

유럽연합 인공지능법안에 대한 비판적 고찰*

김 라 은**

차 례

I. 들어가면서

1. 법안의 취지와 목표
2. 주요 내용 인공지능의 개념과 위험 기반 위험도 구분

II. 인공지능법안에 대한 평가

1. 금지된 인공지능과 예외
2. 고위험 인공지능
3. 감독

III. 인공지능과 인권, 민주주의 및 법치주의 원리와의 관계

1. 인공지능과 공적 의사결정에 대한 민주적 참여
2. 인공지능과 민주주의
3. 인공지능과 노동인권
4. 자유, 안전, 공정한 재판, 죄형법정주의 - 유럽인권조약

IV. 나가면서 - 인공지능법안에 대한 최근 유럽연합의 움직임, 우리 법에 주는 시사점

* 본 연구는 한성대학교 학술연구비 지원과제임.

** 한성대학교 미래플러스대학 융합행정학과, 조교수

I. 들어가면서

지난 2021년 4월 유럽연합 집행위원회는 유럽연합 시장에서 판매하고 사용하는 인공지능을 규제하는 법안을 제안하였고 유럽연합 집행위원회와 의회, 회원국 대표들은 2023년 12월 유럽연합 인공지능법안에 합의하였다. 지금까지 내용으로 보았을 때 인공지능법안은 유럽연합 전체에서 적용되는 인권, 민주주의, 법의 지배와 관련하여 인공지능에 대한 규제의 원칙과 비전을 제안하고자 했다는 점에서, 인공지능에 대한 국제표준을 제시하였다는 점에서 큰 의미를 찾을 수 있다. 문제는 유럽연합의 인공지능법안이 다양한 사회적 요구를 포함하고 있음에도 여전히 일정한 한계를 가지고 있다는 점이다. 또한, 최근 유럽연합 내부에서는 인권, 민주주의, 법의 지배에 대한 인공지능의 영향력을 어떻게 평가할 것인가에 대한 의문이 제시되고 있다. 적어도 현재까지는 인공지능 개발자나 사용자들이 인권이나 민주주의, 법치주의 등에 대하여 가질 수 있는 함의나 영향력을 제대로 알고 있다고 보기 어렵기 때문에 일단 유럽연합 인권조약에 규정된 인권과 인권의 사회적, 정치적 의미를 인공지능에 적합한 맥락으로 전환하는 작업 또한 시급하다. 이와 같은 문제점에서 출발한 본 논문은 유럽연합 인공지능법안의 내용을 비판적으로 평가하고 이른바 법 집행, 민주주의, 법의 지배 측면에서 향후 인공지능법안의 방향성에 관하여 고민하여 보고자 한다.

1. 법안의 취지와 목표

유럽연합 집행위원회(이하 집행위원회라 한다)는 2021년 4월 21일 유럽연합 회원국에 통일적으로 적용될 인공지능법안을 제안하였다. 인공지능법안은 이미 집행위원회가 인공지능의 신뢰성과 혁신이라는 원칙을 근간으로 유럽연합 전체에 적용될 기본방향에 대하여 2020년 발표한 백서를 대상으로 제시된 다양한

의견을 바탕으로 확정된 것이다.¹⁾ 여기에서 언급한 다양한 의견이란 주로 유럽 연합 이사회(이하 이사회라 한다)가 2020년 10월 21일 유럽연합 인공지능의 기본 개념에 대하여 내린 결론,²⁾ 유럽연합 의회(이하 의회라 한다)의 다수 결의사항,³⁾ 인공지능에 대한 고위전문가그룹이 제시한 견해를 의미한다.⁴⁾ 인공지능법은 그 이유서에서 향후 유럽에서 인공지능이 가지고 올 수 있는 긍정적 변화와 기회는 물론 위험성에 대해서도 상세하게 분석하고 있는데, 이미 이사회는 인공지능이 가져올 수 있는 편향적 결과 및 인공지능이 학습을 통하여 얻는 인간이 예측하기 어려운 불명확한 결과를 경계하였다. 특히 인공지능은 유럽연합 회원국 국민의 기본권 보장이라는 이념과 목표에 역행할 수 없으며, 동시에 법의 집행을 부당하게 지연하여 혁신을 방해하는 일이 없어야 함을 강조한 바 있다. 집행위원회는 인공지능법안의 몇 가지 목표에 대하여 언급하고 있다. 첫째, 유럽 연합 내에서 출시되고 사용되는 인공지능은 안전하여야 하며, 유럽연합의 가치와 회원국 국민의 기본권을 보장하여야 한다. 둘째, 인공지능에 대한 투자를 촉진하고 기술혁신을 위해서는 법적 안정성이 보장되어야 한다. 셋째, 인공지능과 관련된 기존 법률의 거버넌스와 효율적인 집행 및 기본권 보장, 인공지능의 안전성에 대한 요구를 강화하여야 한다. 넷째, 인공지능을 적법성, 안전성, 신뢰성을 담보하는 방식으로 사용하여 유럽연합 내의 시장을 분열시키지 않고 안정적인 발전을 도모하여야 한다. 집행위원회의 인공지능법안은 인공지능에 대한 세

1) EU, see <<https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52020DC0065&from=EN>>Updated 2020. Accessed 27 Nov 2023.

2) EU, see <<https://data.consilium.europa.eu/doc/document/ST-11481-2020-INIT/de/pdf>>Updated 2020. Accessed 19 Apr 2024.

3) Entschlüsseungen des Europäischen Parlaments vom 20. 10. 2020 mit Empfehlungen an die Kommission zu dem Rahmen für ethische Aspekte von künstlicher Intelligenz, Robotik und damit zusammenhängenden Technologien 2020/2012 (INL), für eine Regelung der zivilrechtlichen Haftung beim Einsatz künstlicher Intelligenz, 2020/2014 (INL), zu den Rechten des geistigen Eigentums bei der Entwicklung von KI-Technologien, 2020/2015 (INI), Berichtsentwürfe über künstliche Intelligenz im Strafrecht und Ihre Verwendung durch die Polizei und Justizbehörden in Strafsachen, 2020/2016 (INI), über künstliche Intelligenz in der Bildung, der Kultur und dem audiovisuellen Bereich, 2020/2017 (INI).

4) EU, see <<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>>Updated 2019. Accessed 19 Apr 2024.

계 최초의 법안이며 매우 정교하고 혁신적인 내용을 담고 있는 것으로 평가하는 것이 일반적인 기류라고 할 수 있다.⁵⁾ 특히 인공지능법안은 민간 부문은 물론 정부가 인공지능을 사용하기 위한 법적 근거를 마련함과 동시에 인공지능의 사용과 관련된 사람들의 기본권 등 법적인 보호를 중점적으로 다루고 있다는 점에서 긍정적이다. 미래는 디지털 사회이며 사회의 디지털 전환은 기업과 소비자 모두에게 혁신과 효율성을 재고한다. 하지만, 향후 디지털의 가속화에 따르는 장점과 더불어 법치국가의 기본적인 가치들과의 신중한 형량이 필요하다.

2. 주요 내용 - 인공지능의 개념과 위험 기반 위험도 구분

인공지능법안은 제3조 제1호에 따라서 인공지능의 개념과 범위를 다음과 같이 광의로 규정하고 있다. “부속서 I에 언급된 하나 이상의 기술과 접근방식으로 개발되고, 또한 인간이 정의한 목적에 따라, 상호작용하는 환경에 영향을 미치는 콘텐츠, 예측, 추천, 또는 의사결정과 같은 결과물을 생성할 수 있는 소프트웨어를 말한다.” 인공지능 시스템에서 처리하는 정보 역시 제3조 제29호 내지 제33호에서 “학습정보, 타당성 검증정보, 시험정보, 입력정보 및 생체정보”를 포함하는 넓은 개념을 사용하고 있다.

인공지능법안은 인공지능의 위험도에 따라서 수용 불가, 고위험, 저위험 및 최소위험으로 인공지능을 분류하고 있으며, 수용 불가 인공지능은 생산과 사용이 원칙적으로 금지된다(제5조). 이 중에서 인공지능법안이 집중적으로 규정하고 있는 것은 고위험 인공지능인데(제6조 이하), 여기에 대해서는 위험관리시스템(제9조) 구축을 위한 학습정보, 타당성 검증정보, 시험정보, 입력정보 및 생체정보에 대한 데이터 거버넌스 및 정보행정절차(제10조), 기록 및 표시의무(제11조, 제12조), 인공지능의 작동과 결과에 자동적인 표시의무(제12조), 투명성 요건(제13조) 등이 적용되며, 인간에 의한 감시(제14조) 정확성, 견고성 및 안전성

5) 김광수, “인공지능 알고리즘 규율을 위한 법제 동향: 미국과 EU 인공지능법의 비교를 중심으로”, 『행정법 연구』, 제70호(2023), 181면 이하.

이 보장되어야 한다(제15조). 인공지능 제공자는 다양하고 엄격한 의무의 준수, 일정한 기록의무의 이행과 더불어 품질관리시스템을 설치하여야 한다(제17조 이하).⁶⁾ 고위험 인공지능의 경우 인공지능 생산자, 대리인, 수입자, 판매자, 소비자, 유럽연합 규격마크(CE) 의무 등에 적용되는 제조물책임법의 규정과 상당 부분 유사한 내용이 적용된다.⁷⁾ 고위험 인공지능을 제외한 제한된 위험, 최소위험과 같은 이른바 기타 인공지능의 경우에는 인간과의 교류, 감정인식, 생체정보 및 이른바 딥페이크 등을 가능하게 하려면 투명성 의무를 준수하여야 한다(제52조).⁸⁾

II. 인공지능법안에 대한 평가

1. 금지된 인공지능과 예외

인공지능법안은 특정인을 대상으로 사회적 활동이나 공개 또는 예견된 개인적 특성 등을 바탕으로 사회적 신뢰도를 구분하는 사회적 점수평가 인공지능을 원칙적으로 금지하고 있다. 그런데 사회적 연관 관계에 비추어 애초의 정보 생성과 맥락을 달리하여 특정 자연인을 차별하거나 불리하게 대우하는 경우가 아닌 경우, 특정 자연인의 사회적 행태나 활동 범위를 부당하고 과도한 방식으로 판

6) 인공지능 공급자 내지 제공자의 의무와 관련된 상세한 논의는 김진우, “유럽연합 인공지능법안에 따른 고위험 인공지능 시스템공급자 등의 의무”, 『과학기술과 법』, 제12권 제2호(2021), 151면.

