

국민권익위원회 청렴윤리경영 브리프스 10

Ⅰ AI와 청렴윤리경영

2023 October

Vol. 130

AI거버넌스와 청렴윤리경영

연세대학교 객원교수 조신

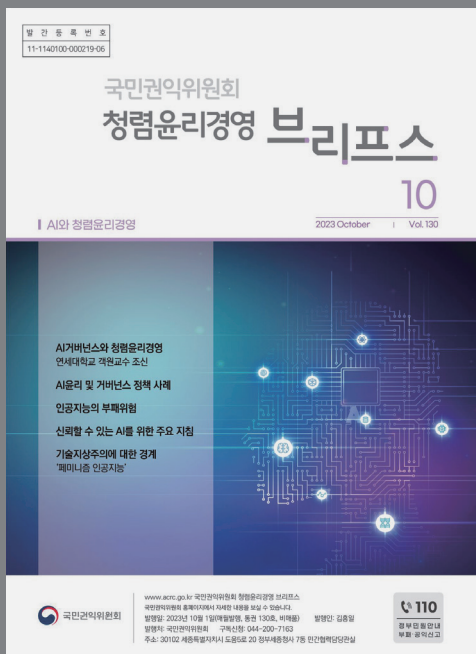
AI윤리 및 거버넌스 정책 사례

인공지능의 부패위험

신뢰할 수 있는 AI를 위한 주요 지침

기술지상주의에 대한 경계

‘페미니즘 인공지능’



COVER STORY

AI(인공지능) 기술이 활용됨에 따라 새로운 위험과 부작용에 대한 우려가 함께 논의되고 있다. 특히, AI의 알고리즘 코드와 데이터의 기밀성으로 인해 감지하기 힘든 반면, 알고리즘이 비도덕적으로 이용되는 경우 넓게는 수백만 명에게도 영향을 미칠 수 있어 부패분야에 있어서도 이에 대한 리스크 관리가 필요하다. 국제사회는 AI의 윤리적 사용을 위해 EU AI ACT, UN AI행동규범 등 규제를 마련하고 있으며, 이번 2023년 G20 정상회의에서도 AI국제 거버넌스 마련이 논의되기도 하였다. 이러한 흐름에 따라 기업은 지배구조 강화 외에도 개발 및 배포 프로세스 단계에서의 투명성, 책임성 강화 등을 요구받고 있다. 그러므로 이번 호에서는 기업의 시기술윤리와 거버넌스의 구축 방안에 대해 알아보고자 한다.

01	전문가 코칭	04
	AI거버넌스와 청렴윤리경영 연세대학교 객원교수 조신	
02	사례돌보기	06
	AI윤리 및 거버넌스 정책 사례	
03	보고서리뷰	12
	인공지능의 부패위험 TI(Transparency International), THE CORRUPTION RISKS OF ARTIFICIAL INTELLIGENCE(2022)	
04	행동하는 윤리경영	18
	신뢰할 수 있는 AI를 위한 주요 지침	
05	문화 속 기업윤리	22
	기술지상주의에 대한 경계 '페미니즘 인공지능(Artificial Unintelligence)'	
06	뉴스클립	23
	국민권익위원회 청렴윤리경영 동향 국내외 동향	
07	웹툰 윤리네컷	26
	AI채용과 공정	
08	행사소식	27
09	퀴즈	28



전문가 코칭

AI거버넌스와 청렴윤리경영

조 신 객원교수
연세대학교



이번 호에서는 연세대학교 부설 바른ICT연구소 객원교수 조신 교수님과의 인터뷰를 통해 AI 윤리와 규제 동향 및 기업의 거버넌스에 대한 고견을 들어보고자 한다.

Q1> AI 윤리에 대한 규제가 논의 및 마련되고 있습니다. 그 이유는 무엇이며, 거버넌스 측면에서 기업들이 주의해야할 점 또는 시사점이 있다면 무엇이 있을까요?

인공지능(AI, Artificial Intelligence)은 이미 경제·사회 전반에 큰 변화를 가져오고 있으며 앞으로의 발전 가능성은 가늠하기조차 어렵다. 하지만 AI는 이런 가능성 못지않게 많은 리스크를 안고 있기도 하다.

먼저 현존하는 데이터가 인종, 성, 종교 등에 대해 편향성을 가지고 있기 때문에 이 데이터로 학습한 AI가 편향적인 결과를 내놓는 경우가 있다. 그리고 데이터나 알고리즘의 문제로 AI가 잘못된 결과를 내놓을 가능성이 항상 존재하는데, 그 오류가 자율주행 자동차처럼 위험도가 큰 경우에는 더욱 심각하다. 또한 AI가 결과를 내놓는 과정이 블랙박스화 같아서 AI가 왜 이런 결과를 내놓는지 설명할 수 없다는 점이 투명성과 공정성 관점에서 볼 때 문제다. 마지막으로, AI를 활용하여 가짜 뉴스를 만든다거나 보이스 피싱 등 범죄에 활용하는 등 AI를 악용할 여지도 크다.

ESG 경영 관점에서 볼 때, 이 같은 AI 리스크들은 곧 기업의 사회적 책임에 해당하는 “소셜(Social)”

이슈들이다. 인증 차별, 정보 오류에 따른 소비자 피해, AI의 블랙박스 특성으로 말미암아 소비자 불만에 투명하고 공정하게 대처하지 못해서 발생하는 문제, 고객 정보를 잘못 관리함으로써 발생하는 프라이버시 문제 등 기업이 대응해야 할 이슈 리스트는 끝이 없다.

이러한 이슈들에 제대로 대응하지 못했을 때 발생할 비용, 명성 및 규제 리스크가 매우 심각할 수 있다는 점을 감안하면, AI 이슈는 “거버넌스(Governance)” 관점에서 적극적으로 관리해야 할 문제다. 특히 기업 지배구조의 정점에 있는 이사회가 AI 거버넌스 정비에 관여할 필요가 있다.

Q2> 앞으로 기업의 AI 기술 도입 및 활용은 어떤 방향으로 이루어져야 할까요? 또는 기업 거버넌스 측면에서 AI 기술의 윤리적 사용을 위해 무엇을 할 수 있을까요?

그러나 기업들은 AI 기술이 야기하는 윤리, 공정성, 투명성 문제 등에 제대로 대처할 준비가 되어 있지 않다. 예컨대 딜로이트가 미국 기업들을 대상으로 AI 거버넌스 실태에 대해 조사한 바에 따르면, AI 정책·행동규범 등을 갖춘 기업은 13%에 불과하고, AI 이슈 대응을 위해 프라이버시, 리스크 관리, 기록 보관 등 기업 방침을 바꾼 기업도 9%뿐이다. 또한 이사회에서 주기적으로 AI 관련 주제를 다루고 있는 기업은 15%에 지나지 않았다.

AI의 잠재적 리스크가 큰 데도 불구하고 이처럼 기업들의 준비가 부족하다는 것은 최고경영진의 관심이 미흡한 데서 비롯된 측면도 있다. 따라서 문제 개선을 위해서 이사회의 관여가 필요하다. 먼저 이사회는 경영진에게 AI 관련 보고를 요구하고 이를 이사회에서 심도 있게 토의할 필요가 있다. 다음 단계로, 이사회는 (1) 이사회에서 AI 이슈를 주기적으로 논의하기로 함으로써 AI가 이사회의 아젠다임을 공식화하고, (2) AI 문제를 담당할 임원 및 조직을 갖추도록 하며, (3) AI 관련 사고나 조사 결과를 주기적으로 이사회에 보고하도록 할 필요가 있다.

