https://www.scu.edu/ethics/focus-areas/technology-ethics/resources/six-approaches->

to-making-ethical-decisions-in-cases-of-uncertainty-and-risk/(검색 일자: 2024. 2. 1.).

<https://www.holisticai.com/blog/regulating-ai-the-horizontal-vs-vertical-approach>(검색 일자: 2024. 2. 1.).

https://www.europarl.europa.eu/doceo/document/TA-9-2021-0405_EN.html(검색 일자: 2024. 2. 3.).

<https://www.europarl.europa.eu/news/en/headlines/society/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>(검색 일자: 2024. 2. 3.).

[https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI\(2021\)698792_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI(2021)698792_EN.pdf)(검색 일자: 2024. 2. 3.).

https://www.europarl.europa.eu/doceo/document/TA-9-2023-0236_EN.html(검색 일자: 2024. 2. 3.).

https://hai.stanford.edu/sites/default/files/2021-06/HAI_Issue-Brief_The-European-Commissions-Artificial-Intelligence-Act.pdf(검색 일자: 2024. 2. 7.).

<https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC>(검색 일자: 2024. 2. 7.).

PROFILE



양 천 수

영남대학교 법학전문대학원 교수

Email: yang1000soo@hanmail.net / Tel: 053-810-2622

학 력

- 고려대학교 법과대학 법학사
- 고려대학교 대학원 법학과 법학석사
- 독일 프랑크푸르트대학교 법과대학 법학박사

주요경력

- 전) 미국 워싱턴대학교 로스쿨 방문연구원
- 전) 영남대학교 법무감사실장
- 전) 영남대학교 법학전문대학원 부원장
- 전) 고려저축은행 사외이사
- 현) 영남대학교 법학연구소장

주요연구

- 새로운 보안체계 구축에 관한 글로벌 규범 현황 및 법제 대응 방안 연구(한국법제연구원, 2023)
- 법과 시간(한국연구재단, 2022)

주요 논문

- “법과 시간: 판례 변경의 필요성과 기준을 예로 하여”, 『영남법학』제57호(2023. 12)
- “인공지능을 활용한 행정의 이론적 문제와 대응: 체계이론의 관점에서”, 『법학논총』(국민대) 제36권 제2호(2023. 10)
- “자율주행자동차에 관한 법정책의 방향: 의의·위험·규제”, 『모빌리티연구』제3권 제2호(2023. 9)
- “자율규제와 법다원주의”, 『법과 사회』제73호(2023. 6)
- “법이 규율하는 사회 개념의 변화: 민사법을 예로 하여”, 『법학논총』(조선대) 제30집 제1호(2023. 4)
- “와가츠마 사카에(我妻榮)의 법진화론”, 『안암법학』제65호(2022. 11)
- “기초법학의 의의와 필요성”, 『법철학연구』제25권 제1호(2022. 4)
- “실정법철학의 가능성과 방법: 민법학을 예로 하여”, 『법학논총』(조선대) 제29집 제1호(2022. 4)
- “현대사회의 구조변혁과 법규범의 대응 방향”, 『인간연구』제46호(2022. 4)