7) 유럽연합 내에서 제조물책임과 관련된 규정에 대한 명확한 이해와 통일적이고 일관된 적용에 관해서는 EU lex, <[https://eur-lex.europa.eu/legal-content/DE/TEXT/PDF/?uri=CELEX:52016XC0726\(02\)&from=SV](https://eur-lex.europa.eu/legal-content/DE/TEXT/PDF/?uri=CELEX:52016XC0726(02)&from=SV)>, Updated 2016, Accessed 19 Apr 2024 참고.

8) 유럽연합 인공지능법안의 구체적인 내용과 평가에 대해서는 홍석한, “유럽연합 ‘인공지능법안’의 주요 내용과 시사점”, 『유럽헌법연구』, 제38호(2022), 245면; NIA 지능정보원, “EU 인공지능법(안)의 주요 내용과 시사점”, (2021), 15면.

단하기 위한 목적이 아닌 경우에는 예외적으로 이 같은 인공지능도 허용될 수 있다는 입장을 취하였다. 때문에, 이 예외 조항으로 말미암아 원칙적으로 금지되는 사회적 점수평가를 목적으로 하는 인공지능의 범위를 현저하게 제한하는 결과를 가져올 수 있다는 문제점이 지속적으로 제기되었다. 초안과 의회안을 통틀어 나타난 또 하나의 문제는 사회적 점수평가를 위한 인공지능 금지는 국가, 공공부문 대상으로만 적용된다는 점이다. 인공지능법안 제5조 제1항 d)호에 의하면 형사소추를 위하여 실시간으로 개인의 생체정보를 확인하는 인공지능은 원칙적으로 금지하지만 여기에 대해서도 역시 상당한 예외를 인정함으로써 원칙이 상대화하는 경향을 보여준다. 물론 인공지능법안 제5조 제2항은 실시간 개인의 생체정보를 확인하는 인공지능을 예외적으로 허용하는 경우도 과잉금지원칙을 엄격하게 준수할 것을 요구하고 있지만 앞서 언급한 예외적 요건에 따라서 회원국들이 인공지능을 활용할 수 있다면 공공장소에서 개인에 대한 감시 행위는 더욱 강화될 것이 분명하다.⁹⁾ 따라서, 이와 같은 예외 조항을 계속하여 인공지능법안에 유지할 것인가에 대해서는 향후 심도 있는 논의가 필요하다.

2. 고위험 인공지능

인공지능법안이 규정하는 고위험 인공지능은 그 사용 목적에 따라서 사람의 건강, 안전 또는 기본권 침해에 대한 고도의 위험을 야기시킬 수 있는 경우를 말한다. 인공지능법안은 제6조 제1항에서 고위험성이 있는 인공지능의 판단 기준으로서 부속서 II에 규정된 제조물책임법 규정에 해당하는 제품을 기준으로 하고 있으며, 이러한 제품에 대해서는 역시 부속서 II가 규정하는 적합성 평가의 대상이 되는 기준을 말한다. 예를 들어 인공지능을 활용한 기계류, 장난감, 스포츠용 선박, 승강기, 무전장치, 의약품 등이 포함된다. 더 나아가 고위험 인공지능

9) Spindler, Gerald. "Der Vorschlag der EU-Kommission für eine Verordnung zur Regulierung der Künstlichen Intelligenz (KI-VO-E)—Ansatz, Instrumente, Qualität und Kontext." *Computer und Recht*, Vol.37 no.6 (2021).

에는 인공지능법 제6조 제2항과 부속서 III에 규정된 것으로서 생체정보 확인, 개인별 분류, 민감한 기반시설의 운영과 관리, 일반 및 직업교육, 고용, 인사관리, 민간 및 공직에 대한 취업, 부속서 III 제6호의 법집행 내지 형사소추와 관련된 이민, 정치적 망명, 국경통제, 사법 시스템 등이 포함된다. 여기에서 특히 눈여겨 보아야 할 분야는 마지막에 언급한 1. 형사소추, 2. 이민, 정치적 망명, 국경통제, 3. 사법 시스템이다. 당해 내용은 애초 초안에는 포함되어 있지 않았으나 후에 의회안에서 등장하였는데 당시에는 금지된 인공지능으로 분류 되었다. 이후 최종적으로 협상안에는 그 내용이 제외되었으며 일괄적으로 고위험 인공지능으로 분류되었다. 사실상 이와 같은 영역은 인공지능을 활용하는 것이 금지되어야 하는 것임에도 불구하고 이를 지금과 같이 고도 위험을 수반하는 인공지능으로 본다면 유럽연합 회원국들이 향후 최종적인 인공지능법이 제정되었을 때 인공지능을 적극적으로 활용할 가능성을 배제할 수 없다. 그렇다면 이러한 세 가지 영역에서 고도 위험을 수반하는 인공지능의 활용 가능성에 대해서는 보다 진지한 논의가 필요하다.

(1) 사법 시스템

인공지능법안 부속서 III 제8호는 사법행정청, 독립 행정청, 형사소추청 등을 지원하기 위하여 운영하는 인공지능을 고위험군으로 규정하고 있다. 여기에서 말하는 사법행정청이란 일차적으로는 법원을 가리킨다. 다만, 법의 적용과 사실관계의 포섭은 행정청도 수행하고 행정청의 법해석은 사법부의 판결에도 적지 않은 영향을 미칠 수 있으므로 행정청의 결정도 포함되는 것으로 보아야 한다. 그런데 이와 같은 영역에서 고위험을 수반하는 인공지능의 활용이 일정 수준 가능하다고 하더라도 이는 사실관계와 법령에 대한 조사와 해석, 사실관계에 대한 법의 적용을 위한 사법행정청을 지원하는 인공지능에 제한되고 사람의 결정이나 판단을 전적으로 대신하여 사법시스템에 관련되는 결정을 스스로 내리는 인공지능은 규정의 해석상 불가한 것으로 보아야 한다. 특히 사법부의 판결이나 그에 준하는 행정청의 결정은 관련되는 사람들의 운명에 결정적인 영향을 미칠 수

있다는 점에서 당해 인공지능, 인공지능법안에서 금지되는 인공지능과 사실상 같은 취급을 받아야 한다. 따라서 사법부의 판결이나 행정청의 결정을 “단순 지원”하기 위하여 인공지능을 사용할 수 있다고 하더라도 재판청구권 등 개인의 기본권 실현에 고도의 위험성을 야기할 가능성이 있으므로 충분하고 효과적인 개인 보호, 독립적이고 공정한 재판, 무죄추정, 좋은 행정의 원칙 등을 침해할 수 없다는 한계를 분명하게 설정하여야 한다. 미국 연방대법원장 존 로버츠 판사는 지난 2023년 연말보고서¹⁰⁾에서 AI가 변호사와 변호사가 아닌 사람 모두를 중요한 정보에 접근할 수 있도록 하여 상호간의 불평등을 해소하는 순기능을 할 수 있다는 측면이 있음을 강조하면서도 법적 결정은 여전히 인간의 판단을 적용하여야 하는 회색 영역을 포함하기 때문에 인간 판사만이 공정성과 신뢰성이 담보된 결정을 할 수 있다고 강조하였는데 존 대법원장의 말도 모두 이와 같은 맥락에서 이해할 필요가 있다. 비록 인공지능이 자연인 법관에 비하여 소송 당사자가 합리적이고 바람직한 판결을 내린다고 볼 수 있는 경우에도 인공지능 시스템은 결코 자연인 법관을 대신하여 종국적인 판결을 할 수 없다. 인공지능의 역할은 어디까지나 재판의 준비 과정이나 판결을 내리기 위한 기술적 작업에 국한되어야 한다. 이는 행정청의 결정에도 마찬가지로 적용된다. 행정청에서 일정한 결정을 내리는 공직자의 자격 역시 헌법과 법률에 따라서 규정되기 때문에 인공지능이 국민 전체와 개인의 운명에 중대한 영향을 미치는 의사를 스스로 결정할 수 없으며, 불확정개념을 해석하거나 정책적 성격이 강한 재량권 행사 등을 할 수 없는 것으로 보아야 한다. 결국 인공지능은 헌법이나 법률상의 사법권, 행정권 등 일종의 고권적 권력을 행사하는 것은 원천적으로 금지되는 것으로 보아야 한다. 헌법과 법률에 따라서 일정한 자격을 가지는 자연인으로서 법관과 또는 공직자에게만 권한 행사를 허용한 것은 헌법상 민주주의의 원리, 민주적 정당성과 결을 함께 하는 것으로 보아야 한다. 때문에, 독일 연방헌법재판소는 개인이 사법절차의 객체가 아니며, 사법부의 판결이 있기 전까지 판결의 절차와 결과

10) Supreme Court of United States, “2023 Year End Report on the Federal Judiciary”, 존 로버츠 대법원장의 언급과 관련하여 자세한 내용은 <<https://www.supremecourt.gov/publicinfo/year-end/2023year-endreport.pdf>>, 2면 이하 참고.