그리고 이사회의 지휘 하에, 다음 네 가지 측면에서 AI 거버넌스를 체계적으로 설계·실행해야 한다. 첫째, 먼저 기업의 AI 윤리 원칙을 정립하고 여기에 맞춰 업무 규정 및 지침을 수립한다. 둘째, AI 전략 및 통제를 전담할 조직을 신설하고, 전담 조직, AI 기술 조직, 현업 조직 간의 역할을 정립한다. 셋째, AI 거버넌스 구축을 위해, 기획 및 설계, 개발, 평가·검증, 운영 단계별로 거버넌스 프로세스를 상세 설계하고 역할을 지정한다. 넷째, AI 모델에 대한 검증은 성능, 편향성, 설명 가능성 기준으로, 데이터에 대한 검증은 적합성 및 정확성 관점에서 검증한다.

새로운 기술의 안착 여부는 기술의 완성도보다 기업 문화나 인센티브 등 제도적 요인에 달려있다는 것은 널리 알려진 사실이다. 이런 관점에서 AI 시대의 성공적 정착에는 좋은 AI 거버넌스 정립이 필수적이다.



AI윤리 및 거버넌스 정책 사례

사례돌보기



전 산업 분야에 걸쳐 다양한 AI기술이 적용됨에 따라 동시에 최근 AI윤리 문제가 논란이 되고 있다. AI윤리 문제로는 편향성, 인격권 침해, 사생활 침해, 신뢰성과 투명성 상실, 책임 문제, 검증되지 않은 정보제공, 거짓 정보 등이 있다. AI기술의 안전성과 신뢰성을 높이기 위해 정부와 기업은 AI윤리 원칙 또는 법·제도 제정을 위한 연구 등으로 활발하게 대응 방안을 모색하고 있다. 특히 일부 기업들은 기업의 사회적 책임 차원에서 AI기술과 관련된 자체적인 원칙 수립, AI거버넌스 정책을 운영하고 있다. 이러한 사례를 통해 AI윤리의 중요성을 일깨우고 관련 윤리 의식의 제고에 도움이 되고자 한다.

1. IBM

IBM은 하이브리드 클라우드 및 AI, 비즈니스 서비스를 제공하는 기업으로, 자사가 개발하는 AI기술은 ‘인간의 지능을 대체하는 것이 아닌 인간의 능력과 잠재력을 강화하고 확장하는 것이 목적’이라고 밝히고 있다.

IBM은 AI의 장점을 보다 극대화하여 AI를 효율적으로 활용할 수 있도록 지원한다. IBM은 개인과 조직이 책임감 있게 AI를 사용하도록 장려하고 있으며, AI를 사용하는 과정에서 윤리 원칙을 도입해야만 신뢰를 바탕으로 시스템을 구축할 수 있다고 말한다.

IBM은 AI가 모두의 업무 능력을 향상시켜야 하며, AI 시대의 이점이 소수 엘리트가 아닌 많은 사람에게 영향을 미쳐야 한다는 보편성을 추구한다. 이를 위해서 자체적인 ‘AI윤리 원칙’을

제정하였으며, 이 원칙은 IBM이 AI윤리를 보다 투명하고 공정하게 지킬 수 있도록 중심을 잡아준다. 그 내용은 다음과 같다.

〈IBM의 AI윤리 원칙〉

순번	항목	내용
1	설명 가능성	모든 AI시스템은 특정 결론에 도달한 방법과 이유를 설명하고 맥락화할 수 있어야 한다.
2	공정성	적절하게 보정된 AI는 인간이 더 공정한 선택을 할 수 있도록 도울 수 있다.
3	견고성	시스템이 중요한 결정을 내릴 수 있도록 AI는 보안이 유지되고 견고해야 한다.
4	투명성	투명성은 신뢰를 강화하며, 투명성을 증진하는 가장 좋은 방법은 공개이다.
5	개인 정보 보호	AI시스템은 고객의 개인 정보 및 데이터 권한에 우선순위를 두고 보호해야 한다.

출처: IBM 공식 홈페이지

또한 IBM은 책임감 있고 신뢰할 수 있는 AI윤리 문화를 조성하기 위해 AI윤리 위원회를 두고 있다. AI윤리 위원회는 IBM의 내부 직원들로 구성이 되어 있으며, IBM의 윤리 정책, 커뮤니케이션, 연구, 제품 및 서비스를 위해 의사 결정 프로세스를 지원하는 역할을 한다.

2. 세일즈포스(Salesforce)

세일즈포스는 미국의 서비스용 소프트웨어 기업으로, AI기술을 통해 의사결정을 가속화하고 생산성을 향상함으로써 기업이 비용을 절감하는데 도움을 주고, 신뢰할 수 있는 서비스를 제공하는데 활용하고자 한다. 또한 세일즈포스는 AI가 사람들이 생활하고 일하는 방식을 변화시킬 수 있는 만큼 AI기술을 올바르게 사용하기 위해서는 책임감이 우선시되어야 하며, AI기술의 이점이 모든 사람에게 제공되어야 한다고 강조한다. 이러한 책임감 하에, 세일즈포스는 직원, 고객 및 파트너가 AI를 안전하고 윤리적으로 개발하고 사용하도록 장려하고 있다. 세일즈포스는 '생성형 AI' 기술을 통해서 고객 데이터를 관리하는 프로그램을 출시하였으며, AI 기술이 접목된 이 프로그램을 운영하는데 있어서 다음의 '생성형 AI 관련 다섯 가지 원칙'을 제시하였다. 이 원칙에서는 데이터 수집과 사용에 있어서 편향되지 않고, 유해한 결과를 만들지 않도록 하는 기본적인 윤리 기준 등을 다루고 있다.

〈세일즈포스의 생성형AI 관련 원칙〉

순번	항목	내용
1	정확성 (Accuracy)	고객이 자체 데이터로 모델을 훈련시킬 수 있도록 함으로써 해당 모델에서 정확성, 정밀도, 재현율을 검증 가능한 결과로 균형 있게 제공해야 한다.
2	안정성 (Safety)	모든 AI 모델과 마찬가지로 편향성, 유해한 결과물 및 유독성을 완화하기 위해 평가 및 레드팀 ¹⁾ 테스트(red teaming)를 수행해야 한다.
3	정직성 (Honesty)	모델을 훈련하고 평가하기 위해 데이터를 수집할 때 데이터의 출처를 존중하고 데이터 사용에 대한 동의를 확보해야 한다.
4	권한 부여 (Empowerment)	일부 경우에는 프로세스를 완전히 자동화하는 것이 가장 좋을 수 있지만, AI가 인간을 보조해야 하거나 인간의 판단이 필요한 경우도 있다.
5	지속가능성 (Sustainability)	탄소 배출량을 줄이기 위해 가능한 한 적절한 크기의 모델을 개발해야 한다.

출처 : 세일즈포스 공식 홈페이지

세일즈포스는 AI가 인간의 능력과 결합되어 사람들이 더 나은 결정을 내릴 수 있도록 할 때 가장 잘 활용되며, AI를 통해 다양성과 평등을 촉진할 수 있는 만큼 모든 사람들이 해당 원칙과 윤리를 지키는 것이 중요함을 강조한다.