에 대하여 일정한 영향력을 미치는 주체라는 사실을 강조하고 있다.¹¹⁾ 행정절차에 참여하는 이해 관계인 역시 행정청이 결정을 내리는 과정에서 필요한 자료를 제시하고 의견을 진술하는 수동적인 입장에 그치는 것이 아니라 행정청의 결정에 영향을 미치기 위하여 행정청을 설득하고 소통하는 기본권 행사의 주체라는 점이 중요하다. 그런데 이와 같은 사법절차의 진정한 기능이나 다수인이 행정절차에 참여하여 다양한 이해관계를 형량하는 과정은 사실상 인공지능에서는 실현하기 어렵다. 아무리 단순한 사건이라고 하더라도 자연인 법관을 대신하여 인공지능이 스스로의 판단에 근거하여 판결할 수 있도록 한다면 헌법과 법률에 근거하여 법관에 의한 재판을 받을 권리는 파괴되고 인공지능이라는 무자격 사인에 의하여 재판받는 결과로 이어질 것은 분명하다. 인공지능이 가지고 있는 편의성과 경제성이라는 장점에도 불구하고 민사 및 행정소송절차에서 판결의 확정력 등을 고려할 때 비록 쟁점이 없는 단순한 판결의 경우라도 자연인 법관이 내리는 경우만 수용할 수 있다고 보아야 한다. 이러한 측면에서 사법부의 판결이나 행정청의 결정을 지원하는 인공지능을 허용하기는 하지만 이를 고위험군으로 분류하고 특히 인간에 의한 통제 가능성을 규정한 인공지능법안 제14조에 따라서 다양한 조건을 부과하는 것은 합리적인 규율방식으로 보인다. 그런데 자연인 법관 역시 이와 같은 인공지능에 대한 근거 없는 긍정적 편견을 가질 수 있으며, 심리 과정에서 무의식적으로 당사자의 의견을 충실하게 청취하는 과정을 간과할 수 있다. 소송 과정에서 당사자는 법관과는 달리 재판을 지원하는 인공지능과 세부적인 사항에 대하여 상호 반응할 수 없으며, 인공지능을 이용하는 사용자도 아니기 때문에 인공지능법안 제13조와 제52조에 따르는 투명성 의무 역시도 그들에게 큰 도움이 되지 못한다. 최종안에 따라 유럽연합 시민들의 경우 고위험 시스템에 기반한 결정에 대하여 설명을 요구하고 민원을 제기할 수 있는 권리를 보장받았다고 하더라도 보호수준은 여전히 충분하다고 보여지지 않으며 궁극적으로 소송당사자는 재판 과정에서 어떠한 인공지능이 재판을 지원하기 위하여 사용되었는지 충분하게 설명을 요구할 권리를 행사할 수 있어야 하

11) BVerfG, NJW 1984, S. 113.

며, 이를 법적으로 보장할 필요가 있다. 따라서 앞서 언급한 바와 같이 소송당사자들에게 재판 과정을 지원하는 인공지능의 종류와 지원 내용에 대하여 정보를 제공할 의무를 법관 등 사법부의 인공지능 사용자에게 부과하는 방안을 입법적으로 적극적으로 고려할 필요가 있다.

(2) 고위험 인공지능과 형사소추

인공지능법안 부속서 III, 제6호에서는 형사소추 관청이 고위험 인공지능을 사용할 수 있는 경우를 1. 특정인이 범죄를 범하거나 재범을 저지르거나 범죄의 희생자가 될 수 있는 위험성을 평가하기 위하여 인공지능을 사용하는 경우 2. 거짓말 탐지기 또는 이와 유사한 수단이나 특정인의 감정상태를 조사하기 위하여 인공지능을 사용하는 경우 3. 인공지능법안 제52조 제3항의 이른바 뎁페이크를 발견하기 위하여 인공지능을 사용하는 경우 4. 범죄행위의 조사나 소추 과정에서 증거물의 신뢰성을 평가하기 위하여 인공지능을 사용하는 경우 5. 형사소추 관청이 범죄행위의 예방, 수사, 형벌의 집행 등을 위한 개인정보 활용으로부터 관계인을 보호하기 위하여 유럽연합 의회와 집행위원회가 제정한 유럽연합 지침 2016/680 제3조 제4항을 근거로 실제적이거나 혹은 잠재적인 범죄의 발생 및 재범을 예측하기 위한 경우 또는 개인이나 집단의 범죄행태, 특성, 개인적 성향 등을 평가하기 위하여 인공지능을 사용하는 경우 6. 유럽연합 지침 2016/680 제3조 제4항을 근거로 범죄의 발견, 수사 및 소추를 목적으로 특정인의 프로필을 생성하기 위하여 인공지능을 사용하는 경우 7. 특정인의 범죄를 분석하거나 형사소추 관청이 다양한 情報源으로부터 수집한 대량의 복잡한 연관 정보 및 다양한 정보를 수색하기 위하여 인공지능을 사용하거나 이를 통하여 익명의 情報典範을 발견 또는 이와 같은 정보와 숨은 관계에 있는 정보를 수집하기 위하여 인공지능을 사용하는 경우. 로 나열하고 있다. 형사소추를 위하여 인공지능을 사용하는 과정에서는 형사소추 관청과 피의자 등과의 힘의 불균형으로 인하여 형사소추 관청은 범죄의 혐의가 있는 특정인을 일방적으로 감시할 수 있으며 경우에 따라 그의 신체적 자유를 박탈할 수 있는 경우까지 있을 수 있다. 때문에,

유럽연합 권리장전에 규정된 개인의 기본권을 침해할 우려가 매우 높다. 인공지능이 사용되거나 운영되기 이전에 정보에 대한 학습이 부족하거나 정확성과 견고성을 충족하지 못한 경우, 인공지능 시스템이 적절하게 구성되지 못하거나 시험이 이루어지지 않은 경우 등에는 특히 개인의 기본권을 침해할 우려와 위험성이 커질 수 있다. 결국 이와 같은 감시 과정을 통하여 범죄행위의 의심이 있는 특정인을 차별하거나 형사사법절차에서 고도 위험 인공지능을 무분별하게 사용하는 경우 특정인의 효과적인 사법적 권리구제, 공정한 재판을 받을 권리, 무죄추정 및 방어권이 훼손될 수 있다. 여기에 더하여 인공지능이 애초 입법자들이 요구하는 수준과 비교하여 충분하게 투명하지 못하거나 설명 가능성이 없고 관련 기록이 부실한 상태에서 형사소추를 위한 다양한 목적으로 사용되는 경우는 정확성, 신뢰성, 투명성을 요건으로 하는 인공지능의 기능적 요건은 심하게 훼손될 가능성이 있다. 때문에, 이를 고위험군으로 분류하고 있음을 충분히 인지하여 현장에서 나타날 수 있는 부정적 효과를 차단하고 공공의 신뢰를 회복하며 그 책임성과 관계인들의 효과적인 권리보장을 제고하는 것이 중요하다.

(3) 이민, 정치적 망명, 국경통제 영역과 인공지능

인공지능법안 부속서 III 제7호는 이민, 정치적 망명, 국경통제 영역에서 관련 행정청이 사용하는 인공지능을 다음과 같은 경우에 고위험군으로 분류하고 있다. 1. 거짓말 탐지기 또는 이와 유사한 수단이나 특정인의 감정 상태를 조사하기 위한 목적으로 인공지능을 사용하는 경우 2. 특정인의 불규칙한 이민행태 등으로 인하여 발생하는 안전에 대한 위험을 평가하거나 회원국의 영토에 입국하거나 입국할 것을 의도하는 특정인의 건강상태에 대한 위험을 평가하기 위하여 인공지능을 사용하는 경우 3. 특정인의 여행서류 및 증명서류의 진정성을 검사하거나 안전성 기준에 따라서 부정한 서류를 발견하기 위하여 인공지능을 사용하는 경우 4. 정치적 망명이나 비자 및 체류자격 및 그에 대한 이의를 신청하는 특정인의 자격을 확인하기 위한 작업을 지원하기 위하여 인공지능을 사용하는 경우. 이민, 정치적 망명, 국경통제 등의 대상이 되는 사람들은 대부분 그 지위

가 불안한 상태에 있으며, 이러한 업무를 수행하는 행정청의 일방적인 결정에 의존한다. 따라서 이러한 영역에서 사용되는 인공지능은 특히 정확하고 비차별적이어야 하며 투명성 요건을 충족하여야 한다. 자국에서 탄압받거나 심대한 위협을 피할 목적으로 정치적 망명을 신청하는 경우 거짓말 탐지기 또는 이와 유사한 수단이나 특정인의 감정상태를 조사하기 위한 목적으로 인공지능을 사용하는 것은 사실상 국제법적으로 승인된 정치적 망명권을 침해화시킬 우려가 있다. 특히 이와 같은 위험하고 불안한 지위에 있는 특정인이 인공지능의 결정에 대하여 실질적으로 효과적인 이의를 제기하는 등 불복하기 어렵다. 따라서 이러한 상황에 놓인 정치적 망명 신청자 등에게 거짓말 탐지 등을 위한 인공지능을 사용하는 것을 금지하는 것이 바람직하다. 최근 유럽연합에서 정치적 망명의 신속화와 공정 제고를 위하여 이른바 사전 판단 절차를 도입하려는 입법적인 상황을 고려할 때 정치적 망명 신청자의 진의를 심사하기 위하여 거짓말 탐지기 등 인공지능 사용금지를 진지하게 고려할 필요가 있다.

(4) 투명성 의무의 함정

앞서 언급한 바와 같이 인공지능법안 제13조 및 제52조에 따라서 고위험 인공지능 및 인간과 상호작용하는 특정 인공지능의 경우에 적용되는 투명성 의무는 안전하고 사후 검증 가능한 인공지능을 유지하기 위한 핵심적인 부분이다. 그런데 이와 같은 투명성 의무는 인공지능법안 제13조 제1항과 제52조 제1항에 따라서 원래 개발된 목적으로 인공지능이 사용되는 경우에 국한하여 적용된다. 다시 말하자면 인공지능이 인공지능법안 제52조가 규정하는 특정 용도로 개발된 것이 아니라 단순한 기능과 형태로 시작하였다고 하더라도 머신러닝 등 스스로의 학습과정을 통하여 다른 목적으로 전환되거나 확대되지 않는 경우 투명성 의무가 적용되지 않는다. 때문에, 이는 현행 법안상 커다란 입법적 흠결이라고 할 수 있다. 원래의 목적에만 국한되지 않고 인공지능의 사용 목적이 전환, 확대되는 경우라도 투명성 의무를 적용할 수 있도록 입법적인 개선이 필요하다. 그리고 인공지능법안 제52조에 의하면 투명성이 보장되지 않는 특정 인공지능이 제

공하는 재화나 용역에 대하여 관계인이 이를 거부할 수 있는 권리를 인정하고 있지 않는데, 이 역시 중요한 입법적 흠결로 판단된다. 관계인에게 투명성이 다소 부족한 특정 인공지능이 제공하는 재화와 용역을 거부할 수 있도록 하는 기회를 부여하는 규정을 보완할 필요가 있다.

3. 감독

인공지능법안 제56조¹²⁾ 이하에 규정된 감독체계와 제71조¹³⁾, 제72조¹⁴⁾의 재규정은 유럽연합 개인정보 보호지침의 관련 규정과 매우 유사한 형태와 내용으로 구성되어 있다. 인공지능법안은 유럽연합 차원에서 인공지능위원회를 설치하고 인공지능 운영 전반에 걸쳐서 집행위원회를 조언하고 권고하는 기능을 가지게 되는데, 특히 집행위원회와 회원국의 인공지능 감독기구와의 조정, 지원, 협력을 통하여 유럽연합 전체 회원국에서 통일적으로 인공지능법을 적용할 수 있도록 집행위원회를 돕는 것이다. 개인정보 보호지침과는 달리 인공지능법안은 회원국 감독기구가 인공지능을 감독할 수 있도록 하고 있는데, 회원국 감독기구

12) 유럽 인공지능위원회를 설립한다(제1항). 유럽 인공지능위원회는 다음과 같이 집행위원회에게 자문의견을 제공하고 그 활동을 지원한다. (a) 본 법안의 적용을 받는 사항에 대해 국내의 감독기관과 집행위원회가 효과적으로 협력하도록 지원한다 (b) 본 법안의 적용을 받는 사항이 국제시장에서 현안으로 대두되는 경우 이에 대한 집행위원회와 국내 감독기관, 그 외 관할기관이 그에 대한 자문과 분석을 하는 것을 조율하고 지원한다. (c) 국내 감독기관, 집행위원회를 지원하여 본 법안이 일관성 있게 적용되도록 한다.

13) 회원국은 본 법안이 규정한 바에 따라서 본 법안의 규정을 위반하는 경우 벌금을 비롯하여 제재처분에 대한 규칙을 제정하고 동 규칙을 적절하고 효과적으로 시행하기 위한 필요한 모든 조치를 취하여야 한다. 제재처분은 실효성이 있어야 하고 과도하지 않아야 하며 유사한 위반행위를 방지하는 효과가 있어야 한다. 제재처분을 할 때에는 특히 소규모 제공자의 이해관계와 스타트업의 경제적 활력성을 고려하여야 한다(제1항).

14) 유럽 데이터보호 감독기구는 본 법안의 규율대상에 속하는 유럽연합 관청과 기관, 기구에 대해 벌금을 부과할 수 있다. 벌금 부과 여부와 건별 벌금 액수를 결정할 때에는 구체적인 상황과 관련된 모든 정황을 고려하되, 특히 다음 사항을 감안하여야 한다. (a) 위반행위의 성격, 경중, 지속기간, 파장 (b) 침해로 시정하고 그 침해로 향후 발생할 수 있는 악영향을 완화하기 위하여 유럽 데이터보호 감독기구에 협력한 이력, 유럽 데이터보호 감독기구가 동일한 사안에 대해 유럽연합 관청과 기관, 기구에 명했던 조치를 준수한 경우, 이 내용도 포함된다. (c) 해당 유럽연합 관청, 기관 또는 기구가 과거 유사한 위반행위를 한 이력(제1항).

는 인공지능에 대한 인증 및 시장에서의 감독권한을 가지고 있으며, 회원국들은 복수의 감독기구를 설치할 수 있다. 회원국 감독기구는 자신의 권한 행사와 업무 집행에 있어서 최대한 객관성과 공정성을 유지하여야 하지만, 업무상 완전한 독립성이 요구되는 것은 아니다.

집행위원회는 인공지능법안에 대한 白書を 발간하면서 각 분야별로 규제 내지 감독기구를 설치하여 인공지능법안의 관련 규정들을 효과적으로 집행할 것을 제안하였고 실제로 인공지능법안은 제63조 제1항¹⁵⁾에서 이를 고려한 조항을 도입하였다.¹⁶⁾ 우선 고도위험 인공지능 시스템은 제63조 제3항¹⁷⁾에 의하여 시장 감독규정에 따라서 감독대상이 되므로, 시장감독청은 인공지능 및 인공지능 시스템을 고도위험에 적용되는 제품으로 보고 위 규정을 적용한다. 그리고 이러한 시장감독규정은 금융감독의 경우 규제대상인 금융기관이 출시, 운영 또는 사용하는 인공지능에만 적용되며, 이에 대한 금융감독은 관할 시장감독청이 행한다. 형사소추, 이민, 정치적 망명, 국경통제 영역에서 사용되는 인공지능의 경우 별도의 감독기구를 설치할 수 있으며, 이러한 감독기구는 유럽연합 개인정보 보호 지침, 경찰 및 형사사법청에 대한 개인정보 보호지침 및 개별 회원국의 관련 규정에 따라 이미 업무를 집행하고 있는 경우에는 동시에 인공지능에 대한 감독기구로 기능한다(제63조 제5항). 제63조 제6항은 기타 인공지능법안의 적용범위에 해당하는 유럽연합의 기관, 시설 및 기타 행정청에 대해서는 유럽연합 정보 보호관이 시장감독청으로서 감독권한을 행사하는 것으로 규정하고 있는데, 이는 제59조 제8항과 동일한 내용이다.

15) 본 법안을 적용하는 인공지능 시스템에는 유럽연합 규정 2019/1020을 적용한다. 다만, 본 법안을 실효성 있게 집행하기 위하여 다음 사항을 적용한다. (a) 유럽연합 규정 2019/1020에서 말하는 경제적 운영자는 본 법안 제3편 제3장에 명시된 모든 운영자를 포괄하는 것으로 이해한다. (b) 유럽연합 규정 2019/1020에서 말하는 제품은 본 법안의 범위에 속하는 인공지능 시스템을 모두 포괄하는 것으로 이해한다.

16) EU lex, see <[https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52016XC0726\(02\)&from=SV](https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52016XC0726(02)&from=SV)>, Updated 2016, Accessed 19 Apr 2024.

17) 부속서 II 제A절에 명시된 법안의 내용을 적용받는 고위험 인공지능 시스템의 경우, 본 법안의 목적에 따르는 시장감시 기관이 본 법안에서 규정한 시장감시 업무를 담당하는 기관이 된다.

II. 인공지능과 인권, 민주주의 및 법치주의 원리와의 관계

1. 인공지능과 공적 의사결정에 대한 민주적 참여

앞서 언급한 유럽연합의 인공지능 고위 전문가 그룹은 연구기관, 민간 부문의 전문가, 관련 학자 및 시민사회의 대표자들을 모두 망라하도록 구성하여 인공지능의 의사결정 과정에 발생하는 왜곡 현상과 더불어 긍정적인 기회와 문제들을 모두 논의할 수 있도록 하였으며, 특히 아래와 같이 인공지능으로 인한 정치과정 에 대한 민주적 참여의 역할을 아래와 같이 제시하고 있다.¹⁸⁾ 1) 인공지능은 인간이 행하던 의사결정을 자동화되고 정보를 탑재한 기술로서 일반화하는 것을 의미한다. 문제는 이는 정부 등 공적 의사결정에서 일반 국민이나 이해관계인들을 배제할 위험성을 내포하고 있다는 점이다. 정부의 공적 의사결정을 인공지능이 대체하는 경우 일반 국민은 인공지능을 생산하거나 사용하는 사람들이 아니기 때문에 정작 이들의 일상에 영향을 미치는 공적 결정에 대한 민주주의 사회에서 국민의 참여와 여론 형성의 기능이 질식되는 결과가 발생할 수 있다. 2) 인공지능을 학습시키거나 시험하는 정보들은 때로는 특정 사회의 중위층을 과다 대표하거나 일부 계층을 배제하는 등의 경향성을 보이기 때문에 인공지능으로 결정하는 사항은 일정한 사회적 편견을 반영하는 경우가 상당 부분 존재한다. 이러한 편향적이고 편견에 치우친 의사결정은 인공지능의 투명성과 신뢰성을 떨어뜨리고 사회적 영향력에 대한 대표성을 훼손할 가능성이 상존한다. 3) 인공지능의 개발자와 사용자 간의 쌍방향 소통이 존재하지 않기 때문에 사용자는 오히려 개발자에게 필요한 정보를 제공하는 역할에 그치는 문제가 발생한다. 인공지능은 인간이 하는 결정을 대신하도록 개발·설계되었기 때문에 인공지능을 사

18) Community Reference Meeting : Challenges and Opportunities of Regulating AI. Report, see <<https://mau.diva-portal.org/smash/get/diva2:1695574/FULLTEXT01.pdf>> Updated 2022. Accessed 27 Nov 2023.