참고

- WEF, IBM Responsible Use of Technology: The IBM Case Study(2021.9)
- IT조선, “AI 발전 만큼 AI 윤리 병행돼야”...카카오, 새로운 알고리즘 윤리헌장 발표 (2019.08.30)
<https://it.chosun.com/news/articleView.html?idxno=2019083002528>
- 디지털조선일보, [칼럼] 인공지능 윤리 표준화가 시급하게 필요하다. (2023.03.22)
https://digitalchosun.dizzo.com/site/data/html_dir/2023/03/22/2023032280216.html
- 카카오 ESG Report 2022
https://t1.kakaocdn.net/kakaocorp/kakaocorp/admin/esg/report/ESGReport2022_KOR.pdf?download
- 카카오 공식 홈페이지 | <https://www.kakaocorp.com/page/>
- IBM 공식 홈페이지 | <https://www.ibm.com/kr-ko>
- 세일즈포스 공식 홈페이지 | <https://www.salesforce.com/kr/>

1) 조직의 전략을 점검, 보완하기 위하여 조직 내 취약점을 발견, 공격하는 역할을 부여받은 하위 조직



인공지능의 부패 위험

■ 보고서: TI(Transparency International), THE CORRUPTION RISKS OF ARTIFICIAL INTELLIGENCE(2022.9.9)

국제투명성기구에서 발간한 보고서 ‘The Corruption Risks Of Artificial Intelligence(2022)’에 따르면 기존의 보고서들은 대부분 AI 기술이 활용됨에 따라 의도치 않게 발생하는 부작용들을 다루어 왔다. 그에 따라 AI와 같은 새로운 기술들이 악의적으로 부패의 ‘조력자’로 남용되어 발생하는 영향은 상대적으로 덜 조명되었다. 이에 국제투명성기구는 이 보고서를 통해 “위탁받은 권력을 가진 자가 사적 이익을 위해 AI 시스템을 남용”하는 것을 의미하는 ‘부패 AI(corrupt AI)’개념을 소개하고 AI가 부패를 구성하는 방식으로 설계, 조작 또는 적용될 수 있는 방식을 정리한다.

이번 보고서 리뷰에서는 해당 보고서를 통해 기업 및 담당자가 AI 기술과 연관되는 부패 위험을 식별하고 관리하는데 참고할 만한 내용을 알아보고자 한다.

AI시스템의 주요 기능

보고서는 먼저 AI기술에 따른 잠재적 부패 위험을 모델링하기 위해 AI시스템의 특징과 부패에 취약한 부분을 설명한다. 우리가 직면한 AI의 부패 위험을 이해하려면 인공지능 중에서도 인공 협소 지능(artificial narrow intelligence)²⁾에 초점을 맞춰야 한다. 이 시스템은 내부적으로 규칙기반, 머신러닝 등 다양한 알고리즘에 의해 운영된다.

〈인공 협소 지능(artificial narrow intelligence)의 알고리즘과 부패 취약성〉

규칙기반 알고리즘	프로그래머가 작업의 모든 부분에 대해 코드로 규칙을 작성하는 알고리즘으로 작성된 코드와 알고리즘 운영 간의 관계가 예측 가능하다. - 취약성: 규칙이 의도적으로 특정 개인이나 그룹에 유리하게 설정된 경우 부패할 수 있음.
머신러닝 알고리즘	데이터 또는 자체 ‘경험’을 통해 학습하여 의사결정을 내리는 알고리즘을 말하며, 민간 및 공공 부문에서 관심과 활용이 확대되고 있다. 세부적인 의사결정과정의 불투명한 경향이 있어 프로그래머도 완벽히 이해하지 못하는 경우도 있다. - 취약성: 불투명성 때문에 부정한 목적으로 악용하려는 사람들에게 취약함.

2) 인간 수준에서 단일하게 수행할 수 있는 작업(또는 관련된 작업의 작은 집합)

AI 부패 위험 파악을 위해 사회 기술적 구현에 대한 면밀한 조사 또한 요구된다. AI시스템 개발 및 구현 과정에는 일반적으로 ①기업 최고책임자(commissioner) 또는 소유주, ②데이터 과학자, ③감사인, ④시스템 사용자의 네 주체가 참여한다.

먼저 개발 과정을 살펴보면, 제도적 맥락에 대한 전문 지식을 갖춘 기업의 최고책임자가 새로운 AI 개발 프로젝트를 시작하여 시스템 목표를 설정하고 데이터 과학자 및 프로그래머를 배정한다.

최고책임자는 데이터 과학자 등과 함께 ‘시스템의 핵심 목표 설정 - 윤리적 가치와 선호도 합의 - 신뢰성 있고 유효한 데이터 식별 - 적합한 예측모델 설정’의 단계를 진행한다. 이 단계들은 피드백 과정을 반복하고 지속적으로 업데이트하여 시스템 보안과 공정성을 보장하고 의도되지 않았으나 잠재된 결과를 완화한다.

AI시스템에서 중요한 단계는 독립적인 감사인의 엄격한 감사다. 감사인은 알고리즘이 악용될 취약성 등 잠재적 문제를 평가하고 인사이트를 최고책임자와 데이터 과학자 및 프로그래머에게 피드백해야 한다. 또한 이상적으로는 의도치 않은 편향성 등 문제 식별 및 광범위한 제도적, 사회적 맥락 평가를 수행해야 한다. 현실적으로 AI의 윤리보다 기술적 취약성에 대해 감사하고 있어 주로 전문기술지식을 가진 데이터 과학자로 구성된다.

그리고 구현 과정에는 AI 의사결정의 직접적 영향을 받는 사용자가 (자신이 AI의사결정의 대상이 된다는 사실을 인식하지 못하더라도) 주체로 포함된다. 이러한 일련의 과정에는 부패에 취약한 부분도 존재한다. 이를 정리하면 다음의 표와 같다.

〈AI시스템 개발 및 구현 과정의 주체별 부패 위험〉

기업 소유주/최고책임자	개인적 이익을 위해 비도덕적 AI 이용
데이터 과학자/프로그래머	개인적 이익을 위해 AI시스템 이용 가능성
감사	뇌물을 받고 AI로 인한 잠재적 피해를 외면할 경우 잠재적인 부패 유발
사용자	사용자가 자신의 이익을 위해 AI시스템의 단점을 이용

부패 AI 유형

보고서는 ‘부패 AI’라는 용어를 ‘(위임받은)권력자가 사적 이익을 위해 AI시스템을 남용하는 행위’로 정의한다. 이때 ‘권력자’란 다른 사람이 필요로 하거나 가치 있게 여기는 자원에 대한 접근권한을 가진 사람을 의미한다. 디지털화 되는 사회에서는 코드와 데이터를 가진 사람들에게 더

강력한 힘(권력)이 있다. 이러한 요소 때문에 비도덕적 AI는 기존의 권력 불균형을 공고히 하고 더 악화시키거나, 더 관리하기 어려운 형태의 범죄가 될 수 있다. 보고서는 비도덕적 AI가 특히 문제가 되는 이유를 다음과 같은 몇 가지 유형 및 사례를 통해 설명하고, AI가 부패에 사용될 때 기존의 부패 유형이 어떻게 변형되는지 보여준다.

1. 부패 AI의 설계

사적 이득을 위해 의도적으로 거짓정보를 생성 및 확산시켜 정보를 조작하도록 설계된 AI시스템을 말한다. 예를 들면 정치인 및 기타 권력자들이 자신의 권력을 공고히 하려는 목적으로 상대 후보자나 반대자를 비방, 위협 및 평판을 훼손하고자 딥페이크 영상을 제작하는 것이 있다. AI시스템에 대한 가장 큰 우려사항은 거짓정보가 정확한 정보보다 빠르게 확산되며 특히 소셜미디어에 깊숙이 침투하는 경향이다.

2. 부패 AI 조작

AI시스템이 악의적으로 설계되지 않았더라도 시스템의 취약점을 악용하여 발생하는 부패를 말한다. 크게는 ‘알고리즘 캡처(algorithmic capture)’와 사용자가 AI를 교란하는 것으로 분류할 수 있다.