용하는 사람은 자신이 살아가는 사회를 건전하고 지속 가능한 방식으로 개선하는 일에서 소외감, 무력감을 느끼거나 무관심으로 일관하는 경향성을 띠게 된다. 4) 반면에 인공지능은 빅데이터 등을 이용하여 기존의 전통적인 민주주의 원리에 따르는 공적 의사결정에서 소외된 계층에게 문호를 열어 주는 긍정적인 측면도 존재한다. 다시 말하자면, 인공지능은 기존의 민주주의 시스템에 존재하는 구조적인 차별 문제를 발견하고 중요한 의사결정에서 소외될 수 있었던 계층의 견해를 포함시키는 길을 열어 줄 수도 있다. 이를 통해, 보다 민주주의 원리에 적합한 결정을 내릴 수도 있다. 또한, 인공지능은 공적 의사결정에 참여하는 사람들에게 양질의 자료와 정보를 제공하고, 복잡한 의사결정 과정에서 서로 간의 소통을 강화함으로써 민주적 의사결정의 폭을 넓힐 수 있다는 장점도 있다.

그러나 이것이 인공지능을 통한 공적 의사결정에 잠재하는 다양한 사회적 가치의 갈등 상황을 무리 없이 조절함이 용이하다는 의미는 아니다. 물론 인공지능은 공적 의사결정의 일관성과 공정성을 제고하기 위하여 긍정적인 기여를 할 것으로 기대하고 있지만, 의사결정 과정에 일반 국민이나 이해관계인의 실질적 참여가 가능하도록 디자인하지 않고 인공지능이 머신러닝이나 딥러닝 등 기계적인 학습에 과도하게 의존하는 경우 민주주의의 근간을 훼손할 가능성은 매우 크고 현실적이다. 특히 공적 의사결정 과정에 배제된 사회적 계층의 경우 정치적 의사결정이 단순히 기계적인 것으로 대체되었다고 느끼는 순간 그들이 갖게 되는 분노와 상실감은 민주주의의 근간을 흔들 위험성이 있다. 때문에, 민주적 의사결정 과정에 광범위한 인공지능의 사용이 이루어지는 경우 국민의 참여권과 의사소통을 위한 보호장치들을 마련하는 것이 중요하며, 특히 인간의 존엄과 가치라는 헌법의 최고이념과 가치를 어떠한 방식으로 보호하고 실현할 것인가를 법적, 제도적으로 고민하여야 할 필요가 있다. 인공지능을 활용하여 지식을 생산하고 확산함에 있어 인공지능을 디자인하고 활용하는 단계에 이르기까지 공동체 내의 모든 이해 관계인 간의 사회적 네트워크를 형성하고 각각의 행위 주체들을 서로 연결하는 것이 성패의 관건이다. 그러기 위해서는 인공지능을 보다 협력적이며 배분적인 지식과 정보의 원천으로 활용하는 것이 중요하다.

예를 들어 인공지능을 활용한 원격 비대면 진료가 이루어지는 경우 지금까지

지리적, 사회적 여건 등으로 인하여 지식과 정보에 접근이 어려웠던 사회적 계층과 상호 소통하여 진료와 그에 필요한 데이터의 질을 향상시킬 수 있다. 이와 같은 이른바 참여형 인공지능은 이해관계인들 간 보다 활발한 소통을 필요로 하며, 인공지능의 장단점에 대한 교육도 중요하다. 인공지능 개발자, 판매자, 유통자 등은 상호 유기적으로 협력하여 인공지능을 실제로 이용하는 사용자들이 인공지능이 만들어 내는 결과와 기술적 용어에 쉽게 접근할 수 있도록 공동의 목표를 설정하여, 이러한 목표가 대표성을 가질 수 있어야 한다. 여기에서 대표성이란 인공지능과 관련된 모든 결정이나 행위가 이해관계인들의 집합적인 노력의 결과가 되어야 함을 의미한다. 만약, 이러한 대표성과 집합적인 노력이 결여되었다면 이것은 결국 인공지능 개발자와 일상적 사용자 간 소통이 부재하다는 의미이며, 보이지 않는 배제가 존재한다는 것을 암시하는 것이다.

2. 인공지능과 민주주의

(1) 사회적, 정치적 담론의 형성, 정보에의 접근, 투표행태

민주주의가 제대로 작동하기 위해서는 민주적 의사결정 과정에 참여하는 시민들이 필요한 정보를 충분히 가져야 하며, 개방적인 사회적, 정치적 담론을 형성할 기회가 부여되어야 한다. 또한 민주주의의 핵심적인 제도로서 선거를 조작하거나 부정적인 영향력이 배제되어야 한다. 지식정보화 사회에서 시민들은 일반적으로 차고 넘치는 정보의 홍수 속에서 자신의 관심을 끄는 제한된 정보와 지식을 취사하여 선택하는 경향이 있다. 그런데 인공지능을 활용하는 검색엔진이나 인터넷 사이트들은 정보의 생성과 이를 필요한 시민에게 전달하면서 정보의 내용을 사전에 자신들에게 유리한 방향으로 결정한다. 물론 이러한 과정을 통하여 긍정적인 효과를 전적으로 배제할 수 없는 것은 아닌데, 예를 들어 번역기 등을 통하여 외국어에 능통하지 않은 이용자들에게 국내에서는 쉽게 접근하지 못하는 정보를 얻을 수 있도록 함으로써 민주적 참여 능력이 향상될 수도 있

다.¹⁹⁾ 하지만 동시에 인공지능은 자신이 원하는 방향으로 생성한 정보와 지식을 시민들에게 일방적으로 전달하고 특정 정보를 억압하는 행태를 보이는 경향을 볼 수 있는데, 이를 통하여 정보와 지식은 상당 부분 왜곡되고 일정한 편견으로 경도되는 위험성을 가지고 있다. 물론 유럽연합 내 월 평균 활성 서비스 이용자 수가 4500만 이상인 소셜 미디어 플랫폼 이른바 거대 온라인 플랫폼은 고위험 인공지능 시스템으로 분류되어 특별한 감시와 통제 대상이 되지만 오늘날 인터넷 공간을 이용한 크고 작은 플랫폼의 사회적 영향력과 전파력을 고려한다면 인공지능이 상업적 목적이나 특정한 정치적 경향성을 강하게 반영하는 경우 기술이 지배하는 우리 사회의 정보와 지식의 구조가 의사 형성을 극렬하게 분열시키고 특정 개인에게 영향력을 행사할 위험성은 여전히 상존한다.²⁰⁾ 이러한 현상은 결국 민주주의의 건강성을 나타내는 징표라고 할 수 있는 상호 간의 이해와 존중을 통한 사회적 통합을 방해한다. 인공지능이 민주적 의사형성 자체를 마비시키고 개인의 의사결정에 일방적인 영향을 미치는 경우 유권자들은 선거에서 자신의 고유한 정치적 견해나 판단에 따르는 선거권이나 투표권을 행사하기 어렵다.²¹⁾ 인공지능은 선거 과정 또한 왜곡시킬 수 있는 위험성이 있다. 선거운동에 인공지능을 사용하는 경우 정치인들은 특정 유권자들을 대상으로 자신에게 유리한 선전이나 정책을 일방적으로 홍보하는 과정에서 정치적 책임이나 윤리적 측면을 무시한 채 지속적으로 자신의 메시지를 전달한다.²²⁾ 물론, 이러한 정치선전과 일방적인 메시지 전달을 중심으로 하는 목표설정 전략이 선거 결과에 어느 정도 효과적인가에 대해서는 갑론을박이 있을 수 있다.²³⁾ 하지만, 각종 오디오와 비디오, 딥페이크 등을 사용하여 지속적으로 특정 유권자를 대상으로

19) Schroeder, R., *Social Theory after Internet*. UCL Press, 2018.

20) Zuboff, Shoshana. "The age of surveillance capitalism." *Social Theory Re-Wired*. Routledge, 2015.

21) Floridi, Luciano, et al. "How to design AI for social good: seven essential factors." *Ethics, Governance, and Policies in Artificial Intelligence* (2021)p. 751-752.

22) Bradshaw, Samantha, and Philip N. Howard. "Social media and democracy in crisis." *Society and the internet* (2019).

23) Bond, Robert M., et al. "A 61-million-person experiment in social influence and political mobilization." *Nature* 489.7415 (2012), pp. 297.

하는 표적 선거 전략은 이미 보편화되었으며 선거를 앞둔 선거운동 기간 내내 유권자들의 판단 능력을 흐리게 함으로써, 민주적 담론 형성과 자유롭고 자율적인 의사결정을 상당 부분 방해한다는 것은 실증적인 자료가 없다고 하더라도 충분히 경험적으로 입증할 수 있는 사실이다.

(2) 차별과 분리, 체계화된 위험성, 디지털 권력 집중

인공지능을 활용하여 의사결정의 투명성과 합리성을 제고하겠다는 것은 이론의 여지가 없다. 하지만, 인공지능을 이른바 “블랙박스 알고리즘”이라고 부르는 이유는 그 작동과정 자체를 전혀 알 수 없는 상태에서 이미 사회적으로 정보와 지식이 유통되는 시장에서 소외되거나 차별받는 계층의 차별과 분리를 더욱 촉진하는 위험성이 있기 때문이다.²⁴⁾ 사회의 특정 계층을 차별하고 분리하는 것은 건전한 민주주의 사회가 존속하기 위한 최소한의 경제적, 사회적 평등의 기반을 잠식하는 것으로서 세심한 주의가 필요하다. 민주주의 사회에서 인간이 하는 공적 결정에는 사회의 다양한 계층에 속하는 다수 이해관계인이 직접, 간접적으로 참여할 수 있는 반면에, 인공지능이 그 기능을 대신하는 경우, 소수의 제한된 주체나 그들의 견해를 주로 반영하는 알고리즘이 그들만의 일방적인 결과를 발생시킬 수 있다. 이러한 위험성을 흔히 인공지능이 야기하는 체계화된 위험성(systemic risks)이라고 부른다. 자본시장에서 인공지능이 상호 반응하여 초고속으로 자본의 흐름을 극대화하고 짧은 시간에 시장 혼란을 야기할 수 있다면 위에서 언급한 체계화된 위험성의 대표적인 예로 들 수 있는 사안이다.²⁵⁾ 또한 원자력이나 수력발전소, 지능형 교통시스템, 병원 등이 인공지능에 대한 의존도가 증가하는 경우 한 번의 오작동이나 시스템적 위험발생으로 사회 전체에 치

24) Eubanks, Virginia. Automating inequality: How high-tech tools profile, police, and punish the poor. St. Martin's Press, 2018; Lassegue, Jean. "Weapons of Math Destruction. How Big Data Increases Inequality and Threatens Democracy." Crown Publishing Group, 2017.

25) Wellman, Michael P, and Uday Rajan. "Ethical issues for autonomous trading agents." Minds and Machines Vol.27 (2017), pp. 619-620.

명적인 결과를 가져올 수 있다. 일단 인공지능을 활용하여 매우 효율적인 운영시스템이 도입되면 이러한 치명적 위험을 알면서도 점점 그 편이성에 익숙하게 되어 이를 다시 다른 것으로 대체하거나 효과적으로 대처한다는 것이 현실적으로 어렵다. 또한, 인공지능이 군사활동이나 국가의 안보 문제를 결정할 위치에 있게 될 경우 국제적인 안전과 평화가 위협받는 사정이 발생할 가능성 또한 배제할 수 없다.²⁶⁾

3. 인공지능과 노동인권

인공지능은 작업이 어렵고, 위험하며, 불결하고, 반복적이며 힘든 작업에 투입되는 경우 상당한 효과를 기대할 수 있다. 하지만, 반면에 인공지능은 고용 현장과 사업장에서 근로자들의 행태를 확인, 기록, 모니터링하고 인간의 간여를 배제한 상태에서 업무를 배분함으로써 장차 특정 근로자의 고용과 해고 등에 결정적인 영향을 미칠 수 있다. 물론 최종안에서는 부속서 3에서 근로자 관리를 고위험 인공지능으로 분류하고 있고, 위험관리 시스템의 구축, 실행 유지(제9조), 사용자에 대한 투명성 확보 및 정보 제공(제13조), 인적감독 보장(제14조), 정확성과 견고성(제15조)등의 다양한 의무를 부과하고 있다. 하지만, 결국 이것은 인간에 대한 감시를 인간이 아닌 인공지능으로 대체하는 것이며, 앞서 언급한 인공지능 특유의 편견과 결합하는 경우 유럽인권조약 제2조, 제3조 등이 규정하는 공정하고 안전하며 건강한 근로환경에서 일할 권리, 근로자들의 존엄권과 단결권 등을 침해할 수 있다. 다시 말하자면, 사용자가 인공지능을 활용하여 근로자들을 감시하는 경우 헌법상 보장된 단결권 등 근로자의 기본권이 무력화될 위험성이 존재한다. 그리고 인공지능이 특정 근로자의 근로상태나 그 결과를 확인하고 평가하여 향후 노동생산성 등을 예측함으로써 이를 임금이나 고용조

26) Nemitz, Paul. "Constitutional democracy and technology in the age of artificial intelligence." *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences* 376.2133 (2018).

건 등을 포함하는 근로계약이나 해고 등에 반영한다면 인공지능에 사용되는 표준정보 때문에 발생하는 특유의 편견으로 인하여 성별에 따르는 차별이나 고용상의 평등한 처우 등이 문제될 수 있다. 결국 당해 사안은 고위험 인공지능보다는 금지된 인공지능의 범주에 포함되어야 하는 것이라고 이야기 할 수 있다.

당장은 인공지능이 스스로 행하는 작업이나 그들의 평가에 의존하는 근로자들의 노동생산성이 향상될지는 모르지만 장기적으로는 인공지능이 선호하는 근로방식이나 근로자 상이 정형화, 표준화됨으로써 특수한 작업에 소요되는 숙련 노동자들을 찾아볼 수 없게 되어 결과적으로는 사업장 전체의 노동생산성이 하락하는 결과를 가져올 수도 있다.²⁷⁾ 숙련노동자들이 사라진다는 의미는 생산성의 하락에 그치지 아니하고 숙련노동이 요구되는 작업에 지속적으로 새로운 근로자들이 투입됨으로써 산업재해가 증가하는 부작용이 발생할 가능성이 결코 가벼이 여겨져서는 안된다.

4. 자유, 안전, 공정한 재판, 죄형법정주의

- 유럽인권조약: Convention for the Protection of Human Rights and Fundamental Freedoms. ECHR 제5, 6, 7조

인공지능이 기존의 편견을 고착화시키고 확대시킬 수 있는 대표적인 영역으로 법집행 및 사법 분야를 들 수 있다. 예를 들어 인공지능이 상습범의 위험을 예측하고 그에 따르는 가중처벌 등을 결정하는 경우와 같이 개인의 신체적 자유와 안전이 문제되는 사안에서 헌법상 보장된 공정한 재판을 받을 권리가 침해될 위험성이 존재한다. 인공지능이 흔히 블랙박스로 불리는 이유는 판사, 검사 및 변호인들의 입장에서 인공지능이 내리는 결정이 어떠한 법적 추론이나 논리적 연관관계 하에서 이루어지는 것인지를 쉽게 이해할 수 없으며, 그러한 판단에 대한 상소를 어렵게 하기 때문이다. 이러한 문제는 체포나 구금과 같은 법집행 영

27) Brynjolfsson, E./McAfee, A., The Second machine Age : Work, Progress and Prosperity in a Time of Brilliant Technologies, 2014, W. W. Norton & Company.

역에서도 발생할 수 있는데, 자의적인 의심이나 철폐, 구금을 받지 않을 신체의 자유 등이 역시 인공지능의 결정에 따라서 상당한 위협을 받을 수 있다. 특히 상습범의 경우 인공지능이 재범의 위험성을 예측하면서 구체적이고 개별적인 위협에 근거한 범죄나 비위사실에 초점을 맞추기 보다는 예를 들어 주소, 소득, 국적, 채무상태, 고용, 평소의 행태, 친구나 가족구성원과의 관계 등 이른바 정형적이고 표준화된 추상적 의심에 근거한다면 법집행이 불합리한 결과를 낳을 수 있다. 그리고 앞서 언급한 정형적이고 표준화된 위험 표지들이 어떻게 인공지능의 판단에 기초적인 자료로 활용되었으며, 각 표지의 비중은 어떠한 것인지 등등의 문제는 통상적으로 잘 알 수도 없다. 물론 같은 법집행 영역이라고 하더라도 실종된 사람을 찾기 위하여 인공지능을 활용하여 사진 속 인물의 나이 등을 정확하게 예측하거나 공항에서 의심되는 화물을 스캔하여 위험물을 탐지하는 등 과학적이고 합리적인 방식으로 의사결정을 내릴 수도 있다.

IV. 나가면서

- 인공지능법안에 대한 최근 유럽연합의 움직임 우리 법에 주는 시사점

2023년 유럽의회, 유럽연합 집행위원회, 유럽연합 이사회는 이른바 제3자 협상을 통하여 인공지능법안에 대한 최종안에 합의하였다. 합의된 법안 초안은 유럽의회와 회원국들의 공식 승인 과정을 거쳐야 하며, 승인 후 완전히 발효되기까지는 약 2년여의 시간이 소요될 것으로 예상되고 있다. 이미 언급한 바와 같이 인공지능법안은 위험 구분형 접근방식을 취하고 있으며, 이는 개발자와 사용자로 하여금 인공지능이 야기할 수 있는 위험성에 바탕을 두고 위에서 언급한 원칙을 따라야 할 것을 요청하는 것이라고 할 수 있다. 집행위원회가 제안한 인공지능법안과 함께 의회 역시 일정한 인공지능은 개발이나 사용 자체를 원천적으