- **알고리즘 캡처(algorithmic capture):** 부패에 연루된 데이터 과학자가 특정 그룹에 체계적으로 유리하도록 AI모델의 알고리즘을 조작하는 것을 말한다. 그 예로 AI 채용 시스템의 알고리즘이 원어민 억양을 선호하도록 설계되어 해당언어 능력이 필요하지 않은 직무에도 이민자를 배제하는 것을 들 수 있다. 그 밖에 온라인 구매조달 시스템 또는 사기 탐지 알고리즘의 조작과 관련된 부패도 있다. 이러한 조작은 쉽게 확장되어 대규모의 사람들에게 영향을 미칠 수 있다.
- **사용자의 AI 교란:** 사용자가 사적 이익을 위해 AI를 속이는 경우를 말한다. 의사나 병원직원이 수익 극대화를 위해 사진 픽셀이나 촬영 각도를 변경해서 이미지 분류 알고리즘과 보험 소프트웨어를 교란하는 것을 예로 들 수 있다. 이러한 취약점은 이미지 분류 외에도 대부분의 머신 러닝 알고리즘에서도 입증되었는데, 악용될 시에는 또 다른 부패가 발생할 수 있다. 예를 들면 채용 담당자가 뇌물을 받고 AI채용 시스템의 취약점을 파는 행위가 있다. 뇌물수수는 기존에 존재하던 부패유형이지만 AI 알고리즘 시스템의 불투명성과 결합하면 탐지가 더 어려워진다.

3. 부패 AI 활용

선의의 목적으로 만들어졌으나 부패에 이용되는 AI시스템을 말한다. 권력자들이 사적 이익을 위해 AI시스템을 악용하여 라이벌을 감시, 위협, 협박함으로써 자신의 권력을 공고히 하는 것을 예로 들 수 있다. 연관된 사례로 2016년 페가수스 스파이웨어 사건이 있다. 페가수스 소프트웨어는 이스라엘의 보안기업 NSO가 공식적으로는 대테러를 사용 목적으로 제작했으나, 모로코 정부가

이를 여러 나라의 언론인, 인권 운동가, 정부 관계자들의 스마트폰을 해킹하는 데 사용했음이 국제 언론기관들의 보도되었다.

부패의 주요 원인

SI기술의 특성에서 비롯되는 부패 요인들에 대해 정리하면 크게는 기술적 요인과 인적 요인으로 나뉜다. AI의 기술적 요소는 여러 인적 요인을 유발하는 원인이 되기도 한다.

〈AI 기술과 부패 위험 요인〉

기술적 요인	자율적으로 행동하는 AI 능력	자율성은 AI시스템 자체에 더 큰 책임감을 부여한다, 이런 특성으로 AI시스템 자체가 부패 행위자가 되는 독특한 부패 위험이 발생할 수 있다.
	불투명한 작동방식	알고리즘의 내부적 작동과 결정은 쉽게 설명하거나 관찰하기 어려워 ‘블랙박스’ 알고리즘이라고 부른다. 또한 알고리즘 학습에 사용되는 데이터는 대중의 감시가 닿지 않는 경우가 많아 사적으로 알고리즘이 조작될 가능성이 있다.
	수신자에 대한 대규모 개인화	AI는 ‘마이크로 타겟팅(microtargeting) ³⁾ ’을 통해 특정 수신자에게 맞춤형 콘텐츠를 전달할 수 있으며, 낮은 비용으로도 확장 가능하다. 이러한 기술이 악의적으로 사람을 속이는 데 사용되는 경우, 강력한 조작력을 가진다.
인적 요인	AI시스템을 통한 책임의 분산	사람들은 그동안 동료에게 책임을 분산했으나 이제 AI시스템에도 책임을 전가한다. 책임 분산은 부패의 적발 및 제재 가능성을 감소시키므로 부패의 제약도 줄어든다.
	AI가 연루된 부패의 과실 입증 어려움	데이터나 알고리즘 조작 여부와 주체를 탐지하는 것은 어렵거나 불가능한 경우가 있다. 때문에 부패 행위자는 악의적 조작행위를 부인할 가능성이 높다.
	피해자와의 심리적 거리 증가	AI시스템을 부패에 이용하면, 피해자는 더욱 모호해지고 부패 행위자와 피해자 간 심리적 거리는 더 멀어진다.

권장 사항

보고서는 SI기술과 관련된 새로운 부패 위험에 대응하는 방법을 규제영역, 기술영역, 인적영역으로 나누어 설명한다. 규제영역은 SI부패와 관련된 법 제정을 촉구하는 내용을 다루고 있어 여기서는 제외하였다.

3) 식별된 선호도에 따라 소수의 사람들을 대상으로 광고를 사용하는 것을 말한다.

기술영역: 데이터 및 코드의 투명성 강화, 언어의 상호 운용성⁴⁾ 보장으로 감사 용이성 제고

감사를 수행하려면 데이터와 코드가 공개적으로 사용 가능해야 한다. 데이터 및 코드의 투명성은 책임성을 확보하고 AI부패를 방지하기 위한 중요한 기초 단계다. '알고리즘 저스티스 리그(the Algorithmic Justice League)' 또는 '알고리즘 워치(Algorithm Watch)' 등 시민단체를 통한 독립적인 감사도 실시해 볼 수 있다. 이를 통해 알고리즘이 규제 가이드라인에 명시된 윤리적 원칙을 준수하도록 설계되었는지 확인할 수 있다. 감사의 용이성 제고를 위해 데이터 코드의 투명성 외에도 AI시스템의 지속적인 품질 검사와 머신러닝 프로그램의 상호 운용성 촉진, AI악용에 대응하는 알고리즘의 테스트 및 개발도 중요하다.

인적영역: 부패방지 인식 강화

프로그래머, 데이터 과학자와 코드 감사자는 AI시스템을 구현하는데 중요한 이해관계자가 되었으나 이들에게는 부패 방지를 위한 구체적 행동 강령이 마련되어 있지 않다. 전문 교육(윤리), 행동 강령, 규정 준수 지침을 마련하여 데이터 과학자 및 코드 감사자가 AI부패에 민감하게 하는 것은 부패 방지 및 윤리적이고 책임감 있는 AI보장을 위한 출발점이 될 수 있다.

또한, AI도입으로 내부고발자 역량이 감소하고 있다. AI의 인력 대체로 제보 가능한 사람의 절대적 인원 감소와 AI의 부정행위 가능성을 의심하지 않는 AI기술에 대한 긍정적 인식이 그 원인이다. 또한 AI 프로세스의 불투명성은 내부고발능력을 더욱 감소시킨다. 사람들이 AI부패에 대해 말하게 하려면 먼저, 내부고발능력 감소에 대한 인식을 제고해야 한다.