로 금지할 것을 당부하고 있는데, 특히 의도적으로 조작적 기술을 사용하여 특정인들의 취약점을 악용하는 사례로서, 소비패턴 등 사회적 행태, 사회적, 경제적 지위, 성향 등을 기준으로 이들에게 낙인효과를 부여하거나 악의적으로 분류하는 이른바 사회적 점수평가가 여기에 해당한다. 이와 같은 이유에서 의회는 다음과 같은 인공지능의 개발과 사용을 엄격하게 금지하는 수정안을 의결하였다. 상당 부분은 집행위원회의 인공지능법안을 확인하는 내용이다. 의회는 사람의 건강, 안전, 기본권 및 환경에 영향을 미치는 이른바 고위험군 인공지능의 범위를 인공지능법안보다 다소 확장하였다. 이 외에도 의회는 인공지능 분야 기술 발전을 의식하여 인공지능을 개발하거나 출시하는 자들에 대하여 인권, 건강, 안전, 환경 및 민주주의와 법치주의를 강력하게 보호할 의무를 부과하였다. 이들에게는 인공지능으로부터 발생하는 위험을 평가, 완화하고 법안이 요구하는 디자인, 정보, 환경요건에 대한 준수, 유럽연합 데이터베이스에 대한 등록 의무 등이 부과된다. GPT와 같은 이른바 생성형 기초모델의 경우 일정한 콘텐츠가 인공지능에 의하여 생산, 발생되었다는 점을 반드시 공개하고 이러한 모델이 불법적 콘텐츠나 인공지능에 대한 학습용 데이터를 제작하기 위하여 사용된 정보 중 저작권에 해당하는 데이터의 개요를 공개하는 것을 금지하도록 요청하는 내용 등이 포함되었다. 한편, 의회는 인공지능의 혁신을 촉진하기 위하여 인공지능을 연구목적으로 개발하거나 그 부품에 대하여 사전 승인을 받은 경우 위에서 언급한 여러 가지 규제를 받지 않게 함으로써 인공지능 혁신을 위한 예외를 인정하고 있다. 또한 의회는 인공지능법안이 규제샌드박스를 지향하여 혁신을 촉진하는 동시에 인공지능을 실제로 사용하기 이전에 관할 행정청이 이를 시험 운영하도록 하여 이른바 “통제된 인공지능 환경”을 유지하도록 하였다. 더 나아가 의회는 인공지능을 실제로 사용하는 시민들이 용이하게 민원을 제기하는 방법을 강구하고 시민의 권리에 심대한 영향을 미치는 고위험군 인공지능이 내린 결정에 대하여 설명을 요구할 수 있는 방안을 검토하도록 요청하였다. 이를 통하여 전 세계적으로 인공지능에 대한 기준점 격인 법안이 세계 최초로 마련되었으며, 인권에 대한 강력하고 효과적인 보장, 인간중심의 민주적 통제와 동시에 인공지능 관련 산업에 대하여 예측 가능성을 부여함으로써 인공지능과 관련된 기

솔혁신을 이룩할 토대가 제공된 것으로 평가되고 있다. 최근 인공지능이 여러 분야에 활용되며 우리나라에서도 인공지능의 편익과 위험성에 관한 논의가 이어지고 있다. 최근 인공지능에 관한 기본법으로서의 성격을 가진 ‘인공지능법’이 국회 상임위원회 법안소위 문턱을 넘었는데 이 때문에 우리나라에서도 멀지 않은 시일 내에 인공지능법이 제정될 것으로 기대를 모으고 있다. 그러나 제정안의 내용은 유럽연합과 비교하였을 때 인간에 대한 안전성이나 신뢰성을 높여야 한다는 것보다는 인공지능 이용의 활성화, 관련 연구개발 뒷받침이나 국가 역량 집중 투자 등을 위한 제도 마련에 주된 관심을 두고 있다는 점에서 아쉬움을 남긴다. 유럽연합 인공지능법안에서 주목하여야 할 점은 유럽연합의 경우 인공지능으로 인한 건강, 안전 및 기본적 인권에 대한 위험으로부터 사람들을 보호하고 신뢰성 있는 인공지능에 대한 혁신을 도모하는 목적을 가지고 있다는 점이다. 물론 앞서 검토하였던 유럽연합 인공지능 법안의 경우 지나치게 포괄적이고 강한 규제를 통해 첨단 기술 발전이나 혁신을 저해하는 부작용이 있을 수 있음을 우려하는 목소리가 많지만, 인공지능이란 경험해보지 못한 새로운 기술이 실시간으로 발전하며 우리 사회를 변화시키고 있는 와중에 순수히 사람의 선의나 자율성에 기댄 통제의 원칙이 향후 기술 발전 과정에서 과연 지켜질 수 있는 것인지, 마치 지동설 시대에도 중세의 윤리와 종교를 가지고 천동설을 고수하겠다는 시대착오적 논리에 불과한 것인지는 좀 더 깊은 논의와 고민이 필요한 부분이다.

유럽연합이 말하는 기본적 인권은 크게 다음과 같은 4가지 유형으로 구분하여 볼 수 있다. 1. 인간존엄(정직성, 프라이버시, 공정한 재판). 2. 자유(표현, 정보, 집회의 자유 등). 3. 공정성(차별금지, 동등한 처우). 4. 사회적 권리(교육, 건강한 일터, 사회적 편익 등)가 그것이다. 인공지능법안은 인공지능으로 인한 이와 같은 개인의 기본적 인권에 대한 침해의 위험성을 단계별로 구분하여 금지된 인공지능, 고위험 인공지능, 기타 인공지능으로 각각 규제의 엄격성을 차별적으로 적용하고 있다. 이에 따라서 인공지능법안은 인공지능을 통하여 일정한 사실을 조작할 위험성이 있는 인공지능, 사회적 등급·점수제, 법집행을 위한 공공장소에서의 실시간 원격 생체정보 확인제 등을 위한 인공지능 등을 금지하고 있다. 인공지능법안은 고위험 인공지능에 대해서는 그 사용을 허용하지만 동시

에 다양하고 엄격한 규제를 통하여 앞서 언급한 기본적 인권에 대한 보호를 철저하게 시행하기 위한 체제를 지향하고 있으며, 그 외 기타 인공지능에 대해서는 투명성 원칙을 요구하고 있다. 일반적으로 인공지능은 특정 개인에 대한 확인, 프로필 형성, 조사 등을 하는 것을 포함하여 개인의 행태를 스크린하거나 그에 대한 영향력 행사, 너지 등의 작업을 수행하기 때문에 도덕적 진실성 등 인격권에 대한 잠재효과를 미칠 뿐만 아니라 표현이나 정보의 자유에 대한 냉각효과를 가질 수 있다. 다시 말하자면, 특정인을 인식하기 위하여 인공지능을 활용하는 경우 개인의 의견표명 등 표현의 자유를 위축시킬 수 있으며, 특히 일정한 집단에서 특정인의 확인이 가능하다면, 그 집단 전체에 대한 익명성 역시 보장될 수 없다. 안면인식이 광범위하게 이루어지는 경우 사람들은 다수인이 참여하는 집회나 시위에 참여하는 것을 꺼리는 결과를 가져온다.²⁸⁾ 인터넷과 소셜미디어를 통하여 집회와 결사의 자유를 보다 효과적으로 실현할 수 있는 기회가 열린 것은 민주주의 사회에서 상당한 긍정적인 효과 중의 하나라고 할 수 있다. 하지만, 동시에 인공지능은 안면인식 기술 등을 활용하여 이와 같이 확장된 집회와 시위의 자유를 위축시키고 냉각시킬 수 있다는 점은 우리의 경우에도 유럽연합과 마찬가지로 깊이 고려할 필요가 있다.²⁹⁾ 수년 전 홍콩에서는 집회와 시위 과정에서 참가자들이 이와 같은 안면인식 기술을 두려워한 나머지 복면을 하거나 안면인식과 채증을 위한 카메라에게 레이저빔을 발사하는 등의 행태를 보이면서 안면인식을 위한 첨단 기법의 활용과 그에 대한 저항은 현실적인 문제로 제기되고 있다. 또 다른 중요하게 고려하여야 할 점은 인공지능이 특정인의 정보를 광범위하게 수집하고 이를 분류, 평가하여 다른 사람들이나 집단과 차별화하는 목적으로 사용하는 경우, 예를 들어 금융기관으로부터 대출을 받는 정보, 고용과 퇴직 등과 관련되는 정보가 그 대상이 되는 경우에도 개인의 표현이나 일반적인 행동자유권은 상당한 정도로 제한된다는 점이다. 이 경우, 개인들은 공개적

28) Privacy Impact Assessment Report for the Utilization of Facial Recognition Technologies to Identify Subjects in the Field, 30. June 2011, p. 18.

29) Council of Europe, Algorithms and Human Rights, Study on the Human Rights Dimensions of Automated Data Processing Techniques and Possible Regulatory Implications, 2018.

으로 자신의 의견을 표명하거나 심지어 특정 언론 매체를 구독, 시청, 접근하는 것을 상당한 정도로 꺼릴 수 있다. 개인이 타인의 의견을 청취하려 하는 경우 자신의 행태를 끊임 없이 확인하여 이를 체계화하고 분류하는 인공지능이 있다는 점을 의식할 수 밖에 없다면 과연 자신이 마주하는 정보나 의견이 스스로의 인격적 결정에 토대를 두고 있는 것인지 의심스러우며, 오히려 인공지능이 제공하는 너지효과에 익숙해질 위험성도 있다. 결국 인공지능의 광범위한 사용으로 인하여 사상과 표현의 자유가 인간의 자유로부터 기계의 자유로 대체되어 가는 것은 아닌지 자유의 본질에 대한 근본적인 문제와 마주하는 상황에 이르게 된다는 것을 항상 염두에 두어야 할 필요가 있다.³⁰⁾

실제로 오늘날 많은 비즈니스 모델이 온라인상의 광고 수입에 토대를 두고 있다는 점을 고려한다면 이 같은 현실이 일상화 되어가는 모습을 우리 스스로가 겪고 있다는 점과 무관하지 않다. 결국 인공지능이 지배하는 검색엔진이나 알고리즘화된 정보와 광고 플랫폼이 제공하는 정보와 뉴스 등이 객관성, 진실성, 다양성 등을 담보하고 있는지와는 관계 없이 사람들을 특정 플랫폼에 묶어두는 행태는 더욱 빈번해지고 일상화하는 경향을 나타내고 있다. 또한 인터넷 댓글 조작 등 불법 선거운동을 일상적으로 겪고 있는 한국의 상황을 고려하면 인공지능과 알고리즘이 지배하는 정보와 광고 플랫폼이 단지 상업적 목적으로만 사용되는데 그치지 않고 민주주의의 핵심사항인 여론과 정치적 의사 형성 과정에서 특정한 정치적 견해나 주장을 시민들에게 강요함으로써 자유로운 여론과 의사 형성을 왜곡시키는 위험성을 가지고 있다. 그리고 인공지능을 활용한 이른바 딥페이크 기술은 특정인의 이미지와 정체성을 왜곡하여 그릇된 판단을 유도하는 상황 역시 우리에게 전혀 낯설지 않은데, 이러한 일들은 앞으로도 계속하여 더욱 교묘해지고 지능화될 것으로 예상할 수 있다. 결국 인공지능이 가지고 올 혁명적인 변화와 함께 표현과 사상의 자유, 자유롭고 민주적인 의사 형성을 방해할 수 있는 각종 유형과 부작용을 고려하여 여기에 효과적으로 대처할 수 있는 인공지능법을 촘촘하게 입안하고 시행하는 것이 우리 경우에도 무엇보다도 중요

30) Burrel, J. How machine 'think' : Understanding Opacity in Machine Learning Algorithms. Big Data & Society, 2016, 3(1), 2053951715622512.