참고

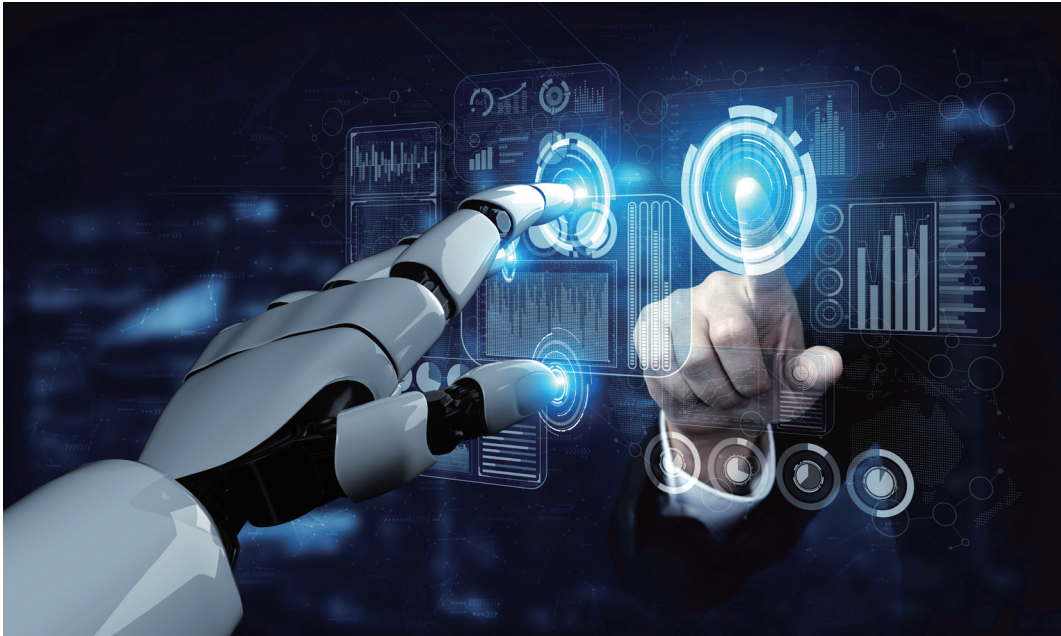
- TI(Transparency International), The Corruption Risks Of Artificial Intelligence(2022.9.9)
- Transparency International Blog, "Bribes For Bias: Can AI Be Corrupted?"(2023.2.17) | <https://www.transparency.org/en/blog/bribes-for-bias-can-ai-be-corrupted>
- 주간경향, "[IT칼럼]AI의 복잡성과 블랙박스 문제"(2023.9.11) <https://weekly.khan.co.kr/khnm.html?mode=view&code=114&artid=202309011056331>
- 글로벌 이코노믹, "마크롱 프랑스 대통령, 스파이웨어 '페가수스' 탈릴까 스마트폰 교체"(2021.07.23) http://www.g-enews.com/ko-kr/news/article/news_all/2021072307434492556336258971_1/article.html
- 네이버 지식백과 IT용어사전(검색일: 2023.09.21) <https://terms.naver.com/entry.naver?docId=1915471&cid=50373&categoryId=50373>
- 네이버 지식백과 AI용어사전(검색일: 2023.09.22) <https://terms.naver.com/entry.naver?docId=6643514&cid=69974&categoryId=69974>

4) 같은 기종 또는 다른 기종의 기기 간에 통신할 수 있고, 정보 교환이나 일련의 처리를 정확하게 실행할 수 있는 것



행동하는
윤리경영

신뢰할 수 있는 AI를 위한 주요 지침



AI 기술의 도입으로 생산성 향상과 업무 지원 등의 긍정적 효과가 기대되는 한편, AI기술로 야기될 가능성이 있는 부작용을 최소화하기 위해 국제사회는 규제 가이드라인 및 규범을 개발하고 있다.

국제투명성기구(TI)의 ‘The Corruption Risks Of Artificial Intelligence(2022.9.9)’ 보고서에 따르면 EU 집행위원회, OECD, 유네스코의 AI 전문가 그룹은 윤리적 AI의 규제 가이드라인들을 제시하고 있는데 그러한 가이드라인의 절반 이상에서 공통적으로 5가지 핵심 원칙(투명성, 정의 및 공정성, 무해성(non-maleficence 또는 ‘do no harm’), 책임성, 개인정보보호)이 언급되고 있다.

한편, 글로벌 회계·재무 컨설팅 그룹 KPMG와 퀸즈랜드대학교가 함께 개발한 AI 활용 체크리스트 ‘Navigating AI(2023)’는 기업이 비즈니스에 AI를 도입 및 사용할 때 신뢰할 수 있는 AI시스템을 운영할 수 있도록 지침을 제공한다. 이번 행동하는 윤리경영에서는 해당 체크리스트(a checklist to provide and use of AI)를 5가지 핵심 영역을 중심으로 살펴보고 기업 및 담당자가 윤리적 AI 정책 구성을 하는데 도움을 주고자 한다.

1 주요 영역별 체크리스트

윤리적 AI를 위한 글로벌 가이드라인들이 공통적으로 제시하고 있는 각 원칙들의 의미와 지침은 다음과 같이 정리할 수 있다.

투명성

‘유네스코 인공지능 윤리 권고’와 EU 인공지능 특별위원회(CAHAI, Ad hoc Committee on Artificial Intelligence)의 ‘AI시스템 법적 프레임워크’에 따르면 AI시스템의 투명성은 다른 원칙과 권리를 보장하는 필수 선결 조건이다. AI시스템이 어떤 기준에 따라 작동하는지에 대한 투명성이 없다면 AI시스템의 영향력 여부나 어떤 영향을 미치는지에 대해 확인하기 어렵거나 불가능할 수 있다. 이는 AI시스템이 수행한 조치에 대한 이의 제기도 어렵게 할 뿐만 아니라, 시스템이 피해를 야기하더라도 개선, 수정이 불가능하게 만든다. 따라서 투명성은 특히 AI시스템이 인권, 민주주의 또는 법치주의에 영향을 미칠 수 있다는 맥락에서 필수적이다.

반면, 투명성의 수준은 항상 맥락과 영향에 따라 적절한 수준이어야 하는데, 이는 투명성과 프라이버시, 안전 및 보안과 같은 다른 원칙 사이에 균형을 이루어야 할 필요가 있기 때문이다. 투명성과 관련하여 KPMG의 보고서에서는 기업이 다양한 이해 수준을 고려해 AI시스템의 알고리즘과 그 결과와 관련하여 다음과 같은 지침들을 제시한다.

- 알고리즘의 기술적 특징과 설계를 문서화하여 작동 방식과 결과에 도달하는 방법을 이해할 수 있도록 했는가?
- 알고리즘의 동작과 설계 논리에 대한 명확한 설명이 있는가?
- 이해관계자에게 어떤 데이터가 수집되고 어떤 프로세스가 자동화되는지 명확하게 알려주고 있는가?
- 이해관계자가 시행결과에 접근 및 확인할 수 있는가?

책임성

CAHAI의 보고서에 따르면 공공 및 민간 조직 등을 AI시스템을 설계, 개발, 배포 또는 평가하는 인공지능 행위 주체는 시스템에 대해 궁극적인 책임과 함께 각 영역에서 법적 규범이 지켜지지 않거나 최종 사용자 또는 다른 사람에게 부당한 피해가 발생했을 경우, 그때마다 책임을 져야 한다. KPMG의 보고서는 시스템 개발 목표 설정부터 배포 및 유지보수 단계까지 전 과정에서 책임과 의무를 명확하게 정의할 것을 권장하며 다음과 같이 지침을 제공한다.

- AI시스템에 책임이 있는 사람이 설계, 개발 및 배포 과정 전반에 걸친 자신의 책임 및 통제권을 명확하게 알고 있는가?
- 이를 적절히 문서화하여 모든 직원이 이용할 수 있도록 했는가?
- 책임, 거버넌스 구조 및 위험 관리 프로세스가 지역 및 국제적 상황에 부합하는지 평가하고 있는가?

개인정보보호

프라이버시는 인간의 존엄성, 자율성, 활동을 보호하는데 핵심적 권리로 AI시스템 수명주기 전 영역에서 반드시 존중, 보호, 증진되어야 한다. 데이터 보호 프레임워크 및 메커니즘은 개인 정보의 수집 사용 공개 및 데이터 주체의 권리 행사와 관련된 국제 데이터 보호 원칙 기준을 참고해야 하며, 동시에 개인 데이터 처리에 대해 인지동의를 비롯한 타당한 법적 근거와 합법적인 목표를 확고히 해야 한다. 인공지능 행위 주체는 개인 정보를 AI시스템의 전 영역에서 보호하고 있음을 보장하고 그들이 AI시스템의 설계 및 구현에 책임을 진다는 것을 보장해야 한다.

KPMG에 따르면 개인정보 영향 평가 및 절차는 법률 준수 및 이해관계자의 윤리적 개인정보 보호 기대치를 충족하도록 마련되어야 하며, 관련하여 다음의 사항들을 고려해 볼 수 있다.

- AI시스템에 특정 데이터 액세스가 필요한 이유를 파악하여 액세스해야 하는 데이터의 양을 줄일 수 있는가?
- AI시스템에 적용할 수 있는 개인정보 영향 평가 메커니즘이 있는가?
- 적절한 데이터 통제 거버넌스가 마련되어 있는가?
- 사용 중인 데이터(예: 개인 데이터)의 취약성을 이해하고 있는가? 사용자에게 민감한 개인 정보가 있는가?
- 시스템의 수명 주기 동안 시스템에서 생성되는 개인에 대한 인사이트와 개인 정보를 동적으로 모니터링할 수 있는 방법이 있는가?

공정성

CAHAI 보고서는 AI시스템이 차별적이거나 부당한 편견과 유해한 고정관념을 지속시키고 증폭시킬 수 있다는 사실을 지적하고 있다. 편향된 AI시스템에 대한 의존은 개인뿐만 아니라 사회 전체에 불평등 심화, 사회결속 약화 등의 부정적인 영향을 미칠 수 있다. 인공지능 행위 주체는 차별적이거나 편향된 결과가 강화되거나 영속화하지 않도록 최대한 합리적인 노력을 기울여서 AI시스템 수명 주기 전 영역에서 시스템의 공정성을 보장해야 한다. KPMG의 지침에 따르면 기업

및 담당자는 AI시스템의 결과를 정기적으로 모니터링하여 다양한 이해관계자를 포용할 수 있도록 불공정한 편견 및 차별이 없고 공정하게 설계되었는지 확인해야 한다.

- AI 모델에 사용된 데이터에 포함된 의도하지 않은 편향성에 대한 완화 조치를 취했는가?
- AI시스템으로 인해 취약 계층이 받을 수 있는 영향을 고려했는가?
- AI시스템이 차별 관련 법 등의 규정을 준수하고 있는가?

무해성(non-maleficence)

CAHAI와 유네스코의 권고에 따르면 AI시스템에서 비롯될 수 있는 개인, 사회, 환경에 대한 피해 방지는 필수적이다. 원치 않은 피해(안전 위험)와 공격에 대한 취약성(보안 위험)은 인공지능 수명 주기 내내 지양되어야 하며 해결, 예방, 제거되어야 한다. 만약, 어떤 결정이 돌이킬 수 없거나 번복하기 어려운 영향력을 가지는 것으로 판단되는 경우나 생사 결정을 논하는 상황이라면 사람이 최종 결정을 내려야 한다. 특히, AI시스템은 사회적 점수 평가(social scoring) 또는 대중에 대한 감시 목적으로는 사용되지 않아야 한다.

KPMG의 보고서에서도 AI시스템의 위험, 의도하지 않은 결과 및 피해 가능성은 배포 전, 배포 중에 완전히 평가되고 완화되어야 한다고 제시한다. 특히 인권과 취약한 이해관계자에 대한 각별한 주의를 요구하고 있다.

- AI시스템에 대한 인권 평가를 수행했는가?
- 안전 설계 원칙을 준수하고 있는가? 시스템 설계 시 취약 계층을 고려하고 보호하고 있는가?
- 기업 차원에서 시스템이 유해한 결과를 초래할 가능성을 매핑하고 적절한 위험 완화 조치를 마련했는가?
- AI시스템이 창출하는 긍정적인 가치나 영향의 타당성을 증명할 수 있는가?

참고

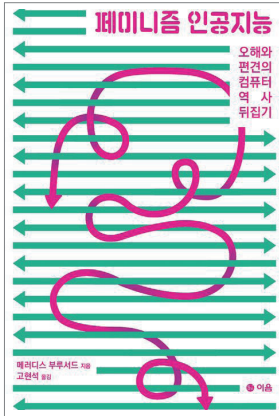
- CAHAI – Ad hoc Committee on Artificial Intelligence, A legal framework for AI systems(2021)
- Deloitte, AI governance for a responsible, safe AI-driven future How Deloitte can help you accelerate your AI adoption(2021)
- Deloitte, Building trust in AI: How to overcome risk and operationalize AI governance
- KPMG, Navigating AI(2023)
- TI(Transparency International), The Corruption Risks Of Artificial Intelligence(2022.9.9)
<https://www.coe.int/en/web/artificial-intelligence/cahai>
- 유네스코 한국위원회-한국법제연구원, AI 윤리의 쟁점과 거버넌스 연구-부록: 유네스코 인공지능 윤리 권고(2021)



문화 속
기업윤리

기술지상주의에 대한 경계,

도서 ‘페미니즘 인공지능(Artificial Unintelligence)’



*이미지출처 예스24

전직 소프트웨어 개발자였던 저자는 도서 ‘Artificial Unintelligence’을 통해 AI 알고리즘의 편향과 위험을 논리적으로 설명하고 있다. 도서의 제목을 직역하면 ‘인공무(無) 지능’이라고 할 수 있으나 내용에서 알고리즘에 반영된 성차별도 다루기 때문에 국내에서는 제목이 ‘페미니즘 인공지능’으로 번역되었다. ‘인공무지능 (Artificial Unintelligence)’이란 복잡한 사회문제 등을 풀기 위해 컴퓨터에만 의존하는 행위로 정의된다. 저자에 따르면 사람들은 컴퓨터 기술을 이용한 의사결정의 허점을 모르거나 지나치게 신뢰하여 적합하지 않은 부분까지도 컴퓨터 기술을 활용한다.

“데이터 과학”이라는 단어는 분석에 기반하고, 객관적이고, 중립적인 느낌을 준다. 저자는 이러한 데이터 과학이 알고리즘을 만드는 사람들에게 의해서 오해와 편견이 반영될 수 있음을 설명한다. 예를 들자면, ‘좋은’ 셀카 사진은 어떤 사진일까? 초점과 앵글이 잘 잡힌 사진이 좋은 셀카일 수 있지만,

알고리즘에 의해 잘 찍었다는 평가를 받은 사진들의 특성을 분석한 결과, 거의 다 젊은 백인 여성의 사진이었다는 점에 주목한다. 사회에서 관습적으로 정의된 ‘매력’의 좁은 범주에 들어맞는 결과는 알고리즘을 만드는 인간의 편향이 반영되어 생성되는 것이다.

의도를 가지고 그렇게 하는 것은 드물겠지만 의도 여부와 상관없이 데이터 과학자들에게는 책임이 있고, 잘못될 가능성에 대해 항상 비판적인 태도와 경계심을 지녀야 한다. 만약 ‘차별’이 사회의 기본값이라면 사회의 시스템은 평등의 개념이 작동하는 방향으로 설계되어야 함을 도서에서는 강조한다. 그러기 위해서 컴퓨터 과학자들은 자신들이 만들고 있는 선례나 설계와 관련된 작은 결정들의 의미를 깊이 생각할 필요가 있다.