하다.³¹⁾

생각건대 인공지능의 경로의존성이 불가피한 만큼 유럽연합 정보 지침과 같은 영향력을 가진 규제와 입법의 필요성은 시대와 사회를 막론하고 필연적인 부분이다. 지금 시점에서 중요한 것은 인공지능에 대한 규제 가부를 다투는 것보다 규제는 필요하되, 이러한 규제가 향후 인공지능의 발전에 있어서 어떠한 영향을 미칠 것인가를 면밀하게 검토하여 개인과 기업의 권리보호 간 균형점을 마련하는 합리적 규제를 고민하는 것이다. 이러한 연장선 안에서 향후 우리 역시 유럽연합의 경우와 마찬가지로 규제의 목표와 관련하여 다음과 같은 점들이 심도 있게 고려되어야 한다. 1. 인공지능에 대한 규제가 어떠한 목적을 가지는 것인지를 명확하게 설정하지 않는다면 규제를 위한 입법은 무의미하며, 규제의 목적을 달성할 수도 없다. 2. 인공지능 규제에 대한 논의는 일종의 과포화 상태이기 때문에 규제의 다양한 방향성을 통합하고 균형을 유지하여 규제목표와 방향을 일치시키는 노력이 필요하다. 3. 인공지능 규제에 대한 초기 설계와 집행 단계부터 있어서 법, 기술, 사회연구 전문가 등 이해관계자들을 과정에 참여시키는 것이 매우 중요하다. 미래 사회와 그 변화에 대비하여 첫째, 인공지능을 통한 혁신 제고, 둘째, 사회적 경제적 변화에 대한 대응, 셋째, 적절한 법적, 윤리적 프레임 워크 마련이 앞으로 우리의 과제로 남았다.

투 고 일 : 2024. 03. 06.

심사종료일 : 2024. 03. 18.

게재확정일 : 2024. 03. 20.

31) UN Report of the Special Rapporteur on the Promotion and Protection of the Right to Freedom of Opinion and Expression, A/73/348.

[참 고 문 헌]

- 김광수, “인공지능 알고리즘 규율을 위한 법제 동향: 미국과 EU 인공지능법의 비교를 중심으로”, 『행정법 연구』, 제70호(2023).
- 김진우, “유럽연합 인공지능법안에 따른 고위험 인공지능 시스템공급자 등의 의무”, 『과학기술과 법』, 제12권 제2호(2021).
- 홍석한, “유럽연합 ‘인공지능법안’의 주요 내용과 시사점”, 『유럽헌법연구』, 제38호(2022).
- NIA 지능정보원, “EU 인공지능법(안)의 주요 내용과 시사점”(2021).
- Eubanks, Virginia. Automating inequality: How high-tech tools profile, police, and punish the poor. St. Martin's Press, 2018.
- Lassegue, Jean. "Weapons of Math Destruction. How Big Data Increases Inequality and Threatens Democracy." Crown Publishing Group, 2017.
- Schroeder, R., Social Theory after Internet. UCL Press, 2018.
- Staben, Julian. Der Abschreckungseffekt auf die Grundrechtsausübung: Strukturen eines verfassungsrechtlichen Arguments. Mohr Siebeck, 2017.
- Zuboff, Shoshana. "The age of surveillance capitalism." Social Theory Re-Wired. Routledge, 2015.
- Barrett, Lisa Feldman, et al. "Emotional expressions reconsidered: Challenges to inferring emotion from human facial movements." Psychological science in the public interest, Vol.20 no.1 (2019).
- Bond, Robert M., et al. "A 61-million-person experiment in social influence and political mobilization." Nature 489.7415 (2012).
- Bradshaw, Samantha, and Philip N. Howard. "Social media and democracy in crisis." Society and the internet (2019).
- Burrell, Jenna. "How the machine 'thinks': Understanding opacity in machine learning algorithms." Big data & society Vol.3 no.1 (2016).
- Colonna, Liane, et al. "Community Reference Meeting: Challenges and Opportunities of Regulating AI." (2022).
- Enders, Peter. "Einsatz künstlicher Intelligenz bei juristischer Entscheidungsfindung." Juristische Arbeitsblätter (2018).
- Floridi, Luciano, et al. "How to design AI for social good: seven essential factors." Ethics, Governance, and Policies in Artificial Intelligence (2021).
- Spindler, Gerald. "Der Vorschlag der EU-Kommission für eine Verordnung zur Regulierung der Künstlichen Intelligenz (KI-VO-E)—Ansatz, Instrumente,

- Qualität und Kontext." Computer und Recht, Vol.37 no.6 (2021).
- Mohapatra, Seema. "Use of facial recognition technology for medical purposes: balancing privacy with innovation." Pepp. L. Rev. Vol.43 (2015).
- Nemitz, Paul. "Constitutional democracy and technology in the age of artificial intelligence." Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences 376.2133 (2018).
- Wellman, Michael P., and Uday Rajan. "Ethical issues for autonomous trading agents." Minds and Machines Vol.27 (2017).
- Council of Europe, See<<https://edoc.coe.int/en/internet/7589-algorithms-and-human-rights-study-on-the-human-rights-dimensions-of-automated-data-processing-techniques-and-possible-regulatory-implications.html>>Updated 2018. Accessed 27 Nov 2023.
- Community Reference Meeting : Challenges and Opportunities of Regulating AI. Report, see <<https://mau.diva-portal.org/smash/get/diva2:1695574/FULLTEXT01.pdf>> Updated 2022. Accessed 27 Nov 2023.
- EU, see <<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>>Updated 2019. Accessed 19 Apr 2024.
- EU lex, see <[https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52016XC0726\(02\)&from=SV](https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52016XC0726(02)&from=SV)>, Updated 2016, Accessed 19 Apr 2024.
- EU, See<<https://eur-lex.europa.eu/legal-content/DE/TXT/PDF/?uri=CELEX:52020DC0065&from=EN>>Updated 2020. Accessed 27 Nov 2023.
- EU, see <<https://data.consilium.europa.eu/doc/document/ST-11481-2020-INIT/de/pdf>>Updated 2020. Accessed 19 Apr 2024.

■ 국문 초록

유럽연합 인공지능법안에 대한 비판적 고찰*

김락은**

지난 2021년 4월 유럽연합 집행위원회는 유럽연합 시장에서 판매하고 사용하는 인공지능을 규제하는 법안을 제안하였고 유럽연합 집행위원회와 의회, 회원국 대표들은 2023년 12월 유럽연합 인공지능법안에 합의하였다. 지금까지 내용으로 보았을 때 인공지능법안은 유럽연합 전체에서 적용되는 인권, 민주주의, 법의 지배와 관련하여 인공지능에 대한 규제의 원칙과 비전을 제안하고자 했다는 점에서, 인공지능에 대한 국제표준을 제시하였다는 점에서 큰 의미를 찾을 수 있다. 문제는 유럽연합의 인공지능법안이 다양한 사회적 요구를 포함하고 있음에도 여전히 일정한 한계를 가지고 있다는 점이다. 또한, 최근 유럽연합 내부에서는 인권, 민주주의, 법의 지배에 대한 인공지능의 영향력을 어떻게 평가할 것인가에 대한 의문이 제시되고 있다. 적어도 현재까지는 인공지능 개발자나 사용자들이 인권이나 민주주의, 법치주의 등에 대하여 가질 수 있는 합의나 영향력을 제대로 알고 있다고 보기 어렵기 때문에 일단 유럽연합 인권조약에 규정된 인권과 인권의 사회적, 정치적 의미를 인공지능에 적합한 맥락으로 전환하는 작업 또한 시급하다. 이와 같은 문제점에서 출발한 본 논문은 유럽연합 인공지능법안의 내용을 비판적으로 평가하고 이른바 법 집행, 민주주의, 법의 지배 측면에서 향후 인공지능법안의 방향성에 관하여 고민하여 보고자 한다.

주제어 : 유럽연합, 인공지능, 법치주의, 민주주의 원리, 규제

* 본 연구는 한성대학교 학술연구비 지원과제임.

** 한성대학교 미래플러스대학 융합행정학과, 조교수

■ Abstract

A critical review of the EU artificial intelligence bill*

Ra eun, Kim**

In 2021, the European Commission proposed legislation to regulate artificial intelligence(AI). The European Commission, parliament, and member state representatives agreed on the EU AI bill in December 2023. The AI Bill presents an international standard for AI in that it seeks to propose principles and vision for regulation of AI in relation to human rights, democracy, and the rule of law. The problem is that although the EU's AI bill includes various social demands, it still has certain limitations. Additionally, questions have recently been raised within the EU about how to evaluate the influence of AI on human rights, democracy, and the rule of law. Starting from these problems, this paper critically evaluates the contents of the EU's AI bill and considers the future direction of the AI bill in terms of so-called law enforcement, democracy, and rule of law.

KeyWord : EU, AI, Rule of Law, Democratic Principals, Regulation.

* This research was financially supported by Hansung University

** Hansung University