이러한 사회적·윤리적 판단은 ‘윤전’ 등 생명과 연관된 분야에 인공지능이 접목될 때 더욱 중요하다. 인공지능의 윤리적 판단에 자주 등장하는 트롤리 딜레마(Trolley dilemma)라는 사고(思考) 실험이 있다. 달리는 기차가 선로를 바꾸지 않으면 5명이 죽게 되고, 선로를 바꾸면 바꾼 선로의 1명이 죽게 된다. 어떤 선택을 하는 것이 윤리 관점에서 올바른 선택이라고 할 수 있는가? 윤리적 판단에 논쟁이 있을 수 있지만, 이 판단이 자율주행 차량에 적용이 된다면 단순한 논쟁거리만은 아닐 것이다. 자율주행 자동차 회사 사무실의 기술자들이 각자 내린 결정이 사회 전체의 이익을 뒷받침하리라는 보장이 없다. 그 의사결정을 누가할 것인지를 생각해본다면, 컴퓨터과학자, 소프트웨어 개발자, 엔지니어들의 윤리적 판단이 정교하게 작용해야 하고, 사회적 합의 과정을 거쳐야 할 것이다. 또한 지속적으로 테스트, 논의, 평가, 검증하여 사회적 편견이 제거될 수 있도록 하는 절차를 통해 지각없이 컴퓨터에 의존하는 ‘인공 무지능(Artificial Unintelligence)’을 경계하는 인식이 필요한 때이다.

참고

- 문화일보, “〈AI 글로벌 최전선을 가다〉AI 알고리즘 편향성·위험성… 풍부한 사례 통해 고발” (2019.8.1)
<https://www.munhwa.com/news/view.html?no=2019080101031803009001>



뉴스클리프

국민권익위원회 청렴윤리경영 동향

내년 예산안 1,116억 원 편성, “신고자 보상·포상금 대폭 확대”



국민권익위원회는 내년도 예산으로 올해 예산보다 166억여 원 늘린 1,116억 원을 편성해 국민권익 보호와 사회 전반의 공정성 제고에 집중하기로 했다. 예산안은 국회 심의·의결을 거쳐 올해 12월에 최종 확정될 예정이다.

국민권익위는 어려움을 겪고 있는 국민의 권익을 신속하고 실질적으로 구제하기 위해 다음과 같은 핵심사업을 추진한다. 먼저, 공정과 상식을 저해하는 부패와 이권 카르텔에 의한 공공재정 낭비·누수를 막기 위해 부패행위나 공공재정 부정수급 신고자에

대한 보상·포상금을 대폭 확대한다. 청렴성을 제고하고 청렴·공정 문화를 확산시키기 위해서는 공공기관 종합청렴도 평가 대상기관에 모든 지방의회를 포함한다. 국민의 편의와 신속, 공정한 민원 해결을 위해 원스톱 행정심판 시스템을 구축해 행정심판 시스템을 일원화하고 증가하는 집단민원을 중점 관리한다.

■ 국민권익위원회 2023년 9월 14일

https://www.acrc.go.kr/board.es?mid=a10402010000&bid=4A&act=view&list_no=46809

국민권익위-인도네시아 반부패 협력 양해각서(MOU) 재체결



9월 25일 오후 3시, 정부세종청사에서 국민권익위는 인도네시아 부패방지위원회와 반부패 협력에 관한 양해각서(MOU)를 재체결하여 반부패 관련 협력과 지원을 강화해 나가기로 했다(‘한-인도네시아 반부패 협력 양해각서’ 체결). 국민권익위는 2006년 12월 인도네시아 부패방지위원회와 반부패 협력 MOU를 최초로 체결했다. 이를 바탕으로 양 기관 간 활발한 반부패 관련 협력 활동을 실시해 왔으나 코로나19 등으로 한동안 교류가 중단됐다. 이번 양해각서는 양국 간의 지속적인 반부패 교류·협력

필요성에 따라 인도네시아측 요청으로 체결하게 됐다. 이를 통해 양국은 향후 3년간 양국의 부패 예방 및 척결 관련 정책과 경험 등을 공유하고, 반부패 제도에 대한 역량강화 및 기술지원 등을 협력할 계획이다. 특히 인도네시아에 소개되어 시행 중인 ‘청렴도 평가’, ‘부패영향평가’ 등을 비롯한 한국의 반부패 정책이 보다 효과적으로 실시되고 발전해 나갈 수 있도록 지속 지원할 예정이다.

■ 국민권익위원회 2023년 9월 25일

https://www.acrc.go.kr/board.es?mid=a10402010000&bid=4A&act=view&list_no=46916

국내외 동향

ESG 선도기업 주가 수익률, “후발기업 보다 50% 더 뛰어나”

기업 지배구조와 리스크 및 투명성을 전문으로 하는 금융자문사 크롤(Kroll)의 최근 연구보고서에 의하면 ESG(환경·사회·지배구조) 등급이 높은 기업이 낮은 기업보다 일반적으로 주가가 더 많이 올랐다. 2013~2021년 9년 동안 배당금과 자본가치 상승을 모두 합친 기업의 총 주식 투자수익률과 ESG 등급 평가회사 MSCI가 발표한 ESG 등급 간의 관계를 살펴보면 ESG 등급이 높은 기업의 주가가 낮은 기업의 주가에 비해 ‘일반적으로 더 높은 수익률(generally outperformed)’을 나타냈다. 크롤은 보고서에서 “전 세계적으로 ‘ESG 선도기업’은 최근 9년 동안 연평균 12.9%의 수익률을 기록한 데 비해 ‘ESG 후발기업’의 수익률은 8.6%에 그쳤다”고 언급했다. 이는 최고 등급의 ESG 기업이 상대적 주가 수익률 측면에서 약 50%의 프리미엄을 더 누린다는 뜻으로 해석된다. 크롤의 평가 디지털 서비스 그룹 칼라 누네스(Carla Nunes) 상무는 “ESG 등급에 대한 규제가 강화되면 이 분야에 어느 정도 통일성을 가져올 것으로 보인다”고 말했다. 또한 “연초 발효된 유럽연합(EU)의 지속가능 금융 공시 규정과 최근 승인된 유럽의 ‘지속가능성 공시 기준(ESRS)’ 같은 의미있는 법안 때문에 이러한 추세가 계속될 것으로 본다”고 전망했다.

■ Bloomberg 2023년 09월 13일

<https://www.bloomberg.com/news/articles/2023-09-13/kroll-s-message-for-haters-esg-makes-money-for-investors#xj4y7vzkg>

■ ESG경제 2023년 9월 19일

<http://www.esgeconomy.com/news/articleView.html?idxno=4640>

구글, ‘위치 정보 편법 수집’으로 캘리포니아주에 합의금 9,300만 달러 지불

9월 14일 캘리포니아주는 구글이 위치 데이터 사용과 관련해 이용자를 기만했다는 이유로 9,300만 달러의 합의금을 지불키로 했다고 밝혔다. 캘리포니아주 법무부는 구글이 개인의 위치데이터를 수집, 저장, 사용하는 방식에서 사용자를 속였다고 주장하며 소송을 제기했다. 이용자들이 ‘위치 히스토리’라는 기능을 비활성화하는 경우 구글은 관련 데이터를 저장하지 않는다고 했지만 안드로이드와 아이폰에 탑재된 검색 엔진을 통해 이용자들의 위치 정보를 계속 수집하고 저장했다는 것이 그 이유다. 구글은 이같은 위치 데이터를 활용해 지역 타깃 광고를 노출한 것으로 알려졌다. 구글은 작년 11월 비슷한 이유로 미국의 코네티컷주 등 40개 주와 3억 9150만달러에 합의하기도 했다. 이번 합의와 관련하여 구글은 합의금 외에 위치 관련 설정 과 실제 위치 추적에 대한 투명성을 강화하기 위해 추가 정보를 사용자에게 공개하기로 했다.

■ 조선일보 2023년 9월 15일

https://www.chosun.com/economy/tech_it/2023/09/15/SFJEZL5L7NBZJF3XGOFMK2BN34/?utm_source=naver&utm_medium=referral&utm_campaign=naver-news

■ 컴플라이언스 위크 2023년 9월 16일

<https://www.complianceweek.com/regulatory-enforcement/google-to-pay-93m-in-california-location-data-settlement/33565.article#toggle>

자산운용업계, 유럽서 ESG 라벨 없이는 펀드 팔기 힘들어져

9월 5일 골드만삭스 그룹 애널리스트들의 분석 결과 유럽에서 펀드를 판매하려는 자산운용사들은 해당 상품이 ESG펀드로 등록되어 있지 않을 경우 판매하기가 점점 더 어려워지는 상황에 직면하고 있다. 이 같은 분석은 유럽연합이 ESG공시 규정집인 '지속가능금융공시규제(Sustainable Finance Disclosure Regulation, SFDR)'를 시행한 지 2년이 지나면서 나왔다. 해당 규제는 2018년 3월 투자 자산의 지속가능성을 둘러싼 위험 및 해당 투자가 사회와 지구에 미치는 영향에 대한 정보를 금융업계가 공시하도록 의무화한 것으로 전 세계적으로 적용된다. 골드만은 SFDR이 처음 채택된 2019년 이후 8조(條)와 9조 두 가지 ESG펀드로의 투자금 유입액이 비(非) ESG 카테고리인 6조 펀드로의 유입액에 비해 3.4배나 증가했다고 분석했다. 6조는 모든 자산운용사에게 지속가능성 위험의 통합과 그것이 그들이 제공하는 금융상품의 수익에 미칠 수 있는 영향에 대한 공시를 요구한다. 8조는 지속가능성 특성을 촉진하는 펀드가 계약 전 공시에서 환경이나 사회적 특성 내지 둘의 조합을 촉진하는 방법과 펀드가 투자하는 기업이 모범적인 거버넌스 관행을 따르는 방식을 명시하도록 요구한다. 9조 역시 지속가능성 목표를 가진 펀드가 계약 전 공시에서 그러한 목표를 달성하는 방법과 지수가 참조 벤치마크로 지정되었는지 여부를 명시하도록 한다. 블룸버그의 해석에 따르면 펀드 업계가 가능한 한 많은 펀드에 ESG 라벨을 붙여야 한다는 압박을 받게 됐고, 그렇지 못하면 고객의 투자금을 유치하기 힘들 수 있다.

■ Bloomberg 2023년 09월 5일

<https://www.bloomberg.com/news/articles/2023-09-04/goldman-analysts-say-funds-without-esg-struggle-to-get-clients?smd=green-esg-investing#ix47vzkg>

■ ESG경제 2023년 09월 5일 <http://www.esgeconomy.com/news/articleView.html?idxno=4530>

AI시대…가짜뉴스 방지·지식재산권 강조 '디지털 권리장전' 발표

정부가 '디지털 공동번영 사회의 가치와 원칙에 관한 헌장: 디지털 권리장전'(이하 디지털 권리장전)을 발표했다. 디지털 권리장전은 디지털 심화 시대에 맞는 국가적 차원의 기준과 원칙을 제시하고 보편적 디지털 질서 규범의 기본 방향을 담은 헌장으로, 전문과 함께 총 6개 장, 28개 조가 담긴 본문으로 구성됐다. 먼저, 디지털 권리장전은 제1장 제3조 '안전과 신뢰의 확보'와 제2장 제7조 '디지털 표현의 자유', 제4장 제20조 '건전한 디지털 환경 조성' 등에서 가짜 뉴스의 확산 방지에 대한 원칙을 제시한다. 헌장은 또한 디지털 기술과 서비스가 "개인과 사회의 안전에 위협이 되지 않도록" 하는 한편, "타인의 명예나 권리 또는 공중도덕이나 사회윤리를 침해하지 않도록 책임 있게 이루어져야 한다"고 강조한다. 특히 "허위 조작 및 불법·유해 정보의 생산·유통이 방지되는 등 건전한 디지털 환경이 조성되어야 하고, 디지털 환경에서 발생하는 범죄로부터 피해자를 보호하기 위한 실효적인 수단과 절차를 마련되어야 한다"고 명시한다. 제3장 제13조 '디지털 자산의 보호'에서 지식재산권 보호에 대해서도 다룬다. 디지털 자산은 정당한 보호를 받아야 하고, 거래에 관한 계약이 공정하고 자유롭게 체결될 수 있도록 보장되어야 한다는 내용이 담겼다. 정부는 이를 토대로 '인공지능법', '디지털 포용법' 등 디지털 법령을 마련하고 권리장전에 따른 관계부처의 정책·제도 정비에 지원할 방침이다.

■ 연합뉴스 2023년 9월 25일

<https://www.yna.co.kr/view/AKR20230925060951017?input=1195m>



AI채용과 공정

웹툰

윤리네컷





행사소식

UNGC(유엔글로벌콤팩트), Korea Leaders Summit 2023

국제기구 고위급 인사, 국내외 지속가능성 이슈 전문가 및 기업 대표들이 연사로 참여해 기업 지속가능성을 내재화하려는 비즈니스 리더들을 위해 UNGC가 추구하는 인권, 노동, 환경, 반부패, 유엔 지속가능발전목표(SDGs), ESG 전반을 아우르는 현안과 인사이트를 공유하는 서밋.

주최 유엔글로벌콤팩트 한국협회

일정 2023년 11월 16일(목)

장소 그랜드 하얏트 서울 그랜드볼룸

참고 <http://unglobalcompact.kr/%ec%86%8c%ec%8b%9d/?mod=document&uid=2514>

제6회 ESG경영과 지속가능성 국제 컨퍼런스

한국기업들이 글로벌 ESG 규제에 대응하는 것을 넘어 새로운 규칙 제정에 적극적인 역할을 할 수 있도록 ESG 관련 글로벌 석학과 실무자 그룹으로 구성된 글로벌 지식 네트워크를 구축하여 ESG 경영의 현재와 미래에 대한 집단지성의 기초를 제공하기 위한 포럼.

주최 국제ESG협회, 고려대학교 ESG 연구원, 한태평양대학협회 지속가능폐기물관리

일정 2023년 11월 28일(화) - 30일(목)

장소 고려대학교 SK미래관

참고 <https://www.globalesgforum.org/ko/>





Q. 다음 중 AI와 관련된 부패의 원인과 가장 먼 것은?

- ① 데이터 및 시스템의 편향성에 대한 완화 조치
- ② AI시스템에 책임 전가
- ③ AI시스템 취약점의 사적인 이용
- ④ ‘마이크로 타겟팅’을 이용한 정보 조작

퀴즈 응모 2023년 10월 24일(화)까지

- (1) ‘응모하기’ 페이지(<https://quiz.assist.ac.kr>)에서 응답하시거나
- (2) 국민권익위원회 민간협력담당관실(esg@assist.ac.kr) 앞으로
정답과 성함, 연락처(휴대폰 번호)를 보내주세요.

정답을 보내주신 분 중 5명을 추첨하여 모바일 기프티콘을 보내드립니다.

* 수집된 개인정보는 상품 발송을 위한 정보로만 활용되며, 추첨 이후 파기됩니다.



지난 호 정답자는

윤○숙님, 이○은님, 박○지님, 서○욱님, 정○형님입니다.
축하합니다!!

민원상담 신청부터 결과확인까지 온라인으로 한번에 끝내기

법령·제도·행정 등 민원상담은 ▶ 국민신문고 www.epeople.go.kr

- | | |
|----------|--------------------------|
| ① 상담신청 | 국민신문고 ▶ 민원상담·안내 |
| ② 상담관지정 | 민원 내용에 따른 상담관 지정 |
| ③ 사실관계조사 | 관계 법령 검토 및 자료조사로 사실관계 확인 |
| ④ 결과확인 | 국민신문고 ▶ 민원상담 신청결과 |

